

José Hierro S. Pescador

Principios de Filosofía del Lenguaje

Alianza Editorial

Principios de Filosofía del Lenguaje

Primera edición en «Alianza Universidad Textos», 1986

© José Hierro S. Pescador
© Alianza Editorial, S. A., Madrid, 1986
Calle Milán, 38, 28043 Madrid; teléf. 200 00 45
ISBN: 84-206-8106-7
Depósito legal: M. 17.932-1986
Compuesto en Fernández Ciudad, S. L.
Impreso en Hijos de E. Minuesa, S. L.
Ronda de Toledo, 24. 28005 Madrid
Printed in Spain

INDICE

Presentación	7
Cap. 1. Filosofía del lenguaje: esbozo de justificación	9
1.1. La filosofía, 9.—1.2. La filosofía del lenguaje, 14.—1.3. Las ciencias del lenguaje, 17.—Lecturas, 20.	
Cap. 2. Signos, signos, signos	23
2.1. Definición y clasificaciones, 23.—2.2. Elementos, 30.—2.3. Nueva clasificación, 31.—2.4. El lenguaje, 36.—2.5. La semiótica, 39.—Apéndice. La teoría de los signos de Occam. 40.—Lecturas, 47.	
Cap. 3. ¿De lo abstracto a lo concreto o de lo concreto a lo abstracto?	48
3.1. Lenguaje, lengua y habla, 48.—3.2. Sistema y norma, 52.—3.3. Competencia y actuación 54.—3.4. La creatividad del lenguaje, 61.—3.5. Recapitulación terminológica, 65.—Lecturas, 66.	
Cap. 4. <i>Ars grammatica</i>	67
4.1. La definición del lenguaje, 67.—4.2. Los universales lingüísticos, 72.—4.3. La gramática como sistema generador, 79.—4.4. El modelo chomskiano de gramática transformacional, 87.—4.5. Desarrollos posteriores de la gramática transformacional, 114.—4.6. Características formales y universalidad de la gramática. 122.—4.7. Justificación de una gramática, 131.—Lecturas, 134.	

Cap. 5. Ideas nonnatas	136
5.1. La teoría chomskiana sobre la adquisición del lenguaje, 136.—	
5.2. Las críticas a la teoría de Chomsky, 139.—5.3. El conocimiento del lenguaje, 145.—5.4. Lenguaje y conocimiento, 149.—5.5. Biología y lenguaje, 156.—5.6. ¿Es específicamente humana la facultad del lenguaje?, 159.—5.7. La tesis de la relatividad lingüística, 168.—Lecturas, 171.	
Cap. 6. A la busca del lenguaje perfecto	173
6.1. Preámbulo, 173.—6.2. Connotación y denotación, 174.—6.3. Sentido y referencia, 177.—6.4. El atomismo lógico, 189.—6.5. Hechos y proposiciones, 196.—6.6. Denotación y descripciones, 201.—6.7. Algunos inconvenientes de la doctrina de Russell, 208.—6.8. El lenguaje como representación figurativa en Wittgenstein, 213.—6.9. Teoría de la proposición, 221.—6.10. La estructura de la realidad, 229.—6.11. De lo que puede hablarse, 240.—6.12. Balance del <i>Tractatus</i> , 253.—Lecturas, 264.	
Cap. 7. Los abusos del uso	268
7.1. Prolegómeno, 268.—7.2. Significado y uso en el segundo Wittgenstein, 272.—7.3. La crítica de los lenguajes privados, 286.—7.4. La Filosofía como descripción de los usos lingüísticos, 296.—7.5. La herencia de Wittgenstein, 302.—7.6. Cómo hacer cosas con las palabras, 311.—7.7. Actos de habla, 320.—7.8. Tipos de discurso, 329.—7.9. Una teoría pragmática del significado, 340.—7.10. La implicación pragmática y la implicación contextual, 349.—Lecturas, 358.	
Cap. 8. Desde un punto de vista lógico	360
8.1. La teoría verificacionista del significado en Carnap, 360.—8.2. El modo material y el modo formal, 371.—8.3. De la sintaxis lógica a la semántica formal: el concepto semántico de la verdad, 377.—8.4. Extensión e intensión: ontología y semántica, 386.—8.5. La teoría del significado en Quine y la crítica al concepto de analiticidad, 401.—8.6. La indeterminación de la traducción radical, 411.—8.7. La regimentación lógica del lenguaje y el criterio de compromiso óntico, 416.—8.8. Crítica a Quine y defensa de la analiticidad, 429.—8.9. Significado y verdad, 436.—8.10. El significado como función, 441.—8.11. Una teoría de los nombres propios, 446.—Lecturas, 456.	
Epílogo. Hacia una teoría unitaria del significado	458
Apéndice. Ideología y lenguaje	470
Bibliografía	478
Índice analítico	494

PRESENTACION

Este libro aspira a ofrecer al lector una visión general y sistemática de un campo teórico que, poseyendo gran actualidad, se halla todavía escasamente ordenado y resulta difícilmente abarcable. Aunque en los últimos años han venido apareciendo algunas obras generales sobre esta materia, que sin duda han de contribuir a familiarizar al lector interesado con los problemas que la filosofía se plantea hoy acerca del lenguaje, no podía dejar de hacer mi aportación a esa tarea. La justificación principal es mi dedicación habitual durante los últimos quince años a la enseñanza de esta disciplina. He venido enseñando Filosofía del Lenguaje, primero en la Universidad Autónoma de Madrid, luego en la Universidad Complutense, y actualmente de nuevo en la Universidad Autónoma. Con la experiencia que dan estos años de docencia, con las sugerencias, aportaciones y críticas de mis alumnos de sucesivas promociones, y, en fin, con ampliaciones que son producto de lecturas siempre renovadas, ofrezco aquí el contenido básico de mi curso.

Con ello aspiro a suministrar un instrumento que facilite el trabajo y que dé la información básica necesaria para poder entender las obras de carácter monográfico y proseguir con un estudio más profundo y creativo. A estos efectos, espero que cualquier lector algo familiarizado con la filosofía y con la lógica pueda servirse con utilidad de esta obra para orientarse en este complejo campo. Todo ello no me ha impedido en absoluto orientar mis explicaciones de acuerdo con mis personales convicciones en materia filosófica, pero estoy seguro de que el lector sabrá siempre obtener del libro la información que necesite sin sentirse obligado a compartir mis juicios críticos.

Esta obra apareció originariamente en dos volúmenes separados, aun cuando el propósito era unirlos en uno, como ahora se ha hecho. El primer volumen contenía los capítulos 1 a 5, bajo el título «Teoría de los

Signos, Teoría de la Gramática, Epistemología del lenguaje», y apareció en 1980. El volumen segundo incluía los capítulos 6 a 8, junto con el epílogo y el apéndice, bajo el título «Teoría del significado», y salió en 1982. Ahora aparecen los ocho capítulos juntos, pero el lector, claro está, notará diferencias de enfoque entre unos y otros. Los cinco primeros capítulos tratan ciertos temas de semiótica, de lingüística, y de epistemología del lenguaje que constituyen, a mi entender, el prolegómeno necesario para cualquier tratamiento filosófico del lenguaje a la altura de nuestros tiempos. Los tres capítulos restantes desarrollan la teoría filosófica del significado, lo que involucra, como se verá en su momento, el problema de las relaciones que tiene el lenguaje con la lógica, con la realidad y con el sujeto hablante. Se trata de estos temas siguiendo un orden que es con preferencia histórico, aunque al tiempo se intenta resaltar los rasgos de carácter sistemático. En el epílogo, articulo entre sí de manera abreviada los conceptos que me parecen básicos para una teoría completa y unitaria del significado. El apéndice, por último, estudia el problema de las relaciones entre la ideología y el lenguaje, tema en el que el pensamiento dialéctico ha hecho su mejor aportación a la filosofía del lenguaje.

Huelga añadir que esta obra tiene un débito muy considerable para con un gran número de alumnos de distintas promociones, en discusión con los cuales se han ido conformando mis ideas. Dar aquí los nombres sería tedioso para el lector, y a la postre injusto para los que mi infiel memoria pasara por alto. Únicamente nombraré a algunas personas que, siendo actualmente colegas míos (y algunos de ellos, antiguos alumnos) han mantenido conmigo continuada relación, y han estado más estrechamente vinculados a la redacción de este libro. Con Daniel Quesada, de la Universidad de Barcelona, he tenido muy esclarecedoras conversaciones sobre teoría lingüística, a la vez que recibía de él muy notable ayuda bibliográfica. Juan José Acero, de la Universidad de Granada, además de enviarme regularmente sus trabajos, ha prestado inmerecida atención a los míos; otro tanto puedo decir de Eduardo Bustos, de la Universidad Nacional de Educación a Distancia. Mi actual colaborador, y gran amigo, Anastasio Alemán, ha discutido conmigo muchos de los problemas que aquí se estudian, obligándome a extremar el rigor; algo semejante puedo decir de mi antiguo colaborador, Jaime Sarabia, y de mi alumno Jorge Rodríguez Marqueze. En fin, la confección del índice analítico la realizaron en su momento dos antiguos alumnos, Manuel Casal y Alfonso Bravo. Por último, debo la corrección de pruebas de lo que fue el segundo volumen de esta obra, a mi hermano Liborio Hierro y a mis antiguos alumnos Alfonso Bravo, José Luis Colomer y Eusebio Fernández.

En cuanto al estilo del libro, el lector notará lo que sigue. He procurado evitar las notas, tanto que sólo he recurrido a ellas para algunas adiciones de última hora. Por esta razón, las referencias bibliográficas aparecen dentro del texto entre paréntesis. No creo que esto resulte más molesto, distrayente o tedioso que consignarlas a pie de página o al final de capítulo. Lo realmente tedioso son las referencias bibliográficas, vayan

donde vayan, pero en una obra de esta índole son inevitables. Aunque casi todas las obras citadas las he consultado en su edición original, siempre que existía traducción castellana he hecho referencia a ésta, a fin de hacer notar al lector la existencia de dicha traducción. Tanto para originales como para traducciones, los datos de edición se encontrarán en la bibliografía final. A cada capítulo sigue un apéndice de lecturas, donde sugiero las más convenientes para una ampliación o complemento de los temas discutidos en el capítulo. Por tratarse de un primer nivel de ampliación, me he limitado, con muy pocas excepciones, a obras en castellano, de las que, por otra parte, contamos ya con una bibliografía bastante amplia sobre estos temas.

EL AUTOR

Madrid, septiembre de 1985.

Capítulo 1

FILOSOFIA DEL LENGUAJE: ESBOZO DE JUSTIFICACION

Las palabras pueden matar todo, incluso el amor.
(Lawrence DURRELL, *Justine*.)

1.1 La filosofía

La filosofía es la forma más general y abstracta de la teoría. Como tal, incluye dentro de sí todo tipo de objeto o de problema siempre y cuando se considere el objeto o se plantee el problema en función de sus conexiones totales, y en una perspectiva que alcance más allá de las limitaciones particulares de una metodología específica. Por ello, la filosofía incluirá no sólo todos los tipos de objetos y de problemas, sino que también aspirará a dar razón de perspectivas concretas sobre esos objetos o problemas, perspectivas tales como la científica, la religiosa, la artística, etc. En última instancia, la filosofía aspira a formular las condiciones más generales de toda teoría y de su articulación con la praxis, dando razón, finalmente, de sí misma. La filosofía es la única forma de teoría que se autojustifica (Heidegger) o se autoelimina (Wittgenstein).

Como toda forma de teoría, la filosofía es histórica, y lo es en grado extremo a causa de su generalidad. Todos los cambios en el pensamiento, tales como la sucesiva constitución de los saberes llamados científicos, el desarrollo de la técnica, el descubrimiento de las relaciones entre pensamiento e interés de clase, la pérdida de vigencia de la religión, etc., han tenido repercusiones en la estructura y en el contenido de la teoría filosófica, y en otros casos han sido, en variable medida, consecuencia de cambios en la actitud filosófica. La filosofía, contra lo que pensaba Carnap en los primeros años del Círculo de Viena, sí tiene objeto, y no solamente el lenguaje de las ciencias: también los propios objetos de éstas en cuanto se consideren en sus interconexiones totales, más allá de los límites metodológicos de cada ciencia, y en el contexto de una teoría global que aspire

a dar razón de todo. Esta concepción no coincide exactamente con la concepción de la filosofía como *geometría de las ideas*, que Gustavo Bueno ha sustentado (*El papel de la filosofía en el conjunto del saber*), pero pienso que la incluye. Pues, en efecto, un tratamiento teórico de nivel superior, metacientífico (en el sentido etimológico de más allá de la ciencia), de los objetos mencionados puede presentarse como una geometría, es decir, como una construcción estructural, de las ideas de esos objetos y de las relaciones entre ellas. Aquí hay que mencionar que, dentro de un cierto contexto filosófico, puede haber lugar para el tratamiento de ideas de escaso contenido teórico; nada se opone en principio a que dentro de una filosofía de la técnica se trate de la idea de mesa, o dentro de la filosofía del arte, de la idea de marco (por citar dos ideas de rango un tanto inferior, que Bueno menciona; y recuérdese la «Meditación del marco» de Ortega en el volumen III de *El espectador*). E incluso puede imaginarse algún contexto en el que tenga un sentido analizar ideas como la de suciedad o la de pelo que, por razones epistemológicas y ontológicas muy específicas, tanto azoraban a Platón (*Parménides*, 130 c). Es indudable que, por este lado, el peligro es la caída en la trivialidad y la ampliación del concepto de filosofía hasta convertirlo en un concepto puramente negativo bajo el que subsumir toda actividad intelectual para la que no tengamos otro nombre mejor. Por ello debe insistirse en la necesidad de que la consideración filosófica de un objeto, cualquiera que éste sea, forme parte de un contexto teórico general que cumpla con las condiciones mínimas que vamos a ver.

Una concepción de la filosofía tan amplia como la que aquí estoy presentando tiene el propósito de dejar fuera el menor volumen posible de lo que históricamente se ha hecho bajo el nombre de filosofía. Por ello, mi concepción no ha equiparado la filosofía a una forma de conocimiento; sin embargo, el término «teoría» es lo suficientemente amplio como para incluir una dimensión cognoscitiva, incluso en el caso de que se quiera concebir la filosofía como una ciencia empírica con facultades de predicción refutable (Kuznetsov, «Pero la filosofía es una ciencia»), cuánto más si el tipo de conocimiento que se reserva para la filosofía es una peculiar intuición intelectual (Husserl, *La filosofía como ciencia estricta*). En mi opinión, lo más peculiar de la filosofía hoy no es el conocimiento, el cual, a menos que nos resignemos a dejar su concepto sumido en la más completa confusión, debe quedar reservado para las ciencias empíricas. Lo peculiar de la filosofía es la construcción de teorías del más alto nivel abstracto y general, entendiendo aquí por teoría toda forma conceptualmente elaborada de interpretación de la realidad y de nuestros modos de trato con ella. Algo, por tanto, que excede de la ciencia y del conocimiento en la medida en que una interpretación no tiene por qué ser validada (confirmada o falsada) recurriendo a la experiencia. En esta medida, la filosofía es una forma del pensamiento, aunque no necesariamente del conocimiento. La filosofía, así entendida, incluye también la sabiduría (*sagesse*) que Piaget (*Sabiduría e ilusiones de la filosofía*, p. 231) considera, a diferencia del conocimiento, propia de la filosofía, y que concibe como

síntesis razonada entre las creencias y las condiciones del saber. Lo que me interesa subrayar es el carácter interpretativo de la teoría. Incluyendo o no la dimensión cognoscitiva, tratando de esencias o de ideas, aceptando lo transcendental o negándolo, razonando mecánicamente o dialécticamente, sujetándose a la lógica o intentando sobrepasarla, ocupándose de la ciencia o centrándose en la literatura, atendiendo al lenguaje o dándolo por supuesto, *la filosofía es una interpretación que aspira a ser total, razonada y autónoma*. Interpretación quiere decir representación en la que se asigna a cada parte de lo representado un sentido, una función o un puesto dentro del todo. Total, porque no hay nada que no pueda, en algún aspecto o condición, ser objeto de esa interpretación. Razonada (o como dirían otros: racional) porque requiere razones, esto es, porque la única justificación de esa interpretación es el propio ajuste de las partes en el todo y la utilidad de la propia teoría en cuanto instrumento de orientación conceptual en el laberinto de la realidad. Autónoma, porque su justificación última no le viene de ninguna otra forma de pensamiento o tipo de discurso. Como ya he dicho, la filosofía da razón de todo, incluso de sí misma. Se autojustifica o, llegado el caso, se autoelimina.

Pienso que, básicamente, esta concepción coincide con la de Waisman cuando caracteriza la filosofía como visión («How I See Philosophy», sección VII), aunque su caracterización parece especialmente adecuada para la metafísica y es efectivamente utilizada para una reivindicación de ésta. Yo insistiría en el aspecto subjetivo de la visión atribuyéndole así un carácter interpretativo. En definitiva, las conexiones entre las cosas las pone el filósofo por medio de los conceptos que constituyen su teoría. No son conexiones entre conceptos simplemente, sino conexiones entre cosas, partes o aspectos de la realidad en cuanto interpretados a través de nuestra teoría. O, como ha escrito Gustavo Bueno, una conexión (o como él la llama, *symploké*) «de las cosas por medio de las ideas» (*op. cit.*, p. 232).

La filosofía se distingue claramente de las ciencias por su aspiración de totalidad que le permite incluir dentro de sí, como uno de sus temas, el de la explicación y justificación lógica, metodológica y epistemológica de las propias ciencias. A cambio de esto, la filosofía renuncia a ocuparse de las cosas con el detalle, la minuciosidad y bajo las específicas exigencias metodológicas de las ciencias, y renuncia, por consiguiente, a conseguir el conocimiento necesario para un manejo exitoso de la naturaleza, que constituye, en cambio, el empeño de las ciencias. La filosofía no está, por eso, obligada a hacer predicciones comprobables o falsables, aunque a veces pueda hacerlas. Naturalmente que si, aceptando esto, se insiste en afirmar que la filosofía es un saber riguroso (por ejemplo, de esencias), y que, en consecuencia, merece el nombre de ciencia, aunque se trate de una ciencia especial, no empírica (Husserl), o se insiste en que la filosofía, debido a su carácter general y fundamentante con respecto a las ciencias, es también una ciencia, sólo que de un tipo general, con respecto al cual las demás ciencias son especializaciones (Kuznetsov), lo único que se está haciendo es ampliar el sentido del término «ciencia» de una manera que me

parece doblemente peligrosa. Peligrosa de un lado porque tiende a ocultar el diferente nivel en el que se desenvuelve la filosofía con respecto a las ciencias y, sobre todo, su independencia y autonomía con respecto a las exigencias metodológicas propias de ellas (exigencias que serán diferentes para las ciencias empíricas y para las ciencias formales, pero ninguna de las cuales tiene por qué cumplir necesariamente la filosofía). Y, de otro lado, peligrosa también porque tiende a reducir el amplio campo de los objetos filosóficos asemejando implícitamente la filosofía a las ciencias existentes, semejanza que viene sugerida por el uso del término «ciencia», por muchos calificativos que a continuación queramos añadirle. Esta es, por cierto, la confusión en la que se mueve la discusión de Husserl sobre la falta del carácter científico en la filosofía, al comienzo de *La filosofía como ciencia estricta*, confusión que aquí viene favorecida por la visión altamente idealizada que Husserl tiene de la ciencia (una actitud muy «fin de siglo»), junto con el hecho de referirse conjuntamente a las ciencias empíricas y a las ciencias formales sin establecer las radicales diferencias epistemológicas entre ambas ni sacar de ellas las obligadas consecuencias. Cuando Husserl afirma, en tono de consternación, que la filosofía «no es todavía una ciencia», tiene razón; pero habría que añadir: ni tiene por qué serlo. Es decir, que sobra el «todavía» y el tono de lamentación. Y no deja de ser irónico que, al intentar hacer de la filosofía una ciencia, Husserl la eximiera de las más elementales exigencias científicas de carácter metodológico y lanzara a la filosofía por una de las sendas más apartadas y divergentes del camino científico que ha conocido el pensamiento moderno. Esto es algo que, por esa época, Wittgenstein tenía muy claro: «La filosofía no es una ciencia natural; la palabra 'filosofía' tiene que referirse a algo que está o por encima o por debajo de las ciencias naturales, pero no junto a ellas» (*Tractatus*, 4.111; Wittgenstein no menciona las ciencias formales porque tiene también perfectamente claro que éstas no tratan de la realidad). Aunque esto no implica que la filosofía haya de reducirse, como Wittgenstein pretende, a una tarea aclaratoria, por lo demás, de justificación metodológica muy problemática desde los supuestos del *Tractatus*.

La filosofía, como teoría total, razonada y autónoma, puede incluir, y ha incluido en diferente proporción según las épocas y los pensadores, todo lo siguiente.

En primer lugar, una concepción general de la realidad, que naturalmente puede empezar por decidir qué es lo que se va a aceptar como real y lo que no, concepción que para algunos equivale a una teoría general de los objetos, esto es, de lo que hay, de lo que se acepta como objeto del pensamiento, o del lenguaje, y que llamarán ontología, aunque otros preferirán llamarla metafísica, reservando acaso el término «ontología» para una teoría acerca de cierta dimensión de lo que hay, o de parte de lo que hay, a saber, la dimensión que consiste en ser. En el pensamiento contemporáneo más atento y próximo a los saberes científicos se tiende a vincular la concepción de la realidad al contenido del conocimiento científico de dos maneras distintas. Por una parte, se sugiere que la metafísica, en cuanto

teoría máximamente general acerca de lo real, puede suministrar hipótesis a las teorías de las diferentes ciencias (Russell, «Logical Atomism»); de otra, se recomienda que una concepción de la realidad se base sobre los resultados de las ciencias y los sintetice, supliendo de esta forma la división y la parcelación del conocimiento que ha traído el desarrollo moderno de las ciencias (y hay que subrayar que una metafísica así entendida era admitida por el propio Carnap, cuyo famoso artículo «Superación de la metafísica por medio del análisis lógico del lenguaje» únicamente iba dirigido contra la metafísica entendida como conocimiento de esencias; cfr. la primera de sus notas de 1957 sobre el artículo citado en la recopilación de Ayer, *El positivismo lógico*).

En segundo lugar, una teoría sobre el hombre y su relación con el mundo (tema que desde el punto de vista académico se estudia bajo el epígrafe de «antropología filosófica»), que incluye una consideración de lo que es propio del individuo en cuanto sujeto de experiencia, y que da lugar a interpretaciones como la del carácter personal del individuo, o su formulación como una estructura integrada por yo más circunstancia, o se desarrolla como una teoría de la conciencia o del alma, o suministra razones para considerar al hombre como un mecanismo. Todo lo último cae bajo la psicología filosófica, también llamada filosofía de la mente. La relación entre estos temas y su tratamiento científico, que tiene lugar en la antropología biológica y en la psicología empírica, no tiene por qué diferir de lo que acontece en el caso más general de los temas del apartado anterior, y es obvio que una filosofía que pretenda estar a la altura de los tiempos no puede dejar de tener en cuenta lo que los antropólogos y los psicólogos científicos han aportado a nuestro conocimiento del hombre. Aquí se puede agregar también lo que se refiere a los dos tipos de valoración que poseen mayor alcance en relación con el tema del hombre, a saber, la valoración ética y la valoración estética. Y, asimismo, todo cuanto tiene que ver con ciertos aspectos peculiares de la vida humana, como son los aspectos social, político y jurídico. Y, en última instancia, la historia.

Tercero, una teoría sobre el carácter, condiciones, alcance y límites de las diferentes formas de conocimiento, y muy en particular, por lo que se refiere a la época moderna, del conocimiento científico.

Cuarto, una elaboración de las formas de razonamiento y de los requisitos de su respectiva validez, que es la tarea de la lógica. En conexión con ello, un estudio del carácter y fundamentos de los sistemas o ciencias formales, como la matemática.

En quinto lugar, una interpretación de la función del lenguaje, que puede incidir en cada uno de los apartados anteriores, según trate del lenguaje en cuanto incorpora una interpretación de la realidad, o bien del lenguaje en cuanto contribuye a la constitución de la conciencia o en cuanto es el medio primordial en el que se expresa el conocimiento, las diferentes valoraciones y cierta considerable porción de los productos culturales, o bien en cuanto el lenguaje puede compararse hasta cierto punto con un cálculo lógico.

Sexto y último, la filosofía puede, como ya he mencionado, tratar de sí misma (metafilosofía), dando razón de su propia validez y justificación y estableciendo sus propios límites. Esto es lo que hace que las presentes páginas sean también estrictamente filosóficas. Y por esta razón las tendencias filosóficas a menudo difieren no sólo en la manera de plantear ciertos temas o en el modo de solucionar tales problemas, sino más aún en la concepción que tienen de su propia tarea. Ciertamente, un buen número de filósofos rechazarían la amplia caracterización que en estos seis apartados he hecho del contenido de la filosofía, eliminando uno o varios de esos aspectos y reduciendo o reformulando el alcance de otros. Como ya he advertido, esta caracterización es a propósito omnicomprendensiva y pretende estar históricamente fundada. Pero yo mismo no estaría dispuesto a dar la misma importancia ni idéntico alcance, en la situación actual, a todos los aspectos reseñados. Añadiré que este aspecto metafilosófico que ahora brevemente estamos viendo incluye esa terapéutica lingüística que consiste en eliminar problemas filosóficos mostrando que su planteamiento es debido a errores de lenguaje, práctica que Wittgenstein puso de moda en el pensamiento contemporáneo y que, en una especie de autoaniquilación nihilista a la que ya me he referido antes, le llevó, a mi juicio equivocadamente (y la equivocación consistió en considerar la filosofía como conocimiento y no como interpretación), a defender la imposibilidad teórica de la filosofía en el sentido de cualquiera de los cinco apartados precedentes.

1.2 La filosofía del lenguaje

¿Qué pasa, entonces, con la filosofía del lenguaje? Es claro que, en un primer aspecto, una filosofía del lenguaje no tiene otra justificación que la que pueda tener una filosofía del hombre, o de la sociedad, o del derecho, o de la naturaleza, o de la historia, o del arte, o de la religión...

En un segundo aspecto, sin embargo, la filosofía del lenguaje tiene hoy una actualidad y un puesto central dentro del sistema de los saberes filosóficos que, como el que en otros tiempos tuvieron la filosofía del ser, la filosofía del conocimiento o la filosofía moral, obedece a muy específicas condiciones, tanto externas al propio desarrollo filosófico como internas a éste. Con la inevitable exageración que comporta este tipo de afirmaciones, se ha dicho que la filosofía del siglo xx ha descubierto el lenguaje. Por imperativo de sobriedad, limitémonos a reconocer que el lenguaje ha sido uno de los grandes temas de la filosofía de nuestro siglo, y más todavía en su segunda mitad. En términos muy generales, por los dos tipos siguientes de razones. Razones externas, como el hecho de que solamente en nuestro siglo se haya constituido una ciencia del lenguaje, que desde las aproximaciones de Saussure ha llegado al considerable grado de desarrollo que presenta en la obra de Chomsky y sus continuadores. Razones internas, que pueden resumirse en el hecho de que, por primera vez en su historia, la filosofía ha cobrado conciencia de que su medio natural y único de expre-

sión, el lenguaje, puede, por su propia estructura en unos casos, por la manera como era usado en otros, haber estado condicionando el planteamiento y solución de ciertos problemas filosóficos. O de todos. A través del análisis del lenguaje la filosofía ha tomado distancia de sí misma y se ha puesto en cuestión. En un caso extremo el resultado ha sido la negación del sentido de sus proposiciones y su autoeliminación como discurso (Wittgenstein). En otro, se ha venido a reconocer que sin una reflexión suficiente sobre el lenguaje nunca sabremos verdaderamente qué es filosofía (Heidegger).

El grado de la atención dedicada al lenguaje, la metodología empleada y los resultados obtenidos han variado ampliamente según la tendencia filosófica de que se trate. En mi opinión, las corrientes filosóficas características del siglo xx en Occidente y con mayor vigencia actual pueden dividirse en tres grandes grupos: el enfoque especulativo, el enfoque dialéctico y el enfoque analítico.

Dentro de las tendencias de tipo especulativo incluyo todas aquellas formas de filosofar que se vinculan a alguna doctrina metafísica, reconociendo, por tanto, la legitimidad epistemológica del discurso metafísico. Por lo que respecta al lenguaje, estas tendencias tienden a estudiarlo en el contexto de una antropología filosófica o de una filosofía de la conciencia, y hacen lo que, para abreviar y a reserva de matizaciones que haré ulteriormente, podemos llamar una *teoría transcendental* del lenguaje. Sus episodios más sobresalientes, y salvando todas las diferencias que existen entre unos y otros, son la teoría de la significación de Husserl, las reflexiones de Heidegger sobre el lenguaje, especialmente en sus últimos escritos, la visión del lenguaje en el contexto de la hermenéutica, en Lipps y Gadamer, el estudio de los símbolos en Cassirer y Ricoeur, el pensamiento de Merleau-Ponty y la gramatología de Derrida. En este ámbito filosófico la preocupación por el lenguaje ha sido más tardía que entre las tendencias analíticas y ha tenido también menos importancia en general. Qué fecundidad tenga este enfoque y qué utilización podamos hacer de él es algo que veremos en su momento ulteriormente.

Incomparablemente mayor ha sido la atención dedicada al lenguaje en la filosofía analítica, y también más decisiva para los resultados de este grupo de tendencias, hasta el punto de que la preocupación por el lenguaje y el análisis del mismo es precisamente la característica más peculiar de este enfoque filosófico. Aquí, el estudio del lenguaje se ha realizado en el contexto del desarrollo de la lógica simbólica y de la filosofía de la ciencia, por lo que ha tenido especial relieve el análisis de las características lógicas o formales del lenguaje, su relación con los cálculos lógicos y, por lo que toca a la relación entre el lenguaje y el mundo, las consecuencias de la verdad y la falsedad y lo que implican estas categorías. Los momentos fundamentales de esta evolución son la teoría del significado de Frege, el atomismo lógico de Russell y el primer Wittgenstein, la filosofía del lenguaje corriente que preludia Moore, se inicia con el segundo Wittgenstein, y continúan Austin, Ryle y Strawson, el pensamiento de Carnap, que va

desde la sintaxis lógica a la semántica formal, las teorías semánticas de Quine, Davidson y Lewis, la teoría del significado de Grice y la teoría del lenguaje de Montague. Aquí hay que incluir asimismo las derivaciones filosóficas de la lingüística, tan patentes en Chomsky y en Katz. También pasando por alto diferencias y precisiones que en su momento haré, puede resumirse el carácter genérico de este enfoque diciendo que desde él se aspira a elaborar una *teoría formal* del lenguaje.

Las filosofías dialécticas son las que parecen haber llegado más tarde a una consideración temática y particularizada de los problemas del lenguaje, probablemente porque su atención hacia la praxis social y política y el papel que en estas tendencias juega la relación entre teoría y praxis, ha contribuido a ocultar la importancia que tiene el lenguaje para la teoría, así como la función que desempeña en la formación y transmisión de la ideología. Sin embargo, hay ya en los clásicos de estas tendencias, Marx y Engels, consideraciones sumamente agudas sobre la relevancia del lenguaje para la filosofía, y hay algún pensador como Voloshinov quien, en época todavía temprana como es el comienzo de los años treinta, tenía ya una clarísima idea de la importancia que una filosofía del lenguaje tiene para un correcto y fecundo planteamiento del problema de la ideología (pero debido a los avatares de la política, concretamente a las purgas estalinianas, Voloshinov ha estado perdido para el pensamiento occidental durante cuarenta años; cfr. mi artículo «Lenguaje, ideología y clases sociales»). Dentro de este enfoque dialéctico, que encierra todas las formas de filosofía marxista desde el leninismo a la filosofía crítica del grupo de Frankfurt, hay que mencionar como aportaciones más destacadas, además de la teoría semiótica de la ideología de Voloshinov, la semántica de la comunicación de Schaff, la discusión sobre el carácter clasista del lenguaje en Nicolás Marr y Stalin, y la concepción del lenguaje como trabajo y como mercado en Rossi-Landi, quien constituye, en mi opinión, el primer intento sistemático de construir una teoría marxista del lenguaje. Para resumir también en este caso el carácter más típico de esta forma de pensamiento, podemos decir, con análogas salvedades a las ya hechas, que su propósito es hacer una *teoría social* del lenguaje.

Hay que reconocer que ciertas manifestaciones de las escuelas y tendencias anteriores son formalmente incompatibles con las de otras, tanto en sus planteamientos más generales, metodológicos y epistemológicos, por ejemplo, como en su doctrina sobre el lenguaje. Por poner un ejemplo extremo, la filosofía del lenguaje del primer Wittgenstein parece claramente incompatible con todas las demás (¡incluida la del segundo Wittgenstein!). De otra parte, la del segundo Wittgenstein es a primera vista incompatible tanto con teorías marxistas sobre el lenguaje como con concepciones especulativas en general; aunque la inversa no siempre es cierta, pues recuérdese que Rossi-Landi ha defendido la posibilidad de una utilización dialéctica del segundo Wittgenstein («Per un uso marxiano di Wittgenstein»). E igualmente, dentro de un mismo enfoque, no todos los pensadores estarán de

acuerdo en rechazar otras perspectivas. Así, mientras que Ponzio (*Producción lingüística e ideología social*) considera ideológica toda teoría formal del lenguaje, Schaff, en cambio («Sobre la necesidad de una investigación lingüística marxista»), se conforma con exigir la introducción de una perspectiva sociológica que completaría la consideración formal propia del enfoque analítico. En la investigación que sigue pretendo mostrar las aportaciones más características y más iluminadoras que los principales protagonistas de cada tendencia han realizado al tema del lenguaje, juzgando las supuestas incompatibilidades entre ellos en función de una voluntad de integración para la que lo decisivo a la hora de la crítica serán únicamente las propias dificultades, incoherencias y oscuridades de cada doctrina. Pero en principio espero evitar el prejuicio de aceptar una única perspectiva, y espero asimismo poder ofrecer una visión, una interpretación, una teoría del lenguaje, lo bastante rica para no dejar fuera completamente ninguna perspectiva suficientemente original, rigurosa y responsable.

1.3 Las ciencias del lenguaje

Una última palabra sobre otras disciplinas teóricas que se ocupan del lenguaje. En primer lugar, la propia ciencia del lenguaje, la lingüística. En cuanto ciencia, y por tanto conocimiento y explicación, de algo empíricamente dado, como es el lenguaje, la lingüística está sujeta a todas las limitaciones metodológicas propias de las ciencias y, como cada una de ellas, posee un campo suficientemente delimitado, sin que sea de su competencia ni el asunto de las relaciones entre el lenguaje y otros fenómenos distintos, ni la integración de sus resultados en un sistema teórico más amplio. La lingüística puede suministrar un conjunto de afirmaciones sobre el lenguaje, la mayor parte de las cuales serán en alguna medida empíricamente contrastables, y sobre cuya base se aspirará a hacer predicciones válidas. Como es natural, esta caracterización es lo bastante amplia como para permitir todas las diferencias existentes entre las distintas concepciones de la lingüística. Aquí hay, además, que tener en cuenta que la lingüística es probablemente la ciencia de constitución más reciente. Desde la antigua filología histórica del siglo pasado hasta la lingüística transformacional hay una línea cuya primera inflexión importante es Saussure, y que conduce hasta Chomsky a través de nombres bien conocidos, como Jakobson, Trubetzkoy, Sapir, Hjelmslev, Bloomfield, Harris, etc. Hay razones para afirmar que, en un sentido pleno, la lingüística sólo se constituye como ciencia con Chomsky, y por ello cualquiera que se interese hoy día por el lenguaje desde cualquier punto de vista tiene que habérselas con él. Desde luego que para los pensadores que juzgan que la filosofía constituye una peculiar forma de conocimiento que alcanza más allá de la experiencia sensible, y que por tanto traspasa los límites de la ciencia, lo que ésta tiene que decir sobre el lenguaje puede resultar insuficiente o irrelevante, o ambas cosas.

Tales pensadores propenderán a desarrollar su reflexión sobre el lenguaje al margen de la lingüística, con independencia de ella y, frecuentemente, en ignorancia de sus resultados. Esta actitud es, además de peligrosa, injustificada. Tiene el peligro de que puede conducir al filósofo a descubrir mediterráneos, a obtener conclusiones que, en la soledad de su gabinete, le parezcan profundas verdades sobre el lenguaje (no hay propiedad filosófica más peligrosa que la profundidad, pues con frecuencia se reduce a mera oscuridad), cuando son ya triviales y antiguas aseveraciones que la lingüística, o alguna otra de las disciplinas científicas que se ocupan del lenguaje, vienen confirmando tiempo ha. Y tiene también el peligro de que puede llevarle a un pantano de irrelevancias, en el que las pretendidas verdades no sean sino un conjunto de proyecciones subjetivas y de extrapolaciones arbitrarias tan ajenas al desarrollo del conocimiento que resulten inútiles y estériles por completo. Ya veremos algún ejemplo de esto en su momento. Y es también injustificada esta actitud, porque en una época tan caracterizada por el desarrollo de los conocimientos científicos como la presente, prescindir de ellos a la hora de hacer filosofía sobre algo que es también objeto de la ciencia, equivale a colocarse en una situación precientífica, y por lo mismo arcaica; es, simplemente, no estar a la altura de los tiempos. Naturalmente esto no solamente es válido para el tema del lenguaje, sino que lo propio habría que recomendar para una filosofía del tiempo y el espacio, de la materia, de la sociedad, etc.

En conformidad con la caracterización precedente de la filosofía, hay que notar que la relación entre la filosofía del lenguaje y la lingüística no se limita a que la primera haya de tener en cuenta los resultados de la segunda (y probablemente también a la inversa). La lingüística es tema a su vez de la filosofía del lenguaje en lo que se refiere a los problemas metodológicos y epistemológicos que aquella presenta. Para una concepción estrictamente neopositivista de la filosofía del lenguaje, y por tanto exageradamente cientifista, la filosofía del lenguaje, especialmente si se entiende como del lenguaje natural (dejando, por tanto, aparte el estudio de lenguajes formalizados), tendería a reducirse a ese estudio lógico, metodológico y epistemológico de la lingüística. Resulta irónico que esta posición haya sido defendida precisamente por un chomskiano, y como tal muy crítico del neopositivismo, como es Katz («What's Wrong with the Philosophy of Language?», escrito junto con Fodor). Años después, no obstante, el propio Katz ha modificado su posición para reconocer a la filosofía del lenguaje un campo netamente distinto al de la filosofía de la lingüística, campo que tiene como tema el conocimiento conceptual, y que él define como «aquel campo en el que se aspira a aprender lo que pueda aprenderse sobre el conocimiento conceptual a partir del modo en que tal conocimiento se expresa y comunica por medio del lenguaje» (*Filosofía del lenguaje*, capítulo 1). De hecho, a juzgar por el libro de Katz, semejante tarea se reduce a extraer de la lingüística chomskiana las consecuencias relevantes para ciertos problemas filosóficos tradicionales cuya solución puede verse

afectada por el análisis del lenguaje, como son el problema de la analiticidad, el de las ideas innatas o el de las categorías. Pero la teoría del lenguaje aplicada sigue siendo exclusivamente una teoría científica del lenguaje. Mi principal diferencia con Katz (al margen de importantes diferencias respecto al valor y solidez de las conclusiones filosóficas extraídas de la teoría chomskiana, que ya expondré con detalle ulteriormente en otro capítulo) es que considero que la filosofía de la lingüística es también parte de la filosofía del lenguaje, y que pienso que ésta incluye, además de aquella, no sólo esa tarea de entendimiento conceptual, tan restrictivamente concebida por Katz, sino, asimismo, todo lo necesario para una síntesis de todos los elementos disponibles (y no sólo los suministrados por la lingüística) en una interpretación general del fenómeno lingüístico.

Por lo que respecta a aquellas disciplinas que, dentro de otras ciencias, tratan del lenguaje, como son la biología, la psicología y la sociología del lenguaje, el razonamiento anterior es igualmente aplicable en su primer punto. El estar mínimamente al tanto de los resultados alcanzados en esas disciplinas no puede por menos de enriquecer, apoyar y dar solidez a cuanto el filósofo quiera decir sobre el lenguaje, por mucho que pretenda trascender las limitaciones propias de las mismas. Y sobre todo puede evitarle el bochorno, escándalo todavía frecuente en la Academia y causa importante del desprestigio que la filosofía tiene en ella, e incluso en la calle, de caer en ingenuos errores o crasas trivialidades propias de una época en la que el pensamiento, falto de métodos y de instrumentos rigurosos y científicos, sólo tenía como solución, frente a la superstición y el mito, la defensa de una pura especulación no por racional menos apriorística (en el sentido de empíricamente incontrastada). Precisamente porque la filosofía es histórica, cosa que suelen pasar por alto las tendencias historicistas, no se puede hacer filosofía al margen de aquello que, en el ámbito del conocimiento y de la teoría, es más característico de la época que vivimos, y determina en mayor medida nuestras posibilidades individuales y nuestras relaciones sociales: la ciencia.

Todo lo anterior es incompatible con la opinión que Chomsky ha expresado no hace mucho de que la distinción entre lingüística, psicología del lenguaje y filosofía del lenguaje parece estar acabando (*El lenguaje y el entendimiento*, cap. 1). Como esta afirmación sólo es comprensible y tiene sentido desde los supuestos de la teoría chomskiana sobre el lenguaje, no la voy a examinar ahora; esperaremos hasta que llegue el momento de considerar en detalle los conceptos fundamentales de la lingüística transformacional. Me limitaré a señalar ahora que esa distinción no ha desaparecido todavía a pesar de la influencia portentosa, y merecida, de la obra de Chomsky, y que su posición se basa tanto en una idea excesivamente restringida de la psicología y de la filosofía del lenguaje como en una expectativa excesivamente optimista sobre el alcance de su teoría lingüística. Pero esto es algo que no quedará claro hasta un momento posterior de esta investigación.

Lecturas

El concepto de filosofía, en relación directa con el problema de su enseñanza académica, dio lugar hace unos años a una interesante polémica entre los profesores Gustavo Bueno y Manuel Sacristán. Sus obras representativas son, respectivamente, *El papel de la filosofía en el conjunto del saber* (Ciencia Nueva, Madrid, 1970) y *Sobre el lugar de la filosofía en los estudios superiores* (Nova Terra, Barcelona, 1968). Mis opiniones quedan del lado de Gustavo Bueno, aunque mi admiración es idéntica hacia ambos. Un interesante y perceptivo diagnóstico sobre la situación de la filosofía en el mundo actual es el que ofrece Emilio Lledó en *La filosofía, hoy* (Salvat, libros GT, Barcelona, 1975). El libro de Pedro Cerezo, *Metafilosofía* (Labor, Temas de Filosofía, Barcelona, en preparación) constituye un excelente estudio histórico y sistemático de gran alcance sobre la evolución del concepto de filosofía y sobre su planteamiento actual.

Entre las obras extranjeras traducidas me parecen dignas de mención dos obras de introducción a la filosofía, que son: *¿Qué es filosofía?*, de S. Körner (Ariel Quincenal, Barcelona, 1976), y, casi con el mismo título, *Qué es filosofía*, de A. Danto (Alianza, El Libro de Bolsillo, Madrid, 1976). Ambas son obras de orientación analítica, que pueden compensarse con una obra tan clásica y tan aguda, pero para mi sorpresa tan poco conocida, como *Marxismo y filosofía*, de Karl Korsch (Ariel, Barcelona). Un libro muy sugerente, que todo filósofo de fuerte propensión especulativa debería conocer, es *Sabiduría e ilusiones de la filosofía*, de Piaget (Península, Barcelona, 1970). Por último, las relaciones entre ciencia y filosofía son objeto de un atractivo debate en los artículos de Ayer, Gellner y Kuznetsov recogidos bajo el título de *Filosofía y ciencia* (Cuadernos Teorema, Universidad de Valencia, 1975).

Por lo que respecta a obras generales de filosofía del lenguaje en castellano, la situación no era muy buena en los últimos años; actualmente, y por lo que se verá a continuación, la situación está mejorando. *La Filosofía del lenguaje* de Alston (Alianza Universidad, Madrid, 1974) es una pequeña y fragmentaria introducción, útil para algunos temas, pero muy simplista para otros, y en conjunto muy insuficiente. Aunque con parecidas limitaciones, resulta más esclarecedora *El laberinto del lenguaje*, de Max Black (Monte Avila, Caracas, 1969), que es también más extensa. *La Filosofía del lenguaje* de Katz (Martínez Roca, Barcelona, 1971) es fundamentalmente una exposición y ardiente defensa de la teoría chomskiana del lenguaje, de la que se obtienen algunas consecuencias filosóficas no muy agudas, y a la que precede una crítica de la filosofía analítica. Aunque el libro no es muy bueno, lo peor, con mucho, es la versión castellana. Las *Indagaciones sobre el lenguaje*, de Ferrater Mora (Alianza, El Libro de Bolsillo, Madrid, 1970), son un ensayo sobre muy variadas cuestiones filosófico-lingüísticas, y contienen una gran cantidad de información, pero creo que entenderlas bien requiere cierto conocimiento previo del trasfondo sobre el que discurren las consideraciones del autor.

Cuando escribo esto acaba de aparecer en castellano la *Filosofía del lenguaje*, de Kutschera (Gredos, Madrid, 1979), obra que es, sin duda, la que más se parece, por su concepción y por su estructura, al presente libro. Las diferencias entre mi libro y el de Kutschera son, sin embargo, numerosas, como podrá comprobar el lector que tenga la curiosidad de compararlos. Así, por ejemplo, la atención con que examinaremos posteriormente la teoría chomskiana de las ideas innatas o la filosofía del lenguaje del primer Wittgenstein no tienen parangón en la obra de Kutschera, en la que además apenas hay nada sobre temas como la sintaxis lógica del lenguaje de Carnap o la filosofía del lenguaje de Russell, temas que también veremos con cierto detenimiento. A cambio hay, naturalmente, otros temas en los que Kutschera se extiende más de lo que lo hará la presente obra, como son, por ejemplo, la teoría de Humboldt, la hipótesis de Sapir y Whorf o la teoría de la gramática lógica. En fin, el lector encontrará asimismo notables diferencias por lo que hace a la ordenación de los temas y al estilo.

Otra obra parecida a la anterior y de semejante nivel es la publicada por tres colegas míos, y dos de ellos antiguos alumnos, Juan José Acero, Eduardo Bustos y Daniel Quesada, con el título *Introducción a la Filosofía del Lenguaje* (Cátedra, Madrid, 1982). El campo que cubren coincide bastante con el cubierto por la obra de Kutschera y por el presente libro, pero ellos dedican una mayor atención a los aspectos técnico-formales y al tratamiento detallado de algunos desarrollos recientes, como la semántica de la computación o la teoría pragmática del significado. Por su estructura y por su estilo, es obra que se diferencia claramente de las otras dos. Aquí, trabajar en colaboración les ha permitido concentrarse cada uno en temas de su especialización, pero el lector puede tener algunas veces dificultad para percibir con claridad el ajuste de los diferentes temas entre sí. Para desarrollos recientes como los mencionados, el lector podrá ciertamente ampliar con el libro de mis compañeros lo que pueda encontrar en la presente obra.

Con las dos publicaciones anteriores, la bibliografía de obras generales de filosofía del lenguaje ha experimentado, sin duda, una mejora cualitativa de gran trascendencia para el lector interesado en estos temas, quien encontrará en ellas un útil instrumento para iniciarse en esta disciplina. Por lo que respecta a la lingüística, la situación ya era más satisfactoria desde tiempo antes. Me limitaré a citar dos obras excelentes de referencia que cubren muy bien el amplio espectro de la ciencia del lenguaje, la *Lingüística estructural*, de Rodríguez Adrados (Gredos, Madrid, 2.^a edición, revisada y aumentada, 1974), y la *Introducción en la lingüística teórica*, de Lyons (Teide, Barcelona, 1971). Quienes prefieran iniciarse directamente en la lingüística generativa tienen el manual de Ruwet, *Introducción a la gramática generativa* (Gredos, 1974). Para otras disciplinas que tratan del lenguaje recomendaré dos obras, la *Introducción a la psicología del lenguaje*, de Herriot (Labor, Barcelona, 1977), y los *Fundamentos biológicos del lenguaje*, de Lenneberg (Alianza Universidad, Madrid, 1975). Una reco-

pilación excelente de trabajos de diferentes autores en muy variadas perspectivas (pero con cuidadosa exclusión de la perspectiva filosófica, el autor sabrá por qué) es la *Presentación del lenguaje*, realizada por Francisco Gracia (Taurus, Madrid, 1972).

En cuanto a obras auxiliares, como diccionarios, pueden utilizarse con provecho el *Diccionario de términos filológicos*, de Lázaro Carreter (Gredos, Madrid, 3.ª edición, 1974); el *Diccionario enciclopédico de las ciencias del lenguaje*, de Ducrot y Todorov (Siglo XXI, Buenos Aires, 1974); el *Diccionario de lingüística*, de Mounin (Labor, Barcelona, 1979), y con idéntico título, el *Diccionario de lingüística*, de Dubois y otros (Alianza, Madrid, 1979). Aunque sea un diccionario filosófico general, por la atención que concede a ciertos temas de filosofía del lenguaje puede ser útil consultar el *Diccionario de filosofía*, de Ferrater Mora (6.ª edición, Alianza, Madrid, 1979). Para algunos temas de filosofía analítica del lenguaje se encontrará información en la *Enciclopedia concisa de filosofía y filósofos*, dirigida por Urmson (Cátedra, Madrid, 1979).

Capítulo 2

SIGNOS, SIGNOS, SIGNOS

El mundo era tan reciente que muchas cosas carecían de nombre, y para mencionarlas había que señalarlas con el dedo. (GARCÍA MÁRQUEZ, *Cien años de soledad*.)

2.1 Definición y clasificaciones

Al comienzo de su libro *Signo*, Umberto Eco nos ha recordado, por medio de una pequeña historieta, que vivimos inmersos en signos. Y así es. Si un signo es *todo cuanto representa otra cosa en algún aspecto para alguien*, entonces la vida humana no es concebible sin signos. ¿Pero es concebible sin signos alguna forma de vida? En el lugar indicado, Eco ha escrito que los fenómenos naturales no dicen nada por sí mismos, que se vive en un mundo de signos porque se vive en sociedad (p. 11). Según esto, los fenómenos sýgnicos, los fenómenos de significación, serían característicos de los seres humanos porque viven en sociedad, y formarían parte de los códigos que rigen las relaciones sociales entre ellos, o como otros preferirían decir, de los usos sociales.

Sin embargo, es cierto que los animales emiten y perciben diferentes clases de signos o señales. En primer lugar, producen determinados tipos de sustancias químicas por medio de las cuales dan a conocer ciertos estados de su organismo o determinadas condiciones del entorno. Tales sustancias se denominan, por ello, semioquímicas, y funcionan, bien entre individuos de especies diferentes, bien entre los pertenecientes a la misma especie. En este último caso reciben el nombre de feromonas. Según Antonio Gallego, «en los insectos, donde han sido muy bien estudiadas, actúan como señales de alarma, dan lugar a la agregación o dispersión de los individuos de la colonia, regulan su conducta sexual y condicionan su organización social. En los mamíferos, las feromonas participan en la organización jerárquica de grupos, en la delimitación del territorio que ocupan, marcado de individuos y en la conducta sexual» («Feromonas», p. 4). Como era de

esperar, y el autor subraya, las respuestas desencadenadas por estas sustancias son estereotipadas en los insectos y relativamente flexibles en los mamíferos, como corresponde a las diferencias en la evolución del sistema nervioso, llegando en el hombre al extremo de que hace dudoso que se pueda hablar en su caso de feromonas (*op. cit.*, p. 8). Tal vez lo que hay que pensar es que la utilidad de éstas decrece en el hombre ante medios de comunicación más eficaces, de tal manera que el sistema feromonal queda atrofiado. Es igualmente cierto, de otra parte, que existen además, entre los animales, señales físicas, como el conocido baile de las abejas por medio del cual se dan a conocer unas a otras la dirección y distancia aproximada de la fuente alimenticia respecto a la colmena (cfr. Von Frisch, *La vida de las abejas*, cap. 11). ¿Constituye todo esto ejemplos de signos? Se notará, naturalmente, que hay en estos casos una comunidad de individuos en la que funcionan las señales mencionadas; pero entre esas comunidades y las sociedades humanas hay toda la distancia que separa a la naturaleza de la cultura. El uso de tales medios de comunicación y de significación, su producción e interpretación, no es algo propiamente aprendido; constituye el producto directo del desarrollo espontáneo de las capacidades biológicas de la especie. Es exclusivamente natural. En este sentido, y en contra de Eco, sí hay fenómenos naturales que digan algo por sí mismos (aunque siempre para algún organismo), y no veo ningún inconveniente en considerarlos como signos, puesto que tienen un significado.

Claro está que ni Eco ni nadie pretende negar los hechos semióticos naturales mencionados. La cuestión consiste entonces en recurrir a otra categoría que los cubra y los distinga de aquellos procesos o fenómenos semióticos típicamente humanos, y que como tales son convencionales y requieren un código. Para tales efectos Eco utiliza la categoría de señal, y caracteriza entonces el signo así: «hay un signo cuando, por convención previa, cualquier señal está instituida por un código como significante de un significado» (*Signo*, secc. 5.3). La categoría de señal es, por consiguiente, más amplia; los signos son señales que cumplen con esas condiciones.

Esta concepción difiere de la concepción clásica, originada en Peirce, quien definió el signo como «algo que está para alguien en lugar de algo en algún respecto o capacidad» (*Collected Papers*, secc. 228). De las varias clasificaciones que hace Peirce para los signos (y cuyos detalles pueden encontrarse en la obra de Eco citada, secc. 2.11), la más conocida, y única relevante para nosotros ahora, es la que atiende a la relación entre el signo y el objeto significado. En su virtud, Peirce distingue tres clases de signos. Iconos o signos icónicos; son aquellos que se refieren a un objeto en razón de sus caracteres propios, lo que quiere decir que algunos de tales caracteres corresponden a los del objeto, y por tanto que entre el signo y el objeto existe una relación de semejanza. Son ejemplo de estos signos las fotografías, planos, diagramas, etc. En segundo lugar, índices, indicios o signos indéxicos, en los cuales hay una relación de efecto a causa, en el sentido de que tomamos algo como signo de otra cosa en la medida en que ha sido causalmente afectado por ella; por ejemplo, como ocurre cuando

tomamos el humo como signo de fuego, la huella como signo de la presencia de un ser humano (no por su semejanza con el pie que la imprimió, pues entonces se trataría de un signo icónico) o la luz roja que se enciende automáticamente en el cuadro de mandos del automóvil como signo del bajo nivel de aceite en el motor. Por último, símbolos o signos simbólicos, que son aquellos cuyo carácter de signo obedece sólo o principalmente al hecho de ser así utilizados, y los cuales carecen, por tanto, de relación propia con el objeto significado. Es lo que acontece con la luz roja que indica un peligro (cuando, a diferencia de lo que ocurría en el ejemplo anterior, no hay una relación causal entre el peligro y la luz), y es lo que acontece en general con las palabras. Los símbolos se distinguen claramente de los otros tipos de signos por cuanto solamente adquieren su carácter de signos en el proceso de la comunicación, y por lo tanto son signos en cuanto que hay reglas que rigen su uso como tales. Esto es lo que se quiere decir cuando se afirma que los símbolos son signos por convención.

Los trabajos de Peirce, de cuya complejidad, riqueza y dificultades internas no puede dar idea el breve y parcial resumen anterior, tuvieron una gran influencia en Morris, quien, medio siglo después, desarrolló algunas de aquellas ideas en el intento de crear una ciencia general de los signos, una semiótica científica, que, en la medida en que se veía obligada a tomar en cuenta los procesos y relaciones en los que aparecen los signos (es decir, los fenómenos semióticos) era de orientación absolutamente conductista. Morris ha caracterizado el signo así: «Si algo (A) rige la conducta hacia un objetivo en forma similar (pero no necesariamente idéntica) a como otra cosa (B) regiría la conducta respecto de aquel objetivo en una situación en que fuera observada, en tal caso (A) es un signo» (*Signos, lenguaje y conducta*, cap. 1, secc. 2). La conducta a la que aquí se alude es la de cualquier organismo, y por consiguiente la categoría de signo abarca a los medios de comunicación animal. Nótese que la anterior no es, en la intención de Morris, propiamente una definición, pues Morris deja abierta la posibilidad de que haya signos que no cumplan con esas condiciones, y hay que subrayar que esta posición todavía la mantiene en un libro muy posterior a la obra clásica citada (cfr. *La significación y lo significativo*, cap. 1, secc. 2). Así entendidos, los signos se dividen para Morris en dos categorías fundamentales, señales y símbolos. Un símbolo es «un signo que produce el intérprete para que actúe como sustituto de algún otro signo del cual es sinónimo» (*Signos...*, cap. 1, secc. 8); según esto, las palabras, en general, son símbolos. Una señal es cualquier signo que no sea símbolo; por ejemplo, el pulso es señal de un cierto estado del organismo.

No muy diferente es la clasificación básica de los signos ofrecida por Schaff, quien, tomando también como categoría más general la de signo, distingue entre signos naturales (síntomas) y signos artificiales (*Introducción a la semántica*, p. 180 y ss.). Estos últimos los clasifica a su vez de manera un tanto complicada e introduciendo ya considerables diferencias terminológicas con respecto a Morris. En primer lugar contrapone, dentro

de los signos artificiales, los signos verbales a todos los demás; en segundo lugar, distingue, en los signos artificiales no verbales, entre los que denomina señales y los que llama signos sustitutivos. El criterio de la distinción es la función que cumplen; las señales tienen la función de influir directamente en la conducta humana, mientras que los signos sustitutivos actúan sustituyendo o representando un objeto, situación o acontecimiento. Una señal es, por ejemplo, una luz verde que da paso o la sirena de una ambulancia que lo pide. Por su parte los signos sustitutivos se dividen en símbolos y los que no son símbolos (que Schaff denomina signos sustitutivos en sentido estricto). Son símbolos aquellos signos sustitutivos que representan nociones abstractas: la balanza es símbolo de la justicia, la paloma lo es de la paz, el color rojo simboliza peligro, etc. Son signos sustitutivos no simbólicos los que representan algo material, como ocurre con las pinturas, fotografías y demás signos de tipo icónico.

Como se ve, la clasificación de Schaff introduce importantes variantes terminológicas con respecto a la de Morris, ya que afecta a términos tan utilizados como «señal» y «símbolo». Mientras que en la tradición de Peirce y Morris, que ha influido ampliamente en la filosofía analítica y en la semiótica, el lenguaje se categoriza como un sistema de símbolos (aunque Morris no acaba de decidirse sobre esto), Schaff separa completamente los signos lingüísticos de los símbolos. En esto hay que reconocer que se encuentra, también, en una importante tradición, la de la lingüística, que se origina en Saussure. En efecto, Saussure distinguió entre el signo lingüístico y el símbolo de la siguiente manera (*Curso de lingüística general*, páginas 129 y ss.): el signo lingüístico lo consideraba como una entidad psíquica compuesta de concepto o significado e imagen acústica o significante, y lo caracterizaba por ser arbitrario, en el sentido de que no hay vínculo interno ni necesario que una significado y significante para constituir el signo. Justamente lo contrario de lo que ocurre en el símbolo tal como Saussure lo concibe, en el cual hay siempre un rudimento de vínculo natural entre significante y significado, y de aquí que el símbolo no sea nunca totalmente arbitrario. Saussure menciona la balanza como símbolo de la justicia, señalando que no vale cualquier otro objeto indistintamente para cumplir esa función simbólica; la balanza tiene algo, la posición de equilibrio que se pretende conseguir entre sus brazos, que se asemeja al contenido básico del concepto de justicia.

La influencia de las definiciones de Saussure ha sido muy profunda en la teoría lingüística, y particularmente dentro de la cultura francesa, y puede encontrarse presente, por ejemplo, en el artículo dedicado al signo en el *Diccionario Enciclopédico de las Ciencias del Lenguaje*, dirigido por Ducrot y Todorov. Por cierto que aquí se distingue el signo no sólo del símbolo, sino también de la señal, de la cual se afirma que provoca una reacción, pero que no implica ninguna relación de significación (p. 125; el autor del artículo es Todorov). Dicha influencia se acusa igualmente en Piaget (*Psicología de la inteligencia*, cap. V), que distingue entre signo y

símbolo en los mismos términos que Saussure, aunque acentuando aún más el carácter icónico del símbolo. Piaget añade, además, a ambos conceptos el índice y la señal, definiendo el primero de estos a la manera de Peirce, y considerando la señal como parte antecedente de un proceso de conducta artificialmente provocado en condiciones experimentales. A las señales les reconoce un carácter semántico semejante al de los índices.

Hay que recordar aquí, también, la curiosa forma en que Wittgenstein utiliza los términos «signo» y «símbolo» en el *Tractatus Logico-Philosophicus*, pues vienen a corresponder respectivamente a lo que Saussure llamaba significante y signo. Según Wittgenstein, el signo es lo que puede percibirse del símbolo, a saber, los sonidos o formas gráficas, a los cuales hay que añadir el modo de significar para que se constituya el símbolo; de aquí que un mismo signo (secuencia de sonidos, rasgos gráficos, etc.) pueda constituir al mismo tiempo diferentes símbolos según su manera de significar (cfr. especialmente proposiciones 3.32 a 3.326). Si despojamos a las definiciones de Saussure de su carácter mentalista, lo que Wittgenstein llama «signo» viene a coincidir bastante bien con la idea de significante, e igual que el signo (lingüístico) es para Saussure un compuesto de significante y significado, el símbolo para Wittgenstein es un compuesto de signo y significado. Como puede apreciarse, el tema del signo es un tema en el que no parece haber límites para la variedad terminológica.

La oposición entre signo y símbolo, que en una u otra forma aparece en casi todas las clasificaciones anteriores, se ha utilizado también para distinguir lo propiamente humano, el mundo de la cultura, de lo puramente animal, la naturaleza. Así se encuentra en Cassirer (*Antropología filosófica*, caps. II y III). Cassirer, que es muy sensible a ciertos usos del lenguaje como el uso emotivo y poético frente a usos privilegiados por el interés filosófico como el uso lógico y científico, prefiere definir al hombre como animal simbólico más bien que como animal racional, en la medida en que la racionalidad no abarca a todas las formas de la cultura, y el simbolismo sí. De aquí que Cassirer vea conveniente distinguir entre los signos, propios de los procesos semióticos animales, y de hecho reducidos por él a señales, y los símbolos, característicos del universo humano. Las señales son parte del mundo físico del ser; los símbolos lo son del mundo humano del sentido (p. 57). Late aquí, como puede apreciarse, una división radical entre naturaleza y cultura, división característica del pensamiento neokantiano en el que Cassirer hunde sus raíces. Hay que notar que caracterizaciones parecidas pueden encontrarse asimismo en obras de antropología científica. Por ejemplo, en un reciente manual de la materia, se afirma que «el género Homo se distingue por la posesión de instrumentos y el empleo de símbolos» (Valls, *Introducción a la antropología*, cap. 1).

En un sentido distinto, y en el contexto del análisis hermenéutico de la obra de arte, Gadamer ha distinguido también entre el signo, cuya esencia consiste en referirse o apuntar a algo, y el símbolo, cuya esencia es reemplazar o estar en lugar de otra cosa (*Verdad y método*, pp. 202 y s.).

Creo que las principales divergencias de interpretación de los conceptos de signo y símbolo pueden conectarse a alguna de las fuentes que he mencionado. Para facilitar la referencia, resumo estas clasificaciones en el cuadro siguiente.

PEIRCE:	Signos	Iconos Indices Símbolos			
MORRIS:	Signos	Señales Símbolos			
SCHAFF:	Signos	Naturales (síntomas) Artificiales	Verbales No verbales	Señales Signos sustitutivos	Símbolos No simbólicos
SAUSSURE:	Signos lingüísticos Símbolos				
TODOROV:	Señales Signos Símbolos				
PIAGET:	Señales Indices Signos Símbolos				
CASSIRER:	Signos Símbolos				
GADAMER:	Signos Símbolos				

En las diferencias terminológicas que hemos comprobado es posible que concurren las específicas influencias teóricas propias de cada autor juntamente con el contexto específico particular en el que introduce sus términos y la necesidad de definirlos y delimitarlos con los propósitos clasificatorios que son patentes. Casi todas las clasificaciones y teorías sobre los signos parecen ser sensibles a una dualidad básica entre lo que, provisionalmente, podemos denominar lo natural y lo convencional. De otro lado, suele reconocerse asimismo una diferencia entre el signo y el símbolo, pero sobre ello se perfilan dos posiciones contrapuestas: mientras que para los autores conectados con la semiótica y la filosofía analítica los símbolos son una subclase de los signos (Peirce, Morris, Schaff), para aquellos situados al margen de esa dirección, los símbolos constituyen una clase contrapuesta a la de los signos (Cassirer, Gadamer, Saussure; para este último, contrapuesta específicamente a los signos lingüísticos). Como otra subclase de los signos aparecen a veces las señales, salvo en algún caso en que son contrapuestas a los signos y a los símbolos como una tercera clase, o bien equiparadas a los signos.

Creo que es muy claro que estos términos, principalmente «signo», «señal» y «símbolo», no tienen en el lenguaje común límites del todo precisos, y que en esta medida su definición y delimitación a efectos teóricos ha de resultar por fuerza un tanto artificiosa e incongruente con el uso ordinario. Pero por lo mismo, tampoco puede pretenderse, y más de un autor lo pretende, que una definición determinada o una específica manera de clasificación haya de ser la única correcta y acordada con el uso corriente de esos términos. No estará de más, por todo ello, que empecemos por echar un vistazo al uso del castellano.

Según el *Diccionario de uso del español*, de María Moliner (edición de 1967), «señal» es, en su primera acepción, «cualquier cosa que sirve para indicar algo», y con este sentido es sinónimo de «signo» y también de «seña». Por su parte, «signo», en su primera acepción, es «cualquier cosa, acción o suceso que, por una relación natural o convencional, evoca otra o la representa». Hay por lo pronto una coincidencia semántica fundamental entre «seña», «signo» y «señal», más clara aún entre estos dos últimos términos, pues el primero tiene acepciones un poco más restringidas. «Símbolo», en cambio, aparece especificado en relación con los términos anteriores, que serían su género inmediato; su primer sentido dice así: «cosa que representa convencionalmente a otra». Es decir, que mientras que en el signo puede darse una representación natural o convencional, en el símbolo solamente cabe esta última. Hay indicios para pensar que la autora tiene en la mente una restricción aún mayor del concepto de símbolo, pues en los ejemplos que da, lo representado es siempre algo abstracto y el símbolo siempre un objeto material; los ejemplos son estos: «La azucena es el símbolo de la pureza. El olivo es el símbolo de la paz. El papel moneda es un símbolo del valor de las cosas.» El contraste, pues, entre signo o señal, en general, y símbolo, como lo particular, resulta claro. Las acepciones correspondientes a estos términos en el *Diccionario de la Real Academia Española* (edición de 1970) tienen un parecido sentido a lo anterior, pero son, en mi modo de ver, mucho más inexactas y no corresponden en tanto grado a mis intuiciones semánticas de hablante nativo del castellano, razón por la que prefiero pasarlas por alto.

Para los amantes de las etimologías, recordaré que «signo» viene del latín *signum*, que Cicerón define como «*quod sub sensum aliquem cadit et quiddam significat*» (según se recoge en el *Dictionnaire étymologique de la langue latine*, de Ernout y Meillet, edición de 1951); «seña», por su parte, procede del plural de la palabra anterior, *signa*, y «señal», de un adjetivo formado tardíamente sobre el sustantivo anterior, *signalis*, que significa «que sirve de signo» (según el *Diccionario crítico etimológico de la lengua castellana* de Corominas). El significado más general y primario, tanto de *signum* como del correspondiente término griego *σημεῖον* (y también de *σην*), era el de «marca distintiva por la que algo es conocido». En cuanto a «símbolo», proviene del latín *symbolum*, que significaba «signo de reconocimiento», por tanto un particular tipo de signo. El término latino transcribe el griego *σημαῖον*, que significaba primeramente cualquiera

de las dos mitades de un objeto previamente partido y dividido entre dos personas celebrantes de un contrato; cada una de ellas conservaba una mitad para servir de prueba ulterior de su identidad como parte contratante. El término, que estaba por ello conectado con el verbo συνήμι, «unir», significaba también cualquier contraseña que probara la identidad, así como, por extensión, cualquier garantía o contraseña en general.

Para frustración de quienes, por una influencia desmesurada y acrítica de Ortega y Heidegger, esperan grandes revelaciones de la etimología de las palabras, no me parece que esta breve incursión por los arcanos de nuestra lengua nos suministre sorpresa alguna, antes al contrario parece apoyar plenamente las dos ideas básicas a las que implícitamente habíamos arribado: que los símbolos son una clase de signos, y que la idea de signo comporta la idea de una relación entre dos entidades, según la cual una de éstas remite a la otra. Es de subrayar, no obstante, la desconexión originaria, etimológica, entre «signo» y «símbolo».

2.2 Elementos

Como se recordará, he definido el signo al comienzo de este capítulo como *todo cuanto representa otra cosa en algún aspecto para alguien*. A fin de que la definición resulte suficientemente general y exacta, conviene tener en cuenta las siguientes precisiones. En primer lugar, hay que entender el término «representar» en su sentido más primario, a saber, como «hacer presente», y no en el sentido, más restringido y derivado, de «sustituir o hacer las veces de». En segundo lugar, hay que tener en cuenta que la cosa representada o evocada por el signo puede ser tanto una cosa propiamente, es decir, un objeto material, cuanto una idea abstracta, una propiedad de un objeto, un sentimiento, un contenido proposicional, etc. En tercer lugar, el término «alguien» se refiere a cualquier organismo capaz de utilizar signos. Aunque en lo sucesivo vamos a ocuparnos fundamentalmente de los seres humanos y de un particular sistema de signos, el lenguaje verbal, aceptaré por definición que solamente los organismos vivos utilizan signos.

Los elementos fundamentales del signo son los siguientes. En primer lugar, lo que sirve de signo, significante, que debe ser algún objeto perceptible por los sentidos. En principio, y a la vista de las dificultades de la física actual con el concepto de materia, no hay que pensar que tal objeto haya de ser necesariamente material. Queda también abierta la cuestión de cómo definir la percepción y de si incluir en ella supuestos sentidos internos, como los propioceptores y los interoceptores (cfr. Pinillos, *Principios de psicología*, cap. 3). Nada, en principio, se opone a ello. Hay que tener en cuenta, asimismo, que para los organismos muy inferiores los sentidos se difuminan en receptores muy simples y primarios. Todo esto son problemas para una teoría de los signos de alcance biológico general, pero a nosotros no nos afectan más que de una forma tangencial, por lo

que basta dejar constancia de ellos. Nótese que este concepto de significante no coincide con el de Saussure, primero porque se refiere a cualquier tipo de signo y no sólo al lingüístico, y segundo porque pretende evitar cualquier connotación mentalista. A este respecto, consideraré como significante propio para cualquier signo lingüístico a la reproducción material, hablada o escrita, de ese signo en cada utilización concreta, o sea, tomando el signo como acontecimiento individual y concreto, lo que se ha llamado «*token*» en la literatura anglosajona, o «*sinsigno*» (cfr. Eco, *Signo*, 2.7). Por ejemplo, la palabra «ejemplo» en la frase anterior. A diferencia del signo acontecimiento, el signo tipo (*type*) o legisigno, no es más que una abstracción perteneciente a esa otra abstracción que es el sistema. Así, y por lo que respecta al sistema de la lengua castellana, la palabra «ejemplo» en general, en cuanto *lexema* infinitamente replicable en sus diversas utilizaciones como signo acontecimiento. Estas consideraciones excluyen implícitamente cualquier sistema de signos puramente mentales, como es el llamado lenguaje mental, de antigua tradición en la filosofía (se remonta por lo menos a Aristóteles) y de importante función, pues a él se subordinaba el lenguaje oral y escrito, como ocurre en la interesante teoría de los signos de Occam (véase el apéndice a este capítulo).

El segundo elemento del signo es el significado, a saber, aquella función que hace de algo un signo. El significado no es el objeto que el signo representa, evoca o hace presente. Todo lo más, y en los casos en los que la función significativa se agote en la referencia, podremos decir que el objeto es *lo* significado por el signo. Así, podemos pensar que el fuego es lo significado por el humo, o que cierto edificio de Grecia, actualmente ruinoso, es lo significado por la expresión «el Partenón». Pero *el* significado, que también podríamos llamar «la significación», del humo o de la expresión citada, es una cierta función que consiste en remitirnos respectivamente al fuego o a cierto edificio. El tercer elemento fundamental es el intérprete, aquel organismo para el cual el signo es signo, o sea, el receptor, que es quien realiza el paso del signo a lo significado haciendo operativa la conexión entre ambos. En muchas especies, incluyendo la humana por lo que respecta al lenguaje verbal, todos los individuos normales son receptores no sólo de las emisiones ajenas sino también de las propias, de manera que todo emisor es al mismo tiempo receptor, y por tanto intérprete.

2.3 Nueva clasificación

Los intentos más ambiciosos de clasificar los signos parecen pecar o de inexactos o de fútiles. Son tantos los criterios que pueden entrar en una clasificación general de los signos, que el resultado final es inexacto y confuso. Por ejemplo, en la clasificación de Schaff, resulta poco convincente oponer las señales a los signos sustitutivos dentro de la categoría de los signos no verbales. De una parte, también las señales sustituyen a otra cosa; la sirena de la ambulancia hace presente a la propia ambulancia por

un medio más fácilmente perceptible que esta última, a saber, su sonido y su luz (aunque, por supuesto, su propósito último sea conseguir paso libre), y la luz verde en el semáforo hace presente la autorización para pasar. En este último caso, por cierto, no se ve cuál sea la conducta que se pretende originar, modificar o detener, ya que nadie está obligado a pasar; cuando Schaff escribe: «la aparición de la luz verde en la esquina de una calle es una señal para que los peatones crucen» (p. 186), esto es inexacto tomado literalmente; la luz verde no es una señal para cruzar, porque la autoridad municipal no pretende que nadie cruce por ese lugar en ese momento; la luz verde es un signo de autorización o permiso para que el que quiera cruzar por ese lugar en ese momento lo haga. Otra cosa es lo que acontece con la luz roja. Aquí sí hay un propósito de modificar la conducta, a saber, la de todos aquellos que pretendan cruzar por un lugar cuando el semáforo está rojo. De otra parte, hay signos a los que Schaff llama sustitutivos, cuyo propósito es desencadenar un comportamiento. Por ejemplo, el retrato robot de un delincuente distribuido a todas las comisarías de una región, pues no persigue sino facilitar la búsqueda y captura del sujeto en cuestión. No pretendo negar, con esto, que haya signos cuya utilización tiene, en general, como finalidad originar, modificar o detener el comportamiento. Lo que quiero decir es que no veo en ello un criterio apropiado para distinguir esos signos de otros dentro de una clasificación general. Pienso que tal propósito constituye una característica completamente externa y puramente accidental con respecto al carácter propiamente semiótico de los signos.

Otra perplejidad semejante, también referida a la clasificación de Schaff, es la que produce la caracterización de los símbolos como signos sustitutivos materiales que representan ideas abstractas en virtud de una convención aceptada (p. 190). Esta caracterización coincide muy exactamente con la primera acepción de «símbolo» recogida en el diccionario de María Moliner que acabamos de ver. Es curioso que Schaff declare que «a la luz del uso lingüístico existente es dudoso que los signos matemáticos y lógicos, por ejemplo, puedan considerarse símbolos (aunque se dice con frecuencia que pueden serlo)» (p. 191). Es curioso, porque a los signos matemáticos se les llama símbolos desde muy antiguo, y la lógica contemporánea, por operar con ciertos tipos peculiares de signos, es llamada simbólica. Por consiguiente, si hay algo que justifique el admitir como símbolos los signos de la lógica y la matemática es precisamente el uso lingüístico. Y desde el punto de vista de Schaff tampoco parece completamente incorrecto. Pues los signos lógicos y matemáticos son materiales (rasgos gráficos sobre una superficie lisa, y derivadamente sonidos cuando las fórmulas son leídas; aunque cabe pensar verosímilmente que no es esto lo que Schaff considera objetos materiales), representan ideas abstractas (igualdad, funciones veritativas diversas, deductibilidad, tipos de cuantificación, etc.), y lo hacen ciertamente sobre la base de una convención previamente aceptada. No se trata de negar que una cosa es simbolizar la paz por medio de una paloma blanca o el luto por medio del color negro o la justicia por una balanza,

y otra, un tanto diferente, simbolizar la cuantificación universal mediante una uve mayúscula invertida, o la estructura condicional por una flecha horizontal apuntando hacia la derecha. Pero resulta por lo menos llamativo que se quiera negar el carácter simbólico de estos últimos ejemplos, y más aún que esto se presente como algo intuitivo y basado en un uso lingüístico.

Pienso que una clasificación general y unitaria de los signos no es viable debido a la diversidad de los criterios que son utilizables y que se entrecruzan. Propondré, por ello, una pluralidad de clasificaciones parciales según los principios de división que me parecen más relevantes.

En primer lugar, según el intérprete al que le son propios, los signos pueden ser humanos y no humanos. Esta clasificación puede ser tachada de antropocéntrica, pero en definitiva esta investigación trata del lenguaje humano, y es en esa clasificación donde éste primeramente adquiere relevancia.

En segundo lugar, según que los signos se den en el ámbito de la naturaleza o en el de la cultura, podemos distinguir entre signos naturales y culturales. Los signos naturales lo son por tener con lo significado una relación puramente natural, esto es, fruto espontáneo de la manera de ser y comportarse los objetos, sin intervención ni mediación de convenciones ni reglas interpretativas propias de una cultura. Esta división está muy próxima a la anterior, pero no coincide con ella totalmente. La razón es que, si bien todos los signos no humanos son, por definición, naturales, los signos humanos pueden ser tanto naturales como culturales. Las feromonas sexuales en los diferentes animales, la danza de las abejas o el grito de alarma entre ciertos primates no humanos son signos naturales producidos por las especies respectivas para la comunicación entre sus individuos, y a veces también entre individuos de especies distintas. El humo como signo del fuego es igualmente un signo natural, y lo es tanto para los hombres como para aquellos animales para los que el humo tiene esa función semiótica. En cuanto a los signos producidos por el ser humano y utilizados en su comunicación, es patente que hay tanto signos naturales como signos culturales. El olor a sudor en un lugar es signo natural de la presencia allí inmediatamente anterior de un ser humano; el olor de *Eau de calandre*, en cambio, es un signo cultural que nos remite asimismo al paso de un ser humano, e incluso, más concretamente, de una dama (el lector de Durrell recordará al protagonista de *Justine* evocándola al aroma de un perfume maravillosamente llamado *Jamais de la vie*). Es claro que sudar es natural, pero perfumarse es cultural. Por lo mismo, diversas manchas y coloraciones en determinadas partes del cuerpo son signos naturales de distintas situaciones anormales de partes del organismo, como es signo natural de cierto tipo de sensaciones llorar; pero es un signo cultural de esta misma sensación vestir de negro. Con esto no quiero negar que la cultura influye en la atrofía o represión de ciertas manifestaciones naturales o incluso en algunas características accidentales de las mismas, pero la distinción básica entre signos naturales y culturales no me parece negable, si se toma el signo con la generalidad con que aquí lo hemos tomado.

En tercer lugar, según la estructura de los signos podemos distinguir entre signos verbales y no verbales. Los signos verbales constituyen siempre un sistema de posibilidades de combinación por lo menos en dos dimensiones, según el medio material (por ejemplo, la dimensión fonológica en el lenguaje oral, cuyas unidades son los fonemas) y según la significación (cuyas unidades, en el ejemplo del lenguaje oral, son los morfemas o monemas). A los signos que carecen de tales características los consideraremos no verbales. No hace falta subrayar que son signos verbales no sólo los del lenguaje oral, sino también los del lenguaje escrito y los realizados, por ejemplo, por medio de los dedos cuando cada signo corresponde a un fonema o grupo de fonemas con independencia de la significación. ¿Cómo se relaciona esta clasificación con las anteriores? Por lo pronto, todos los signos no humanos son igualmente no verbales. En cuanto a los signos humanos, es claro que tanto los que son naturales como los que son culturales pueden ser no verbales, como lo son todos los ejemplos mencionados hace un momento. Pero los signos verbales, a pesar de que la capacidad para tenerlos es una característica central de la naturaleza humana, son, como atestigua su diversidad y su variabilidad histórica, un producto de la cultura. Nótese que al definir los signos verbales en función de la doble articulación (concepto introducido por Martinet) eliminamos la posibilidad de hablar con sentido de un lenguaje verbal de carácter puramente mental, pues en tal lenguaje faltaría el modo material de articulación.

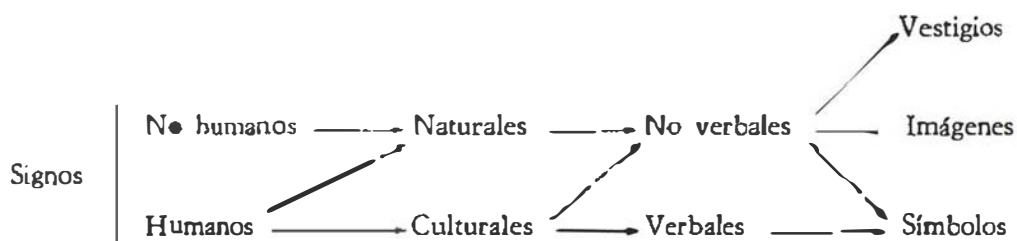
Por último, desde el punto de vista de la relación entre el signo y lo significado, basta reproducir la división de Peirce que ya hemos visto, pero cambiaré la terminología levemente, sustituyendo algunos de los términos de Peirce, un tanto técnicos y esotéricos, por términos que, además de ser más castellanos, tienen más tradición, pues traducen términos usados hace seis siglos en latín por Occam (véase el apéndice a este capítulo). Primero, lo que llamaré vestigios, aquellos signos que hacen algo presente en virtud de haber sido causalmente afectados por ello. Así, el humo es un vestigio del fuego, el olor, un vestigio del organismo que lo produjo, etc. Segundo, las imágenes, o signos que evocan algo por asemejarse a ello en alguna medida, como las fotografías, pinturas realistas, etc. Finalmente, los símbolos, signos cuya relación con lo significado es arbitraria; esto quiere decir, exclusivamente, que entre el signo y el objeto significado no hay relación alguna ni de causalidad ni de semejanza. Tomado el símbolo con esta generalidad, hay que reconocer que los símbolos no son exclusivos del mundo de la cultura ni de los seres humanos, aunque ciertamente adquieran entre éstos una complejidad y una importancia que no tienen en la naturaleza. Pero existen en ésta, pues no puede negarse que, por ejemplo, la danza circular por medio de la cual las abejas indican que existe alimento en las inmediaciones de la colmena, es un signo arbitrario. Hay que añadir que en el proceso comunicativo de las abejas hay también elementos no simbólicos. Así, la danza del coleteo, que resulta orientada de acuerdo con la dirección en la que se halla, respecto a la colmena y respecto a la posición del sol, la fuente de alimento, parece un signo fundamentalmente

icónico. El olor del que se ha impregnado la abeja, y el cual delata el tipo de flores al que se refiere la comunicación, es en cambio un signo de tipo indécico (véase von Frisch, *Vida de las abejas*, cap. 11). Las formas de comunicación entre los mamíferos son aún más claramente simbólicas. Así, no es otra que simbólica la relación existente entre el gruñido por el que el chimpancé comunica el descubrimiento de comida y dicha comida, como simbólicos son los tres tipos distintos de grito de alarma con los cuales ciertos monos africanos comunican respectivamente la proximidad de un águila, de un leopardo y de una cobra; para cada uno de estos animales tienen un tipo particular de grito de alarma, pero naturalmente no parece haber relación alguna ni icónica ni indécica entre cada especie de grito y cada tipo de animal (Marler, «Animal Communication», p. 34). Por lo mismo son también símbolos las palabras, si se exceptúan casos muy marginales como las onomatopeyas, que serían, al menos en parte, imágenes, puesto que se refieren a un objeto en la medida en que se asemejan al ruido que produce ese objeto. Con mayor razón serán símbolos los signos lógicos y matemáticos, y en general, aunque no siempre, los emblemas, banderas, etc. Por lo que hace a los vestigios y a las imágenes, es aún más claro que se den en la naturaleza tanto como en el contexto humano y cultural.

Lo que no aparece en estas clasificaciones son las señales. La razón es que, si lo característico de una señal es desencadenar una conducta, pararla o modificarla, esto no tiene nada que ver con la definición de signo que hemos aceptado, y si bien es posible dividir los signos entre aquellos que pueden funcionar como señales y aquellos que no, tal división es totalmente extrínseca al carácter de signo, esto es, a su función significativa y a su capacidad semiótica. La tendencia a dar una excesiva importancia al comportamiento del intérprete y, más concretamente, del receptor del signo, y, por tanto, la tendencia a hacer de los signos señales, es una consecuencia del conductismo que ha empapado buena parte de la semiótica contemporánea, particularmente la doctrina de Morris, y que ha extraviado a autores como Schaff, quienes, no sabiendo bien cómo deshacerse de ellas, han acabado por hacer de las señales un grupo aparte, sin advertir que no responden al criterio estrictamente semiótico propio de una clasificación así. En resumen, las divisiones que he propuesto son éstas:

Por el intérprete		Signos humanos Signos no humanos
Por el ámbito en el que se dan		Signos naturales Signos culturales
Por su estructura		Signos verbales Signos no verbales
Por su relación con lo significado		Vestigios Imágenes Símbolos

Reunidas en un solo cuadro, esas categorías muestran las siguientes relaciones:



2.4 El lenguaje

He hablado hasta ahora de signos, incluyendo los signos verbales, pero he evitado cuanto me ha sido posible hablar de lenguaje. Aunque el lenguaje verbal será, al menos en su aspecto más general y abstracto, tema de los capítulos restantes, procede, antes de cerrar éste, situar el concepto de lenguaje en el contexto de la teoría de los signos.

En términos genéricos, se habla de lenguaje siempre que hay *una pluralidad de signos de la misma naturaleza cuya función primaria es la comunicación entre organismos*. Así se habla del lenguaje de los animales (en rigor: del lenguaje de cada grupo o especie animal), del lenguaje del arte (o mejor: del lenguaje de la música, de la pintura, etc.), del lenguaje de los gestos, del lenguaje de las flores (de lo que convencionalmente expresa, en una sociedad determinada, cada clase de flores), del de los colores, etcétera. ¿Qué es lo que tienen en común todos estos supuestos lenguajes? No mucho. Simplemente que en todos ellos hay una pluralidad de significantes a los que se asigna de forma en cierto grado arbitraria una pluralidad de funciones significativas a efectos de una relación de comunicación entre sus intérpretes. Esto es, hay, de un lado, un sistema de significantes, como son las diferentes clases de flores, los distintos colores, la variedad de los gestos del cuerpo humano, determinados movimientos de la abeja sobre una superficie plana, las obras de arte y, naturalmente, la diversidad de los fonemas que integran una lengua. Y de otro, un sistema de funciones significativas asignadas, como la expresión de emociones, de prohibiciones, de permisos, de advertencias, de designaciones, de descripciones, etc. A esta asignación de funciones significativas al sistema de significantes lo llamaremos *código*. Hay que notar que, al hablar aquí de funciones significativas, no se toma el término «función» en su sentido lógico, puesto que muchos lenguajes, y particularmente el lenguaje verbal, tienen un alto grado de ambigüedad, de manera que la misma función significativa aplicada a un significante puede tener valores diversos. Por poner un ejemplo muy simple: el significante «gato», en cuanto signo verbal de la lengua castellana, es objeto de la asignación de una función designativa que tiene como valores, indistintamente, una clase de animales, una clase de mecanismos ele-

vadores, una clase de bolsos y alguna otra clase más de objetos, como puede comprobarse fácilmente con cualquier diccionario. Las funciones significativas, que resultan de la relación *ser el significado de* cuando el argumento es un significante determinado (nótese que digo «ser el significado de», y no «ser lo significado por»), son, por tanto, funciones de valores múltiples y no de valor único, y, en consecuencia, no son funciones en el estricto sentido lógico del término (cfr. Quine, *Lógica matemática*, secc. 40).

Usualmente (por ejemplo, Eco, en *Signo*, 3.5) se llama código a una supuesta asociación entre el sistema de significantes y el sistema de significados, pero esto es extremadamente impreciso por las siguientes razones. Primero, porque «significado», como ya he señalado, es ambiguo entre «el significado» y «lo significado». Y segundo, porque al hablar de los significantes y los significados como de dos conjuntos de elementos con los que ya se cuenta desde un principio y cuya asociación se establece por medio de un código, se dan por resueltos, o al menos por suficientemente claros, importantes problemas ontológicos y epistemológicos sobre esas extrañas entidades llamadas «significados», y de hecho se deja abierto el camino hacia un mentalismo ingenuo y acrítico, con toda la carga de arcaísmo que esto supone.

Es comprensible que el tipo de signos que, por su relación con lo significado, se presta a una mayor riqueza de funciones significativas y es más apto para constituir un lenguaje, sean los símbolos. La mayor parte de los lenguajes, desde el de las abejas al de los semáforos y el de la música, son simbólicos. Lo que no quiere decir que no haya otros tipos de lenguaje. Han existido, por ejemplo, lenguajes escritos de tipo pictográfico, y, por consiguiente, de carácter fundamentalmente icónico, puesto que cada signo representaba algo sobre la base de una relación de semejanza. Todavía hoy está en uso entre los habitantes de la provincia de Yunán (unos ciento cuarenta mil), al sur de China, un lenguaje, el *nakhi*, cuya representación escrita es pictográfica (se encontrará una representación del mismo en la página 217 de *The Languages of the World*, de Katzner). Por lo demás, resulta casi ocioso recordar que la mayor parte de los lenguajes son humanos y culturales, y que de todas las formas de lenguaje conocidas, parece ser el lenguaje verbal el que, especialmente en su forma escrita, ha adquirido mayor riqueza y complejidad de funciones significativas. En algunos de los pensadores que hemos considerado anteriormente, como Saussure y Schaff, hemos podido comprobar una clara resistencia a considerar los signos lingüísticos verbales como símbolos. Esto, naturalmente, depende de cómo se defina el símbolo. Tal y como aquí lo hemos definido, siguiendo a Peirce, ambas categorías no pueden contraponerse, sino que, por el contrario, los signos lingüísticos quedan categorizados como una subclase de símbolos. Es importante notar, empero, que ni en el uso del castellano ni en la etimología de esas palabras hay, como ya hemos visto, fundamento suficiente para contraponer signo lingüístico y símbolo, y, si se considera de algún interés recurrir a la tradición filosófica, recordaré también que en uno de los lugares clásicos de esta discusión, el comienzo de *Sobre la interpretación*,

Aristóteles se refiere a los signos vocales llamándolos tanto símbolos (σύμβολα) como signos (σημεῖα; 16 a, 4 y 6).

En el contexto de su teoría semiótica, a la que ya hemos hecho alguna alusión, y que a su vez se inscribe dentro de la metodología conductista, Morris ha propuesto una definición de lenguaje que podemos comparar brevemente con las consideraciones precedentes. Morris (*Signos, lenguaje y conducta*, cap. 2, secc. 2) propone cinco criterios para la definición semiótica de lenguaje. Primero, una pluralidad de signos, de los que se componga el lenguaje; segundo, que cada signo tenga un significado común a cierto número de intérpretes; tercero, que los signos puedan ser producidos por los intérpretes y que tengan el mismo significado para el productor y para el receptor; cuarto, que los signos sean plurisituacionales, esto es, que puedan utilizarse con el mismo significado en situaciones distintas; y quinto, que los signos estén interrelacionados entre sí, formando un sistema, el cual permita ciertas combinaciones entre ellos pero no otras.

El primero de estos criterios es trivial y no tiene ningún problema; de hecho está literalmente recogido en mi anterior descripción. Los criterios segundo y tercero son aspectos del mismo requisito: que los signos sean utilizados como vehículo de comunicación por una comunidad de intérpretes implica que tengan el mismo significado, al menos fundamentalmente, para todos ellos, pues de otro modo la comunicación no sería posible. Todo ello está implícito en mi anterior alusión a la comunicación como función primaria de los signos en un lenguaje. Y de aquí puede también deducirse, según pienso, el criterio cuarto de Morris, pues si los signos no fueran plurisituacionales, esto es, si cambiaran de significado en cada nueva situación, entonces propiamente no serían los mismos signos, serían, en todo caso, los mismos significantes constituyendo signos nuevos en cada situación, y por tanto un nuevo lenguaje; pero que sea posible para cualquier grupo de intérpretes constituir un nuevo lenguaje en cada nueva situación es algo que no parece compatible con las leyes naturales que rigen el desarrollo de los organismos y las relaciones entre ellos. En cuanto al último criterio, me parece demasiado fuerte exigir que los signos formen un sistema en el sentido de estar relacionados entre sí según reglas que permitan ciertas combinaciones pero no otras. Mientras que esto es cierto de ejemplos paradigmáticos de lenguaje como el verbal en todas sus formas, no me parece, sin embargo, aplicable a otros casos, ya mencionados, en los que asimismo se habla de lenguaje, como el del arte, el de los colores, el de los gestos, etc., pues aquí cada signo (obra de arte, color, gesto, etc.) no se combina con los demás propiamente. Este quinto criterio de Morris restringe extraordinariamente su definición del lenguaje haciéndole solamente aplicable a los lenguajes constituidos por signos de tipo verbal, y excluyendo de su aplicación ejemplos como los que acabo de citar. En definitiva, y resumiendo, creo que todos los criterios imprescindibles para que un conjunto de signos constituya un lenguaje, en el amplio sentido en el que cotidianamente se habla de lenguaje, que es justamente un sentido puro y vagamente semiótico, se reducen a dos: que esos signos sean de la

misma naturaleza y que sirvan primariamente a la comunicación de un grupo de organismos entre sí.

2.5 La semiótica

Para cerrar este capítulo, recordemos brevemente las tres partes principales que usualmente se distinguen en la semiótica o estudio de los signos. En la formulación más breve se dice que la pragmática considera las relaciones entre los signos y sus intérpretes o usuarios, la semántica se ocupa de las relaciones entre los signos y los objetos denotados por ellos, y la sintaxis estudia exclusivamente las relaciones de los signos entre sí (Morris, *Fundamentos de la teoría de los signos*, secc. 3; Carnap, *Introduction to Semantics*, secc. 4). En Carnap se presentan estos tres campos como tres niveles sucesivos de abstracción, pues se menciona la semántica como un estudio de los signos en el que se realiza abstracción de sus usuarios, y la sintaxis como un estudio en el que se hace abstracción «también de los objetos denotados» (p. 9). Posteriormente, Morris, en su obra principal (*Signos, lenguaje y conducta*, cap. 8, secc. 1) ha definido de forma más elaborada y explícita estos conceptos, de la manera siguiente: la pragmática trataría del «origen, usos y efectos de los signos dentro de la conducta en que se hacen presentes»; la semántica estudiaría «la significación de los signos en todos los modos de significar»; y la sintaxis se ocuparía de «las combinaciones entre signos, sin atender a sus significaciones específicas o a sus relaciones dentro de la conducta en que aparecen». Frente a Carnap, que establece estas distinciones con referencia a un lenguaje, Morris, más genuinamente preocupado por una teoría general de los signos, se esfuerza en formular sus definiciones en términos de signos, y sin prejuizar que éstos hayan de constituir un lenguaje. Esto, empero, no es sino una diferencia de detalle cuya única consecuencia interesante que se me alcanza es la de que, según la concepción de Morris, que en esto me parece la correcta, el estudio sintáctico podría ser inexistente cuando se trate de signos únicos o de un conjunto de signos los cuales no sean susceptibles de combinarse entre sí (por ejemplo, las obras de arte). Más importante es la mayor elaboración de las definiciones de Morris por lo que respecta a la semántica y a la pragmática. Pues, en efecto, el estudio de las relaciones entre los signos y la realidad (no necesariamente material) que los signos representan no puede reducirse a relaciones de denotación o referencia entre los signos y sus objetos (aunque éstos no hayan de ser necesariamente materiales). El problema, sin embargo, es que si la semántica da entrada a todos los modos de significar, puede no ser fácil distinguir algunos de éstos de ciertos usos y efectos que los signos tengan para los intérpretes, con lo que la delimitación entre semántica y pragmática resulta cuestionable. Como este orden de problemas habremos de tratarlo en detalle más adelante, especialmente por lo que toca al lenguaje verbal, no lo proseguiremos aquí. Baste dejar indicado que la distinción entre las tres partes

clásicas de la semiótica no es cosa tan clara como su aparente simplicidad podría hacer pensar.

Tales partes son, naturalmente, más que tres subdisciplinas o campos de estudio, tres enfoques del mismo tema. Cómo se conciban estos tres enfoques es relevante para los problemas que puedan plantearse en cada caso. No es lo mismo pensar que la pragmática constituye un estudio, por así decirlo, completo del lenguaje con respecto al cual la semántica y la sintaxis representarían abstracciones sucesivas, que pensar que los tres enfoques constituyen tres perspectivas independientes entre sí e igualmente abstractas que conjuntamente proporcionan una visión completa de un sistema de signos. Da la impresión de que esta última sería la concepción de Morris, con los problemas que acabo de indicar, mientras que Carnap sería partidario de lo primero. La cuestión es que no parece fácil expurgar de la semántica toda referencia a los intérpretes de los signos si la semántica se concibe con un mínimo de amplitud, como plantea igualmente gran cantidad de problemas el intento de estudiar las relaciones de los signos entre sí al margen de sus funciones significativas, al menos para los lenguajes verbales que, por oposición a los lenguajes formalizados o técnicos, se denominan naturales (como son las lenguas; aunque aquí, desde el punto de vista de la semiótica, hemos considerado a todos los lenguajes verbales como compuestos de signos culturales). Pero como todos estos son problemas que habrán de surgir con mayor relieve en momentos ulteriores de esta investigación, es preferible no entrar ahora en mayores detalles.

APENDICE

La teoría de los signos de Occam

El tema que hemos discutido en este capítulo tiene un precedente ilustre y clarificador en Occam.

En el primer capítulo de su *Summa Logicae*, dedicado a la definición y clasificación de los términos, Occam distingue dos sentidos del término *signum*. En su acepción más general, llama signo a «todo aquello que, cuando es aprehendido, hace conocer otra cosa, aun cuando no se trate de un primer conocimiento, sino que requiera un previo conocimiento habitual». En este sentido, la palabra hablada significa naturalmente a la manera como el efecto significa la causa. En un sentido más restringido, empero, en el cual toma él el término, signo es «aquello que hace conocer algo y que está destinado a suponer por ello» (a sustituirlo), o sea, la palabra significativa, así como «lo que puede añadirse a lo anterior en la proposición» (Occam se refiere con esto a aquellas palabras que carecen de significado completo por sí solas), e igualmente «lo que se compone de lo anterior», o sea, la proposición misma. En este sentido restringido, la palabra hablada (*vox*) no es signo natural de nada.

Con esta distinción se relaciona estrechamente la división más general de los signos en Occam. De un lado están aquellos signos que no dan un

conocimiento primario de la cosa significada, y no lo dan porque exigen un conocimiento habitual de la relación significativa que liga al signo con el objeto. Estos signos son signos en la primera de las dos acepciones anteriores, pero no en la segunda. A estos signos Occam los llama representativos, y su manera de significar, puesto que requiere el conocimiento anterior de la cosa significada, se basa en el recuerdo y es, por consiguiente, una función rememorativa (*Comentario al libro primero de las Sentencias*, dis. 3, q. 9 y 10). De otro lado, están los signos lingüísticos, que son signos en la segunda de las acepciones citadas, pues no requieren conocimiento previo de lo significado, y dan, por tanto, un conocimiento primero del objeto. Su característica principal es que tienen como función peculiar sustituir al objeto significado (suponer por él) en el contexto de la proposición.

Los signos rememorativos o representativos a su vez son divididos por Occam en dos tipos, el *vestigium* y la *imago*. La imagen es signo de algo en cuanto que guarda con ello una relación de semejanza, y corresponde, por consiguiente, al icono de Peirce. El vestigio lo es en cuanto que es efecto del objeto significado, de manera que corresponde al índice de Peirce. Dentro de los signos lingüísticos, que Occam llama igualmente términos, distingue entre el signo hablado, el signo escrito y el signo mental. El signo hablado o *terminus prolatus* es la parte de una proposición que ha sido pronunciada con la boca y está destinada a ser oída. El signo escrito o *terminus scriptus* es la parte de una proposición inscrita en algún material y la cual puede ser vista. Finalmente, el signo mental o conceptual, *terminus conceptus*, es «una intención o pasión del alma que significa o consigna algo de manera natural, y está destinado a ser parte de una proposición mental y a suponer por eso que significa» (*Summa Logicae*, c. 1, 13-21). La diferencia fundamental la establece Occam entre el signo lingüístico conceptual y los otros dos tipos de signo lingüístico, el escrito y el hablado, pues mientras que el primero «significa naturalmente aquello que significa», en cambio el signo escrito o hablado «no significa nada a no ser por voluntaria institución» (*loc. cit.*, 47-49), y esto tiene como consecuencia que estos dos tipos de signo puedan cambiar su significado *ad placitum*, lo que no ocurre en el caso del signo conceptual.

Tenemos aquí, pues, un conjunto de signos lingüísticos que, a diferencia de los sonidos y de sus representaciones gráficas, tienen significado natural. Estos términos y proposiciones conceptuales son aquellas palabras mentales (*verba mentalia*) de las que, como recuerda Occam, decía San Agustín que no pertenecen a ninguna lengua, porque «se encuentran sólo en la mente y no pueden ser proferidas al exterior», aunque sí lo sean los sonidos a ellas subordinados. ¿En qué consiste esta manera natural de significar? Este es el problema más grave para la teoría occamista del lenguaje, puesto que la significación convencional, llamémosla así, está subordinada a la significación natural o, lo que es lo mismo, los signos hablados y escritos lo están a los conceptos o signos mentales. Cuando en otro lugar de la *Summa Logicae* (cap. 14, 53 ss.), Occam habla del universal, distin-

que entre el universal natural y el universal de institución voluntaria. El primero es aquel signo que, de manera natural, se predica de muchas cosas, de modo parecido a como el humo significa naturalmente el fuego, o el gemido significa el dolor del enfermo, o la risa significa la alegría interior. Estos son tres ejemplos de significación natural. Ahora bien, como se habrá observado, los tres son casos en los que el fundamento de la significación es una relación de causalidad, y, por tanto, se trata de signos del tipo de lo que Occam ha llamado vestigios, y Peirce, índices. El humo es signo del fuego en cuanto que es su efecto natural; e igualmente por lo que respecta a la risa con relación a la alegría y al gemido en relación con el dolor. Puesto que los conceptos o términos mentales significan de manera natural, parece obligado concluir que habrá que buscar alguna relación causal que explique el carácter de signos naturales que aquéllos tienen. Esta es la posición de dos calificados estudiosos de Occam, como Hochstetter y Boehner (citados en Teodoro de Andrés, *El nominalismo de Guillermo de Ockham como filosofía del lenguaje*, p. 97). La relación de causalidad es el fundamento del significado natural tanto en el caso de los términos mentales como en otros casos tan dispares como los citados. Boehner concretamente ha señalado que el concepto sería en parte efecto del objeto significado y en parte efecto del propio entendimiento, a lo que el propio Boehner añade el siguiente elemento de semejanza: el concepto se asemeja al entendimiento por ser inmaterial, y al objeto por ser imitación del mismo. Nótese que este último punto es particularmente oscuro, pues mientras no se nos explique de qué manera imita el concepto al objeto (lo cual es todo menos evidente) no sabremos hasta dónde llega la semejanza entre ellos. La imposibilidad de explicar este punto es reconocida explícitamente por Boehner, quien afirma que «especificar esta semejanza ulteriormente parece imposible, pues nos encontramos ante un hecho último de la psicología cognitiva» (en Teodoro de Andrés, *loc. cit.*).

En resumen, la clasificación de los signos y de los modos de significar en Occam queda así:

Signos	Rememorativos	{ Vestigios Imágenes Mentales Orales Escritos }	Natural	Modo de significar
	Lingüísticos		Convencional	

Si la manera natural de significar es propia de los vestigios y de las imágenes, resulta difícil, a falta de otras razones en contra, no buscar en los signos mentales aquellas relaciones causales e icónicas que son respectivamente características de unos y de otras. Y es lo que hace Boehner. Esto tiene una inmediata consecuencia que parece rechazable: aproxima los signos lingüísticos mentales a los signos no lingüísticos de tal manera que la distinción, en apariencia importante, entre signos rememorativos y signos lingüísticos tiende a difuminarse. Pues si los conceptos o signos mentales

significan a la manera como significan los signos no lingüísticos, ¿en qué medida puede seguir manteniéndose la distinción entre éstos y los lingüísticos en general? Esta es la razón que lleva a Teodoro de Andrés a rechazar la interpretación de Hochstetter y Boehner, y, sin negar la función que la relación de causalidad tiene en la teoría del conocimiento occamista, a proponer la siguiente interpretación alternativa.

Apoyándose especialmente sobre los ejemplos de la risa y del gemido, Teodoro de Andrés sugiere (*ob. cit.*, p. 98 y ss.) que aquí lo fundamental no es la relación de causalidad, aunque exista, ya que no agota el fundamento de la significación, sino que ésta se basa en lo que él llama «una especie de preordenación estructural, fundada a su vez en la estructura misma psico-somática del hombre» (p. 99). Es decir, que el quejido significa dolor y la risa significa alegría, no sólo porque sean efecto lo uno de lo otro, sino además, y fundamentalmente, porque la risa y el gemido están estructuralmente preordenados a ser expresión, respectivamente, de la alegría y del dolor. De igual manera, el signo mental descansará en una «preordenación estructural del hombre a abordar significativo-lingüísticamente la realidad exterior» (p. 101), y constituirá una reacción espontánea del entendimiento frente a la realidad.

En mi opinión, estamos aquí ante una solución puramente verbal. Para reforzar la caracterización independiente del signo lingüístico, que Occam ha puesto en peligro al reconocerle, en su manifestación mental, una forma natural de significación, se ofrece simplemente una nueva expresión: «preordenación estructural». Pero ¿qué explica esta expresión? Todo depende de lo que signifique, y esto es lo que no se nos ha explicitado. ¿En qué consiste esta preordenación estructural? Aparentemente en que, dada la constitución del ser humano, ciertas experiencias o estados internos como la alegría o el dolor son causa de ciertos efectos como la risa o el gemido. Si hay algo más que causalidad en esa preordenación estructural, como pretende Teodoro de Andrés, no lo sabemos ni se ve en qué pueda consistir. (Y naturalmente estoy empleando los términos «causalidad» y «causa» con una liberalidad que no me permitiría si estuviera desarrollando mi propia reflexión; pero intento mantenerme dentro de los límites del pensamiento y de la época de Occam.) Parece que cualquier ejemplo de una relación causal podría también justificarse aludiendo a una supuesta preordenación estructural. ¿Pues no es una cierta preordenación estructural la que les permite a los cuerpos combustibles, en ciertas condiciones, producir humo? Decir que el hombre está estructuralmente preordenado a abordar significativo-lingüísticamente la realidad exterior y que el signo es la reacción espontánea del entendimiento frente al objeto, no es distinto, en principio, y a menos que se explique debidamente, de decir que la concurrencia del objeto y del entendimiento producen como efecto conjunto el concepto o signo mental. Que es lo que claramente dice Occam, según muestra la interpretación de Boehner. La interpretación de Teodoro de Andrés tiene el plausible propósito de suministrar una justificación propia para la naturalidad de los signos lingüísticos, pero me temo que vale lo

mismo para cualquier otro tipo de signos, y por tanto no prueba nada. ¿Por qué negar que el humo exige una preordenación estructural en el combustible o que el gemido del animal la supone en su organismo?

Rechazada así esa interpretación, ¿qué nos queda para distinguir los signos lingüísticos de los rememorativos? Dos características sin duda muy fundamentales para la teoría de Occam y muy propias del lenguaje. En primer lugar, el hecho de que, dicho en términos modernos, el signo lingüístico se articula en un conjunto complejo que es la oración; dicho en términos de Occam: que el signo lingüístico tiene como función suponer dentro de la proposición (y en el caso de los signos lingüísticos mentales, la proposición será asimismo una proposición mental). En segundo lugar, el hecho de que los signos lingüísticos no aportan simplemente un recuerdo de lo significado, porque no requieren un conocimiento previo de esto, sino que tales signos dan un conocimiento primario del objeto. Esto basta para distinguir radicalmente a los signos mentales de los signos rememorativos, sin que sea necesario recurrir a la manera de significar. Esta sólo puede servir para distinguir a los signos orales y escritos, los cuales significan convencionalmente, de los signos rememorativos, que significan naturalmente. Los signos mentales, por significar también naturalmente, se aproximan en este aspecto a los signos rememorativos y se diferencian de los demás signos lingüísticos. Late aquí una tensión, que va a sobrevivir en las clasificaciones contemporáneas de los signos (como ya hemos visto), entre el carácter natural que parece tener el lenguaje humano, y en lo cual no se distinguiría del conocimiento en general, y el carácter convencional, e incluso arbitrario, que parece tener la significación en las diferentes lenguas. Habría así, para Occam, un lenguaje natural y, por tanto, único para la especie humana, que sería el lenguaje mental (*mentalia verba*), y una pluralidad de lenguajes convencionales constituidos por sonidos y por las representaciones gráficas de éstos. Chomsky, que tanto empeño ha puesto en rastrear precedentes de su concepción mentalista del lenguaje en la historia de la filosofía, tiene aquí un filón más rico probablemente que los que él, con tan poca fortuna, ha intentado beneficiar.

Queda ahora la cuestión de cómo se relacionan ambos lenguajes. La posición de Occam es que a todo término mental corresponde uno hablado y escrito, pero no viceversa (*Summa Logicae*, cap. 3). Así, a los términos orales o escritos que son sinónimos entre sí solamente corresponde un término mental, mientras que para cada término oral o escrito que sea ambiguo habrá tantos términos mentales como significados o acepciones quepa distinguir en él. La razón es que la variedad de sinónimos no obedece a necesidades de significación (*non est propter necessitatem significationis inventa*) sino que sirve a la mayor riqueza del discurso o a algo de este tipo. Y por ello, Occam reconoce que ha de haber términos mentales que correspondan a los nombres, a los verbos, a los adverbios, a las conjunciones y a las preposiciones, y que han de tener igualmente su contrapartida en la proposición mental accidentes como el caso, el número, y, por lo que respecta al verbo, el modo, el tiempo, la persona y la voz, pues todo esto es

necesario para las diferencias de significado. Duda en cambio sobre los participios y los pronombres, puesto que, en cuanto a los primeros, no parece que la necesidad de la significación requiera una forma como el participio, la cual, junto con la forma adecuada del verbo «ser», corresponderá siempre a una forma determinada del verbo original; y en cuanto a los segundos, sin duda porque pueden siempre sustituirse por el nombre al que reemplacen. Y duda también respecto a los comparativos, aunque no sobre el género y la declinación, accidentes que claramente sólo afectan a los términos hablados y escritos, puesto que nombres que son sinónimos entre sí pueden tener géneros diversos o pertenecer a distintas declinaciones, y puesto que el valor de verdad de la proposición no cambia ni por el género ni por la declinación, aunque sí puede cambiar por una modificación del caso o del número.

Esta alusión al valor veritativo de la proposición (*loc. cit.*, 56 y ss.) ha sido tomada por Teodoro de Andrés como otro principio, el de la *veritas propositionis*, que completaría al principio de la *necessitas significationis* que, como vemos, rige para la relación entre el lenguaje mental y el lenguaje externo. No parece, sin embargo, que Occam le conceda tanta relevancia, por una parte, pues su referencia no excede de una alusión al paso. Por otra, a mi modo de ver, no es sino una manifestación del propio principio de la exactitud y economía de la significación, ya que, si lo que éste exige es que en la proposición mental haya todo aquello, y solamente aquello, que es necesario para la significación, es patente que, desde el punto de vista de ese ingrediente del significado que es la función referencial, en la proposición mental habrá de haber todo cuanto es relevante para determinar la verdad o falsedad de la proposición. De todas formas, una cierta cautela al interpretar la doctrina de Occam sobre estas cuestiones no está de más, pues bien pudiera ocurrir que al proyectar sobre sus afirmaciones categorías actuales estuviéramos dándole un rigor que nunca tuvo. La conveniencia de esta actitud viene recomendada por hechos tan llamativos como que Occam, cuando quiere contrastar una proposición verdadera con una falsa, da como ejemplo de la primera «El hombre es un animal», y como ejemplo de la segunda, «El hombre es los animales». Proposición esta última que, estando mal formada y siendo antisintáctica, no calificaríamos hoy de falsa propiamente, pues sólo de una oración bien formada, y además con sentido, podemos preguntarnos si es verdadera o falsa. No conviene, por ello, apresurarse incautamente a interpretar a los clásicos con categorías actuales. Los excesos de Chomsky con los racionalistas deberían servirnos de recordatorio.

La idea de un lenguaje mental en contraste con el lenguaje oral y escrito es ciertamente una idea antigua. Su discusión se remonta al *De Interpretatione* (I, 16) de Aristóteles, que a través de la interpretación de Boecio, originó la doctrina escolástica clásica de que el lenguaje oral y escrito significa inmediatamente el lenguaje mental y sólo a través de éste se refiere a la realidad exterior. Occam, siguiendo aquí una idea que puede encontrarse sugerida por Duns Escoto, alteró esa doctrina sosteniendo que tanto

el lenguaje mental como el lenguaje oral y escrito significan directa e inmediatamente las cosas, aunque, siendo el lenguaje mental un lenguaje natural por su modo de significación y el otro, en cambio, un lenguaje convencional, éste esté subordinado al anterior y sea secundario con respecto a él. El carácter exacto de esta subordinación y la estricta relación existente entre ambos tipos de lenguaje, es algo que, a pesar de los esfuerzos interpretativos realizados para aclararla (cfr. Teodoro de Andrés, pp. 144 y ss.), Occam deja en la mayor oscuridad.

Lo que está claro es que el lenguaje mental es una suerte de lenguaje universal, común a la especie humana, y semánticamente perfecto, pues contiene todo y nada más que aquello que es necesario para las necesidades de significación. De aquí que se aluda a este lenguaje como un lenguaje mejor construido que el oral-escrito, al cual este último puede reducirse con propósitos aclaratorios (Pinborg, *Logik und Semantik im Mittelalter*, p. 128). Esto constituye, sin duda, un claro precedente de las finalidades y procedimientos típicos de algunos filósofos analíticos contemporáneos, como, por ejemplo, el reconstruccionismo de Russell y el intento de encontrar la forma lógica del lenguaje. Por esto se ha comparado también la doctrina de Occam a los intentos contemporáneos de construir lenguajes ideales sobre la base de la lógica (Trentman, «Ockham on Mental»). Este paralelismo es inexacto si se entiende el lenguaje ideal como un lenguaje artificialmente construido para un propósito particular. Pero es exacto si el lenguaje ideal es, simplemente, aquella estructura básica común a todas las lenguas y, por tanto, universal. Es decir, veo a Occam próximo al Russell de la teoría de las descripciones, al Wittgenstein del *Tractatus*, o a Chomsky. Pero no al Russell que habla de un lenguaje lógicamente perfecto como diferente del lenguaje natural en «La filosofía del atomismo lógico», ni menos aún a autores como Carnap. Geach lo ha entendido correctamente y ha podido por ello acusar a Occam de transferir al lenguaje mental los rasgos gramaticales del latín, para intentar luego explicar la existencia de estos rasgos en latín por correspondencia con el lenguaje mental, incurriendo en evidente círculo vicioso (Geach, *Mental Acts*, cap. 23). La crítica me parece correcta con tal que se amplíe a otras lenguas además del latín, pues es patente que los rasgos atribuidos por Occam al lenguaje mental no son exclusivos del latín, ni ésta era la única lengua que Occam conocía bien. Lo que esto ilustra en todo caso, y Geach tiene razón al subrayarlo, son los peligros de investigar lo mental con el modelo del lenguaje o, lo que es lo mismo, hacer una teoría general del lenguaje sobre la base de la estructura de un cierto tipo de lenguas. Son los peligros que hemos visto acechar en los últimos años a la utilización chomskiana de los universales lingüísticos, y sobre los que Quine ha advertido repetidamente (cfr. mi libro, *La teoría de las ideas innatas en Chomsky*, secc. 5.3). Pero sobre todo esto ya habremos de tratar más extensamente en un momento ulterior de esta investigación.

Lecturas

Una obra breve que resume, aunque con falta de claridad notable muchas veces, la mayor parte de lo que puede decirse sobre el signo, es *Signo*, de Umberto Eco (Labor, Temas de Filosofía, Barcelona, 1976). Una ampliación de esta obra, más elaborada, y con adición de temas nuevos, es el *Tratado de semiótica general* del mismo autor (Lumen, Barcelona, 1977), pero no se encontrarán en ella precisiones mucho mayores sobre el concepto de signo. Otra síntesis interesante, pero a mi juicio menos conseguida que la de Eco, es la *Teoría de los signos* de Bertil Malmberg (Siglo XXI, México, 1977).

La obra clásica de la semiótica contemporánea es *Signos, lenguaje y conducta*, de Morris (Losada, Buenos Aires, 1962; la edición original es de 1946), y es lástima que obra tan ambiciosa se encerrara en moldes conductistas tan estrechos. Esta obra había sido preludiada por un escrito más breve de Morris, *Fundamentos de la teoría de los signos* (Universidad Nacional de México, 1958), cuyos dos primeros apartados se encontrarán en la recopilación de Gracia, *Presentación del lenguaje* (Taurus, Madrid, 1972).

Sobre la teoría de Occam puede verse con provecho *El nominalismo de Guillermo de Ockham como filosofía del lenguaje*, de Teodoro de Andrés (Gredos, Madrid, 1969).

Capítulo 3

¿DE LO ABSTRACTO A LO CONCRETO O DE LO CONCRETO A LO ABSTRACTO?

El lenguaje comienza allí donde la comunicación está en peligro. (Henry MILLER, *Sexus*.)

3.1 Lenguaje, lengua y habla

Situado el tema del lenguaje en las coordenadas semióticas tal y como se ha hecho en el capítulo anterior, pasamos ya al lenguaje verbal, característica fundamental de la especie *homo sapiens* y sistema simbólico que es, sin duda, el más poderoso de cuantos se conocen y el que hace posible la tradición, la historia y la cultura. En lo sucesivo nos limitaremos al estudio de este tipo de lenguaje, y sólo accidentalmente se hará referencia a los lenguajes animales o a los lenguajes humanos de carácter no verbal. El lenguaje verbal, como ya se ha insinuado, es primariamente oral, consiste en sonidos producidos por medio de los órganos fonadores; derivadamente, es también escrito, en cuanto los sonidos, y ciertas características de éstos, como la intensidad, las pausas entre ellos, etc., son representados por medio de marcas visibles sobre algún material. No todos los lenguajes escritos son, desde luego, de este tipo. No hay que olvidar que hay lenguajes escritos de tipo pictográfico, como el nakhí ya citado o el chino primitivo, cuyos caracteres representan tipos de objetos o situaciones sobre la base de una relación fundamentalmente icónica, como los hay de tipo ideográfico, y tal es el chino actual, cuyos caracteres representan contenidos significativos o ideas, pero sin correspondencia con un sistema fonológico determinado. Por cierto, esto es lo que permite que los hablantes de distintos dialectos chinos cuyas diferencias fonéticas les impiden entenderse entre sí, puedan, en cambio, entenderse por medio del lenguaje escrito. Como curiosidad adicional puede notarse que el japonés escrito incluye tanto caracteres silábicos como caracteres ideográficos (que originariamente fueron tomados del chino y, por tanto, eran pictográficos). Es patente que los lenguajes

escritos de tipo ideográfico, y *a fortiori* los de tipo pictográfico, constituyen recursos tan especiales que se desvían en aspectos importantes de las características propias del lenguaje verbal, más aún que los lenguajes escritos de tipo alfabético. Por ello, cuanto digamos, en general, sobre el lenguaje humano irá básicamente referido al lenguaje oral.

Vale la pena recordar aquí que el lenguaje hablado no es la única realización sonora del lenguaje humano. Hay también el lenguaje silbado. El lenguaje silbado ha sido caracterizado como «una versión acústicamente simplificada del lenguaje hablado», del cual conserva el vocabulario y la sintaxis, y en una gran medida la fonología (Busnel y Classe, *Whistled Languages*, p. 108). Los lenguajes silbados constituyen realizaciones fonéticas del lenguaje verbal más simples que el lenguaje hablado, respecto al cual resultan complementarios. Por ello, los lenguajes silbados parecen haberse desarrollado preferentemente en lugares de geografía quebrada y entre pastores, que encuentran en ellos un medio más apto para la comunicación a distancias medias (hasta un par de kilómetros). Aunque se ha señalado la existencia de unos treinta lenguajes silbados distintos, los que mejor se conocen son el mazateco, que se emplea en una región de México por los indios del mismo nombre; el que se utiliza en Aas, pueblo del Pirineo francés; el de la región de Kusköy, en Turquía, y el llamado silbo de la isla Gomera, que antaño estaba extendido por todas las Canarias. Naturalmente, lo que se silba en cada lugar es el lenguaje de la zona, quiero decir, entre los mazatecos la lengua mazateca, en Aas el dialecto francés local (que tiene influencias del español), en Kusköy el turco y en la Gomera el castellano. Las diferencias más dignas de mención son que el mazateco silbado se basa en el tono, es un lenguaje tonal, mientras que los otros tres ejemplos citados se basan en la articulación, son lenguajes articulados. Según parece, el grado de inteligibilidad de estos lenguajes para sus usuarios es normal, teniendo en cuenta el papel que juega el contexto extralingüístico en cada episodio de comunicación, pues es, por otro lado, patente que, debido a su simplicidad fonética, estos lenguajes tienen un grado de ambigüedad mucho más alto que el que puede encontrarse en los correspondientes lenguajes hablados (véase la obra citada de Busnel y Classe y *El silbo gomero*, de Ramón Trujillo).

Hay que mencionar, finalmente, que se ha hablado también de un lenguaje mental; ya hemos tenido contacto con este concepto al discutir la teoría de Occam, y se trata indudablemente de una categoría que, bajo un nombre u otro, ha tenido vitalidad bastante para estar presente en la filosofía desde Aristóteles, por lo menos, hasta nuestros días. Pero trátase también de una categoría tan polémica y espinosa que no me parecería aceptable dar por supuesto su utilidad ni su capacidad explicativa, de manera que nada de lo que en lo sucesivo se diga debe entenderse referido al lenguaje mental a menos que este concepto haya sido explícitamente introducido. De hecho, una buena porción de los problemas que habremos de discutir giran en torno a esta supuesta forma de lenguaje, cuya función es

primordial para el problema de las relaciones entre el pensamiento y el lenguaje.

Hablando ya, por tanto, del lenguaje verbal, y principalmente de su forma oral, lo primero que hay que decir es que tal lenguaje puede entenderse en principio como una facultad biológica y psicológica que poseen única y exclusivamente los individuos de la especie humana, y que por ello es característica de ésta. Se trata, pues, de algo totalmente natural. El lenguaje, en cuanto facultad, pertenece a la estructura biopsíquica del ser humano, y es producto, en definitiva, de determinados caracteres alcanzados en el curso de la evolución, entre los que hay que subrayar la complejidad del cerebro y la peculiar conformación de los órganos fonadores. Todo lo referente a la facultad lingüística es tema de la biología y de la psicología del lenguaje (se dirá algo sobre estos temas en el capítulo quinto).

Pero esta facultad, con ser común a todos los hombres y característica de su naturaleza, ha dado lugar a una considerable variedad de sistemas verbales o lenguas, que, aunque podrían estar todas ellas últimamente relacionadas entre sí lógicamente e históricamente, constituyen maneras muy diferentes de realizarse esa facultad, las cuales son convencionales en el sentido de que están determinadas no por la naturaleza, sino por la cultura y por la historia. Esto es lo que Saussure llamó lengua (*langue*), a diferencia de lenguaje (*langage*) entendido como facultad. La lengua para él era «un producto social de la facultad del lenguaje y un conjunto de convenciones necesarias adoptadas por el cuerpo social para permitir el ejercicio de esa facultad en los individuos» (*Curso de lingüística general*, p. 51). Este producto social es un sistema de signos, un sistema gramatical, que está «virtualmente existente en cada cerebro, o, más exactamente, en los cerebros de un conjunto de individuos, pues la lengua no está completa en ninguno, no existe perfectamente más que en la masa» (*op. cit.*, p. 57). A diferencia de la lengua, el habla (*parole*) está constituida por el conjunto de actuaciones lingüísticas individuales en las que se actualizan esas convenciones que constituyen la lengua. Al distinguir lengua y habla, afirma Saussure, se distingue lo social de lo individual, lo esencial de lo accesorio (*loc. cit.*).

Bajo su aparente simplicidad, esta distinción encierra, sin embargo, grandes problemas. Que hay que hacer distinciones de este tipo y que el uso ordinario del término «lenguaje» tiende a ocultarlas, es cosa clara. Decimos que el lenguaje distingue al hombre del animal (facultad), hablamos del lenguaje castellano (lengua) y afirmamos de alguien que en cierta ocasión utilizó un lenguaje muy rebuscado (habla). Es patente la equivocidad de la palabra «lenguaje». Pero la forma en que Saussure explicó las diferencias entre estos aspectos de la realidad lingüística ha originado un sinnúmero de encontradas interpretaciones y ha suscitado críticas, algunas de ellas justificadas. En un importante trabajo sobre este tema, Coseriu ha pasado revista a una representativa muestra de tales interpretaciones señalando los puntos débiles de la dicotomía saussuriana que son responsables de la confusión engendrada («Sistema, norma y habla», recogido en *Teoría del lenguaje y lingüística general*). Según Coseriu, los tres aspectos fundamentales

del concepto de lengua corresponden a tres clases diferentes de oposición y constituirían, en realidad, más bien tres conceptos distintos. En primer lugar, la lengua como realidad psíquica, puesto que solamente existe en la mente (aunque Saussure hable más bien del cerebro), y en cuanto se opone al habla, que es psicofísica. En segundo lugar, la lengua en cuanto realidad social, puesto que sólo existe completa en la masa y se impone al individuo; se opone así a lo individual, que es el habla. Y en tercer lugar, la lengua como sistema funcional de posibilidades que se realiza en el habla, lo que, aparentemente, correspondería a la oposición entre lo abstracto (lengua) y lo concreto (habla). Aquí hay que recordar que para Saussure la lengua era algo tan concreto como el habla (p. 59 del *Curso*), pero en esto, sin duda, se equivocaba y con ello introducía un elemento de incoherencia que sólo podía enturbiar la perfecta comprensión de lo que es una gramática. Saussure afirma de los signos lingüísticos que «no por ser esencialmente psíquicos son abstracciones», y que el conjunto de asociaciones que constituye la lengua «son realidades que tienen su asiento en el cerebro», y que dichos signos son, por decirlo así, tangibles, puesto que «la escritura puede fijarlos en imágenes convencionales» (*loc. cit.*).

Pero es patente que nada de esto hace a la lengua concreta, pues si los signos son signos tipo, entonces la lengua es una abstracción por mucho que esos signos tipo puedan representarse por escrito, ya que cada representación no será otra cosa que un episodio de habla (es decir, de *parole*, aunque escrita); y si se trata de signos acontecimiento, entonces ya no estamos en el ámbito de la lengua, sino claramente en el del habla. Es posible que la ubicación de la lengua en el cerebro, sobre la que Saussure insiste, contribuya a crear esta ilusión de concreción, otro aspecto más de ese confuso mentalismo que ha convertido a Saussure en fácil pasto para ciertas tendencias especulativas de la filosofía francesa.

Por la oposición entre lo individual y lo social es justamente por donde antes se quiebra la construcción de Saussure. Lo social queda situado en el plano de lo funcional, y por tanto, y a pesar de la intención de Saussure, en el plano de lo abstracto, que es el plano de la lengua, mientras que lo único concreto es el fenómeno puramente individual, el habla, cuya articulación con lo social y con lo funcional, y en definitiva con la lengua, queda sumido en la más completa oscuridad. Falta un plano intermedio entre lengua y habla en el que los actos individuales de habla queden formalizados con algún grado de abstracción, a la vez que se muestre operativamente la conexión entre lo individual y lo social. Por esta brecha precisamente se ha precipitado también, y con razón, la crítica desde el campo marxista. Así, Ponzio (*Producción lingüística e ideología social*, pp. 187 ss.) ha subrayado la insuficiencia del concepto saussuriano de lo social buscando su origen en la perspectiva propia de la sociedad burguesa en la que acriticamente Saussure se sitúa y que le impide manejar las categorías dialécticas necesarias para el análisis de lo social. Al colocar lo social exclusivamente en el nivel de la lengua, lo social queda asumido como aquello ajeno al individuo, regido por leyes independientes y manifiesto esencial-

mente a través de la constricción impersonal. El habla, puesto que es producida por el individuo, no tiene, por lo mismo, nada de social, excepto el servirse de la lengua como instrumento. Y así, todo lo referente al acto de comunicación queda encerrado en el ámbito del individuo y desconectado de lo propiamente social. En ello ve Ponzio (p. 193) una muestra de la separación entre lo público y lo privado, característica de la sociedad burguesa. El habla, categorizada como fenómeno individual, oculta ideológicamente el hecho de que las relaciones de comunicación, como todas las relaciones entre individuos, son sociales y como tal constituyen manifestación de una estructura que, si el análisis dialéctico es correcto, es una estructura de dominación y por consiguiente es alienante. La posición de Saussure queda facilitada por el hecho de que la lengua, tal y como él la considera, y a pesar de sus declaraciones en contra, es una entidad abstracta.

3.2 Sistema y norma

Para resolver esa insuficiencia de la dicotomía saussuriana y paliar su rigidez, Coseriu ha propuesto desdoblar el concepto de lengua en dos conceptos distintos, pero íntimamente conectados, los cuales, a su parecer, ya están implícitos en el propio Saussure, y sin duda venían siendo de hecho utilizados en la lingüística estructural: los conceptos de sistema funcional y de sistema normal, abreviadamente, sistema y norma. Ambos son grados sucesivos de formalización y abstracción realizados sobre la realidad lingüística concreta que es el hablar.

El primer grado de abstracción es la norma, que contiene todas aquellas estructuras, sean fonológicas, morfológicas, sintácticas o semánticas, que, permitidas por la lengua, son tradicionales y caracterizan a la comunidad, a un subgrupo de la misma, o simplemente al individuo. Pues Coseriu admite tanto una norma propiamente social como una norma individual. Es, sin embargo, la primera la que más interés tiene y desempeña un papel más decisivo en la teoría lingüística, y la que el mismo Coseriu especialmente toma en consideración, como cuando afirma: «La norma es un sistema de realizaciones obligadas, de imposiciones sociales y culturales, y varía según la comunidad. Dentro de la misma comunidad lingüística nacional y dentro del mismo sistema funcional pueden comprobarse varias normas (lenguaje familiar, lenguaje popular, lengua literaria, lenguaje elevado, lenguaje vulgar, etc.), distintas, sobre todo, por lo que concierne al vocabulario, pero a menudo también en las formas gramaticales y en la pronunciación» («Sistema, norma y habla», p. 98 de *Teoría del lenguaje...*).

El sistema, por su parte, constituye un grado ulterior de abstracción en el estudio del lenguaje, y contiene todos aquellos elementos que son «esenciales e indispensables» en la lengua, esto es, el conjunto de «oposiciones funcionales» en que ésta consiste. Al pasar del habla a la norma se prescinde de todo aquello que es puramente individual, ocasional y momentá-

neo. Al pasar de la norma al sistema, se abandona todo cuanto es pura repetición y hábito individual (norma individual), así como todo lo que sea costumbre y tradición del grupo o subgrupo al que el individuo pertenece (norma social).

Distinguir qué aspectos pertenecen al sistema y cuáles a la norma puede ser más o menos fácil, según los casos. Parece claro que las diferencias de pronunciación que distinguen a unos grupos regionales de otros o a diversos estratos sociales entre sí son diferencias que pertenecen a la norma, como también lo son, sin duda, las preferencias características de un grupo social frente a otro a la hora de elegir entre términos sinónimos. También pertenecerían a la norma, según Coseriu, ciertas formaciones irregulares, como «anduve» por «andé», y ciertas construcciones, como «se me» frente a «me se», la cual es a veces tenida por incorrecta.

En todo caso, no es competencia nuestra entrar en los detalles de la aplicación de la distinción de Coseriu, ni tampoco se aprecia en ello interés filosófico. Lo único que ahora interesa es definir cuidadosamente una serie de categorías que son útiles para el análisis del fenómeno lingüístico y de las cuales hemos de hacer uso en el curso ulterior de esta investigación. Concretamente, la diferencia entre sistema y norma me parece importante para corregir las insuficiencias del concepto saussuriano de lengua, y por tanto para satisfacer críticas como las que ya hemos visto. Habiendo tenido, de otra parte, los conceptos de Saussure tanta influencia en la filosofía del lenguaje, y no sólo en la ciencia lingüística, la aportación de Coseriu me parece digna de puntual mención. Pienso que el concepto de norma representa un nivel del análisis lingüístico en el que es fácil plantear toda una serie de problemas de índole estrictamente social sobre los que no es posible alcanzar, en cambio, claridad con la mera distinción entre lengua y habla. Por ejemplo, y por lo que nos puede interesar más cercanamente en una perspectiva filosófica, todos los problemas que se refieren a las relaciones entre el lenguaje y las clases sociales, y que giran en torno al funcionamiento del lenguaje como vehículo de la ideología. Desde luego, no se encuentra en el trabajo de Coseriu ninguna sugerencia en este sentido, y no cabe atribuirle paternidad alguna en el alcance que propongo dar a su concepto de norma. Pero no se ve tampoco ninguna razón por la que no sea lícito dárselo. La lista de normas que como ejemplo menciona Coseriu invita a añadir nuevas categorías, como podría ser la de *lenguaje de clase*. Si cabe distinguir entre un lenguaje familiar, un lenguaje vulgar y un lenguaje literario, con mayor razón podría distinguirse un lenguaje de clase, esto es, una norma lingüística conectada al carácter de clase de un grupo social, y cuya manifestación en el habla exprese la pertenencia del individuo a esa clase. Tomando el término «clase» en su sentido marxista, la norma lingüística de clase sería fundamentalmente de dos tipos, según caracterizara a las clases dominantes o a las clases dominadas. La aspiración a imponer la norma lingüística de las clases dominantes como el lenguaje total de la sociedad equivaldría precisamente a utilizar el lenguaje como vehículo de la ideología. De esta manera, lo que es en principio neutral e igual-

mente apto para todas las clases sociales, el sistema de la lengua, resultaría ideológicamente limitado y deformado al utilizarse para la expresión de los intereses de las clases dominantes, dando lugar a un peculiar tipo de norma lingüística. Este tipo de norma tendría carácter semántico, especialmente, y se manifestaría tanto en la selección del léxico «acceptable» como en los sentidos y cargas semánticas asignados «normalmente» al vocabulario, y en particular a la porción del vocabulario más apta para la expresión de la ideología (por ejemplo, los términos de evaluación moral, los predicados abstractos expresivos de cualidades individuales, etc.). A este nivel puede, según estimo, plantearse con relativa facilidad toda una serie de problemas sobre el lenguaje que resultan extremadamente complicados dentro de una dicotomía como la de lengua y habla. Son el tipo de problemas que preocupan a un autor como Ponzio y que le hacen sentirse tan incómodo con la distinción de Saussure. Pero sobre todo ello habremos de volver con más detalle en un capítulo ulterior.

3.3 Competencia y actuación

Paralela a la distinción de Saussure entre lengua y habla, y parcialmente coincidente con ella, es la distinción entre competencia (*competence*) y actuación (*performance*) introducida por Chomsky como parte importante de su teoría del lenguaje (*Aspectos de la teoría de la sintaxis*, cap. 1, secc. 1). La actuación, en cuanto uso real del lenguaje en situaciones concretas, parece coincidir suficientemente con el habla en el sentido de Saussure. No así la noción de competencia, pues mientras que la lengua, como vimos, era para Saussure un sistema de signos y de convenciones regulativas de éstos, la competencia lingüística para Chomsky no es tanto el sistema, en cuanto inventario o repertorio inerte, sino la interiorización por el hablante de dicho sistema, entendiendo éste, además, fundamentalmente como un mecanismo generador de todas las posibles expresiones correctas de la lengua. Con esto aspira Chomsky a entroncar, por encima de Saussure, con la concepción de Humboldt, que entendía la lengua como fuerza creadora (*energeia*) más que como producto acabado (*érgon*).

La competencia es aquello que hace posible la actuación o comportamiento lingüístico, y que se manifiesta en éste. Esta manifestación, sin embargo, es imperfecta, pues en condiciones reales la competencia y la actuación no coinciden. Cualquier fragmento de habla mostrará, según Chomsky, titubeos, cortes e irregularidades características de todo tipo, fruto de las condiciones peculiares del sujeto, tales como el alcance de su memoria, falta de concentración, distracciones, errores, etc. ¿Cuándo coincidirán la competencia y la actuación? Solamente en el caso de un hablante que no se halle afectado por estas circunstancias, que pertenezca a una comunidad lingüística homogénea y que conozca perfectamente la lengua que habla (*op. cit.*, pp. 3-4). Pero individuo tal no existe en la realidad. Como afirma Chomsky, se trata de un hablante-oyente ideal, y éste es jus-

tamente el objeto de la teoría lingüística. ¿Por qué? Porque como sólo en este caso ideal coinciden la competencia y la actuación, sólo aquí aparecerá claramente el sistema de la lengua. Por eso añade Chomsky que «la gramática de una lengua aspira a ser la descripción de la competencia intrínseca del hablante oyente ideal». Por cierto que esto, literalmente, es inexacto, o por lo menos confundente, en la medida en que puede dar a entender que no es la competencia del hablante real la que le interesa al lingüista. En rigor, sin embargo, y si se atiende a lo que parece ser la intención de Chomsky, no hay diferencias entre la competencia de un hablante real y la de un hablante ideal. La diferencia estaría en la actuación. Mientras que la de un hablante real no correspondería fielmente a su competencia, la de ese supuesto hablante ideal, en cambio, sí correspondería. Y como lo único que tenemos accesible de forma directa en nuestra experiencia es el comportamiento de los hablantes reales, puesto que este comportamiento no manifiesta fielmente su competencia, hay que concluir que debemos dejar a un lado a éstos y ocuparnos únicamente de ese hablante ideal en el que no hay divergencias entre su competencia y su actuación.

La competencia, a diferencia de la actuación y en cuanto que la hace posible, es una realidad mental, y de aquí que la teoría lingüística, para ser propiamente científica, haya de ser mentalista. Tal mentalismo consiste simplemente en tomar el comportamiento lingüístico como dato para el estudio de la competencia, haciendo de ésta, y no de aquél, el objeto de la teoría del lenguaje. A reforzar esta actitud mentalista tiende la frecuente referencia de Chomsky a la competencia como un tipo de conocimiento, a saber, el conocimiento inconsciente que el hablante tiene de su lengua. La competencia lingüística consiste, según esto, en el conocimiento tácito que el hablante tiene de la gramática de su lengua. Esto es importante porque comporta ya una manera específica de categorizar lo que en principio sólo es una facultad o capacidad lingüística general. Esta facultad es un tipo de conocimiento, aunque sea tácito, y con todo lo que ello pueda implicar. El contenido de este conocimiento es la gramática de la lengua de que se trate, y conocerla tácitamente consiste, al parecer, en tenerla interiorizada. Por eso Chomsky pasa en muchos contextos de identificar la competencia con el conocimiento tácito de la gramática, a identificarla con la gramática misma. Así, cuando escribe que hay que considerar la competencia lingüística, el conocimiento de una lengua, como un sistema abstracto que subyace a la conducta, y que está constituido por reglas que conjuntamente determinan la forma y significado de un número potencialmente infinito de oraciones, y que tal sistema constituye una gramática generativa (*El lenguaje y el entendimiento*, p. 62 de la primera edición inglesa).

En todo esto hay, por lo pronto, dos aspectos fundamentales que distinguir cuidadosamente. De un lado, el haber separado de esa forma la competencia y la actuación lingüísticas reservando a la primera el carácter de tema de la teoría del lenguaje. De otra parte, el haber categorizado la competencia como conocimiento inconsciente del sistema gramatical de la

lengua. Las consecuencias principales de la teoría chomskiana para la psicología y la filosofía del lenguaje derivan de ambos aspectos y de las recíprocas conexiones entre ellos.

Es claro que, al separar así la competencia de la actuación subrayando el carácter mental de la primera, Chomsky ha psicologizado la teoría lingüística, dejando planteados importantes problemas para los propios psicólogos del lenguaje. Pero como los psicólogos (y me refiero, naturalmente, a los psicólogos científicos) venían ocupándose fundamentalmente del comportamiento humano, bien porque, a la manera de los conductistas extremados, hicieran de él el objeto único de la psicología, bien porque consideraran que el comportamiento es la única prueba empírica suficientemente fiable en psicología, la posición de Chomsky ha producido enorme perplejidad entre ellos. Así, por ejemplo, en una reciente *Introducción a la psicología del lenguaje*, el autor, Peter Herriot, menciona el carácter peculiar de la doctrina de Chomsky, y señala dos rasgos que distinguen el modelo chomskiano «de cualquier teoría psicológica reciente» (cap. 3, p. 73), a saber: en primer lugar, que no recurre a ninguna prueba relativa al comportamiento para justificar sus componentes y, en segundo término, que tal prueba no es lógicamente posible, porque si la actuación nunca responde fielmente a la competencia, entonces no es posible inferir ésta a partir de aquélla. De lo que concluye el autor que el modelo de Chomsky no es empíricamente verificable en principio y que la competencia, así separada de la actuación, no es accesible a la psicología científica.

Aunque son bastantes los que piensan de esta manera, no quiero decir que esa opinión sea unánime. Hay que recordar que algunos psicólogos y biólogos del lenguaje, como Miller y Lenneberg, han acogido con entusiasmo el planteamiento de Chomsky (véase, por ejemplo, *Nuevas direcciones en el estudio del lenguaje*, recopilación de Lenneberg, y la obra de éste *Fundamentos biológicos del lenguaje*). Pero, en mi opinión, la actitud de estos últimos se ha debido a la significación que la posición de Chomsky tuvo como parte de la crítica al conductismo radical de Skinner, y ha pasado por alto las deficiencias teóricas de la metodología chomskiana. La verdad es que reacciones como la de Herriot hay que atribuir las al hecho de que Chomsky ha presentado su teoría del lenguaje como una teoría de la competencia lingüística, concebida ésta como realidad mental, y ha dado así la impresión de que lo dado como objeto para la lingüística sería esa capacidad mental y no el comportamiento lingüístico. De hecho, sin embargo, la competencia no es, dentro del sistema de Chomsky, sino un concepto teórico presuntamente necesario para explicar no otra cosa que el comportamiento de los hablantes, y más concretamente el hecho de que cualquier hablante produce una variedad en principio ilimitada de nuevas oraciones, que establece entre ellas diferentes tipos de relaciones, etc. Justamente por ello hay que suponer que, en condiciones ideales, la competencia coincide con el comportamiento, a saber, en aquellas condiciones en las que se prescinde de la influencia que sobre la actuación ejercen los elementos extralingüísticos como la emotividad del hablante, el alcance de

su memoria, su grado de cultura, posibles afecciones de sus órganos fonadores, etc. Es decir, que el lingüista, como cualquier otro científico, abstrae de aquella porción de lo dado que le interesa lo que le resulta relevante, desentendiéndose de lo demás. No de otra manera que lo hace el físico cuando construye sus conceptos de movimiento y aceleración abstrayendo estas características de otras que, en los cuerpos reales, interfieren con ellas. Pero Chomsky, como es el caso de muchos otros científicos en diferentes disciplinas, posee ideas muy confusas sobre las características metodológicas y epistemológicas de su propia actividad. Por partir de su oposición a Skinner, que además no es un lingüista, sino un psicólogo, y por haber aceptado la contraposición entre mente y conducta (contraposición nefasta si la hay), Chomsky ha creído que la teoría lingüística tiene como objeto lo mental, aproximándola así de modo peligrosamente confundente a la psicología del lenguaje, hasta el punto de presentar a la primera como una «rama particular de la psicología cognitiva» (*El lenguaje y el entendimiento*, p. 1 de la primera edición inglesa).

El hecho es que Skinner, como psicólogo, ha estudiado el comportamiento lingüístico en su conocida obra *Verbal Behavior*, mientras que Chomsky, como lingüista, ha estudiado, y con gran agudeza y penetración, el sistema de la lengua, específicamente el de la lengua inglesa, y derivadamente los elementos que pudieran ser comunes a todas las lenguas, lo que se ha llamado gramática universal, y que consideraremos con más detenimiento en un momento ulterior. Ahora bien, el mero estudio del sistema lingüístico, en cuanto realidad abstraída a partir del comportamiento, no obliga a colocar el sistema en ningún sitio, y en particular no autoriza a situarlo precisamente en la mente humana, como si éste fuera un espacio de tipo peculiar, ni permite otorgarle una realidad propiamente mental. La decisión sobre este punto habrá que tomarla por razones internas a la psicología científica, y la mera consideración, por muy profunda que sea, de esa abstracción que es el sistema, no puede dar razones suficientes para esa decisión.

Chomsky, por su parte, ha intentado suministrar tales razones a través del requisito de adecuación explicativa, que consiste en lo siguiente. Una teoría lingüística, en cuanto hipótesis sobre la estructura del sistema de una lengua, es descriptivamente adecuada cuando sus descripciones corresponden a la intuición del hablante nativo para un conjunto significativo de casos cruciales (con independencia de que el hablante sea o no consciente de esa intuición); y es además explicativamente adecuada cuando consigue seleccionar una gramática descriptivamente adecuada sobre la base de los datos lingüísticos primarios, esto es, sobre la base de los datos que el niño encuentra en su experiencia (*Aspectos de la teoría de la sintaxis*, cap. 1, secc. 4, pp. 24-26 de la edición original). Así, pues, justificar explicativamente una teoría lingüística resulta muy parecido a justificar por qué el niño adquiere precisamente el sistema de la lengua a la que está expuesto, aunque no sea lo mismo. Chomsky, sin embargo, no tiene el menor inconveniente en dar el pequeño salto necesario para afirmar que el problema

de la adecuación explicativa «es esencialmente el problema de construir una teoría sobre la adquisición del lenguaje, una explicación de las capacidades innatas específicas que la hacen posible» (*op. cit.*, p. 27 de la edición original). Es decir, que no tendremos una teoría lingüística completamente justificada a menos que contemos con una teoría paralela sobre la adquisición del lenguaje, y además una teoría que recurra a capacidades innatas de tipo específico (lo que se refiere a este último aspecto lo veremos en el capítulo quinto). La psicologización, particularmente en un sentido mentalista, queda consumada de una manera evidentemente hábil, pero no por ello convincente. Pues al concebir así el tipo más fuerte de justificación para una teoría del lenguaje, se proyectan sobre la lingüística exigencias que son ajenas al mero estudio del sistema de la lengua y que son más bien propias de una ciencia como la psicología, que posee su metodología y su marco teórico específicos. En definitiva, una teoría sobre la adquisición del lenguaje no es más que una teoría sobre el desarrollo de ciertos aspectos del comportamiento del niño en relación con su medio, a lo que se puede añadir, naturalmente, una teoría sobre las facultades específicas innatas que lo hacen posible; por tanto, una parte de la teoría del aprendizaje, algo que es totalmente distinto de la construcción de una gramática. Diferencia que, por cierto, Chomsky tiende a oscurecer cuando afirma que el niño que está adquiriendo su lengua nativa está construyendo inconscientemente la gramática de esa lengua y, por tanto, inconscientemente formulando una teoría lingüística (*Aspectos*, cap. 1, secc. 4, p. 25 de la edición original).

Con esto no pretendo negar ni la íntima relación que la psicología del lenguaje tiene con la lingüística ni el hecho de que el mayor impulso para esa relación ha venido de la obra de Chomsky. Y ello es cierto hasta tal punto que el término «psicolingüística», que había comenzado a usarse para una psicología del lenguaje especialmente vinculada y atenta a los desarrollos de la lingüística, se ha extendido por efecto de la obra de Chomsky y se aplica especialmente a cuantos psicólogos están en alguna medida bajo su influencia (véase J. Greene, *Psycholinguistics*, introducción). Tampoco hay que enrender lo anterior como prejuzgando un juicio determinado sobre la crítica de Skinner por Chomsky. Aparte de su exceso de carga emotiva y de posibles malentendidos en cuestiones de detalle, la crítica tenía básicamente buen sentido y exteriorizaba una laudable sensibilidad hacia la unilateralidad y la desmesura del conductismo radical skinneriano. Pero rechazar las extrapolaciones doctrinales realizadas por Skinner a partir del método conductista no obliga a retornar a un mentalismo precientífico ni exime de emplear algún tipo de control empírico cuando se hacen afirmaciones sobre la realidad. No es ahora momento de entrar en los detalles de esta polémica, que constituye, en definitiva, una discusión sobre la metodología apropiada para el estudio psicológico del lenguaje, pero el lector debe recordar, a favor de Chomsky, que el conductismo radical de Skinner deja fuera de su consideración importantes porciones de nuestra experiencia e implica una concepción sumamente restringida y anticuada de lo que es la ciencia (véase, por ejemplo, el artículo de Angel Rivière «El

análisis experimental de la conducta y el conductismo radical como filosofía); y a favor de Skinner debe tener en cuenta que la crítica de Chomsky está realizada sobre la base de una extrapolación (como tal, injustificada) desde los resultados del análisis lingüístico formal a los requisitos de una teoría sobre la adquisición del lenguaje, para lo que no ofrece apoyo empírico alguno (puede verse sobre ello los artículos recogidos bajo el título *¿Chomsky o Skinner? La génesis del lenguaje y la introducción del copilador*, Ramón Bayés).

Sigamos con nuestro tema, que es el de los conceptos de competencia y actuación en cuanto categorías generales para el análisis del fenómeno lingüístico. Las ambigüedades y vacilaciones que hemos visto en Chomsky a la hora de caracterizar la competencia han originado una serie de propuestas en el sentido de distinguir varios tipos de competencia, cuando no han llevado pura y simplemente a rechazar la utilidad del concepto, así como la de su correlato, el de actuación, sobre la base de que ambos términos son hoy más un estorbo que una ayuda (Sampson, *The Form of Language*, páginas 203-204). Entre las propuestas mencionadas puede citarse la de Greene (*Psycholinguistics*, pp. 97-98), que distingue entre un sentido débil y otro fuerte de la competencia, que aproximadamente corresponden a la adecuación descriptiva y a la adecuación explicativa; la de Derwing, que cree encontrar en las obras de Chomsky tres conceptos diferentes de competencia, a saber: como modelo idealizado de la actuación, como componente central de tal modelo y como entidad abstracta independiente y aparte de la actuación (*Transformational Grammar as a Theory of Language Acquisition*, cap. 8), y la de Campbell y Wales («The Study of Language Acquisition»), que distinguen también, a la postre, tres conceptos de competencia, en la forma que vamos a ver con mayor detenimiento.

Campbell y Wales comienzan igualmente distinguiendo un sentido débil y un sentido fuerte de la competencia. En el primer sentido, la competencia es aquella capacidad que puede oponerse a la actuación y encontrarse sólo imperfectamente reflejada en ésta. Pero las diversas capacidades o competencias de los seres humanos pueden estar sujetas a ciertas limitaciones, algunas de las cuales pueden no ser específicas de la capacidad de que se trate. Por ejemplo (y son ejemplos de los autores), la capacidad para el cálculo aritmético está limitada por la cantidad de información que el cerebro humano puede manejar en un tiempo dado, pero esta limitación no es específica de la capacidad de cálculo, sino genérica para todas las capacidades mentales; o también, de manera semejante, nuestra capacidad para caminar está limitada por la cantidad de tiempo que podemos pasar sin comer o descansar, pero tal limitación es general para todas las capacidades físicas y no afecta solamente a la capacidad de andar. Por ello podemos excluir dichas limitaciones a la hora de dar una explicación teórica de estas capacidades, explicación que en esa medida resultará idealizada, y que nos proporcionará un sentido más fuerte de esas capacidades o competencias. De forma análoga podemos, al explicar la capacidad lingüística, prescindir de limitaciones no específicas, como el alcance de la memoria,

el carácter de las capacidades sensomotoras que afectan a la producción y recepción del habla, etc., llegando así a un concepto fuerte de la competencia lingüística.

Pues bien, Campbell y Wales mantienen que Chomsky, aunque aparentemente interesado en caracterizar la competencia en este sentido fuerte (es decir, la capacidad humana *específicamente* lingüística), de hecho está operando con una competencia más restringida, a saber, la que queda cuando a la competencia en sentido fuerte se le resta la capacidad para producir y entender expresiones no tanto gramaticales cuanto «apropiadas para el contexto en el que se emplean» (*op. cit.*, p. 247). El término «contexto» incluye aquí tanto el contexto propiamente lingüístico como la situación. Al prescindir del contexto, particularmente del contexto de situación, la competencia queda reducida a la capacidad para producir y comprender expresiones gramaticales, lo que constituye una reducción desafortunada del objeto de la psicología del lenguaje (o psicolingüística, que para este caso es lo mismo), reducción que de hecho realiza Chomsky a pesar de haber reparado alguna que otra vez en la importancia de esa adecuación entre el uso del lenguaje y la situación (véase *Aspectos*, p. 6 de la edición original, y *El lenguaje y el entendimiento*, p. 10, también de la edición inglesa). En resumen, tenemos, según Campbell y Wales, tres conceptos de competencia lingüística. Primero, entendida como simple capacidad que subyace al comportamiento lingüístico. Segundo, entendida como la capacidad anterior menos las limitaciones inespecíficas; su característica central sería la aptitud para emplear expresiones adecuadas al contexto. A esta capacidad la denominan competencia de comunicación. Y tercero, la capacidad anterior, exceptuando justamente esa característica, con lo que quedaría lo que ellos llaman una competencia gramatical, esto es, la capacidad para emplear expresiones bien formadas. De ellas es precisamente la capacidad de comunicación la que ofrece a la psicología del lenguaje un instrumento más útil para el estudio del comportamiento lingüístico.

Posteriormente, Chomsky, después de quejarse de que el término «competencia» es confundente porque sugiere «capacidad» (*ability*), ha recogido una distinción muy semejante a la anterior. De un lado, lo que llama «competencia gramatical», a saber, el estado cognoscitivo que abarca todos aquellos aspectos de la forma y del significado que pueden asignarse propiamente al subsistema de la mente humana encargado de relacionar entre sí las representaciones de la forma y del significado. Sigue llamando a este subsistema «facultad del lenguaje», aun reconociendo que esta denominación es un tanto confundente. Por otra parte, la «competencia pragmática», que es la que subyace a la capacidad de usar el conocimiento anterior, junto con el resto del sistema conceptual, para tales o cuales propósitos. Y asume que es posible, en principio, tener completa competencia gramatical y carecer de competencia pragmática. (*Rules and Representations*, capítulo 1).

El resumen de la prolija discusión anterior puede ser éste. Primero, que el concepto de competencia sólo tiene buen sentido como concepto

teórico para el estudio del comportamiento lingüístico. Y segundo, que, en esa medida, el concepto de competencia que interesa es el de competencia de comunicación, pues si nuestro tema, lo que encontramos como dado, es la conducta lingüística, no está justificado prescindir del contexto. Naturalmente, la competencia a la que se hace referencia es una competencia de comunicación específicamente lingüística en el sentido humano, esto es, una competencia de comunicación verbal, pues es claro que una competencia de comunicación genérica hay que presumirla al menos en todos los organismos vivos en la medida en que éstos son sujetos de conductas semiósicas, según se apuntó en el capítulo precedente.

En su caracterización de la competencia, Chomsky se ha quedado por debajo de sus pretensiones de radicalidad al desatender la actividad lingüística y sus circunstancias propias (cfr. Sánchez de Zavala, *Indagaciones praxiológicas*, introducción y capítulos 1 y 2). Pero tan pronto como se atiende a ella se repara en que la competencia exige una caracterización más compleja en el sentido apuntado. E incluso hay razones para distinguir en ésta dos aspectos que, además, no coinciden genéticamente, a saber: la capacidad de entender las expresiones verbales y la capacidad para producirlas; o como las ha llamado Sánchez de Zavala (*loc. cit.*), la cuasicompetencia de recepción y la cuasicompetencia de producción, de las cuales aparece antes en el desarrollo infantil la primera.

3.4 La creatividad del lenguaje

Chomsky ha afirmado que el uso del lenguaje se caracteriza por ser creativo, y que esta creatividad se manifiesta en los hechos siguientes (*El lenguaje y el entendimiento*, cap. 1, p. 10 de la edición inglesa). Primero, en que el uso del lenguaje es innovador, en cuanto que gran parte de las expresiones que pronunciamos son totalmente nuevas y no constituyen una repetición de las ya escuchadas con anterioridad. Chomsky añade, incluso, que ni siquiera son «semejantes en estructura» (*similar in pattern*) a éstas, al menos en un sentido útil de «semejante» y «estructura» (sentido que Chomsky no ilustra ni aclara). Este hecho sería, según piensa Chomsky, incompatible con la concepción conductista del uso del lenguaje como conjunto de hábitos, y con el recurso a la analogía para explicar la aparición de expresiones aparentemente nuevas. No hace falta decir que el interés de Chomsky por señalar esta supuesta creatividad de la conducta lingüística obedece a su polémica actitud frente al conductismo skinneriano.

Chomsky, sin embargo, parece oscilar en la formulación de su tesis, pues pasa de afirmar, como acabamos de ver, que mucho de lo que decimos es nuevo, o de afirmar que el número de estructuras o pautas (*patterns*) es mucho mayor que el número de segundos en una vida, a decir que es potencialmente infinito (*loc. cit.*; véase también el comienzo del cap. 1, secc. 3 de *Aspectos*). Pero hay aquí dos cuestiones distintas que resultan confun-

didas. Una es la afirmación de hecho, y por tanto comprobable por observación de la conducta lingüística de los seres humanos, de que normalmente utilizamos expresiones que no habíamos escuchado antes. Otra es la afirmación de que el número de expresiones que pueden formarse correctamente en un lenguaje es potencialmente infinito. Esta última es una afirmación que enuncia una característica de los sistemas lingüísticos más bien que del uso de tales sistemas, y para comprobarla basta recurrir a las reglas del sistema, sin que sea relevante ocuparse del comportamiento. De todas formas, las consecuencias de esta afirmación no tienen por qué ser especialmente interesantes. Téngase en cuenta que para que esa afirmación sea verdadera basta con que no haya límite para la longitud de una oración, de manera que siempre sea posible formar una oración más larga que otra dada (y, por consiguiente, aumentar el número de oraciones bien formadas). Que esto es posible en las lenguas que conocemos parece claro, si se tiene en cuenta que tales lenguas cuentan con recursos como las oraciones de relativo, la conjunción y todos los demás modos de unión de oraciones simples que pueden aplicarse reiterada e indefinidamente a cualquier expresión por larga o compleja que sea. Naturalmente, por encima de cierto orden de longitud o complejidad, una oración se torna en mayor o menor medida improbable, pero esto es ya otro asunto, si bien desde el punto de vista de un estudio empírico puede que sea el asunto que nos interese. Piénsese que, a estos efectos, lo que nos interesa no es tanto las oraciones gramaticalmente correctas, como tales, sino, de éstas, aquellas que empíricamente pueden ser entendidas, y cuya utilización tiene, por ello, un grado apreciable de probabilidad. Pues es patente que una oración de cierta complejidad o longitud, por ejemplo, con diez cláusulas de relativo en su interior, no es comprendida por un cerebro normal, y por tanto es altamente improbable que aparezca en ningún contexto. Es evidente que lo que interesa para un estudio del uso del lenguaje son aquellas oraciones que, además de ser gramaticalmente correctas, pueden ser entendidas. Siguiendo un uso común, podemos llamar a tales oraciones «aceptables». Aquellas oraciones que, aun siendo correctas en un sentido formal, gramatical, no son aceptables, no nos interesan para un estudio como el presente, en el que lo que cuenta es el uso del lenguaje, del cual se está predicando la creatividad. De hecho, cabe esperar que nos interesen más aquellas oraciones que, aun no siendo gramaticalmente correctas, son normalmente utilizadas y entendidas. En resumen, con vistas a subrayar la creatividad del uso del lenguaje, es relevante el hecho de que frecuentemente entendemos y pronunciamos oraciones que, literalmente al menos, no habíamos encontrado antes; pero no es relevante que el número de oraciones correctas sea infinito, pues sólo un limitado subconjunto de éstas pueden empíricamente ser comprendidas y son, por tanto, utilizadas (cfr. Sampson, *The Form of Language*, pp. 67-70, y Cooper, *Knowledge of Language*, cap. 6).

A la creatividad en el ambiguo sentido anterior, Chomsky añade un segundo sentido, aquel en el que puede afirmarse que el uso del lenguaje

está «libre del control de estímulos detectables, externos o internos» (*loc. cit.*, p. 11). Esta es, aún más que la anterior, si cabe, una tesis íntimamente ligada al ataque contra el conductismo y a la polémica con filósofos empiristas como Quine, a la que me referiré en un capítulo ulterior. Lo único que Chomsky quiere decir es, simplemente, que el uso del lenguaje no está unívocamente determinado por los estímulos, tanto externos como internos. Ahora bien, esto es igualmente cierto para otros comportamientos del ser humano, e incluso para ciertos tipos de conducta animal, especialmente en animales superiores, y Chomsky mismo reconoce que ni con éste ni con el primer sentido de la creatividad hemos ido más allá de los límites de la explicación mecánica. Por ello propone pasar a un tercer sentido, que es el que tiene la creatividad cuando consiste en que el uso del lenguaje resulta ser «coherente y apropiado para la situación» (*loc. cit.*). Chomsky añade que, aunque no podemos decir de manera clara y definida en qué consiste esto, no hay duda de que tales conceptos son significativos, puesto que podemos distinguir entre el uso normal del lenguaje y las expresiones erráticas de un demente o el producto de un ordenador con un componente de azar. El lector no puede dejar de pensar, en este punto, que es perfectamente posible suministrar mayores precisiones sobre esos conceptos, y que la cuasidogmática forma en que Chomsky se limita a formular el aspecto decisivo de la creatividad deja a este concepto sumido en una vaguedad total e injustificada.

Tenemos, en resumen, que el uso del lenguaje es creativo porque es coherente y apropiado con respecto a la situación en la que se da tal uso, y que esto diferencia al uso normal del lenguaje de las manifestaciones lingüísticas de un demente o de un ordenador. La cuestión es que otro tanto puede afirmarse de otras manifestaciones y comportamientos no lingüísticos, así del hombre como de diversos organismos animales. No sólo el uso normal del lenguaje, sino también el comportamiento normal de un ser humano, la utilización de sus miembros, la posición y movimiento de su cuerpo, los gestos, etc., son coherentes y apropiados para la situación en la que se encuentra. Esto, y no únicamente las manifestaciones lingüísticas, constituye el conjunto de los criterios para distinguir a la persona normal del demente o de la máquina (más claramente en el primer caso; lo de la máquina es distinto porque se trata de una entidad de tipo diferente al del hombre). Y esto tampoco es privativo del hombre. Dentro de sus propias condiciones de vida y de sus específicos recursos genéticos, cualquier organismo animal se caracteriza por un repertorio de formas de conducta que podemos considerar normales, puesto que son coherentes y apropiadas para la situación; no sería, en cambio, normal el comportamiento de ese organismo cuando no fuera coherente ni apropiado, por ejemplo, porque se hallara bajo los efectos de determinadas sustancias, estimulaciones eléctricas, etc., según ha podido comprobarse en experimentos con ratas, simios y otros animales. Por consiguiente, la creatividad en este tercer sentido no caracteriza particularmente al uso del lenguaje; más bien caracteriza en

general a todo comportamiento humano, e incluso animal, que estemos dispuestos a considerar normal. Puesto que tampoco en los dos primeros sentidos suministra un criterio suficiente para caracterizarlo, ¿qué queda del concepto de creatividad? Lo único digno de señalarse es la característica formal, que el lenguaje natural comparte con los lenguajes artificiales, de no haber límite para el número de oraciones correctas distintas que puede construirse aplicando las reglas del sistema. En este sentido, la creatividad del comportamiento lingüístico no es distinta de la creatividad que se manifiesta en la actividad del matemático cuando resuelve problemas de aritmética o en la actividad del lógico cuando deduce teoremas en su cálculo (cfr. Sampson, *The Form of Language*, p. 53). Se trata de la creatividad que puede atribuírsele al uso de cualquier sistema bien definido (sobre este concepto, véase la sección 4.6). Si efectivamente los lenguajes naturales son sistemas bien definidos, como parece pretender Chomsky (hay una interesante discusión de este punto por Hockett en *El estado actual de la lingüística*), parece razonable pensar que esto nos suministra una pista importante para llegar a averiguar cuáles son las características del modo de funcionar el cerebro humano. Pero no parece, en cambio, ofrecer ningún interés especial para un estudio empírico del uso del lenguaje, y es indudable que, en este contexto, las afirmaciones de Chomsky sobre la creatividad desenfocan el tema totalmente.

Es curioso que la sobrevaloración de la creatividad por parte de Chomsky se produce en conexión con su crítica a la analogía cuando se recurre a ésta para explicar el hecho de que seamos capaces de producir y entender oraciones nuevas. Y es curioso justamente porque la analogía puede dar lugar a otro sentido de la creatividad que Chomsky no considera, y que se encuentra más cercano del sentido que usualmente tiene ese término en el lenguaje ordinario. Es la creatividad propia de la actividad artística, y que consiste en que los productos futuros de tal actividad pueden no encajar del todo en una definición construida sobre la base de ejemplos anteriores. Así, por ejemplo, cuando apareció la pintura abstracta, ésta habría sido incluida bajo el concepto de pintura por analogía con la pintura usual, y a pesar de que tal concepto no sería del todo adecuado para esta nueva forma de entender la pintura. De manera semejante, el uso del lenguaje es creativo en el sentido de que continuamente extendemos y ampliamos el sentido de las palabras y expresiones a nuevos objetos, fenómenos y situaciones para los cuales, en principio, no eran totalmente adecuadas. Según Sampson (*op. cit.*, pp. 54 s.), esto corresponde a la concepción del segundo Wittgenstein acerca del lenguaje como un conjunto de usos entre los cuales no habría más que semejanzas parciales (parecidos de familia). Según esto, la concepción de Wittgenstein en su segunda etapa, que ya veremos con detalle en otro capítulo, implicaría un concepto de creatividad del uso lingüístico mucho más claro y menos problemático que el concepto chomskiano. Esta creatividad de orden semántico es paralela al proceso de creación de nuevas formas sintácticas, e incluso fonéticas, que caracteriza la evolución de la lengua. Pero, como es sabido, la teoría lingüística de Chomsky

está ligada a una consideración sincrónica de la lengua que hace abstracción de sus procesos de modificación.

3.5 Recapitulación terminológica

El término «lenguaje» se emplea con un alto grado de ambigüedad, como ya habrá podido apreciarse.

En primer lugar, se llama lenguaje a cualquier sistema de signos que cumpla con unas condiciones mínimas, como las enunciadas en la sección 2.4, anteriormente.

En segundo lugar, y limitándonos al lenguaje verbal humano, se puede llamar lenguaje a la facultad específica humana (especie *homo sapiens*) de comunicarse por medio de sonidos articulados. Los fundamentos y presupuestos biológicos de esta facultad los estudia la biología del lenguaje, y su desarrollo y patología, la psicología del lenguaje.

En tercer lugar, puede llamarse lenguaje a un producto particular de la facultad lingüística, es decir, a una lengua, como el latín o el castellano. De las lenguas trata la lingüística aplicada.

En cuarto lugar, puede hablarse del lenguaje, en general, para referirse a aquello que es común a todas las lenguas, lo cual recibe también el nombre de gramática universal, y es objeto de la teoría lingüística.

En quinto lugar, se habla de lenguaje, añadiéndole algún calificativo, cuando se hace referencia a un modo particular de usar la lengua, es decir, a lo que Coseriu llama norma lingüística. Se habla así de lenguaje literario, vulgar, científico, político, etc. En su aspecto más general, esto es tema para la teoría lingüística, y en cuanto a sus aspectos concretos, especialmente por lo que se refiere a aquellos tipos de lenguaje vinculados a grupos sociales característicos, para la sociología del lenguaje.

En sexto y último lugar, se emplea también el término «lenguaje» para referirse a un acto individual de habla, como cuando decimos «Tu lenguaje en esa ocasión fue grosero».

Por lo que se refiere al lenguaje verbal humano, los sentidos anteriores pueden reducirse a tres fundamentalmente distintos. El lenguaje como facultad, el lenguaje como sistema de signos y el lenguaje como conjunto de episodios individuales. Recordando la distinción entre signo tipo y signo acontecimiento (introducida en la sección 2.2), el lenguaje como sistema constituye un conjunto de tipos y el lenguaje como episodios de habla constituye un conjunto de acontecimientos.

La ambigüedad señalada es corriente en los escritos de todo tipo sobre el lenguaje, y muy especialmente en las obras de carácter filosófico, por lo que debe tenerse en cuenta a fin de evitar posibles confusiones. En esta investigación, y aun cuando en muchas ocasiones utilizaré términos específicos como «lengua», «norma», «sistema», «uso», «habla», etcétera, seguiré empleando, por razones así de estilo como de brevedad, el término «lenguaje» en cualquiera de las acepciones relacionadas, pero siempre estará claro, por el contexto, de qué acepción se trate.

Lecturas

Un excelente estudio de los conceptos de lengua y habla en Saussure, y de las interpretaciones que de ellos se han dado en la lingüística, es el trabajo de Coseriu, ya citado, que bajo el título «Sistema, norma y habla» se halla incluido en su libro *Teoría del lenguaje y lingüística general* (Gredos, Madrid, 1967). No hace falta añadir que este trabajo contiene además, naturalmente, una clara exposición del concepto de norma.

Para los conceptos de competencia y actuación en Chomsky, puede verse su libro *El lenguaje y el entendimiento* (Seix Barral, Barcelona, 1971), que sin duda constituye el mejor resumen hecho por el autor de las implicaciones filosóficas de su teoría del lenguaje. Se encontrará un planteamiento más técnico y más riguroso de esos conceptos en el primer capítulo de su obra principal, *Aspectos de la teoría de la sintaxis* (Aguilar, Madrid, 1970). El breve libro de John Lyons, *Chomsky* (Grijalbo, Barcelona, 1974), es, no obstante su brevedad, muy claro y riguroso, y trata estas cuestiones en el capítulo octavo.

Capítulo 4

RS GRAMMATICA

La palabra. esa lana marchita...
La palabra, esa arena machacada...
La palabra. la palabra, la palabra, qué torpe vientre
[hinchado.

(Vicente ALEIXANDRE, *Espadas como labios.*)

4.1 La definición del lenguaje

Si en el curso de un viaje por tierras desconocidas topáramos de súbito con una comunidad de seres humanos y advirtiéramos que éstos producían con la boca ciertos sonidos, que se dirigían tales sonidos unos a otros y que a la audición de dichos sonidos seguían diferentes formas de comportamiento por parte de los presentes (incluida la producción de nuevos sonidos), sin duda inferiríamos que estos seres se comunicaban por medio de un lenguaje, aun cuando de este supuesto lenguaje no nos fuera dado, por el momento, entender nada. En principio, y para cualesquiera seres inequívocamente reconocibles como humanos, no parece difícil decidir si tienen un lenguaje, esto es, si emplean una lengua. Es más, lo que nos extrañaría, y encontraríamos difícil de explicar, es encontrar seres aparentemente humanos sin lenguaje; si tal ocurriera, nuestra reacción más probable consistiría en buscar razones adicionales para negar a esos seres la condición de humanos, esto es, de pertenecientes a la especie *homo sapiens*.

Siendo esto así, parecerá cuestión bizantina la de definir el lenguaje humano verbal, y sin duda en parte lo es. Pero no se olvide que buscar definiciones es uno de los empeños filosóficos más antiguos, y aunque nosotros no pensemos, a diferencia de algunos de nuestros más ilustres predecesores, que conseguir una exacta definición equivale a penetrar de manera cuasimística en el interior de lo definido, si seguimos creyendo, al modo socrático, que una buena definición proporciona claridad conceptual definitiva a la hora de argüir sobre lo definido. Pero hay otra razón de orden más práctico y de interés más inmediato. Resulta razonable esperar que, puesto que todos los seres humanos poseemos una organización biológica

fundamentalmente idéntica, que es la que determina la posibilidad de hablar científicamente de una especie humana, todas las lenguas humanas tengan unas características comunes y en las que coincidan. Es patente que la explicitación de estas características constituye un tema de interés teórico, tanto por lo que puede contribuir a nuestro conocimiento de las lenguas particulares como por lo que puede aclarar y completar nuestro conocimiento de la naturaleza humana. Pues bien, de estas características universales para todas las lenguas hay algunas tan generales, básicas e inmediatas que no puede dudarse de que vayan a estar presentes en una lengua determinada, pues si no lo estuvieran, sólo por esto habría que cuestionar si se trataba realmente de una lengua. Tales características son las que podemos considerar como integrantes de una definición del lenguaje. Son características que, como acontece con todo conocimiento empírico y de lo natural, hemos obtenido por inducción, más o menos rigurosa y explícita, sobre la base de las lenguas que conocemos.

¿Qué ocurriría si halláramos una comunidad humana cuyo sistema de comunicación mostrara la ausencia de alguno de los rasgos incluidos en nuestra definición de lenguaje? Hay dos alternativas. O bien negar a esta forma de comunicación la condición de lenguaje o bien decidir que nuestra definición era errónea y modificarla de manera acorde. Es claro que la elección de una de estas alternativas dependerá de la importancia que demos al rasgo en cuestión, así como del grado de semejanza que tal sistema de comunicación guarde con las lenguas conocidas. En todo caso, la contemplación de posibles ejemplos de solución dificultosa no pasa de ser, me parece, un entretenimiento para aficionados a la ciencia ficción. La identificación de un sistema humano de comunicación como lenguaje verbal y su diferenciación de otros posibles sistemas de comunicación, humanos o animales, no ofrece por el momento dificultad apreciable, aunque ello no impide que muchas de las definiciones existentes de lenguaje resulten insatisfactorias.

En un interesante trabajo, Hockett ha sugerido una lista bastante completa de características definitorias para el lenguaje humano («The Problem of Universals in Language»). De las que Hockett señala, recogeré y comentaré a continuación las que me parecen más claras y más básicas, modificando en algunos casos su formulación y reordenándolas de manera más sistemática; también incluiré algún rasgo de cuyo carácter definitorio Hockett duda. La lista que ofrezco aquí es más breve y más sistemática que la que presenté hace tiempo en otro lugar siguiendo más literalmente a Hockett (*La teoría de las ideas innatas en Chomsky*, sección 3.4).

Un lenguaje verbal (como es el humano) es un sistema de signos que posee las características siguientes:

1. Por lo que se refiere al medio empleado para la comunicación y a los órganos por los que ésta se realiza, se emplean sonidos y se utiliza el *canal vocal-auditivo*. Esto tiene como consecuencia que las señales lingüísticas se transmitan de forma difundida, se reciban en una dirección

determinada (la dirección en la que está el emisor) y desaparezcan con rapidez. A diferencia del lenguaje, ciertos sistemas de comunicación animal no son sonoros (la danza de las abejas) o aunque sean sonoros no son vocales (las señales de los grillos). El presente rasgo excluye también de la categoría de lenguaje verbal en sentido estricto otros sistemas humanos de comunicación que podemos considerar como *no naturales*, por ejemplo, la escritura, el morse, las señales de humo o las producidas por medio de banderas, tambores, etc.

2. Por lo que se refiere a su estructura, consta de dos subsistemas, el subsistema fonológico y el subsistema gramatical (que incluye tanto elementos sintácticos como semánticos). Por esta razón, se denomina *dualidad* a esta característica. Esto corresponde a lo que Martinet ha llamado doble articulación, a saber, el hecho de que el lenguaje verbal se articula al nivel fonológico en unidades no significativas (fonemas) y al nivel gramatical en unidades significativas (morfemas o monemas). A veces, como ocurre en la teoría de Chomsky, se distinguen tres subsistemas, el fonológico, el sintáctico y el semántico. Esto puede tomarse como el resultado de una subdivisión del subsistema gramatical, de manera que, en rigor, no tiene por qué ser incompatible con la dualidad. No se olvide, además, que la distinción entre sintaxis y semántica es uno de los puntos más disputados en la lingüística reciente, y en particular entre algunos de los discípulos de Chomsky.

3. Por lo que se refiere a su potencia, esto es, a los límites de su utilización, es un sistema abierto, o, lo que es lo mismo, posee *creatividad*, y esto en todos los sentidos que ya hemos visto. En primer lugar, desde el punto de vista formal o sintáctico, porque no hay límite para el número de expresiones correctas distintas que pueden formarse aplicando las reglas del sistema. En segundo lugar, desde el punto de vista semántico, porque es posible expresar nuevos contenidos semánticos (ideas, designaciones, descripciones, valoraciones, emociones...) por medio de los elementos existentes utilizando los necesarios procedimientos analógicos. Y finalmente, porque es incluso posible llegar a modificar las reglas del sistema, semánticas, sintácticas o morfofonémicas, por medio de una práctica desviada suficientemente prolongada. En sentido estricto, sólo la primera forma de creatividad sería una característica o rasgo del sistema. Las dos formas restantes serían más bien características de su utilización o uso, aunque no menos centrales para una definición del lenguaje. Piénsese que el rasgo número uno no era tampoco propiamente un rasgo del sistema cuanto del medio y de los órganos a través de los cuales éste se emplea.

4. Por lo que se refiere a las relaciones entre el sistema y la realidad, en el sistema se establecen una serie de relaciones convencionales o arbitrarias entre algunos de sus elementos y ciertos rasgos o partes de la realidad. Podemos llamar a este rasgo *convencionalidad*. En el segundo capítulo lo vimos subrayado por Saussure para el signo lingüístico, y tomado

por Peirce y otros como rasgo característico de los signos simbólicos, entre los cuales colocan el lenguaje. Es, por consiguiente, uno de los rasgos definitorios señalados con más frecuencia, aunque no hay que olvidar que es un rasgo en el que coinciden con el lenguaje verbal otros sistemas de signos. Muchos sistemas de comunicación entre animales superiores, como diferentes tipos de mamíferos, emplean también signos que son arbitrarios desde el punto de vista semántico. Incluso la danza de las abejas es simbólica en parte (y en parte icónica en la medida en que la dirección de la danza indica la dirección en que se encuentra el referente, a saber, la fuente de alimento, pero se trata de una iconicidad muy restringida, como puede apreciarse).

Hockett distingue la convencionalidad de otro rasgo, que llama semantividad, y que consiste en el hecho de que al menos algunos de los signos del sistema poseen denotación, refieren a elementos o aspectos de la realidad en virtud de las relaciones convencionales mencionadas. La semantividad es mencionada por muchos autores como uno de los tres o cuatro rasgos definitorios más típicos y fundamentales para definir el lenguaje. Es patente, sin embargo, que la semantividad no distingue al lenguaje de ningún otro sistema de signos, ya que, por definición de signo (como vimos en el capítulo segundo), no hay signo sin significado, ni por consiguiente sin remisión a la realidad. Por tal razón, no me parece interesante distinguir la semantividad de la convencionalidad como dos rasgos diferentes, ya que, estando además implicada la primera por la segunda, queda de hecho recogida al hablar de ésta.

5. Por lo que hace a la utilización del sistema en un contexto determinado, es posible emplear el lenguaje para referirse a objetos o aspectos de la realidad lejanos respecto del lugar y momento de la comunicación. A este rasgo se le llama a veces *desplazamiento*. También la danza de las abejas lo posee. Otros sistemas de signos, en cambio, sólo pueden emplearse para la comunicación acerca de objetos que quedan en ese momento dentro del campo perceptivo de los organismos, como acontece con muchos insectos y mamíferos.

6. Por lo que toca a la relación entre los signos y sus usuarios, el lenguaje verbal se caracteriza por el hecho de que estos últimos son indistintamente emisores y receptores, y de que el emisor es siempre al mismo tiempo receptor de su propia emisión o mensaje. Ciertos sistemas de comunicación entre insectos no cumplen con estas condiciones, pues en ciertas especies sólo los machos emiten señales, y en otras el emisor no puede percibir las señales que emite él mismo.

7. El lenguaje puede ser ámbito de referencia de sí mismo, esto es, puede utilizarse reflexivamente dando lugar a metalenguajes. La *reflexivi-*

dad, así entendida, distingue sin duda al lenguaje verbal de todos los sistemas de comunicación no humanos, aunque no lo distingue de otros sistemas de comunicación humanos no verbales; sistemas de comunicación como las banderas o el morse podrían utilizarse metalingüísticamente, y no es de extrañar puesto que derivan del lenguaje verbal. Esta característica de la reflexividad procede sin duda de la más peculiar característica psicológica de la especie humana, a saber, la capacidad para responder a sus propias respuestas, reduplicando así progresivamente el ámbito de sus experiencias.

Las características anteriores son lo bastante generales, y a primera vista imprescindibles, como para que sea razonable pensar que no pueden faltar en ninguna lengua humana, y conjuntamente parecen suficientes para distinguir el lenguaje verbal de cualquier otro sistema de comunicación humano o animal (se encontrará una comparación entre el lenguaje humano y diferentes sistemas de comunicación animal sobre la base de los rasgos de Hockett en el artículo de Thorpe, «The Comparison of Vocal Communication in Animals and Man»). Pero las características o rasgos que son comunes a todas las lenguas, y por tanto universales, no se reducen a éstas. Es posible señalar una serie de características que, aunque considerablemente más concretas que los rasgos definitorios mencionados, aparecen en todas las lenguas conocidas y son, por tanto, también universales, al menos en relación con los límites de nuestro conocimiento actual de las lenguas humanas. A estas características se les llama *universales lingüísticos*, aunque no son más universales que los *rasgos definitorios*. En ocasiones, se toma la categoría de universales lingüísticos como una categoría general dentro de la que se distinguen aquellos universales que, por ser más generales y abstractos, se van a utilizar como características definitorias, y aquellos otros que no van a tener este carácter. A fin de marcar más claramente la distinción, aquí emplearé regularmente la primera terminología.

¿En qué se diferencian los universales lingüísticos de los rasgos definitorios? En que, por ser más concretos, su existencia universal en todas las lenguas no resulta tan inmediata ni tan clara, y en tal medida cabe razonablemente dudar de ella. Esto tiene las siguientes consecuencias. Primero, que cada vez que comprobamos que un universal aparece en una lengua, esto aumenta efectivamente nuestros conocimientos sobre el lenguaje humano, en un sentido en el que no puede decirse lo mismo cuando meramente comprobamos que una lengua se ajusta a nuestra definición de lenguaje. Y segundo, que por ello mismo los universales constituyen hipótesis de trabajo cuya confirmación o falsación es parte de la investigación lingüística. En terminología tradicional, podría decirse que los rasgos definitorios son necesarios mientras que los universales son contingentes o accidentales. Pero esta terminología puede sugerir imágenes y plantear cuestiones que, siendo más falsas que reales, es preferible rehuir. La única diferencia controlable entre características necesarias y accidentales es que nos parece altamente improbable que un lenguaje carezca de alguna de aquéllas (y de hecho no somos capaces de imaginar cómo sería un lenguaje

así), mientras que nos parece más fácil que pudiera faltar alguna de las segundas (y de hecho podemos imaginar que así acontezca).

4.2 Los universales lingüísticos

En el trabajo ya citado («The Problem of Universals in Language») Hockett ha suministrado un variado conjunto de universales lingüísticos, de los que vamos a ver los más relevantes a efectos filosóficos. De entrada, dejaremos fuera de nuestra consideración los de tipo fonológico, y limitaremos esta exposición a los que Hockett llama gramaticales, que son en su mayoría de orden semántico. También en este caso reordenaré la lista de Hockett para hacerla más sistemática, y en algún caso modificaré su formulación según indicaré en su momento.

1. En toda lengua hay elementos cuya denotación cambia dependiendo de ciertos rasgos elementales de la situación, esto es, del contexto extralingüístico. Tales son los elementos que se denominan deícticos (del griego *deixis*, acción y efecto de mostrar o señalar) o indéxicos. Son, por ejemplo, los pronombres personales y demostrativos, ciertos adverbios como «ahí», «ahora», etc.

2. Entre los elementos anteriores hay en toda lengua uno que denota o se refiere al sujeto que habla y otro que denota al sujeto al cual se habla; es decir, que los pronombres de primera y segunda persona del singular son, en principio, universales.

3. Todas las lenguas poseen elementos que no denotan nada y cuya función consiste en relacionar entre sí los elementos denotativos.

Con terminología tradicional podríamos decir que toda lengua tiene elementos sincategoremáticos, además de elementos categoremáticos o denotativos (recuérdese que la existencia de estos últimos es un rasgo definitorio de todo sistema semiótico, a saber, la semanticidad, y lo es *a fortiori* del lenguaje verbal; de hecho está implicado por el rasgo de convencionalidad como hemos visto en el número cuatro de la sección precedente). Hockett da como ejemplos preposiciones y conjunciones (aunque en el caso de aquéllas señala que poseen alguna denotación). A efectos de la terminología tradicional mencionada, debe tenerse en cuenta que se suelen considerar como sincategoremáticos los adjetivos y los adverbios, además de las preposiciones y conjunciones, pues se estima que los términos sincategoremáticos o sincategoremas significan algo en cuanto modificación o determinación de otra cosa. Los términos categoremáticos, o categoremas, en cuanto que significan algo por sí solos, se reducen a nombres y verbos. La distinción escolástica, por consiguiente, no coincide con la que estamos viendo, aunque tiene un sentido parecido. La distinción de Hockett, por atenderse a la función denotativa, es susceptible de interpretaciones diferen-

tes según la teoría de la referencia que se acepte, y deja abiertas, por ello, numerosas cuestiones.

4. Toda lengua tiene nombres propios, esto es, elementos cuya función se reduce a denotar algo, pero sin connotar ninguna propiedad cuya posesión por el objeto denotado justifique por sí sola la aplicación del nombre.

De hecho, puede ocurrir que en una lengua la mayor parte de los nombres propios de cierto tipo connoten alguna propiedad. Así, en castellano, la mayor parte de los nombres de pila de personas connotan o bien masculinidad o bien feminidad, y en razón de esto se aplican respectivamente a hombres o a mujeres, pero ello no ocurre absolutamente siempre («Trinidad» sería un ejemplo) y, por tanto, prueba que no es necesario que ocurra. El tema de los nombres propios tiene una interesante tradición dentro de la filosofía analítica y ha originado recientemente problemas en relación con el concepto de necesidad, todo lo cual veremos en su momento.

5. En toda lengua hay elementos gramaticales que no pertenecen a ninguna de las categorías mencionadas, es decir, que ni son elementos deícticos, ni son elementos no denotativos, ni son nombres propios.

Esta formulación resulta excesivamente débil, pues parece razonable esperar que pudiera especificarse qué otras categorías cabe señalar hipotéticamente como universales.

6. Aparte de las tres categorías mencionadas, ninguna lengua posee un vocabulario gramaticalmente homogéneo, y en éste puede siempre hacerse una distinción del tipo de la distinción entre nombre y verbo.

Este universal viene a completar al anterior, aunque en la lista de Hockett ambos están en lugares apartados entre sí, lo que dificulta percibir la relación entre ambos.

7. Toda lengua distingue entre predicados monádicos y predicados poliádicos, o sea, entre predicados con un argumento y predicados con más de un argumento. Dicho de otra forma, toda lengua distingue entre propiedades y relaciones.

Hockett se limita a formular este universal en términos de la distinción entre predicados monádicos y diádicos, es decir, con uno y con dos referentes. La formulación anterior, por su mayor generalidad, me parece preferible, y no imagino qué razones pueden hacer más probable una hipótesis más restringida como la de Hockett.

8. Toda lengua posee un tipo de cláusula de estructura bipartita cuyos componentes pueden denominarse tema y comentario. Esto es, una cláusula en la que, por ejemplo, se menciona aquello de lo que se va a hablar y a continuación se dice algo sobre ello.

Parece que se trataría de la distinción entre sujeto y predicado. Naturalmente, el orden en el que aparecen dentro de la cláusula varía según las lenguas.

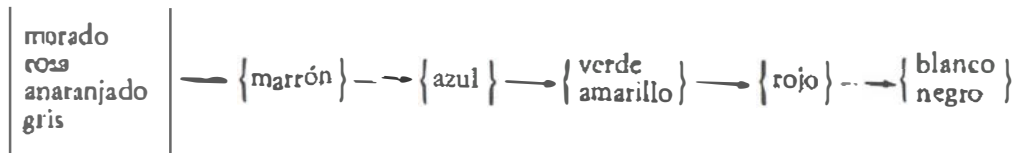
9. Toda lengua tiene por lo menos dos órdenes básicos de estructuración (*patterning*) gramatical. En el caso de que sean solamente dos, corresponderían a lo que tradicionalmente se llama morfología y sintaxis. Hockett señala, no obstante, que para ciertas lenguas como las de tipo chino parece que es preferible prescindir de la pluralidad de estructuraciones.

Este universal resulta, en todo caso, particularmente debatible si se tiene en cuenta la tendencia actual a disolver la morfología entre la fonología y la sintaxis, tendencia muy acusada en la lingüística transformacional.

La lista anterior constituye una buena ilustración de lo que se entiende por universales lingüísticos, y de su diferencia con los rasgos definitorios. Para nosotros es en particular útil por cuanto casi todos sus ejemplos son semánticos y tienen relevancia para el problema de las relaciones entre la lógica y el lenguaje. En la literatura sobre el tema se encuentran igualmente listas de universales tanto de tipo fonológico (por ejemplo, los sugeridos por Hockett, y a los cuales se ha hecho alusión antes) como de tipo sintáctico. Entre estos últimos, cabe mencionar una interesante lista de características universales presentada por Greenberg que en su mayoría afectan al orden de las palabras en la frase («Some Universals of Grammar with Particular Reference to the Order of Meaningful Elements»). Una gran parte de estos ejemplos consisten en universales condicionales, esto es, del tipo de: «Si una lengua tiene la característica x, entonces tendrá también la característica y» (aunque no necesariamente viceversa). Véase como muestra el universal número trece de la lista de Greenberg: «Si, en una lengua dada, el objeto nominal precede siempre al verbo, entonces los verbos subordinados preceden siempre al verbo principal». O por poner un ejemplo que no afecta al orden sintáctico sino a las categorías morfológicas generales, considérese el universal número treinta y seis de la lista citada: «Si una lengua posee la categoría de género, entonces tendrá también la categoría de número». Prestar más atención a los universales sintácticos nos apartaría con exceso y sin justificación del propósito central de este estudio.

Hablando de universales condicionales (o implicativos, como también se los llama), es interesante notar que pueden encontrarse ejemplos igualmente en semántica. Tal es el caso de la hipótesis de Berlin y Kay (*Basic Colour Terms*) acerca de la estructura de los términos para colores en las diferentes lenguas. Sobre la base de haber comparado cerca de un centenar de lenguas, de estructura y grado de evolución muy diversos, Berlin y Kay han sugerido que existen once categorías básicas para clasificar los colores, y que, aunque las lenguas examinadas difieren entre sí extraordinariamente según el número y variedad de las categorías que utilizan (desde dos hasta

las once, en diferentes combinaciones), hay entre ellas relaciones de condicionamiento en el sentido de que cualquier lengua que posee ciertas categorías posee también ciertas otras, de acuerdo con el siguiente esquema:



Esto significa que cualquier lengua que contenga una de estas categorías tiene asimismo todas las que estén a su derecha en la tabla. Es decir: si una lengua tiene un término equivalente a «rojo», tiene entonces también términos equivalentes a «blanco» y «negro», esto es, términos que designen los colores que nosotros designamos así en castellano; si tiene un término para designar el color morado, entonces tiene igualmente términos para designar los colores marrón, azul, verde, amarillo, rojo, blanco y negro, etcétera. Cuando dentro de un grupo hay varias categorías significa que entre ellas no hay preeminencia alguna. Por ejemplo: si una lengua cuenta con un término para el amarillo, entonces tiene también términos para el rojo, el blanco y el negro, y lo propio acontece si tiene un término para el verde o los dos, para el verde y el amarillo. Las lenguas consideradas por Berlin y Kay van desde el jalé, hablado en Nueva Guinea, que solamente cuenta con términos para el blanco y el negro, hasta el inglés, que naturalmente posee las once categorías básicas, pasando por lenguas como el hanunoo de Filipinas, que tiene términos para el blanco, el negro, el rojo y el verde. Berlin y Kay han sugerido además que el esquema anterior corresponde a los estadios por los que, en su desarrollo, atraviesa la lengua de un pueblo en lo que respecta a los términos para colores, estadios que van desde el más simple, caracterizado por la posesión de dos términos («blanco» y «negro»), hasta la adquisición de las once categorías básicas. Con esto, la hipótesis de estos autores se convierte en un universal sobre la evolución del lenguaje.

El interés del estudio de Berlin y Kay consiste en que, aplicando la relación de condicionamiento, ha puesto de manifiesto una importante uniformidad que, mientras el conocimiento de otras lenguas no venga a demostrar lo contrario, preside la estructuración de las categorías de colores en las diversas lenguas; y hay que tener en cuenta que éste era un campo en el que no parecía fácil detectar uniformidad alguna ni hallar base para hipótesis universalistas. Es importante notar que la hipótesis no prejuzga la posibilidad o imposibilidad de traducir con exactitud unos por otros los términos para colores de diferentes lenguas. Es de esperar que una lengua que no tenga términos como «rosa» y «anaranjado», pero sí, en cambio, uno como «rojo», dé a este último un alcance mucho mayor que el que tendrá en una lengua que disponga de las tres palabras, pues a falta de otros términos como «rosa» y «anaranjado» sin duda tenderá a incluir en el ámbito de aplicación de «rojo» muchos objetos que en la otra lengua se incluirían

bajo «rosa» o «anaranjado». En esa medida, los términos para el rojo de ambas lenguas no serán exactamente intertraducibles. ¿Por qué afirmar que ambos términos representan la misma categoría? Porque, a pesar de que su aplicación difiera en una serie de casos periféricos, ambos tienen en común como campo de referencia una serie de casos centrales en los que se usan de manera coincidente. Y es esta coincidencia central la que permite comparar ambos términos. Naturalmente, los problemas de esta índole pueden multiplicarse si pasamos de las once categorías básicas a categorías secundarias como «escarlata», «gualdo», etc., pero en principio estas categorías pueden caracterizarse como subtipos de alguna de las categorías básicas (escarlata con respecto al rojo, gualdo con respecto al amarillo, etc.). Por ello, la hipótesis de Berlín y Kay opera únicamente con términos básicos. Hay que añadir que algunos de los datos considerados no se adaptan completamente a la hipótesis. Así ocurre con el ruso, donde no hay un término para el azul, sino dos, uno para el azul claro y otro para el oscuro, y con el húngaro, donde ocurre lo propio respecto al rojo. Tales lenguas tendrían doce términos básicos y no once. (Para lo anterior puede verse el resumen que ofrece Leech en el capítulo oncenso de su *Semántica*, donde se hallará asimismo un resumen de las investigaciones sobre los términos de parentesco, campo en el que también podrían encontrarse, aunque no tan claramente, interesantes ejemplos de universales semánticos.)

Hay que mencionar que este tipo de universales condicionales es, en un sentido trivial, más fácil de ejemplificar que los que consisten en la atribución de propiedades. Téngase en cuenta que un universal condicional no consiste en atribuir a todas las lenguas una cierta propiedad *x* o *y*, sino en postular que, si una lengua tiene la propiedad *y*, entonces tendrá la propiedad *x*. Naturalmente, un universal de este tipo es confirmado por cualquier lengua que cumpla con cualquiera de las tres condiciones siguientes: que posea las características *x* e *y*; que posea *x* pero no *y*; que no posea ninguna de las dos. O dicho a la inversa: un universal de ese tipo solamente queda falsado por una lengua que tenga la propiedad *y* y que carezca de *x*. El filósofo de la ciencia puede aquí replantear, si lo desea, todos los problemas relativos a la llamada paradoja de la confirmación, pero para nosotros no es el momento de entrar en detalles de metodología científica (para un tratamiento elemental y general del tema puede verse el capítulo cuarto de la *Introducción a la filosofía de la ciencia*, de Lambert y Brittan) *.

Los universales lingüísticos desempeñan dentro de la teoría de Chomsky una peculiar función que excede del carácter de meras hipótesis generales que poseen en la teoría lingüística no chomskiana. Se trata de que, como parte de su concepción mentalista acerca del lenguaje, Chomsky, por razones que veremos con mayor detalle más adelante, ha defendido la necesidad

* Como puede comprenderse, una investigación de tan amplio alcance como los cuatro volúmenes recopilados por Greenberg bajo el título *Universals of Human Language*, sólo es posible sobre la base de un concepto de universal lingüístico muy débil y relativo aunque no por ello carente de interés.

de atribuir al niño que está aprendiendo su lengua un conocimiento tácito de esos universales, conocimiento sin el cual no podría explicarse, al parecer, el aprendizaje de una lengua nativa (*Aspectos*, cap. 1, secc. 5).

Chomsky distingue dos tipos distintos de universales, los formales y los sustantivos. Una teoría de los universales sustantivos afirmará que los elementos lingüísticos de cierto tipo deben, para cualquier lengua, ser extraídos de un conjunto fijo de tales elementos. Es decir, los universales sustantivos no son conjuntos de elementos todos los cuales hayan de aparecer en todas las lenguas, sino conjuntos de elementos a los cuales pertenecen los elementos de ese tipo que aparecen en todas las lenguas. O dicho de otra forma: para cada conjunto de elementos, si éstos son efectivamente universales, no habrá lengua alguna en la que haya elementos que no pertenezcan al conjunto, pero puede haber elementos del conjunto que no se encuentren en algunas lenguas. Digámoslo aún de otra manera: para todo conjunto de elementos de cierto tipo, este conjunto es un conjunto de universales sustantivos si, para cualquier lengua, los elementos de ese tipo que en ella aparecen constituyen un subconjunto de aquéllos. Veamos los ejemplos que Chomsky menciona. En cuanto al componente fonológico, serían universales sustantivos el conjunto de los rasgos fonéticos por medio de los cuales, según la teoría de Jakobson, se puede caracterizar los fonemas de todas las lenguas. Estos rasgos se determinarían con independencia de cualquier lengua, recurriendo a las características acústicas y articulatorias de los sonidos vocales. Ejemplos de tales rasgos son la cualidad de vocal o de consonante, la nasalidad, la cualidad de sordo o sonoro, etc. Para el componente sintáctico, serían ejemplos de universales sustantivos las categorías que tradicionalmente se vienen utilizando para el análisis gramatical, como son las de nombre, verbo, etc. En el componente semántico, los universales sustantivos incluirán aquellas categorías bajo las cuales se pueden agrupar y distinguir los lexemas de una lengua según, por ejemplo, su función designativa; tales distinciones son las que se pueden hacer entre los lexemas según designen personas, sentimientos, formas de conducta, diferentes clases de objetos, etc.

Es patente, por lo dicho, que los universales sustantivos son universales solamente en un sentido débil, particularmente manifiesto en el hecho de que la ausencia de un universal en un determinado conjunto de lenguas no nos permite asegurar que tal elemento no sea universal en este sentido. Por su parte, los universales formales consisten en cumplir con determinadas condiciones abstractas que se refieren al carácter de las reglas gramaticales de toda lengua. A diferencia de los universales sustantivos, que hacen referencia, en frase de Chomsky, al vocabulario para la descripción de una lengua, los universales formales incluirán las características más abstractas de la gramática de cualquier lengua. Puesto que Chomsky piensa que su teoría lingüística es la correcta, y que es por ello aplicable a toda lengua, serán en principio universales formales las características de la gramática generativo-transformatoria de tipo chomskiano. Así, en el componente sintáctico, la condición de que existen reglas transformatorias que conviertan

las estructuras profundas en estructuras superficiales, el que las transformaciones dependan de la estructura de las oraciones, etc. En el componente fonológico serán universales formales exigencias como la de que al menos algunas de las reglas fonológicas se apliquen cíclicamente, desde los elementos más simples a los más complejos de la oración. Por último, por lo que hace al componente semántico, serían ejemplos de universales formales la condición de que los nombres propios, y otros términos designadores de objetos, sólo puedan designar objetos espacio-temporalmente continuos (*sic* en Chomsky, aunque por la nota que añade se ve que quiere decir «continuos»); e igualmente, la condición de que, en todas las lenguas, los términos para colores dividan el espectro en segmentos continuos, o la de que los artefactos sean definidos en términos de propósitos y necesidades humanas, y no en términos de propiedades físicas solamente. (Para todo lo anterior, además de la sección quinta del capítulo uno, de *Aspectos*, puede verse el capítulo dos de *Lenguaje y entendimiento*, así como «La naturaleza formal del lenguaje».)

Como se ve por los escasos ejemplos aducidos por Chomsky, los universales, particularmente los de índole formal, están en su teoría íntimamente ligados a su concepción de la gramática. Esto puede resultar negativo en la medida en que ciertos aspectos de ésta son tan discutibles, y de hecho tan controvertidos, que puede ser preferible desligar de ellos la teoría de los universales lingüísticos, y construir ésta, en cambio, sobre la base de características gramaticales sobre las que pueda haber mayor acuerdo entre los lingüistas. Juzgo que esto es lo que ocurre con los ejemplos de universales ofrecidos por Hockett. Piénsese que hasta el propio concepto de transformación ha sido eliminado de la descripción gramatical por algunos lingüistas postchomskianos, como veremos posteriormente. En un sentido general, sin embargo, está claro que la teoría de los universales lingüísticos va ligada al desarrollo de la teoría lingüística en su conjunto, puesto que la primera no hará sino ir recogiendo aquellas afirmaciones generales sobre las lenguas humanas que vayan siendo progresivamente mejor confirmadas a la vez que muestren un rendimiento explicativo mayor.

Precisamente para evitar la excesiva vinculación de la teoría de los universales lingüísticos a un tipo específico de lenguas o a un modelo particular de gramática, así como para eludir las fuertes implicaciones mentalistas de la concepción chomskiana, se ha propuesto la alternativa de considerar los universales lingüísticos, no como propiedades comunes a todas las lenguas, sino como características propias de la estructuración (*patterning*) de las variaciones existentes de unas lenguas a otras con respecto a una propiedad dada (Keenan, «The Logical Diversity of Natural Languages»). La idea básica que subyace a esta concepción es que, contra lo que parecen pensar los chomskianos, cualquier lengua sólo realiza en una mínima medida lo que es universalmente posible. Esto es, que las condiciones y exigencias con las que cumple la gramática de una lengua son mucho mayores que las que son universalmente válidas, y, por tanto, que más que buscar los universales en cada lengua hay que buscarlos en las diferencias que,

respecto a una cierta característica, las separan. Es interesante constatar que uno de los tipos generales de universal lingüístico que ejemplifica adecuadamente esta propuesta es el tipo de los universales condicionales a los que hemos hecho alusión más arriba, puesto que un universal condicional no supone la existencia de una cierta propiedad *x* en todas las lenguas, sino simplemente afirma que, si una lengua tiene otra cierta propiedad *y*, entonces tiene también *x*. Y para determinar esto son relevantes todas aquellas lenguas que se distinguen entre sí por tener las propiedades *x* e *y*, por tener *x* pero no *y*, o por no tener ninguna de las dos. Por consiguiente, lo que cuenta aquí es cómo se relacionan entre sí las diferencias de unas lenguas a otras.

Antes de abandonar este tema vale la pena recordar que Greenberg (*Language Universals*) ha puesto de relieve la universalidad de una distinción que, por su generalidad, es común a todos los componentes de la gramática, y excede en consecuencia las diferencias entre fonología, sintaxis y semántica. Se trata de la distinción entre categorías marcadas (*marked*) y categorías no marcadas (*unmarked*), que podríamos llamar también categorías débiles y fuertes respectivamente. Entre dos categorías opuestas entre sí, es fuerte o no marcada aquella que puede tener dos valores: o bien el valor parcial que se agota en ser término de oposición a la otra categoría, o bien el valor total que integra ambos términos de la oposición. Así, cuando se trata de dos lexemas contrapuestos entre sí, es característico que el fuerte o no marcado se emplee como genérico además de como específico. Por ejemplo, en la contraposición hombre-mujer el término «hombre» es la categoría fuerte, pues se usa tanto con el valor de «varón», en cuanto opuesto a «mujer», como con el valor genérico de «ser humano», cuando se prescinde de la contraposición que se basa en la diferencia de sexo. Greenberg ha mostrado ejemplos de la distinción entre categorías marcadas y no marcadas en los diferentes componentes gramaticales, y por lo que respecta al semántico, a propósito particularmente de los términos de parentesco.

4.3 La gramática como sistema generador

Hemos visto con anterioridad, en la sección 3.4, que uno de los sentidos en que puede entenderse la creatividad del uso del lenguaje es aquel que tiene cuando se afirma que el número de oraciones correctas que pueden formarse en una lengua es potencialmente infinito. Esta infinita variedad de las oraciones correctas de una lengua es, como ya se insinuó, una característica del sistema de ésta, y tal característica deriva del hecho de que las reglas de la gramática de una tal lengua se aplican de forma repetida a una expresión dada por larga o compleja que sea, produciendo así, con cada nueva aplicación, una nueva oración más larga o más compleja (o ambas cosas a la vez). Dicho de otra forma: las reglas gramaticales se aplican recursivamente, y así generan todas las posibles oraciones correctas de la lengua en cuestión. O formulado de otro modo: cualquier expresión

correcta de una lengua puede obtenerse por aplicación recursiva de las reglas de su gramática.

Ha sido lo anterior precisamente lo que ha hecho que la teoría lingüística viniera a concebir la gramática de una lengua como un sistema generador de todas las posibles expresiones correctas en esa lengua, y lo que justifica que una gramática así considerada reciba el nombre de «gramática generativa». A la capacidad de generar así las oraciones correctas se le denomina «capacidad generativa débil». Pero si se considera como lo generado no ya simplemente las oraciones, sino sus descripciones estructurales, entonces se habla de «capacidad generativa fuerte». En este último sentido, una gramática generará un conjunto de descripciones, cada una de las cuales especificará una oración de la lengua de que se trate. Como vamos a ver a continuación, algunos tipos de gramática generativa son defectuosas desde el punto de vista de su capacidad generativa débil, por cuanto ni siquiera generan todas las oraciones de una lengua; así ocurre, por ejemplo, con la llamada «gramática de estados finitos». La gramática transformatoria de tipo chomskiano clásico que examinaremos en la sección 4.4 probablemente posea capacidad generativa débil, pero es más discutible que tenga capacidad fuerte, es decir, que genere el conjunto adecuado de descripciones estructurales para las oraciones de una lengua. De aquí que se hayan propuesto modificaciones posteriores, algunas de las cuales veremos en la sección 4.5. Una gramática generativa es susceptible de una descripción meramente formal en términos matemáticos y lógicos, particularmente dentro de la llamada teoría de autómatas, puesto que éstos no son sino un cierto tipo de sistemas formales (cfr. Quesada, *La lingüística generativo-transformacional: supuestos e implicaciones*, caps. 3 y 6; y Bach, *Teoría sintáctica*, caps. 2 y 8). Aquí prescindiré, sin embargo, de una presentación técnica de este tipo, ya que sería irrelevante para el uso posterior que podamos hacer de estos conceptos, al tiempo que nos obligaría a demorarnos con exceso en este tema. Escogeré, según está siendo la norma en este trabajo, una presentación más bien intuitiva e informal, de la que sólo me apartaré en aquellos momentos o para tales detalles que me parezcan indispensables.

El tipo más simple de gramática generativa es el llamado sistema de reescritura sin restricciones. Consiste en un sistema de reglas de aplicación recursiva que especifican cómo es una oración correcta para un vocabulario dado. Formalmente se trata, sin duda, de una gramática generativa, pero desde el punto de vista lingüístico, esto es, en cuanto a su aplicación a la descripción de las lenguas, carece por completo de interés, y la razón es que esos sistemas de reescritura no establecen condiciones respecto a la estructura interna de la oración, esto es, respecto a los agrupamientos que, de hecho, se realizan con los elementos de una oración en cualquier lengua, agrupamientos tales como sintagma o frase nominal, sintagma verbal, sujeto, complemento, etc. Por esta razón, un sistema de reescritura carente de restricciones es un sistema demasiado potente, pues es capaz de generar, no sólo cualquiera de las lenguas conocidas (quiérese decir: el conjunto de

las oraciones correctas de cualquier lengua conocida) sino también lenguajes aparentemente tan poco naturales como uno en el que todas las oraciones hubieran de tener un número par de palabras; e igualmente podría generar cosas como el conjunto de jugadas posibles en un juego o de ecuaciones en un determinado sistema. Un sistema generativo de este tipo es, por consiguiente, inadecuado, por exceso, para caracterizar y describir los lenguajes naturales. Podemos decir, por ello, que ser un sistema de reescritura sin restricciones constituye una condición necesaria, pero no suficiente, para que un sistema sea una gramática generativa en sentido estricto, esto es, aplicable a la descripción de los lenguajes naturales (véase Bach, *op. cit.*, sección 2.6, de donde he tomado, con adaptaciones, algunos ejemplos). Puesto que todo lo que hay en un sistema no restringido existe asimismo en cualquiera de los que veremos a continuación, lo anterior quedará más claro por lo que sigue.

Una gramática generativa que, incluyendo naturalmente la capacidad para generar un conjunto infinito de oraciones, establece mayores limitaciones sobre su funcionamiento, y que, por tanto, resulta más ajustada al carácter de los lenguajes naturales, es la llamada gramática de estados finitos. Consiste en un sistema que genera las oraciones a base de añadir sucesivamente los distintos elementos, que son elegidos de entre todos los que integran el vocabulario o léxico de la lengua en cuestión (léxico que es finito). La elección de cada elemento se hace en función, no sólo del léxico disponible, sino también del último elemento previamente añadido a la cadena oracional. Así, la generación de la oración castellana:

Los conspiradores fueron descubiertos

se explicaría como el producto de un proceso que comenzaría por seleccionar «los» de entre todas las palabras que, perteneciendo al léxico del castellano, puedan aparecer al principio de una oración. A continuación se habría elegido «conspiradores» entre todas aquellas que pueden aparecer a continuación de «los» (lo cual excluye, en principio, todos los sustantivos y adjetivos femeninos o en singular y todos los verbos). A su vez, «fueron» sería una de las posibles elecciones entre todas las palabras que pueden ser correctamente insertadas a continuación de «los conspiradores» (lo cual excluye, por ejemplo, cualquier artículo). E igualmente para «descubiertos». Nótese que hay alternativas a las elegidas que afectan al resto de la oración y otras que no. Si hubiéramos elegido para comenzar la oración la palabra «aquellos», el resto de la oración podría haber sido el mismo. Si, en cambio, se hubiera seleccionado una palabra como «el», las elecciones subsiguientes también habrían de haber sido diferentes.

Este tipo de gramática generativa recibe su nombre del hecho de que puede representarse como un mecanismo (materializable, si se quiere, como una máquina) que fuera atravesando una serie de estados sucesivos en cada uno de los cuales realiza la oportuna selección léxica a la vista del vocabulario disponible y de la elección realizada en el estado precedente. Por

esta razón se dice que, en este sistema, la generación se produce de izquierda a derecha y palabra por palabra. Para apreciar la riqueza del sistema, así como su plausibilidad, es conveniente tener en cuenta que hay muchos elementos (palabras) que pueden aparecer indistintamente en diferentes estados, y que ciertos elementos, como los adjetivos, son indefinidamente reiterables, esto es, pueden colocarse unos a continuación de otros indefinidamente sin merma de la gramaticalidad de la oración con tal de que se respeten las reglas de concordancia sintáctica y semántica. La generación de infinitas oraciones queda asegurada con tal de que se prevea la generación de oraciones compuestas. En el caso que hemos visto como ejemplo, la diferencia consistiría en que el «mecanismo» generativo, en lugar de detenerse a continuación de la palabra «descubiertos», lo que supone que la oración está completa, prosigue a un estado posterior en el que elige algún elemento como «y», «pero», «aunque», etc., para continuar sucesivamente hasta completar otra, u otras, oraciones adicionales.

Una gramática de este tipo, que formalmente es muy simple, como puede apreciarse, resultaría intolerablemente prolija para dar cuenta de todas las reglas que rigen las diferentes posibilidades de combinación de todos los diversos lexemas y sus formas para constituir la infinita variedad de oraciones de una lengua. Pero además, este tipo de gramáticas es inadecuado en cuanto que no da cuenta, como ha mostrado Chomsky, de determinadas características estructurales de las lenguas que conocemos (véase el capítulo tercero de *Estructuras sintácticas*, de Chomsky, y el resumen de Lyons, en el capítulo quinto de su libro *Chomsky*). Se trata del hecho de que en muchas oraciones se dan relaciones de dependencia entre elementos que no van seguidos, sino que antes bien se encuentran entre sí separados por otros elementos que constituyen una cláusula distinta. Es, en definitiva, el fenómeno llamado autoincrustación, que ejemplifican ciertas oraciones de relativo. Así, la oración:

Los conspiradores, a los que quienes se decían leales ayudaban,
fueron descubiertos

se caracteriza por una relación entre «conspiradores» y «fueron», que, siendo la misma que había en el ejemplo precedente, no descansa, como es patente, en una sucesión de estados, puesto que ambas palabras se encuentran separadas por una serie de estados intermedios que son irrelevantes para la relación que las une. O dicho en otro modo: para la elección de «fueron» no cuenta, contra lo que exigiría una gramática de estados finitos, la palabra «ayudaban», que es la inmediatamente precedente, sino la palabra «conspiradores», que pertenece a un estado muy anterior. Ello justifica la búsqueda de otro tipo de gramática generativa más apropiado.

Este nuevo modelo podría ser el que se conoce como gramática de estructura sintagmática o, en traducción literal, gramática de estructura de frases. Una gramática así consigue, por supuesto, todo lo que puede conseguir una gramática de estados finitos, al tiempo que supera sus insuficien-

cias, o al menos algunas de ellas. Por lo pronto, la gramática de estructura sintagmática toma en consideración, y de ahí su nombre, las frases o sintagmas que pueden discernirse en la oración agrupando sus palabras según los criterios funcionales que el análisis gramatical ha aplicado tradicionalmente. Las palabras, que a estos efectos pueden tomarse como constitutivos últimos de la oración, no son ahora simplemente estados sucesivos de una cadena, sino que se agrupan de forma que revelan la existencia de peculiares relaciones entre sí, relaciones que dependen de la función que las palabras tienen en la oración en cuestión. Consideremos el siguiente ejemplo:

La anarquía aumenta la frustración

Una gramática generativa del tipo que estamos considerando comienza por descomponer la unidad que es esa oración en dos partes, el sintagma o frase nominal, que realiza la función tradicionalmente llamada de sujeto, y representada aquí por «la anarquía», y el sintagma o frase verbal, cuya función es la de predicado, que en el ejemplo consiste en la cadena «aumenta la frustración». A su vez, estos agrupamientos o sintagmas pueden ser ulteriormente descompuestos, de manera que distingamos, por lo que respecta a la porción nominal, entre el artículo y el nombre, y por lo que toca a la parte verbal, entre el verbo y un nuevo sintagma nominal, que corresponde a lo que el análisis tradicional llamaba complemento, y que es «la frustración». Por su parte, y finalmente, este grupo puede analizarse distinguiendo en él de nuevo el artículo y el nombre que en él aparecen. La gramática de estructura sintagmática suministra, pues, una serie de categorías sintácticas en las que descomponer cada unidad sucesivamente, a lo que se añade una especificación de los elementos léxicos que pertenecen a cada categoría simple y que, por tanto, pueden aparecer en el lugar oportuno dentro de la oración. Según esto, la oración que acabamos de analizar es generada por una gramática de estructura sintagmática que consista en las siguientes reglas:

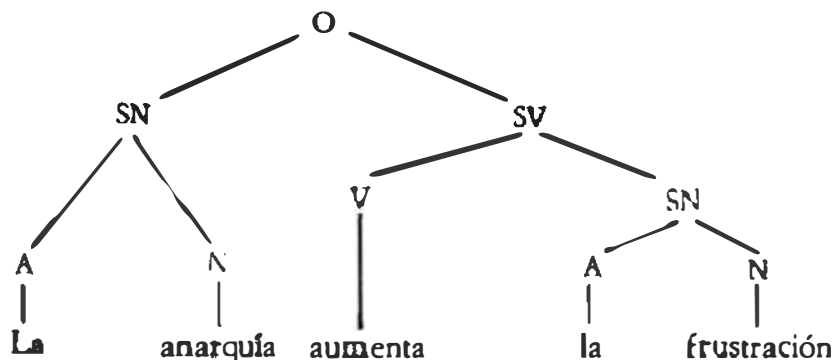
- (1) $O \rightarrow SN + SV$
- (2) $SN \rightarrow A + N$
- (3) $SV \rightarrow V + SN$
- (4) $A \rightarrow \{la, una\}$
- (5) $N \rightarrow \{anarquía, introspección, frustración, psicóloga\}$
- (6) $V \rightarrow \{practica, aumenta, cura\}$

Como se ve, las tres primeras reglas, que pertenecen al tipo de las llamadas reglas de reescritura, justifican las operaciones que acabamos de realizar con nuestra oración. El valor de las abreviaturas es claro: O = oración; SN = sintagma nominal; SV = sintagma verbal; A = artículo; N = nombre; V = verbo. Estas categorías reciben el nombre de vocabulario auxiliar. Las tres reglas restantes especifican el léxico, también llamado vocabulario terminal, que esta gramática acepta para cada categoría básica. Por razones de brevedad se ha escogido un fragmento mínimo del léxico cas-

tellano, y se ha prescindido de complicaciones ulteriores como la flexión, consignando los verbos en un tiempo, modo, persona y número determinados, y los artículos y nombres en singular. No hay que subrayar que se trata de una mínima y supersimplificada gramática que sólo pretende servir de ejemplo. Además de la oración considerada, esta minigramática genera también oraciones como:

La introspección cura la frustración
Una psicóloga practica la introspección
Una frustración aumenta la introspección
La psicóloga practica la anarquía, etc.

Cada una de estas oraciones es el resultado (técnicamente denominado cadena terminal) de una generación, derivación o deducción que procede por aplicación sucesiva de las seis reglas indicadas. No hay que olvidar que cada regla puede ser aplicada más de una vez, en distintos momentos del proceso. Es fácil comprobar que, para cada una de las oraciones que pueden ser generadas por nuestro ejemplo, la regla dos se aplica dos veces, primero al sintagma nominal que hace de sujeto y posteriormente al que hace de complemento dentro del sintagma verbal. Igualmente, las reglas cuatro y cinco se aplican también dos veces, puesto que en cada oración intervienen dos artículos y dos nombres. Este sistema de exhibir la estructura sintagmática de una oración puede representarse gráficamente por medio de diagramas en forma de árbol, del tipo de los que la obra de Chomsky ha popularizado. Para la primera de las oraciones que hemos considerado, su estructura sintagmática vendría dada por el siguiente árbol:



Podría igualmente representarse por medio de paréntesis o corchetes que agruparan oportunamente los distintos elementos de la oración, consignando junto a cada paréntesis la abreviatura correspondiente a la categoría sintáctica de que se tratara. Este recurso se emplea menos porque resulta más dificultoso de leer, especialmente para oraciones largas o complejas. A cualquiera de ambas representaciones gráficas se le denomina, por la función que cumplen, marcador sintagmático o de frases.

Una de las simplificaciones de la minigramática anterior consiste en que los nombres y los verbos que forman parte de su vocabulario terminal cum-

plen las reglas de concordancia en cuanto al género y al número. No habiendo nombres ni verbos que estén en plural, no se plantea la posibilidad de elegir un nombre en singular con un verbo en plural, ni viceversa, un nombre en plural con un verbo en singular. También se han conservado las reglas de concordancia para género y número entre artículo y nombre, puesto que se han evitado, tanto para el nombre como para el artículo, el plural y el masculino. Esto quiere decir que el contexto lingüístico en el que aparezcan los términos de nuestro vocabulario carece de relevancia a la hora de elegir un término u otro. A una gramática de este tipo se le llama gramática libre de contexto. Pero piénsese que, según es el caso en lenguas como el castellano, tuviéramos entre nuestro vocabulario terminal, junto con «la» y «una», palabras como «el», «un», «los», «unas», etc. E igualmente podríamos tener indistintamente nombres y verbos en singular y plural. Es patente que necesitaríamos, entonces, reglas que determinaran que «un» solamente puede emplearse ante un sustantivo masculino y singular, que «anarquía» únicamente puede usarse precedida de un artículo singular y femenino y seguida por verbo en singular, etc. Nuestra gramática sería ahora una gramática de las llamadas sensibles al contexto. Las reglas que tienen en cuenta el contexto de esta manera se simbolizan así: $X \rightarrow Y / V \rightarrow W$; lo que significa: X se sustituye por Y , o se reescribe como Y , cuando aparece entre V y W . A diferencia de este tipo de reglas, las que constituirían nuestra minigramática tenían simplemente esta forma: $X \rightarrow Y$, que quiere decir: sustitúyase X por Y , o reescribese X como Y . Lo anterior no significa que en el vocabulario terminal hayan de estar presentes todas las palabras del léxico en sus diferentes flexiones. Más bien tendríamos reglas contextuales que especificaran la formación del plural, las flexiones verbales, etc. Todas las oraciones que pueden generarse con una gramática libre de contexto pueden también ser generadas por una gramática sensible a éste, pero no viceversa. El poder generativo de las gramáticas sensibles al contexto es, por consiguiente, mayor que el de las gramáticas libres del mismo, y por ello se considera que éstas constituyen una subclase de las primeras; las gramáticas libres de contexto serían, en esta consideración, aquellas gramáticas sensibles al contexto en cuyas reglas del tipo citado $X \rightarrow Y / V \rightarrow W$, las variables contextuales V y W son vacuas y el contexto queda, por tanto, sin determinar.

Teniendo en cuenta que los sistemas de reescritura sin restricciones son excesivamente potentes para que resulte iluminador tomarlos como sistemas generadores de las lenguas humanas, y que las gramáticas de estados finitos tienen tan poco poder que no dan cuenta de ciertas agrupaciones estructurales que caracterizan a nuestras lenguas, puede, a la luz de lo que hemos visto, establecerse la siguiente jerarquía de sistemas generativos:

Sistemas de reescritura sin restricciones.

Gramáticas de estructura sintagmática sensibles al contexto.

Gramáticas de estructura sintagmática libres de contexto.

Gramáticas de estados finitos.

Lo propio de esta jerarquía es que, según se desciende por ella, el sistema se hace más restrictivo y disminuye su capacidad generativa. En consecuencia, todos los lenguajes que pueden ser generados por un tipo de sistemas pueden serlo también por cualquier tipo superior, pero no viceversa. De manera que, al pasar de un sistema a otro inferior en la jerarquía, se reduce la variedad de lenguajes que pueden ser generados. Por lo que respecta a los lenguajes naturales, o lenguas humanas, ya hemos visto que todas las razones inclinan a pensar que los sistemas adecuados para su generación son las gramáticas de estructura sintagmática sensibles al contexto. Ello significa que nuestras lenguas son *también* generadas por cualquier sistema de reescritura sin restricciones, pero, como antes vimos, estos sistemas carecen de interés suficiente en la medida en que, por razón de su gran potencia, generan asimismo conjuntos muy distintos de los lenguajes humanos, cosas como sistemas de ecuaciones, jugadas posibles en un juego, etc.

La gramática de estructura sintagmática sensible al contexto constituye, sin duda, un sistema generativo de alto poder que en principio parece capaz de generar todas las oraciones de lenguas, como el inglés o el castellano. Su principal inconveniente se halla en que asigna a cada oración un marcador sintagmático único de acuerdo con la estructura sintáctica de aquélla; pero no suministra ninguna explicación sobre las relaciones de carácter semántico que los hablantes perciben entre oraciones de estructura sintáctica distinta, y por consiguiente con marcadores sintagmáticos diferentes; como tampoco explica las diferencias semánticas, igualmente perceptibles, que median a veces entre oraciones con idéntico marcador de frases. Considérense los siguientes pares de oraciones del castellano:

Me preocupa que venga
Su venida me preocupa

A continuación él la besó apasionadamente
A continuación ella fue besada por él apasionadamente

Vi surcar el cielo un objeto brillante
El objeto que vi surcar el cielo era brillante

Aunque de forma diversa y con marcadores sintagmáticos distintos, es patente que las oraciones que forman cada uno de estos tres pares tienen entre sí una relación semántica tan íntima que casi diríamos que significan exactamente lo mismo y que sólo se distinguen por poner el énfasis en uno u otro elemento. Lo contrario acontece con los pares siguientes:

Le prometí estar aquí
Le recomendé estar aquí

Es una cuestión fácil de responder
Es un hombre ansioso de vivir

Aquí tenemos dos pares en los que las oraciones tienen respectivamente una estructura semejante, y por tanto un marcador sintagmático idéntico, a pesar de las profundas diferencias semánticas que separan entre sí a las oraciones de cada par, y que quedan manifiestas en las siguientes paráfrasis respectivas:

Le prometí que yo estaría aquí
Le recomendé que estuviera aquí

Es una cuestión a la cual es fácil responder
Es un hombre el cual está ansioso de vivir

Este tipo de relaciones son las que restan sin explicar en una gramática de estructura sintagmática, y lo que justifica que se haya buscado un nuevo modelo de gramática generativa que dé razón de tales características del lenguaje. Este nuevo modelo es la gramática transformatoria o transformacional cuyo desarrollo actual es obra fundamental, aunque no exclusiva, de Chomsky. (Para lo anterior véase Bach, *Teoría sintáctica*, secc. 5.1; Quesada, *La lingüística generativo-transformacional*, cap. 4; Lyons, *Chomsky*, cap. 6, y Chomsky, *Estructuras sintácticas*, caps. 4 y 5.)

4.4 El modelo chomskiano de gramática transformacional

a) *Esquema general*

Ya en *Estructuras sintácticas*, obra de 1957, Chomsky había presentado el modelo transformatorio de gramática generativa como un modelo que permite superar las insuficiencias de la gramática de estructura sintagmática dando cuenta de esas relaciones que, como acabamos de ver, pueden vincular íntimamente entre sí oraciones aparentemente distintas, o bien separar oraciones que en apariencia poseen idéntica estructura. Desde entonces hasta la fecha Chomsky y sus discípulos han continuado trabajando sobre las características e implicaciones de ese modelo, que ha recibido su formulación más típica en *Aspectos de la teoría de la sintaxis* (1965), obra en la que se completa y, para algunos aspectos, se modifica, lo que se había presentado en *Estructuras sintácticas*. Por esta razón, y aunque posteriormente no sólo algunos de sus discípulos, sino incluso el propio Chomsky, han introducido modificaciones, la versión de la gramática transformacional que vamos a ver a continuación irá básicamente referida a *Aspectos*; de las innovaciones posteriores se hará en algún caso breve mención de pasada, y de los cambios más radicales daré noticia en la sección próxima.

La idea básica de la gramática transformatoria es que no basta dar un marcador sintagmático para cada oración correcta de una lengua, por mucho que ese marcador exhiba claramente la aplicación de las oportunas reglas de reescritura y de léxico, a la manera que hemos comprobado en el

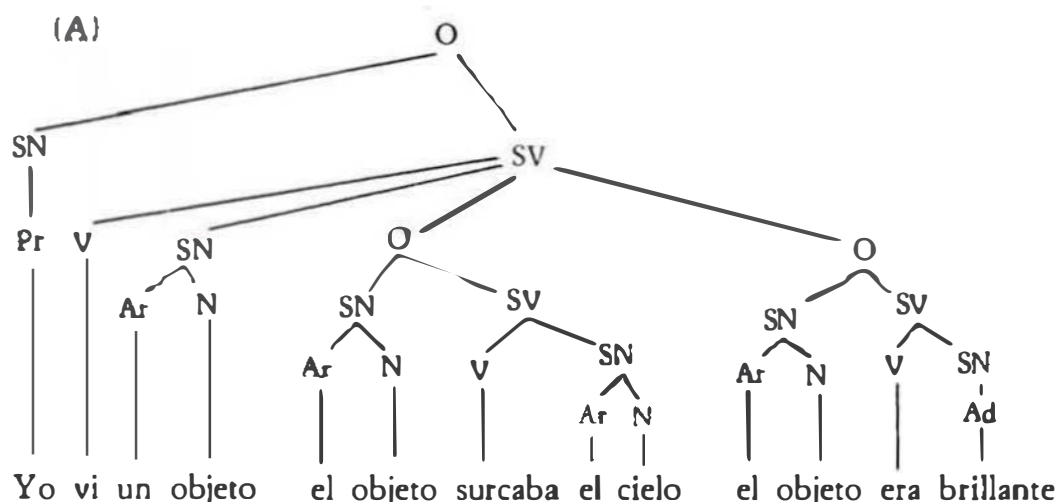
ejemplo de la sección anterior. Si queremos dar razón de las relaciones que unen o separan a oraciones como las que componen los pares que acabamos de ver, y por tanto si hemos de lograr una descripción de la lengua más cercana aún a las intuiciones del hablante, hemos de encontrar reglas que correspondan a tales relaciones, y esto supone reconocer que la forma aparente de una oración puede no ser un fiel trasunto de su auténtica estructura gramatical. Lo primero, por consiguiente, es distinguir entre el marcador sintagmático que corresponde a la apariencia de la oración y el que nos daría su estructura interna. El primero de estos dos tipos de marcador nos dará lo que Chomsky llama la estructura superficial de la oración en cuestión, mientras que el segundo nos ofrecerá su estructura profunda. Teniendo en cuenta que ambas estructuras pueden ser muy diferentes una de otra, ¿cómo se relacionan entre sí? Se relacionan de tal manera que la estructura superficial es el resultado de someter la estructura profunda a una o varias transformaciones, de acuerdo con unas reglas peculiares, inexistentes en los otros modelos de gramática generativa, que son las reglas de transformación. La gramática transformacional se caracteriza y distingue de la gramática de estructura sintagmática precisamente porque cuenta no sólo con el tipo de reglas, léxicas y de reescritura, que ya hemos visto en la gramática de estructura de frases, sino además con reglas de transformación. Son estas reglas las que explican el paso de la estructura profunda a la estructura superficial. Esta última, por tanto, procede de la primera a través de transformaciones. Y la estructura profunda, ¿de dónde procede? Esta es la directamente generada por la gramática, y su generación se explica recurriendo a reglas del tipo de las que son propias de una gramática de estructura sintagmática. De esta forma, el componente sintáctico de la gramática, que es el responsable de la estructura de la oración, queda dividido en dos partes o subcomponentes. Uno, llamado subcomponente de base, que es el propiamente generativo, y lo que genera son estructuras profundas. Este subcomponente es casi idéntico (hay algunas diferencias de detalle) a una gramática de estructura sintagmática. Y el otro, llamado subcomponente transformacional, que transforma las estructuras profundas en superficiales. Hay que consignar que, para ciertas oraciones, no interviene transformación alguna, y en tales casos, por consiguiente, coinciden la estructura profunda y la superficial. Como puede imaginarse, ello ocurre cuando se trata de oraciones de estructura muy simple, como en el ejemplo que analizamos en la sección anterior: «La anarquía aumenta la frustración.» El marcador que esbozamos para la estructura de esta oración serviría tanto para su estructura superficial como para su estructura profunda, pues entre ambas no habría diferencias apreciables. Veamos a continuación algunos ejemplos en los que sí pueden señalarse diferencias de interés entre la estructura superficial (también llamada patente) y la estructura profunda (o subyacente).

Consideremos, en primer lugar, el caso en el que oraciones con estructura superficial distinta proceden, por diferentes transformaciones, de una estructura profunda común. Esto acontecería respectivamente en los tres

pares de oraciones mencionados en primer término al final de la sección precedente. Así, las oraciones:

- (1) Vi surcar el cielo un objeto brillante
- (2) El objeto que vi surcar el cielo era brillante

tienen estructuras superficiales claramente distintas que serían exhibidas por sus correspondientes marcadores. Sin embargo, hay entre ellas una relación que, de manera informal e ingenua, podríamos formular afirmando que ambas dicen, más o menos, lo mismo. Pues bien, esta relación es la que la gramática transformacional caracteriza mostrando que ambas oraciones pueden derivarse, por medio de las oportunas transformaciones, a partir de una estructura profunda común a ambas, que podría representarse por el marcador siguiente:



Se notará que el paso de la estructura anterior a las oraciones que estamos examinando incluye transformaciones como las siguientes. Para la oración (1), la elisión del pronombre «yo», la sustitución de la oración «el objeto surcaba el cielo» por «surcar el cielo» y la sustitución de la última oración «el objeto era brillante» por «brillante» a continuación de la única intervención del sintagma «un objeto», una vez que éste ha sido colocado a continuación de «surcar el cielo». Para la oración (2), las transformaciones incluyen asimismo la elisión del pronombre «yo» y la sustitución de «el objeto que surcaba el cielo» por «surcar el cielo», añadiendo además la transformación de relativo que permite poner como sujeto principal de la oración «el objeto» y a continuación el pronombre relativo «que», para finalizar con el sintagma verbal de la última oración «era brillante», elidiendo el sintagma nominal, cuya repetición es innecesaria. Esto es una explicación totalmente superficial y abreviada de lo que ha ocurrido para que sea posible pasar de esa estructura profunda a las estructuras superficiales de (1) y (2). Explicado en términos técnicos y rigurosos sería más largo

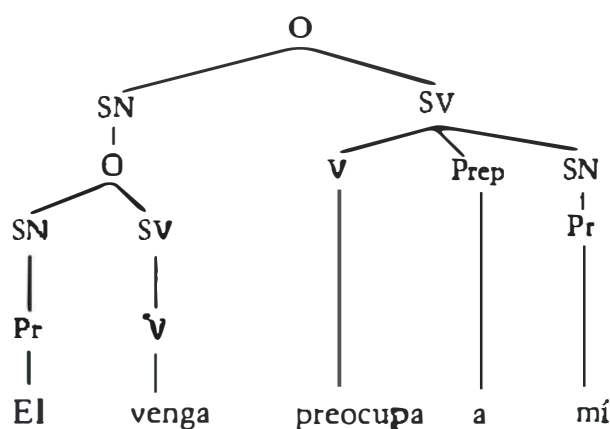
y más complicado, pero lo anterior es suficiente para dar idea de lo que implica la noción de transformación sintáctica.

Veamos aún otro ejemplo de oraciones con estructuras superficiales diversas a las que cabe atribuir una estructura profunda común. Recuerdese las oraciones citadas:

- (3) Me preocupa que venga
(4) Su venida me preocupa

A ambas podemos atribuirle la estructura profunda siguiente:

(B)



Para pasar de aquí a las oraciones (3) y (4) se aplican, por lo pronto, transformaciones como la que consiste en sustituir «a mí» por «me», anteponiéndolo al verbo principal, transformación que se aplica, como puede apreciarse, en ambas oraciones. Ulteriormente, ambas se diferencian en que para obtener (3) se ha elidido el pronombre «él», pasando el verbo de la oración que constituye el sintagma nominal, «venga», a continuación de «me preocupa», insertando entre medias «que» para formar una oración subordinada. En cambio, para obtener (4) simplemente se ha transformado el sintagma nominal de la oración principal, conservándolo en el mismo orden, es decir, delante del sintagma verbal transformado «me preocupa». Nótese que «su venida» es ambiguo entre «que venga» (hecho posible) y «que ha venido» (hecho realizado). El marcador presentado sólo es válido para la primera de estas alternativas. Para la segunda, el marcador habría de ser distinto, puesto que también lo sería la estructura profunda. En este caso, es decir, si (4) se refiriera a un hecho que ha tenido lugar, su estructura profunda tendríamos que representarla más bien por esta cadena: «El ha venido preocupa a mí.» Lo único que cambia, naturalmente, es la oración que ocupa el lugar de sintagma nominal en la oración principal, esto es, el SN más alto en el árbol. Nótese igualmente que he asumido siempre que el sujeto de «venga» es «él», pero podría igualmente ser «ella»,

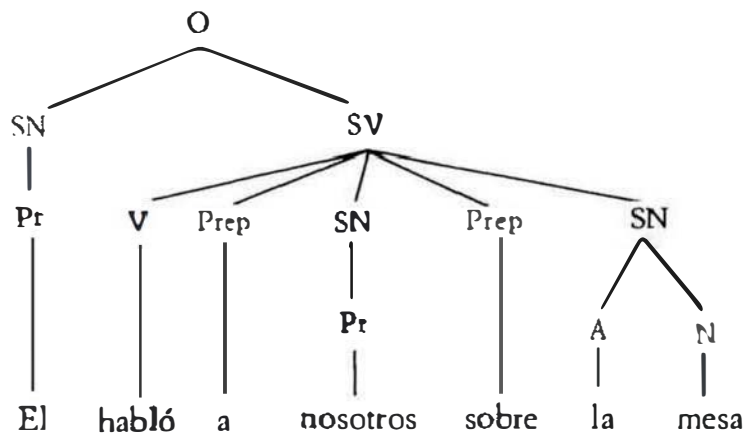
ya que a este respecto (3) y (4) son ambiguas. También, como para los ejemplos (1) y (2), subrayaré que esta breve descripción de los procesos transformatorios que conducen a (3) y (4) es técnicamente insuficiente, pero tan sólo pretende aclarar las nociones básicas de la gramática transformacional.

Acabamos de comprobar, a propósito de (4), la posibilidad de que una misma oración tenga estructuras profundas diferentes según sus diversos sentidos. Esto significa que la gramática transformacional da cuenta de la ambigüedad como una característica que puede tener una estructura superficial en la medida en que, por ser más oscura que la estructura profunda, puede no reflejar todos los aspectos de ésta. Así, la oración:

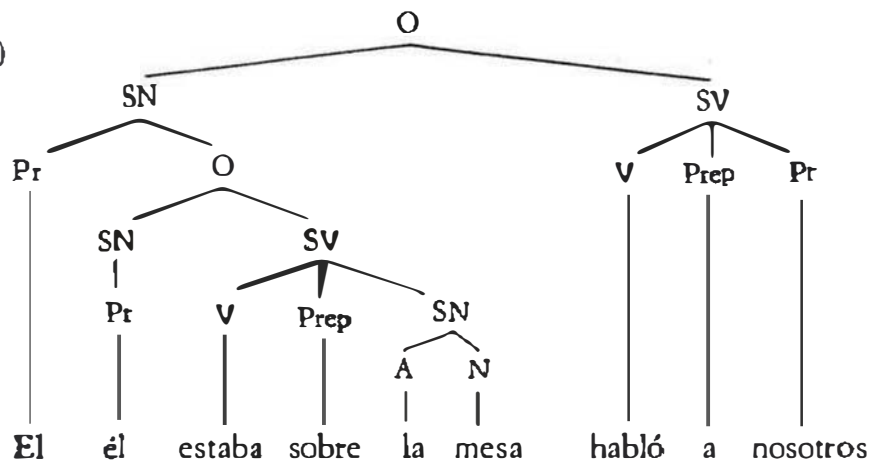
(5) Nos habló sobre la mesa

posee una ambigüedad que consiste en que a ella puede llegarse por medio de las transformaciones oportunas a partir de cualquiera de las dos estructuras profundas siguientes:

C)



(D)

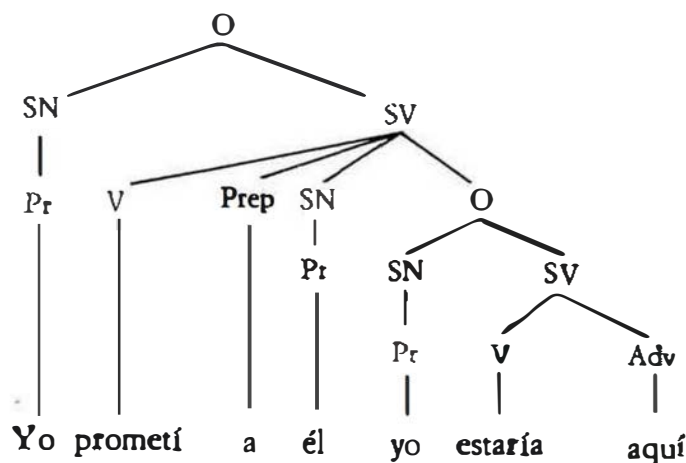


Por la misma razón que a una oración, según sus sentidos, pueden corresponder diferentes estructuras profundas, también dos oraciones con la misma estructura superficial pueden tener estructuras profundas diversas. Son ejemplo de ello los dos pares de oraciones mencionados en último lugar al final de la sección anterior. Consideremos uno de ellos, tal como el formado por las oraciones:

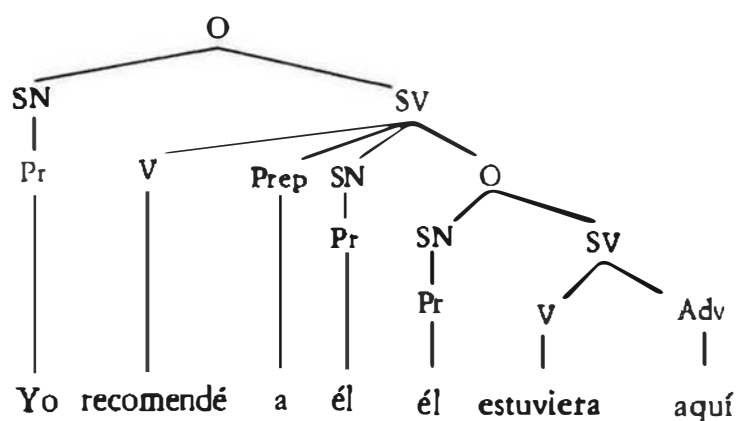
- (6) Le prometí estar aquí
(7) Le recomendé estar aquí

Sus respectivas estructuras profundas vendrán dadas por los marcadores siguientes:

(E)



(F)



Como ampliación de los ejemplos considerados citaré a continuación algunas reglas de transformación para el castellano tomadas de las que Hadlich ha estudiado en su *Gramática transformativa del español*. Mencionaré las que me parecen más claras y sencillas, y acerca de las cuales puede esperarse más fácilmente un acuerdo entre los especialistas.

1. La transformación que da lugar a una estructura pasiva a partir de la correspondiente estructura activa.
2. La elisión de los pronombres personales en función de sujeto, que desaparecen en la estructura superficial, como hemos visto en el caso de las oraciones (1) y (2): *ella llegó ayer* → *llegó ayer*.
3. La introducción del pronombre en función de objeto de interés, colocado entre el sujeto y el verbo: *ella dio la noticia a él* → *ella le dio la noticia*.
4. La sustitución de «le» por «se» cuando el objeto directo es un pronombre; así, la transformación de la estructura recién mencionada en: *ella se la dio*.
5. La introducción del pronombre reflexivo: *ella miraba ella* → *ella se miraba*.
6. La formación de estructuras negativas a partir de las correspondientes afirmativas mediante la introducción del adverbio «no».
7. La formación de cláusulas de relativo con la introducción del correspondiente pronombre.
8. La transformación de cláusulas de relativo en adjetivos: *una persona que es sensata es pacifista* → *una persona sensata es pacifista*.
9. La introducción de la conjunción «que» para formar estructuras subordinadas, del tipo de la que ya vimos: *me preocupa que venga*.
10. La introducción de la preposición «a» delante del complemento directo cuando éste es de persona.

Todo esto no es más que una breve muestra de lo que pretende explicar el subcomponente transformacional. Como ya he indicado, los marcadores de estructura profunda (A) a (F), que han servido como ilustración de las anteriores consideraciones, están muy simplificados, y en ellos se han pasado por alto problemas como los relativos a la concordancia y a la flexión, cuya manifestación en el marcador correspondiente lo haría más complejo. Piénsese, por ejemplo, que la categoría V (verbo), para estar suficientemente detallada debería ir acompañada de la mención de las características de aspecto, tiempo, etc. Pero lo más importante en cuanto a lo inadecuado de esos marcadores, es el hecho de que bajo cada una de sus categorías terminales hay una palabra. En rigor, ello es totalmente inexacto, pues lo que genera el componente sintáctico, en una gramática chomskiana, son estructuras sintácticas, y lo que resulta de la operación de las reglas transformadoras sobre éstas son otras estructuras sintácticas. Por consiguiente, no han aparecido todavía los fonemas, y esas estructuras están sin materializar, y en consecuencia tampoco pueden representarse por palabras escritas. Y aún más: tales estructuras carecen de significado.

La materialización de una estructura sintáctica en fonemas es el resultado de obrar sobre aquélla el componente fonológico. La asignación de significado a una tal estructura es, en cambio, efecto de la actuación del componente semántico. Lo característico del modelo clásico de Chomsky, presentado en *Aspectos*, reside en que ambos componentes funcionan en

nivel distinto. El componente fonológico actúa en el nivel superficial, y por la aplicación de sus reglas interpreta las estructuras superficiales asignándoles representaciones fonéticas. El componente semántico, por su parte, interpreta las estructuras profundas para asignarles representaciones semánticas, esto es, significados. El componente sintáctico es, así, generativo y transformativo, mientras que los componentes fonológico y semántico son interpretativos. Esto implica claramente que, para cualquier oración, su significado viene dado con su estructura profunda, conservándose a través de todas las transformaciones hasta la estructura superficial. Esta es una importante tesis acerca del análisis semántico de las oraciones y tiene honda semejanza con algunas teorías filosóficas sobre el significado, como en su momento comprobaremos. Se trata también de una de las primeras tesis que los discípulos de Chomsky pusieron en duda, y sobre la que éste mismo introdujo correcciones, como vamos a ver.

Examinemos brevemente el funcionamiento de cada uno de los componentes gramaticales (su esquema general está resumido en la figura de la página siguiente, que debe leerse de abajo arriba).

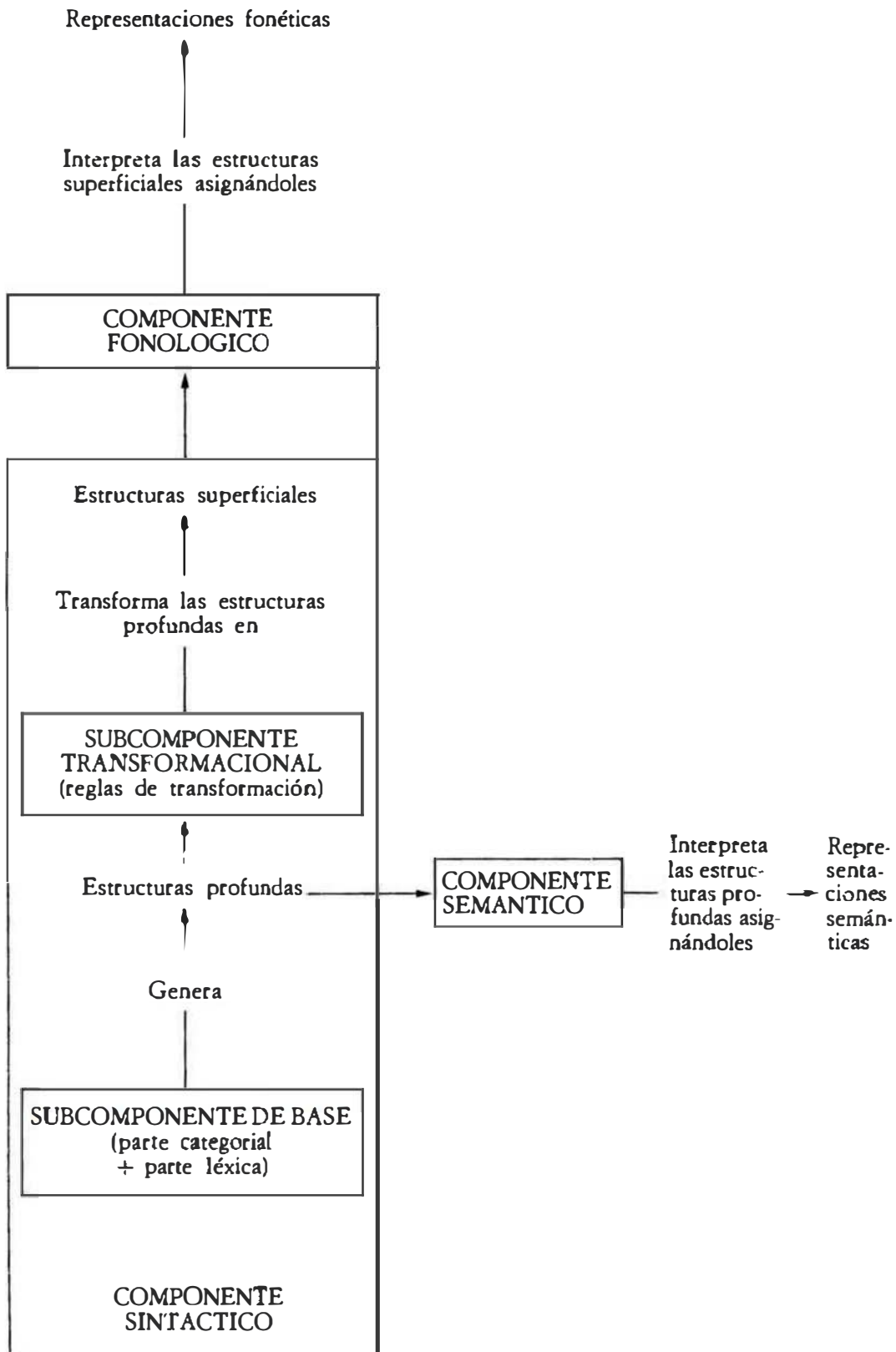
b) *El componente sintáctico*

Dentro del componente sintáctico, como ya hemos visto, se distingue entre el subcomponente de base, que genera las estructuras profundas, y el subcomponente transformativo, que convierte éstas en estructuras superficiales por la aplicación de las reglas de transformación oportunas. A su vez, dentro del subcomponente de base hay que diferenciar una parte categorial de una parte léxica. La parte categorial consiste en un conjunto de reglas de reescritura como las que son propias de una gramática de estructura sintagmática, y de las cuales se presentaron como ejemplo las reglas (1) a (6) en la sección anterior. Tales reglas pueden ser de dos tipos. Por un lado, reglas de ramificación, que son aquellas que especifican las posibilidades de sustitución entre categorías sintácticas. Son de este tipo, entre las reglas mencionadas, las reglas (1) a (3), que, como se recordará, eran:

- (1) $O \rightarrow SN + SV$
- (2) $SN \rightarrow A + N$
- (3) $SV \rightarrow V + SN$

Son, pues, reglas que establecen en términos de qué categorías ha de analizarse cada categoría, justificando así diagramas como los que hemos visto utilizados en forma de árbol. De otro lado están las reglas de subcategorización, que son las que determinan las características léxico-sintácticas de cada categoría. Estas reglas pueden representarse esquemáticamente así:

- (4) $A \rightarrow \Delta$
- (5) $N \rightarrow \Delta$
- (6) $V \rightarrow \Delta$



Estas son, por tanto, reglas que todavía no asignan a las categorías realizaciones léxicas, sino que se limitan a señalar la posibilidad de estas realizaciones, o dicho de otro modo, que se limitan a indicar los lugares en los que se insertarán las unidades léxicas. Como todavía no tenemos estas unidades (al contrario de lo que ocurría en el caso de las reglas (4), (5) y (6) de la sección anterior), el lugar que ocuparán esas unidades o realizaciones léxicas viene señalado por un símbolo comodín como el triángulo, que es el usualmente empleado. Esto significa que la parte categorial del subcomponente de base genera unas cadenas abstractas, llamadas cadenas preterminales, compuestas por comodines y por características gramaticales como «determinado» (para el artículo), «pasado», «presente», «perfecto» (para el verbo), etc. Los comodines incorporan las características léxico-sintácticas mencionadas. Si se trata del comodín para un nombre como «niño», tales características serían, por ejemplo: «común», «contable», «animado», «humano», «que no ha llegado a la adolescencia», etc.

Como puede apreciarse, la parte categorial es muy semejante a una gramática de estructura sintagmática, con la importante diferencia de que excluye el léxico. Dentro del modelo chomskiano que estamos considerando, el léxico viene dado por la otra porción del subcomponente de base, la parte léxica. Esta consiste en un conjunto de entradas léxicas, cada una de las cuales es, en rigor, una matriz de rasgos fonológicos (rasgos tales como vocálico, consonántico, palatal, bilabial, sordo, sonoro, etc.) propios de un morfema, esto es, de una unidad significativa de la lengua de que se trate. A cada entrada léxica corresponde un conjunto de rasgos sintácticos y semánticos que la caracterizan. Así, por ejemplo, a una cierta entrada le podrían corresponder los rasgos sintácticos siguientes: verbo, transitivo, con sujeto animado...; y rasgos semánticos tales como: indicativo de actividad, actividad de tipo nutritivo... Rasgos de este tipo caracterizarían a la matriz fonológica correspondiente al término castellano «comer». Es usual representar esto gráficamente de la siguiente manera:

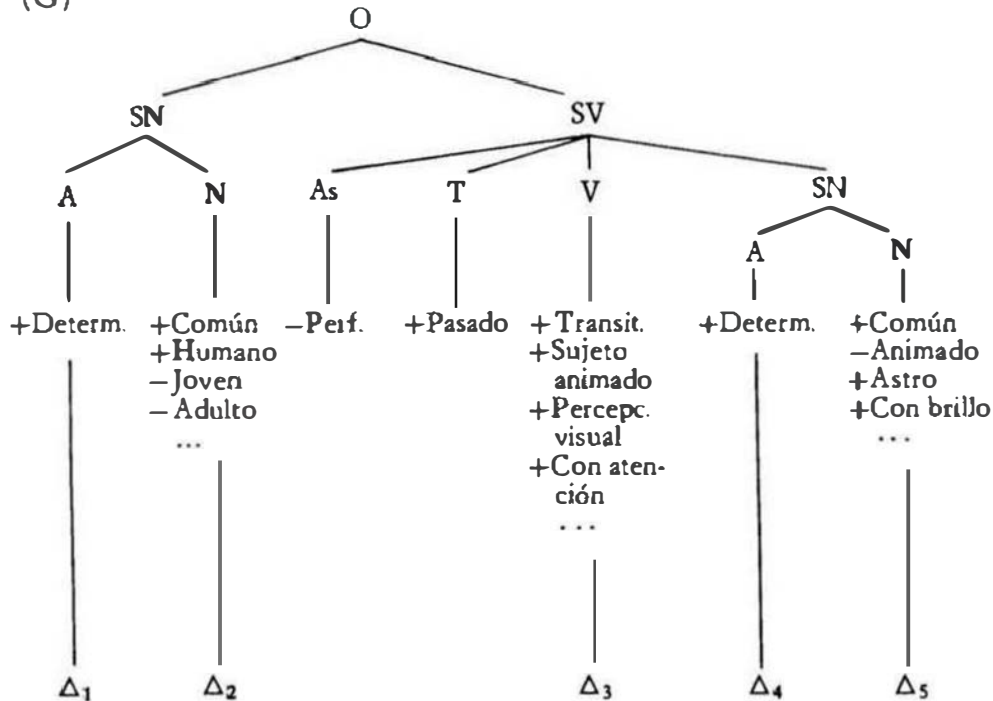
comer (+ V) (+ trans) (+ sujeto animado) (+ actividad) (+ nutrición)...

Naturalmente, esto no es exacto, en la medida en que la entrada está representada por la palabra «comer» (o mejor dicho: por su representación escrita) y no por la matriz de los rasgos fonológicos correspondientes a ella, o más rigurosamente, al morfema básico «corn-».

¿Cómo se relacionan la parte categorial y la parte léxica del subcomponente de base? Por medio de reglas llamadas de inserción léxica, las cuales regulan la sustitución de los comodines que integran las cadenas preterminales por las entradas léxicas que componen el léxico. Estas reglas permiten la sustitución de un símbolo comodín por un elemento léxico cuando las características léxico-sintácticas incorporadas en el comodín corresponden a las que caracterizan a ese elemento léxico.

Veamos un ejemplo. Considérese la oración «El niño observaba las estrellas». El fragmento categorial del subcomponente de base de la gramática castellana generaría una estructura abstracta del siguiente tipo:

(G)



Las características léxico-sintácticas que los comodines incorporan para el caso de esta oración van mencionadas debajo de cada categoría terminal, indicando los puntos suspensivos que puede haber más características, ya que las señaladas lo son sólo a modo de ejemplo y sin pretensiones de constituir un análisis completo. Así, el diagrama indica que se trata de un artículo determinado en primer lugar, seguido por un nombre común de persona humana que no es ni joven ni adulta, seguido por un verbo transitivo, de sujeto animado, que expresa percepción visual atenta, y que en esta cadena tiene además las características gramaticales de aspecto no perfecto y de tiempo pasado, seguido a su vez por un artículo determinado, al que finalmente sigue un nombre común, de objeto no animado, que es astro brillante. Se observará que este árbol es, por lo que respecta al verbo, más detallado que los que hemos examinado anteriormente, puesto que se separa el aspecto y el tiempo de lo que constituiría el propio morfema verbal. Para los efectos de lo que estamos viendo era conveniente introducir esta pequeña complicación a fin de mostrar que lo que va a sustituir al correspondiente comodín, que es el número tres, no es la matriz de algo que corresponda a «observaba», sino simplemente a «observar», o acaso más exactamente, a su raíz. Esta es también la razón por la que no han sido mencionadas las características de género y número. Según *Aspec-*

tos (cap. 4, secc. 2.2), las reglas que gobiernan la flexión y la concordancia son de tipo transformacional. Como es sólito, se ha recurrido tanto a la caracterización positiva, por medio de rasgos poseídos, que se indican con el signo «+» antepuesto, como a la caracterización negativa, que recurre a la ausencia de rasgos, aludida mediante el signo «—». Así, se ha señalado la niñez por contraposición a la juventud y a la condición de adulto, el aspecto imperfectivo por contraposición al perfectivo, y el carácter de objeto inanimado como lo contrapuesto a la condición de animado. Los sucesivos conjuntos de características léxico-sintácticas que corresponden a los elementos del marcador anterior, han sido representados por comodines en la manera que ya se comentó, numerándose éstos para evitar confusiones. Las reglas de inserción léxica autorizarán ahora la sustitución de esos comodines por aquellas entradas del léxico cuyos rasgos correspondan a los que caracterizan a aquéllos. Autorizará, por ejemplo, la sustitución de Δ_1 por la entrada que posea esos rasgos, que será la matriz fonológica propia del morfema «observar», o más exactamente de «observ-».

c) *El componente semántico*

Veamos a continuación cómo opera el componente semántico. Para ello completaremos las consideraciones de Chomsky, más escasas, con las de Katz, que habiendo dedicado más atención a esta parte de la gramática transformacional, ha acabado por establecer lo que puede considerarse como la forma canónica de la semántica chomskiana.

Como ya se ha mencionado, el componente semántico posee la función de asignar representaciones semánticas, esto es, significados, a las estructuras profundas generadas por el subcomponente de base dentro del componente sintáctico. La idea básica para construir las operaciones del componente semántico es la de que el significado del todo es función a la vez del significado de las partes y de la estructuración sintáctica en que éstas se encuentran. Por eso se mantiene que el componente semántico consta, primeramente, de un diccionario o léxico, que contiene los significados correspondientes a todos los *morfemas*, o unidades significativas (en la terminología de Martinet, *monemas*), de la lengua en cuestión. Esto hay que restringirlo a los morfemas con significado independiente, que, tal y como aparecen en un diccionario, y por tanto fonológicamente realizados y ortográficamente representados, reciben el nombre de *lexemas*, y se distinguen así de los morfemas llamados gramaticales, que no tienen significado por sí solos, sino únicamente unidos a otros. Por ejemplo, son de este tipo los prefijos, como, entre otros, los de iteración (re-) y privación (a-), las terminaciones indicativas de género y número, etc. En cambio, es un lexema del castellano el sustantivo «casa», que por cierto corresponde a un solo morfema, a diferencia, por ejemplo, de «gato», lexema en el que hay que distinguir un morfema principal y otro indicativo del género. Naturalmente, hay muchas palabras que no son propiamente lexemas, como el plural femenino de la anterior, «gatas», en la cual hay tres morfemas (principal, de

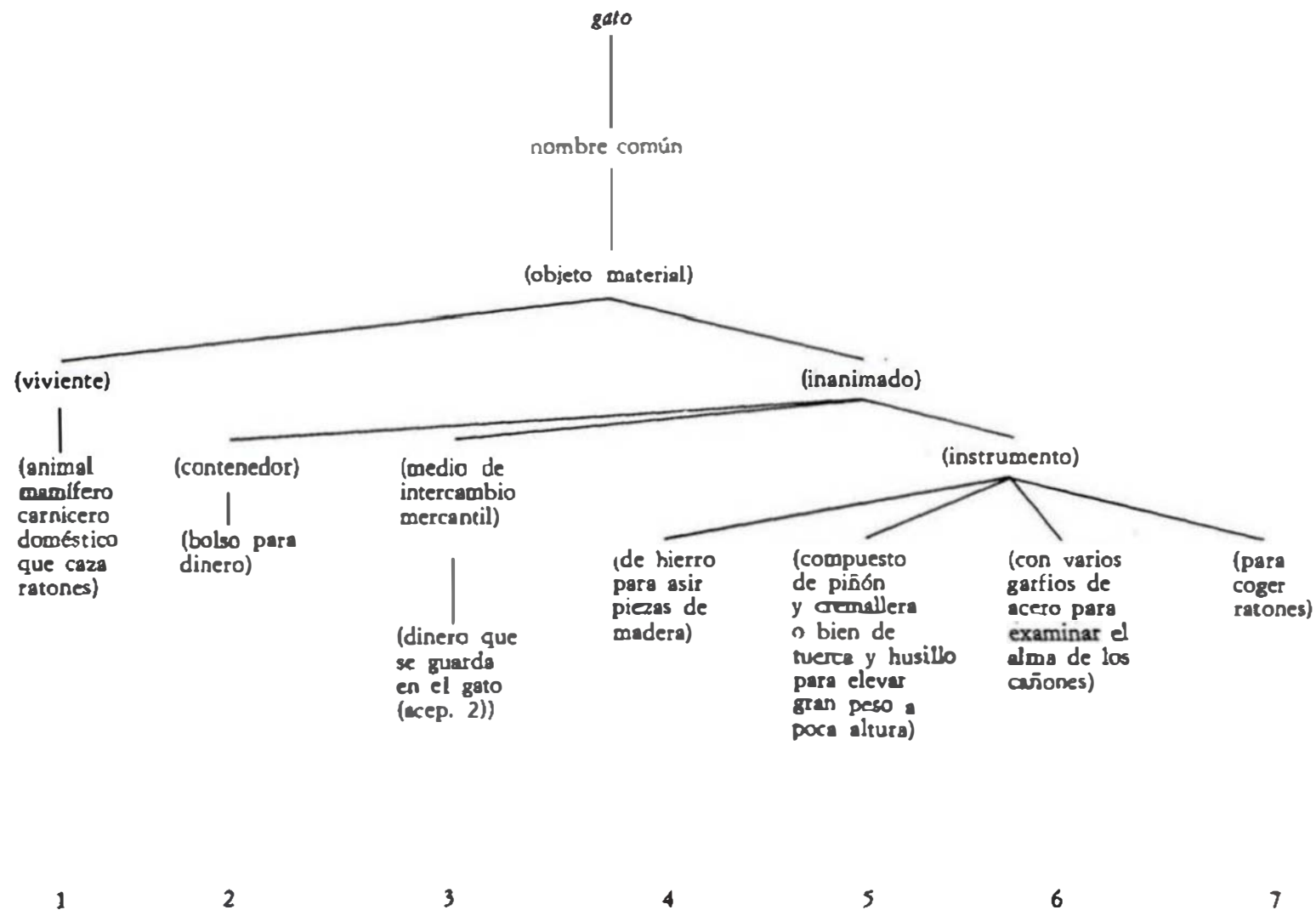
género y de número) o las diferentes formas verbales distintas del infinitivo, como, por ejemplo, «leíamos», en que también pueden distinguirse tres morfemas («le-ía-mos»). En consecuencia, lo más riguroso es llamar lexemas, o entradas léxicas, a los elementos de que se compone el diccionario que forma parte del componente semántico, sin olvidar que se trata de unidades abstractas que carecen, todavía, de realización fonética (puesto que aún no ha operado el componente fonológico). En segundo lugar, el componente semántico contiene un conjunto de reglas, llamadas reglas de proyección, que son las que explican la constitución de un significado complejo, que será últimamente el significado de una oración, a partir de los significados de los lexemas y por la composición de los mismos.

Un lexema puede ser ambiguo, esto es, puede tener varios significados distintos. Katz usa el término «sentido» (*sense*) para referirse a cada uno de tales significados, y el término «significado» (*meaning*) lo reserva para el conjunto de esos sentidos (*Semantic Theory*, p. 36). Esto no es más que una decisión terminológica en favor de la claridad, pero como se aparta decisivamente de lo que es usual en la filosofía contemporánea del lenguaje sólo me atenderé a ella mientras esté exponiendo la teoría de Katz. A la representación semántica de cualquiera de los sentidos de un lexema se le denomina en esta teoría lectura léxica, y a la de cualquiera de los sentidos de una expresión compleja, sea parte de una oración o una oración completa, lectura derivada. Las lecturas derivadas, por consiguiente, sean o no finales, esto es, correspondan o no a oraciones completas, se forman por composición de lecturas léxicas mediante la aplicación de las oportunas reglas de proyección. Cada uno de los posibles sentidos de un lexema puede a su vez ser dividido en constitutivos conceptuales que reciben el nombre de marcadores, o indicadores, semánticos. Estos pueden ser tanto simples como complejos. En el primer caso constituirán los elementos más simples, o atómicos, en los que puede descomponerse el sentido de un lexema. De aquí que pueda decirse que cualquier lectura, léxica o derivada, no es más que un conjunto de indicadores semánticos. Tales indicadores parecen estar próximos, como se ha subrayado en esta dirección, a lo que tradicionalmente se ha llamado conceptos, así como la lectura derivada para una oración parece asemejarse a lo que los filósofos han denominado proposición. La condición ontológica de los indicadores semánticos es, sin embargo, cuestión en la que Katz, obrando aquí como buen científico, prefiere no entrar (*op. cit.*, p. 39); tendremos ocasión de volver sobre esto más abajo.

En una primera exposición madura de esta teoría semántica, Katz y Fodor ofrecieron un análisis de los sentidos del término inglés *bachelor* que se ha hecho clásico y ha sido reproducido en numerosos lugares (Katz y Fodor, «La estructura de una teoría semántica»). Este análisis descomponía los cuatro sentidos que ese término posee en inglés, mostrando en qué indicadores semánticos coinciden unos sentidos con otros y en cuáles se diferencian. En la página 103 ofrezco, a modo de ejemplo castellano, un análisis semejante para el lexema «gato».

Lo primero que aparece, siguiendo el método de Katz, es un indicador gramatical que suministra la categoría sintáctica a la que pertenece la entrada léxica en cuestión. En nuestro ejemplo, es la categoría de nombre común, o sustantivo. A continuación, y encerrados entre paréntesis, al modo de Katz, vienen los distintos indicadores semánticos que caracterizan a una o varias de las acepciones o sentidos del lexema que estamos considerando. En el caso que nos ocupa hay un indicador común a todos los sentidos, el de «(objeto físico)». Esto significa que, en todas sus acepciones, «gato» designa siempre un objeto físico. A partir de este momento, los sentidos se distribuyen en dos grupos. El primero, que contiene un solo sentido, el señalado con el número uno, se opone a los demás por poseer el marcador «(viviente)», mientras que los demás tienen el marcador «(inanimado)». Como la oposición está ya marcada en este nivel, los demás indicadores o marcadores del sentido número uno, a saber: «(animal)», «(mamífero)», etc., están agrupados, y conjuntamente caracterizan el sentido que tiene «gato» cuando designa un animal de cierta clase. Los demás sentidos, del dos al siete, coinciden en la posesión del indicador «(inanimado)» y se diferencian entre sí de la siguiente manera: los sentidos cuatro a siete coinciden en el marcador «(instrumento)», que los unifica, y los opone tanto al sentido dos como al tres. Frente al indicador «(instrumento)» habría que oponer, para el sentido dos, un indicador tal como «(contenedor)», puesto que, en esta acepción, «gato» designa un contenedor de cierto tipo, a saber, un bolso o talego para guardar dinero. Como en el sentido uno, los últimos indicadores, «(bolso)», y «(para guardar dinero)», son ya indicadores que carecen de oposición en otros sentidos, y por ello se agrupan conjuntamente. El sentido tres, patentemente relacionado con el precedente, es el que tiene «gato» cuando designa precisamente el dinero que se guarda en el gato (en la acepción dos). Se me ha ocurrido señalar como indicador distintivo de este sentido y contrapuesto a «(contenedor)» por un lado, y a «(instrumento)» por otro, el indicador de «(medio de intercambio mercantil)», pensando en otros medios para tal intercambio distintos del dinero, como los talones bancarios. Por último, los sentidos cuatro a siete, que coinciden todos en el indicador «(instrumento)», se distinguirían por diversos y variados indicadores semánticos, más o menos numerosos y detallados según los casos, tal y como aparecen resumidos en la figura.

Añadiré que los sentidos enumerados de «gato» están tomados del *Diccionario de la Academia*, edición decimonovena, de 1970, y la numeración que he consignado coincide con la que ahí se hace de esos sentidos. Por cierto que la acepción número siete está calificada como en desuso en el *Diccionario ideológico* de Casares (edición de 1971). También hay que añadir que he dejado fuera, a fin de no complicar mi ejemplo innecesariamente, los sentidos que el *Diccionario de la Academia* da como figurados y familiares, y que para «gato» son: nueve, ladrón que roba con astucia; diez, hombre astuto; once, nacido en Madrid. Tampoco he tenido en cuenta el sentido que «gato» tiene en Argentina cuando designa un cierto baile popular y la música que lo acompaña.



El conjunto de todos esos sentidos de «gato» es lo que constituye, en la terminología de Katz, el significado de «gato», o como otros dirían, el *semema* de «gato» (quienes utilizan esta terminología llamarían *semus* lo que Katz denomina indicadores semánticos). Se notará que este análisis nos suministra una especie de espectro semántico para cada lexema tomado en un momento de su evolución, como en un corte sincrónico, pero que, por no contemplar el lexema en su aspecto diacrónico, en su evolución, en su historia, deja fuera relaciones internas entre sus diferentes sentidos que pueden ser muy esclarecedoras y que, en un análisis como el de Katz, tan sólo pueden conjeturarse a partir de la aparición de un mismo indicador (o de indicadores análogos) en diferentes sentidos. Así, por lo que hace a nuestro ejemplo, es claro que debe de haber una interesante relación histórica entre los sentidos uno y siete, como sin duda la hay entre los sentidos dos y tres, y probablemente entre los sentidos cuatro, cinco y seis, amén de otras posibles relaciones aún más interesantes que podrían no asomar siquiera en nuestra figura.

A pesar de todo, de ser viable, en general, para el léxico de cualquier lengua, este análisis resultaría importante y acaso nos condujera, como pretende, a una determinación más satisfactoria que las tradicionales de las categorías generales y básicas del pensamiento humano. Aunque ahora no es el momento de entrar en los detalles de esta cuestión, que procede aplazar para un momento ulterior de esta investigación cuando nos hallemos en contexto más estrictamente filosófico, conviene señalar desde ahora algunas dificultades, por lo demás obvias, que habrá de encontrar cualquier intento de llevar hasta el final la descomposición de los significados y de los sentidos en indicadores semánticos. La principal es la que toca al criterio de dicha descomposición. Aun suponiendo que el diccionario nos dé todos los sentidos que componen el significado de un lexema, el análisis de un sentido puede realizarse de diversos modos alternativos, y por tanto señalando diferentes conjuntos de indicadores semánticos, conjuntos que, en el mejor de los casos, no tienen por qué coincidir entre sí más que parcialmente. Sobre cuál de esos conjuntos sea preferible nada puede decir la gramática, y la elección entre ellos no puede sino ser consecuencia de otros factores, tales como la concepción del mundo que se tenga. Tampoco hay que olvidar que las caracterizaciones léxicas son frecuentemente incompletas, o circulares, o ambas cosas. El tema es de mayor interés filosófico de lo que por el momento puede resultar, si se tiene en cuenta que está ampliamente extendida la doctrina de que una oración analíticamente verdadera es una oración que es verdadera en función de lo que significan sus términos, y que, por lo mismo, las relaciones de deducibilidad se explican recurriendo al significado de las oraciones de que se trate. De aquí la conexión entre el análisis de los significados que la teoría de Katz propugna y lo que desde Carnap se ha llamado postulados de significado, como veremos en su momento (cfr. Barbara Partee, «Possible Worlds Semantics and Linguistic Theory»).

La insuficiencia del análisis de Katz, pero al propio tiempo su interés filosófico, se manifiesta además en el hecho de que indicadores semánticos como los sugeridos para el ejemplo anterior pueden a su vez ser analizados en un intento de llegar a los componentes últimos del significado. Así, el concepto de objeto puede analizarse como «organización de partes espacio-temporalmente contiguas que forman un todo estable con una orientación en el espacio» (*Semantic Theory*, p. 40). Dado que esta definición está ofrecida para el indicador o concepto de objeto con independencia de su cualificación como físico, es bien patente que no sirve para el concepto de objeto que se maneja cuando se habla de «objeto abstracto» o de «objeto ideal», como se hace frecuentemente en filosofía, tal cual ocurre en la teoría de los objetos de Meinong o en la semántica de Frege. Todavía para el concepto de objeto físico podría haber valido la definición de Katz, pero en el ejemplo que está considerando él separa los conceptos o indicadores «(Objeto)» y «(Físico)», y propone su análisis con exclusiva referencia al primero. Resulta digno de mención que, a estos efectos, la definición de «objeto» que puede encontrarse en el *Diccionario de la Academia* sea mucho más exacta, pues en él se da como primera acepción para «objeto», «todo lo que puede ser materia de conocimiento o sensibilidad» (el *Diccionario* de Casares dice «... de conocimiento intelectual o sensible»). Sea como fuere, si algo hay claro en todo esto es que el análisis del significado en sus elementos últimos, como el análisis del pensamiento en sus conceptos más simples, es tarea que puede realizarse de modos alternativos acerca de la elección entre los cuales ningún criterio definitivo es posible encontrar dentro del lenguaje. Señalar esto es tanto más importante cuanto que el análisis de los componentes semánticos puede utilizarse, y así es utilizado por Katz, para explicar las relaciones de deducibilidad semántica (*entailment*, que algunos traducen como «entrañamiento»). Así, de la oración «He perdido el gato» se deduce «He perdido un animal», si «gato» se tomaba en el primer sentido de los señalados antes, o bien se deduce «He perdido el bolso», si se empleaba en la segunda acepción, o bien «He perdido el dinero», si se utilizó en la tercera, etc. Naturalmente, esto tiene también una aplicación a la explicación de la verdad analítica, como hace un momento indiqué, pues si una oración analítica es una oración que es verdadera exclusivamente en función de lo que significan sus términos, el análisis riguroso del significado será un requisito previo para una efectiva determinación del ámbito de las oraciones analíticas. De esta manera, la oración «Todos los solteros están sin casar» será analítica (a pesar de los argumentos de Quine en contra, que veremos en su momento) por atribuir al sujeto, los solteros, una propiedad que constituye uno de los componentes semánticos del sentido del término «soltero». Esto invita inmediatamente a plantear la cuestión siguiente. Si eso es así, entonces la oración «Todos los gatos son animales domésticos» también será analítica, puesto que ser animal doméstico es parte del significado de «gato» (en su primer sentido). Ahora bien, mientras que en el primer caso el propio lenguaje nos impide dedicarnos con sentido a la tarea de buscar solteros casados.

en el segundo ejemplo, en cambio, no parece que así sea. ¿Por qué no podría haber algún tipo de gatos que no fueran domésticos? Ello prueba que, para el análisis semántico que se pretende, no basta cualquier definición léxica. Habría que distinguir, probablemente, entre rasgos propiamente definitorios y generalizaciones hipotéticas, al modo como se ha hecho para el propio concepto de lenguaje al comienzo de este capítulo. Distinción que sin duda se sale de los límites de la mera gramática, a pesar de lo que Katz pretende cuando discute el tema de la analiticidad en la obra citada (*Semantic Theory*, cap. 6, secc. 2; Katz distingue entre una definición léxica, como las usualmente suministradas por los diccionarios, y una definición teórica, como la que suministraría el tipo de teoría semántica que estamos viendo, pero en mi opinión una definición tal es una construcción extragramatical que sólo es posible como parte de una determinada concepción del universo; volveremos sobre este tema al hablar de Quine en un capítulo posterior).

Ya hemos visto que se denomina lectura léxica a la representación semántica de cada uno de los sentidos de un lexema, y lectura derivada a la representación semántica de cada uno de los sentidos de una expresión compleja, sea una oración completa o parte de una oración. La formación de lecturas derivadas la realiza el componente semántico combinando las lecturas léxicas oportunas mediante la aplicación de las correspondientes reglas de proyección. Estas reglas autorizan determinadas combinaciones de sentidos, como, por ejemplo, las que dan lugar al sentido de la oración:

(8) Las grandes ideas surgen raramente

Pero, al mismo tiempo, el proceso de combinación semántica debe contar con un mecanismo que evite la formación de expresiones complejas carentes de sentido, aunque formadas por expresiones más simples que sí lo tengan. Tal mecanismo es la restricción selectiva, que limita el ámbito de posibilidades combinatorias de cada sentido para formar sentidos más complejos, evitando así oraciones semánticamente anómalas como:

(9) Las verdes ideas duermen furiosamente

Es de notar que las oraciones que infringen reglas de selección, como la anterior, pueden recibir a veces una lectura metafórica (Chomsky, *Aspectos*, cap. 4, secc. 1.1). En el ejemplo anterior podría interpretarse «verdes ideas» como «pensamientos inmaduros, todavía imprecisos o vagos», y «duermen furiosamente» como «están inéditos esperando sus autores con ansia el momento de darlos a conocer». Que el lenguaje cotidiano abunda en expresiones hechas de ese tipo, tales como:

(10) El camino del infierno está empedrado de buenas intenciones

es perfectamente claro. Y puesto que estamos considerando ahora las características del componente semántico siguiendo las directrices de Katz, no

está de más recordar que, en *Aspectos*, Chomsky coloca las reglas de selección en el componente sintáctico, aunque dudando si tendrían mejor acomodo en el semántico (cap. 4, secc. 1.2; sobre esto, Juan Luis Tato, *Semántica de la metáfora*, cap. 3).

Puesto que las restricciones de selección dependen de los indicadores semánticos que integran cada sentido, ellas se relacionan con la clasificación de lexemas de la siguiente manera. Por un lado, todos aquellos lexemas a los que es aplicable la misma regla de selección para alguno de sus sentidos formarán parte de una misma clase o categoría definida por dicha restricción selectiva. Así, los términos de colores constituyen una categoría definida de esta manera. Pero, de otra parte, la misma regla delimita la clase de todas aquellas cosas que pueden ser coloreadas, que pueden tener algún color, o lo que es lo mismo, la clase de todos aquellos sentidos (de lexemas) que pueden combinarse con los sentidos (literales) de los lexemas que expresan colores.

Las restricciones de selección suelen escribirse a continuación de los indicadores semánticos y encerradas entre ángulos, a fin de distinguirlas de aquéllos, como en el ejemplo que sigue:

- verde* [Adj.] 1. (de color semejante al de la hierba fresca) <(objeto material)>
 2. (que conserva savia), <(vegetal)>
 3. (recién cortada), <(leña)>
 4. (fresca), <(legumbre)>
 5. (aún no maduro), <(fruto)>
 10. (primeros años, años de juventud), <(tiempo de la vida humana)>
 11. (al principio de su desarrollo), <(organismo), (proceso), (idea)>
 12. (indecente), <(obra artística), (discurso hablado o escrito)>
 13. (con inclinaciones galantes impropias de su edad o estado), <(humano)>

Los números remiten a las correspondientes acepciones del *Diccionario de la Academia*, de donde se han resumido, como muestra, los anteriores indicadores semánticos y restricciones selectivas, ocasionalmente completados por el autor (las acepciones 10 a 13 aparecen como sentidos figurados). De ello se obtiene que en la expresión «sillón verde», «verde» sólo puede estar tomado en el primer sentido, significando un color, puesto que las restricciones selectivas de los demás sentidos excluyen la combinación de «verde» con «sillón». Por lo mismo, si encontráramos «verde» predicado de una persona, el esquema anterior muestra que este término tendría que estar tomado solamente en alguno de los sentidos primero, oncenso o tredecimo. (y suponiendo, naturalmente, que la semántica de «verde» esté correctamente recogida en el *Diccionario de la Academia*). Nótese, por cier-

to, que la undécima acepción (que es figurada) justifica la lectura de «verdes ideas» que, como metafórica, se ha sugerido hace un momento. Siendo esto así, es discutible que esta expresión pueda ser considerada, en rigor, como anómala desde el punto de vista semántico, pero esta cuestión de la anomalía semántica es, en cualquier caso, cuestión entre las más disputadas y disputables.

La presentación de las lecturas léxicas puede simplificarse y abreviarse recurriendo a las llamadas reglas de redundancia semántica. Estas reglas permiten prescindir de una serie de indicadores semánticos cuando éstos están ya semánticamente implicados en otro indicador que aparezca en la misma lectura. Así, para las lecturas cuarta a séptima del lexema «gato», que ya vimos, podrían eliminarse los indicadores «(objeto material), (inanimado)», pues éstos están incluidos, por definición, en el indicador «(instrumento)». Por idéntica razón podrían suprimirse los indicadores «(objeto material), (viviente)» en la primera lectura, pues están implicados por el indicador «(animal)».

Así pues, la asignación de significado que el componente semántico realiza para la estructura profunda de una determinada oración, es una operación compleja que consiste en partir de los sentidos o lecturas léxicas de los elementos simples de la estructura, sentidos que vendrán dados por el diccionario, para pasar, por la aplicación de las reglas de proyección oportunas, a asignar significados o lecturas derivadas a las porciones sucesivamente más complejas de la estructura, hasta llegar a la oración completa. En la fase de partida, cuando tenemos el conjunto de las lecturas léxicas asignadas a los elementos de una estructura profunda, se dice que tenemos un indicador sintagmático subyacente interpretado léxicamente; una vez que han actuado las reglas de proyección y tenemos una lectura derivada final para la oración como tal, se dice que tenemos un indicador sintagmático subyacente interpretado semánticamente.

Hay que tener en cuenta que las reglas de proyección difieren según el tipo de relaciones sintácticas a las que corresponden, pues se aplicará una regla u otra según que entre los elementos, cuya lectura derivada se trata de producir, haya una relación u otra, por ejemplo: según que se trate de sujeto y verbo, o de verbo y objeto directo, o de nombre y adjetivo, etc. Naturalmente, la operación de las reglas de proyección comporta la aplicación de las restricciones selectivas que correspondan en cada momento, de tal manera que ciertas lecturas puedan combinarse entre sí mientras que la combinación de otras quede excluida de la manera que se ha visto en los ejemplos anteriores. Por todo ello puede afirmarse que «las reglas de proyección asocian tipos de combinación semántica con relaciones gramaticales», y que de este modo «especifican las implicaciones semánticas de las relaciones gramaticales definidas en la teoría sintáctica» (Katz, *Semantic Theory*, p. 47).

En resumen: la importancia y la significación filosófica de esta concepción de la semántica de una lengua, está en su pretensión de alcanzar, por

medio exclusivo del estudio del lenguaje, una determinación de los rasgos semánticos universales, y, por tanto, comunes a todas las lenguas, o lo que tanto vale, independientes de las peculiaridades de cada lengua particular; tarea que el propio Chomsky (*Aspectos*, cap. 4, secc. 1.3) ha formulado, con filosófica expresión, como la determinación de «el sistema de conceptos posibles». Un estudio más profundo del lenguaje, y especialmente de las relaciones de su componente semántico, habrá de contribuir decisivamente —así lo espera— a una mejor comprensión de las relaciones entre significado y conocimiento, esto es, entre los sistemas semánticos y los sistemas de creencias (*loc. cit.*, secc. 1.2).

A lo largo de las páginas precedentes he anotado marginalmente ciertas dudas sobre las posibilidades de llevar a buen fin esa tarea sin salirse de los límites de la gramática. Las mejores razones de esas dudas saldrán a la luz más adelante, y tienen que ver, en gran parte, con los problemas metodológicos que presenta el estudio de la gramática universal, y que consideraremos en el capítulo próximo. De todas formas, es de justicia subrayar que esas dudas se refieren sobre todo a la teoría de Katz, pues Chomsky ha sido siempre mucho más cauto (y también mucho más vago) acerca de las relaciones entre semántica y representación del mundo, y de hecho ha llegado a desvincularse explícitamente de la posición de Katz sobre la separabilidad de ambas (cfr. Ronat, *Conversaciones con Chomsky*, pp. 188-192).

Este problema está a la base, por cierto, de muchas de las críticas que se le han hecho a Katz. Piénsese en la diferencia de connotación que existe entre «llevar a un niño al zoológico» y «llevar a un león al zoológico» (el ejemplo está tomado de Eco, *Signo*, secc. 3.8). Lo primero connota la idea de instrucción y entretenimiento, mientras que lo segundo connota la idea de encierro. Parece claro que entender ambas frases incluye percibir estas diferencias, y en tal medida esas connotaciones deben ser parte del significado respectivo. Pues bien, se ha sugerido que el significado de cada una de esas frases, su lectura derivada, en la terminología de Katz, no parece que pueda derivarse simplemente de las lecturas léxicas de sus elementos más simples (más las reglas de proyección correspondientes) si ha de incluir la señalada connotación, y que por consiguiente aquí se entremezclan con las relaciones semánticas cuestiones sobre nuestras ideas acerca del mundo, concretamente, y en este caso, acerca de los propósitos usuales por los cuales llevamos al zoológico a los niños y a los animales salvajes. Tal vez no sea éste, sin embargo, el mejor ejemplo que puede ofrecerse de esta dificultad, pues podría responderse que una definición completa de «(parque) zoológico», o lo que es lo mismo, una descomposición completa del significado de este término, es suficiente para explicar las mencionadas diferencias de connotación según que la idea de parque zoológico se combine con la idea de llevar a un niño o con la idea de llevar a un animal salvaje. Para otros ejemplos, en cambio, esta respuesta no valdría. Las connotacio-

nes que poseen expresiones como «guateque de solteros» o «gobierno de tecnócratas», especialmente en ciertas circunstancias, escapan al análisis en componentes propios de la teoría de Katz. Con esto entramos claramente en el campo de la influencia que el contexto extralingüístico ejerce sobre la interpretación del significado de nuestras locuciones, y aquí sí que la teoría de Katz se queda corta. Su única salvación, según creo, estaría en reconocer sus propios límites y declarar que su exclusivo objeto es el significado que pudiéramos llamar «gramatical», quedando fuera cualesquiera otros aspectos semánticos que, en definitiva, siempre habrían de añadirse tomando como base aquél. Volveremos sobre esta cuestión cuando, muy posteriormente, discutamos los ingredientes de una teoría unificada del significado.

d) *El componente fonológico*

Como ya vimos antes, brevemente y de pasada, el componente fonológico es, como el semántico, de carácter interpretativo, y opera sobre las estructuras superficiales para asignarles representaciones fonéticas, esto es, para materializar esas estructuras sintácticas significativas convirtiéndolas en sonidos. No diré mucho más ahora, puesto que esta parte de la lingüística carece, hasta el momento, de consecuencias filosóficas dignas de mención. Lo que voy a indicar es a simple título de información general y para completar la visión que en esta sección he querido ofrecer del sistema gramatical.

Se recordará que ya dentro del componente sintáctico, en la parte léxica del subcomponente de base, aparecen rasgos fonológicos constituyendo las matrices en las que consisten las entradas léxicas que integran dicha parte del subcomponente citado. Dichas matrices de rasgos fonológicos corresponden fundamentalmente a los morfemas básicos, esto es, a las raíces de las palabras; por ejemplo, «com» para «comer», «comerá», «comía», etc. Esto era tradicionalmente objeto de estudio por lo que se llamaba morfología derivacional, a diferencia de la llamada morfología inflexional, que trataba de las diferentes terminaciones, tanto verbales (de número, persona, aspecto, tiempo, etc.) como de género y número para los nombres y adjetivos o de caso para ciertos pronombres. Es característico de la gramática transformatoria negar que haya dos niveles distintos de abstracción, uno morfológico y otro fonológico, y absorber ambos dentro del componente fonológico, lo que no impide que algunos lingüistas de esta escuela utilicen a veces el término «morfofonémica» para dejar claro lo que incluye este componente (así, Hadlich, *Gramática transformativa del español*, parte segunda).

Puesto que los morfemas básicos, o raíces de las palabras, forman parte del subcomponente de base, ya que constituyen el léxico, las reglas fonológicas operan propiamente para asignar representaciones fonéticas a las flexiones correspondientes (que son por su parte objeto de las oportunas

reglas de transformación). Así, por considerar un ejemplo muy simple, la sencilla oración:

Nadábamos

tiene una estructura superficial que, abreviadamente y prescindiendo del árbol usual, puede representarse así:

nad₀ — perf. + pas. + 1.^a pers. + pl

donde al morfema básico «nad-» se añade, como subíndice, la vocal temática «a» que indica la conjugación a la que pertenece, y los indicadores sintácticos de aspecto (carencia de perfectivo), tiempo (pasado), persona (primera) y número (plural). Pues bien, las reglas fonológicas correspondientes determinarán que esa estructura se realice fonéticamente por medio de la secuencia sonora que corresponde a la expresión escrita «Nadábamos». ¿Cuáles son esas reglas? En este caso, dos reglas muy simples. La que establece que el pasado imperfecto de los verbos con vocal temática «a» tiene la realización fonética «aba», y la que determina que la primera persona del plural se realiza fonéticamente como «mos». Estas reglas se denominan a veces reglas de segmentalización. Hay además otras reglas fonológicas, como son las que dan cuenta de las variaciones de la raíz, por ejemplo, en los verbos irregulares («dec-ía», «dij-e», «dig-o»), o las que explican las variaciones de pronunciación del mismo fonema según su diferente posición (por ejemplo, las diferencias de pronunciación de la «n» en: «un peso», «un gato», «un ojo»). A todos los aspectos fonéticos mencionados se los llama características segmentales, para distinguirlas de aquellas características que son independientes de la segmentación fonética, como el acento y la entonación, y que se denominan características suprasegmentales. (Para todo lo anterior puede verse la segunda parte de la obra citada de Hadlich).

Como para el componente semántico, también aquí puede procederse a un análisis de los sonidos de cada lengua particular, fonemas, en términos de unas características comunes a todas las lenguas y por tanto universales. Con la diferencia de que el análisis fonológico no presenta las dificultades que hemos visto a propósito del análisis semántico, puesto que aquí no interfieren nuestras creencias sobre el mundo, y con la peculiaridad, que también hay que señalar para evitar equívocos, de que la determinación de esas características es tarea en la que se viene trabajando desde hace tiempo y al margen de los supuestos específicos de la lingüística transformatoria, y tarea en la que se han alcanzado resultados bien establecidos sobre los que existe un acuerdo muy general entre los especialistas. La idea es que cualquier fonema de cualquier lengua puede descomponerse en unos rasgos distintivos a los que se denomina fonéticos, pues es práctica común distinguir la fonología, en cuanto estudio de los sonidos de una lengua particular, de la fonética, o investigación de los sonidos en cuanto

fenómenos físicos al margen de su funcionalidad lingüística. Los rasgos fonéticos constituyen un conjunto del cual puede encontrarse una porción o subconjunto presente en el sistema fonológico de cada lengua. Tales rasgos son por ello, y como se ha visto anteriormente en la sección segunda, universales, aunque con la limitación que se acaba de indicar (es decir: que su universalidad no significa que todos ellos estén presentes en cualquier lengua, sino que lo está únicamente algún subconjunto de los mismos).

Como ejemplificación de lo anterior se encontrará más abajo dos tablas en las que se caracterizan respectivamente los fonemas vocálicos

VOCALES

<i>Localización</i>	<i>Anteriores o palatales</i>	<i>Central</i>	<i>Posteriores o velares</i>
<i>Abertura</i>			
Altas	/i/		/u/
Medias	/e/		/o/
Bajas		/a/	

CONSONANTES

<i>Punto de articulación</i>		<i>Labiales</i>	<i>Labio- dentales</i>	<i>Inter- dentales</i>	<i>Den- tales</i>	<i>Alveo- lares</i>	<i>Pala- tales</i>	<i>Velares</i>
<i>Modo de articulación</i>								
Oclusivas	Sordas	/p/			/t/			/k/
	Sonoras	/b/			/d/			/g/
Africadas	Sordas						/tʃ/	
Fricativas	Sordas		/f/	/θ/		/s/		/x/
	Sonoras						/y/	
Nasales	Sonoras	/m/				/n/	/ɲ/	
Laterales	Sonoras					/l/	/ʎ/	
Vibrantes	Simple					/r/		
	Múltiple					/ʀ/		

y los fonemas consonánticos del castellano en términos de sus rasgos distintivos (con leves variantes, en general tendentes a la simplificación, ambas tablas están tomadas de las páginas 279 y 299 de la *Gramática española*, de Alcina y Blecha; tablas semejantes pueden verse en el lugar citado de la obra de Hadlich). El lector debe tener en cuenta que los símbolos de esas tablas no son signos del alfabeto ortográfico castellano, sino signos del alfabeto fonético recomendado por la *Revista de Filología Española* (hay un alfabeto fonético internacional levemente distinto). Como es usual, los símbolos, cuando representan fonemas (en este caso, del castellano), se escriben entre barras oblicuas; cuando representan simplemente sonidos, que pueden ser también pronunciaciones alternativas, o alófonos, de fonemas, se suelen escribir entre corchetes.

Los sonidos que representan los signos de las tablas son en muchos casos patentes por coincidir éstos con signos ortográficos del castellano. Para los casos más oscuros, las equivalencias aproximadas son las siguientes:

- /θ/ para el sonido de la «c» en «medicina».
- /tʃ/ para el sonido de la «ch».
- /ɲ/ para el sonido de la «ñ», y también para el de la «n», por ejemplo, en «concha».
- /l/ para el de la «ll».
- /k/ para el de la «c» en «médico» y para el de la «q».
- /g/ para el de la «g» en «ganar».
- /x/ para el de la «g» en «gemir» y para el de la «j».

Ahora debe quedar ya definitivamente claro qué es lo que contienen las matrices fonológicas que aparecían como entradas léxicas en el subcomponente de base. Contienen, para cada morfema básico (esto es, sin flexión ni afijos), una caracterización de cada uno de los fonemas que lo componen en términos de rasgos fonéticos distintivos.

Los rasgos pueden ser como los que se han utilizado en las tablas anteriores, o algo distintos, según la teoría fonética que se elija. A este respecto, hay que mencionar la teoría propuesta por Jakobson y otros, que reducen, como hipótesis, todos los rasgos posibles a doce oposiciones binarias, entre las que cada lengua escogería las que van a caracterizar sus propios fonemas. Esas doce oposiciones son, por tanto, universales en el sentido débil que ya vimos en la sección segunda de este capítulo. Los doce rasgos distintivos señalados por Jakobson son los siguientes (véase Jakobson y Halle, *Fundamentals of Language*; se encontrará en útil resumen en Alcina y Blecha, *Gramática española*, pp. 234 y ss., y en Lepschy, *La lingüística estructural*, cap. 6):

1. Vocálico-no vocálico. Los sonidos vocálicos se caracterizan porque en ellos sale el aire libremente a través de la cavidad bucal.

2. Consonántico-no consonántico. Es característico de los consonánticos, en cambio, el que la salida del aire encuentre alguna obstrucción.
3. Denso-difuso. En los sonidos densos predomina la cavidad bucal sobre la faríngea, como en las consonantes velares y palatales, mientras que en los difusos ocurre al revés, como es el caso de las consonantes dentales y labiales.
4. Tenso-relajado. En los sonidos tensos hay mayor difusión de energía y mayor duración de la misma.
5. Sonoro-sordo. Frente a los sonidos sordos, los sonoros se caracterizan por la vibración de las cuerdas vocales.
6. Nasal-oral. La nasalidad deriva del uso de la cavidad nasal como resonador.
7. Interrupto-continuo. Frente a los sonidos continuos, los interruptos se caracterizan por rápidas interrupciones que se acusan en el espectrograma (representación gráfica de las características acústicas de una emisión de voz) por zonas de silencio. Son interruptas las consonantes oclusivas y vibrantes.
8. Estridente-mate. Los sonidos mates se caracterizan por una onda sonora más uniforme, y son así las consonantes bilabiales y velares, mientras que, por ejemplo, las labiodentales son estridentes.
9. Recursivo-infraglotal. Los sonidos recursivos se producen por una oclusión de la glotis, que tiene como consecuencia mayor descarga de energía en menor tiempo.
10. Grave-agudo. Son graves los sonidos cuya energía se concentra en la zona baja de frecuencias (consonantes labiales y velares), y agudos aquellos cuya energía lo hace en la zona alta (dentales y palatales).
11. Bemolizado-normal. En los sonidos bemolizados se da un descenso de tono junto con aumento del resonador bucal por velarización.
12. Sostenido-normal. En cambio, en los sonidos sostenidos se da elevación del tono junto con palatalización.

Si la hipótesis de Jakobson es correcta, el sistema anterior de oposiciones es suficiente para caracterizar el sistema fonológico de cualquier lengua, en el sentido, que ya hemos visto, de que no habría ningún fonema de ninguna lengua que no fuera caracterizado en términos de algunos de los rasgos incluidos en la lista, aunque podría haber algunos de esos rasgos que carecieran de ejemplificación en una lengua determinada. Aunque se han ofrecido, de paso, algunos ejemplos del sistema fonológico castellano en la lista precedente, se encontrará resumida toda la caracterización fonética de los fonemas españoles en la siguiente tabla de Alarcos, que como puede apreciarse solamente recurre a siete de los doce rasgos distintivos enumerados (*Fonología española*, p. 179; se encontrará una tabla parecida en términos de doce rasgos en la segunda parte de la *Gramática transformativa del español* de Hladich). Nótese que, para cada oposición o pareja de rasgos, la primera denominación está representada por el signo «+», y la segunda, por el signo «—».

1. Vocal/No vocal
2. Consonante/No consonante
3. Denso/Difuso
4. Grave/Agudo
5. Nasal/Oral
6. Continuo/Interrupto
7. Sonoro (flojo)/Sordo (tenso)

	o	a	e	u	i	l	l	r	ʀ	g	x	k	ŋ	y	s	ç	m	b	f	p	n	d	θ	t
1. Vocal/No vocal	+	+	+	+	+	+	+	+	+	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-
2. Consonante/No consonante	-	-	-	-	-	+	+	+	+	+	+	+	+	+	+	+	+	+	+	+	+	+	+	+
3. Denso/Difuso	+	+	+	-	-	+	-			+	+	+	+	+	+	+	-	-	-	-	-	-	-	-
4. Grave/Agudo	+	+	-	+	-					+	+	+	-	-	-	-	+	+	+	+	-	-	-	-
5. Nasal/Oral										(-)	(-)	(-)	+	-	-	-	+	-	-	-	+	-	-	-
6. Continuo/Interrupto						+	+	-	-		+	-			+	-			+	-			+	-
7. Sonoro (flojo)/Sordo (tenso)								+	-	+	(-)	-		+	(-)	-		+	(-)	-		+	(-)	-

Por consiguiente, quedamos en que la parte léxica del subcomponente de base, dentro del componente sintáctico, contendrá una caracterización fonética de cada morfema básico en términos de una serie de rasgos distintivos universales como pueden ser los de la lista y tabla anteriores. De acuerdo con esto, la matriz fonológica de rasgos distintivos del morfema «com-», correspondiente a todas las formas del verbo «comer», sería así:

1. Vocal/no vocal	—	+	—
2. Consonante/no consonante	+	—	+
3. Denso/difuso	+	+	—
4. Grave/agudo	+	+	+
5. Nasal/oral	—		+
6. Continuo/interrumpido	—		
7. Sonoro/sordo	—		

Tabla en la que cada una de las tres columnas caracteriza sucesivamente los fonemas /k/, /o/, /m/, que por este orden componen el morfema «com-». Adviértase que para el establecimiento de la matriz fonológica de un morfema no necesitamos ni una representación ortográfica del morfema ni una representación fonológica del mismo, es decir, que podemos prescindir tanto del alfabeto ortográfico como del alfabeto fonético por lo que se refiere al morfema en cuestión, según se desprende de la tabla anterior, en la cual no aparecen para nada expresiones escritas como «com-» o [kom].

Como ya se ha indicado al comienzo de este apartado, sobre estas representaciones fonológicas subyacentes de los morfemas básicos, tomadas del léxico, operarán las reglas fonológicas de segmentalización para añadir las matrices fonológicas correspondientes a las diferentes terminaciones según los casos. Así, por ejemplo, la regla oportuna añadirá la matriz fonológica correspondiente al morfema «ía» cuando se trate de expresar la acción de comer en primera o tercera persona de singular del pasado imperfecto de indicativo. De esta forma, las reglas del componente fonológico asignan una representación fonética a las estructuras superficiales, las cuales están compuestas de elementos léxicos (correspondientes a los morfemas básicos o raíces de las palabras) que poseen una determinada matriz fonológica, más los elementos gramaticales que corresponden a las diferentes terminaciones (caso, género, número, persona, tiempo, aspecto, etc.). Tanto el sistema de representaciones fonológicas subyacentes como el conjunto de las reglas fonológicas de una lengua caracterizan, según los lingüistas de la escuela chomskiana, la competencia propia de un hablante en virtud de la cual éste puede formar y pronunciar correctamente las palabras de esa lengua, aun cuando no sea consciente de dichas reglas ni representaciones.

4.5 Desarrollos posteriores de la gramática transformacional

El modelo chomskiano clásico o estándar que acabamos de examinar ha tenido un extraordinario impacto en la teoría lingüística de los años

sesenta, tanto que, con evidente apresuramiento, se le ha considerado como un nuevo paradigma en lingüística (en el oscuro pero importante sentido que el término «paradigma» tiene desde Kuhn). Particularmente en Estados Unidos, el volumen de investigaciones y trabajos que ha provocado es muy considerable y de gran importancia, pero lo más llamativo ha sido que de la discusión desarrollada han surgido, no ya modificaciones de detalle o revisiones menores, sino formulaciones alternativas de la propia gramática transformacional que introducen modificaciones radicales en el modelo de Chomsky. O más bien deberíamos decir: en el modelo de *Aspectos*. Porque, como vamos a ver brevemente, el propio Chomsky ha acabado por aceptar una de esas enmiendas fundamentales a su modelo clásico. Aquí examinaremos, con obligada brevedad, las dos alternativas que tienen más alcance teórico, y que por cierto afectan a la manera como se entienda el funcionamiento del componente semántico y su relación con el sintáctico.

La primera alternativa es la que se conoce como «semántica generativa», y sobre ella han trabajado lingüistas como Lakoff, Ross, McCawley y Fillmore. Esta tendencia estaba ya incoada en 1965, el año de publicación de *Aspectos*, en la tesis doctoral de Lakoff, que llevaba por título *On the Nature of Syntactic Irregularity*, y que fue publicada cinco años después bajo el título *Irregularity in Syntax*. La característica fundamental de la posición que se desarrolló a partir de la tesis de Lakoff consiste en negar la clara y tajante distinción entre sintaxis y semántica, que hemos visto en el modelo chomskiano, atribuyendo a las estructuras profundas, en consonancia, una naturaleza mixta sintáctico-semántica. El nombre de «semántica generativa» alude al hecho de que, en esta concepción, la semántica es un componente tan generativo como la sintaxis, cosa patente por lo demás, puesto que se hace borrosa y cuestionable la propia distinción entre ambos componentes. El subcomponente de base del modelo clásico queda ahora sustituido por dos sistemas de reglas generativas: uno, que define la clase de las representaciones semánticas posibles, y otro, que restringe la clase de las estructuras superficiales posibles (McCawley, prólogo a la obra citada de Lakoff). De esta forma, la noción de una estructura profunda de carácter puramente sintáctico desaparece. Varias características del modelo chomskiano clásico son las que, sometidas a rigurosa crítica, han tenido una relación directa con esa desaparición y con la crisis del modelo. El caso más claro fue probablemente el problema de la inserción léxica. Como hemos visto, las reglas de inserción léxica operan en el nivel de la estructura profunda, autorizando la sustitución de esos elementos abstractos, que representábamos por medio de triángulos, por entradas del léxico. Esto significa que la inserción léxica se produce antes de cualquier transformación de la estructura profunda. Pues bien, con posterioridad al libro citado, que en este punto sigue el modelo chomskiano, Lakoff ha mantenido que esto no es cierto para una parte del léxico, cuya inserción sólo es explicable sobre el supuesto de que haya habido algunas transformaciones previas, y Lakoff ha documentado su posición con interesantes ejemplos de la lengua inglesa (puede verse su artículo «Sobre la semántica ge-

nerativa»). Otra característica que también fue criticada, y que afecta directamente a la separación entre sintaxis y semántica, es la extraña duplicidad que supone tener un léxico dentro del componente sintáctico, exactamente en el subcomponente de base, y otro dentro del componente semántico, lo que sin duda hace sospechar que la distinción entre sintaxis y semántica sea artificiosa (esta crítica puede encontrarse en Weinreich, *Explorations in Semantic Theory*).

Se ha subrayado que la semántica generativa no ha llegado nunca a constituir una escuela, y que aunque todos sus representantes compartían al menos algunas ideas básicas sobre deficiencias del modelo chomskiano como las que se acaban de aludir, sus trabajos fueron en su mayor parte independientes y les condujeron ulteriormente por caminos separados. Se ha señalado asimismo, y el propio Chomsky lo ha hecho (*Conversaciones con...*, pp. 199 y ss.) que acabaron recayendo en alguna forma de descriptivismo taxonómico, cuyo objeto serían las estructuras semánticas, resultando incapaces de suministrar una teoría alternativa rigurosa. La razón tal vez sea que ni se pretendía ni se podía formular una alternativa teórica porque la teoría seguía siendo la chomskiana, y de lo que se trataba simplemente era de llegar a una mayor claridad sobre el carácter y los requisitos del análisis semántico, pero sin salir de los límites más generales de la teoría lingüística tal y como aparecían trazados en *Aspectos*. De esta manera, se ha podido presentar la semántica generativa como una suerte de desarrollo «natural» del modelo clásico (Galmiche, *Sémantique générative*, conclusión). Merece subrayarse, también, que algunos de los estudios de semántica generativa comportan un cierto acercamiento a la lógica formal, en la medida en que se recurre al simbolismo de la lógica estándar como el mejor medio de describir las representaciones semánticas, tal y como, por ejemplo, ocurre en McCawley («¿De dónde proceden los sintagmas nominales?»), o en la medida en que se piensa que una gramática tiene como función, además de generar las oraciones de la lengua, relacionarlas con sus formas lógicas respectivas, entendiendo por tales aquellas estructuras determinadas por la *lógica natural* y en términos de las cuales pueden caracterizarse las inferencias válidas en la lengua en cuestión, como lo piensa Lakoff («Linguistics and Natural Logic»). La lógica natural así entendida es el punto de encuentro de la teoría lingüística con los últimos desarrollos de la lógica en ámbitos tales como la lógica modal, la lógica trivalente, en relación con el concepto de presuposición, etc., es decir, una lógica considerablemente más rica y más cercana al lenguaje ordinario que la clásica lógica de predicados. La lógica natural se presenta así como el estudio lógico-lingüístico de la naturaleza del lenguaje y del razonamiento.

La otra línea principal de desarrollo ha tenido también como motivo básico el problema de la semántica, y se origina igualmente hacia 1965, el año de *Aspectos*. Como ha reconocido el propio Chomsky (*Conversaciones*, p. 202), fue Jackendoff quien, por esta época, mostró que la interpretación semántica no puede tener como único objeto la estructura profunda, sino que, para muchos casos, hay que tener en cuenta también la estructura

superficial. Esta posición está ya explícitamente aceptada por Chomsky en «Estructura profunda, estructura superficial e interpretación semántica» (1968). La consiguiente revisión de su modelo clásico que ello supone ha sido bautizada por él mismo con el nombre de «teoría estándar ampliada».

Por lo que respecta a Jackendoff, la evolución de su pensamiento ha tenido un primer resultado importante y de gran influencia en su libro *Semantic Interpretation in Generative Grammar* (1972). Jackendoff cuestiona aquí la posición clásica de *Aspectos*, formulada igualmente por Katz y Postal (*An Integrated Theory of Linguistic Descriptions*, 1964), de que la única información sintáctica relevante para la interpretación semántica sea la contenida en la estructura profunda, y en consecuencia de que el significado se conserve sin cambios a través de las transformaciones. A esta idea, que hemos visto desarrollada en la sección precedente, Jackendoff opone una concepción de la gramática en la que el componente semántico opera sobre todos los estadios del proceso de transformación de las estructuras sintácticas, desde las subyacentes hasta las superficiales, para producir las representaciones semánticas. Una gramática transformacional, así entendida, consta de los siguientes componentes (*op. cit.*, secc. 10.1):

Un léxico, conteniendo todos los elementos significativos de la lengua, con sus propiedades sintácticas, semánticas y fonológicas.

Un componente de base, que contiene una gramática de estructura sintagmática, más un conjunto de reglas de inserción léxica. Estas reglas tienen, en Jackendoff, una forma algo distinta de la que presentan en *Aspectos*, pero no entraremos aquí en estos detalles.

Un componente transformatorio, del mismo tipo que el de *Aspectos*.

Un componente fonológico, que funciona también como el que ya hemos visto a propósito de Chomsky.

Y, por último, un componente semántico, que se compone de cuatro subcomponentes que corresponden a los cuatro aspectos siguientes de la representación semántica: la estructura funcional, la estructura modal, las relaciones de coreferencia, y el foco más la presuposición.

Cada uno de estos subcomponentes actúa en determinados niveles de la estructura sintáctica. El subcomponente primero deriva la estructura funcional a partir de la estructura profunda y de las propiedades semánticas de los elementos del léxico. La estructura funcional representa las relaciones que, en la oración, dependen de los verbos, e incluye nociones como agente, movimiento y dirección.

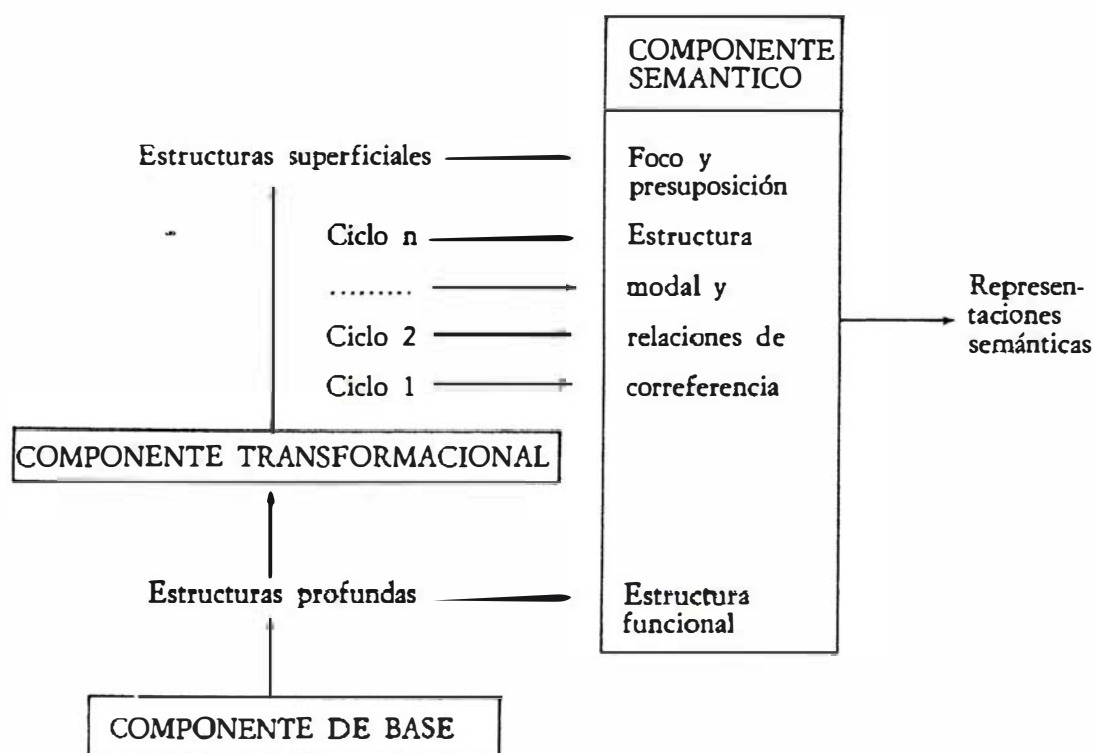
El segundo subcomponente desarrolla la estructura modal a partir de las propiedades léxicas de los operadores modales y de la configuración, tanto de las estructuras profundas como de las superficiales. La estructura modal especifica las condiciones en las cuales una oración pretende corresponder a la realidad, e incluye aspectos como la negación y los cuantificadores.

El tercer subcomponente establece las relaciones de coreferencia sobre la base de la estructura final con la que concluye cada ciclo de transfor-

maciones. Relaciones de correferencia son las que se dan entre diferentes sintagmas nominales de una oración cuando se refieren al mismo objeto.

El cuarto subcomponente opera sobre la estructura superficial para derivar lo que se llama el foco y la presuposición. El foco es la información transmitida por la oración, y que el hablante asume, que no es compartida por el oyente sino que es nueva para éste. Suele venir marcado por un énfasis en la entonación. La presuposición es aquella información, contenida en la oración, que el hablante asume que es compartida por el oyente (para lo anterior véanse las secciones 1.2 y 1.5 de la obra citada).

La idea básica es que la interpretación semántica es el resultado de un proceso de adición en el que cada regla añade más información, teniendo en cuenta la estructura sintáctica en un determinado nivel de la derivación y, en ciertos casos, también la porción de interpretación semántica ya obtenida. En contra de la tesis clásica de que el significado se mantiene a través de las transformaciones, Jackendoff afirma, por tanto, que más bien «hay que esperar que las transformaciones creen nuevas posibilidades semánticas o eliminen ambigüedades» (secc. 10.3, p. 385). De conformidad con ello, y sobre la base del diagrama ofrecido por él (p. 4), podemos representar la estructura de una gramática en los aspectos que hemos discutido, de la manera siguiente:



Por contraste con la semántica generativa, la tendencia ejemplificada por Jackendoff recibe el nombre de «semántica interpretativa», puesto que aquí el componente semántico sigue siendo interpretativo, a diferencia del

componente sintáctico, que es el único generativo, y del cual continúa distinguiéndose claramente.

Por lo que toca al propio Chomsky, éste ha pasado, de la idea de que también la estructura superficial aporta algo a la interpretación semántica, a la hipótesis de que solamente dicha estructura es el objeto de la interpretación semántica (*Reflections on Language*, 1975, p. 96). Pero ahora se trata de una noción de estructura superficial que difiere en caracteres importantes de la noción del modelo clásico. La diferencia más llamativa es la teoría de las huellas de las reglas de movimiento (*trace theory of movement rules*). Según esta teoría, cuando una transformación mueve un determinado sintagma *S* desde una posición *x* a una posición *y*, deja en la posición *x* una huella ligada a *S* (*op. cit.*, p. 95). La huella es un elemento fonológicamente nulo, y con este sentido ha podido decirse que representa una especie de memoria de la estructura profunda en la estructura superficial (*Conversaciones con Chomsky*, p. 219). Parangonando un ejemplo de Chomsky (*Reflections on Language*, p. 97), ejemplo cuya traducción castellana conserva, me parece, las características que se trata de señalar, podemos decir que solamente recurriendo a la estructura superficial es posible dar cuenta de la diferencia de significado que existe entre las dos oraciones siguientes:

- (1) Los tordos construyen nidos
- (2) Los nidos son contruidos por tordos

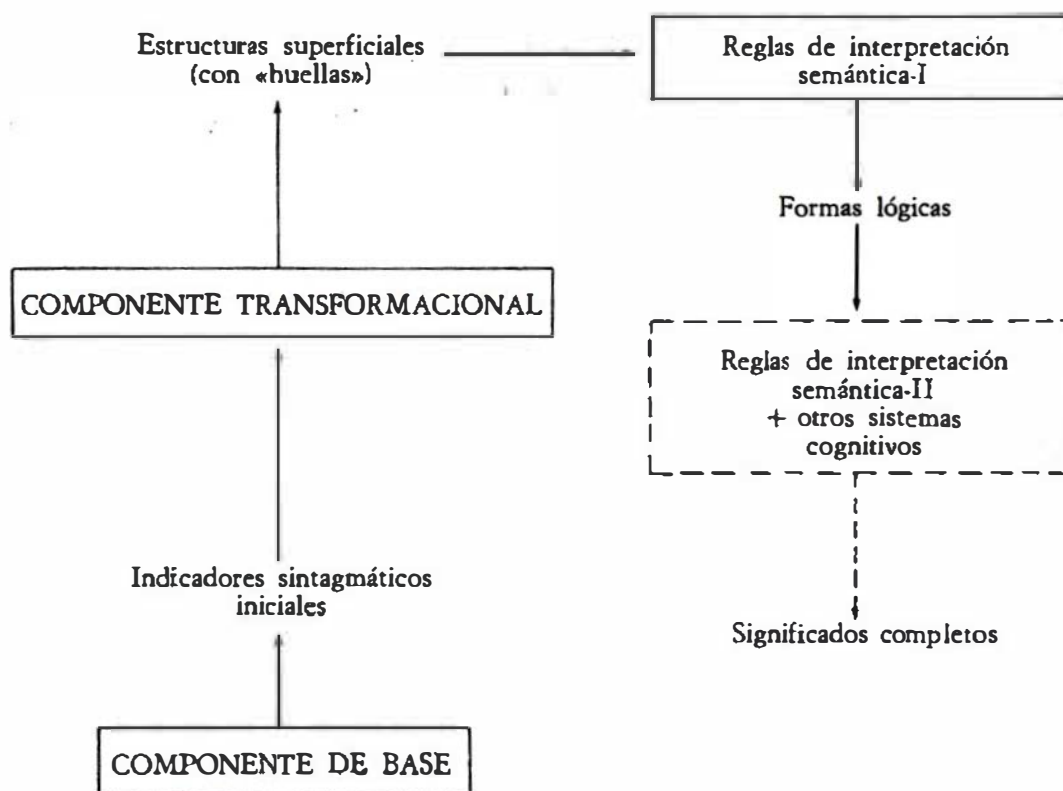
Pues mientras que la primera afirma una propiedad de los tordos en general, a saber, que construyen nidos, la segunda oración afirma, en cambio, una propiedad de los nidos en general, el ser contruidos por tordos, y es claro que las condiciones de verdad de ambas oraciones difieren, y que mientras que la primera es presumiblemente cierta, la segunda patentemente no lo es (otras muchas aves también construyen nidos). Lo fundamental, por consiguiente, es tener en cuenta que la regla de movimiento según la cual el objeto de la oración activa (1) pasa a sujeto de la oración pasiva (2) no debe encubrir el hecho de que el nuevo sujeto sigue siendo, no obstante, el objeto semántico del verbo también en la oración (2), es decir, que la relación semántica entre «nidos» y «construir» no ha cambiado. Esto es lo que, con arreglo a la teoría de las huellas, estaría indicado en la estructura superficial de (2), donde habría una huella *h* ligada por el sujeto «los nidos», huella cuya misión sería recordar que «los nidos» es el objeto semántico de «construir». Según lo cual, la estructura superficial de (2) sería:

- (3) Los nidos son [contruidos *h* por tordos]

Como habrá podido apreciarse, lo que en definitiva distingue a (1) de (2) es el tipo de cuantificación al que implícitamente están sometidos los nidos. En (1) se trata de algunos nidos (otros los construyen otras aves);

en (2), por el contrario, al pasar a sujeto gramatical, se está haciendo una afirmación general sobre todos los nidos. Chomsky ha subrayado que, en efecto, la estructura superficial parece ser relevante para la interpretación de la cuantificación, de las partículas lógicas y de otros aspectos de lo que se puede llamar «forma lógica» de la oración (*op. cit.*, p. 98).

La forma lógica por su parte no es todavía el resultado final de la interpretación semántica. Sobre la forma lógica operan otras reglas semánticas que proporcionan una representación final, más completa, del significado de la oración. Estas otras reglas semánticas, según Chomsky, pueden operar juntamente con otros sistemas cognitivos, y quedan ya fuera de lo que cabe llamar estrictamente «gramática». Partiendo del diagrama ofrecido por Chomsky (*op. cit.*, p. 105), podemos representar las relaciones entre sintaxis y semántica, según lo anterior, del modo que sigue:



He encerrado entre líneas punteadas lo que se refiere al segundo conjunto de reglas semánticas y a otros sistemas cognitivos para indicar que no pertenecen a la gramática de la oración en sentido estricto. Se notará también que lo generado por las reglas de la base no recibe aquí el nombre de «estructuras profundas» sino el de «indicadores sintagmáticos iniciales»; la noción es la misma, pero el cambio de nombre obedece, por parte de Chomsky, al propósito de evitar en lo posible ciertas connotaciones de la expresión «estructura profunda» que han inducido a algunos a atribuirle

la opinión de que tales estructuras son innatas, o la de que son idénticas para todas las lenguas (cfr. *Conversaciones...*, p. 228). Aunque ciertamente la primera de estas afirmaciones es una grave confusión, pues es claro que lo único de lo que Chomsky ha postulado un conocimiento innato son los universales lingüísticos, esto es, el contenido de la gramática universal (como veremos en el próximo capítulo), hay algo que decir, sin embargo, acerca del segundo punto. Y es que Chomsky ha solido ser poco explícito al respecto, y la manera como muchas veces ha hablado del subcomponente de base ha conducido a muchos a pensar que es una conclusión inevitable de su teoría la de que las estructuras profundas sean virtualmente idénticas para todas las lenguas. Así, Chomsky ha escrito que «en gran medida puede que las reglas de la base sean universales» (*Aspectos*, cap. 3, p. 141 del original, y en el mismo sentido, p. 117). Aunque esto hay que entenderlo referido fundamentalmente a las reglas de reescritura y de inserción léxica, si a ello se añaden las pretensiones universalistas de Katz respecto a los últimos elementos semánticos, no se ve cómo pueda evitarse concluir la universalidad de las estructuras profundas. Aun cuando Chomsky nunca la defendiera explícitamente. (y por lo que respecta a las pretensiones de Katz ya hemos visto en la sección anterior que ha acabado por desligarse de ellas tajantemente), no cabe duda que constituye como un desarrollo sistemático *natural* del modelo clásico completado con la teoría semántica de Katz.

Justamente porque Chomsky quiere subrayar ahora su apartamiento de la actitud de Katz, y dejar claro que no es posible separar nítidamente la representación semántica y el conocimiento del mundo, es por lo que el significado completo de una oración viene dado no sólo por las reglas semánticas *gramaticales*, sino además por otras reglas externas ya a la gramática de la oración (*sentence grammar*) que interactúan con sistemas cognitivos extralingüísticos. Curiosamente, lo que resulta ahora de la aplicación de las reglas semánticas gramaticales recibe la denominación de «formas lógicas». Y es curioso porque lo que anteriormente se había llamado así en la lógica y en la filosofía analítica fue muchas veces (por ejemplo, en Russell) más bien semejante a las estructuras profundas del modelo clásico chomskiano. Las formas lógicas ahora, en la teoría ampliada, son el resultado de interpretar semánticamente las estructuras superficiales, atendiendo especialmente a ciertos aspectos lógicos como la negación, la cuantificación, la referencia compartida, etc. Se habrá notado, así, una doble coincidencia muy general entre la semántica generativa y la semántica interpretativa (incluyendo aquí la teoría ampliada de Chomsky): en ambas direcciones el problema original es conseguir una hipótesis más satisfactoria sobre el funcionamiento del componente semántico, y ambas direcciones acaban, en el curso de su desarrollo, por volverse a los recursos y planteamientos a los que la lógica había venido recurriendo en su tratamiento del lenguaje *.

* Conservando básicamente el modelo anterior, Chomsky ha expuesto posteriormente su teoría con algunas modificaciones, especialmente de terminología (*Rules and Representations*, cap. 4, 1980). Aquí explica que la base genera estructuras pro-

4.6 Características formales y universalidad de la gramática

Vimos en la sección 3.4 que uno de los sentidos en los que puede entenderse la creatividad del lenguaje es en cuanto característica del sistema de la lengua, consistente en que dentro de ésta pueda formarse un número de oraciones correctas potencialmente infinito. Como allí se explicó, para que eso sea posible basta con que no existan límites a la longitud posible de una oración, de modo tal que sea siempre posible construir una oración más larga que otra dada. Y esto es lo que ocurre de hecho en las lenguas que conocemos, puesto que recursos como las diferentes maneras de coordinación y subordinación pueden aplicarse reiteradamente a cualquier oración dada por larga o compleja que sea. Que de hecho haya límites a partir de los cuales la excesiva longitud o complejidad de una oración la hacen ininteligible y, por tanto, inútil para la comunicación, constituye un rasgo empírico de nuestro uso del lenguaje que tiene que ver con ciertas propiedades del organismo humano tal y como lo conocemos, cuales son el alcance de la memoria, la cantidad de información que el cerebro puede procesar por unidad de tiempo, etc. Pero no excluye que el sistema tenga la propiedad de generar un número infinito de oraciones distintas por aplicación iterada de ciertas reglas.

Tal propiedad se denomina recursividad. De acuerdo con ello, podemos considerar una lengua como un conjunto potencialmente infinito de oraciones determinables por recursión. Como la recursividad consiste en la aplicación sucesiva de las reglas a cada resultado de la aplicación anterior, quiérese decir que, para averiguar si una determinada oración es correcta gramaticalmente, esto es, pertenece al conjunto de las oraciones de la lengua de que se trate, tendremos siempre una sucesión finita de pasos por medio de los cuales probar si pertenece o no. En otras palabras: probar que una oración es correcta, y por tanto pertenece a la lengua, equivale a reproducir el proceso de su generación a partir del sistema de esa lengua. Un procedimiento que puede descomponerse en una serie finita de pasos es un algoritmo, y el tipo más general de algoritmos es el que se conoce como «máquina de Turing», aunque no sea propiamente una máquina en el sentido material sino simplemente un sistema formal.

fundas, que denomina estructuras-D (por *deep structures*). Estas, a su vez, son convertidas, por la operación de transformaciones, en estructuras-S. Tales transformaciones pueden reducirse, según expone, a una regla general que se aplica repetidamente, y que es una regla de movimiento que establece simplemente «Muévase α », siendo α cualquier categoría sintagmática. Los movimientos regidos por esta regla dejan una huella, en el sentido que ya hemos visto. Una huella es —afirma— «un elemento real de la representación mental», que aunque carece de contenido fonético parece poseer reflejos fonéticos. El movimiento deja vacante la categoría movida, esto es, la priva de realización fonética, pero no la elimina. Las estructuras-S obtenidas por aplicación de la regla de movimiento son a su vez convertidas en estructuras superficiales por aplicación de diversos tipos de reglas, que implican, fundamentalmente, la asignación de representación fonética y la asignación de una forma lógica. Las reglas que asocian las estructuras-S con las formas fonética y lógica son las que a veces reciben el nombre de «reglas de interpretación».

Ahora bien, con respecto a una oración por cuya corrección nos preguntamos pueden ocurrir dos cosas. O bien que pueda probarse que lo es, y que por consiguiente pertenece a la lengua en cuestión. O bien que pueda probarse que no es correcta, y que por tanto no pertenece a la lengua. Si las lenguas humanas son tales que puede probarse efectivamente que las oraciones correctas de una lengua pertenecen a ella, se dice que el lenguaje humano es recursivamente enumerable. Si además pudiera probarse que las oraciones incorrectas en una lengua lo son y que no pertenecen a ella, habría que decir que el lenguaje es recursivo. Veamos con más detenimiento en qué consiste la diferencia.

Recuérdese que estamos considerando una lengua como un conjunto de oraciones (correctas, claro está). Pues bien, se dice que un conjunto es recursivamente enumerable cuando hay un procedimiento para probar que un elemento determinado pertenece al mismo; si además puede probarse que un cierto elemento no pertenece a él, dicho conjunto es recursivo. La diferencia es que, para un conjunto recursivamente enumerable, tenemos un algoritmo, por ejemplo una máquina de Turing, que especifica los elementos que son miembros del conjunto. Por ello, para un elemento determinado, si pertenece al conjunto, averiguarlo es simplemente una cuestión de tiempo; tarde o temprano nuestro algoritmo nos lo dirá. ¿Pero y si no pertenece al conjunto? Esto puede que no lo averigüemos nunca, pues nuestro algoritmo nos informa, con el debido tiempo, de que tal o cual elemento pertenece al conjunto, pero no nos dirá nunca de un elemento que *no* está en el conjunto. Esto si se trata de un conjunto recursivamente enumerable. Si el conjunto es propiamente recursivo, entonces nuestro algoritmo puede decirnos de cualquier elemento tanto que pertenece como que no pertenece al conjunto dado. Aquí es cuestión de aplicar el algoritmo durante el tiempo necesario para encontrar la prueba pertinente. Sabemos de antemano que siempre acabaremos por averiguar del elemento que nos interesa si pertenece o no al conjunto en cuestión. Puesto que para cualquier conjunto su complemento está formado por todos los elementos que no pertenecen a aquél, decir que un conjunto es recursivo equivale a decir que tanto él como su complemento son recursivamente enumerables. O lo que es lo mismo: que podemos probar para cualquier elemento que pertenece o bien al conjunto o bien a su complemento.

¿Qué significa lo anterior trasladado al caso del lenguaje? Si éste es un conjunto recursivo, significa que por medio de un procedimiento mecánico, esto es, de una sucesión finita de pasos, podremos decidir acerca de cualquier serie de signos si constituye o no una oración correcta del lenguaje. Si solamente es recursivamente enumerable, quiere decir que, por un procedimiento del tipo indicado, tan sólo podremos decidir que la serie de signos constituye una oración, pero nunca tendremos la posibilidad de decidir que no lo es. Naturalmente, tal procedimiento pertenece a la aplicación de nuestra teoría del lenguaje, es decir, de nuestra gramática, y por consiguiente las decisiones mencionadas están en principio abiertas al teórico. al

especialista. Si el comportamiento real de los hablantes correspondiera exactamente a la actuación del teórico, o en otras palabras, si se diera una estrecha correlación entre la teoría de la competencia y la teoría de la actuación, podríamos poner lo anterior en términos de esta última y afirmar que si el lenguaje es recursivo, entonces un hablante siempre será capaz de determinar si una cierta serie de signos es o no una oración de su lengua (dentro de los límites impuestos por su organismo, y, por tanto, excluyendo las secuencias excesivamente largas o complejas); mientras que si el lenguaje es nada más que recursivamente enumerable, entonces el hablante podrá determinar que una serie de signos es una oración de su lengua, si lo es, pero nunca podrá determinar que no lo es.

Parece razonable suponer que las lenguas humanas son sistemas por lo menos recursivamente enumerables. Si no lo fueran no se ve cómo podrían tener creatividad, en el sentido indicado, ni podrían ser descritas por una gramática. Tampoco resulta fácil imaginar, en este caso, cómo podría explicarse su rápido aprendizaje por parte del niño. Es probable que, además, sean sistemas recursivos. Desde luego, si la gramática adecuada para ellas es una gramática generativa de estructura sintagmática, como la que hemos visto en la sección 4.3, y que está a la base de la gramática transformatoria, entonces las lenguas humanas son recursivas, pues el lenguaje generado por un sistema de estructura sintagmática es recursivo (véase Bach, *Teoría sintáctica*, 8.3). Examinaremos algunas dudas sobre este punto un poco más abajo.

Se ha mencionado más arriba la máquina de Turing como el tipo más general de algoritmo. Veamos en qué consiste. Una máquina de Turing es un sistema formal que incluye lo siguiente (sigo aquí la exposición de Bach, *loc. cit.*):

1. Un conjunto finito de estados internos: $\{E_0, E_1, \dots, E_n\}$
2. Un alfabeto finito: $\{a_0, a_1, \dots, a_m\}$
3. Un conjunto determinado de estados iniciales tomados del conjunto de estados internos; convencionalmente se representa el estado inicial por E_0 .
4. Un conjunto determinado de estados finales tomados igualmente del conjunto de estados internos.
5. Un conjunto finito de instrucciones, siendo cada una de ellas una cuádrupla de alguna de las formas siguientes:

(a) $E_i \ a_j \ a_k \ E_l$

(b) $E_i \ a_j \ D \ E_l$

(c) $E_i \ a_j \ I \ E_l$

Cada una de estas instrucciones supone que el sistema se halla en el estado E_i ante el signo (del alfabeto) a_j . Entonces las instrucciones son respectivamente en cada uno de los tres casos: (a) pasar al estado E_l susti-

leyendo el signo a_i por el signo a_k ; (b) pasar al estado E_i moviéndose un lugar a la derecha; (c) pasar al estado E_i moviéndose un lugar a la izquierda.

Uno de los signos del alfabeto puede utilizarse para indicar un lugar en blanco y marcar con él el comienzo y el final de la expresión que está siendo computada. Para estos efectos suele utilizarse el signo «#». En cada momento del proceso el sistema o «máquina» puede ser descrito dando el estado en que se encuentra, el signo ante el que se halla y el resto de la expresión o serie de signos al que pertenece este último. Dado el curso de esta explicación se entenderá fácilmente que sea usual representar una máquina de Turing como una cinta dividida en cuadros, en cada uno de los cuales está uno de los signos de la expresión que se quiere computar (incluyendo el signo «#» al comienzo y al final); a lo largo de la cinta se mueve una cabeza lectora cuyo cambio de estado consiste en el paso de un lugar a otro a lo largo de la cinta. En rigor, esto no es más que una representación intuitiva de lo que es una máquina de Turing; un mecanismo material que funcionara de esa manera no sería, en definitiva, más que una posible realización material, entre otras muchas, de una máquina de Turing.

Cuando la máquina llega al final de la expresión y encuentra por última vez el signo «#», se para. Se dice entonces que ha aceptado o reconocido la expresión en cuestión, esto es, la cadena de signos que ha recorrido su cabeza lectora, trátase de una fórmula lógica, de una oración de una lengua, o de lo que quiera que sea. Esto nos proporciona una primera representación intuitiva, y en nuestro caso muy esquemática, de conceptos como los de computación, algoritmo, decisión mecánica, etc., que fueron definidos de una manera técnica y formal por Turing («On Computable Numbers with an Application to the *Entscheidungsproblem*» 1936-37).

Nótese que todo conjunto recursivamente enumerable puede ser definido por una máquina de Turing, y adviértase también la estrecha conexión entre un sistema de este tipo y las gramáticas generativas que hemos estudiado en la sección 4.3. Ello no es nada sorprendente si se tiene en cuenta la extrema generalidad y alcance de las máquinas de Turing. De aquí lo que se conoce como tesis de Church, que reza así: «Cualquier proceso que pueda especificarse por medio de una serie de pasos explícitos, puede ser desarrollado por una máquina de Turing o por cualquiera de los sistemas formales equivalentes a ella». Como ejemplificación de esta tesis, Bach (*loc. cit.*) ha mostrado que todo conjunto de oraciones especificable por una gramática transformatoria lo es también por una máquina de Turing (añadiendo algunas complicaciones ausentes de mi exposición; para más detalles puede verse Quesada, *La lingüística generativo-transformacional*, capítulo 6, y sobre las máquinas de Turing, Gross y Lentin, *Nociones sobre las gramáticas formales*, cap. IV; Gross, *Modelos matemáticos en lingüística*, cap. II, y la obra clásica de Davis, *Computability and Unsolvability*, cap. 1).

Esta correspondencia entre una gramática transformacional y una máquina de Turing parecería que debe ser saludada en la medida en que nos

asegura sobre la posibilidad de decidir al menos sobre las oraciones correctas de una lengua. Sin embargo, la teoría transformacional comporta también, desde el punto de vista formal, graves inconvenientes, como han mostrado los trabajos de Peters y Richtie. Un importante resultado de estos estudios es la prueba de que una gramática transformatoria como la clásica de Chomsky en *Aspectos* es demasiado potente, tan potente que hay una gramática de ese tipo absolutamente para cualquier conjunto de oraciones que sea recursivamente enumerable (Peters y Richtie, «On Restricting the Base Component of Transformational Grammars»). Teniendo en cuenta que este resultado se mantiene para cualesquiera restricciones que se introduzcan en la base, la conclusión inevitable es que el exceso de potencia se encuentra en el subcomponente transformatorio. Vale la pena mencionar que la demostración de Peters y Richtie incluye como parte central la formulación de una regla transformacional que reproduce el comportamiento de una máquina de Turing. Este exceso de poder generador en las gramáticas transformatorias significa que éstas no constituyen una teoría suficiente para delimitar el concepto de lenguaje humano.

Lo anterior tiene asimismo otra consecuencia metodológica de sumo interés. Se ha hecho mención en la sección precedente de la sugerencia de Chomsky en el sentido de que tal vez las reglas de la base sean en gran medida universales (*Aspectos*, cap. 3), y muchos lingüistas trabajando en esta dirección han asumido esta hipótesis, relegando al subcomponente transformacional las reglas peculiares de cada lengua que daban así cuenta de las diferencias entre sus estructuras superficiales y las de otras lenguas. Como vimos en la sección 4.4, el subcomponente de base contiene una serie de categorías, como oración, verbo, sintagma nominal, etc., un conjunto de reglas de reescritura, y un léxico que asigna miembros a cada una de esas categorías. Pues bien, una hipótesis sobre la universalidad de la base incluirá en primer lugar la hipótesis de que categorías como las citadas, o al menos un subconjunto de ellas, aparecen en todas las lenguas. Son lo que Chomsky llama, según comprobamos en la sección 4.2, universales sustantivos sintácticos. Una primera cuestión es cuántas de estas categorías hacen falta realmente y son irreducibles a las demás. Sobre ello no parece haber acuerdo todavía. Aunque los lingüistas suelen, por comodidad y claridad, recurrir a una variedad de ellas, ha habido numerosas propuestas de reducción. Así, Lakoff ha defendido que los verbos y los adjetivos tienen tantas características comunes que deben representarse por una sola categoría, y Bach y Lyons han mantenido, por razones semejantes, la reducción de verbos, nombres y adjetivos a una categoría única. Por su parte, y desde la perspectiva filosófica, Putnam ha mostrado que bastan las categorías de nombre y oración para derivar las categorías restantes, cosa que por lo demás ya era sabida al menos desde que en 1935 Ajdukiewicz presentó una gramática categorial, idea que ha sido reutilizada recientemente por autores como Lewis (véase Bach, *Teoría sintáctica*, 11.3; Putnam, «The 'Innateness Hypothesis'...»; y Lewis, «General Semantics»).

En segundo lugar, una hipótesis sobre la universalidad de la base habrá de incluir determinadas restricciones sobre las reglas de este subcomponente. Esas restricciones, no obstante, no pueden determinarse con cierta seguridad a menos que el subcomponente transformacional reciba limitaciones en su capacidad para reagrupar de diferentes maneras los constituyentes de la oración. En todo caso, una hipotética restricción universal sobre las reglas de la base sería la que, para oraciones transitivas, aceptara como posibles estructuras subyacentes cualquiera de las siguientes:

sujeto-verbo-objeto
 verbo-objeto-sujeto
 sujeto-objeto-verbo
 objeto-verbo-sujeto

excluyendo, en cambio:

verbo-sujeto-objeto
 objeto-sujeto-verbo

puesto que estas dos últimas estructuras rompen la unidad del sintagma verbal al introducir el sujeto entre el verbo y el objeto, lo que es contrario a los marcadores sintagmáticos que la teoría de Chomsky asigna a las estructuras profundas de este tipo (véase Bach, *loc. cit.*). Finalmente, una base presuntamente universal habría de contener también restricciones sobre la asignación de rasgos sintácticos y semánticos a las categorías del léxico que, como se recordará, forma parte del subcomponente de base en el modelo transformacional clásico.

El problema es que los estudios de Peters y Richtie conducen a la conclusión, de gran importancia metodológica, de que la hipótesis que afirma que la base de las lenguas humanas es universal no puede ser falsada por ninguno de los medios de prueba habituales en lingüística. Por tanto, carece de confirmación posible. La razón es que, existiendo una gramática transformacional para cualquier posible conjunto de oraciones que sea enumerable recursivamente, y siendo esto así para cualquier restricción que se introduzca en la base (esto es: tanto si la base es sensible al contexto, libre de contexto e incluso de estados finitos), es siempre posible construir una base universal para la gramática en cuestión. O dicho de otra forma: puesto que el exceso de potencia de las gramáticas transformatorias está en el subcomponente transformacional, se sigue lógicamente que hay una base universal a partir de la cual puede ser generado cualquier lenguaje natural por medio de una gramática transformatoria. Y, por tanto, que para cualquier lengua hay una gramática de tipo chomskiano con una base universal (véase Peters, «Why There Are Many 'Universal' Bases»). La hipótesis de la existencia de una base universal común a todas las lenguas no puede ser confirmada en ningún grado y carece, en consecuencia, de contenido empírico. En este aspecto, la situación, según Bach, no es mejor en las revisiones de la teoría estándar que vimos en la sección anterior: siendo

tanto la semántica generativa como la semántica interpretativa y la teoría ampliada de Chomsky aún más potentes que la teoría clásica de *Aspectos*, todo lo dicho es *a fortiori* aplicable a ellas.

Esta deficiencia formal y metodológica de la lingüística transformacional ha sumido a muchos en el escepticismo y ha sido utilizada a veces como una prueba del total fracaso del paradigma chomskiano en todas sus versiones. Actitud semejante es desmesurada. De hecho, hay caminos alternativos por los que los especialistas conscientes de las limitaciones señaladas se han aventurado. Fundamentalmente, las alternativas parecen ser dos. De un lado, buscar confirmaciones empíricas para la gramática que vayan más allá de los simples juicios que como hablantes nativos podemos hacer sobre la gramaticalidad o no de las presuntas oraciones de nuestra lengua. Esta vía conducirá a encontrar, en caso de éxito, razones psicológicas por las cuales los lenguajes humanos han de tener las características universales y comunes que suponemos que tienen. Esta vía nos saca de la lingüística formal para introducirnos en la psicología del lenguaje (véase Bever, «The Cognitive Basis for Linguistic Structures»). La otra alternativa, que nos mantiene dentro de la teoría formal de la gramática, y es por ello la que particularmente han proseguido los lingüistas interesados en estos problemas, consiste en buscar restricciones que limiten el poder de las transformaciones, limitando así, no la clase de las gramáticas transformatorias, que permanece igual, puesto que la teoría sigue siendo la misma, sino la clase de los lenguajes generados por dichas gramáticas. Esta restricción supone claramente un paso adelante en la tarea de definir el concepto de lenguaje humano. Se pueden encontrar ejemplos de esta posición en Bach (*Teoría sintáctica*, 11.5), en Wasow («On Constraining the Class of Transformational Languages»), y en los propios Peters y Richtie («Nonfiltering and Local-filtering Transformational Grammars»).

Toda la discusión anterior ha asumido, como es sólo entre los especialistas, que las lenguas humanas son por lo menos recursivamente enumerables, y tal vez recursivas. Por ello, antes de cerrar esta sección, tiene interés mencionar que no todos aceptan ese presupuesto, y entre las excepciones destacan un lingüista, Hockett, y un filósofo, Hintikka.

Hockett ha negado que un lenguaje natural constituya un sistema bien definido, esto es, un sistema definible por una función algorítmica del tipo de una máquina de Turing (*El estado actual de la lingüística*, caps. 3 a 5). La base de la argumentación de Hockett es que la aplicación de una función recursiva no conserva la gramaticalidad de la estructura de partida para una nueva estructura de complejidad n cuando n es un número relativamente elevado. Es decir, que aunque, en el caso más simple, tengamos una serie como la siguiente formada con la oración *A*:

A
A&A
A&A&A

cuyas primeras cadenas sean gramaticales, ello no significa que cualquier cadena de este tipo, por larga que sea, sea gramatical. Pretender que una estructura así con *A* repetida mil veces, por ejemplo, es gramatical, resulta para Hockett «empíricamente absurdo» (*op. cit.*, 5.1). Naturalmente, Hockett no quiere decir que sea posible determinar exactamente a partir de qué cadena hemos salido de la gramaticalidad; piensa, más bien, que hay una especie de restricciones flexibles (*flexible constraints*) características de los sistemas mal definidos como el lenguaje. Si esto es así, entonces no tiene sentido buscar un algoritmo que nos dé el conjunto de oraciones correctas de una lengua, como tampoco lo tiene concebir tal conjunto como un conjunto enumerable recursivamente. Si la posición contraria, que hemos considerado y aceptado en esta sección, fuera correcta, habría que mostrar, según Hockett, cómo es posible un sistema bien definido, como el lenguaje, en un sistema mal definido, como el universo físico, tarea muy difícil en apariencia. Si es Hockett quien está en lo cierto, entonces lo que resta por explicar es cómo a partir de un sistema mal definido, como el lenguaje natural, puede construirse sistemas bien definidos, como ciertos cálculos lógicos y matemáticos, tarea claramente más simple que la anterior, y que Hockett desarrolla al menos en esquema (cap. 6).

En mi opinión, la crítica de Hockett no muestra propiamente deficiencia alguna en la teoría chomskiana. Deficiencias, y serias, son las que pusieron de manifiesto Peters y Richtie (cuyos trabajos, desde luego, son posteriores al libro de Hockett). La posición de éste se apoya más bien en un cierto concepto de gramaticalidad y de lenguaje que se aparta del concepto de Chomsky. La idea de que el lenguaje es un sistema mal definido no es tanto en lo que concluye la argumentación de Hockett cuanto su punto de partida: una lengua no puede ser un sistema bien definido, esto es, un conjunto recursivamente enumerable, porque entonces tendríamos que admitir que oraciones de longitud y complejidad tan grandes que no son aptas para la comunicación pueden ser, sin embargo, gramaticales. Y esto le parece rechazable por absurdo. En otras palabras: que la gramaticalidad no es una propiedad puramente formal, sino que tiene que ver con las condiciones empíricas de nuestro uso del lenguaje. Y la cuestión es: ¿por qué preferir este concepto de gramática al chomskiano, que, por cierto, es el propio de la lingüística matemática? Las razones para elegir una u otra alternativa deberían guardar relación con su rendimiento explicativo y con la fertilidad de sus resultados, y, según creo, estos aspectos no están presentes en la argumentación de Hockett de forma apreciable. Sus consideraciones sobre la calificación de sistema mal definido que merece el universo físico, suponiendo que sean correctas, no parecen en exceso relevantes. La cuestión es que el sistema de la lengua puede abstraerse de aquellas condiciones empíricas no específicas que restringen su uso entre los hablantes. Y esta abstracción nos permite construir una teoría formal de la lengua aplicando la teoría de la recursión. Si este tipo de teoría es o no fecundo, tal vez no está aún del todo claro, pero son trabajos como los que se han citado anteriormente en esta sección los que pueden contribuir

a disipar nuestra ignorancia. Argumentos como los de Hockett son, desgraciadamente, irrelevantes.

Más interesante ha sido la aportación de Hintikka, que ha intentado mostrar, con referencia al inglés, que sus reglas no determinan un conjunto recursivamente enumerable de oraciones correctas. Hintikka ha centrado su análisis en el ejemplo del término *any*, y de aquí que su posición se conozca como «tesis sobre *any*» (*any-thesis*; véase «On the Limitations of Generative Grammar»). Su argumento ha subrayado que el intento de ofrecer una explicación generativa de la aparición del término *any* tropieza con grandes dificultades. Su tesis la enuncia así («Quantifiers in Natural Languages: Some Logical Problems», secc. 11):

«La palabra *any* es gramaticalmente aceptable en un contexto dado $X-any\ Y-Z$ si, y sólo si, sustituyendo *any* por *every* resulta una expresión gramatical que no es idéntica en significado a $X-any\ Y-Z$.»

Lo que la tesis enuncia, por tanto, es el hecho bien conocido de la imposibilidad de intercambiar *any* y *every* en los contextos lingüísticos ingleses del tipo citado sin que varíe el significado de la oración. De entre los varios ejemplos que Hintikka cita, y que se pueden encontrar en cualquier gramática del inglés, son bien claros los siguientes:

If any member contributes, I'll be happy
(Si cualquier miembro contribuye, estaré contento)
If every member contributes, I'll be happy
(Si cada miembro contribuye, estaré contento)

Es patente la diferencia de significado entre ambas oraciones. La conclusión que saca Hintikka es que nos hallamos ante un caso en el que la aceptabilidad gramatical de una oración depende de sus propiedades semánticas, y por consiguiente no es el resultado puro y simple de un proceso generativo como pretendería la teoría lingüística. Esta consideración se puede ver con mayor claridad si se tiene en cuenta que oraciones subordinadas como la primera de las dos anteriores no pueden generarse a partir de la correspondiente oración simple. Véase este otro ejemplo de Hintikka («On Any-Thesis and the Methodology of Linguistics»):

If John found any mistake of Bill's, I'll be surprised
(Si Juan encontró cualquier error de Guillermo, me sorprenderé)

Esta oración no puede haber sido generada por transformación a partir de la oración simple

John found any mistake of Bill's

pues, a diferencia de aquélla, esta última es incorrecta, ya que donde dice *any* debería utilizar *some*. Parece que la única solución es aceptar que *any* queda introducido al final de la derivación de la oración condicional, pero

no en ninguno de los estadios previos. Sin embargo, esto lo rechazarían casi todos los lingüistas generativos, pues infringe el conocido principio de la aplicación cíclica de las reglas, según el cual toda regla aplicable a una estructura compleja determinada es aplicable a cualquier estructura más simple de la que aquélla provenga. Por tanto, queda en entredicho este principio y el carácter de recursividad de las oraciones correctas que él implica. Hintikka ha intentado probar esta conclusión de forma más técnica mostrando que hay ejemplos típicos de oraciones inglesas traducibles a fórmulas de la lógica de primer orden que no son decidibles, esto es, para las cuales no hay procedimiento recursivo que pruebe su validez.

La tesis de Hintikka es importante y compleja, pero entrar en detalles nos sumergiría en una discusión técnica sobre ciertos aspectos de la lengua inglesa que no es de este lugar. Sus puntos más discutibles son, en mi opinión, la posibilidad de traducir exactamente ciertas oraciones complejas a fórmulas lógicas en la manera como él lo hace, y la justificación que pueda haber para ciertas consideraciones sintácticas y semánticas que hace a propósito de algunos de sus ejemplos, consideraciones que a veces parecen ser muy subjetivas. Quede la cuestión abierta, recordando al tiempo que algunos, como Wasow, no están convencidos de que la tesis sobre *any* sea correcta («On Constraining the Class of Transformational Languages», nota 9). Chomsky, por su parte, con una salida muy suya, ha afirmado que, incluso si el conjunto de oraciones gramaticales no es recursivamente enumerable, como mantiene Hintikka, esto únicamente muestra que dicho conjunto no está determinado por una gramática generativa, pero no muestra que los seres humanos no tengamos precisamente ese tipo de gramática representado en la mente (*Rules and Representations*, cap. 3). Uno no puede dejar de extrañarse de que la gramática que los seres humanos interiorizan y se representan mentalmente, y que, por consiguiente, forma el contenido de su competencia lingüística, sea una gramática incapaz de determinar el conjunto de las oraciones gramaticales. Como no es posible dejar de preguntarse qué pruebas puede haber de que la gramática interiorizada sea de esta suerte y no de otra distinta. La cuestión es: ¿cuándo es adecuada una gramática?

4.7 Justificación de una gramática

Una gramática, para una lengua, puede ser aceptable o adecuada, por lo pronto, en dos sentidos diferentes. Se dice que es adecuada observacionalmente cuando presenta de forma correcta los datos que por simple observación pueden obtenerse sobre tal lengua. Es decir, cuando las oraciones y expresiones que la gramática considera corresponden fielmente a lo que se puede observar, leer o escuchar, inspeccionando fragmentos diversos de dicha lengua. Es patente que, desde un punto de vista teórico, este grado de aceptabilidad o adecuación constituye un requisito sumamente débil que no otorga a la gramática interés apenas (Bach, *Teoría sintáctica*, 10.2).

Más allá de esto, una gramática es adecuada descriptivamente en la medida en que las varias descripciones estructurales que asigna a las oraciones de la lengua corresponden a las intuiciones (a menudo tácitas) de los hablantes nativos. Puesto que, como vimos en la sección 3.3, esas intuiciones forman el contenido de la competencia lingüística, puede igualmente afirmarse que una gramática es adecuada descriptivamente en tanto que corresponde a dicha competencia y que puede tomarse como una descripción de la misma (Chomsky, *Aspectos*, cap. 1, secc. 4).

Estos dos grados de justificación, observacional y descriptiva, pueden predicarse no sólo de una gramática como tal, sino también de una teoría del lenguaje, o metateoría gramatical. Puesto que una teoría así será una teoría que especifique una clase de gramáticas, la teoría será observacional y descriptivamente adecuada cuando las gramáticas definidas por ella lo sean. Pero hay un tercer grado de justificación, el más exigente y el más interesante desde el punto de vista teórico, la justificación explicativa. Por su propio carácter, la justificación explicativa tan sólo es predicable de la teoría lingüística, no siéndolo, en cambio, de una gramática como tal. Pues, en efecto, hay que decir que una teoría del lenguaje es explicativamente adecuada cuando, sobre la base de los datos de observación relativos a la lengua de que se trate, es capaz de seleccionar para ésta una gramática descriptivamente adecuada de entre aquellas posibles gramáticas que no lo son o que lo son menos. Según Chomsky (*Aspectos*, cap. 1, secc. 6), una teoría que aspire a esa justificación debe contener todo lo siguiente:

- 1.º Una enumeración de la clase de las oraciones posibles.
- 2.º Una enumeración de la clase de las descripciones estructurales posibles.
- 3.º Una enumeración de la clase de las gramáticas generativas posibles.
- 4.º La especificación de una función tal que, dados como argumentos una oración y una gramática determinadas, el valor de la función sea una cierta descripción estructural asignada a dicha oración en dicha gramática (a estos efectos, y para el caso de oraciones ambiguas, hay que entender por «oración» cada una de las interpretaciones posibles, de tal modo que, por ejemplo, una oración con triple ambigüedad tendrá tres descripciones estructurales distintas).
- 5.º La especificación de una función tal que, dada una gramática como argumento, el valor de la función sea un número entero que represente el valor de esa gramática.

Nótese que los cuatro primeros requisitos únicamente aseguran que la teoría suministre gramáticas adecuadas, y por consiguiente hacen a la teoría adecuada tan sólo en el sentido descriptivo. Para que la teoría quede justificada asimismo en el sentido explicativo ha de poseer la quinta característica, pues ésta es la función que introduce una jerarquía entre las gramáticas posibles y que, por tanto, permite seleccionar entre ellas la me-

jor, la más satisfactoria teóricamente. Se trata de una función de evaluación formal que jerarquizará las gramáticas posibles. ¿En función de qué criterios? En función de criterios que asignen a una gramática un valor más alto cuanto menos se aparte de la forma general del lenguaje y cuanto más exactamente dé cuenta de los rasgos universales del lenguaje, pues una teoría lingüística poseerá mayor alcance explicativo cuanto más clara, directa y completa sea la relación que establezca entre la gramática de una lengua y los caracteres generales del lenguaje humano, y por consiguiente cuanto menos reglas particulares y supuestos *ad hoc* haya de introducir en dicha gramática. Dada la situación de cambio y flujo en que se encuentra la lingüística transformatoria, de acuerdo con lo que hemos estudiado en las secciones precedentes, no hace falta subrayar que no se ha obtenido todavía ninguna gramática completa descriptivamente adecuada y mucho menos una teoría del lenguaje que esté explicativamente justificada.

Hasta el momento hemos considerado la justificación explicativa de una teoría lingüística como la posibilidad de seleccionar una gramática de entre todas las que la teoría autorice y que resulten descriptivamente adecuadas. En principio, esto parece ser cuestión que afecta al especialista y que se refiere a la formulación y contrastación de la teoría lingüística. Sin embargo, es característico de Chomsky y de los lingüistas vinculados a él haber extendido el planteamiento de estos problemas hasta zonas que hasta ahora habían quedado fuera de los límites estrictos de la lingüística, y más bien confinadas dentro de la psicología y de la teoría del conocimiento. Se trata de lo siguiente.

Puesto que la gramática mejor adecuada en el sentido descriptivo será aquella que más fielmente y en su totalidad describa la competencia de los hablantes, esto es, sus intuiciones, declaradas o tácitas, acerca de las características gramaticales de su lengua, justificar la selección de esta gramática frente a otras posibles es un proceso paralelo al de explicar por qué el niño que aprende su lengua materna adquiere la competencia que corresponde a esa mejor gramática y no a otra cualquiera que fuera descriptivamente adecuada. En otras palabras: la selección que el lingüista hace de una gramática como la más adecuada para la lengua que está estudiando tiene como correlato ontogenético la selección que hace el niño de esa gramática como la más adecuada para la lengua que está aprendiendo. Con esta diferencia: el lingüista obra conscientemente y de propósito; el niño actúa inconscientemente y siguiendo un proceso natural. Se apreciará fácilmente que esto es una consecuencia de cómo se ha definido la competencia lingüística. Dado que ésta constituye un conocimiento tácito, una interiorización, de la gramática de la lengua, es sencillo concluir que el niño realiza en el proceso de aprendizaje de su lengua una selección gramatical análoga a la que lleva a cabo el lingüista que está intentando formular la gramática de esa lengua. Y de aquí no hay más que un paso a defender que el problema de la justificación explicativa «es esencialmente el problema de construir una teoría de la adquisición del lenguaje» (*Aspectos*, cap. 1, secc. 4), dando así entrada en la teoría lingüística a problemas estrictos de psico-

logía. De hecho, esta conmixtión de disciplinas, que como veremos en el próximo capítulo incluye también a la teoría del conocimiento, es característica de la posición de Chomsky, y aunque acorde con la progresiva importancia que los estudios interdisciplinarios van adquiriendo, no ha contribuido, en mi opinión y por lo que veremos, a clarificar las cuestiones del lenguaje. Por lo que toca a Chomsky, el paso resulta dado con más facilidad en la medida en que se recurre a expresiones metafóricas, como cuando afirma que el niño que está aprendiendo su lengua «construye una gramática», y que para ello debe poseer como condiciones previas, primero, «una teoría lingüística que especifique la forma de la gramática», y segundo, «una estrategia para seleccionar una gramática de la forma apropiada que sea compatible con los datos lingüísticos primarios» (*loc. cit.*). Esta ambigüedad sistemática de los términos «gramática» y «teoría», que Chomsky reconoce, y en cuya virtud se aplican tanto a la explicación del lingüista como a las intuiciones lingüísticas implícitas del hablante, facilita el tránsito desde una teoría formal del lenguaje a una teoría psicológica mentalista del mismo, pero en mi opinión carece por completo de justificación.

En todo caso, y puesto que ha sido el propio Chomsky quien ha equiparado la justificación explicativa de la teoría lingüística a la tarea de suministrar como parte de ésta una teoría de la adquisición del lenguaje, ello nos suministra excelente excusa para abandonar la teoría de la gramática y pasar a la epistemología del lenguaje.

Lecturas

La mejor obra que conozco en castellano de introducción a la teoría de la gramática es la *Teoría sintáctica* de Bach (Anagrama, Barcelona, 1976). Los primeros capítulos del libro de Quesada, *La lingüística generativo-transformacional: supuestos e implicaciones* (Alianza Universidad, Madrid, 1974), proporcionan una inestimable ayuda para entrar en estos temas, y se desarrollan en un nivel más técnico que el propio de mi exposición.

Por lo que toca a Chomsky, su mejor obra, con mucho, me parece que sigue siendo *Aspectos de la teoría de la sintaxis* (Aguilar, Madrid, 1970), y aunque no ofrezca su última palabra sobre la cuestión, por su valor clásico y su amplio influjo todavía constituye el paradigma de la teoría transformacional. Para una versión más reciente de su doctrina puede verse *Reflexiones sobre el lenguaje*, cuya edición en castellano está anunciada cuando escribo esto (Ariel, Barcelona). En su más temprana obra, *Estructuras sintácticas* (Siglo XXI, México, 1957), se encontrarán interesantes consideraciones generales sobre las gramáticas generativas. Hay, además, en castellano, dos útiles resúmenes de la teoría chomskiana. Uno, más técnico, del que son autores Chomsky y Miller, es *El análisis formal de los lenguajes naturales* (Alberto Corazón Editor, Madrid, 1972); el otro, más sencillo e informal, es el artículo de Chomsky «La naturaleza formal del lenguaje», apéndice a la obra de Lenneberg, *Fundamentos biológicos del*

lenguaje (Alianza Universidad, 1975), y que se encuentra recogido también en la recopilación de Gracia, *Presentación del lenguaje* (Taurus, Madrid, 1972). Es de gran ayuda la pequeña obra, ya citada en otro capítulo anterior, *Chomsky*, de Lyons (Grijalbo, Barcelona, 1974).

Para la evolución posterior de la lingüística transformacional, y especialmente para la concepción generativista de la semántica, puede verse una útil monografía de Galmiche, *La semántica generativa*, cuya traducción castellana está anunciada (Gredos, Madrid). Algunos de los artículos comentados por Galmiche, junto con otros muchos igualmente representativos de las diferentes polémicas que han tenido lugar en este contexto teórico, se encuentran recogidos en la extensa y cuidada recopilación de Sánchez de Zavala, *Semántica y sintaxis en la lingüística transformatoria*, en dos volúmenes (Alianza Universidad, Madrid, 1974 y 1976, respectivamente).

Para una aplicación de la gramática transformatoria al castellano, el único libro de conjunto que conozco es la *Gramática transformativa del español*, de Hadlich (Gredos, Madrid, 1975), obra breve y, al decir de los especialistas, muy insuficiente. Trabajos monográficos de diferente alcance e interés se encontrarán en la recopilación de Sánchez de Zavala, *Estudios de gramática generativa* (Labor, Barcelona, 1976), así como en los libros de María Luisa Rivero, *Estudios de gramática generativa del español* (Cátedra, Madrid, 1977), y de Violeta Demonte, *La subordinación sustantiva* (también en Cátedra, Madrid, 1977; ambos como primeros títulos de la colección «Gramática Generativa Transformacional del Español»). El reciente libro de Francesco D'Introno, *Sintaxis transformacional del español* (Cátedra, Madrid, 1979), contiene un pormenorizado estudio de muy diversos aspectos de la sintaxis castellana desde el punto de vista chomskiano.

Es también de reciente aparición en castellano la *Teoría semántica* de Katz (Aguilar, Madrid, 1979), que, aunque no coincida en todo con las últimas posiciones de Chomsky, constituye el producto más completo de la escuela chomskiana en materia semántica.

Capítulo 5

IDEAS NONNATAS

Todo cuanto los hombres hemos inventado cabe en nuestras palabras y aun nos sobran. (TORRENTE BALLESTER, *Fragmentos de apocalipsis*.)

5.1 La teoría chomskiana sobre la adquisición del lenguaje

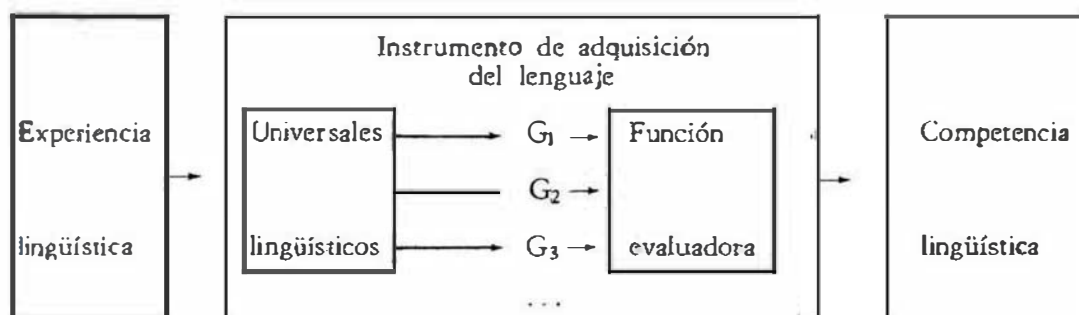
Chomsky parte de la base de que existe una gran diferencia cuantitativa y cualitativa entre el contenido de la experiencia lingüística del niño y el contenido de lo que resulta de su aprendizaje del lenguaje, a saber, su competencia lingüística (*El lenguaje y el entendimiento*, cap. 2). Los datos con los que el niño cuenta son de número y variedad muy reducidos —piensa Chomsky—; el niño sólo tiene ocasión de escuchar a unas pocas personas, que son las que usualmente le rodean, y por consiguiente sólo puede haber tenido experiencia de unos pocos tipos de expresiones, que en su mayoría se repiten. De otra parte, el discurso hablado ordinario se caracteriza por su pobre calidad sintáctica y semántica: se dan cortes, construcciones incorrectas, confusiones de significado, etc. ¿Cómo puede explicarse que sobre base tan pobre el niño adquiriera una capacidad tan rica y compleja como la capacidad de producir y comprender todas las posibles oraciones correctas de su lengua que son potencialmente infinitas? En otras palabras: ¿cómo puede explicarse que con una experiencia lingüística tan reducida se adquiriera la competencia lingüística? Para Chomsky, sólo hay una respuesta: el secreto está en lo que el niño aporta. El niño debe contar con recursos que le permitan construir la gramática de su lengua con tan escasos datos, y puesto que no se conoce que exista más facilidad o mayor predisposición para aprender una lengua que otras, tales recursos deben capacitarle para adquirir indistintamente cualquier lengua y deben ser compatibles con todas las lenguas posibles. Los recursos citados constituyen en definitiva un cierto instrumento para adquirir el lenguaje, y habrán de incluir, como se ha visto al final del capítulo anterior, y empleando tér-

minos que más parecen metafóricos que otra cosa, tanto una teoría lingüística que especifique la forma de la gramática de cualquier lengua humana posible como una estrategia para seleccionar una gramática de la forma apropiada y que sea compatible con esos datos lingüísticos originales y primarios (*Aspectos*, cap. 1, secc. 4). Considerado con más detalle, el instrumento de adquisición del lenguaje habrá de poseer un contenido que sea el correlato de lo que contenga una teoría lingüística que aspire a la justificación explicativa. La razón —como vimos en la sección 4.7— es el paralelismo que Chomsky asume entre el lingüista que está intentando formular la gramática de una lengua y el niño que se halla en el proceso de aprenderla. Si se comparan los cinco puntos enumerados en aquel lugar con los cinco aspectos siguientes del instrumento de adquisición del lenguaje, se verá que corresponden uno a uno. Este instrumento debe, en suma, contener lo siguiente (*Aspectos*, cap. 1, secc. 6):

- 1.º Una técnica para representar las señales recibidas.
- 2.º Una manera de representar la información estructural sobre esas señales.
- 3.º Alguna delimitación inicial de una clase de posibles hipótesis sobre la estructura del lenguaje.
- 4.º Un método para determinar lo que cada una de tales hipótesis implica con respecto a cada oración.
- 5.º Un método para seleccionar una de las posibles hipótesis permitidas según el punto tercero y que sean compatibles con los datos lingüísticos primarios (puede presumirse que las hipótesis así permitidas son infinitas).

Todo esto es contenido del instrumento de adquisición del lenguaje y, por consiguiente, todo esto lo posee, al menos tácitamente, el niño. En esto consiste su aportación al aprendizaje del lenguaje. Y en ello queda esquematizado el proceso de este aprendizaje. Se supone que el niño considerará, inconscientemente, las diferentes hipótesis aceptables sobre la estructura del lenguaje humano (o, si son infinitas, algunas de ellas); de éstas él seleccionará las que sean compatibles con los datos lingüísticos primarios tal y como vienen representados según los requisitos primero y segundo, determinando, de acuerdo con el punto cuarto, las consecuencias de cada hipótesis para los datos en cuestión. De las hipótesis así seleccionadas, las cuales constituyen gramáticas potenciales para la lengua que está aprendiendo, seleccionará finalmente la mejor aplicando la medida de evaluación con la que cuenta según el punto quinto. El niño ha descubierto, inconscientemente, la gramática de su lengua. O dicho de otra manera: ha adquirido su lengua nativa. Y puesto que lo aportado por el niño es indiferente a la lengua que esté aprendiendo y vale para cualquier lenguaje posible, hay que concluir que es común a todas las lenguas. Son los requisitos mínimos que todas ellas han de cumplir y en los que todas ellas coinciden. Es la gramática universal, esto es, el conjunto de los universales

lingüísticos, sobre los que debatimos en la sección 4.2. El instrumental para adquirir una lengua que hay que suponer innato en todo ser humano contiene, pues, fundamentalmente, los universales lingüísticos, más esa función evaluadora recogida en el punto quinto que permite seleccionar la mejor de las gramáticas posibles. Podemos representar esto de la siguiente manera:



En el esquema anterior, G_1 , G_2 , G_3 , ... representan las posibles gramáticas que, estando de acuerdo con los universales lingüísticos o gramática universal, son asimismo compatibles con los datos obtenidos por la experiencia lingüística. De ellas, la función evaluadora seleccionará la gramática correcta, cuya interiorización o conocimiento tácito constituye la competencia lingüística (véase Derwing, *Transformational Grammar as a Theory of Language Acquisition*, sec. 3.4). Por tanto, el instrumento de adquisición del lenguaje equivale a un procedimiento para descubrir una gramática.

Como puede apreciarse, la posesión innata, e inconsciente, de los universales lingüísticos o gramática universal, desempeña un papel central en este modelo. Es la gramática universal la que especifica la forma de la gramática de cualquier lengua humana posible, pues es «el sistema de principios, condiciones y reglas que son elementos o propiedades de todas las lenguas humanas, no por mero accidente, sino por necesidad, ... por necesidad biológica, no lógica; así, la gramática universal puede tomarse como expresiva de 'la esencia del lenguaje humano'» (Chomsky, *Reflections on Language*, cap. 1, p. 29). Por tratarse de algo biológicamente necesario, su estudio es una tarea empírica; de modo diferente, una gramática universal que expresara condiciones lógicas o conceptualmente necesarias no sería propiamente objeto de una investigación empírica. Sobre este último tipo de estudio Chomsky se ha manifestado escéptico, y ha dejado bien claro que el sentido en que toma la expresión «gramática universal» es el primero, sentido que resultaría, naturalmente, vacío en el caso de que llegara a descubrirse que la adquisición del lenguaje no obedece a un instrumento específico, sino que es el resultado de procedimientos genéricos de aprendizaje (*Rules and Representations*, cap. 1).

La gramática universal es, por consiguiente, algo que el niño posee de modo innato, un conocimiento inconsciente y previo a toda experiencia, algo que aporta a su proceso de aprendizaje del lenguaje y que lo hace posible (*Reflections*, p. 118). Esta hipótesis es, como se ve, sumamente

fuerte en cuanto que resalta la importancia del aparato innato, de lo que el niño aporta al proceso de adquisición del lenguaje, dejando en un segundo lugar su experiencia. De aquí que el modelo chomskiano se presente como mentalista y haya sido el caballo de batalla más destacado en la polémica de los últimos años con el empirismo. ¿Pero pueden aceptarse sin más las razones invocadas por Chomsky en favor de su modelo? ¿Se trata de una hipótesis aceptable para la explicación del proceso de adquisición del lenguaje? ¿Ha sido empíricamente confirmado en alguna medida? Esto es lo que trataremos de ver a continuación.

5.2 Las críticas a la teoría de Chomsky

Por la propia índole del planteamiento chomskiano las críticas filosóficas, psicolingüísticas y estrictamente lingüísticas se han producido en una estrecha interconexión que en ocasiones hace difícil distinguir unas de otras y sopesar su importancia relativa. En lo que sigue resumiré las críticas que me parecen más acertadas poniendo el énfasis en las que, por su carácter epistemológico o metodológico, poseen mayor relevancia filosófica.

1. En primer lugar, se ha puesto en duda que tenga sentido afirmar, como hace Chomsky, que el niño aprende su lengua nativa con gran facilidad y rapidez. La cuestión es: ¿con qué puede compararse el aprendizaje del lenguaje para considerarlo fácil y rápido? Así, Putnam («The 'Innateness Hypothesis' and Explanatory Models in Linguistics») ha señalado que, desde el punto de vista de una teoría conductista del aprendizaje, el número de horas que un niño está expuesto a su lengua nativa y la continuidad de esta exposición bastan para explicar su adquisición del lenguaje sin que sea necesario recurrir a una hipótesis innatista tan fuerte como la de Chomsky. Además, el lenguaje no es el primer sistema de signos que el niño adquiere. Como ha recordado Goodman («The Epistemological Argument»), el niño cuenta ya, antes de aprender una lengua, con un sistema de signos prelingüísticos que incluye, por ejemplo, gestos y otros acontecimientos perceptibles. De aquí que mientras un seguidor de Chomsky como McNeill se asombra de la velocidad con la que el niño aprende su lengua («The Creation of Language by Children»), Fraser, en cambio, considerando que al menos durante cinco años el niño dedica la mayor parte de su tiempo y de su actividad a esa tarea, no ve en ello nada particularmente asombroso (en la discusión del trabajo de McNeill).

2. Se ha cuestionado también que la experiencia lingüística del niño sea tan pobre y escasa como Chomsky pretende. La escasez de estudios empíricos suficientemente detallados en este aspecto no permite afirmaciones de gran alcance, pero me parece notable que Brown, que ha realizado una importante investigación empírica sobre las primeras etapas del lenguaje infantil, haya criticado a los lingüistas chomskianos en este punto, señalando que el discurso cotidiano de los adultos es más gramatical de lo que

aquéllos suponen, y especialmente cuando se trata de la comunicación con niños pequeños («Development of the First Language in the Human Species»). Incluso es posible que los procesos de comunicación entre el niño y sus padres, o personas con las que convive, no sean la porción más importante de su experiencia lingüística, sino que lo sea, por ejemplo, la comprobación de los éxitos y fracasos que siguen respectivamente a sus actuaciones lingüísticas correctas e incorrectas; es decir, puede que sea más decisiva la retroalimentación a partir del propio comportamiento lingüístico del niño, como han sugerido Campbell y Wales («The Study of Language Acquisition»).

3. Otra objeción, estrechamente ligada a la primera, es la que se dirige contra el supuesto de que la adquisición del lenguaje constituye un modo específico de aprendizaje distinguible de los demás y con requisitos y presupuestos propios. Así, Morton ha subrayado que otras habilidades que el niño adquiere por la época en que aprende su lengua, como el reconocimiento de estructuras o la coordinación senso-motora, no se han estudiado lo suficiente como para poder situar en su contexto propio y por relación con ellas la adquisición del lenguaje («What Could possibly Be Innate?»); la moraleja de su consideración es que, si hay que recurrir a componentes innatos, éstos no son específicos del lenguaje, y por tanto no son universales lingüísticos, sino en todo caso comunes a cualquier adquisición de conocimientos. Por lo que sabemos sobre la relación entre la adquisición del lenguaje y la de otros conocimientos (véanse las seccs. 5.4 y 5.5), esta crítica tiene más fundamento del que Chomsky le concedería. En un argumento semejante abunda Herriot (*Psicología del lenguaje*, capítulo V, secc. 1). Parecería que Chomsky no se halla lejos de este enfoque cuando afirma que el problema de determinar el contenido del instrumento de adquisición del lenguaje es análogo al problema de estudiar los principios innatos que hacen posible para un pájaro adquirir el conocimiento que se expresa en la construcción del nido o en la producción del canto («Recent Contributions to the Theory of Innate Ideas»). Pero nótese que si empleáramos en este caso la terminología chomskiana habríamos de afirmar que el pájaro posee de manera innata e inconsciente una teoría de la construcción de nidos o una teoría de la armonía, cosa a todas luces absurda. De otra parte, ni siquiera en la conducta animal los mecanismos innatos son todo. Existen pájaros que, habiendo crecido en aislamiento, desarrollan su canto de modo anormal, pero después de haber convivido con un adulto de su especie desarrollan su habilidad vocal del modo usual (Marler, «Animal Communication»).

4. Un fundamental punto de apoyo para la teoría de Chomsky es, como hemos visto, la aparente existencia de universales lingüísticos. ¿Hasta dónde llega nuestro conocimiento de éstos? Después de lo que vimos en la sección 4.2 no parece que podamos sentirnos muy optimistas al respecto. Los ejemplos con los que contamos, o tienen un carácter muy relativo o son más bien triviales. Naturalmente, para Chomsky son universales todas las características formales y sustantivas que su teoría lingüística especifica

para cualquier lengua. Pero, a pesar de todo, el valor que la teoría gramatical de Chomsky posee, valor que nunca le hemos negado en el capítulo anterior, desgraciadamente no ha alcanzado todavía un grado apreciable de confirmación empírica, y las propias modificaciones que ha sufrido durante los últimos años (que brevemente revisamos en la sección 4.5) no contribuyen precisamente a que podamos sentirnos cercanos a un conocimiento adecuado de los universales lingüísticos.

Pero además hay aquí una cautela metodológica digna de consideración. Como ha señalado Putnam («The 'Innateness Hypothesis'...»), es fácil exagerar los hechos por lo que se refiere a los universales lingüísticos. ¿Cuántas categorías sintácticas, por ejemplo, son universales? Siguiendo el análisis propio de las gramáticas categoriales, puede mostrarse que todas las categorías son reducibles a dos: nombre y oración. Pues un sintagma nominal no es sino una expresión que ocupa el lugar de un nombre; un sintagma verbal, es la frase que, combinada con un nombre o con un sintagma nominal, da como resultado una oración; y así sucesivamente. Y esto suponiendo que sepamos permanecer sensibles al carácter hipotético de los universales lingüísticos. Porque podría bien ocurrir que los introdujéramos por definición inadvertidamente y nos resistiéramos a aceptar contraejemplos, eventualidad sobre la que ha avisado Black («Comment»). Hasta podría darse el caso, mencionado por Quine («Methodological Reflections on Current Linguistic Theory»), de que, buscando confirmación de nuestros hipotéticos universales en otras lenguas, los impusiéramos en ellas al traducirlas a la nuestra; es decir, podría acontecer que al investigar otras lenguas las interpretáramos desde las categorías apropiadas para la nuestra.

5. Lo anterior no son sino objeciones periféricas en comparación con la cuestión central metodológica que expresa la siguiente interrogación: ¿es confirmable empíricamente el modelo chomskiano de adquisición del lenguaje? Goodman (*op. cit.*) lo ha negado, como también Quine (*op. cit.*). Pues, en definitiva, el único control empírico para las afirmaciones de la teoría lingüística es el que procede a través del requisito de adecuación descriptiva. Lo único que se puede comprobar es que la teoría define gramáticas que son descriptivamente adecuadas en cuanto que corresponden a las intuiciones que los hablantes poseen sobre su lengua; pero que estas intuiciones, para desarrollarse, requieran la posesión innata de la gramática universal no se ve cómo pueda comprobarse. Pues supóngase que tenemos dos gramáticas distintas que son descriptivamente adecuadas para una lengua determinada, ¿cómo saber cuál de ellas seleccionan los hablantes de dicha lengua? Puesto que ambas corresponden a las intuiciones de éstos, tampoco habrá nada en su comportamiento lingüístico que nos suministre la prueba que buscamos (la sugerencia es de Quine). De otra parte, ¿qué restricciones impone la gramática universal? Por los trabajos de Peters y Richtie, comentados en la sección 4.6, sabemos que con cualquiera de los modelos transformatorios corrientes puede construirse, para toda lengua, una gramática que tiene una base universal. En consecuencia, el postulado de una gramática universal ha de confirmarse empíricamente por separado

y al margen de la teoría transformatoria, al menos mientras no tengamos gramáticas que impongan restricciones tales sobre la potencia del subcomponente transformatorio que eviten las consecuencias denunciadas por Peters y Richtie. En términos parecidos, esta objeción puede encontrarse en Herriot (*op. cit.*, cap. III, secc. 1) y en Derwing (*op. cit.*, secc. 8.2). Pese a tratarse de una objeción de principio, diversos psicólogos y lingüistas han emprendido experimentos de varios tipos con el propósito de comprobar en qué medida podrían encontrarse mecanismos psicológicos que correspondieran a características de la gramática universal, mecanismos tales como el recurso a transformaciones o la identificación de estructuras profundas. Por el momento, y hasta donde llega mi conocimiento, las conclusiones obtenidas son tan escasas, parciales y confusas que en modo alguno suministran un apoyo apreciable al modelo chomskiano. Herriot, por ejemplo, después de revisar un cierto número de estos experimentos, se inclina a pensar que la mejor manera de explicar la adquisición del lenguaje y del pensamiento es atender a las interconexiones entre los esquemas lingüísticos y no lingüísticos, lo que implica situar la conducta lingüística en el contexto de una teoría general del comportamiento (*op. cit.*, cap. VII, sección 3). En conformidad con ello, concluye también que no hay mecanismos específicos y localizados en lugares particulares del cerebro que sean responsables del proceso de adquisición del lenguaje durante la niñez, pues la localización en el hemisferio cerebral izquierdo tiene lugar con la pubertad (cap. V, secc. 7). Lo que puede identificarse, por tanto, en la niñez son mecanismos generales de aprendizaje, algunos de los cuales sin duda son particularmente relevantes para el aprendizaje del lenguaje, como la ordenación jerárquica o las operaciones de transformación; estas últimas, por cierto, son patentes cuando el niño descubre una «estructura profunda» común a estructuras perceptibles aparentemente desemejantes (*loc. cit.*). Sobre este o parecido fundamento tiende a aceptarse que puede haber alguna realidad psicológica que corresponda a la distinción entre estructura superficial y estructura profunda, así como a las operaciones transformatorias, pero que éstas o aquélla tengan, en cuanto entidades psicológicas, las mismas características que les confiere el modelo clásico de Chomsky o alguna de sus versiones posteriores no parece haber sido confirmado en absoluto (véase también, en este sentido, Greene, *Psycholinguistics, Chomsky and Psychology*). Una de las mayores dificultades de los experimentos ideados con este propósito ha sido, naturalmente, la imposibilidad de hacer comprobaciones sobre la competencia del sujeto prescindiendo de los factores de su actuación.

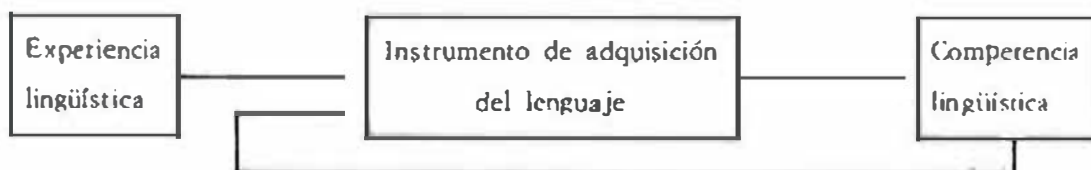
Aunque el libro de Herriot es de 1970 y el de Greene de 1972, no parece que la situación haya cambiado en los años posteriores por lo que se refiere a nuestro tema. Al menos Herriot, en la introducción añadida a la edición de 1976, se ratifica en el enfoque general de su obra y hace referencia a su apartamiento del enfoque innatista. Y ya en 1978 un buen conocedor del tema, Wasow, se ha referido a la falta de pruebas psicolin-

güísticas convincentes en favor de la realidad psicológica de la gramática («On Constraining the Class of Transformational Languages», p. 85).

6. Por último podemos preguntarnos si el modelo de Chomsky representa realmente lo que pretende representar. Pues lo que pretende representar es la adquisición del lenguaje por el niño, y esto lo representa como el resultado de procesar, a través del instrumento de adquisición, los datos lingüísticos que el niño obtiene de su entorno por experiencia. Ese resultado es su competencia o conocimiento lingüístico. Según esto, parecería que el lenguaje se adquiere de golpe y como efecto de dos factores, lo que el niño escucha y lo que lleva dentro. Ahora bien, el niño no se limita a escuchar, sino que también habla, y lo que dice es, en cada momento a lo largo de su aprendizaje, función de lo que lleva aprendido. ¿Es que sus datos no incluyen la experiencia de su propio comportamiento lingüístico? ¿Es que no cuenta el hecho de que tal aprendizaje se desarrolle a lo largo del tiempo durante varios años?

En honor de la verdad hay que decir que Chomsky ha reconocido en algunos lugares la legitimidad de estas consideraciones. En un principio fue particularmente en una nota de *Aspectos* (nota 19 del capítulo 1). Aquí aceptaba que una teoría del aprendizaje del lenguaje habría de afrontar la cuestión del crecimiento continuo de la habilidad lingüística, que puede proseguir incluso después de haber llegado a dominar la forma básica del lenguaje, y aceptaba igualmente la posibilidad de que se apliquen esquemas cada vez más detallados en los sucesivos estadios del proceso. Posteriormente, en *Reflections on Language* (cap. 3, pp. 119 ss.), ha hablado de manera más explícita sobre esos diferentes estadios, admitiendo que para la construcción de la gramática en cada estadio cuentan no sólo los datos de experiencia, sino también el grado de desarrollo alcanzado por la gramática en el estadio anterior, y mostrándose sensible al gran número de problemas que pueden plantearse sobre la forma de ese proceso y sobre las relaciones entre las diferentes etapas del mismo. Es característico, no obstante, de los escritos clásicos de Chomsky, haber presentado su modelo de adquisición del lenguaje en los términos que hemos visto, y, por tanto, como una suerte de modelo «instantáneo», que describe, por decirlo en sus propias palabras, «una idealización en la que sólo se considera el momento de la adquisición de la gramática correcta» (*Aspectos, loc. cit.*)

La influencia que la competencia lingüística alcanzada en cada estadio probablemente tiene sobre el proceso de adquisición en el estadio siguiente es algo que ya se venía señalando en los últimos años. Así, por ejemplo, McCawley, en un trabajo que, como no publicado, cita Derwing (*Transformational Grammar...*, secc. 3.6), proponía completar el esquema clásico de Chomsky del siguiente modo:



Según esto, cada estadio del proceso consistirá en modificar la gramática que el niño tiene ya adquirida, extendiéndola a hechos que hasta ahora no cubría. Lo que el instrumento de adquisición recibe en cada estadio estará integrado por una gramática más un conjunto de nuevos hechos, y lo que resulta será una modificación de la gramática recibida. Esta modificación del modelo chomskiano no es accesoria, pues tiene consecuencias importantes para la concepción de la función evaluadora. La función evaluadora, según McCawley, decidirá, no tanto entre gramáticas completas, cuanto entre posibles modificaciones que pueden efectuarse en una gramática para adaptarla a nuevos hechos de la experiencia lingüística del niño.

Cuando se habla de los datos lingüísticos primarios que integran la experiencia lingüística del niño, hay que tener en cuenta, por consiguiente, que no incluyen únicamente las emisiones verbales que él escucha en su entorno (dirigidas o no a él mismo), sino también sus propios actos de habla, y además las consecuencias de éstos (éxitos o fracasos), así como las reacciones de los demás a su propio comportamiento cuando éste sea respuesta a las emisiones verbales dirigidas a él. No hay que subrayar que el grado de su competencia lingüística en cada momento determinará tanto lo que puede asimilar y procesar de cuanto oye hablar a su alrededor como el nivel de complejidad y corrección de sus propias emisiones. Ante esto cabe preguntarse qué sentido tiene conservar el concepto de instrumento de adquisición del lenguaje y si la importancia de éste no decrece a medida que se desarrolla la competencia lingüística. Pienso que el mero plantearse estas cuestiones indica ya que el modelo chomskiano, incluso tras revisiones como la de McCawley, resulta totalmente inadecuado para lo que pretende representar. Parece mucho más aceptable concebir el aprendizaje lingüístico como el proceso de maduración de la facultad del lenguaje que se producirá como resultado del intercambio comunicativo entre el niño y su medio. En dicho proceso pueden distinguirse, más o menos artificialmente y a efectos de su investigación, diversos estadios. Para el estudio de cada etapa serán relevantes tanto la experiencia lingüística (y asimismo no lingüística) del niño como el grado de la competencia ya alcanzada. Pero esta competencia no será algo adquirido *por medio* del instrumento de adquisición del lenguaje, sino que será el estado final de madurez de la propia facultad lingüística al término de su proceso natural de desarrollo en determinadas condiciones culturales y sociales. Y más que preguntarse por la diferencia entre lo que el instrumento de adquisición recibe (los datos lingüísticos) y lo que produce (la competencia), tendrá sentido preguntarse por la diferencia entre el estado inicial y el estado final en el desarrollo de la facultad lingüística. Esto no supone en modo alguno que la facultad del lenguaje haya de estar vacía en su estado inicial. De hecho no prejuzgamos este punto. La facultad lingüística puede contener todo aquello que nos parezca necesario postular para explicar su maduración en función de la experiencia lingüística; pero dado el papel progresivamente más importante que ésta desempeña a medida que la facultad se va desarrollando, hay que desconfiar en principio de que resulte

necesario atribuir a esta última como contenido la gramática universal toda. En principio es, desde luego, posible que su contenido no sea específicamente lingüístico. En todo caso se trata de un problema empírico mucho más complejo y oscuro de lo que el modelo chomskiano da a entender, y el cargar la cuenta sobre el componente innato a base de hipótesis *ad hoc* no nos pone más cerca de su resolución.

En qué medida la hipótesis que acabo de esbozar tenga apoyo en la psicolingüística y en la biología del lenguaje recientes lo veremos brevemente en una sección posterior. Basten de momento las leves referencias hechas a la obra de Herriot. Añadiré que a las críticas anteriores ha respondido Chomsky repetidamente y en diferentes escritos. Habiendo examinado en otro lugar con cierto detalle estas respuestas, y no pareciéndome que Chomsky haya aportado en ellas ninguna argumentación más clara o más rigurosa en favor de su tesis que las ya aludidas, me limitaré aquí a remitir al lector interesado a ese lugar (*La teoría de las ideas innatas en Chomsky*, caps. 5-6).

5.3 El conocimiento del lenguaje

La competencia lingüística de un hablante equivale al conocimiento que éste tiene de su lengua; el contenido de este conocimiento es una gramática, la gramática propia de esa lengua (*Aspectos*, cap. 1, secc. 1; *El lenguaje y el entendimiento*, pp. 19, 27 y 62 de la primera edición inglesa). Según Chomsky, la habilidad que un hablante tiene de emplear correctamente su lengua implica que conoce de modo tácito las reglas de su gramática. Este conocimiento, naturalmente, es adquirido. Es el resultado del proceso de adquisición del lenguaje, sobre el que ya hemos discutido. Según hemos visto, hay otro tipo de conocimiento lingüístico, éste innato, a saber, el conocimiento que, según Chomsky, todo ser humano posee de la gramática universal y que constituye su aportación al aprendizaje de una lengua. En esta sección discutiremos únicamente acerca del conocimiento que cada hablante tiene de su lengua nativa, puesto que lo que se refiere a la posesión innata de los universales lingüísticos ya lo hemos visto. Pero quede constancia de que Chomsky utiliza el término «conocimiento» (*knowledge*) en ambos casos (*Reflections on Language*, pp. 118 y 164).

¿Qué tipo de conocimiento es el que de su lengua tiene un hablante? Prescindiremos aquí de anticuadas teorías del conocimiento, con frecuencia formuladas en términos más metafóricos que otra cosa y ajenas al moderno desarrollo de la psicología y de las ciencias de la naturaleza. Esto dicho, podemos distinguir tres formas generales de conocimiento. En primer lugar, el conocimiento de mero trato o familiaridad, producto de la experiencia directa de las cosas y de las personas. Es el conocimiento que tiene alguien de una ciudad en la que ha estado, de una persona a la que ha tratado, de una pintura que ha visto, de una composición musical que ha escuchado, etc. Es patente que el conocimiento de la lengua incluye esa

forma de conocimiento, aunque no se reduce a ella. La razón es que se puede tener experiencia directa de una lengua, porque se haya oído hablarla o se la haya visto escrita, pero ser incapaz de utilizarla y entenderla.

En segundo lugar, hay el conocimiento práctico, que consiste en saber hacer cosas, en saber comportarse de tal o cual manera. Es el conocimiento de quien sabe montar en bicicleta, dar volteretas, bailar, hablar en público, tocar la flauta dulce, etc. En principio, parecería que el conocimiento de la lengua es de esta clase. Pues la competencia lingüística implica saber hacer algo: saber producir las secuencias sonoras apropiadas para conseguir lo que uno pretende en cada situación o proceso comunicativo, y ser capaz de entender las producidas por otros. Por ello, Ryle ya consideraba como ejemplos de este conocimiento el hablar gramaticalmente y el comprender las expresiones ajenas (*El concepto de lo mental*, cap. II, seccs. 3 y 9).

Y hay, finalmente, el conocimiento teórico, o de hechos, el conocimiento que consiste en saber que las cosas son de tal o de cual manera. Algunos de estos conocimientos se obtienen a partir de un conocimiento directo o de un conocimiento práctico. Así, quien ha visto una pintura sabe qué tamaño aproximado tiene, qué colores predominan y qué es lo pintado, aunque acaso ignore a qué estilo pertenece o en qué año fue pintada; y quien sabe montar en bicicleta sabe que hay que mantener firme el manillar, aunque puede no saber cómo se explica la relación entre su esfuerzo muscular y el movimiento de las ruedas. De otro lado, es posible tener gran cantidad de conocimientos teóricos sobre algo que no se conoce directamente o que no se sabe hacer. Se puede saber casi todo sobre una pintura que no se ha visto nunca o sobre la manera de bailar un baile que uno es totalmente incapaz de bailar. Y así, se pueden tener muchos conocimientos sobre una lengua que se es incapaz de hablar y entender. Sin embargo, ¿hay algún conocimiento de este tipo teórico que sea parte del conocimiento de la lengua nativa? Aquí hay que tomar en consideración que la adquisición de la competencia lingüística no requiere propiamente estudio o instrucción. Pienso, por ello, que una primera respuesta, ingenua, pero justificada, sería negativa. Cabría decir: el conocimiento que el hablante tiene de su lengua es un conocimiento práctico, consiste en saber hacer algo, como podía esperarse tratándose de la maduración de una facultad innata sometida a los estímulos adecuados. Y esto es lo que la mayor parte de los filósofos y psicólogos del lenguaje venían diciendo. Pero el planteamiento de Chomsky ha venido a complicar las cosas. Porque asigna a ese conocimiento un contenido proposicional como son las reglas de la gramática. Y porque, coherentemente, califica a ese conocimiento de tácito o inconsciente.

El hablante —afirma Chomsky— conoce implícitamente las reglas de la gramática de su lengua, puesto que las aplica de manera inconsciente a diario en sus actos de habla. Podemos todavía preguntarnos en qué medida esto representa una novedad específica del lenguaje. Tal vez tenga buen sentido decir que quien sabe bailar conoce inconscientemente las reglas del baile, o que quien sabe montar en bicicleta tiene un conocimiento

implícito de las reglas cuya aplicación le permite mantenerse en equilibrio conduciendo la máquina en la dirección deseada. No habría aún motivo para renunciar a caracterizar la competencia lingüística como un buen ejemplo de saber hacer. No lo habría si Chomsky no hubiera forzado una estéril polémica en la que se negaba a aceptar que el conocimiento de la lengua pudiera considerarse o como un saber hacer o como un conocimiento teórico («Linguistics and Philosophy», «Knowledge of Language»). De paso negaba también que ambas formas de conocimiento constituyeran categorías exhaustivas para el análisis epistemológico, aunque nunca propuso una categoría nueva. Su posición se resumía en las razones siguientes. El conocimiento de la lengua no es un saber hacer porque no es una habilidad ni un conjunto de hábitos, y no es un saber teórico porque el hablante no tiene por qué ser capaz de enunciar las reglas de su gramática. Esta consideración provocó una discusión en la que hubo gran número de aportaciones, algunas de considerable sutileza y enrevesamiento. En otro lugar la he relatado y he mediado en ella (*La teoría de las ideas innatas en Chomsky*, cap. 2). La ecléctica posición que allí mantuve era la de que el conocimiento del lenguaje es una forma compleja de conocimiento en la que participa el saber práctico, presente en la aplicación inconsciente de las reglas gramaticales, y el saber teórico, patente en la capacidad para distinguir entre oraciones correctas y oraciones incorrectas. El sentido de mi argumentación venía a ser éste: puesto que no parece que la competencia lingüística sea claramente ni un conocimiento teórico ni un conocimiento práctico, y ya que no se ve de qué otras categorías pueda disponerse, tal vez se trata de un conocimiento con elementos de ambos tipos. La cuestión era encontrar un contenido para el aspecto teórico (el práctico está más claro). Lo único que hallé fue la distinción entre oraciones correctas e incorrectas, pues evidentemente forma parte de la competencia del hablante poder distinguir entre unas y otras, es decir, *saber que* ciertas oraciones o expresiones son correctas y que otras no lo son. Y esto me parecía que podía categorizarse como un conocimiento de hechos gramaticales, y por tanto como un saber teórico. Presumo que me influía aquí el hecho de que Chomsky caracterice la competencia lingüística como un sistema de creencias («Linguistics and Philosophy»), caracterización que ahora no me parece más que una mala metáfora.

Justifiqué entonces ese elemento teórico razonando que se trata de una forma de conocimiento típica de quien obra según reglas. El sujeto sabe, en estos casos, hacer algo, y sabe además que ciertos resultados de esa clase de actividad son aceptables y que otros no lo son. Quien sabe tocar un instrumento sabe cuándo una ejecución es correcta y cuándo no; y lo propio puede afirmarse de quien sabe montar en bicicleta, bailar, razonar lógicamente o hablar una lengua. Mas este argumento lleva ya a la vista su defecto. Pues si todo saber hacer comporta la capacidad de distinguir cuándo se han aplicado bien las reglas y cuándo no, entonces no se trata de una característica típica del conocimiento lingüístico. O dicho de otra manera: de acuerdo con mi argumentación no habría conocimientos prácticos;

todos los supuestos ejemplos de este tipo de conocimiento serían en realidad conocimientos mixtos. Con lo cual no habríamos conseguido distinguir el conocimiento de la lengua de otras habilidades prácticas. En resumen: pienso que mi argumentación de entonces no consigue su propósito, y que el conocimiento de una lengua no puede caracterizarse sino como un conocimiento práctico, como un saber hacer. Eso sí, todo lo peculiar y complejo que se quiera.

Otro aspecto que entonces me parecía también teórico es el conocimiento del significado que se expresa mediante afirmaciones como «Sé lo que esa expresión significa». Pero es claro, según pienso ahora, que saber lo que una expresión significa es, desde el punto de vista de un hablante, simplemente saber utilizarla en las situaciones y para los propósitos apropiados (otra cosa es el conocimiento, propiamente teórico, que sobre ello tenga un lingüista).

No veo, por consiguiente, otra manera de caracterizar el conocimiento de la lengua que como un saber práctico adquirido por la maduración, en cierto contexto social y cultural, de la facultad lingüística. Que este saber consiste en aplicar unas reglas, nada tiene de particular, y es cosa que tiene en común con otros saberes prácticos. Que la observación de esas reglas sea inconsciente y no requiera conocimiento teórico de las mismas tampoco es privativo del lenguaje. Que ello comporte saber distinguir qué oraciones son correctas y cuáles no tampoco es sorprendente, pues forma parte de distinguir entre ejecuciones lingüísticas correctas e incorrectas.

A fin de evitar los problemas aludidos, Chomsky ha propuesto, en escritos más recientes, la utilización de términos técnicos para referirse al conocimiento del lenguaje. Los términos elegidos para sustituir a *know* y *knowledge* son, respectivamente, *cognize* y *cognization* (*Reflections on Language*, cap. 4, pp. 164 ss.; *Rules and Representations*, cap. 3). Como los anteriores, los nuevos términos los aplica tanto a la competencia lingüística como al contenido del instrumento de adquisición del lenguaje, es decir, tanto a la gramática particular de una lengua en cuanto interiorizada por el hablante, como a la gramática universal en cuanto innatamente poseída. Así, afirma Chomsky de un hablante que éste *cognizes* su lengua, y por tanto los principios de su gramática, así como los principios de la gramática universal. No es que Chomsky piense que las cosas cambian radicalmente porque modifiquemos la terminología, no. Más bien se trata de un recurso aclaratorio. Puesto que hablar de conocimiento del lenguaje plantea tantos problemas respecto que a qué clase de conocimiento sea, y puesto que sin duda se trata de una forma de conocimiento muy peculiar, utilicemos otro término y así no tendremos que plantear esos problemas.

El recurso es útil en la medida en que manifiesta la peculiaridad que posee la interiorización de las reglas gramaticales, sea de una gramática particular, sea de la gramática universal, pero no resuelve nada. Pues lo que se requiere es explicar qué significa el nuevo término y, en consecuencia, qué es lo que implica decir de un hablante que *cognizes* una lengua o una gramática. A la vista de la nueva sugerencia, podría decirse que lo que

Chomsky estaba haciendo, al hablar del conocimiento del lenguaje en la forma que hemos visto, era una propuesta terminológica. Chomsky habría estado, de hecho, proponiendo una nueva acepción para los términos «conocer» y «conocimiento» cuando éstos se aplican al lenguaje; su posterior recurso a términos técnicos así lo probaría. Curiosamente, esta interpretación fue ya propuesta por Quesada hace tiempo (*La lingüística generativo-transformacional...*, cap. 10). A ella respondí, en su momento, primero, que en tal caso Chomsky estaría dando un extravagante sentido al término «conocimiento» por mucho que luego le añadiera el calificativo de «tácito» o «inconsciente», y segundo, que Chomsky nunca había dado a entender que estuviera haciendo una propuesta terminológica (*La teoría de las ideas innatas...*, p. 57). Ambas afirmaciones me siguen pareciendo correctas, pero es indudable que la reciente sugerencia chomskiana de recurrir a neologismos que sustituyan a los términos ordinarios en este contexto teórico prueba que Quesada tenía razón, y, por consiguiente, que tenía una idea mucho más clara que la del propio Chomsky sobre lo que éste estaba haciendo. En todo caso, y como acabo de indicar, acuñar términos técnicos no resuelve nada mientras no se especifique con claridad lo que significan y se compruebe cómo funcionan en la teoría. Y sobre esto no sabemos más de lo que se ha resumido en las páginas anteriores.

5.4 Lenguaje y conocimiento

Hemos visto en su momento que exigir de una teoría lingüística la justificación explicativa implicaba —según Chomsky— completar la teoría de la gramática con una teoría sobre la adquisición del lenguaje. Y hemos aceptado esto como excusa para examinar en este capítulo la teoría chomskiana al respecto. La razón por la que Chomsky formula su teoría debe, pues, estar clara. Que la teoría haya de tener el extremado carácter innatista y mentalista que ha recibido de Chomsky puede estar menos claro. En mi opinión, y por lo que hemos visto en las secciones anteriores, no lo está en absoluto. Aceptar la teoría transformacional de la gramática, en cualquiera de sus versiones, no obliga, a mi modo de ver, a mantener una posición innatista ni mentalista sobre la adquisición del lenguaje. En rigor, no creo que obligue a mantener ninguna posición al respecto, porque pienso que el requisito de justificación explicativa, tal y como Chomsky lo ha propuesto para la teoría lingüística, debe rechazarse, y por consiguiente que hay que renunciar a extraer, a partir de ésta, consecuencias para el problema de explicar la adquisición del lenguaje. Mi razón fundamental es que este problema sólo puede ser adecuadamente planteado e investigado en el contexto de un estudio psicológico empírico sobre el desarrollo de las facultades cognoscitivas, sobre el aprendizaje de las diferentes pautas de comportamiento y, en especial, del comportamiento verbal; porque sólo aquí existen controles metodológicos mínimos sobre las condiciones de la investigación. La vagarosa psicolingüística chomskiana tiene el peligro, va

suficientemente comprobado, de sumirnos en el reino de las fantasías. Pero el caso es que Chomsky ha dejado los problemas así planteados, y la gran influencia y mérito de su pensamiento obliga hoy, a cualquiera que se aventure por la teoría del lenguaje, a asomarse al tema que ahora nos ocupa, y por ello a introducirse, aunque sea levemente, en el ámbito de la psicología. Esta es la razón por la que ahora aludiremos, aunque sin entrar en grandes detalles ni precisiones, a otras formas alternativas de enfocar nuestro problema con las cuales contrasta vivamente el modelo chomskiano que ya hemos estudiado.

Uno de los primeros escritos de Chomsky, y probablemente el primero en el que trata específicamente temas de psicolingüística, es la crítica que publicó en 1959 sobre el libro de Skinner, *Verbal Behavior*, que había aparecido dos años antes. La crítica de Chomsky hizo fortuna entre sus discípulos y, en general, entre los pensadores opuestos al conductismo skinneriano, especialmente cuando eran ajenos al campo de la psicología. El hecho de que prácticamente hasta 1970 no se produjera ninguna respuesta de conjunto a esa crítica tal vez contribuyó a reforzar su prestigio, al menos entre quienes no eran psicólogos. Creo firmemente que, por poco que coincidamos con ellos, debemos estar agradecidos a Skinner y sus discípulos por no haber respondido; en otro caso, el resultado hubiera sido una de esas estériles polémicas que Chomsky parece especialista en provocar. De continuar en el tono con que Chomsky la había iniciado, la polémica no hubiera resultado tampoco un buen ejemplo de los usos que deben regir en la comunidad académica. Pues, como ha puesto de manifiesto la ponderada contracrítica de MacCorquodale («Sobre la crítica de Chomsky...»), y en ello coincide la aún menos sospechosa consideración de Richelle («Análisis formal y análisis funcional del comportamiento verbal»), la crítica chomskiana, además de claramente teñida de emotividad, desfiguraba los propósitos de la obra de Skinner, y más bien que mostrar que el proyecto de éste fuera inviable, se embarcaba en disquisiciones generales de metodología científica, ora inadecuadas, ora irrelevantes.

No entraremos aquí en los detalles del libro de Skinner, que no viene a cuento para nuestro tema. Baste recordar que se trata de un análisis causal del comportamiento verbal en términos de relaciones de control que incluyen las motivaciones del hablante, sus estimulaciones cotidianas y sus refuerzos. Naturalmente, Skinner recurre a categorías propias de su concepción del comportamiento como un proceso de condicionamiento, las cuales son en principio aplicables a la explicación del aprendizaje del lenguaje. Así, por ejemplo, pueden considerarse formas de respuesta lingüística operante lo que Skinner llama *mand*, una petición que resulta reforzada por ir seguida de una consecuencia característica; el *tact*, una respuesta consistente en nombrar, describir o de algún modo referirse a los objetos; y la respuesta ecoica, que consiste en una repetición del estímulo verbal, esto es, de las palabras oídas, a modo de eco (*Verbal Behavior*, secciones 3-5). Es indudable que todas estas formas de respuesta lingüística tienen su lugar dentro del proceso de adquisición del lenguaje y en relación

con los diferentes tipos de estímulos, tanto internos como externos. Pero es también patente que su consideración no basta para explicar el proceso de adquisición del lenguaje, para lo que hay que tener en cuenta además los complejos problemas que se refieren a la estructura sintáctica, semántica y fonológica de las expresiones. No es, por eso, de extrañar que los textos de psicolingüística completen las referencias al análisis funcionalista de Skinner con consideraciones estructuralistas (véanse, por ejemplo, los de Herriot y List).

Indudablemente, el conductismo radical de Skinner no es más que una forma, particularmente extremosa, de la posición empirista en general y del conductismo en particular. Tanto que, para algunos como List, no puede tomarse como representativo del empirismo (*Introducción a la psicolingüística*, p. 116). De aquí que convenga tener en cuenta que las críticas de Chomsky a Skinner, en cuanto que se hallan estrechamente vinculadas a las categorías de la metodología skinneriana, no afectan de por sí, y en principio, a otros psicólogos conductistas o simplemente no mentalistas. Hay otros enfoques alternativos del problema de la adquisición del lenguaje que, totalmente alejados del planteamiento de Skinner, son igualmente incompatibles con el ingenuo innatismo chomskiano. Por ejemplo, la teoría de Piaget.

Es característico de Piaget subordinar el lenguaje al pensamiento y, con un enfoque también funcionalista, considerar la función verbal como parte de la función simbólica general, explicando la aparición de aquella sobre la base de un nivel de conocimiento previamente adquirido (*Psicología del niño*, III.VI; *Seis estudios de psicología*, 3.I; *Problemas de psicología genética*, 6). El nivel de la inteligencia que es previo al lenguaje y que lo hace posible es el que corresponde al primero de los cuatro períodos en que Piaget ha dividido el desarrollo psicológico, a saber, el período sensomotriz, que va desde el nacimiento a los dos años. Este período se caracteriza por la progresiva diferenciación entre el sujeto y los objetos a base de la manipulación y trato directo con ellos, y tiene como fuentes de conocimiento fundamentales la coordinación de percepciones y movimientos. Dentro de este período pueden distinguirse hasta seis estadios, que van desde la mera consolidación de los esquemas sensomotrices innatos hasta la interiorización de estos esquemas por medio de combinaciones mentales que permiten al sujeto salir de los límites de la pura acción. Recordemos que los otros períodos son: el preoperacional, que va de los dos a los siete años y en el que se distinguen dos estadios, el preconceptual, en el que el niño mantiene una actitud egocéntrica y hace generalizaciones incorrectas a causa de una parcial incapacidad de abstraer, y el intuitivo, en el que, alcanzada esta capacidad, el niño utiliza conceptos de clases y relaciones, a la vez que el egocentrismo deja paso a la socialización; el período de las operaciones concretas, caracterizado por el uso de operaciones lógicas como la reversibilidad, la clasificación y la seriación, y que alcanza de los siete a los once años, y, finalmente, el período de las operaciones

formales, en el que se integran las operaciones concretas en un sistema combinatorio que supone el pleno acceso al ámbito de lo abstracto.

Pues bien, el lenguaje, según Piaget, aparece en el período preoperacional y supone ya algo sumamente importante, el nivel de conocimientos obtenidos en la fase anterior sensomotriz y el desarrollo consiguiente de la inteligencia que se ha alcanzado al final de la misma. Pero el lenguaje no es, según Piaget, lo único que, a diferencia de la percepción y del movimiento, caracteriza al segundo período del desarrollo. El lenguaje, más bien, se da dentro de un contexto general de trato simbólico con la realidad. Antes de haber llegado a dominar medianamente el conjunto de signos intersubjetivos y convencionales que es el lenguaje, el niño posee ya un conjunto de significantes individuales de carácter simbólico. Piaget distingue los siguientes tipos fundamentales de conductas simbólicas (*La formación del símbolo en el niño*; un breve resumen en *Seis estudios de psicología*, cap. 3). En primer lugar, la imitación diferida de algo o alguien cuando el modelo no está presente; aquí la conducta imitativa opera como signifiante diferenciado del modelo imitado. En segundo lugar, el juego simbólico o imaginativo, en el que el niño finge que hace algo, dormir, lavarse, etc. Luego, a medio camino entre este último y la imagen mental, el dibujo, y, finalmente, la imagen mental (véase también la *Psicología del niño*, cap. III). De esta manera, se va desarrollando toda una función simbólica o representativa que se manifiesta claramente, y se ejerce fácilmente, en los casos citados. Sobre ella puede así producirse la adquisición del lenguaje, puesto que éste no es sino una forma particular, y más complicada, de dicha función. (Se observará algo que ya vimos en el capítulo segundo: Piaget, siguiendo la tradición de Saussure, llama «signos» a las palabras, y «símbolos» a los significantes de tipo icónico; paso ahora por alto, de otro lado, su consideración de las imágenes mentales como símbolos, pues criticarlo no es relevante para los propósitos de esta sección. Mi resistencia a aceptar significantes de carácter mental habrá quedado bien manifiesta en el capítulo indicado, así como en otros lugares de este libro.)

La función representativa permite al niño evocar situaciones no actuales y traspasar los límites del aquí y del ahora, ir más allá del alcance de su percepción y abandonar definitivamente el período sensomotriz. Pero la información sobre la realidad obtenida en este período está presupuesta por el uso de los símbolos, así como éste es a su vez un presupuesto de la adquisición del lenguaje. En esta medida, el lenguaje no es suficiente para explicar el pensamiento, incluso en sus formas más abstractas y complicadas, puesto que éste, como lo expresa gráficamente Piaget, «hunde sus raíces en la acción y en mecanismos sensomotrices más profundos» (*Seis estudios...*, cap. 3). Pero cuanto más complejas se hacen las estructuras del pensamiento, más necesario es el lenguaje para éste; para las operaciones lógicas, admitirá Piaget, el lenguaje constituye una condición necesaria aunque no suficiente, y concluirá: «Entre el lenguaje y el pensamiento existe un círculo genético tal que cada uno se apoya necesariamente en el otro, en formación solidaria y en perpetua acción recíproca; pero, en defi-

nitiva, los dos dependen de la inteligencia, que es anterior al lenguaje « independiente del mismo» (*loc. cit.*, p. 124).

Nótese que el niño, según esto, no aporta a la adquisición del lenguaje capacidades específicamente lingüísticas, las cuales supondrían, como ha señalado Richelle, estructuras operatorias de las que aquél aún se halla lejos (*L'acquisition du langage*, VI.2). Lo que aporta es una función más general, la función representativa o simbólica, de la que ya está haciendo un uso más amplio. De acuerdo con la terminología elegida en el capítulo segundo, tendríamos que denominar a tal función más bien «función semiósica». Es comprensible, por ello, que a Piaget, instalado en esta perspectiva que toma como contexto general el desarrollo del conocimiento, el recurso chomskiano a las ideas innatas le parezca inútil y una simple remisión del problema a la biología (*El estructuralismo*, secc. 16; *La epistemología genética*, 2.II; *Seis estudios de psicología*, 5.I). Pues cuanto haya que considerar innato de una forma mínimamente controlable ya viene dado por los esquemas sensomotrices, y a éstos habrá que añadir el resultado de las coordinaciones y autorregulaciones producidas por su aplicación. En qué medida, sin embargo, basten los presupuestos que Piaget señala, no creo que pueda considerarse cuestión resuelta. Es claro que, como recuerda Richelle (*loc. cit.*, p. 141), Piaget no ha seguido nunca en su detalle la evolución de las estructuras lingüísticas del niño, ni se ha interesado por la génesis del lenguaje en sí mismo considerado y en toda su complejidad. Un estudio más detallado y profundo podría conducir a postular capacidades más específicas que constituirían el contenido de la facultad lingüística a la que me referí en la hipótesis general formulada en la sección precedente.

Otra tesis típica de Piaget, por lo que respecta a la evolución del lenguaje infantil, es la distinción entre una etapa egocéntrica y una etapa posterior socializada (*El lenguaje y el pensamiento en el niño*). La actitud egocéntrica, que caracteriza, dentro del período preoperacional, a la fase pre-conceptual, supone que el niño es poco sensible a la función comunicativa del lenguaje, que prescinde de su interlocutor, y que más bien habla para sí mismo y utiliza el lenguaje como un medio auxiliar para estructurar su actividad. Piaget ha señalado tres tipos de utilizaciones egocéntricas: las repeticiones ecológicas, los monólogos y los monólogos colectivos. En el primer caso, el niño habla aparentemente por el placer de hablar y de escucharse a sí mismo; en el monólogo, es como si pensara en voz alta; en el monólogo colectivo, parece asociar a otros a su actividad, pero sin pretender comunicarse con ellos. Falta en todos los casos el punto de vista del interlocutor, al cual el niño en esta etapa parecería insensible. Con la socialización, la función comunicativa del lenguaje iría adquiriendo cada vez mayor importancia. Sobre la base de los resultados cuantitativos obtenidos por Piaget en sus experimentos, parece que es posible afirmar que la proporción de actos verbales egocéntricos iría desde más del 75 por 100 antes de los tres años, hasta menos del 25 por 100 del total después de los siete años.

Esta tesis, que Piaget expuso en una de sus primeras obras (1923), recibió pocos años después una importante crítica por parte de Vygotsky. El hecho de que la obra principal de Vygotsky, *Pensamiento y lenguaje*, publicada póstumamente el año de su muerte, 1934, estuviera retirada de la circulación durante veinte años, hizo que su influjo y recepción en los países occidentales haya sido tardío; el propio Piaget no pudo conocer dicha obra hasta su traducción inglesa de 1962. La posición de Vygotsky es que no tiene sentido distinguir entre un lenguaje egocéntrico y un lenguaje socializado posterior. Una mayor atención a la estructura de las expresiones lingüísticas empleadas, le conduce a encontrar tales diferencias entre el lenguaje comunicativo y el lenguaje egocéntrico que no cabe pensar que aquél derive de éste. Tanto en el niño como en el adulto, dirá Vygotsky (*op. cit.*, cap. 2), la función primaria del lenguaje es la comunicación, el contacto social, y en este sentido, las formas más primitivas del lenguaje infantil son también sociales. Pero pronto se distinguen en él dos usos diferentes, el comunicativo y el egocéntrico, ambos, desde luego, igualmente sociales. El discurso egocéntrico aparece cuando el niño transfiere las formas propias del comportamiento social al ámbito de sus funciones psíquicas internas. Así, el niño que antes había conversado con otros, comienza a hablar consigo mismo, como pensando en voz alta. El discurso egocéntrico conducirá ulteriormente al lenguaje interior, que en el adulto posee gran importancia y sirve tanto a propósitos autísticos como al pensamiento lógico. La línea del desarrollo no es, pues, desde el lenguaje individual al lenguaje social, como parecería serlo en Piaget, sino desde el lenguaje social al lenguaje individual. El lenguaje egocéntrico es, por tanto, «un fenómeno de transición desde el funcionamiento interpsíquico al funcionamiento intrapsíquico, esto es, desde la actividad social y colectiva del niño a su actividad más individualizada» (*op. cit.*, capítulo 7, p. 133 de la traducción inglesa). La evolución del discurso refleja, así, la evolución psicológica general, que va en el sentido de una individualización progresiva. El uso egocéntrico del lenguaje tiene, por ello, una función semejante a la del lenguaje interior: sirve a la orientación mental y a la comprensión consciente. En contra de lo que defiende Piaget, el lenguaje egocéntrico no declina y desaparece, sino que más bien aumenta en importancia y evoluciona para convertirse en lenguaje interior (*loc. cit.*).

Tal vez las posiciones de Piaget y Vygotsky no sean últimamente irreconciliables. Se podría afirmar, como hace Richelle, que la concepción de Piaget deriva de que su interés primordial es el desarrollo del conocimiento y no la evolución del lenguaje, mientras que Vygotsky está más preocupado por este último tema y en tal medida atiende más a las estructuras lingüísticas y a su funcionamiento con independencia de los aspectos cognoscitivos (Richelle, *L'acquisition du langage*, V.3, p. 125). Pruebas empíricas parece que las hay tanto en favor de una tesis como de la otra. Pues, en definitiva, ambas podrían combinarse en una visión unitaria como la siguiente: el lenguaje, originariamente social y aprendido por el niño en los procesos de comunicación, es utilizado por éste durante una primera

etapa infantil, aproximadamente entre los tres y los seis años, con propósitos que no son frecuentemente comunicativos, pues resulta ser él mismo al tiempo emisor y destinatario de sus palabras (todo lo cual no es muy llamativo si se repara en lo que al aprendizaje del lenguaje puede ayudar la ecolalia y el monólogo); adquirido hacia los siete años un dominio suficiente de la lengua, las formas del lenguaje egocéntrico perderían interés para el niño y darían paso progresivamente al lenguaje interior. El desacierto de Piaget estaría en haber sugerido con su terminología que la fase egocéntrica no es, desde el punto de vista lingüístico, social, mientras que sí lo es la fase posterior. De otra parte, es cierto que la teoría de Piaget no ha concedido atención suficiente a los factores culturales para los que el lenguaje constituye el medio más idóneo de transmisión, ni a la aportación de aquéllos a la formación de los conocimientos y a la estructuración de la inteligencia, que parecería resultar exclusivamente de procesos individuales de adaptación.

Una teoría que intenta superar estas deficiencias conservando los aciertos básicos de la escuela de Piaget, es la que ha sugerido Bruner (en su aportación a *Studies in Cognitive Growth*). La progresiva independencia que, con respecto a los estímulos inmediatos, caracteriza al desarrollo psicológico del niño, depende —según él— de procesos representativos que se dan en tres niveles sucesivos. En el primero, se prolonga la acción motriz; las respuestas motoras adquiridas son reproducibles por el sujeto, aunque los estímulos originales no estén presentes; a esto denomina Bruner «representación enactiva». En el segundo nivel se prolonga la organización perceptiva, dando lugar los datos sensibles a imágenes, sin necesidad de recurrir a la acción; a esta forma de representación le llama «icónica». En el tercer nivel se recurre a símbolos, y se trata por tanto de una representación simbólica, que incluye el uso del lenguaje. Característicamente, Bruner ha subrayado que, mientras que los dos primeros niveles suponen simplemente la relación del organismo con el medio, el tercer nivel requiere la cultura, que es la que suministra los símbolos, y especialmente el lenguaje. Una vez que aparece éste, el desarrollo del conocimiento no puede considerarse separado de él, y, por consiguiente, tampoco separado de la cultura, pues ésta va a operar justamente a través del lenguaje sobre el desarrollo individual. Como ha señalado Richelle, la teoría de Bruner tiene un doble mérito: no separa la adquisición del lenguaje del resto del desarrollo cognoscitivo, y al propio tiempo inserta éste en el fenómeno cultural del cual es parte el lenguaje (*L'acquisition du langage*, VI.5, p. 159).

Teorías como las anteriores tienen, a mi juicio, la importante ventaja sobre la de Chomsky de considerar la adquisición del lenguaje en un contexto más amplio, psicológico y cultural, en el que sin duda aquélla tiene lugar. Podría defenderse a Chomsky manifestando que presta más atención a los aspectos estructurales propiamente lingüísticos; pero entonces es obligado añadir que, al declarar innatos la mayor parte de estos aspectos, se ha renunciado propiamente a explicar su adquisición. Si la gramática universal

es innata y la estrategia para seleccionar la gramática adecuada a una lengua también lo es, entonces no es mucho lo que resta por explicar. Tal es la impresión que produce una obra, más detallada que cuanto Chomsky haya podido escribir sobre el tema, y fielmente representativa de su posición, como es *The Acquisition of Language*, de McNeill. Aquí, el desarrollo del lenguaje parece concebirse, como ha indicado Richelle (*op. cit.*, p. 156), *in vacuo*. Otros aspectos de la evolución del conocimiento, anteriores o coetáneos a la adquisición del lenguaje, brillan por su falta de mención, y la función del contexto cultural queda reducida al mínimo. Son los males del innatismo, que soluciona todo pero no explica nada. Porque si se toma con rigor, el innatismo lo único que hace es remitir los problemas a la biología. ¿Y qué hay que esperar que nos diga la biología en favor del innatismo de la gramática universal?

5.5 Biología y lenguaje

Desgraciadamente, para los propósitos de Chomsky la biología tiene por ahora poco que decir. Porque si consideramos la excelente sistematización de este tema realizada por Lenneberg, que además, y por si hubiera dudas, se declara desde un principio decidido innatista y entusiasta simpaticante de Chomsky, no encontraremos nada que pueda tomarse como apoyo, siquiera indirecto, para la tesis chomskiana de las ideas innatas, esto es, para la afirmación de que el niño tiene innato el conjunto de las categorías y reglas de la gramática universal.

En Lenneberg encontramos, primeramente, una alusión a los mecanismos fisiológicos innatos que controlan la utilización del lenguaje, su producción y recepción (*Fundamentos biológicos del lenguaje*, cap. 5, secc. V). Cosas tales como la modulación característica de las células nerviosas, la estructura de las cadenas neuronales, las características oscilatorias de las actividades endógenas, y, en general, todo lo que es hoy conocido sobre la estructura y funcionamiento del sistema nervioso en los aspectos presumiblemente más relacionados con la actividad verbal. Y no es de pasar aquí por alto la cautela que manifiesta el autor cuando dice: «Cómo interactúen estos fenómenos para elaborar el lenguaje sigue siendo un misterio» (*loc. cit.*, p. 221 de la edición original).

Encontramos en segundo lugar, y principalmente, una teoría biológica general sobre el lenguaje que, sobre la base de los datos biológicos que se poseen, lo pone en relación con las capacidades psicológicas generales que caracterizan a la especie humana. Llamando «matriz biológica» al conjunto de características innatas que determinan, para cada especie de organismos, el resultado de someterlos a las estimulaciones del medio, se trataría de averiguar cómo hay que definir la matriz biológica humana con respecto al lenguaje (*op. cit.*, 9.V). Lenneberg parte aquí de algunas consideraciones generales de las que las más relevantes son las siguientes. Existe para cada

especie una función cognoscitiva típica, llamando así a las funciones cerebrales que median entre las estimulaciones sensibles y la respuesta motriz. La neurofisiología de esta función es en gran parte desconocida, pero pueden señalarse como sus correlatos de comportamiento la propensión a categorizar de modos específicos (que se basa en la extracción de semejanzas), la capacidad para resolver problemas, la tendencia a generalizar en ciertas direcciones y la facilidad para memorizar ciertas condiciones. La interacción de estas capacidades produce en cada especie sus particularidades cognoscitivas, que determinan, en última instancia, la peculiar visión de la realidad propia de cada una (cap. 9.I). Las capacidades cognoscitivas se diferencian espontáneamente con la maduración, dada, naturalmente, la concurrencia del medio adecuado. Sobre esta base, Lenneberg expone su concepción biológica del lenguaje en los siguientes términos (cap. 9.II).

El lenguaje es una manifestación de las tendencias cognoscitivas (*cognitive propensities*) propias de la especie humana, es decir, viene determinado por la función cognoscitiva que caracteriza a aquélla. Esta función es una adaptación del proceso, que puede notarse en los animales superiores, de categorización y extracción de semejanzas. Según Lenneberg, esta última no opera solamente sobre los estímulos físicos sino también sobre las categorías, puesto que las palabras, con excepción de los nombres propios, no designan objetos directamente sino más bien clases y, por tanto, procesos de categorización.

Ciertas especializaciones de la fisiología y anatomía periférica explican algunos de los rasgos universales de las lenguas humanas, pero no por ello pueden tomarse como una explicación de la filogénesis del lenguaje. Ya en otro lugar Lenneberg había subrayado que el tamaño del cerebro no es determinante para la capacidad de aprender una lengua con tal que se trate de un cerebro humano, puesto que ciertos enanos, cuyo cerebro es considerablemente más pequeño de lo normal, no por eso carecen de dicha capacidad («A Biological Perspective of Language», y también en *Fundamentos biológicos...*, 2.III). Ciertos rasgos, sin embargo, como el desarrollo de un tracto suprararíngeo dividido en dos conductos curvos, es considerado por otros como «el estadio final crucial en la evolución del lenguaje humano» (Lieberman, *On the Origins of Language*, cap. 11, p. 159). Esto habría permitido desarrollar un lenguaje enteramente dependiente del componente vocal, relegando los gestos a una función auxiliar, paralingüística. Presumiblemente, los australopitecinos poseían un lenguaje compuesto tanto de gestos como de elementos vocales; a partir de ellos, algunos homínidos habrían retenido ese tipo de lenguaje, mientras que otros habrían evolucionado, sobre la base de una modificación del tracto suprararíngeo, hacia un lenguaje enteramente vocal (véase una clasificación de los restos fósiles por razón de esta característica, en Lieberman, *op. cit.*, p. 160, y en su artículo «Un enfoque unitario de la evolución del lenguaje», p. 187). Tal modificación permitiría definitivamente producir señales acústicas óptimas, aumentando el repertorio de sonidos producibles y haciéndolos más efica-

ces. De aquí la significación que Lieberman le concede, cuando afirma al final de su libro: «... la forma particular que ha adoptado el lenguaje humano parece ser el resultado de la evolución del aparato vocal supralaríngeo; esto y los cambios producidos en la morfología del cráneo, que distinguen al *Homo sapiens* de los demás animales, fueron en los últimos estadios de la evolución homínida factores tan importantes como la dentición y la postura bípeda en los estadios primeros» (p. 181). Afirmación tan comprometida con la importancia filogenética de los rasgos anatómicos ciertamente no se encuentra en Lenneberg (sus consideraciones, mucho más vagas, sobre la funcionalidad de los cambios en el tracto vocal, se encontrarán en el cap. 2, secc. II, apartado (4) de su obra).

La maduración del organismo humano —sigue Lenneberg en el resumen de su teoría— conduce los procesos cognoscitivos a un punto que denomina de «preparación para el lenguaje». Los datos lingüísticos aportados por la experiencia del niño actuarán como materia prima con la que éste actualizará su estructura lingüística latente convirtiéndola en estructura lingüística realizada. Según Lenneberg, aquélla puede considerarse como el correlato biológico de la gramática universal, y la estructura realizada, como el correlato de la gramática particular. Se notará que la llamada «estructura lingüística latente» es un nuevo elemento que no viene justificado ni exigido por el resto de la teoría; y se notará asimismo que no se nos indica ningún método biológico que nos permita averiguar cuál es el contenido de esa estructura. Lo único claro es que se trata de un proceso de maduración que incluye aspectos orgánicos y cognoscitivos, y que cuando Lenneberg tiene que referirse al elemento universal y común a todas las lenguas, se limita a indicar que lo que es universal es «nuestro modo de calcular con categorías», pero no las categorías mismas ni las operaciones realizadas con ellas (cap. 9.II, ap. (10)).

A la vista de lo anterior, pretender que la biología ofrece hoy pruebas en favor de un innatismo como el chomskiano es forzar a los biólogos a decir lo que no dicen. Ni Lenneberg, que ya desde un principio está dispuesto a decir cuanto pueda en favor de Chomsky, puede de hecho ir más allá de postular, aparte de las condiciones anatómicas y neurofisiológicas, una función cognoscitiva común a la especie humana, que se manifiesta en la identificación de semejanzas, en la generalización y en la resolución de problemas. Capacidades éstas que, por cierto, también Quine, a veces calificado de conductista, y con frecuencia criticado por Chomsky, acepta como innatas y como necesarias para la adquisición del lenguaje («Reflexiones filosóficas sobre el aprendizaje del lenguaje»). Sobre el contenido y desarrollo ontogenético de la función cognoscitiva, y sobre su relación con el lenguaje, ciertamente hay sugerencias más interesantes y complejas en Piaget que en todas las doctrinas innatistas. En definitiva, no se ha probado que tengamos innata otra cosa que ciertas capacidades de conocimiento. Y esto no es mucho. Porque el innatismo es sólo la porción más fácil y más reducida de una teoría epistemológica.

5.6 ¿Es específicamente humana la facultad del lenguaje?

En todo lo anterior hemos estado asumiendo que el lenguaje es algo que caracteriza a la especie humana y la distingue de las demás especies animales. Si se atiende a las realizaciones lingüísticas, a las lenguas, es patente que ninguna otra especie animal conocida posee un sistema de comunicación que pueda compararse con el lenguaje verbal humano. ¿Pero significa esto que no posea en absoluto la capacidad para tenerlo? Es decir, ¿puede afirmarse que sólo la especie humana posee la facultad lingüística? Es característico de Chomsky haberlo afirmado así: «Cualquiera que se ocupe del estudio de la naturaleza y capacidades humanas tendrá que abordar el hecho de que todos los seres humanos normales adquieren el lenguaje, mientras que ya la adquisición de sus más elementales rudimentos se encuentra mucho más allá de las capacidades de un primate por inteligente que sea» (*Language and Mind*, cap. 3, p. 59). Y en el mismo tenor, Lenneberg ha escrito: «No hay prueba alguna de que un organismo no humano tenga la capacidad para adquirir ni siquiera los estadios más primitivos del desarrollo lingüístico» («A Biological Perspective of Language», p. 67). Pues bien, diferentes experimentos realizados a partir de los años sesenta han despertado algunas dudas sobre la posible justificación de estas afirmaciones, y han aportado, en todo caso, nuevos datos sobre los límites de la capacidad de aprendizaje de los primates superiores.

El intento de enseñar a un chimpancé a comunicarse con los seres humanos tiene algunos antecedentes un poco más antiguos. Los más conocidos son el del matrimonio Kellogg y el del matrimonio Hayes. Los Kellogg, por los años treinta, criaron en su casa a un chimpancé hembra llamado Gua, al que hablaban en su lengua. Se trataba, por tanto, de una enseñanza pasiva, con la que se consiguió que el sujeto del experimento llegara a responder de forma adecuada (y en esta medida podemos decir: llegara a *entender*) a unas setenta expresiones distintas. Los Hayes, por su parte, en los años cuarenta, criaron un chimpancé, también hembra, de nombre Viki, al que intentaron enseñar un comportamiento lingüístico más activo, pero todo lo que consiguieron fue que llegara a pronunciar cuatro sonidos algo semejantes a cuatro palabras inglesas (tomo la información de la introducción de Brown a *A First Language*).

La habilidad de Gua, de carácter pasivo, no parece que difiera mucho de la de ciertos mamíferos superiores como el perro, y ciertamente no constituye en modo alguno ni el más remoto inicio de aprendizaje lingüístico. Tampoco la reducida habilidad de Viki, por activa que sea, supone un contraejemplo para las afirmaciones que hemos recordado al principio. Que en ciertas condiciones de domesticidad, algunos mamíferos superiores, determinadas clases de aves, o animales marinos como los delfines, sean capaces de responder a un reducido sistema de señales, comportándose de manera adecuada o incluso emitiendo, a su vez, otros signos es, todo lo más, prueba de una capacidad semiótica y comunicativa (y de hecho apa-

reció debidamente recogida en el capítulo segundo), pero queda todavía a enorme distancia de una capacidad lingüística comparable con la humana.

Los grandes avances, teóricos y experimentales, en la metodología conductista han permitido, sin embargo, y en tiempos más recientes, llevar a cabo experimentos mucho más rigurosos, y de resultados más llamativos, que los anteriores. En estos experimentos ha sido una idea básica la de que había que sortear las dificultades que pudieran proceder del hecho de que la distinta conformación de los órganos fonadores del chimpancé le impide producir sonidos articulados. Se ha recurrido en todos ellos a algún lenguaje, o sistema de signos, que, conservando características centrales del lenguaje verbal humano, no requiera el uso de órganos fonadores y pueda, en cambio, ser empleado por el chimpancé. De esta forma se evitaban las patentes dificultades anatómicas de Viki, al tiempo que se podía aspirar a un comportamiento activo, y no meramente pasivo como el de Gua. Todos estos experimentos han tenido lugar en los Estados Unidos.

El primer experimento fue el que iniciaron los esposos Gardner en 1966 con un chimpancé, también hembra, llamado Washoe, y sobre cuyos resultados ya hay varios informes (en castellano, «Cómo enseñar el lenguaje de los sordomudos a un chimpancé»). Los Gardner razonaban que, dada la importancia que tiene el uso de las manos en la conducta de los chimpancés, un lenguaje que podría resultar de fácil aprendizaje para ellos es un lenguaje de gestos manuales, y en consecuencia eligieron uno de los lenguajes empleados por los sordomudos en Norteamérica. Estos tienen a su disposición dos tipos diferentes de lenguaje. En uno, los signos realizados con las manos representan letras del alfabeto; es, por tanto, una manera de deletrear manualmente las palabras, y depende directamente del lenguaje escrito. El otro, que es el que eligieron los experimentadores, consiste en representar, por medio de gestos manuales, contenidos significativos completos, y en esta medida cada signo corresponde, en principio, a un lexema. Unos signos son totalmente arbitrarios, mientras que otros tienen un carácter claramente icónico. A diferencia del primer tipo de lenguaje manual, este último permite a los sordomudos entenderse entre sí con independencia de cuál sea la lengua de su país, y constituye propiamente un lenguaje distinto. Aprovechando las grandes facultades de imitación de los chimpancés y utilizando este lenguaje para acompañar a una serie de actividades cotidianas que desarrollaban con Washoe, los Gardner y sus ayudantes consiguieron, no sólo que aquélla llegara a entender el significado de cierto número de signos, sino más aún, que los realizara por sí misma de manera espontánea en las situaciones apropiadas para expresar ciertos deseos o referirse a determinados objetos. En principio, el método utilizado, aunque sometido a las condiciones propias de un experimento, no parece distinto del que permite a un niño aprender su lengua cuando vive en un medio normal. Simplemente, el lenguaje lo usan aquellos que le rodean para comunicarse entre ellos y con él, y el uso que éste haga será reforzado con reacciones aprobatorias y con el propio éxito de la comunicación. A esto, los experimentadores añadían también diversas técnicas

de condicionamiento operante. Con este sistema, Washoe había llegado a emplear, tres años después de iniciarse el experimento, unos ochenta y cinco signos distintos, y para 1971 había alcanzado ciento treinta y dos, aunque el número de signos tal vez no es lo más importante aquí. Entre los signos aprendidos están los que significan *más, arriba, dulce, abrir, ir, fuera, oír-escuchar, por favor, divertido, comida-comer, perro, tú, dentro, olor, gato, llave*, etc. Hay que notar que los signos no recogen las diferentes características verbales (como tiempo, aspecto, etc.) y que muchos signos funcionan indistintamente como verbos y como substantivos; tales diferencias vienen indicadas por la situación en la que se emplean. Los signos eran utilizados no sólo aislados sino asimismo en secuencias, algunas de ellas inventadas por el simio, y que llegaron a tener hasta cinco signos. Lo que es realmente importante notar es que los signos aprendidos (y sólo se aceptaban como tales los que eran empleados correctamente con cierta asiduidad determinada) eran utilizados por Washoe espontáneamente y no sólo como respuesta a sus experimentadores; que utilizaba esos signos en situaciones nuevas y para designar nuevos objetos; y que los combinaba espontáneamente.

¿Puede afirmarse que Washoe haya aprendido un lenguaje? En el artículo citado antes, los Gardner declaran que responder a una pregunta así les resulta difícil porque se trata de una pregunta ajena al espíritu de su investigación. Y lo es en cuanto que implica una distinción entre una conducta comunicativa que podamos considerar lingüística y otra que no, cuestión en la que ellos han preferido no entrar. Pero en la medida en que nos interesa averiguar si los primates superiores poseen la capacidad para adquirir un lenguaje, no parece que podamos eludir la pregunta anterior. Contestarla supone, naturalmente, tener una teoría sobre lo que hace de un sistema de comunicación un lenguaje en el sentido de las lenguas humanas, y si vamos a recurrir a un conjunto de características definitorias como las que aceptamos en la sección 4.1, entonces probablemente resulte fácil, tal vez demasiado fácil, responder nuestra pregunta negativamente. Por ello pienso que es más constructivo asumir por hipótesis que el sistema de comunicación que Washoe ha aprendido *es* un lenguaje (en definitiva lo utiliza una porción de los seres humanos) y preguntarnos más bien: primero, ¿con qué grado de perfección lo ha aprendido nuestro chimpancé?; y segundo, ¿hay distinciones relevantes entre su comportamiento lingüístico y el de los infantes humanos?

En primer lugar hay que decir que, en condiciones experimentales bien controladas, se ha demostrado que Washoe, en el uso de signos designativos, alcanza un índice de corrección que se encuentra entre el 50 y el 55 por 100 de los casos (Brown, *A First Language*, p. 40). Un niño que ha aprendido un término designativo lo usa, en cambio, con un acierto del 100 por 100 (Pinillos, «Comunicación animal y lenguaje humano», p. 15). Se ha señalado igualmente que, aunque las secuencias de signos empleadas por Washoe corresponden bastante bien a las estructuras lingüísticas propias del primer nivel del desarrollo lingüístico infantil (alcanzado entre los

dieciséis y los veintisiete meses, según Brown), existe la siguiente diferencia. Washoe no otorga relevancia al orden de los signos en la secuencia, orden que modifica de forma arbitraria. Mientras que, según ha mostrado Brown (*loc. cit.*, p. 41), al menos para la lengua inglesa y de acuerdo con los datos que ha obtenido experimentalmente, el orden de las palabras ya en esa primera etapa es algo a lo que el niño concede una significación. Claro que aquí cabría observar que esto es cosa más del propio lenguaje enseñado a Washoe que del uso que ésta hace de aquél. Pues si es correcta la indicación de Sampson (*The Form of Language*, p. 123), el lenguaje de signos manuales utilizado en este experimento es un lenguaje que carece de sintaxis y en el que el orden de las palabras o bien es arbitrario o bien tiene una función particular cuasi-icónica, por ejemplo, para indicar el orden temporal de las partes de un proceso. El propio Sampson menciona experimentos que mostraban que los usuarios humanos de un lenguaje así no podían expresar por medio de sus signos la diferencia entre un chico dando una manzana a una chica y una chica dando una manzana a un chico. Si esto es así, entonces la falta de orden en las secuencias de Washoe no hay que tomarla como una característica muy peculiar de su comportamiento. A decir verdad, Brown mismo no concede gran importancia a la cuestión del orden de los signos, y por diferentes razones a las que sería prolijo aludir aquí, piensa que el experimento suministra razones suficientes para concluir que Washoe expresa intenciones semánticas (y, por tanto, que su comportamiento no es el mero resultado de un condicionamiento) y que ha alcanzado un nivel semejante al primer nivel del desarrollo lingüístico infantil (*loc. cit.*, p. 43). Como él mismo sugiere, el estudio que tendría realmente valor comparativo sería el de un niño que aprendiera ese lenguaje manual como su primer lenguaje, e informa de que Bellugi-Klima tiene uno en curso de realización (desgraciadamente no conozco ningún informe sobre este estudio).

Todo lo más que podría afirmarse sobre Washoe es, pues, que posee capacidad para llegar al primer nivel del desarrollo lingüístico humano. Esto no prueba, sin embargo, que pueda llegar más allá. Brown piensa que es una capacidad suficiente para justificar un considerable grado de cultura, y si es así, ello sugiere que los chimpancés no han hecho uso de toda su capacidad mental. Aunque acaso esto no deba extrañarnos: probablemente tampoco los seres humanos hemos llegado a utilizar totalmente nuestra capacidad. El caso es que, debido al crecimiento de Washoe, y por la amenaza que podía suponer para los ayudantes más jóvenes e inexpertos, su mantenimiento en condiciones experimentales se hizo tan difícil que hubo que trasladarla a una reserva de primates. Aunque había intención de continuar el experimento dentro de lo posible, no conozco ningún informe sobre progresos posteriores a 1971. Lo que sí se encuentran son informes sobre otros experimentos análogos. Así, por ejemplo, Fouts («La adquisición y comprobación de signos gestuales en cuatro chimpancés jóvenes», 1973) enseñó diez signos del lenguaje de Washoe a cuatro chimpancés, dos machos y dos hembras, a fin de comparar el ritmo de adquisi-

ción entre ellos y Washoe. Los resultados muestran que unos signos son más fáciles de adquirir que otros, y que unos chimpancés tienen más habilidad o capacidad para adquirir signos que otros. La comparación con Washoe prueba que el caso de ésta no es en modo alguno excepcional. Otro experimento semejante es el iniciado por los propios Gardner en 1972 («Signos tempranos de lenguaje en el niño y en el chimpancé», 1975). La principal novedad es que los chimpancés sometidos a experimentación (dos, según el informe que acabo de citar) eran recién nacidos, y estuvieron sometidos al experimento desde el comienzo de sus vidas. Los ayudantes del experimento eran, por su parte, o sordomudos o hijos de sordomudos, y por ello mucho más expertos y fluidos en el uso del lenguaje de signos. El informe relata que los simios comenzaron a usar signos cuando contaban unos tres meses, y a los seis meses usaban ya unos quince signos. Para estimar la diferencia con Washoe, piénsese que ésta, a los seis meses, sólo utilizaba dos signos. La diferencia, claro está, no se debe a diferencias de capacidad en los nuevos sujetos del experimento, sino al mayor dominio del lenguaje manual que tienen los experimentadores y a la más temprana edad en que iniciaron el experimento los chimpancés.

A resultados aún más espectaculares parece haber llegado el chimpancé, también hembra, Sarah, sometido a experimentación por Premack a partir de 1966 (sobre este caso hay varios informes y un voluminoso libro; en castellano, «Algunas características generales de un método para enseñar el lenguaje a organismos que normalmente no lo adquieren»; el libro, de Premack, es *Intelligence in Ape and Man*). El lenguaje enseñado a Sarah difiere notablemente del utilizado en los experimentos que hemos comentado. En este caso se trataba de piezas de plástico, de diferentes tamaños, colores y formas, y con el reverso de metal, lo que permitía adherirlas a una pizarra magnética. Los signos formaban secuencias que se representaban verticalmente sobre la pizarra; la razón de esta escritura vertical, según Premack, es que, desde un principio, el simio pareció preferirla. Propiamente, pues, puede decirse que Sarah aprendió a leer y escribir. Cada pieza de plástico representaba una palabra, y hacia 1972 Sarah utilizaba ciento treinta piezas, con un índice de aciertos entre el 75 y el 80 por 100 (Premack, «Teaching Language to An Ape», p. 92). Lo más llamativo es el amplio alcance semántico de algunas de estas piezas. Por lo pronto había dos que oficiaban como nombres propios para designar, respectivamente, a Sarah y a la experimentadora. Otros símbolos representaban las habituales clases de objetos: *plátano, cubo, plato, manzana...*; algunos verbos como *dar, coger, colocar*; propiedades como *rojo y amarillo, grande y pequeño, redondo y cuadrado*; y términos cuantificadores como *todos, ninguno, uno y varios*. Había, además, piezas que representaban relaciones como *igual a y distinto de*; había un símbolo para la afirmación, otro para la negación y otro para la interrogación, así como otra pieza representando la relación condicional *si... entonces*. Y más interesante aún: para cada uno de los pares de propiedades mencionados existía una pieza que significaba la correspondiente propiedad de segundo orden: *ser-el-color-de, ser-el-tama-*

ño-de, y *ser-la-forma-de*. Había, finalmente, una pieza de carácter claramente metalingüístico, puesto que funcionaba con el significado de *ser-el-nombre-de*.

Con estos elementos, los experimentadores consiguieron que Sarah descifrara y reprodujera «oraciones» de una complejidad semántica muy llamativa. Premack seguía con ella el método skinneriano de las aproximaciones sucesivas por el que iba condicionando su comportamiento a base de premiar sus aciertos parciales, hasta conseguir así un éxito total en la conducta comunicativa. Las diferentes oraciones o secuencias de signos que Sarah consiguió aprender a completar resultan, en esta medida, un tanto inconexas y como no integradas en un sistema conjunto, según ha puesto de manifiesto Brown (*A First Language*, p. 34). Pero ciertamente algunas de estas oraciones, si efectivamente eran «entendidas» por Sarah, suponen una capacidad mental claramente lingüística. Así, por ejemplo, cuando entre dos objetos el experimentador colocaba el signo de interrogación y los signos de *igual a* y *distinto de*, el chimpancé aprendió a sustituir el signo de interrogación por uno de estos dos últimos, según los objetos fueran iguales o diferentes. Aprendió igualmente a colocar la pieza para *ser el nombre de* entre la pieza que significaba *manzana* y una manzana real, añadiendo la pieza *no* si en lugar de una manzana se colocaba otra fruta. Utilizaba también correctamente la pieza que significaba *si... entonces*, pieza que colocaba entre la oración antecedente y la oración consecuente, al modo como se utiliza en lógica el signo correspondiente, y demostró capacidad para transferir esta relación a contextos nuevos. Podía asimismo formar oraciones del tipo de «*amarillo no ser-el-color-de manzana*», en contestación a una pregunta formada por el símbolo de *amarillo*, el símbolo de la interrogación y el símbolo de *manzana*.

¿Qué decir sobre los éxitos de Sarah? Premack piensa que resiste bien la comparación con un niño de dos años en lo que a habilidad lingüística se refiere («Teaching Language to an Ape», p. 99), y que la comparación está justificada por lo que toca a su peculiar lenguaje. ¿Cuándo deja una pieza de plástico de ser una pieza de plástico para convertirse en una palabra? La respuesta de Premack es simple: cuando se usa como una palabra, cuando aparece junto con otras palabras de la clase gramatical apropiada formando oraciones, cuando se usa como respuesta o como parte de una respuesta a una pregunta («Language in Chimpanzee?», p. 820). Como en el caso anterior, discutir si estamos o no ante un auténtico lenguaje sería errar el tiro. Por un lado, las diferencias entre este juego de ciento treinta fichas y el sistema de los fonemas de una lengua son lo bastante claras como para que no tenga interés señalarlas, especialmente desde el punto de vista de sus posibilidades de combinación y del alcance de sus significaciones. La cuestión, una vez más, es lo que representan las habilidades adquiridas por el simio en cuanto expresivas de una posible capacidad lingüística. Y en este aspecto, y a pesar de los excelentes resultados del experimento, no parece que se puedan sacar conclusiones diversas de las obtenidas en el caso de Washoe. De hecho, y en comparación con ésta, se ha notado en el caso de Sarah una mayor pasividad. Como ha señalado

Brown (*loc. cit.*), Sarah no parecía interesarse por las piezas de plástico salvo bajo estimulaciones previas del experimentador, y no demostró capacidad apreciable de iniciar la comunicación. Esto podría ser debido a las condiciones del experimento, que, a diferencia del realizado con Washoe, tal vez no atendió al desarrollo de la iniciativa del chimpancé, pero no deja de ser llamativo, y en este punto Sarah está, desde luego, muy por debajo de Washoe. El índice de aciertos de aquélla, por otra parte y como ya se mencionó antes, aunque tan alto que llegó al 80 por 100 de los casos, no alcanza aún al ideal 100 por 100. Aquí puede uno extrañarse, por cierto, como hace Brown (*loc. cit.*, p. 44), por el hecho de que el índice de aciertos no varíe de acuerdo con la complejidad de las estructuras y de los símbolos usados, de manera que Sarah acertaba lo mismo cualquiera que fuera la aparente dificultad del caso propuesto. Brown sugiere, a causa de ello, que los éxitos de aquélla acaso no obedezcan a una auténtica comprensión semiótica de la función de las piezas de plástico tanto como a su habilidad para desarrollar un cierto comportamiento operante en respuesta a ciertos estímulos que no se agotarían en las preguntas del experimentador sino que incluirían también determinadas características del entorno experimental. Es decir, que podría ser un caso semejante al de los pichones que fueron entrenados para hacer algo parecido al juego de ping-pong, aunque naturalmente no puede decirse que propiamente estuvieran jugando.

En última instancia, parece que lo que se ha ganado en cuanto a complejidad de las estructuras o a dificultad semántica de algunos de los símbolos empleados, se ha perdido en cuanto a la riqueza y autonomía de las formas de comportamiento adquiridas. En mi opinión, el balance final es más favorable a Washoe que a Sarah. Y ello por estas razones. El comportamiento de la primera era más activo y daba la clara impresión de estar más libre de las condiciones experimentales inmediatas. De otra parte, su lenguaje, aunque sintáctica y semánticamente era más sencillo, era más humano, y se utilizaba en más directa vinculación con las situaciones normales de su vida cotidiana (una vida que, en algún grado, constituía un reflejo de la vida humana). No hace falta añadir que Premack tiene una opinión considerablemente más alta de las capacidades de su chimpancé, y que algún autor, como Sampson, piensa incluso que los resultados del experimento de Sarah constituyen un contraejemplo para las afirmaciones de Chomsky reproducidas al comienzo de esta sección (Sampson, *The Form of Language*, p. 126). La justificación para los juicios de Premack y Sampson me parece por demás oscura.

Muy diferente es el experimento desarrollado posteriormente con otro chimpancé hembra, Lana, y sobre el que han informado Rumbaugh, Gill y von Glasersfeld («La lectura y el completado de oraciones realizados por un chimpancé», 1973). Aquí la comunicación del chimpancé se realizaba con una máquina, un ordenador de particulares características. Formaban parte de él consolas provistas de teclas de palabras, así como un dispositivo para suministrar las recompensas (comida, bebidas, juguetes, etc.). Una de

las consolas de teclas era accesible únicamente a los experimentadores, y servía a los fines que vamos a ver. Cada palabra del lenguaje utilizado estaba representada por un lexigrama consistente en un símbolo geométrico en blanco sobre un fondo de color, y al apretar la tecla correspondiente aparecía reproducido el lexigrama sobre un proyector a la vista del simio. Las teclas podían cambiarse de posición a fin de que el orden de colocación no ayudara a recordar su significado. El ordenador estaba programado de manera que, ante una sucesión incorrecta de símbolos, éstos desaparecían de los proyectores, mientras que si la secuencia era correcta y estaba completa, suministraba la recompensa apropiada. El entrenamiento de Lana incluyó, desde apretar teclas sueltas para solicitar determinadas recompensas, hasta formar oraciones apretando las teclas necesarias en el orden apropiado, al tiempo que aprendía que, no teniendo sentido continuar una oración en la que se ha cometido un error, el comportamiento adecuado era borrar el fragmento incorrecto y empezar de nuevo, como así lo hacía. Cuatro experimentos realizados después de seis meses de entrenamiento mostraron que el chimpancé era capaz de completar adecuadamente, con un alto índice de aciertos, las oraciones que los experimentadores iniciaban desde su consola con un primer fragmento válido, y de borrar aquellas expresiones que, por ser incorrectas, no servían para iniciar una oración. En la mayor parte de las pruebas, los aciertos de Lana estuvieron entre el 85 y el 100 por 100 de las contestaciones. Fueron de este orden, por ejemplo, las respuestas al fragmento «por-favor máquina dar» (compuesto de tres lexigramas), que ella podía completar añadiendo los lexigramas correspondientes a «zumo», o bien «trozo de plátano», etc. También aprendió a nombrar objetos, de tal modo que a la pregunta, que los experimentadores le hacían por medio del ordenador, «¿qué nombre-de esto?», Lana respondía de acuerdo con el objeto mostrado, por ejemplo: «plátano nombre-de esto». Se comprobó asimismo que era capaz de emplear los signos aprendidos en situaciones nuevas, y en el curso de intercambios comunicativos con los experimentadores llegó a preguntar por el nombre de objetos cuya denominación desconocía a fin de pedirlos por su nombre (estos últimos datos proceden de un informe posterior, «The Mastery of Language-Type Skills by the Chimpanzee (Pan)», por Rumbaugh y Gill, 1976).

El nuevo experimento ha probado, una vez más, que el chimpancé puede aplicar con corrección unas elementales reglas sintácticas, y que conecta adecuadamente la aplicación de estas reglas con la obtención de ciertas recompensas. Con relación a este vocabulario de lexigramas y a esta sencilla sintaxis, y sin olvidar que los aciertos solamente en algunos casos se acercan al 100 por 100, tal vez pueda afirmarse, con los autores del informe, que Lana ha aprendido a leer y escribir.

El último experimento de que tengo noticia en este campo nos devuelve al lenguaje manual de los sordomudos norteamericanos, pero nos saca del ámbito de los chimpancés. Todos los experimentos relatados tenían como sujetos a chimpancés, que son, desde luego, los primates no humanos usualmente reputados como más inteligentes y diestros. La responsable del nue-

vo experimento, Francine Patterson, ha recordado, sin embargo, que ya los esposos Yerkes habían mencionado, hacia 1929, la posibilidad de que los gorilas mostraran superior inteligencia en comportamientos más complicados que el manejo de utensilios o máquinas. La experiencia de Patterson parece haber venido a confirmar esta idea. Patterson escogió un gorila, también hembra, de nombre Koko, al que, durante seis años, y en condiciones muy parecidas a las creadas por los Gardner, ha estado enseñando el mismo lenguaje que éstos enseñaron a Washoe (Francine Patterson, «Conversations with a Gorilla», 1978). Después de los seis años de entrenamiento, Koko ha llegado a manejar trescientos setenta y cinco signos diferentes, lo que representa un progreso en comparación con los ciento treinta y dos que manejaba Washoe (si bien el entrenamiento de este último sólo duró cuatro años y medio, y habiendo sido el primer experimento de su género, hay que pensar que hubo de vencer mayores dificultades). Las habilidades semánticas y comunicativas de Koko parecen también más evolucionadas: no sólo hace preguntas y da respuestas, sino también afirma estar alegre o triste, se refiere a acontecimientos pasados y futuros, y da definiciones de objetos. Asimismo, muestra sentido del humor, insulta, y en ocasiones miente para evitar una reprehensión. Según testimonio de terceros, el gorila de este experimento, comparado con los chimpancés anteriores, parece hacer un uso más deliberado y consciente de los signos, y parece recurrir a ellos con más frecuencia de manera espontánea y los aplica a un conjunto más amplio de actividades. La experimentadora ha comentado algo que puede ser aún más importante: Koko traduce espontáneamente ciertas frases que oye en inglés (y que, es de esperar, son para ella habituales) a su lenguaje manual. A fin de desarrollar esta habilidad, se le ha construido un ordenador en el que, apretando las teclas correspondientes a ciertas palabras, el ordenador, que posee un sintetizador de voz humana, las replica oralmente. De esta forma, el gorila puede traducir ciertos signos de su lenguaje manual al inglés hablado con sólo apretar las teclas correspondientes. Las teclas representan nombres de objetos, sentimientos y acciones, así como preposiciones y nombres. Hasta el momento del informe, el teclado tiene cuarenta y seis palabras, pero se cuenta con ampliar este vocabulario. Y más todavía: posteriormente se ha incorporado al experimento un gorila macho más pequeño, y al parecer Koko toma parte activa en la enseñanza del lenguaje manual al nuevo compañero y se comunica con él por este medio.

El valor de estos experimentos salta a la vista. Con refinamientos cada vez mayores y sobre la base de los resultados de experimentos anteriores, cabe esperar una importante ampliación de nuestros conocimientos acerca de la capacidad semiótica, comunicativa y, en general, intelectual de los primates superiores. Esto puede tener importantes consecuencias para nuestra comprensión del proceso de aparición del lenguaje en los primeros homínidos. Que en condiciones experimentales particularmente ricas, y con la ayuda de seres humanos especialmente entrenados en métodos de condicionamiento de la conducta, y casi enteramente dedicados a estos experi-

mentos durante su tiempo de duración, los primates superiores puedan desarrollar unas habilidades y dar muestras de una capacidad que no aparecen en condiciones naturales, no lo considero sorprendente. Que estas habilidades, y la capacidad correspondiente, no pueden aún compararse con las propias del hombre, es patente. No se trata de negar que los antropoides de los experimentos comentados han aprendido un lenguaje, si por «lenguaje» entendemos un sistema de signos que utilizan para comunicarse con los experimentadores. Entre los lenguajes que se les ha enseñado hay importantes diferencias, así como en las condiciones del experimento. Las fichas de Sarah y las teclas de Lana parecen más semejantes entre sí, y su uso no requiere en tanta medida un proceso de socialización como el lenguaje de signos manuales de Washoe y Koko. Es este último, por ello, el tipo de lenguaje que, probablemente, hace más fácil la comparación con el aprendizaje lingüístico del niño, pues se trata de un lenguaje más natural y se da en un contexto más socializado (como han señalado Terrace y Bever, «What Might Be Learned from Studying Language in the Chimpanzee?», página 585).

No creo, sin embargo, que en ninguno de los casos relatados se pueda encontrar un contraejemplo a las afirmaciones de Chomsky que abrían esta sección. La peculiar clase de lenguaje que, en cada caso, se les ha enseñado, el porcentaje de fallos que parece existir siempre y la directa conexión con el contexto experimental al que se reduce la utilización de esos lenguajes, hace que nos hallemos ante el resultado de unos experimentos del máximo interés, que prueban lo que, en ciertas condiciones no naturales, puede dar de sí la capacidad psicológica de los primates superiores, pero que, por ahora, no prueban en absoluto que posean la capacidad de adquirir un lenguaje comparable, por su complejidad y por su riqueza semántica, a las lenguas humanas. Precisamente por ello, los responsables del experimento con Lana han concluido que los resultados de los estudios anteriores justifican una redefinición del lenguaje en la que, prescindiendo de una caracterización antropocéntrica, se subrayen los procesos psicológicos subyacentes y se adopte una perspectiva comparativa sobre la evolución de esos procesos en los primates superiores no humanos y en el hombre (Rumbaugh y Gill, «The Mastery of...», pp. 567 y 576). Pero lo cierto es que, por mucho que queramos aproximar entre sí esos procesos, no hay por el momento fundamento ninguno para atribuir a aquellos animales una capacidad lingüística comparable a la humana.

5.7 La tesis de la relatividad lingüística

En conformidad con la tendencia de la teoría lingüística y psicolingüística de los últimos decenios, nos hemos ocupado en general más con lo que pueda haber en común a las distintas lenguas que de aquello que las distingue. Así, se ha tratado en el capítulo cuarto de una serie de características y aspectos formales que la teoría actual postula como universales

y comunes para cualquier lengua humana; y hemos discutido en el capítulo presente las condiciones generales del proceso de adquisición de cualquier lengua, asumiendo que si la forma del lenguaje es la misma para todas las lenguas, y las capacidades biológicas y psicológicas son comunes a todos los individuos de la especie humana, no hay razón para esperar que el proceso de adquisición lingüística no siga también una pauta universal, en conexión con el desarrollo, igualmente universal, del proceso de adquisición de los conocimientos.

Hay, sin embargo, un tema que gozó de gran popularidad en los años cincuenta, y que contrasta vivamente con el enfoque anterior. Es el tema de hasta qué punto las diferencias sintácticas y semánticas entre las lenguas inducen diferencias correspondientes entre las representaciones de la realidad que se hacen sus usuarios. La tesis más conocida sobre esta cuestión es la que sugirieron por los años treinta Sapir y Whorf, y que se conoce como «tesis de la relatividad lingüística».

La idea de Sapir, que luego tomó Whorf y aplicó en sus estudios de las lenguas indias americanas, es que la realidad es algo en gran medida construido inconscientemente sobre los hábitos lingüísticos de la comunidad, y por ello, que diferentes comunidades no se limitan a categorizar de modo distinto el mismo mundo, sino que realmente viven en mundos diversos (Sapir, «The Status of Linguistics as a Science», 1929). Los hombres no viven, según esto, en un mundo natural y social objetivos, acerca del cual se comunican por medio de una u otra lengua; más bien ese mundo en que viven es, en gran parte, un resultado de la lengua que hablan, pues ésta condiciona la interpretación que se tiene, la manera que se posee de contemplar el mundo.

Una tesis como ésta, que pone en la cuenta de las peculiaridades lingüísticas aspectos tan importantes como la concepción del mundo propia de una cultura, y que implicaría la imposibilidad de traducir con mínima exactitud de una a otra lengua, tuvo cierto éxito durante algunos años, especialmente entre pensadores de fuerte propensión especulativa. El hecho de que la tesis contara con pruebas empíricas a su favor en un grado realmente ínfimo no parece que impidió su popularidad. El entusiasmo, sin embargo, duró poco. Unos años después de que se publicara, póstumamente, una selección de sus trabajos (*Language, Thought and Reality*, 1956), la tesis de Whorf había recibido ya críticas tan serias que su valor quedó inmensamente reducido y el interés por ella decreció con rapidez. Desde el lado de la filosofía, Max Black hizo un estudio detenido, y sumamente crítico, tanto de la tesis relativista como de otros aspectos de la doctrina y la metodología de Whorf (Black, «La relatividad lingüística: las opiniones de B. L. Whorf», 1959). Desde la parte de la sociolingüística, Fishman publicó una excelente sistematización de esta doctrina que acababa, como veremos, con unas palabras de clara cautela («A Systematization of the Whorfian Hypothesis», 1960). En fin, una obra más reciente de psicología del lenguaje, a la que ya hemos recurrido anteriormente, no concede interés a la tesis que nos ocupa como no sea en una formulación considerable-

mente más matizada y de alcance más reducido (Herriot, *Introducción a la psicología del lenguaje*, VI.8).

La cuestión es: ¿solamente se percibe del mundo aquello que está distintamente reflejado en las categorías sintácticas y semánticas de la lengua? Una posición afirmativa, llevada hasta el final, nos conduciría al absurdo de pretender que los alemanes pueden sentirse *gemütlich* pero no, en cambio, quienes, como los ingleses o los españoles, carecen de una palabra que corresponda a ésta exactamente (la comparación es de Fishman, *loc. cit.*, página 326). En realidad, lo que diríamos es que nosotros tenemos que expresar esa sensación de diferentes maneras según las circunstancias o acaso recurriendo a expresiones más complicadas. Esto no significa que las diferencias lingüísticas no tengan consecuencias. Fishman (p. 327) cita un estudio en el que se demuestra, aportando los textos oportunos, que el hecho de que los franceses sólo tengan una palabra para los dos términos ingleses «*conscience*» y «*consciousness*» es responsable de confusiones conceptuales que se encuentran con frecuencia mucho mayor en los filósofos franceses que en los ingleses. En este sentido, lo que ocurre no es tanto que se interprete la realidad de modo diverso cuanto que se codifique lingüísticamente el conjunto de las experiencias de manera distinta. Estas diferencias de codificación pueden obedecer a rasgos culturales muy genéricos que reproducen la forma básica de relación entre la comunidad de que se trate y el mundo. No es lo mismo una cultura basada en la ciencia moderna que una cultura centrada en el chamanismo. Como no cabe esperar que coincidan totalmente una cultura polar y una cultura tropical. No es de extrañar, por ello, que los esquimales posean una variedad de términos para designar diferentes estados de la nieve, variedad que carece de correspondencia en las lenguas occidentales. Esto no significa que nosotros no podamos percibir esas diferencias, sino que en nuestra cultura no tienen la relevancia que tienen para los esquimales. Pero las diferencias de la interpretación del mundo no son aquí función de las diferencias lingüísticas; más bien, éstas parecen el resultado de una diferente forma de ver el mundo que depende, entre otras cosas, de las condiciones en que se desarrolla la vida y la cultura de la comunidad. De la misma manera puede explicarse que para nuestro término «agua» los indios hopi posean dos palabras distintas según que el agua esté en reposo o en movimiento. O que los navajos tengan una sola palabra para los colores verde y azul y dos palabras distintas, en cambio, para dos clases de lo que nosotros llamamos negro. Absurdo sería pretender que ellos no perciban la diferencia entre el verde y el azul, como lo sería insinuar que nosotros no percibamos la distinción entre esas dos clases de negro o no distingamos el agua en reposo del agua en movimiento. Por lo que respecta a este aspecto de la tesis, Fishman (*loc. cit.*, p. 328) se ve obligado a concluir, a pesar de su simpatía por la obra de Whorf, que no hay pruebas que dependan de un verdadero análisis estructural del lenguaje ni de un análisis completo de otros elementos concomitantes no lingüísticos.

Sus conclusiones para otros aspectos o niveles de la tesis relativista no son mucho más positivos. Así, por lo que se refiere a las relaciones entre la estructura gramatical y la concepción del mundo, las dificultades son de análogo cariz. La ausencia de una clara distinción entre el nombre y el verbo en *nootka* lo relacionaba Whorf con la visión monista de la naturaleza propia de este pueblo, y si bien tal conexión no es, en principio, dudosa, parecería razonable pensar que el lenguaje expresa aquí esa concepción de la realidad y no que ésta sea una consecuencia de esa particularidad gramatical. Queda, no obstante, una cuestión metodológica principal: ¿cuál es la prueba de que una comunidad determinada tiene una cierta cosmovisión? Como Fishman ha reconocido (*loc. cit.*, p. 333), el único tipo de pruebas que suele proponerse consiste en la propia existencia de tales o cuales características lingüísticas. Con lo que, una vez más, se prescinde de las condiciones extralingüísticas, y la propia tesis lingüística de la relatividad se convierte en prueba de sí misma.

La conclusión final de Fishman (p. 337) es muy cautelosa. Según él, la relatividad lingüística, aunque afecte a una porción de nuestro comportamiento cognoscitivo, es solamente un factor moderadamente eficaz y en contra del cual es posible actuar. En un tono parecido y muy claramente se pronuncia también Pinillos: «Lo único que hasta el momento resulta claro es que ciertos hábitos semánticos aparecen asociados a la forma habitual de organización de la experiencia y pueden facilitar la práctica de ciertos comportamientos no verbales» («Lenguaje, individuo y sociedad», página 1120). Esto último está conectado con el hecho de que algunos experimentos, como los de Brown, Lenneberg y Roberts, han mostrado que existe una correlación entre la facilidad para reconocer colores y la posesión de términos diferenciados para designarlos. Es sin duda un mérito de la hipótesis de Whorf el haber dado impulso a este tipo de experimentos, pero los resultados disponibles por ahora son extremadamente parciales. Por eso, recordando unas palabras de Hockett citadas por Fishman (p. 335), hay que decir que las lenguas difieren no tanto en lo que se puede decir en ellas cuanto en lo que es *relativamente* fácil decir en ellas. Espíritu que recoge admirablemente —según pienso— esta breve sentencia de Pinillos: «hay una relativa relatividad lingüística» (*loc. cit.*, página 1123).

Lecturas

Varios de los temas tratados en este capítulo han sido discutidos con más extensión en mi libro *La teoría de las ideas innatas en Chomsky* (Labor, Barcelona, 1976), donde además se consideran otros temas relacionados, como el de la remota conexión que pueda existir entre la doctrina racionalista de las ideas innatas y las afirmaciones de Chomsky. En ese libro, sin embargo, no entré en ciertos aspectos psicológicos del tema que han sido objeto de atención en las páginas anteriores. El lector habrá no-

BIBLIOTEC

0.2

tado, asimismo, que he modificado alguna de las tesis presentadas en dicho libro.

Un excelente libro de conjunto sobre los aspectos psicológicos del aprendizaje lingüístico, y con discusión de las principales teorías, es *La adquisición del lenguaje*, de Marc Richelle (Herder, Barcelona). Para la teoría de Chomsky, además del primer capítulo de *Aspectos de la teoría de la sintaxis* (Aguilar, Madrid), debe verse *El lenguaje y el entendimiento* (Seix Barral, Barcelona), así como el artículo de Chomsky incluido en el *Simposio sobre la teoría de las ideas innatas*, traducido en un número de la revista *Teorema* del año 1973. Este simposio incluye, además, artículos críticos de Goodman y de Putnam. La famosa crítica de Skinner por Chomsky, junto con las réplicas de MacCorquodale y de Richelle, considerablemente más equilibradas, están traducidas en la interesante recopilación de Ramón Bayés *¿Chomsky o Skinner? La génesis del lenguaje* (Fontanella, Barcelona, 1977).

Para los aspectos psicológicos más generales de estos problemas es útil la *Introducción a la psicología del lenguaje*, de Herriot (Labor, Barcelona, 1977), y para la base biológica, los *Fundamentos biológicos del lenguaje*, de Lenneberg (Alianza Universidad, Madrid, 1975). Sobre los experimentos comentados en la sección 5.6 hay una excelente antología de trabajos realizada por Sánchez de Zavala, con el título *Sobre el lenguaje de los antropoides* (Siglo XXI, Madrid, 1976). Para el tema de la relatividad lingüística puede verse con provecho la segunda parte de *Lenguaje y conocimiento*, de Schaff (Grijalbo, México, 1967).

Se encontrará una interesante comparación general de la psicolingüística de Piaget con la de Chomsky en el trabajo de H. Sinclair de Zwart, «Psicolingüística evolutiva», incluido en la recopilación de Juan Delval, *Lecturas de Psicología del niño*, vol. 2, Alianza, Madrid, 1978.

una excelente recopilación de trabajos sobre Chomsky es la de G. Harman: *Sobre Noam Chomsky: Ensayos críticos*, Alianza Universidad, Madrid.

Capítulo 6

A LA BUSCA DEL LENGUAJE PERFECTO

¿Qué es, en efecto, conocer una cosa sino nombrarla? (UNAMUNO, *Amor y pedagogía*.)

6.1 Preámbulo

Se ha dicho que la principal característica de la filosofía analítica consiste en haber trasladado el nivel trascendental desde el conocimiento al lenguaje. El llamado «giro lingüístico» habría consistido en tomar como pregunta filosófica básica la pregunta sobre las condiciones que ha de cumplir necesariamente todo lenguaje. En este sentido ha podido escribir Habermas que «el paradigma del lenguaje ha conducido a una transformación de la forma del pensamiento trascendental» (*Erkenntnis und Interesse*, epílogo, secc. 6). A la pregunta «¿Cómo es posible el conocimiento?», pregunta que, de alguna manera, caracteriza a toda la filosofía moderna desde Descartes, y que recibe su más clara respuesta de Kant, habría sustituido, por obra de la filosofía analítica, esta otra: «¿Cómo es posible el lenguaje?».

Aunque la frase de Habermas va seguida por una referencia a Chomsky y al segundo Wittgenstein, es debatible que esta caracterización sea adecuada para ellos y cuadre bien a algunas de las últimas manifestaciones de la filosofía analítica. Donde, sin embargo, parece cabal y justa, es precisamente a propósito de los primeros filósofos analíticos, los que primeramente concedieron al lenguaje un papel central en la teoría filosófica e iniciaron la construcción de la teoría del significado. El iniciador de todo ello fue Frege. Apenas conocido en su tiempo, y sólo recientemente estudiado con detalle, Frege, uno de los pensadores más originales en el tránsito del siglo XIX al siglo XX, puso en marcha a la vez la lógica simbólica y la filosofía del lenguaje, y su herencia pesó no menos en la teoría del significado que en el desarrollo de los cálculos lógicos o en la filosofía

de la lógica. Por eso ha podido escribirse, con razón, que con Frege comienza una nueva época en la filosofía, una época cuya inflexión originaria fue colocar la teoría del significado como fundamento de la filosofía (Dummett, *Frege. Philosophy of Language*, cap. 19, pp. 669 y 684).

Ahora que comenzamos el desarrollo de la teoría del significado, y por ello el estudio de la filosofía analítica del lenguaje, es, por tanto, obligado empezar con Frege. Pero antes conviene examinar brevemente un precedente histórico importante de los conceptos de Frege que es Stuart Mill. Aunque las ideas de este último no constituyen todavía una teoría del significado con la entidad e influencia que adquirirán posteriormente conceptos parecidos en Frege, tienen por sí mismas interés suficiente y han sido de influencia dilatada.

6.2 Connotación y denotación

Una teoría referencialista del significado es una teoría que tiende a construir el significado de las expresiones primordialmente bajo la forma de una relación como la que hay entre un nombre y lo nombrado. Este tipo de teoría ha tenido gran peso, a menudo implícito, en la tradición filosófica, y acusa su presencia en la parte que al lenguaje dedica Stuart Mill en su gran obra *A System of Logic* (1843).

Según Mill (*op. cit.*, libro I, cap. i, seccs. 2 y 3), una proposición se forma poniendo juntos dos nombres, y afirmando o negando uno del otro. Se trata del modelo clásico de proposición formada por sujeto y predicado, en la que, por ejemplo, se afirma o niega una cualidad de una sustancia; para ello necesitamos un nombre de la sustancia y un nombre de la cualidad, que hagan, respectivamente, de sujeto y predicado de la proposición. El análisis de la proposición requiere, por esto, como prolegómeno indispensable, el estudio de los nombres, al que Mill dedica todo el capítulo segundo del libro primero de su libro.

Los nombres los toma Mill como nombres que se refieren directamente a las cosas, y no como nombres que lo fueran de nuestras ideas o conceptos de aquellas, concepción esta última que critica. De los nombres distingue aquellas palabras que considera partes de nombres (libro I, cap. ii, secc. 2). Como ejemplos de ellas cita preposiciones, adverbios, adjetivos, y los casos de la flexión nominal para los sustantivos que se declinan. Para los adjetivos, cuando éstos aparecen como predicados en las proposiciones, por ejemplo en «La nieve es blanca», Mill piensa que abrevian una expresión en la que aparezca propiamente un nombre, tal como ocurre en «La nieve es una cosa blanca». Estos términos que son partes de nombres eran los llamados en la tradición medieval —según recuerda Mill— términos «sin-categoremáticos», a diferencia de los nombres que eran llamados «categoremáticos». La distinción, por consiguiente, es plenamente tradicional, y aquí, como en otros lugares, Mill se limita a recogerla.

De las varias distinciones que traza entre los nombres, la fundamental para nuestro propósito es la que separa los nombres *connotativos* de los no connotativos (*loc. cit.*, secc. 5). Estos últimos son los que significan únicamente un sujeto o un atributo; los primeros son aquellos que denotan un sujeto e implican un atributo. Así, «Londres» o «Juan», son nombres no connotativos de sujetos; se limitan a denotar o referirse a algo sin implicar ningún atributo de aquello a lo que se refieren. Igualmente acontece con los nombres «blancura» o «virtud», que son nombres no connotativos de atributos; se limitan a denotar un atributo sin más. Nombres, en cambio, como «blanco» o «virtuoso» son connotativos. Y nótese bien lo que esto quiere decir. En primer lugar, que son nombres no de atributos sino de sujetos; «blanco», según esto, denota o se refiere a todas las cosas blancas, y «virtuoso», a todos los seres humanos virtuosos. Y en segundo lugar, que son nombres que implican o connotan un atributo, a saber, respectivamente, la blancura y la virtud. Es a causa de la posesión de estos atributos por lo que podemos llamar «blanca» a la nieve o «virtuoso» a Sócrates. Al decir «Sócrates fue virtuoso» nos referimos a Sócrates bajo dos nombres; primero, con un nombre singular no connotativo, «Sócrates», y segundo, con un nombre general connotativo, «virtuoso». Todos los nombres generales concretos son connotativos, según Mill. Nombres generales son aquellos que pueden aplicarse en el mismo sentido a un número indefinido de cosas; se oponen a los nombres singulares o individuales, que en el mismo sentido sólo se aplican a una cosa (*loc. cit.*, secc. 3). Nombres concretos son los que se refieren a cosas, y se contraponen a los abstractos, que denotan atributos (secc. 4). Nótese que, para Mill, son nombres concretos «Pedro», «hombre» o «blanco», mientras que son abstractos «humanidad» o «blancura». Ahora se entenderá por qué todos los nombres generales concretos han de tener connotación. En cuanto nombres concretos, se refieren a cosas o sujetos, y no a atributos; en cuanto nombres generales, denotan una variedad de tales cosas o sujetos, y esto sólo pueden conseguirlo por connotar determinado atributo o atributos comunes a dichas cosas. Así, «hombre» denota todos los seres humanos connotando aquellos atributos que les son comunes en cuanto tales, por ejemplo, vida animal y racionalidad (suponiendo que éstos sean los atributos que hacen de algo un ser humano).

Pero también los nombres abstractos, esto es, los nombres de atributos, pueden ser connotativos en ocasiones, de acuerdo con Mill. A saber cuando un atributo posee a su vez otro atributo, el nombre de aquel atributo puede connotar este segundo atributo (o atributo de segundo orden). Mill sugiere el siguiente ejemplo. Decir de una cualidad o atributo que es un defecto, en el sentido de «mala cualidad», es predicar de ese atributo una segunda propiedad. Así, al afirmar «La codicia es un defecto» se utiliza «defecto» para denotar la codicia y para connotar valor moral negativo.

Como ejemplos de nombres no connotativos Mill menciona, según acabamos de ver, nombres propios como «Londres» y «Juan». En efecto, lo que en el lenguaje ordinario consideramos como nombres propios

constituyen un grupo de nombres concretos y singulares (o individuales) que, para Mill, carecen de connotación. Con independencia de que haya alguna razón para imponer un nombre determinado a una cosa, la utilización del nombre propio es independiente de su razón, y su función se agota en aquella individualización de la cosa que la hace apta para ser sujeto del discurso. Así, parangonando un ejemplo de Mill, cabe pensar que el pueblo de Ribadesella continuaría llamándose así aun cuando el río Sella se secara o, por causas naturales o artificiales, desviara su cauce y discurriera alejado del pueblo. No utilizamos el nombre «Ribadesella» para connotar ninguna propiedad del pueblo (aunque de hecho la connote) sino para referirnos a él, para denotarlo.

Pero hay otro tipo de nombres singulares que sí se utilizan con una función claramente connotativa. Son los que posteriormente se llamarán «descripciones definidas»: aquellos nombres que, expresando algún atributo o atributos exclusivos de un objeto, por esto mismo denotan su objeto individualmente y sin posible confusión con otro. Por ejemplo: «El primer emperador de Roma», «El padre de Sócrates», «El autor de *El Quijote*», etc.

Para Mill, es la connotación lo que propiamente da el significado de un nombre: «cuando los nombres dados a los objetos suministran alguna información, esto es, cuando tienen *propriadamente algún significado*, el significado se halla no en lo que denotan sino en lo que connotan» (I. I, c. ii, secc. 5). Y puesto que los nombres propios no connotan nada, «carecen, estrictamente hablando, de significación» (*ibidem*). Hay aquí, como puede apreciarse, un importante correctivo a la teoría referencialista. Si bien las palabras significan en la medida en que son nombres o parte de nombres, su significado no consiste propiamente en nombrar, sino en connotar determinadas propiedades de aquello que nombran. Y así se da la paradójica conclusión de que las palabras que son paradigmas de un nombre, a saber, los nombres propios, carezcan de significado. Pues cuando señalamos a algo diciendo su nombre no suministramos al oyente información alguna excepto la trivial información de que tal es el nombre de aquello a lo que señalamos. Por lo que se refiere a los demás nombres, el problema es cómo determinar lo que cada uno connota, reconociendo aquí Mill que, en algunos casos, puede no ser fácil decidirlo, pues muchas expresiones tienen un significado vago, sobre el que la humanidad no ha alcanzado todavía acuerdo, y señalando que muchos malos hábitos de pensamiento proceden precisamente de la imprecisa connotación de tales palabras.

Con unas u otras variaciones, la distinción de Mill entre connotación y denotación va a estar presente en la teoría del significado posterior. En particular, son de subrayar dos doctrinas que sus sucesores rechazarán mayoritariamente. De un lado, la de que el significado consiste propiamente en la connotación; de otro, la de que los nombres propios carecen de significado, puesto que están faltos de connotación. Ambos puntos estarán en el centro de las discusiones posteriores, y sus consecuencias llegan hasta los tiempos más recientes, como puede advertirse en la

teoría de los nombres de Kripke, que en su momento estudiaremos. Es de notar, por cierto, que cuando Mill caracteriza los nombres individuales afirma, como vimos, que un nombre de este tipo no se aplica a más de un objeto en el mismo sentido (*loc. cit.*, secc. 3). Esto podría hacer pensar que Mill atribuye sentido a los nombres propios. En realidad, él mismo se encarga de deshacer de inmediato el posible equívoco: un nombre propio, como «Juan» o «Toledo», no se aplica a un objeto en ningún sentido, y por consiguiente, cuando se aplica a varios objetos, no puede decirse que se aplique *en el mismo sentido*, al contrario de lo que ocurre cuando llamamos a varios objetos «hombre» o «ciudad», nombres que, por ser generales, habremos de aplicar en el mismo sentido a todos los objetos a los que sean aplicables. La obvia objeción es, naturalmente, que si un nombre propio carece de sentido, entonces no puede afirmarse, como hace Mill, que el nombre propio se aplique en el mismo sentido sólo a un objeto, pues si no tiene sentido, entonces no podemos entender qué quiere decir en este contexto «en el mismo sentido». La afirmación de Mill únicamente sería aceptable para los nombres individuales que sí tienen sentido, a saber, las descripciones definidas. Hay aquí, sin duda, una forma defectuosa de expresión por parte de Mill. La teoría de los nombres propios no recibirá una formulación rigurosa hasta nuestros días, con Kripke, aunque la doctrina de este último tenga sus dificultades propias, como ya veremos. Lo único relevante ahora es avisar que tan reciente doctrina entronca directamente con estas páginas de Stuart Mill.

6.3 Sentido y referencia

En su artículo de 1892, «Sobre el sentido y la referencia», Frege formula en esbozo una teoría del significado que habría de ser muy influyente en los autores posteriores, de modo particular en Russell, Wittgenstein y Carnap.

Frege introduce sus conceptos a propósito de un planteamiento de la llamada «paradoja de la identidad». Si decimos que x es idéntico a y , ¿en qué medida difiere esto de afirmar que x es idéntico a x o que y es idéntico a y ? Por ejemplo: si decimos que el autor de la *Ética a Nicómaco* fue el preceptor de Alejandro Magno, queremos decir que las expresiones «el autor de la *Ética a Nicómaco*» y «el preceptor de Alejandro Magno» designan o denotan el mismo individuo, y en consecuencia podremos emplear cualquiera de ambas expresiones para referirnos a él, así como sustituir una por otra sin que varíe la verdad o falsedad de nuestras afirmaciones. Pero si esto es así, entonces, a partir de la afirmación:

- (1) El autor de la *Ética a Nicómaco* es el preceptor de Alejandro Magno

podremos obtener por sustitución esta otra:

- (2) El preceptor de Alejandro Magno es el preceptor de Alejandro Magno

La cuestión es que, mientras que (1) es una afirmación informativa, que, en principio, podría ser falsa, y que, en la medida en que es verdadera, a muchas personas les puede enseñar algo sobre Aristóteles, la afirmación (2), en cambio, no parece que pueda ser falsa, no transmite información alguna y no nos enseña absolutamente nada sobre Aristóteles ni sobre nadie. Mientras que (1) es una verdad empírica, de hecho, cuya constatación enriquece nuestro conocimiento histórico de cierto personaje griego, (2) es una verdad independiente de los hechos, ajena a nuestra experiencia, a nuestros conocimientos históricos, o, como, otros dirían, una verdad necesaria, o analítica. ¿Cómo es posible que de una afirmación empírica obtengamos una verdad analítica empleando expresiones que denotan, en ambas oraciones, el mismo objeto?

A esto responderá Frege: porque las expresiones utilizadas no se limitan a designar algo, sino que lo designan de un modo determinado, y es el modo de designar lo que las hace diferentes; pues si dos expresiones *x* e *y* no sólo designaran lo mismo, sino que además lo designaran de la misma manera, entonces el *valor cognoscitivo* de «*x* es idéntico a *y*» sería esencialmente igual al de «*x* es idéntico a *x*» o «*y* es idéntico a *y*» (en el supuesto, claro está, de que «*x* es idéntico a *y*» fuera verdadero). Tenemos, pues, que expresiones que denotan el mismo objeto o individuo pueden distinguirse por la manera como lo denotan. «El autor de la *Ética a Nicómaco*» denota la misma persona que «El preceptor de Alejandro Magno», pero la denota de modo diferente, así como el punto de intersección de tres rectas, *A*, *B* y *C*, puede ser denotado indistintamente por las expresiones «intersección de *A* y *B*», «intersección de *B* y *C*», o «intersección de *A* y *C*», aun cuando cada una de ellas lo denote de un modo levemente distinto.

A lo designado por una expresión, Frege lo llama «referencia» (*Be-deutung*), y esto lo distingue de lo que llama «sentido» (*Sinn*), «en el cual se halla contenido el modo de darse» la referencia. Esta última explicación del concepto de sentido es, sin duda, oscura; de momento, y para nuestros efectos, consideraremos el sentido como el modo o manera de designar que tiene una expresión. Hay, en los términos que emplea Frege, un pequeño problema de traducción que se debe mencionar. El término *Sinn* no parece que plantee problemas, pero sí el de *Bedeutung*. Una traducción ordinaria de este término debería dar como equivalente «significado», con lo que resultaría que Frege estaría distinguiendo entre sentido y significado de las expresiones. Estando perfectamente claro por sus afirmaciones que Frege entiende por significado o *Bedeutung* lo designado o denotado por una expresión, parece más aconsejable hacer fuerza al término original traduciéndolo como «referencia» que dificultar la comprensión de la doctrina de Frege traduciéndolo literalmente como «significado». Esta es, por otra parte, la forma usual de traducir dicho término. Geach y Black traducen en inglés *reference* (*Translations from the Philosophical Writings of Gottlob Frege*), que es sin duda la traducción más extendida para el término de Frege en esa lengua, aunque Church lo ha

traducido como *denotation* (*Introduction to Mathematical Logic*, secc. 01) y Carnap, más rebuscadamente, como *nominatum* (*Meaning and Necessity*, secc. 28), pero ninguna de estas traducciones ha hecho tanta fortuna, a pesar de que *denotation* es ya un término utilizado por Russell. La mejor traducción castellana de escritos de Frege que conozco, traduce también *Bedeutung* por «referencia» (Frege, *Estudios sobre semántica*, traducidos por Ulises Moulines), aunque en una importante antología de textos de semántica se traduce como «denotación» (*Semántica filosófica: problemas y discusiones*, recopilación de Thomas Moro Simpson, traducción que este último ya había usado, siguiendo a Church, en *Formas lógicas, realidad y significado*). Hay, finalmente, otra recopilación de escritos de Frege que se aparta extrañamente de las traducciones anteriores para traducir *Bedeutung*, literalmente, como «significado» (Frege, *Escritos lógico-semánticos*). Puesto que esta traducción puede resultar muy confundente, y no se ve qué ventajas tenga, aquí seguiremos usando el término «referencia», que es el más extendido, aunque para el verbo emplearemos con frecuencia «denotar».

Frege aplica su distinción, en primer lugar, a las expresiones que denotan un objeto único, las cuales considera, en sentido amplio, nombres propios. Incluyen tanto lo que en el discurso ordinario se llama estrictamente «nombre propio» como lo que, desde Russell, se llamará «descripción definida». Es un nombre propio «Aristóteles», y una descripción definida «El preceptor de Alejandro Magno» (en la cual, por cierto, está contenido otro nombre propio, «Alejandro Magno»). Son descripciones definidas las expresiones «el lucero vespertino» y «el lucero matutino», pero es un nombre propio, que designa el mismo objeto que aquellas, «Venus».

Por lo que respecta a los nombres, o términos de individuos, los conceptos de sentido y referencia funcionan de forma semejante a como funcionan los de connotación y denotación de Mill, pero con una importante diferencia que, en diversas formas, estará presente en autores posteriores. Mill —como hemos visto— había afirmado que los nombres propios del lenguaje ordinario poseen denotación pero carecen de connotación, o más exactamente, que no tienen por qué connotar nada para funcionar como nombres propios, y por lo tanto, que cuando tienen connotación esto es sólo una característica accidental a su condición de nombres propios. Frege, por el contrario, se aparta de la posición de Mill. Para Frege, todo el que conoce un lenguaje conoce el sentido de los nombres que hay en él, y esto se aplica igualmente a los nombres propios. La cuestión es ésta: puede ser relativamente fácil llegar a un acuerdo sobre el sentido de una descripción definida como «el lucero matutino» o «el preceptor de Alejandro Magno», pero ¿cuál es el sentido de un nombre propio como «Aristóteles», «Venus» o «Alejandro Magno»? Frege responderá en la segunda nota de su artículo que se trata de algo sobre lo que puede haber opiniones divergentes. Así, para algunos, el sentido de «Aristóteles» puede venir dado por la expresión «El preceptor de Alejandro Magno»; para otros, por «El

filósofo griego nacido en Estagira»; para otros más, por «El autor de la *Ética a Nicómaco*», etc. Pero mientras la referencia no varíe —dirá Frege— estas diferencias de sentido son tolerables, *aunque no deberían aparecer en un lenguaje perfecto*.

Esta manera de hablar del sentido de los nombres propios puede inducir a cierta confusión sobre la noción de sentido. Parece claro que Frege admite que el sentido que demos a un nombre propio dependerá de nuestros conocimientos sobre el objeto o individuo designado por tal nombre. Mas esto no debe hacer pensar que el sentido consista en, o se confunda con, nuestra representación del objeto. La representación (*Vorstellung*) que cada cual se haga de un objeto, será algo individual y subjetivo, propio de uno mismo y fundado en sus experiencias cognoscitivas y en su memoria. El sentido (de un nombre propio o de una descripción) a través del cual la expresión se refiere al objeto, no es subjetivo ni individual, antes bien, es perfectamente objetivo, en cuanto perteneciente a una realidad objetiva e independiente de la mente individual como es el lenguaje. Pienso que puede afirmarse con exactitud que, para Frege, el sentido es, en definitiva, condición necesaria para que el lenguaje tenga referencia. En esto se distingue claramente de Mill, pues para éste la connotación no es condición para que haya denotación.

Condición necesaria, pero no suficiente. Puesto que una expresión puede poseer sentido, pero carecer de referencia. ¿Por qué? Porque una expresión tiene sentido en cuanto que expresa un modo de designación de un objeto, pero nada se opone a que tengamos maneras múltiples de designar, a las cuales no corresponda en la realidad objeto alguno. La expresión «El asesino de Aristóteles» tiene sentido porque expresa una forma de designar un posible objeto, en este caso, una persona, pero no tiene referencia puesto que, según nuestros conocimientos históricos, nadie asesinó a Aristóteles. Parece, pues, que el ámbito del sentido crea el ámbito para la posibilidad de la referencia. La efectiva determinación de la referencia, sin embargo, es cuestión extralingüística: requiere ir a la realidad y comprobar si hay los objetos a los que nuestros modos de designación aluden. Aquí podemos aprovechar para hacer una importante puntualización sobre lo que Frege entiende por «objeto» en su teoría del significado. No son objetos solamente las realidades físicas, como los organismos, las personas, las cosas, o sus componentes físico-químicos, sino que también son objetos las entidades matemáticas, como los puntos, líneas, figuras, las diferentes clases de número, etc. Incluso la verdad y la falsedad, entendidas como luego veremos, son objetos. Frege contrasta los objetos con las funciones. Los objetos constituyen la referencia de los nombres, y los nombres (en el sentido amplio propio de Frege) son expresiones completas que incorporan un sentido, esto es, una manera de darse la referencia, el objeto. Las funciones, por el contrario, son designadas por expresiones incompletas, o como dice Frege, no saturadas; las funciones incluyen los conceptos y las relaciones (véanse sus artículos «Función y

concepto» y «Sobre concepto y objeto», incluidos, igual que «Sobre sentido y referencia», en *Estudios sobre semántica*).

Otro aspecto en el que puede advertirse la insuficiencia del sentido para la determinación de la referencia es el siguiente. Hasta ahora hemos visto casos en los que distintas expresiones determinaban, a través de sus respectivos sentidos, la misma referencia. Pero ocurre que una misma expresión con un único sentido puede designar objetos distintos, como es el caso de aquellas expresiones que modifican su referencia de acuerdo con el contexto extralingüístico, como por ejemplo: «El abajo firmante», «El que ahora está hablando», etc.

Como acabamos de ver, los conceptos son designados por cierto tipo de expresiones incompletas, a saber, por aquellas expresiones que funcionan como predicados en la oración. También a estas expresiones extiende Frege la distinción entre sentido y referencia (no habla de ello en el artículo principal, que estamos comentando, pero sí en «Sobre concepto y objeto», del mismo año, 1892, y más todavía en un artículo inmediatamente posterior, «Consideraciones sobre sentido y referencia», que se hallaba inédito y ha sido publicado en 1969 en un volumen de escritos póstumos, hallándose incluidos en *Estudios sobre semántica*). La idea de Frege es que una oración asertórica o declarativa puede, según el análisis tradicional, descomponerse en dos partes, sujeto y predicado, y que estas dos porciones se distinguen en que la primera es completa en sí misma, la segunda, el predicado, incompleta o no saturada. Esto quiere decir que el sujeto, un nombre, tiene sentido completo por sí mismo; el predicado, en cambio, lleva consigo un lugar vacío, y sólo cuando un nombre ocupe ese lugar adquirirá un sentido completo. El concepto designado por el predicado es, por ello, una función que tiene como argumento el objeto designado por el sujeto, y que adquiere como valores los dos valores veritativos, verdad y falsedad. Frege suministra el siguiente ejemplo («Función v concepto», p. 32 de *Estudios sobre semántica*).

César conquistó las Galias

El sujeto es «César», nombre que tiene una referencia (histórica) y un cierto sentido (acerca del cual puede haber divergencias, como vimos a propósito de «Aristóteles»). El predicado es «conquistó las Galias», expresión que designa o tiene como referencia un cierto concepto. Este concepto es una función que, teniendo como argumento el objeto designado por el nombre «César», adquiere el valor verdad. Si tuviera como argumento un objeto distinto, por ejemplo, el personaje romano denotado por el nombre «Marco Antonio», la función adquiriría el valor falsedad.

Con esto queda dicho cuál es la referencia de un predicado o término conceptual: es un concepto. Y un concepto es una función de un argumento cuyo valor es un valor veritativo. Pero ¿cuál es el sentido de un predicado? Frege ha dejado este punto sumido en la oscuridad. Por analogía con el sentido de un nombre, podríamos pensar que el sentido de un término con-

ceptual es el criterio que nos permite decidir si el término puede predicarse con verdad o no de cierto objeto. Dicho de otra forma: el criterio que nos permite decidir si, para un argumento determinado, la función designada por el término posee el valor verdad o el valor falsedad. Con relación al ejemplo mencionado, el sentido del predicado «conquistó las Galias» sería el criterio que nos permita decidir que para el nombre «César» la oración es verdadera y para el nombre «Marco Antonio» es falsa. Esto tal vez no sea mucho decir, pero Frege no ha aportado más claridad a esta cuestión.

En qué medida la clasificación semántica de las expresiones y la aplicación de los conceptos de sentido y referencia dependen, para Frege, de la forma de las expresiones y de que funcionen en la oración como sujetos o como predicados, puede comprobarse por una complicación ulterior que Frege no duda en introducir en su doctrina. Acabamos de ver que un concepto es la referencia de un predicado. Así, la referencia del término predicativo «triángulo» es el concepto de triángulo. Pero ¿cuál es la referencia de la expresión «el concepto de triángulo»? Según Frege, no un concepto, sino un objeto («Sobre concepto y objeto», p. 109 de *Estudios sobre semántica*).

Y no puede ser de otro modo puesto que la expresión «el concepto de triángulo» no es una expresión predicativa, no funciona como predicado en una oración, sino que funciona como sujeto, y su propia forma lingüística, ese comienzo con «el», indica que no designa una función. Por consiguiente, al hacer una afirmación sobre el concepto de triángulo, al decir, por ejemplo, «El concepto de triángulo es muy sencillo», predicamos algo, la sencillez en este caso, de un objeto de cierto tipo peculiar. Al llegar aquí, uno puede preguntarse impaciente: Y bien, pero ¿qué es un objeto? A lo cual, Frege sólo puede responder apelando de nuevo al criterio lingüístico. Un objeto es algo que pertenece a una categoría última del análisis, y que, en consecuencia, no puede ser ulteriormente analizado ni admite descomposición lógica. Y por tanto no puede ser definido («Función y concepto», p. 33 en *op. cit.*). Lo único que puede decirse es que objeto es todo aquello que, a diferencia de una función, es designado por una expresión completa, por una expresión que no muestra ningún lugar vacío, por una expresión que funciona como sujeto en una oración. Y éste es ciertamente el caso de cualquier expresión de la forma «El concepto de *F*».

Hemos visto, pues, en qué consiste el sentido y la referencia tanto para el sujeto de la oración, los nombres, como para el predicado, los términos conceptuales. Pero Frege se pregunta, además, por el sentido y la referencia de la oración, como tal. A ello dedicará la segunda mitad de su trabajo «Sobre el sentido y la referencia».

Lo primero que hay que observar aquí, es que las oraciones cuyo sujeto carece de referencia no por eso dejan de ser inteligibles, o de expresar algo. En la afirmación «Don Quijote arremetió contra los molinos de viento», es claro que el nombre que hace de sujeto no tiene referencia (*tal y como Frege ha entendido la referencia*), y sólo por esto hay que negar que la oración, como tal y en conjunto, la tenga. Pero es una afirmación

que entendemos. Sabemos lo que quiere decir, e incluso, sobre la base de otras afirmaciones que conocemos sobre ese personaje, podemos representarnos, mental o gráficamente, el contenido de esa oración, aun cuando a tal contenido nunca haya correspondido nada en la realidad del mundo. La oración tiene un sentido. El sentido consiste —según Frege— en el pensamiento que expresa, teniendo en cuenta que llama «pensamiento», no al acto subjetivo de representarse el contenido de la oración, sino a este contenido que diferentes personas en diferentes momentos pueden representarse en acontecimientos mentales distintos.

¿Cuál es la consecuencia de que el sujeto de esa oración no tenga referencia? Que no podemos preguntarnos si el concepto designado por el predicado «arremetió contra los molinos de viento» se aplica correctamente o no al objeto designado por el sujeto (puesto que tal objeto, que sería la referencia, no existe). Dicho en los términos técnicos de Frege: que la función designada por el predicado carece de argumento, y por consiguiente, también de valor. Y puesto que los valores posibles de esa función serían los dos valores veritativos, esto significa que no podemos preguntarnos por la verdad o la falsedad de esa oración. Una oración cuyo sujeto carezca de referencia no es ni verdadera ni falsa. Es, por tanto, la referencia del sujeto la que nos permite asignar un valor veritativo a la oración, y es, sin duda, esta conexión entre aquella referencia y dichos valores, la que conduce a Frege a completar su teoría del significado estableciendo que la referencia de una oración es precisamente su valor veritativo. Pues, como escribe: «¿Por qué queremos que cada nombre tenga no sólo un sentido sino también una referencia? ¿Por qué no es suficiente el pensamiento? Porque, y en la medida en que, lo que nos interesa es el valor veritativo». («Sobre sentido y referencia», p. 59 de *Estudios sobre semántica*).

La consecuencia es extraña, pero, dentro de su teoría, armónica. Ya que la referencia de una oración es su valor veritativo, todas las oraciones verdaderas tendrán la misma referencia, la verdad; y todas las oraciones falsas poseerán asimismo referencia idéntica, la falsedad. Y puesto que la referencia es, para el sujeto de la oración, el objeto designado por el nombre, Frege no dudará en dar el siguiente paso: los valores veritativos son objetos, y las oraciones son sus nombres. Todas las oraciones verdaderas son nombres de lo verdadero, y todas las oraciones falsas son nombres de lo falso. Pero ¿en qué consisten la verdad y la falsedad como objetos? ¿Qué es un valor veritativo? Frege aquí se limitará a decir: el valor veritativo de una oración es la circunstancia de que sea verdadera o falsa (*op. cit.*, página 60). Y ciertamente esto no es aclarar mucho.

Frege se ocupa de subrayar que su posición cumple con el principio leibniziano de poder sustituir, en una oración, una expresión por otra con la misma referencia, sin que varíe la referencia de la oración, esto es, su valor veritativo. No cambia la verdad de la afirmación «El preceptor de Alejandro Magno escribió la *Ética a Nicómaco*», si en su lugar decimos «El filósofo griego nacido en Estagira escribió la *Ética a Nicómaco*». La sustituibilidad de expresiones *salva veritate* suministra, así, un apoyo in-

directo a la doctrina de Frege. Puesto que el valor veritativo de una oración no varía al sustituir sus expresiones por otras que posean la misma referencia, esto parece confirmar que es correcto considerar el valor veritativo como la referencia de la oración.

Hay, sin embargo, un caso en el que las oraciones no tienen como referencia su valor veritativo: cuando aparecen como oraciones subordinadas en el estilo indirecto. Consideremos el ejemplo de Frege (*op. cit.*, p. 65):

(3) Copérnico creía que las órbitas de los planetas son circulares

Según Frege, la oración subordinada contenida aquí, esto es, la oración:

(4) Las órbitas de los planetas son circulares

no se refiere, tal y como aparece en (3), a su valor veritativo, y la razón es que la verdad o falsedad de la oración conjunta (3) es independiente del valor veritativo de la subordinada, es decir, de (4). Es verdadero o falso que Copérnico creía eso, con independencia de que sea verdadero o falso el contenido de su creencia. Por ello, la sustitución de la oración subordinada por otra con el mismo valor veritativo no garantiza que la oración conjunta conserve su valor veritativo. Para Frege, una oración, en su uso indirecto o subordinado, adquiere como referencia el pensamiento que expresa, de manera que lo que es su sentido en el uso directo o principal, pasa a ser su referencia en el uso indirecto. En el ejemplo anterior, la referencia de la subordinada en (3) es el sentido que esta propia oración tiene cuando se utiliza fuera de un contexto indirecto, como en (4). ¿Y cuál es su sentido cuando está en un contexto indirecto como (3)? Lo que dice Frege es oscuro: el sentido de una oración en el estilo indirecto no es un pensamiento, sino el sentido de las palabras «el pensamiento de que...», frase que no contiene un pensamiento, puesto que no constituye una oración.

Hemos visto antes que los nombres pueden tener sentido y carecer de referencia. Lo propio acontece con las oraciones. Una oración declarativa cuyo sujeto no tenga referencia carecerá de valor veritativo, no será ni verdadera ni falsa, y en consecuencia, estará falta de referencia. Pero esto no significa que no tenga sentido, pues puede muy bien, a pesar de su falta de valor veritativo, expresar un pensamiento. Es lo que ocurre en la ficción literaria. No necesitamos, para comprender y gozar de una novela, poesía u obra teatral, que sus nombres posean referencia, que denoten personas, lugares o cosas existentes, y, por consiguiente, no nos planteamos si las afirmaciones contenidas en la obra literaria son verdaderas o falsas. Lo relevante para el goce estético es el sentido, más otras cosas, evidentemente más subjetivas, como lo que Frege llama alguna vez «la matización o coloración del sentido», que es algo en lo que pueden diferir varias expresiones que, sin embargo, posean idéntico sentido (véase la nota séptima de «Sobre concepto y objeto»). Oraciones que, como

las citadas, carezcan de referencia, simplemente serán irrelevantes para la investigación científica así como para el cálculo lógico, pues en ambos casos es cuestión central la del valor veritativo de las oraciones que se manjean.

Pero hay otro caso más en el que nos tropezamos con oraciones sin referencia, esto es, sin valor veritativo. A saber, cuando estamos ante oraciones que no son declarativas, ante oraciones que no pretenden decir nada sobre los hechos, los objetos, sus propiedades o sus relaciones, sino que poseen diferente función en el lenguaje. Es el caso de los imperativos. Los imperativos no expresan pensamientos, en el estricto sentido que Frege ha dado a este último término. El sentido de un imperativo consiste en otra cosa: puede ser un ruego, una petición, un mandato, una prohibición... Pero nada de esto es un pensamiento, porque no se trata de contenidos que puedan ser verdaderos o falsos («Sobre sentido y referencia», p. 67 de *Estudios sobre semántica*). Por esto añadirá Frege que un mandato, aunque no sea un pensamiento, está al mismo nivel que éste; pues, en efecto, un mandato es el sentido de un cierto tipo de oraciones, como un pensamiento constituye asimismo el sentido de otro tipo distinto de oraciones. Dicho de otra forma, que los sentidos de las oraciones pueden consistir en pensamientos, ruegos, mandatos, etc. Huelga añadir que las oraciones cuyo sentido es de alguna de estas últimas clases (y tal es el caso de las oraciones imperativas), carecen de referencia, puesto que no pueden ser ni verdaderas ni falsas (véase también su posterior trabajo «Der Gedanke», en *Kleine Schriften*).

Para que podamos preguntarnos por la verdad o falsedad de una oración debe, pues, tratarse de una oración declarativa en primer lugar, y segundo, debe ser una oración cuyo sujeto y predicado tengan referencia. El hecho de que en el lenguaje común se utilicen con frecuencia oraciones cuyo sujeto no denota nada, es considerado por Frege como una imperfección lógica que aleja al lenguaje ordinario de lo que sería un lenguaje lógicamente perfecto («Sobre sentido y referencia», *op. cit.*, páginas 69 y ss.). Pues un lenguaje lógicamente perfecto es un lenguaje en el que cada oración tiene un valor veritativo, y esto *presupone* que los nombres que aparecen en la oración tienen referencia. Dicho en palabras de Frege: «Un lenguaje lógicamente perfecto (conceptografía) debe cumplir la condición de que toda expresión gramaticalmente bien construida como nombre a partir de signos ya introducidos, designe realmente un objeto, y que no se introduzca un nuevo signo como nombre sin que se le asegure una referencia» (*loc. cit.*). Más adelante veremos la diferente solución que da Russell al problema de las expresiones sin referencia; muchos años después, criticando a Russell, Strawson mantendrá una posición semejante a la de Frege.

Al final del artículo que estamos comentando, Frege vuelve sobre la paradoja de la identidad, con la que había iniciado su discurso. La solución resulta, ahora, en extremo simple. La afirmación «x es idéntico a y» difiere de «x es idéntico a x» en la medida en que expresan pensamientos distintos, y esto último ocurre en tanto en cuanto «x» e «y», aun teniendo

la misma referencia, tengan sentido diverso. La diferencia entre ambos enunciados es una diferencia en lo que llama Frege «valor cognoscitivo»; es el tipo de diferencia que hay entre las oraciones que tomábamos como ejemplo al principio, a saber:

- (1) El autor de la *Ética a Nicómaco* es el preceptor de Alejandro Magno y
- (2) El preceptor de Alejandro Magno es el preceptor de Alejandro Magno

Ambas oraciones son verdaderas, pero la primera establece una identidad valiéndose de nombres que, aun cuando con referencia idéntica, tienen sentido distinto. No es lo mismo designar a una persona como autor de cierto libro de ética, que designarla como preceptor de determinado emperador. En cambio, (2) establece la identidad sirviéndose de nombres con idéntica referencia e idéntico sentido, con lo que se convierte en una verdad analítica o tautológica, a diferencia de la verdad empírica, de hecho, que aparece en (1). Esta es la diferencia de valor cognoscitivo a que alude Frege, y que procede de los diferentes pensamientos expresados en ambas proposiciones. Resultará patente, no obstante, que este pequeño problema no es sino una excusa para el desarrollo de unos conceptos, sentido y referencia, cuya aplicación va mucho más allá de estos límites, y cuya influencia vamos a comprobar en una gran parte de la teoría del significado a lo largo del siglo xx.

Durante la primera mitad del siglo, sin embargo, esa influencia fue implícita e indirecta; fundamentalmente se ejerció a través de Russell, y así está presente en Wittgenstein. Sólo después de 1940 se encuentran estudios detenidos sobre los conceptos semánticos de Frege, primero en Church, que fue quien rescató el artículo de Frege de su olvido, y luego en Carnap. Estas consideraciones han solido ir acompañadas de críticas constantes, pero éstas han sido a veces extremadamente confusas. Así, la crítica que Russell le hace (en «On Denoting») resulta difícilmente inteligible, y probablemente encierra alguna confusión de fondo por parte de Russell (véase el cap. 2.2 de la obra de Thiel, *Sentido y referencia en la lógica de Gottlob Frege*). Tampoco las críticas de Carnap (en la secc. 30 de *Meaning and Necessity*) parecen todas ellas justificadas, aunque acaso el método de análisis de Carnap sea superior al de Frege; esto lo veremos en su momento, y aquí no entraremos en detalle sobre estas críticas. Resumiré únicamente los aspectos que me parecen claramente más débiles en los conceptos de Frege.

Primero de todo, es claro que la referencia, tal y como aparece en muchas de las afirmaciones anteriores, no es propiamente parte del significado (tomado este término en un sentido amplio). Si la referencia de un nombre es un objeto, y la de un predicado, una función, esto son realidades extralingüísticas a las que conectamos el lenguaje, como lo serían los valores veritativos para la oración. De alguien que no conozca la referencia de un

nombre o que no sepa si una cierta afirmación es verdadera o falsa, no diríamos que no sabe lo que quiere decir el nombre o la oración en cuestión (v. Dummett, *Frege. Philosophy of Language*, pp. 84 y 91). Se puede no saber a quién se refiere el nombre «el preceptor de Alejandro Magno», entendiendo perfectamente el significado de esta expresión. O se puede ignorar si es verdadera o falsa la oración «El preceptor de Alejandro Magno escribió la *Ética a Nicómaco*», sabiendo bien lo que ésta significa. Las entidades designadas por las expresiones lingüísticas, sean aquellas materiales o intelectuales, concretas o abstractas, no son, en sentido propio, parte de lo que las expresiones significan. Lo que sí se puede tomar como parte de su significado es la función que cumplen las expresiones de remitir o referir a entidades extralingüísticas. Algunas de las afirmaciones de Frege pueden, probablemente, entenderse de esta forma, pero es patente que sus palabras ocultan, en general, la distinción que estoy apuntando. Se trata, simplemente, de la diferencia que hay entre la función denotativa o referencial de las expresiones, y lo denotado en el cumplimiento de tal función. Yo preferiría reservar el término «referencia» o «denotación» para esa función, y llamar a las entidades designadas, objetos o lo que que quiera que sea, «lo denotado». Esto evitaría la dificultad que se acaba de apuntar sobre la terminología de Frege, pues ahora resulta obvio que lo denotado no es, ni puede ser, parte del significado de las expresiones, pero no hay por qué negar que lo sea la función denotativa que éstas poseen. Se trata, en suma, de la distinción que, a otro propósito y dentro de la teoría semiótica, hicimos en el capítulo segundo entre *el significado* de los signos y *lo significado* por los signos. Entonces, esta terminología nos bastaba para nuestros fines. A partir de ahora, en que hemos entrado en posesión de los conceptos de Frege, debemos evitarla. En lugar de «lo significado» digamos «lo denotado». Y en el significado empecemos por distinguir ingredientes. De momento tenemos ya la función referencial o denotativa (abreviadamente: la referencia) y el sentido.

Lo que contribuye a determinar la referencia, en mayor o menor medida según los casos, es el sentido. Es curioso, por ello, que este concepto quede tan confuso en la doctrina de Frege. Primero, resulta sumamente problemática su aplicación a los nombres propios en sentido estricto. Según lo que hemos visto, el sentido del nombre «Aristóteles» variará para las personas según lo que cada cual sepa de ese personaje, y de acuerdo con las características que elija para identificarlo como lo denotado por el nombre en cuestión. Que con esto se subjetiviza el sentido y se lo saca del ámbito de lo lingüístico parece patente, a pesar de las intenciones de Frege. En su momento veremos algunos esfuerzos realizados por hacer más plausible la tesis de Frege; como ya explicaré entonces, pienso que sobre este punto hay que darle la razón a Mill, si se quiere ser fiel al modo como funcionan los nombres propios en el lenguaje ordinario. Que para identificar a la persona denotada por el nombre «Aristóteles» cada cual recurra a lo que sepa de ella, se comprende. Preguntado sobre quién fue Aristóteles, uno puede decir: el discípulo de Platón nacido en Estagira; otro: el autor

de la *Ética Nicomachea*; un tercero: el filósofo preceptor de Alejandro Magno; etc. Y quien sepa lo suficiente elegirá, entre estas y otras muchas más, la identificación que más le guste. Pero ¿qué se soluciona o se aclara manteniendo que cualquiera de estas descripciones constituye el sentido del nombre «Aristóteles»? ¿Es que para un especialista en Aristóteles, este nombre posee varios miles de sentidos? ¿Acaso adquirió el nombre «Aristóteles» (o para ser más preciso: su equivalente griego) un sentido que antes no tenía cuando este filósofo hubo escrito el libro que conocemos como *Ética Nicomachea*? Los problemas que plantea esta teoría son más que los que soluciona, y al dar entrada al sentido como ingrediente del significado de los nombres propios, se aparta decisivamente de la función que éstos cumplen en el lenguaje ordinario, que es la de denotar un objeto señalándolo entre los demás. La función de los nombres propios se agota en la referencia. Que muchos nombres propios tengan algún sentido o connoten alguna característica, muchas veces por razón de su etimología, es cierto, pero accidental. Así, por ejemplo, la mayoría de los nombres de pila que usamos en muchas lenguas connotan masculinidad o feminidad, puesto que se aplican a hombres o a mujeres pero no a ambos. Pero esto es totalmente accidental con respecto a la función del nombre, y aquellos nombres que se aplican indistintamente a hombres y mujeres no cumplen su función peor que los otros. Volveremos sobre esta cuestión con más detalle cuando, en un capítulo ulterior, estudiemos la teoría de Kripke.

Por lo que respecta al sentido de los términos conceptuales, la posición de Frege es, no ya objetable, sino —como hemos visto— del todo oscura. Si el concepto es lo denotado por el término conceptual, ¿cuál es su sentido? Frege no da una respuesta clara, y tampoco resulta fácil imaginarla.

Donde más claro resulta el sentido, aparte de las descripciones definidas, es en las oraciones. Que el sentido de una oración sea el pensamiento expresado por ella, o como se dirá posteriormente: la proposición contenida en ella, no parece presentar dificultades. Pero no resulta tan fácil de entender que las oraciones hayan de poseer también una función referencial propia, y menos aún que ésta consista en denotar un valor veritativo. Puesto que, como Frege ya ha aceptado, las oraciones no declarativas carecen de referencia, y puesto que en el mismo caso están las oraciones cuyo sujeto no la tiene, ¿por qué empeñarse en hacer un caso aparte de las oraciones que son verdaderas o falsas? No se ve bien la necesidad de ello, sobre todo si la consecuencia es convertir tales oraciones en nombres de esos extraños objetos que son los valores veritativos. Como vamos a ver, Wittgenstein empezó por negar que las oraciones sean nombres de algo, y Russell reconoce que fue aquel quien le convenció de ello.

Para acabar, subrayemos el interés de Frege por lo que llama un «lenguaje lógicamente perfecto». Encontraremos de nuevo este interés en Russell. Y notemos que un lenguaje así es, para Frege, no ya un lenguaje que cumpla con determinadas reglas formales, sino, más aún, un lenguaje en el que todas sus expresiones tengan referencia, y por tanto, un lenguaje

conectado en todos sus puntos con la realidad. En consecuencia, un lenguaje cuyas oraciones serán todas o verdaderas o falsas.

6.4 El atomismo lógico

El propósito de Russell es semejante al de Frege, y análoga la justificación de su interés por las condiciones que ha de cumplir un lenguaje para alcanzar la perfección lógica. Pero en Russell, la reflexión se da en un contexto filosófico más rico y logra un grado de elaboración más alto. En la doctrina de Russell, tanto los supuestos epistemológicos como las consecuencias metafísicas poseen una riqueza y tienen una explicitación del todo ausentes en Frege. La teoría de Russell es denominada por él, en virtud de las razones que mencionaremos, «atomismo lógico», y alcanza su madurez hacia 1918, año en que pronuncia las conferencias tituladas «La filosofía del atomismo lógico».

Aquí caracteriza su tema como de *gramática filosófica*, y lo justifica así: «Creo que prácticamente toda la metafísica tradicional está llena de errores que se deben a la mala gramática, y que casi todos los problemas y (supuestos) resultados tradicionales de la metafísica se deben a no hacer, en lo que podemos llamar la gramática filosófica, el tipo de distinciones de las que nos hemos ocupado en estas conferencias (*op. cit.*, conferencia VIII). Y unos años después, en un resumen de su teoría, escribiría: «Creo que la influencia del lenguaje en la filosofía ha sido profunda y casi no reconocida. Para que esta influencia no nos extravíe, es necesario que seamos conscientes de ella, y que deliberadamente nos preguntemos en qué medida es legítima. (...) En este aspecto, el lenguaje nos extravía por su vocabulario y por su sintaxis. Debemos estar en guardia sobre ambas cosas para que nuestra lógica no nos conduzca a una falsa metafísica.» («El atomismo lógico», 1924, pp. 330-331 de *Logic and Knowledge*).

En cumplimiento de estas advertencias, Russell desarrollará un tipo de análisis del lenguaje que aspira a poner de manifiesto sus imperfecciones lógicas, contrastándolas con las cualidades de un lenguaje lógicamente perfecto. ¿Cómo es un lenguaje de esta clase? Lo primero que Russell va a decir hace referencia no tanto al lenguaje en sí y a su estructura formal cuanto a la relación entre el lenguaje y la realidad. La primera condición para que un lenguaje sea lógicamente perfecto es una condición semántica: que las palabras de cada proposición correspondan una por una a los componentes del hecho correspondiente. Se exceptúan palabras tales como «o», «no», «si... entonces», las cuales tienen una función diferente, es decir, las cuales carecen de conexión directa con la realidad; son las palabras que expresan modos de componer oraciones, y que pueden traducirse a funtores lógicos, y que, naturalmente, están incluidas en lo que antes hemos llamado «términos sincategoremáticos». Queda así establecido por Russell el principio de isomorfía semántica: «en un lenguaje lógicamente perfecto habrá una sola palabra para cada objeto simple, y todo lo que no

sea simple será expresado por una combinación de palabras...» («La filosofía del atomismo lógico», II, p. 197 de *Logic and Knowledge*). Un lenguaje semejante tiene la ventaja de que muestra a simple vista la estructura lógica de los hechos que afirma o niega. Según Russell, de esta clase pretende ser el lenguaje de los *Principia Mathematica*, con la única diferencia de que este lenguaje posee sintaxis, pero carece de vocabulario: «es el tipo de lenguaje que, si le añadiéramos un vocabulario, sería un lenguaje lógicamente perfecto» (*loc. cit.*, p. 198).

Hay que entender lo que Russell quiere decir. Los *Principia Mathematica*, como todo cálculo lógico, tienen su vocabulario, a saber, el conjunto de signos con los que se componen sus fórmulas en aplicación de sus reglas. Pero lo que Russell quiere dar a entender es que un lenguaje lógicamente perfecto podría ser un lenguaje que, poseyendo un vocabulario, no de signos lógicos, sino de palabras, como las del lenguaje natural, tuviera una sintaxis, unas reglas de estructuración y composición de oraciones, como las de aquel cálculo lógico.

Los lenguajes naturales, las lenguas humanas, no son de esa manera. Y esto, que es una desgracia desde el punto de vista filosófico, es —para Russell— una ventaja a efectos prácticos de comunicación. A diferencia de un lenguaje lógicamente perfecto, el lenguaje ordinario se caracteriza por la ambigüedad de sus palabras, de tal manera que, cuando alguien usa una palabra, no significa por medio de ella la misma cosa que otra persona. Esto, que a primera vista podría parecer un inconveniente, no lo es en realidad, y lo grave sería que todos los hablantes significaran, con sus palabras, las mismas cosas, pues la comunicación resultaría imposible. ¿Por qué? Porque «el significado que uno dé a sus palabras tiene que depender de la naturaleza de los objetos con los que esté familiarizado, y puesto que las diferentes personas están familiarizadas con diferentes objetos, no podrán hablar entre sí a menos que den a sus palabras significados muy diferentes» (*loc. cit.*, p. 195). Así —y el ejemplo es de Russell—, quien ha paseado por Piccadilly, y está por consiguiente familiarizado con esta calle de Londres, da al término «Piccadilly» un significado distinto del que le dará una persona que nunca haya estado allí, por muchas cosas que sepa de ese lugar. Si insistiéramos en un lenguaje carente de ambigüedad, no podríamos hablar de las cosas que conocemos a quienes no las conocen.

Vale la pena detenerse en lo anterior, porque están ahí presentes varios rasgos característicos de la teoría del significado de Russell. En primer lugar, se observará que el significado depende del conocimiento por familiaridad (*knowledge by acquaintance*) o conocimiento directo, que Russell, en otros lugares, ha contrapuesto al conocimiento por descripción (por ejemplo, en el cap. 5 de *Los problemas de la filosofía*). El conocimiento directo excluye la mediación de procesos de inferencia o de conocimiento de verdades. Los datos sensibles constituyen la apariencia de un objeto material, como color, forma, dureza, etc., son ejemplo de algo que se conoce directamente por familiaridad. El conocimiento del objeto como tal es, en cambio, un conocimiento por descripción: supone, no sólo mis

datos sensibles actuales, sino además el recuerdo de otros, junto con el conocimiento de ciertas verdades físicas que están presupuestas por nuestro trato con los objetos materiales. Estos objetos no nos son, pues, conocidos directamente. Lo que conocemos directamente son los datos sensibles que ellos nos producen; los objetos, como tales, son sólo construcciones lógicas que hacemos sobre la base de nuestros datos sensibles, y los conocemos por descripción. El fundamento de nuestro conocimiento está, por consiguiente, en el conocimiento directo, en la familiaridad. Pero ésta no se limita a los datos sensibles (*sense-data*). En el lugar que se acaba de citar, Russell amplía el conocimiento directo a los recuerdos, con lo que la memoria resulta ser, además de los sentidos, una vía para tal conocimiento; e incluye asimismo, en aquél, los estados psicológicos propios, objeto de familiaridad por autoconciencia, aunque duda sobre si incluir también el propio yo. Y no sólo son conocidos directamente estos fenómenos particulares; los conceptos universales son igualmente conocidos por familiaridad, y son un presupuesto para que pueda haber conocimiento por descripción. Del conocimiento directo quedan explícitamente excluidos por Russell los objetos físicos, en cuanto distintos de los datos sensibles que producen, así como los estados psicológicos ajenos. De aquello que conocemos, todo cuanto no es conocido por familiaridad es conocido por descripción, y esto se aplica tanto a los fenómenos particulares como a los conceptos universales. El conocimiento por descripción tiene la importante función de permitirnos sobrepasar los límites de nuestra experiencia personal. Pero el conocimiento por familiaridad es la base de todo conocimiento, y a él es reducible el conocimiento descriptivo, pues «toda proposición que podamos entender debe estar compuesta enteramente de constitutivos con los que estemos familiarizados» (*Los problemas de la filosofía*, cap. 5, final). La razón de esto ya la hemos visto: el significado que demos a nuestras palabras ha de ser algo con lo que estemos familiarizados (*ibidem*).

El peso de la teoría referencialista en las declaraciones de Russell es patente: los significados de las palabras son los objetos de los que tenemos conocimiento directo. Si se trata de un objeto físico, como el designado por el nombre «Piccadilly», el significado de éste, para cada cual, consistirá en los datos sensibles que tenga de ese lugar, así como en sus recuerdos de datos sensibles pasados y en las demás vivencias y sentimientos que dicho lugar le reproduzca. Si consideramos los objetos como integrantes de un hecho, podremos entonces afirmar, con Russell, «que los componentes del hecho que hace a una proposición verdadera o falsa, son los significados de los símbolos que tenemos que entender para poder entender la proposición» («La filosofía del atomismo lógico», II, p. 196 de *Logic and Knowledge*).

Tenemos, pues, que un lenguaje lógicamente perfecto es, para empezar, un lenguaje cuyos términos carecen de ambigüedad, significan siempre lo mismo, a saber, determinadas características de los hechos de las cuales el sujeto posee conocimiento directo. Y esto tiene la inmediata consecuencia de que será un lenguaje privado, en la medida en que el conocimiento

directo es propio y particular de cada cual, y por ello, «todos los nombres que se usen serán privados de un hablante y no podrán entrar en el lenguaje de otro» (*loc. cit.*, p. 198). Esto, desde el punto de vista de su vocabulario. Desde el punto de vista de su sintaxis, la referencia de Russell a los *Principia Mathematica* nos pone en la pista de un rasgo fundamental que no puede faltar en ningún lenguaje perfecto: la extensionalidad, esto es, que todas sus oraciones complejas puedan descomponerse en oraciones simples, de tal modo que la verdad o falsedad de aquellas sea una función de la verdad o falsedad de las últimas, como ocurre en cualquier cálculo lógico estándar. Ello implica que un lenguaje perfecto está constituido por oraciones que pueden ser verdaderas o falsas, esto significa que solamente es candidata a la perfección lógica aquella porción del lenguaje que utilizamos para declarar los hechos, para hablar de lo que acontece, es decir, aquella porción del lenguaje que empleamos en el discurso declarativo o asertórico. Es la misma reducción que —como hemos visto— había efectuado Frege. Por lo que respecta a Russell, podemos decir, siguiendo su terminología, que se trata de un lenguaje compuesto por proposiciones, ya que una proposición es —según Russell— una oración en el modo indicativo, una oración que afirma algo (a diferencia de aquellas oraciones que expresan preguntas, mandatos o deseos); la proposición es el vehículo de la verdad y de la falsedad (*op. cit.*, I, p. 185).

Así pues, las oraciones complejas de nuestro lenguaje perfecto estarán compuestas de oraciones simples unidas por palabras que, como «y», «o», «no», «si... entonces», etc., representan los modos de composición veritativo-funcional. ¿De qué forma serán las oraciones simples? Estas oraciones, que Russell denomina «proposiciones atómicas», describirán el tipo más simple de hecho, lo que, siguiendo la misma analogía, llamará «hechos atómicos». De aquí el nombre de «atomismo lógico» para su teoría: se trata de llegar a los últimos elementos que el análisis lógico del lenguaje pueda encontrar en éste, y puesto que el lenguaje, en lo que es filosóficamente relevante, y de acuerdo con el principio de isomorfía, corresponde estructuralmente a los hechos, por lo mismo llegaremos a los últimos elementos de la realidad. En este sentido, el análisis de Russell va de la lógica a la metafísica a través de la filosofía del lenguaje.

Para Russell, los hechos más simples que se pueda imaginar, los hechos atómicos, son los que consisten en la posesión de una cualidad por una cosa particular, por ejemplo, el hecho descrito por la proposición «Eso es blanco». Aquí tenemos algo, aquello a lo que se refiere el término «eso», y el color que le atribuimos. Una proposición tal es, desde luego, muy diferente de una proposición como «Esa tiza es blanca». En este caso, al considerar algo como tiza, le estamos atribuyendo ciertas propiedades, algunas muy complejas, que sin duda nos llevan más allá de los meros datos sensibles que ahora tenemos del objeto en cuestión. El término «tiza» encierra una complejidad que lo excluye como candidato a constituyente de una proposición atómica. Por eso, y para no prejuzgar nada sobre dicho objeto, nos limitamos a utilizar un pronombre demostrativo como «eso».

Suponemos, de otra parte, que una cualidad como un color es el tipo más simple de cualidades, y por consiguiente no analizable o descomponible ulteriormente. Hay que tener en cuenta que lo relevante aquí es el color en cuanto percibido, y no en cuanto realidad física que puede estudiarse científicamente. Por ello, el que pueda definirse un color en términos de una determinada longitud de onda es irrelevante para el análisis de Russell. Se trata, no de un análisis físico, sino lógico, aunque tomando este último término con una amplitud peculiar, pues, como ya hemos visto, en él es un presupuesto básico el principio de familiaridad. Esto quiere decir que los términos de las proposiciones atómicas poseen significado en cuanto designan objetos de conocimiento directo, y así se explica que los ejemplos que Russell suministra correspondan a hechos atómicos que son, claramente, datos sensibles (*loc. cit.*, pp. 198 y ss.).

El tipo más simple de hecho consiste, pues, en la posesión de una cualidad simple por una entidad particular. Hechos levemente más complejos son los que consisten en relaciones diádicas, como el descrito por la proposición «Eso está junto a aquello». El tipo siguiente será el de relaciones triádicas, como el hecho descrito por «Eso está entre aquello y aquello otro». Y así sucesivamente. Todos estos hechos son atómicos para Russell, y constituyen una jerarquía de complejidad. En todo hecho atómico hay, pues, una propiedad o una relación, más una o varias entidades que son, respectivamente, sujeto de aquélla o ésta. A estas entidades les llama Russell, abreviadamente, «particulares». Se reconocerá aquí la forma de las funciones proposicionales de *cualquier cálculo lógico* estándar: Px , Rxy , $Rxyz$, etc. Un particular es, por tanto, un sujeto de propiedades y de relaciones. Pero ¿en qué consiste? Russell no dirá nada más. La definición de «particular» es puramente lógica y, por consiguiente, nada puede decir más de lo dicho. Al lógico, como tal, no le interesa la cuestión de en qué consiste un particular, ni de si es posible o no encontrar particulares en el mundo. Todo esto son cuestiones empíricas sobre las cuales el lógico, en cuanto tal, carece de competencia: «el lógico nunca da ejemplos, porque es una de las características de una proposición lógica el que no es necesario saber nada acerca del mundo real para poder entenderla» («La filosofía del atomismo lógico», II, p. 199 de *Logic and Knowledge*). Simplemente añadirá que los particulares, como las sustancias en diferentes doctrinas tradicionales, son autosubsistentes y lógicamente independientes entre sí. Que haya uno solo en el mundo o más de uno, y en este caso, cuántos, es ya una cuestión puramente empírica.

Hay que añadir que, en ocasiones, Russell se refiere a los particulares de tal modo que parecería que entiende por particular algo equivalente a un hecho atómico. Así, cuando habla de los particulares como «pequeñas manchas de color o sonidos» (*op. cit.*, I, p. 179), da la impresión de que el particular es un dato sensible, y por lo tanto, un hecho atómico. Sus declaraciones en la segunda conferencia, que acabamos de ver, son tan claras que deben deshacer el equívoco.

Lo que en la proposición corresponde a una propiedad es el predicado. Lo que expresa una relación suele ser un verbo, o a veces, toda una frase. Y lo que corresponde a un particular es el sujeto, y tiene que ser un nombre propio. ¿Por qué? Porque la única manera de hablar de un particular es nombrarlo. Para describirlo, ya mencionamos sus propiedades y sus relaciones utilizando los términos correspondientes; ahora bien, para referirnos a él como sujeto de aquellas, lo único que podemos hacer es nombrarlo. Y puesto que las palabras obtienen su significado de los objetos con los que estamos familiarizados, quiérese decir que tan sólo podemos nombrar lo que es objeto de conocimiento directo y mientras lo es. La primera consecuencia de esta extraña doctrina es que los nombres propios de particulares, tal y como aparecen en una proposición atómica, serán muy distintos de lo que, en el discurso ordinario, llamamos «nombres propios». Palabras como «Sócrates», «Venus», «Madrid», las usamos para referirnos a sus correspondientes objetos aun cuando éstos no estén presentes; de hecho, parece que su utilidad estriba precisamente en ello, pues quien estuviera ante Sócrates o quien se hallara en Madrid probablemente no necesitaría recurrir a esos nombres. Ahora bien, de acuerdo con la doctrina de Russell, no tenemos conocimiento directo de Sócrates, y por consiguiente, no podemos nombrarlo. Por lo mismo, quien nunca haya estado en Madrid, tampoco podrá dar significado a este término, y tampoco podrá dárselo al término «Venus» quien no haya contemplado este planeta. Ello muestra que tales palabras no son en realidad nombres propios, esto es, que no son nombres propios en sentido lógico. ¿Qué son, entonces? Según Russell, se trata de descripciones encubiertas y abreviadas. «Sócrates» es una abreviatura para cualquier descripción correcta que podamos dar de su correspondiente objeto, por ejemplo: «El filósofo griego que fue condenado a beber la cicuta», o «El maestro de Platón», o cualquier otra. Como «Madrid» abreviará, entre otras muchas, la descripción «La capital de España», o «Venus» equivaldrá, entre otras, a «El lucero matutino». En la medida en que estas descripciones se refieren a sus objetos describiendo ciertas propiedades suyas, resulta patente que esos objetos no pueden ser particulares, pues no son simples. Tenemos, pues, que ni los nombres propios del lenguaje ordinario son nombres propios en sentido lógico ni aquello a lo que se refieren son particulares. Por ello puede afirmar Russell: «Hablando estrictamente, sólo los particulares pueden ser nombrados.» («La filosofía del atomismo lógico», VII, p. 267 de *Logic and Knowledge*).

Así pues, Mill había dicho que los nombres propios en sentido ordinario denotan, pero carecen de connotación; Frege defendió que, no sólo pueden tener referencia, sino que además han de tener sentido, el cual puede quedar explícito por medio de alguna descripción como en los ejemplos anteriores; y Russell viene a añadir que, precisamente por eso, tales nombres no son, lógicamente, nombres propios, pues si es posible sustituirlos por alguna descripción, entonces no se limitan a nombrar.

¿En qué consiste un nombre propio en sentido lógico? Según Russell, las únicas palabras que usamos de esta manera son palabras como «esto», «eso» o «aquello» (*this, that*): «Se puede usar 'esto' como nombre de un particular del que se tiene conocimiento directo en el momento» (*op. cit.*, II, p. 201). Así, si decimos «Esto es blanco», llamando «esto» a lo que vemos, empleamos el demostrativo como nombre propio, en sentido lógico, de un supuesto particular que tiene como propiedad la blancura. Pues, en efecto, los demostrativos no nos dicen nada sobre los objetos a los que, por medio de ellos, nos referimos; se limitan a señalarlos, a denotarlos, y eso prueba que son verdaderos nombres propios y que los objetos que denotan son simples, particulares. Esto implica una curiosa propiedad que Russell se ocupa de señalar, y que, aunque paradójica, es coherente con lo que ya hemos visto. De una parte, que el significado de los nombres lógicamente propios estará cambiando todo el tiempo según cambien nuestras percepciones del mundo, nuestros datos sensibles. Y de otra, que su significado será diferente para el hablante y para el oyente, en cuanto que los datos sensibles que ambos tengan del mismo objeto serán distintos. Con lo que volvemos a comprobar el carácter privado de un lenguaje lógicamente perfecto, puesto que sus nombres propios serán inteligibles únicamente para cada cual en función de sus experiencias propias: «para comprender un nombre hay que estar familiarizado con el particular que nombra, y hay que saber que se trata del nombre de dicho particular» (*op. cit.*, III, p. 205 de *Logic and Knowledge*).

Frente al monismo hegeliano con el que Russell venía polemizando desde años antes, la ontología exigida por su análisis consiste, para empezar, en un pluralismo de los hechos simples o atómicos, que se resuelve en un pluralismo de objetos simples o particulares, independientes lógicamente entre sí y subsistentes por sí mismos, con un tipo de subsistencia que recuerda a la de la sustancia, según se ocupa de señalar el propio Russell (II, p. 202). Por lo que respecta a los objetos de la vida cotidiana, éstos son todos complejos, desde las sillas a las personas, y por esto no se les puede dar un nombre propio lógico. Ya tenemos, por consiguiente, los elementos más simples a los que llega el análisis de Russell: son los particulares, sus propiedades y sus relaciones. Y están representados en la oración de esta manera: los particulares, por los nombres lógicamente propios (términos deícticos, como los demostrativos), las propiedades y relaciones por diferentes clases de adjetivos, verbos y adverbios. Como todo elemento de la oración debe corresponder a un elemento del hecho, hay que concluir que en los ejemplos que, tomados de Russell, hemos visto, sobra algo, a saber, la cópula «es», puesto que a ella nada corresponde en los hechos. Los ejemplos de proposiciones atómicas serán, así, menos idiomáticos de lo que eran los anteriores; con rigor, tales proposiciones serán de la forma «Esto blanco», «Eso junto a aquello», etc. Que así tiene que ser lo prueba, claramente, el que nada hay en las funciones proposicionales de un cálculo lógico que represente al «es»: Px , Rxy , etc., solamente contienen términos de individuos (x , y) y términos de predicados (P , R).

Con lo anterior hemos alcanzado unos átomos lógicos, las proposiciones atómicas, a las cuales corresponden unos hechos simples, que cabe calificar, asimismo, como atómicos. ¿Pero pueden reducirse a aquellas todas las demás proposiciones de un lenguaje perfecto?

6.5 Hechos y proposiciones

Las proposiciones atómicas se combinan entre sí por los medios de composición veritativo-funcional que establecen los *Principia Mathematica* y que se recogen en cualquier libro de lógica; formas de composición que, en el lenguaje ordinario, están representadas, con cierta aproximación, por palabras como «y», «o», «no», «si... entonces», etc. A las proposiciones complejas así formadas las llama Russell, prosiguiendo la misma analogía, «proposiciones moleculares». Es, por tanto, característico de un lenguaje perfecto cumplir con el principio de extensionalidad, a saber: que todas sus proposiciones complejas o moleculares puedan descomponerse en otras simples o atómicas de tal manera que la verdad o falsedad de las primeras sea función de la verdad o falsedad de las últimas. De aquí que las proposiciones moleculares, puesto que son meros compuestos de proposiciones atómicas, carezcan de correlato propio en la realidad. No hay, no tiene por qué haber, hechos moleculares. Ya que toda proposición molecular se descompone en proposiciones atómicas, bastan los hechos atómicos para conectar a la primera con el mundo. Un hecho es, simplemente, aquello que hace verdadera o falsa a una proposición (*op. cit.*, I, p. 182 de *Logic and Knowledge*). Pero una proposición molecular no es verdadera o falsa por sí misma, esto es, en virtud de su relación con el mundo, sino en razón de que sean verdaderas o falsas las proposiciones atómicas que la componen; por consiguiente, la única verdad que depende de los hechos es la de estas últimas, y para declarar verdaderas o falsas a las proposiciones atómicas nos bastan los hechos atómicos. Nótese, además, que si postuláramos la existencia de hechos moleculares, nos veríamos forzados a admitir que hubiera en la realidad, como parte de tales hechos, elementos que correspondieran a los modos de combinación, esto es, a la conjunción, a la disyunción, al condicional, etc. Si, tomando dos ejemplos muy sencillos de proposiciones atómicas, afirmamos «Eso (es) blanco y aquello (es) negro», nuestra afirmación será verdadera, según la interpretación que hace de la conjunción cualquier cálculo lógico estándar, solamente cuando ambas proposiciones simples lo sean. Y para esto basta con sus respectivos hechos atómicos: que lo designado por «eso» sea, efectivamente, blanco, y que lo que llamamos «aquello» sea negro. Simplemente con esto, nuestra afirmación será verdadera. No necesitamos para nada postular un hecho complejo en el que, además de algo blanco y de algo negro haya también algún extraño elemento que corresponda al functor «y». Si todas las proposiciones complejas fueran moleculares, y por ello reducibles a proposiciones atómicas, éste sería el fin de la cuestión. En última

instancia no tendríamos más que las proposiciones atómicas en nuestro lenguaje perfecto, y los hechos atómicos en el mundo. El problema, para Russell, es que encuentra proposiciones complejas cuya reducción a proposiciones simples le resulta problemática.

El primer caso es el de las proposiciones negativas que son verdaderas (*op. cit.*, III). El ejemplo sugerido por Russell es:

(1) Sócrates no está vivo

Podríamos decir que hay aquí una proposición simple (en rigor no lo es, desde el punto de vista de la teoría de Russell, pero podemos suponerlo a efectos del ejemplo), que será:

(2) Sócrates está vivo

y a la que se añade una complejidad lógica que es la negación. Puesto que (1) es verdadera, (2) será falsa. Entonces, la cuestión es: ¿cuál es el hecho que hace falsa a (2)? Si no podemos señalar ningún hecho positivo responsable de la falsedad de (2), no tendremos más remedio que aceptar que el hecho buscado es el mismo que hace verdadera a (1). Y por lo mismo habremos admitido que en el mundo hay, además de los hechos atómicos que ya conocemos, hechos negativos.

Russell no veía la manera de evitar esta consecuencia. Ciertamente no le hacía feliz, y más que defender decididamente que hay hechos negativos se limitaba a sugerir que puede haberlos (*loc. cit.*, p. 212 de *Logic and Knowledge*). En todo caso, le repugnaba menos aceptar hechos negativos que asentir a una explicación alternativa que, por ese tiempo, se le había oftecido. Según tal explicación, afirmar una proposición negativa, *no-p*, equivaldría a afirmar que hay una proposición *q* verdadera e incompatible con *p*. Las razones de Russell para rechazar esta explicación no me parecen claras. Opone, en todo caso, que implica tomar la incompatibilidad como un hecho fundamental, y esto le parece rechazable por cuanto obligaría a considerar las proposiciones como hechos, puesto que la incompatibilidad se da entre ellas. Aparte de que el hecho de la incompatibilidad no parece más fácil de aceptar que los hechos negativos. Sin embargo, estos argumentos de Russell no parecen alcanzar al fondo de la cuestión. De entrada, no se ve que haya que considerar la incompatibilidad como un hecho. La incompatibilidad es una relación entre proposiciones y la único que la explicación alternativa pretende expresar es que *no-p* es verdadera, no porque haya un hecho negativo al cual corresponda, sino porque hay un hecho positivo que corresponde a una cierta proposición verdadera *q* la cual es verdadera, no porque se dé el hecho negativo de que Sócrates no está vivo, sino porque hay un hecho positivo que hace verdadera a una proposición como

(3) Sócrates ha muerto

la cual es incompatible con (2). Ese hecho es, naturalmente, la muerte de Sócrates. Para que esta explicación funcione basta con admitir que (1) equivale a (3) en virtud de lo que significan las expresiones «no está vivo» y «ha muerto». O dicho de otro modo: basta con aceptar que la proposición «Sócrates no está vivo si y sólo si ha muerto» es analíticamente verdadera. Por lo mismo, si, señalando a algo que es blanco, digo «Eso no (es) negro», la verdad de esta afirmación no me compromete a pensar que, además del hecho atómico de que eso es blanco, ha de haber una pluralidad de hechos negativos consistentes en que eso mismo no es negro, ni rojo, ni azul, ni... etc. La blancura de ese particular hace verdadera tanto a la proposición «Eso (es) blanco» como a las proposiciones negativas «Eso no (es) negro», «Eso no (es) rojo», etc. Para lo cual basta admitir como analíticamente verdadera la proposición «Ningún objeto puede ser de dos colores distintos al mismo tiempo». Tan sólo podríamos sentirnos forzados a reconocer hechos negativos si hubiera algún tipo de proposiciones negativas que no pudiera reducirse a proposiciones positivas de la manera indicada. Pero los ejemplos que hemos considerado, ciertamente no parecen presentar problemas.

El segundo tipo de proposiciones complejas que, para Russell, no pueden reducirse a proposiciones atómicas son aquellas proposiciones que expresan actitudes proposicionales, esto es, que expresan ciertos fenómenos mentales que implican una proposición. Por ejemplo, las proposiciones que expresan creencias, deseos, comprensión, etc. Así, cuando decimos «Creo que hoy es martes», cuya forma sería: «Creo que p », donde p representa la proposición «Hoy es martes». Como si decimos, «Deseo quedarme solo», cuya forma podemos dar como «Deseo que p », donde p sustituirá a la proposición «Me quedo solo». E igualmente, cuando afirmamos «Comprendo el teorema de Pitágoras», cuya forma será: «Comprendo p », representando p la proposición que expresa el teorema de Pitágoras. A los verbos de esta clase, que tienen como complemento una proposición, les llama Russell «verbos proposicionales» (*op. cit.*, IV, secc. 4).

Es claro que estas proposiciones complejas no pueden descomponerse extensionalmente. Podemos distinguir, por supuesto, dos partes: la parte que expresa la actitud en cuestión, «creo que», «deseo que», «comprendo»; y la parte que expresa el contenido de la actitud, esto es, la proposición que expresa aquello sobre lo que recae la creencia, el deseo o la comprensión. Pero es patente que la verdad o falsedad de la proposición compleja no es función de sus partes. La proposición «Creo que hoy es martes» es verdadera si efectivamente eso es lo que creo, tanto si hoy es martes como si no lo es. Dicho de otra forma: la verdad de la proposición «Hoy es martes» es irrelevante para la verdad de la proposición que expresa mi creencia en ello. Mi creencia no es menos creencia porque yo esté equivocado. Y lo propio puede afirmarse para los otros verbos de este tipo. Estas proposiciones complejas no son, pues, extensionales, sino que constituyen funciones intensionales. Por consiguiente, Russell está esbozando el análisis de aquellas expresiones que se refieren a procesos

mentales, o supuestamente mentales. Su posición en las conferencias que estamos comentando es clara: tales proposiciones corresponden a una clase peculiar de hechos, dentro de la cual podemos distinguir hechos de creencia, hechos de deseo, hechos de comprensión, etc. Aunque todos los ejemplos que da Russell lo son de hechos mentales, o psicológicos, él mismo deja abierta la posibilidad de que pueda haber verbos proposicionales que correspondan a hechos de carácter no psicológico. Ciertamente que los hechos referidos son de muy distinta condición que los hechos atómicos que ya conocemos. Una posibilidad de reducir aquellos a éstos sería analizar los verbos proposicionales en términos conductistas, de tal manera que la creencia, el deseo, la comprensión, etc., quedarán descompuestos en procesos de comportamiento. Aunque años después, en *The Analysis of Mind*, Russell aceptó este análisis tal reducción es rechazada, aunque sin entrar en detalles (*loc. cit.*, secc. 1). Curiosamente, la mejor razón que da estriba en que no puede explicarse cómo es posible utilizar nombres propios lógicos si se prescinde de la conciencia; pues, en efecto, la referencia de «eso» o «aquello», en los ejemplos que ya hemos estudiado, deriva pura y exclusivamente de la intención referencial del sujeto, y no puede reducirse a ningún aspecto de su comportamiento respecto a las cosas. Dicho en otros términos: admitido que la referencia de los nombres propios es asunto privado, no puede prescindirse de la esfera en la que se da esa privacidad, a saber, la conciencia. Hubiera sido incoherente por parte de Russell afirmar que el vocabulario de un lenguaje lógicamente perfecto es privado, y a continuación reducir los procesos mentales a procesos de conducta.

Hemos de considerar, por último, como proposiciones no analizables, las proposiciones cuantificadas, tanto las generales como las particulares. Una proposición general, o universal, es una proposición de la forma de «Todos los hombres son mortales», que afirma una propiedad de todos los miembros de una clase. Podría mantenerse, y Wittgenstein lo hará posteriormente, que una proposición universal puede descomponerse en una conjunción de proposiciones simples que afirmen la propiedad en cuestión de cada uno de los miembros de la clase. El ejemplo mencionado equivaldría, según esto, a una enumeración de todos los seres humanos, afirmando de ellos, sucesivamente, la mortalidad: «Fulano es mortal y Mengano es mortal y Zutano es mortal...». Russell rechazará este análisis invocando como razón que una enumeración nunca nos da el carácter de generalidad. Una vez que hemos acabado la enumeración, ¿cómo sabremos que los casos referidos son todos? Tendríamos que añadir, al efecto, la cláusula: «y éstos son todos los casos». Supongamos que se trata de enumerar todos los hechos atómicos del mundo: «es perfectamente claro, creo, que, cuando se han enumerado todos los hechos atómicos del mundo, es un hecho ulterior acerca del mundo que esos son todos los hechos atómicos que hay en el mundo, y esto es un hecho objetivo del mundo tanto como lo son cualquiera de los hechos atómicos» (*op. cit.*, V, p. 236 de *Logic and Knowledge*). Hay que concluir, por consiguiente, que en el mundo hay hechos generales.

En análogo caso están las proposiciones particulares o existenciales, es decir, aquellas que afirman que hay entidades que tienen tal o cual propiedad, por ejemplo, la afirmación «Hay hombres». Puesto que estas proposiciones tampoco son funciones veritativas de otras más simples, la consecuencia es que tiene que haber un tipo de hechos que las haga verdaderas, a saber, lo que Russell llama «hechos de existencia» (*loc. cit.*). A decir verdad, la ontología de Russell parece aquí un tanto redundante, pues teniendo en cuenta la interdefinibilidad de los cuantificadores por medio de la negación, podríamos prescindir o bien de las proposiciones generales o bien de las existenciales. Así, en la primera alternativa, la proposición «Todos los hombres son mortales» equivaldría a «No hay hombres que no sean mortales», con lo que podríamos prescindir del supuesto hecho general correspondiente a la primera de estas dos proposiciones, y aceptar en nuestra ontología meramente hechos existenciales negativos como el que correspondería a la segunda de ellas (si bien sería un hecho un tanto complejo). De otra parte, teniendo en cuenta que Russell ha mantenido —como vimos— que los objetos de trato cotidiano son ficciones lógicas construidas por nosotros, los ejemplos que da ahora resultan raros. Pues si los hombres son, en definitiva, construcciones lógicas sobre la base de datos sensibles, no se ve muy bien cuál es la justificación para aceptar como un hecho de existencia el que haya hombres. Parece que todo lo más que habría que reconocer como hecho de existencia es la existencia de los particulares.

La argumentación de Russell relaciona el lenguaje y la concepción de la realidad de un modo curioso, muy típico del atomismo lógico. La idea básica es doble. De un lado, que toda proposición es empíricamente verdadera o falsa en virtud de un hecho que la hace así. De otro, que a toda proposición que no pueda descomponerse en una función veritativa de otras más simples, corresponde un tipo peculiar de hecho. De aquí que Russell tuviera que aceptar un mundo compuesto, no sólo por hechos atómicos, sino también por hechos negativos, por hechos generales, por hechos de existencia, y por diferentes clases de hechos de actitudes proposicionales (creencias, deseos, etc.). Aquí, naturalmente, entraban sus particulares opiniones sobre qué proposiciones son analizables extensionalmente. Como se acaba de mencionar, Wittgenstein, y en una primera época, también Ramsey, fueron partidarios de reducir las proposiciones generales a una conjunción de casos, eliminando así la necesidad de admitir hechos generales, y por lo mismo, hechos de existencia (ya que las proposiciones existenciales son interdefinibles con las generales; véase la discusión de este punto en Urmson, *El análisis filosófico*, cap. 5, secc. B). Como puede apreciarse, su teoría de la lógica y su análisis del lenguaje condicionan del todo, para Russell, su concepción del mundo. Y si no estuviera bastante claro, lo vamos a comprobar de nuevo a propósito de un tema íntimamente relacionado con su teoría del significado.

6.6 Denotación y descripciones

Acabamos de ver que los hechos de existencia consisten en la existencia de determinadas clases de objetos (aunque en rigor, tan sólo la existencia de los particulares estaría justificada). Es la existencia que implica las proposiciones cuantificadas por medio del cuantificador particular. Pero Russell no considera como hecho de existencia la existencia de un objeto singular. ¿Por qué?

Para hablar de una entidad determinada y singular necesitamos una expresión que se refiera a ella, y que la represente en el contexto lingüístico de la proposición. Supongamos que deseamos decir algo acerca de la ciudad en la que se está escribiendo el presente libro; por ejemplo, que cada vez tiene más tráfico. Al predicado «cada vez tiene más tráfico» hemos de anteponer una expresión que se refiera al objeto (ciertamente complejo) que es esta ciudad, y que se refiera a ella sin posibilidad de confusión con otra. No puedo limitarme a escribir «Esta ciudad cada vez tiene más tráfico», pues si el lector no conoce en qué ciudad me encuentro, se quedará sin saber de cuál ciudad hago esa afirmación, y nunca podrá llegar a comprobar, con la única información que mi proposición le suministra, si mi afirmación es verdadera o falsa. Mi afirmación tan sólo sería suficientemente informativa para quien, o bien se encontrara conmigo cuando hago la afirmación, o bien supiera, por otros medios, en qué ciudad me hallo. A quien no se encuentre en ninguno de estos casos, solamente puedo hacerle explícito de qué ciudad estoy hablando refiriéndome a ella de forma unívoca. Y esto tengo básicamente dos maneras de hacerlo: o bien empleando un nombre propio (en el sentido ordinario del término), o bien utilizando una descripción unívoca, esto es, definida. Lo primero es lo que hago cuando afirmo: «Madrid cada vez tiene más tráfico». Lo segundo consiste en decir algo como: «La capital de España cada vez tiene más tráfico», o «La ciudad del oso y el madroño cada vez tiene más tráfico», etc.

Sin embargo, y como ya hemos visto, para Russell, ambas alternativas no son lógicamente distintas. Un nombre propio ordinario no es —en su opinión— otra cosa que una descripción definida abreviada, y por consiguiente tanto da emplear el nombre «Madrid» como recurrir a la descripción «la ciudad del oso y el madroño», o a cualquier otra. Lo primero no hace sino abreviar lo segundo. Por las razones que ya conocemos, «Madrid», como «Cervantes», «Sócrates», «Venus», «Rocinante», etc., no son propiamente nombres propios. Nombrar sólo se puede nombrar lo que se conoce directamente, y esos supuestos nombres se usan, muchas veces, precisamente para mencionar entidades de las que no se posee conocimiento directo alguno.

Pero volvamos al tema de la existencia. Supongamos que pretendemos afirmar que la ciudad de la que venimos hablando en los ejemplos anteriores existe. Imaginemos que, aunque no lo es, «Madrid» fuera realmente un nombre propio en sentido lógico. Esto significa que, por serlo, nombra-

ría un objeto, en este caso, dicha ciudad. Si ahora decimos «Existe Madrid», esta proposición resulta ser una tautología, pues es obvio que si Madrid no existiera no podría ser nombrada, y viceversa, que si es nombrada es porque existe. Por consiguiente, es inútil afirmar la existencia de un objeto singular empleando para referirse a él un nombre lógicamente propio, pues el uso de tal nombre ya implica la existencia del objeto: «un nombre ha de nombrar algo o no es un nombre» (*op. cit.*, VI, p. 243 de *Logic and Knowledge*). Por lo tanto, si no existiera el objeto, no podríamos tener un nombre propio para él. Por eso concluye Russell que, si podemos discutir la proposición «Dios existe», esto prueba que el término «Dios» no es un nombre lógicamente propio, sino una descripción encubierta (*loc. cit.*, p. 250).

La manera de hacer una afirmación de existencia realmente informativa, una afirmación que pueda ser, en principio, verdadera o falsa, es utilizando una descripción definida. Sustituyendo los pseudonombres de los ejemplos anteriores por descripciones explícitas, tendremos:

- (1) Existe la ciudad del oso y el madroño
- (2) El ser infinitamente perfecto existe

¿En qué consisten estas afirmaciones? Aparentemente se afirma un predicado, la existencia, de un sujeto. Parece que esas proposiciones tendrían la forma clásica «S es P». Esto, por cierto, es lo que justificaba el argumento ontológico en favor de la existencia de Dios, pues si la existencia es una perfección, entonces sería contradictorio negársela al ser infinitamente perfecto, y en consecuencia, meramente tautológico el atribuírsela. La idea, ya defendida por Kant, de que la existencia no puede tomarse como un predicado más entre otros, reaparece ahora en Russell con una formulación lógica. Decir que existe la ciudad del oso y el madroño es afirmar que hay una entidad, y sólo una, de la que puede afirmarse con verdad que es una ciudad y que en su escudo tiene un oso y un madroño. Enunciar que existe el ser infinitamente perfecto, es decir que hay una entidad, y sólo una, que es infinitamente perfecta. Lo característico del análisis de Russell es que, en la nueva formulación, desaparecen las descripciones definidas. Ya no se dice nada de «la ciudad del oso y el madroño» ni de «el ser infinitamente perfecto». La verdadera estructura lógica de (1) y (2), lo que Russell llamará su «forma lógica», es, respectivamente:

- (3) Hay una entidad, y sólo una, que es una ciudad y en su escudo tiene un oso y un madroño
- (4) Hay una entidad, y sólo una, que es infinitamente perfecta

En símbolos lógicos (empleando una notación actual y no la de Russell), la forma de las proposiciones anteriores será:

- (5) $\forall x (Fx \wedge \wedge y (Fy \leftrightarrow y=x))$

Dicho en palabras: Hay al menos una entidad x que tiene la propiedad F , y para toda entidad y , tiene la propiedad F si y sólo si es idéntica a x . Respectivamente, F representa la propiedad de ser infinitamente perfecto y la propiedad de ser una ciudad que tiene en su escudo un oso y un madroño. Puesto que el cuantificador particular sólo nos asegura que hay *al menos* una entidad con la propiedad F , añadir que toda entidad que tenga esa propiedad es idéntica a x es necesario para conservar la univocidad de la descripción definida, esto es, la idea de que sólo una entidad tiene tal propiedad. ¿Dónde está ahora la afirmación de existencia? En el cuantificador particular, que por eso se ha llamado también «cuantificador existencial». Afirmamos la existencia de una entidad simplemente al atribuirle una propiedad cualquiera. Dicho de otro modo: la existencia de una entidad es un presupuesto para la atribución de propiedades, y no una propiedad más. Por eso tiene sentido, aunque sea falso, decir que no existe la ciudad del oso y el madroño, pues esto equivale simplemente a afirmar que no hay una única entidad que sea una ciudad y tenga en su escudo un oso y un madroño (lo cual podría ocurrir o bien porque no hubiera ninguna o bien porque hubiera más de una).

Este análisis le permite a Russell una fácil crítica de la teoría de Meinong, con la que había polemizado por esta época («Sobre la denotación», 1905, p. 45 de *Logic and Knowledge*). Según la teoría de los objetos de Meinong, toda descripción denotaría un objeto, aun las descripciones de objetos inexistentes en la realidad, como «el actual Presidente de la República Española» o «la montaña de oro», y las descripciones de objetos imposibles, como «el cuadrado redondo». La idea es que, si no hay algún tipo de objetos que corresponda a estas expresiones, ¿cómo podemos hacer afirmaciones en las que éstas oficien de sujeto? ¿Cómo podemos decir que la montaña de oro no existe o que el cuadrado redondo es imposible si no hay algo sobre lo que versen estas afirmaciones? Por cierto que esta doctrina aparece aceptada por Ortega en una conocida obra de 1929:

Entiendo por Universo formalmente «todo cuanto hay» (...). Este «hay», que no es un grito de dolor, es el círculo más amplio de objetos que cabe trazar, hasta el punto que incluye cosas, es decir, que hay cosas de las cuales es forzoso decir que las hay pero que no existen. Así, por ejemplo, el cuadrado redondo, el cuchillo sin hoja ni cacha o todos esos seres maravillosos de que nos habla el poeta Mallarmé —como la hora sublime que es, según él, «la hora ausente del cuadrante», o la mujer mejor, que es «la mujer ninguna»—. Del cuadrado redondo sólo podemos decir que no existe, y no por casualidad, sino que su existencia es imposible; pero para poder dictar sobre el pobre cuadrado redondo tan cruel sentencia es evidente que tiene previamente que ser habido por nosotros, es menester que en algún sentido lo baya. (*¿Qué es filosofía?*, lección IV, p. 86.)

El recurso al *haber* soluciona el problema de la condición ontológica de los objetos inexistentes. No existen cuadrados redondos, y además, es imposible que existan. Pero hay cuadrados redondos, porque si no los

hubiera ¿cómo se explicaría que hablemos de ellos? La solución de Ortega es meramente verbal, y ejemplifica de manera adecuada un comportamiento relativamente frecuente en la historia del pensamiento filosófico. No sirve el término «existir», es cuestionable el término «ser» (por razones propias de la argumentación de Ortega, que ahora no son del caso), pues recurramos a otra palabra que no esté contaminada de teoría. La cuestión es: ¿en qué se distingue «haber» de «existir»? ¿Cuál es la definición de «haber»? No lo sabemos. Se trata simplemente, al parecer, de que estamos dispuestos a decir que *hay* todo aquello que tendría que ser denotado por cualquier descripción que podamos formar sintácticamente en nuestra lengua. O lo que es lo mismo: que hemos tomado la antioccamiana decisión ontológica de aceptar como objeto todo aquello que, *aparentemente*, correspondería a cualquier descripción sintácticamente posible en nuestra lengua. (Nótese que no he dicho «descripción gramaticalmente posible»; la razón es que esta expresión podría incluir la corrección semántica, y yo no creo que una descripción como «el cuadrado redondo» sea semánticamente correcta en castellano.) Podemos llamar a este proceder «la falacia de la referencia». Sus consecuencias para la ontología me parecen del todo indeseables.

El análisis de Russell, de un sobrio malthusianismo lógico, evita esta superpoblación del universo de nuestro discurso. La razón es que las descripciones han desaparecido. En las proposiciones a las que nos lleva el análisis no hay descripciones, sino afirmaciones de existencia. Y así, decir que la existencia del cuadrado redondo es imposible, no es afirmar nada sobre el cuadrado redondo; es afirmar que es imposible que haya una entidad que posea la propiedad de ser al tiempo cuadrada y redonda. Este análisis le permite a Russell, por cierto, declarar falsas todas las proposiciones que tratan sobre objetos inexistentes. Pues supóngase que decimos:

- (6) El segundo satélite natural de la Tierra se encuentra a 800.000 kilómetros de ésta

Teniendo en cuenta que no se conoce más que un satélite natural de la Tierra, la Luna, la afirmación anterior no trata de nada existente y, por consiguiente, parece que habremos de declarar falsa. Pero de acuerdo con cualquier cálculo lógico bivalente como el de los *Principia Mathematica*, la negación de una proposición falsa es una proposición verdadera, y en consecuencia, verdadera habrá de ser la afirmación:

- (7) El segundo satélite natural de la Tierra *no* se encuentra a 800.000 kilómetros de ésta

Mas ¿qué sentido tiene decir que esta proposición es verdadera (se sobreentiende: actualmente)? ¿Cómo puede hacerse afirmaciones verdaderas (y no tautológicas) sobre objetos que no existen? Decir de un cuerpo celeste que no se halla a cierta distancia de la Tierra es —parece— dar in-

formación, aun cuando negativa, sobre él. ¿Cómo puede darse información sobre algo que no existe?

El análisis de Russell evita, ciertamente, estas aporías. También las eludió en su momento, Frege, aunque de modo diverso. Para Frege, como ya vimos, (6) y (7), no son ni verdaderas ni falsas, pues como su sujeto carece de referencia, también la oración está falta de ella (recuérdese que, según Frege, la referencia de una oración es su valor veritativo). La solución de Russell, sin embargo, permite operar lógicamente con esas oraciones puesto que permite asignarles un valor veritativo (en ambos casos, el valor «falso»). En efecto, la forma lógica de (6) será:

(8) Hay una entidad, y sólo una, que es satélite natural de la Tierra y que fue descubierto en segundo lugar, y tal entidad está a 800.000 kilómetros de la Tierra

La forma de (7) será idéntica, excepto que al final se negará que tal objeto está a dicha distancia. Si representamos por F la propiedad compleja de ser segundo satélite natural de la Tierra, y por G la propiedad de hallarse de ella a 800.000 kilómetros, la expresión lógica de (8) en la notación usada anteriormente será así:

$$(9) \quad \forall x (Fx \wedge \exists y (Fy \leftrightarrow y=x) \wedge Gx)$$

Por lo que respecta a (7), su formalización será idéntica, excepto que la última función proposicional aparecerá negada, así:

$$(10) \quad \forall x (Fx \wedge \exists y (Fy \leftrightarrow y=x) \wedge \neg Gx)$$

(Puesto que «ser satélite natural de» y «estar a tal distancia de» son relaciones, las fórmulas anteriores deberían ser más complejas para ser del todo claras, a saber: deberían tener una constante que representara a la Tierra, pero para nuestro propósito podemos prescindir de esa complejidad.)

Ahora se comprenderá por qué tanto (6) como (7) son falsas. Con independencia de que afirmemos o neguemos la propiedad G , la proposición conjunta es falsa, pues es siempre falso su primer miembro, a saber: es siempre falso que hay un x que tiene la propiedad F , esto es: es siempre falso que hay una entidad que es segundo satélite natural de la Tierra.

Russell considera las descripciones definidas como símbolos incompletos, junto con los nombres de clases y con los nombres propios en sentido ordinario. Esto significa que se trata de símbolos, o expresiones, que, aunque parecen ser parte constitutiva de las proposiciones, no lo son en realidad, pues una vez que éstas han sido analizadas, tales expresiones desaparecen. Los nombres comunes o de clases, porque, al analizar las proposiciones en las que intervienen, quedarán sustituidos por nombres de particulares y de propiedades o relaciones simples, a la manera que ya hemos visto en las secciones precedentes. Esto es coherente con la doctrina russelliana de que

todos los objetos de trato cotidiano son el resultado de construcciones lógicas efectuadas sobre la base de datos sensibles. En cuanto a los nombres propios ordinarios, porque, como hemos comprobado, no pueden tomarse más que como descripciones encubiertas o implícitas. Pero también las descripciones son símbolos incompletos, ya que, al analizar debidamente las proposiciones de las que forman parte, desaparecen, quedando sustituidas por funciones proposicionales cuantificadas del modo que acabamos de ver.

De las descripciones definidas da Russell una caracterización que, en apariencia, es meramente sintáctica. Así, al comienzo de «Sobre la denotación» (p. 41 de *Logic and Knowledge*), las presentará como una clase de expresiones denotativas. Pero éstas son únicamente definidas por una especie de enumeración un tanto informal; son expresiones denotativas frases tales como «un hombre», «algunos hombres», «todos los hombres», «cualquier hombre», «el actual rey de Francia», «la revolución de la Tierra alrededor del Sol», etc. Cuando a continuación de esta enumeración Russell añade: «por lo tanto, una expresión es denotativa solamente en virtud de su *forma*» (*loc cit.*, subrayado en el original), uno no puede dejar de preguntarse en qué consiste tal forma. La realidad es que Russell procede a distinguir tres tipos de expresiones denotativas, pero su clasificación es más bien semántica que sintáctica. Las tres clases son: primero, la de las expresiones denotativas que no denotan nada, como «El actual rey de Francia»; segundo, la de las que denotan un objeto definido, como «El actual rey de España»; y tercero, la clase de las expresiones denotativas que denotan ambiguamente, por ejemplo, «un hombre». Parece patente que el criterio de esta distinción es semántico: se trata de que la expresión en cuestión tenga o no referencia, y de qué manera la tenga. A decir verdad, y a falta de un claro criterio sintáctico, el propio concepto de «expresión denotativa» parece claramente semántico: son aquellas expresiones que sirven para referirse a los objetos, sean del tipo que fueren.

De hecho, la contraposición que a Russell le interesaba fundamentalmente establecer es la que hay entre las descripciones definidas y las descripciones ambiguas. Siguiendo, al parecer, con la pretensión de recurrir a un criterio sintáctico, dirá, en diferentes lugares, que una descripción definida es una descripción de la forma «el tal y cual», mientras que una descripción ambigua es una expresión cuya forma es «un tal y cual» («La filosofía del atomismo lógico», VI; *Introduction to Mathematical Philosophy*, cap. XVI). Las limitaciones de esta distinción, así formulada, son obvias: ¿es que una lengua en la que no haya artículos carece de descripciones? Para empezar, es importante tener en cuenta que lo descrito por una descripción puede ser no sólo un objeto individual, sino también un predicado o una relación. De este último caso ya hemos visto un ejemplo hace un momento: «la revolución de la Tierra alrededor del Sol». A mi juicio, es claro que lo que distingue entre sí ambas clases de descripciones es un rasgo semántico: el que puedan referirse a un objeto definido, o puedan tener referencia sólo de manera indefinida. Por supuesto que una

descripción definida no deja de serlo porque carezca de referencia; la expresión «el actual rey de Francia», o «el segundo satélite natural de la Tierra», son descripciones definidas, sólo que vacías, carentes de denotación. Pero Russell parece pensar, de acuerdo con lo que acabamos de ver sobre las expresiones denotativas en general, que puesto que una descripción definida no deja de serlo por carecer de referencia, lo que la hace descripción definida debe ser su forma («La filosofía del atomismo lógico», VI, p. 244 de *Logic and Knowledge*). Hay, sin embargo, una alternativa más rigurosa y exacta: lo que hace que una cierta expresión sea una descripción definida es que, de denotar algo, denotará un objeto determinado, esto es, que su función consiste en referirse a algo definido.

Que los símbolos incompletos, y por consiguiente las descripciones definidas, han de desaparecer tras un análisis de aquellas proposiciones en las que intervienen, le conduce a Russell a afirmar que no tienen significado por sí solos, sino que solamente tienen significado en uso, esto es, en el contexto de la proposición (*op. cit.*, p. 253 de *Logic and Knowledge*, y «Sobre la denotación», pp. 43 y 51 de la misma obra). Y como el significado para él consiste fundamentalmente en la referencia, hay que concluir que una descripción definida, considerada aisladamente, no se refiere a nada (p. 254 de *Logic and Knowledge*). Esto es por demás extraño, y Russell mismo no ha sido siempre, en sus expresiones, congruente con tal idea. Ya hemos visto que distingue, en general, entre aquellas expresiones denotativas que denotan y aquellas que no denotan, y cuando formula esta distinción no menciona para nada que ésta sólo se da en el contexto de proposiciones determinadas. De otro lado, Russell reconoce en algún momento que, en cierto modo, una descripción definida tiene significado, a saber: aquel que le viene dado por lo que significan las palabras que la componen. Así, escribiendo sobre la descripción «El autor de *Waverley*», dice: «Contiene cuatro palabras, y los significados de estas cuatro palabras ya están fijados, y ellos han fijado a su vez el significado de 'El autor de *Waverley*', en el único sentido en el que esta frase tiene algún significado» (*op. cit.*, VI, p. 244). Aquí parece apuntar Russell a algo parecido a lo que Frege había llamado «sentido». Tenemos, pues, que una descripción tiene, por sí sola, significado, en la medida en que éste viene determinado por lo que significan las palabras que la constituyen. Podemos llamar a esto el «sentido» de la descripción. Pero una descripción tan sólo adquiere denotación o referencia cuando es utilizada en el contexto de una proposición verdadera, proposición que, debidamente analizada, resultará ser una proposición cuantificada, como ya hemos visto. Haber usado para ambos casos el término «significado» (*meaning*), sin recurrir a la distinción de Frege entre sentido y referencia, contribuye, sin duda, a oscurecer las afirmaciones de Russell y las hace extremadamente imprecisas en comparación con la doctrina de Frege.

En comparación con lo que se refiere al significado de las descripciones definidas, está perfectamente claro en qué consiste el significado de un nombre propio. El significado de un nombre propio es su referencia o

denotación: «conocer el significado de un nombre es conocer a quién se aplica» (*loc. cit.*, p. 244). Pero no hay que olvidar, aquí, que, en rigor, nombres propios solamente son, para Russell, términos deícticos como los pronombres demostrativos. En cierta medida, se aproximarían a ellos los nombres de propiedades simples y de relaciones igualmente simples, en tanto en cuanto el significado de estos nombres podría reducirse a la denotación de la propiedad o relación en cuestión, la cual siempre podría establecerse directamente por ostensión, puesto que las propiedades y relaciones simples constituyen el contenido de los datos sensibles que componen la experiencia del sujeto —de acuerdo con lo que vimos anteriormente—. Pero lo que con certeza no son nombres propios, en la opinión de Russell, son los así llamados en el lenguaje cotidiano. Esto es muy importante tenerlo en cuenta, ya que, frecuentemente, Russell recurre a nombres propios ordinarios para compararlos con descripciones. Así, compara con la descripción «El autor de *Waverley*», el nombre «Scott» (*loc. cit.*), aunque a veces tiene la cautela de emplear alguna cláusula que evite la confusión del lector, diciendo, por ejemplo, «si usamos 'Scott' como nombre» (p. 252), o «tomando 'Scott' como nombre» (p. 253). En rigor, «Scott» no es un nombre, sino una abreviatura de descripciones, entre otras, de la propia descripción «El autor de *Waverley*».

6.7 Algunos inconvenientes de la doctrina de Russell

Hemos considerado en las tres secciones precedentes las principales tesis del atomismo lógico de Russell acerca de la estructura del lenguaje y de la relación entre el lenguaje y la realidad. Antes de proseguir hacia otra importante versión de esa teoría, hagamos una pausa para echar una mirada crítica sobre los puntos más débiles del atomismo en la versión que ya conocemos.

1. En primer lugar, y siendo el último tema que hemos estudiado, podemos preguntarnos qué se gana al analizar las descripciones definidas a la manera de Russell. Que una proposición de la forma de «El autor de *El Quijote* era castellano» haya de ser sustituida por otra tal como «Hay una única persona que escribió *El Quijote* y era castellana», es lo bastante extraño como para que no le resulte a uno fácil sentirse convencido. Tal análisis implica, como acabamos de comprobar, que las descripciones definidas no tienen por sí solas referencia, sino que únicamente la adquieren en el contexto de una proposición, que siempre será analizable del modo indicado. Los plausibles motivos que inducen a Russell a defender la necesidad de este análisis han sido subrayados, y alabados, anteriormente. Es sin duda útil en extremo contar con un instrumento conceptual que nos ayude a evitar la falacia de la referencia, esto es, la falacia de pensar que siempre que tenemos una expresión del tipo de una descripción definida ha de haber algún objeto al que la expresión se refiera; y más útil aún si

este mismo recurso nos muestra, de paso, que la existencia no es una propiedad más que pueda atribuirse a un objeto, sino precisamente el presupuesto para que se le pueda atribuir al objeto cualquier propiedad, es decir, precisamente el presupuesto para que se pueda hablar de él.

La cuestión es si no podemos conseguir todo esto prescindiendo de un análisis como el anterior. Pues que una descripción definida como «El autor de *El Quijote*», o «El segundo satélite natural de la Tierra», tiene sentido, en la acepción en que Frege toma este término, parece completamente claro, y que de ambas, la primera, pero no la segunda, tiene referencia (aunque referencia histórica, para ser precisos), parece también del todo evidente, pues justamente por eso podemos utilizar esa descripción para hacer afirmaciones, verdaderas o falsas, acerca del personaje que describe, a saber, Cervantes. Y porque la segunda de esas descripciones, en cambio, carece de referencia, no podemos utilizarla para hacer ninguna afirmación, ni verdadera ni falsa, sobre ningún objeto; porque, naturalmente, afirmaciones como «No existe el segundo satélite natural de la Tierra» o «El segundo satélite natural de la Tierra no ha sido descubierto», no son afirmaciones sobre un objeto, sino más bien, la expresión de que no hay tal objeto. Que una descripción definida tenga referencia es un presupuesto para poder hacer afirmaciones, verdaderas o falsas, en las que la descripción aparezca como sujeto (éste es el espíritu de la crítica que, muchos años después, en 1950, le hará Strawson a Russell en «On Referring»). Por eso no me parece necesario mantener que

(1) El segundo satélite natural de la Tierra está a 800.000 kilómetros de ésta

haya de ser sustituido por

(2) Hay una única entidad que es el segundo satélite natural de la Tierra y está a 800.000 kilómetros de ésta

pues la verdad de

(3) Hay una única entidad que es el segundo satélite natural de la Tierra

simplemente constituye el presupuesto necesario para poder hacer una afirmación como (1). Pero puesto que, según nuestros conocimientos actuales, (3) es falsa, (1) por su parte es una afirmación que no puede hacerse, en el sentido de que es vacía, no dice nada ni verdadero ni falso acerca del mundo, carece de valor veritativo. En consecuencia, la negación de (1), a saber,

(4) El segundo satélite natural de la Tierra no está a 800.000 kilómetros de ésta

será igualmente vacía, y carente de valor de verdad. Es cierto que este análisis tiene como consecuencia el que no podamos atribuir valores veritativos a aquellas oraciones que, como (1) y (4), tienen por sujeto un término singular sin referencia, pero tampoco se ve qué necesidad tendríamos de hacer esa atribución ni qué ganaríamos con ello; ¿qué ganamos, por ejemplo, considerando falsas tanto a (1) como a (4), a la manera de Russell? Con afirmaciones que tratan, aparentemente, de objetos inexistentes, no parece verosímil que hayamos de ocuparnos mucho en lógica. (Añadiré que, en tiempos recientes, y bajo el nombre de «lógica libre», se ha elaborado una lógica que acepta términos singulares carentes de denotación, y por consiguiente atribuye valores de verdad a las oraciones en las que tales términos aparecen como sujetos; no es éste lugar para adentrarnos en este tema, pero el lector interesado encontrará una fácil exposición, que por cierto estudia con detenimiento la discusión entre Meinong y Russell, en *Derivation and Counterexample*, de Lambert y van Fraassen, caps. 6, 7 y 10).

2. Los nombres propios ordinarios han de ser sustituidos por descripciones definidas (ya que aquellos son abreviaturas de éstas), y a su vez las descripciones definidas desaparecen dejando en su lugar oraciones cuantificadas. Así, la oración

(5) Miguel de Cervantes era castellano

ha de analizarse como una oración del tipo de

(6) El autor de *El Quijote* era castellano

y ésta por su parte quedará analizada como

(7) Hay una única persona que escribió *El Quijote* y era castellana

El análisis muestra claramente que ni los nombres propios ordinarios ni las descripciones definidas son propiamente nombres, esto es, términos cuya función consista en designar objetos. Esto constituye, sin duda, una notable paradoja. Además, y puesto que cualquier oración declarativa que tenga como sujeto un nombre propio ordinario o una descripción definida resulta equivaler a una proposición existencial compleja, hay que concluir que tales oraciones declarativas no describen ningún hecho simple o atómico. Es decir, (5) o (6), por ejemplo, no describen un hecho; en realidad son proposiciones complejas, cuya complejidad —se supone— queda exhibida en (7).

Esto es coherente con la noción que Russell tiene de un hecho simple. Lo primero que hemos de tener, para poder describir un hecho tal, es un nombre propio auténtico, en sentido lógico, esto es, un término realmente designativo. Las únicas palabras del lenguaje cotidiano que se asemejan a

esta clase de términos son los pronombres demostrativos, primero, porque no dicen nada de los objetos a los que se aplican, y segundo, porque se utilizan fundamentalmente en presencia de los objetos, adquiriendo así todo su significado en el conocimiento directo de los mismos, significado que será siempre comunicable por ostensión de los propios objetos. Es importante darse cuenta hasta dónde influye en la concepción de Russell un presupuesto epistemológico, en cuya virtud los significados de los signos más simples han de ser aquellos objetos de los que tenemos conocimiento directo, familiaridad. Y se advertirá cómo se combina ese presupuesto con una teoría referencialista del significado. ¿Es posible exigir de un lenguaje lógicamente perfecto que sus nombres sean términos que funcionen como los de acuerdo con los dos criterios señalados? En la época en que el atomismo lógico estaba ya en crisis, se señaló que un término nombre como «eso» podría interpretarse como equivalente a la descripción «el objeto al que estoy apuntado», con lo que tampoco sería un nombre lógicamente propio (Wisdom, «Logical Constructions», 1931, citado en Urmson, *El análisis filosófico*, 5.D). Esta consideración, sin embargo, no puede aceptarse como crítica, pues la descripción citada no recurre a ninguna característica del objeto sino tan sólo a la relación entre éste y el sujeto que se refiere a él, con lo que cumple las dos condiciones señaladas, no afirma nada del objeto y se utiliza en su presencia. No parece, pues, imposible tener un lenguaje cuyos nombres sean términos deícticos semejantes a los demostrativos. La consecuencia es que su uso quedará restringido a las experiencias actuales y mientras éstas duren. Es un lenguaje limitado a la expresión de datos sensibles actuales, y por ello, de vocabulario en gran medida privado y, como tal lenguaje, no apto para la comunicación.

¿Pero qué es lo que nombran esos extraños nombres propios? Según Russell, lo que llama «particulares», esto es, aquello que tiene propiedades o entre los que se dan las relaciones, algo, por tanto, parecido a una especie de sustancia dividida. Es claro que los hechos atómicos de Russell implican un dualismo ontológico de sustancias más propiedades (incluyendo en éstas las relaciones). ¿Cuál es la justificación de este dualismo? Los ejemplos de hechos simples suministrados por Russell son datos sensibles: un color, un sonido... Por cierto que, de forma muy confundente, a veces llama «particular» a un dato sensible, como ya hice notar anteriormente (véase «La filosofía del atomismo lógico», pp. 179, 274 y 275 de *Logic and Knowledge*, por ejemplo). Pero si hemos de llevar este empirismo fenomenalista hasta el final, parecería más justificado describir el dato sensible que ahora tengo de algo rojo diciendo «ahí rojez», que afirmando «eso rojo». Pues mi dato sensible, como tal, no contiene más que el color espacialmente localizado. Alguien dirá que todo color lo es de algo, de alguna porción de extensión en el espacio. Ahora bien, una extensión en el espacio no es lo suficientemente simple como para ser un particular en el sentido de Russell, y por consiguiente sería más bien una construcción lógica, como lo son los objetos físicos. En consecuencia: si hemos de atribuir el color a algo, este algo no se ve cómo pueda ser un particular; y si

hemos de prescindir de estas extrañas entidades metafísicas que son los particulares, entonces más vale extremar la fidelidad a los datos sensibles y limitarnos a expresar el contenido de nuestra sensación al tiempo que localizamos espacialmente su estímulo. El dualismo ontológico de Russell parece una incongruencia metafísica dentro de una teoría fenomenalista del conocimiento.

3. La tesis extensionalista, de que toda proposición compleja pueda descomponerse en otras proposiciones simples, de tal manera que el valor veritativo de la primera sea función de los valores de las últimas, es una tesis central en la concepción del lenguaje perfecto, tesis que viene sugerida por la sintaxis lógica de los *Principia Mathematica*. Si se une esta tesis a la idea de que el mundo se compone de hechos, y de que un hecho es aquello que hace a una proposición verdadera o falsa, y si se añade la aparente imposibilidad de descomponer extensionalmente las proposiciones que describen estados psicológicos, así como las proposiciones cuantificadas, sea con el cuantificador general, sea con el cuantificador existencial, tendremos que acabar reconociendo como componentes de nuestro mundo, además de los hechos atómicos, los hechos psicológicos, los hechos generales y los hechos de existencia. A éstos se añadirán, por razones conexas que ya hemos visto, los hechos negativos. Esta ampliación de la ontología, que se produce por razones de lógica, limita de modo drástico los efectos de la navaja de Occam que Russell aplica en otros casos diestramente, por ejemplo, en su teoría de las descripciones. Con el principio de extensionalidad en una mano Russell amplía la ontología, mientras que con la teoría de las descripciones en la otra, la reduce. Pero ni la teoría de las descripciones es necesaria para esta reducción —como he intentado mostrar más arriba—, ni el principio de extensionalidad fuerza a aceptar aquella ampliación, como veremos más adelante a propósito de Wittgenstein.

4. Que el propósito de Russell es, en última instancia, metafísico, está claro. Se trata de llegar a los últimos elementos de los que se compone el mundo, la realidad (p. 270 de *Logic and Knowledge*), pero en teoría, no en la práctica. Se trata de encontrar de qué clase serán esos elementos y cuál su estructura, pero como tal, el lógico, que así es como Russell se considera, no puede decir cuáles son actualmente dichos elementos: «el lógico, como tal, nunca da ejemplos». Que hay bastante más que lógica en la doctrina de Russell, será obvio después de todo lo anterior. Y, sin embargo, no es suficiente. Para hablar como lógico, Russell dice demasiadas cosas sobre el mundo; pero no bastantes para que su doctrina sea, metafísica y epistemológicamente, satisfactoria. El vínculo entre la lógica y la metafísica lo establece la filosofía del lenguaje, que aparece sujeta, por un lado, a la idea de que la estructura del cálculo de los *Principia Mathematica* es la sintaxis propia de cualquier lenguaje que sea lógicamente perfecto, y de otro, a la exigencia de que un lenguaje así esté directamente relacionado con la realidad a base de que sus términos más simples tengan referencia directa en el mundo, lograda a través del conocimiento

directo. Por lo que se refiere a la primera exigencia, el prodigioso desarrollo que tuvo la lógica a partir de la aparición de los *Principia Mathematica*, y la gran variedad de cálculos que vieron la luz, entre ellos, algunos polivalentes, acabó por extinguir la pretensión de encontrar en la lógica la estructura adecuada para un lenguaje perfecto. Por lo que toca a la segunda, ya hemos visto que tenía como consecuencia que el lenguaje perfecto fuera, desde el punto de vista de su vocabulario, en su mayor parte privado, y por ello inútil para la comunicación. Y si el lenguaje perfecto es privado y particular de cada hablante, y además no tiene un claro fundamento lógico, ¿qué ventajas puede tener su investigación?

6.8 El lenguaje como representación figurativa en Wittgenstein

El atomismo lógico de Russell tiene un representante de excepcional significación en Wittgenstein. Aunque en algunos lugares Russell se muestra perplejo o expresa disconformidad con algunas afirmaciones de su discípulo, habla siempre de él, por esta época, con profundo respeto, e incluso se muestra agradecido por alguna sugerencia particular, como la de que las proposiciones no son nombres de los hechos (*Logic and Knowledge*, p. 187). Pero Wittgenstein formuló la doctrina atomista en una obra tan conseguida y brillante, y con un estilo tan personal, que su profundo débito para con Russell y su total falta de originalidad en las ideas básicas de su construcción, han pasado con frecuencia inadvertidos. A ello contribuye también la extrema parquedad de Wittgenstein en las referencias a otros autores, si bien es de notar que Frege y Russell son, con gran diferencia, los autores más citados por él, y que su influencia está ya reconocida en el prólogo de su obra.

Wittgenstein presentó su versión del atomismo lógico en un escrito muy condensado, de párrafos cortos, extrañamente numerados, y de estilo críptico, que apareció en 1921, en el último número que se publicó de la revista alemana *Annalen der Naturphilosophie*. El trabajo llevaba por título «Logisch-Philosophische Abhandlung». Al año siguiente se publicaba en Inglaterra como libro y en edición bilingüe, acompañando al texto alemán la traducción inglesa. Al parecer por sugerencia de Moore, se le puso un título latino: *Tractatus Logico-Philosophicus*.

La forma de numeración de los párrafos del *Tractatus* pretende expresar la importancia lógica que Wittgenstein daba a cada una de sus afirmaciones en relación con las demás. Así, la obra contiene siete afirmaciones principales, numeradas de 1 a 7, y el resto constituyen comentarios sobre éstas, de la siguiente manera: las afirmaciones 1.1, 1.2, etc., son comentarios al párrafo 1, las numeradas 1.11, 1.12, etc., comentan la afirmación 1.1, y así sucesivamente. No creo, sin embargo, que esta numeración, tal y como Wittgenstein la aplica, resulte especialmente aclaratoria, y de hecho sólo se utiliza hoy como un medio, breve y exacto, de citar el *Tractatus*. Las siete aserciones principales contenidas en él son las siguientes:

1. El mundo es todo lo que acontece.
2. Lo que acontece, el hecho, es la existencia de estados de cosas.
3. La representación lógica de los hechos es el pensamiento.
4. El pensamiento es la proposición con sentido.
5. La proposición es una función veritativa de proposiciones elementales. (La proposición elemental es una función veritativa de sí misma.)
6. La forma general de la función veritativa es: $[\overline{p}, \xi, N(\xi)]$. Esta es la forma general de la proposición.
7. Sobre lo que no se puede hablar, se debe guardar silencio.

(Mi traducción del *Tractatus* se apartará en muchas ocasiones de la traducción castellana de Tierno Galván. Justificaré, en su momento, las divergencias más importantes. Como consideración general, nótese que la traducción de Tierno apareció en la *Revista de Occidente* en 1957, una fecha en la que apenas había estudios sobre Wittgenstein o sobre el *Tractatus*. De hecho, y aparte de la primera traducción inglesa, tan sólo una traducción italiana, que salió en 1954, se anticipa a la española —y excluyendo también una traducción china que apareció entre 1927 y 1928—. Pero lo más decisivo a este respecto es que el traductor castellano debió de ayudarse, como es natural y se infiere fácilmente de su versión, con la única traducción inglesa entonces existente, la de 1922, debida a Ogden, que es muy deficiente; tanto, que en 1961 apareció una nueva versión inglesa de Pears y McGuinness muy superior a aquélla, y a cuya luz debería hacerse una seria revisión de la traducción castellana.)

Esas siete escuetas afirmaciones no dan idea exacta del rico y diverso contenido de la obra de Wittgenstein. Simplemente de ellas no puede inferirse que ahí se trate del problema del solipsismo, que se hable de la ética, la estética y la religión, que se considere el lenguaje como un medio de representar figurativamente los hechos, o que se declaren sin sentido todas las proposiciones lógicas y filosóficas, incluidas las del propio *Tractatus*.

La idea básica de Wittgenstein coincide con la de Russell: la lógica conecta con la metafísica a través del análisis del lenguaje. Por lo que, si se considera este último como una simple aplicación de la lógica (y así lo estiman tanto Russell como Wittgenstein), puede afirmarse que la filosofía se compone de lógica y de metafísica. Así lo afirma Wittgenstein en unas notas de 1913 («Notas sobre la lógica»), añadiendo que la primera es la base de la última. Pues, en efecto, es la lógica la que determina la estructura del lenguaje, y en virtud del principio de la isomorfía entre lenguaje y realidad, la que expresa asimismo la estructura de la realidad. Por eso puede decir Wittgenstein que la lógica es la imagen del mundo en un espejo (*Tractatus*, 6.13; en lo sucesivo citaré esta obra exclusivamente con el número del párrafo que corresponda). El *Tractatus* comienza tratando de la estructura del mundo, esto es, empieza por la metafísica, para desarrollar luego la teoría de la proposición, o teoría del lenguaje, y acaba con la teoría de la lógica, que es, fundamentalmente, una teoría de las fun-

ciones veritativas. Este orden puede advertirse con facilidad en las siete aserciones principales reproducidas más arriba. El orden de la obra es, pues, inverso a su orden lógico. Es la filosofía de la lógica la que funda el análisis del lenguaje y éste el que funda la concepción del mundo.

Básicamente, por tanto, el propósito de Wittgenstein coincide con el de Russell. Pero como Wittgenstein dedica mucha más atención al lenguaje, y como su teoría de la proposición resulta ser el centro en torno al cual queda estructurado todo el *Tractatus*, los comentaristas que han querido destacar su originalidad han tendido a formular la tarea que Wittgenstein se propone en términos exclusivos de su teoría del lenguaje. Así, es corriente comparar a Wittgenstein con Kant, y decir que, así como este último intenta responder a la pregunta: ¿Cómo es posible la ciencia?, Wittgenstein estaría intentando dar respuesta a una pregunta semejante: ¿Cómo es posible el lenguaje? A decir verdad, hay fundamento para esta caracterización en las propias palabras de Wittgenstein, pues éste escribe en su prefacio al *Tractatus*: «El presente libro va a trazar un límite al pensamiento, o más bien, no al pensamiento sino a la expresión de los pensamientos, pues para trazar un límite al pensamiento, tendríamos que poder pensar ambos lados de este límite (esto es, tendríamos que poder pensar lo que no se puede pensar). Por tanto, el límite sólo podrá ser trazado en el lenguaje, y lo que se halle más allá del límite será simplemente sinsentido.»

La idea es que el pensamiento, por sí solo, no puede trazarse límites, pues tendría que ser capaz de traspasarlos, por lo que tan sólo en el lenguaje pueden ser puestos tales límites: lo que esté más acá de ellos tendrá sentido, lo que se encuentre más allá será el sinsentido. ¿Qué tiene el lenguaje que no tenga el pensamiento y que permite trazar en aquél el límite que no puede trazarse en éste? Lo veremos más adelante, pero es claro que, si más allá del límite no hay sentido, entonces tampoco el lenguaje puede cruzarlo. Por ello, se ha comparado el intento de Wittgenstein con la tarea de recorrer el límite de una esfera sin salir de ella, desde dentro. Pero además, puesto que el análisis del lenguaje suministra una visión de la estructura de lo real, y ya que aquel se funda en el análisis de la lógica, que es algo que no requiere para nada el recurso a la experiencia, el propósito del *Tractatus* puede también enunciarse como un empeño metafísico: averiguar, sin recurrir a la experiencia y, por consiguiente, por medios *a priori*, cuál es la estructura de lo real. Así lo formula Wittgenstein en una nota de 1915: «El gran problema en torno al cual gira todo cuanto escribo es: ¿hay, *a priori*, un orden en el mundo, y si lo hay, en qué consiste?» (*Notebooks* 1914-1916, p. 53; esta obra recoge diferentes notas preparatorias del *Tractatus*, y fue publicada en 1961, diez años después de la muerte de Wittgenstein).

Aquí no seguiremos, para la exposición del *Tractatus*, ni el propio orden que en éste llevan los temas, desde la metafísica a la filosofía de la lógica, ni tampoco el orden lógico de su justificación, desde la lógica al análisis de la realidad. Seguiremos un camino intermedio. Partiremos de

la teoría del lenguaje, que es, en definitiva, lo que más nos interesa por razón de nuestro tema, y exploraremos desde ella sus consecuencias para la concepción de la realidad y sus conexiones con la teoría de la lógica. Esta presentación no tiene por qué deformar la arquitectura del *Tractatus*, y es, de otro lado, la más coherente con el desarrollo anterior de nuestro trabajo.

Antes de proseguir conviene, sin embargo, advertir algo. He empleado ya varias veces expresiones como «teoría del lenguaje», «filosofía de la lógica», «metafísica», para referirme a diversos aspectos del contenido del *Tractatus*. Al hacerlo así me he limitado a aplicar a la obra de Wittgenstein categorías corrientes en el discurso filosófico. No hay que pensar, empero, que fuera el propósito de Wittgenstein ofrecer en su obra una teoría del lenguaje, de la realidad o de la lógica, al modo usual. La razón es que toda teoría es un discurso con sentido, mientras que las afirmaciones de Wittgenstein en el *Tractatus* son declaradas por él mismo como sin-sentidos, esto es, se hallan al otro lado de la frontera que separa el lenguaje con sentido de lo que no lo tiene. En el comentario que sigue prescindiremos de este aspecto de la metafilosofía de Wittgenstein, pero lo consideraremos explícitamente cuando llegue el momento.

La doctrina sobre el lenguaje no es, en el *Tractatus*, más que una porción distinguida, y particularmente importante desde un punto de vista filosófico, de lo que podemos llamar, en general, la teoría de las representaciones figurativas o isomórficas. De aquí que el mejor modo de entrar en la primera sea considerar lo que Wittgenstein afirma sobre estas últimas. El primer lugar en que habla de este tema es la proposición 2.1, donde dice: «Nos hacemos representaciones de los hechos». Se trata de una afirmación extraña dentro de las coordenadas del *Tractatus*; aparentemente, es una especie de generalización empírica evidente y trivial, y su colocación dentro de la parte que trata de la estructura de la realidad (por su número, sería un comentario a la proposición 2), oscurece el hecho de que con ella se inicia en realidad una nueva parte de la obra, probablemente la más original, a saber, la teoría de la representación figurativa. Aparece ya aquí un término que Wittgenstein va a utilizar continuamente en este contexto y que ha servido para dar nombre a su teoría, el término *Bilder* (en singular, *Bild*), que acabo de traducir como «representaciones». Con este término, Wittgenstein se refiere a aquellas formas de representación de los hechos que tienen con éstos una relación tal que: primero, a cada elemento de lo representado corresponde un elemento en la representación; y segundo, a las relaciones que hay entre los elementos del hecho corresponden relaciones entre los elementos de la representación. Se trata de representaciones isomórficas. ¿Cómo traducir este término?

Bild no es un término técnico filosófico; en su uso cotidiano incluye todo tipo de representaciones gráficas, como retratos, pinturas, cuadros, fotografías, dibujos, y también las imágenes ópticas. El inglés tiene un término de alcance muy parecido que es *picture*, y con el que queda bien traducido el vocablo alemán; tiene además el verbo *depict*, que sirve para traducir con bastante exactitud el verbo *abbilden*, que Wittgenstein utiliza

con frecuencia para referirse a la relación que hay entre una *Bild* y lo representado en ella. Por su parte, la traducción castellana del *Tractatus* recurre al sustantivo «figura» y al verbo «figurar». El significado de estas palabras, sin embargo, no coincide con el de los términos alemanes que traducen; «figura» es palabra que suele hacer referencia a la forma de un cuerpo, y a una representación gráfica en dos dimensiones no solemos llamarle «figura» excepto si aparece como ilustración en un libro, y por oposición al texto escrito, como tampoco llamamos «figura» a una imagen óptica. En conformidad con dicha traducción, las expresiones de Wittgenstein *die Form der Abbildung* y *die abbildende Beziehung* se vierten, respectivamente, como «la forma de figuración» y «la relación figurativa», traducciones de las que la primera resulta bastante forzada y puede llegar a ser confundente a causa de que usualmente se entiende por «figuración» una imaginación o fantasía.

Para traducir *Bild* y *abbilden* han sido propuestos también los términos «pintura» y «pintar», respectivamente; así, Deaño, en su traducción del libro de Kenny, *Wittgenstein*, y Ferrater Mora en «Pinturas y modelos». Creo que esta traducción presenta todavía más problemas que la anterior. El término «pintura» tiene un significado mucho más restringido que el de *Bild*; una pintura es una representación realizada sobre una superficie por medio de líneas y colores, pero no llamamos «pintura» a una fotografía, a una imagen reflejada ni a una representación tridimensional isomórfica, a todas las cuales consideraría Wittgenstein como *Bilder*. De otra parte, al traducir *abbilden* como «pintar», determinadas expresiones del *Tractatus* hay que traducirlas con infracción de las reglas semánticas del castellano, diciendo, por ejemplo, de una pintura, que «pinta la realidad» (traducción castellana del *Wittgenstein* de Kenny, pp. 60 y 69). Pero no son las pinturas las que pintan, sino los pintores; «pintar» requiere un sujeto humano. Este último inconveniente podría eludirse traduciendo *abbilden* por «ser una pintura de», pero el término «pintura», en todo caso, es excesivamente restringido y tiene connotaciones artísticas totalmente ajenas al concepto de representación isomórfica, que es lo que Wittgenstein quiere expresar. De aquí que resulten, asimismo, confundentes, en mi opinión, expresiones como «forma pictórica» y «relación pictórica», que, en coherencia con lo anterior, se ofrecen para traducir, respectivamente, *Form der Abbildung* y *abbildende Beziehung*.

Una tercera alternativa, que solucionaría de un golpe los problemas estilísticos, sería verter los términos y expresiones mencionados respectivamente como «representación», «representar», «forma de representación» y «relación de representación». Esta traducción, en cambio, pecaría por exceso, pues no es a cualquier clase de representación a lo que Wittgenstein llama *Bild*, sino a las que tienen, con respecto a lo representado, una relación de isomorfía. Naturalmente, la traducción sería del todo exacta añadiendo simplemente este último término como calificativo y haciendo nuestras expresiones un poco más complejas que las originales. de esa forma: «representación isomórfica», «representar isomórficamente», etc. Bien es

verdad que, con esto, perdemos la sobriedad y elegancia del estilo de Wittgenstein. Lo que parece claro es que no es posible encontrar una palabra que traduzca adecuadamente al español, como sustantivo y como verbo, esta terminología del *Tractatus*. A la vista de ello, optaré en lo sucesivo por traducir *Bild* alternativamente como «representación isomórfica», «representación figurativa», o, a modo de abreviaturas de estas expresiones, simplemente como «representación» o «figura», pues debe haber quedado ya claro de qué clase de representaciones o figuras habla Wittgenstein; *abbilden* lo traduciré por «representar isomórficamente» o «representar figurativamente», y de forma abreviada, por «representar» o «figurar». A estos efectos debe tenerse en cuenta que Wittgenstein usa en muchos lugares del *Tractatus*, como sinónimo de *abbilden*, el término *darstellen*, que corresponde más estrechamente a «representar»; así, usa ambos términos en la proposición 2.201, pero sigue utilizando sólo *darstellen* en las siguientes, hasta la 2.221; e incluso emplea *vorstellen* en algunos lugares como la 2.11 y la 2.15. Sobre la base de lo anterior, cabría pensar que la expresión *Form der Darstellung*, que Wittgenstein utiliza en 2.173 y 2.174, haya que tomarla como sinónima de *Form der Abbildung*, pero es posible que Wittgenstein introduzca aquí una diferencia, como veremos un poco más adelante.

Lo que hace de algo una representación o figura es que consta de elementos, cada uno de los cuales se refiere a un objeto de la realidad representada, y que esos elementos están entre sí relacionados de manera correspondiente a como lo están los objetos representados, si la representación es correcta (2.13 a 2.15). Tanto la representación como lo representado son, por consiguiente, relaciones, entre las cuales hay una ulterior relación que las correlaciona. La isomorfía no es, en definitiva, sino una relación entre relaciones, y dos relaciones son isomorfas siempre que hay entre ellas una relación de correlación. La relación correlatora pone cada argumento de una relación en conexión con un argumento, y sólo uno, de la otra relación, de tal manera que si, por ejemplo, x e y son dos argumentos correlacionados, y w y z también lo son, y entre x y w hay una cierta relación R , entonces habrá una cierta relación R' entre y y z . Y diremos que R y R' son isomorfas. La isomorfía es, por consiguiente, una relación diádica entre relaciones n -ádicas, esto es, de cualquier grado de complejidad. Estas ideas pueden ser formuladas de modo técnico y riguroso con los recursos de la lógica, pero el hacerlo no sería relevante para el curso ulterior de nuestro comentario, y Wittgenstein tampoco recurre a tecnicismos cuando expone su concepción de las representaciones isomórficas (el lector interesado encontrará una clara exposición elemental de las relaciones isomórficas en la sección 89 de la *Introducción a la lógica y al análisis formal* de Sacristán; en *A Model Theoretic Explication of Wittgenstein's Picture Theory* Stegmüller ha suministrado una formulación técnica para la teoría figurativa de Wittgenstein que, si bien es disputable en qué medida responda con exactitud a su pensamiento, proporciona, no obstante, una reconstrucción lógica de su teoría sumamente rica en implicaciones y matices, y cier-

tamente mucho más clara y completa que cuanto pueda encontrarse en el *Tractatus*).

Las correlaciones de los elementos de la representación con los elementos de la realidad representada constituyen lo que Wittgenstein llama «relación de representación» (*abbildende Beziehung*, 2.1514). Pero para que algo sea una representación, en este sentido, ha de poseer, además, lo que Wittgenstein denomina «forma de representación» (*Form der Abbildung*, 2.15 s.), y que describe como la posibilidad de la estructura de la representación (2.15), o, con otras palabras, «la posibilidad de que las cosas se hallen relacionadas entre sí como los elementos de la representación» (2.151). Para entender esto no está de más recordar el sentido aristotélico del término «forma», como aquello que hace que algo sea lo que es. Lo que hace que algo sea una representación figurativa es que se trata de una estructura de elementos a la que *puede* corresponder una estructura de cosas en el mundo. Lo que importa es, pues, que es *posible* que se dé en el mundo una estructura o relación de objetos como la que hay entre los elementos de la representación. ¿Por qué esta alusión a la posibilidad? Porque una representación puede representar algo correcta o incorrectamente, verdadera o falsamente, según concuerde o no con los hechos (2.21-2.222). Pero una representación falsa no es menos representación que una verdadera; una pintura de la fachada principal de la Cámara de los Diputados en la que los leones aparezcan arriba en la escalinata, flanqueando la puerta principal, es una representación incorrecta, falsa, de ese edificio, pero no deja de ser, por ello, una representación. Pues bien *podría* haber sucedido que, con extraño sentido de la estética, los leones hubieran sido colocados allí arriba y no en su emplazamiento actual.

Este punto es sumamente importante para entender tanto el pensamiento de Wittgenstein sobre el sentido de las proposiciones como su metafísica. Lo que hace de algo una figura o representación es que *es posible* que se dé lo que la representación representa. La forma de representación es, simplemente, una posibilidad, la posibilidad de que la representación sea correcta o verdadera. Y esta posibilidad, que es la forma de representación, es lo común a la figura y a lo representado por ella (2.16-2.17). Pero repárese en que, si una figura o representación es falsa, entonces lo representado, tal y como está en ella representado, no existe. Y si no existe, ¿cómo puede tener algo en común con su representación? Muy sencillo: porque eso que hay de común es la posibilidad de existencia; tal posibilidad es idéntica a la figura y a lo representado en ella, aunque esto último sea inexistente (2.161). Si llamamos «mundo posible» a cualquier conjunto de hechos posibles que sea consistente, entonces podemos decir que a toda representación corresponde un hecho en algún mundo posible, y por ello, que toda representación es verdadera o correcta en algún mundo posible. La doctrina de las representaciones isomórficas en el *Tractatus* estaría adelantando la moderna semántica de los mundos posibles, que más adelante estudiaremos. La forma de representación puede tener diferentes matices según sea la realidad representada; una figura material puede representar

algo material, una figura coloreada puede representar, además, algo coloreado, etc. (2.171). Y puesto que la forma de representación no es sino la posibilidad de que exista lo representado, ello quiere decir que una figura material expresa la posibilidad de que exista algo material, una figura coloreada la de que exista algo coloreado, etc.

Como ya advertí antes, es posible que Wittgenstein pretendiera distinguir entre la forma de representación o forma figurativa (*Form der Abbildung*), sobre la cual acabamos de discurrir, y lo que llama en un par de sitios (2.173-2.174) *Form der Darstellung*. Puesto que dice de ésta que es el punto de vista de la representación, cabe pensar que se trata de aquello que distingue a esta última de la realidad representada, y podríamos entonces traducir dicha expresión como «forma propia de cada representación». Así, en un retrato de una persona, la espacialidad, la forma humana, el colorido, etc., pertenecerían a la forma de representación o forma figurativa en cuanto que es lo que el retrato tiene en común con la realidad representada en él; en cambio, la bidimensionalidad del retrato pertenecería a la forma propia de representación del mismo, puesto que es algo que distingue al retrato de lo retratado, la persona, que es tridimensional. Esta es la interpretación que, por ejemplo, sugiere Kenny (*Wittgenstein*, cap. 4, p. 57 de la edición de Penguin). De todas formas, Wittgenstein concede tan poca atención a la *Form der Darstellung*, que hay que pensar que, o bien no le daba importancia teórica, o bien la tomaba como sinónima de *Form der Abbildung*. En consecuencia, no nos volveremos a ocupar de aquella.

Sea cual sea la riqueza de la forma figurativa, hay algo que, como mínimo, ésta debe poseer: sea o no sea material, sea o no sea coloreada, sea en dos o en tres dimensiones, sea estática o dinámica, una representación ha de tener, para serlo, una forma mínima que es lo que Wittgenstein llama «forma lógica» (*logische Form*, 2.18). Puesto que toda representación ha de tener, como mínimo, esta forma, toda representación es una representación lógica (*logische Bild*, 2.181 s.), y eso, cualquiera que sea su determinación posterior, esto es, con independencia de que se trate de una representación espacial, coloreada, etc. Pero puesto que la forma es aquello en lo que coinciden la representación y lo representado, lo anterior implica que todo aquello que pueda ser representado, en tanto en cuanto puede serlo, es lógico; por eso dice Wittgenstein: «la forma lógica, esto es, la forma de la realidad» (2.18). Con ello queda formulado explícitamente el principio de isomorfía tal y como Wittgenstein lo entiende: la realidad es representable en la medida en que tiene una estructura o forma lógica, justamente el tipo de estructura o forma que posee toda representación por el hecho de serlo. En la forma lógica coinciden nuestras representaciones de la realidad y la realidad en cuanto representada (2.2). Russell nunca había llegado a una fórmula tan explícita ni tan general (nótese que todavía no hemos hablado del lenguaje).

La forma lógica, en cuanto nuda y básica forma de representación, expresa la mera posibilidad de existencia de lo representado sin más deter-

minación, esto es, prescindiendo de toda otra propiedad. Esto se halla conectado con la idea de Wittgenstein de que una figura representa una situación posible en el espacio lógico (2.202). Entiendo que el espacio lógico es el ámbito creado por las reglas de la lógica, que Wittgenstein considerará en una parte posterior del *Tractatus*. En ese ámbito, la forma lógica, esto es, la estructura de toda situación o hecho posible en cuanto posible, permite la representación de este último. El espacio lógico y el ámbito de lo posible son lo mismo, pues la lógica es anterior a la experiencia, es anterior a que los hechos sean tales o cuales (5.552)*. Sólo puede representarse aquello que es posible, y que, de hecho, será existente o no existente (2.11, 2.201). Si lo representado existe, la representación será verdadera; si no existe, será falsa (2.21). Pero sea lo uno o lo otro, la representación, en cuanto representación, tiene un sentido (*Sinn*), que es la situación representada (2.22 s.). Para decidir si es verdadera o falsa tendremos que comparar la representación con la realidad, a fin de comprobar si lo representado existe o no; en consecuencia, no hay representaciones que sean verdaderas *a priori*, con independencia de la experiencia (2.223-2.225). Lo único que puede decirnos la lógica es que toda representación o es verdadera o es falsa, pero no si es lo uno o lo otro. La forma lógica, sin embargo, no es parte del sentido de la representación, ya que, en cuanto que es lo que hace posible el representar, no es, a su vez, representada. La forma de representación, tanto en su aspecto meramente lógico (forma lógica), como en cualquier otra determinación que tenga (forma espacial, coloreada, etc.), no es propiamente representada por la figura, sino exhibida o mostrada por ella (2.172). La representación representa una situación posible y muestra lo que tiene en común con dicha situación, a saber, la forma de representación. Esta doble función tendrá consecuencias muy llamativas cuando Wittgenstein aplique su doctrina de la representación al lenguaje.

6.9 Teoría de la proposición

Después de exponer su doctrina sobre las representaciones o figuras, y antes de empezar a hablar propiamente del lenguaje, Wittgenstein hace unas reflexiones sobre el pensamiento. Su primera afirmación puede malentenderse fácilmente; dice: «La representación lógica de los hechos es

* Es interesante recordar aquí que, en un manuscrito inédito de 1913, que se conoce bajo el título de *Theory of Knowledge*, Russell había mantenido que, para entender el significado de las proposiciones, tenemos que estar familiarizados no sólo con los particulares y sus propiedades, sino también con la forma lógica. Como es claro, esto implica que en la lógica necesitamos la experiencia. Se sabe que Wittgenstein criticó duramente ese manuscrito, el cual Russell, acaso por esta razón, nunca publicó. La posición de Wittgenstein en el *Tractatus*, según acabamos de ver, es contraria a dicha tesis de Russell. Hay una interesante tesis doctoral de María Teresa Iglesias sobre este tema, bajo el título *Linguistic Representation: A Study of B. Russell and L. Wittgenstein 1912-1922* (Universidad de Oxford, 1979).

el pensamiento» (3). Ahora bien, acabamos de ver que toda representación es una representación lógica por definición. ¿Son todas las representaciones pensamientos? Sí, en el sentido siguiente. Lo que la proposición 3 implica, a mi juicio, es esta doble consecuencia: el pensamiento es aquella representación que es meramente lógica, es decir, aquel modo de representar cuya forma de representación es exclusivamente lógica, y que carece de cualquier otra determinación, no es espacial, ni cromática, etc.; y puesto que toda forma de representación incluye la forma lógica, toda representación, sea del tipo que sea, incluye un pensamiento. Por ejemplo: el pensamiento del libro rojo sobre el libro verde que tengo ante mi vista estará incluido en cualquier otra representación (pintura, dibujo, fotografía, escultura, esquema...) que yo pueda hacer de tal hecho. Los demás tipos de representación poseerán propiedades que no tiene el pensamiento, pero el pensamiento de lo representado estará incluido en cualquier representación como su forma lógica. Por eso añade Wittgenstein: que un estado de cosas pueda ser pensado significa que podemos hacernos de él una representación (3.001).

Todo lo demás que, a continuación, afirma sobre el pensamiento no es sino aplicación de lo que ya hemos visto a propósito de las representaciones en general. Lo que puede pensarse es posible (3.02), puesto que sólo lo posible puede representarse, y no podemos pensar nada que infrinja la lógica, ya que es la lógica la que crea el ámbito de lo posible, y por tanto, de lo representable (3.03), así como tampoco podemos decir qué aspecto tendría un mundo ilógico (3.031). Por lo mismo, el conjunto de los pensamientos verdaderos nos da una representación del mundo (3.01), y la verdad de un pensamiento, como la de cualquier representación, depende de cómo sean los hechos, pues no hay pensamientos que sean verdaderos *a priori* (3.04 s.).

Sobre el pensamiento, en cuanto pura figura lógica, no hay mucho más que decir, o, si lo hay, puede ser difícil intentar decirlo, en la medida en que el pensamiento no es un objeto perceptible, abierto a la experiencia intersubjetiva. Pero todo lo que pueda decirse sobre el pensamiento, se puede decir acerca del lenguaje, en el que aquel se materializa y objetiva. Por eso el siguiente paso de Wittgenstein es abordar el lenguaje y formular sus reflexiones tomando como objeto la proposición, pues «en la proposición (*Satz*) se expresa con sentido y de manera perceptible el pensamiento» (3.1; en un contexto actual sería preferible traducir *Satz* como «oración», pero en el contexto del atomismo lógico russelliano, en que se mueve Wittgenstein, es más adecuado «proposición», como suelen traducir los comentaristas del *Tractatus*).

No hay, por consiguiente, otra diferencia entre el pensamiento y el lenguaje que la que procede de que este último consiste en signos externos, signos proposicionales (*Satzzeichen*, 3.12), por medio de los cuales se expresa el primero. Por lo demás, pensamiento y lenguaje son idénticos; «el pensamiento es la proposición con sentido».(4). Y ya en una nota de 1916 había escrito: «el pensamiento es una forma de lenguaje» (*Notebooks*,

p. 82). Sin embargo, de acuerdo con el principio de isomorfía, toda representación debe constar de elementos que se correspondan, uno a uno, con los elementos del hecho representado. Podemos pensar que, en la proposición, y a reserva de las precisiones que haremos ulteriormente, esos elementos son las palabras, o acaso, los morfemas. ¿Pero cuáles son los elementos constitutivos del pensamiento y cuál es su relación con los elementos del hecho pensado? En una carta de 1919, en la que Wittgenstein contestaba a diversos comentarios de Russell sobre el *Tractatus*, escribe: «No sé cuáles son los constitutivos del pensamiento, pero sé que ha de haber constitutivos que correspondan a las palabras del lenguaje. Es irrelevante qué clase de relación haya entre los elementos del pensamiento y los del hecho representado; averiguarlo sería asunto de la psicología (...). '¿Consiste en palabras un pensamiento?' ¡No! Consiste en constitutivos psíquicos que poseen el mismo tipo de relación con la realidad que las palabras. Pero no sé cuáles son esos constitutivos.» (*Notebooks*, p. 129 s.).

El signo proposicional es un hecho (*Tatsache*, 3.14), como lo es cualquier representación (2.141), incluido el pensamiento (*Notebooks*, *loc. cit.*). Lo que lo hace signo es que sus elementos, las palabras, están articulados, relacionados entre sí de cierta manera (3.14 s.). Es importante, y totalmente representativo de la concepción de Wittgenstein, que para aclarar cuál es la esencia del signo proposicional sugiera imaginar una proposición compuesta, no por signos escritos, sino por objetos tridimensionales como mesas, sillas y libros, añade: «la posición espacial recíproca de estas cosas expresará, entonces, el sentido de la proposición» (3.1431). Esto nos advierte ya contra la posibilidad de olvidarnos del principio de isomorfía al analizar el sentido de las proposiciones, y muestra la determinación de Wittgenstein de aplicar su modelo hasta el final: el sentido de una proposición no difiere esencialmente del sentido de cualquier otra representación isomórfica; lo mismo da que usemos palabras o que usemos cosas: el sentido es la correlación estructural que la representación (o la proposición) tiene con lo representado.

Los elementos últimos de la proposición son aquellos signos simples a los que llegamos cuando la hemos analizado del todo. Según Wittgenstein, estos signos son nombres (3.2-3.202). «El nombre significa (*bedeutet*) el objeto, y éste es su significado (*Bedeutung*)» (3.203). He aquí que aflora, al fin, la teoría referencialista que hemos visto en Russell, y que, en todo caso, estaba ya por detrás del principio de la isomorfía de las representaciones. Las proposiciones se descomponen en nombres, sus elementos o signos más simples no son sino nombres, y el significado de éstos es, simplemente, el objeto al que cada uno se refiere. La profesora Anscombe, en la introducción de su excelente *Introducción al Tractatus de Wittgenstein*, ha afirmado que Wittgenstein sigue a Frege en su uso de los términos *bedeuten* y *Bedeutung*, los cuales, por ello, deben traducirse como «referirse a» y «referencia», y no, en cambio, como «significar» y «significado», según acabo de hacer. A decir verdad, no conozco ninguna prueba definitiva en favor de esa tesis, excepto el hecho de que Wittgenstein distingue

en el *Tractatus* entre *Sinn* y *Bedeutung*, aplicando la primera categoría a las proposiciones y la segunda a los nombres, y esto lo hace en contextos en los que está hablando de Frege (por ejemplo, en 5.4733). Que este último había influido en Wittgenstein ya lo hemos subrayado al principio, y está fuera de duda. Es, desde luego, llamativo que Wittgenstein nunca emplee el término *Bedeutung* cuando habla de proposiciones, y esto no resulta fácil de explicar si está empleando el término en su amplio sentido ordinario. Mas téngase en cuenta que, aun cuando tomara el término de esta forma, y no en la acepción de Frege, nada tendría que cambiar en la interpretación del *Tractatus*, si mantenemos a la vista el contexto russelliano de esta obra. Pues para Russell, como vimos, el significado de los nombres lógicamente propios, que son signos simples, se agota en la referencia, y es del todo coherente con la posición de Russell afirmar que el significado de un nombre es un objeto. De cualquier forma, y para dejar aún más clara la posición de Wittgenstein, traduciré en lo sucesivo *bedeuten* y *Bedeutung* como «denotar» o «referirse a», y «denotación» o «referencia», apartándome así de las traducciones usuales (incluyendo la segunda traducción inglesa que da *mean* y *meaning*, respectivamente).

A los nombres de la proposición corresponden los objetos del hecho representado, y a la configuración de aquellos en la proposición corresponde la configuración de los objetos en el hecho (3.21 s.). De aquí que la única manera de hablar de los objetos sea nombrándolos, mientras que los hechos o situaciones no pueden, en cambio, ser nombrados, sino sólo descritos (3.144 y 3.221). Describir es representar la estructura del hecho por medio de la estructura (isomorfa) de la proposición; tal estructura es el sentido (*Sinn*) de la proposición. Nombrar es poner un signo simple en el lugar de la estructura que corresponde a un objeto; un signo es un nombre sólo cuando funciona como tal en el contexto de una proposición. Por eso afirma Wittgenstein: «Sólo la proposición tiene sentido; sólo en la conexión de la proposición tiene referencia un nombre» (3.3). Como puede apreciarse, hay aquí una importante divergencia respecto de Frege, para la cual tiene un importante papel la doctrina de Russell. Según Frege, un nombre tiene, además de referencia, un sentido, que vendrá dado por cualquier descripción definida que podamos hacer del objeto en cuestión; de modo análogo, una proposición tiene, además de sentido, referencia, la cual consistirá en su valor veritativo. Para Wittgenstein, que aquí se limita a seguir a Russell, un nombre, si lo es realmente y en sentido lógico, se reduce a nombrar, y por tanto no puede tener sentido; si tuviera sentido serviría para describir el objeto, y entonces no sería un signo simple, sino que encerraría alguna complejidad. De modo contrario, una proposición tiene sentido, a saber, el hecho posible que representa, pero no puede tener referencia, pues la proposición no es nombre de nada; decir, como dice Frege, que es nombre de su valor veritativo, no es inteligible, porque el valor veritativo no es nada externo a la proposición, y con lo cual pueda ser comparada ésta, sino algo que pertenece a la relación entre la proposición y lo representado en ella; el valor veritativo es la expresión de que lo

representado por la proposición existe o no existe. Los nombres, pues, poseen referencia pero no sentido; las proposiciones tienen sentido, pero no referencia. En qué consista el sentido de las proposiciones queda aclarado por Wittgenstein con su aplicación del principio de isomorfía, aclaración que va, sin duda, mucho más allá de cuanto Frege y el propio Russell habían sugerido. ¿Pero en qué consiste la referencia de los nombres, cómo la obtienen y cómo puede explicarse a los demás?

La respuesta a estas preguntas había obligado a Russell a traer a colación el conocimiento por familiaridad, pasando así al terreno de la teoría del conocimiento. Wittgenstein, por su parte, exhibe a lo largo del *Tractatus* una notable resistencia a entrar en ese terreno y una sorprendente habilidad para evitarlo. De aquí que la conexión entre los nombres y los objetos de la realidad quede sumida en una incómoda oscuridad. Wittgenstein se limitará a decir que los nombres no pueden ser descompuestos ulteriormente por medio de una definición, puesto que son signos simples y, por tanto, primitivos (3.26). Ahora bien, lo que no se expresa en el nombre (a saber, su conexión con el objeto), lo muestra su aplicación (3.262); por eso, la denotación de los nombres, o signos primitivos, puede explicarse por medio de aclaraciones, esto es, por medio de proposiciones que contengan dichos signos (3.263). El uso de estas proposiciones mostrará a qué se refieren los nombres que aparezcan en ellas. Pero como el propio Wittgenstein se cuida de añadir (3.263), tales proposiciones sólo pueden entenderse si se conoce la denotación de sus signos. La conclusión de este aparente círculo vicioso únicamente puede ser que el uso de un lenguaje presupone la conexión entre sus signos simples y los objetos del mundo, y que esta conexión no puede ser propiamente explicada, sino simplemente mostrada, enseñando cómo se usa el lenguaje.

Una proposición, por consiguiente, no es más que una representación figurativa de la realidad, un modelo de la realidad tal y como la concebimos (4.01). Esto puede resultar, a primera vista, extraño. ¿Cómo puede decirse que la mera secuencia de signos gráficos en que consiste una proposición escrita sea una figura de la realidad? ¿No estará hablando Wittgenstein metafóricamente? La comparación, que ya hemos encontrado antes, entre una proposición y una representación tridimensional por medio de cosas, debe apartar de nosotros cualquier tentación de debilitar la doctrina semántica que Wittgenstein está exponiendo. Por si acaso, vuelve a insistir, sugiriendo: «para comprender la esencia de la proposición pensemos en la escritura jeroglífica, que representa figurativamente los hechos que describe» (4.016). Así pues, la esencia del lenguaje se halla manifiesta, paradójicamente, en la forma de escritura que parece menos lingüística, en los jeroglíficos. Hay, incluso, cierto tipo de proposiciones cuya forma resulta particularmente figurativa, por ejemplo, aquellas que son de la forma *aRb*, en las cuales se da una semejanza entre el signo y lo significado (lo significado es, naturalmente, una relación entre dos objetos, 4.012).

Pero una semejanza tal no es necesaria para que podemos reconocer en la proposición una figura. Pues también la notación fonética y el alfabeto

son representaciones del habla, o la notación musical, representación de la música, y ciertamente no podríamos hablar en estos casos de semejanza (4.011). La semejanza no es, por consiguiente, lo que hace de algo una representación figurativa, y ésta es la razón por la que pensar en pinturas puede ser muy confundente. Wittgenstein llegará a decir que también el disco de gramófono, el pensamiento musical, la partitura y las ondas sonoras se hallan entre sí en una «relación interna de representación» como la que hay entre el lenguaje y el mundo; y añadirá a modo de aclaración: todas estas cosas tienen una estructura lógica común (4.014). Hay que evitar, pues, las posibles connotaciones visuales del término *Bild* y no perder de vista lo que implica el principio de isomorfía. Entre el lenguaje, el pensamiento y la realidad hay una correlación de estructuras, como la hay entre el disco, la composición musical y la partitura. Esa correlación de estructuras, o relación de isomorfía, es la que permite obtener la sinfonía a partir de la partitura o a partir del disco, y viceversa; por idéntica razón podemos pasar de un hecho a su expresión lingüística, o de ésta al pensamiento que contiene, y del pensamiento otra vez a la expresión lingüística. En todos estos casos tenemos una regla de traducción que nos permite pasar de lo uno a lo otro, o, como dirá con expresión un poco más oscura, una ley de proyección que proyecta una de esas cosas en términos de la otra (4.0141).

Lo anterior debe corregir, sin duda, ciertas alusiones como la que acabamos de ver a la escritura jeroglífica. Casos como éste, en los que la relación de isomorfía es visualmente perceptible y se convierte, por ello, en semejanza externa, son casos extremos que pueden ayudarnos a captar cómo concibe Wittgenstein la representación figurativa, pero no expresan lo fundamental de su concepción. Son ejemplos especialmente llamativos, pero nada más.

La proposición es la descripción de un estado de cosas (*Sachverhalt*) o situación (*Sachlage*) (4.023, 4.031). Aquí se requiere otra aclaración terminológica. La traducción castellana del *Tractatus* vierte *Sachverhalt* como «hecho atómico», siguiendo la primera traducción inglesa. Parece que Wittgenstein estaba conforme con la traducción del término alemán como *atomic fact*, al menos en la medida en que, tomando esta expresión con el sentido que tiene en Russell, corresponde en principio a lo que Wittgenstein quiere dar a entender por *Sachverhalt*; así, en una carta a Russell (*Notebooks*, p. 129), comenta que *Sachverhalt* es lo que corresponde a una proposición elemental si es verdadera. Es, de todas formas, llamativo que Wittgenstein nunca utilice ninguna expresión alemana equivalente a «hecho atómico», y acaso esto deba sugerirnos que no entendía por *Sachverhalt* exactamente lo mismo que Russell llamaba *atomic fact*. La primera traducción inglesa, y por tanto también la castellana, tienden a dar una imagen demasiado russelliana de la doctrina de Wittgenstein, imagen que, como vamos a ver, en este punto es muy probablemente errónea. De forma coherente, Wittgenstein nunca emplea ninguna expresión como «proposición atómica», sino que usa siempre *Elementarsatz*, «proposición elemental».

Por estas razones, parece más riguroso, como hace la segunda traducción inglesa, verter los términos alemanes más literalmente, y de acuerdo con ella, traduciremos aquí *Sachverhalt* como «estado de cosas». En cuanto a *Sachlage*, es palabra que Wittgenstein utiliza con menos precisión que la anterior, pero en gran número de ocasiones es claramente sinónima de aquélla; la traduciremos como «situación» (el texto castellano del *Tractatus* da, en cambio, para ella, «estado de cosas»).

Una característica de las proposiciones, que Wittgenstein quiere subrayar de acuerdo con su teoría de la representación, es que el sentido de aquéllas es previo a su verdad o falsedad, y por ello una proposición puede ser entendida sin necesidad de saber si es verdadera o falsa (4.022, 4.024). Entender una proposición es captar su sentido, o lo que tanto vale, conocer la situación que representa (4.021, 4.031), y ello implica saber que los hechos serán de esa manera si la proposición es verdadera. Por el mero hecho de comprender una proposición, y antes de saber si es verdadera o falsa, hemos aprehendido una posibilidad. Wittgenstein repetirá sobre la proposición lo que ya había dicho acerca de las representaciones en general, a saber: la proposición determina un lugar en el espacio lógico, la posibilidad de que exista la situación representada en ella, pues «el lugar lógico coincide con el lugar geométrico en que ambos son la posibilidad de una existencia» (3.4-3.411).

Lo representado por una proposición es, en consecuencia, una situación o estado de cosas posible. Y esa situación aparece representada en la proposición por la configuración de los nombres que la componen. Por eso dice Wittgenstein que entendemos una proposición sin que nos expliquen su sentido, pues el sentido queda mostrado en la proposición (4.021 ss.), ya que no es otra cosa que su estructura. Lo que sí necesitamos que nos expliquen es la referencia de sus constitutivos, esto es, de los nombres (4.026), ya que es una relación entre ellos y la realidad, o más exactamente, entre ellos y los elementos simples de esta última. Es, por tanto, esencial a la proposición que pueda comunicarnos un nuevo sentido por medio de viejas expresiones (4.027 ss.), pues con los mismos nombres podemos formar diferentes sentidos, y esto significa en rigor «diferentes proposiciones», a base de combinar o estructurar los nombres de diferentes maneras (4.0311).

La aplicación del principio de isomorfía al lenguaje exige que éste pueda descomponerse finalmente en nombres. El lenguaje *toca* con la realidad a través de ellos. En cuanto figura, una proposición debe ser descomponible, analizable lógicamente, y en ella debe haber tantas partes distinguibles como en la situación que representa (4.0312-4.04). Esto no significa que sus distintas partes aparezcan gramaticalmente separadas; lo que se postula, igual que en Russell, es un análisis lógico. Por eso se encarga Wittgenstein de señalar que una proposición tan simple como *Ambulo* es, en realidad, compleja, siendo su sentido total el resultado, por lo pronto, de la combinación de una raíz con una determinada desinencia (4.032).

Hemos visto que el sentido de una proposición es su estructura, y que lo que representa es una situación o estado de cosas posible. Para que la proposición sea tal, y por tanto una figura, no es necesario que exista la situación representada; esto tan sólo es necesario para que la proposición sea verdadera. Decir que una proposición representa un estado de cosas posible equivale, por consiguiente, a decir que representa la existencia y no existencia de estados de cosas (4.1). Pero para representar la no existencia de un estado de cosas, la proposición debe estar negada, y la negación es una complejidad lógica añadida a la propia estructura figurativa de la proposición. La idea de Wittgenstein, claramente contraria a lo que, según hemos visto, había mantenido Russell, es que a una proposición negativa no puede corresponder ningún hecho peculiar de carácter negativo; un hecho negativo es, simplemente, un hecho inexistente (2.06). Las proposiciones elementales, son las que afirman la existencia de un estado de cosas (4.21), y ésta es la razón por la que dos proposiciones elementales no pueden ser entre sí contradictorias, pues sólo podrían serlo si una de ellas fuera negativa, en cuyo caso ya no sería elemental (4.211). Por la misma razón, tampoco se puede deducir una proposición elemental de otra (5.134), pues la deducción entre dos proposiciones solamente es posible si, al menos, una de ellas es compleja. Respecto a las proposiciones elementales, la lógica tan sólo puede decirnos que debe haberlas; dado el lenguaje, la existencia de proposiciones elementales tiene un fundamento *a priori*; es la lógica la que exige que las proposiciones complejas sean funciones veritativas de proposiciones elementales (principio de extensionalidad, 5.5562). Pero la lógica no puede decir qué proposiciones elementales hay, en qué consisten o cuáles son sus formas, y el intento de responder a estas cuestiones de modo *a priori* conducirá al sinsentido (5.55, 5.5571). Se trata de una cuestión empírica, ya que para obtener el conjunto de las proposiciones elementales habríamos de disponer del conjunto de los diversos nombres que pueden formar parte de ellas (5.55), pero contar con los nombres equivale a tener el conjunto de los objetos que son sus referentes, y esto es claramente una cuestión empírica que va ligada a nuestro progresivo conocimiento del mundo. Wittgenstein dirá que la decisión sobre qué proposiciones elementales hay es un problema que corresponde a la aplicación de la lógica, y que, por ello, la lógica no puede anticiparlo (5.557). De sus palabras puede inferirse sin dificultad que la aplicación de la lógica, en cuanto que consiste en poner las exigencias de la lógica en contacto con nuestro conocimiento del mundo, requiere dar entrada a elementos empíricos ajenos a la pura lógica. Pues para dar las formas de las proposiciones elementales tenemos que recurrir a los nombres, y para tener nombres hemos de tener algún conocimiento de sus referentes, y por consiguiente, de lo que hay en el mundo; pero «en la lógica no podemos decir: en el mundo hay esto y esto, pero no aquello» (5.61). Suena aquí, como puede apreciarse, un eco de la actitud que hemos visto en Russell: el lógico, como tal, no tiene por qué dar ejemplos. E incluso es posible que, cuando Wittgenstein afirma que el intento lógico

de decir qué proposiciones elementales hay conduce al sinsentido, esté criticando implícitamente de modo tácito los ejemplos de proposiciones atómicas suministrados por Russell. Es curioso, no obstante, que en una nota de 1915 recogida en los *Notebooks* (p. 61), Wittgenstein se muestra muy interesado en dar ejemplos de proposiciones elementales, atribuyendo a ineficacia propia la imposibilidad de darlos; dice así: «Ciertamente mi dificultad consiste en esto: en todas las proposiciones que se me ocurren aparecen nombres, los cuales deben desaparecer por un análisis ulterior. Sé que tal análisis es posible, pero no estoy en posición de llevarlo a cabo completamente. A pesar de ello, creo saber que, si el análisis se llevara hasta el final, su resultado tendría que ser una proposición que, de nuevo, tuviera nombres, relaciones, etc. En resumen: parece como si, de esta suerte, yo conociera una forma sin conocer de ella ningún ejemplo.»

En realidad, y como queda comprobado por todo lo anterior, únicamente a las proposiciones elementales les es aplicable el principio de isomorfía. Las proposiciones complejas contendrán, además de nombres, elementos a los que nada corresponde en la realidad, como, por ejemplo, los cuantificadores, diferentes partículas conectivas, etc. Un análisis de estas proposiciones complejas nos conducirá, obvia e inevitablemente, a proposiciones simples (4.221; éste es el supuesto básico del atomismo lógico). Y una proposición simple es, para Wittgenstein, una estructura o concatenación de nombres (4.22). Los símbolos simples son nombres, y las proposiciones elementales son funciones de nombres (4.24).

En síntesis, tenemos ya todo cuanto la lógica puede enseñarnos sobre el lenguaje. La cuestión ahora es: ¿qué puede mostrarnos el lenguaje acerca del mundo?

6.10 La estructura de la realidad

Abordamos ahora justamente las primeras páginas del *Tractatus*, donde Wittgenstein da sucinta cuenta de cómo tiene que ser la realidad para poder ser objeto de representación isomórfica.

Lo primero que encontramos es que el mundo es todo lo que acontece, esto es, el conjunto de los hechos (*Tatsachen*); el mundo, como tal, consiste y se divide en hechos, no en cosas (1-1.21). El acontecimiento, el hecho, es, a su vez, la existencia de estados de cosas (2). Un hecho es, por consiguiente, algo complejo, compuesto de estados de cosas existentes (cfr. 2.034). Puesto que un estado de cosas existente es lo que corresponde a una proposición elemental verdadera, cabría inferir que un hecho será lo que corresponda a una proposición compleja verdadera. La inferencia, sin embargo, no es correcta. Y no lo es por cuanto una proposición compleja debe contener algo más que nombres, según hemos visto, a saber: términos como «todos», «no», «si... entonces», etc.; pero nada puede haber en la realidad que corresponda a estos términos. Por consiguiente, un hecho es

es un conjunto de estados de cosas. En el caso más simple, un hecho será un estado de cosas; en el caso más complejo, que Wittgenstein sólo considera hipotéticamente, un hecho constará de infinitos estados de cosas (4.2211). Pero es importante darse cuenta de que la categoría de hecho en el *Tractatus* no es propiamente una categoría ontológica, pues no se aplica a ninguna entidad distinta de los estados de cosas. Dicho de otro modo: una reunión o conjunto de estados de cosas no es una nueva entidad con caracteres propios. La razón es que entre los estados de cosas no hay ninguna relación interna o necesaria: los estados de cosas son independientes entre sí (2.061), y de la existencia o inexistencia de uno de ellos no puede deducirse la de otro (2.062). Esto corresponde literalmente, como se habrá apreciado, a la tesis de que las proposiciones elementales son lógicamente independientes entre sí, que ya hemos comentado antes.

En la carta a Russell ya citada anteriormente, Wittgenstein intenta aclararle la diferencia que hay entre un hecho y un estado de cosas, pero su explicación, aunque se ha citado muchas veces, es poco rigurosa; dice así: «*Sachverhalt* (estado de cosas) es lo que corresponde a una proposición elemental si es verdadera. *Tatsache* (hecho) es lo que corresponde al producto lógico de las proposiciones elementales cuando este producto es verdadero. La razón por la que introduzco *Tatsache* antes de introducir *Sachverhalt* requeriría una larga explicación» (*Notebooks*, p. 129). Puesto que el producto lógico de proposiciones es la conjunción de las mismas, podría parecer que una conjunción de proposiciones es una representación isomorfa de un hecho. Por lo que se acaba de indicar, no es así; las proposiciones elementales representan isomórficamente estados de cosas, y un hecho no es otra cosa que un conjunto formado por n estados de cosas existentes ($n \geq 1$). Un hecho, como tal, no es isomórficamente representable, y sólo en un sentido derivado y no riguroso puede decirse que a él *corresponda* una conjunción de proposiciones elementales. De otro lado, decir que un estado de cosas es lo que corresponde a una proposición elemental si es verdadera es inexacto; esto sólo puede afirmarse de un estado de cosas *existente*. Pero Wittgenstein habla también en el *Tractatus* de estados de cosas inexistentes (2.05-2.062), que son los representados por las proposiciones elementales falsas, pues, como anteriormente hemos visto, es característico de Wittgenstein subrayar que también las proposiciones falsas tienen sentido. Es, por lo demás, comprensible que las afirmaciones de Wittgenstein en una carta no tengan el grado de precisión que tienen en su obra.

El mundo es, por lo tanto, el conjunto de los acontecimientos, de los hechos, o lo que es lo mismo, de los estados de cosas existentes (2.04; como ya se indicó anteriormente, Wittgenstein usa con frecuencia el término *Sachlage*, «situación», como sinónimo de *Sachverhalt*, «estado de cosas», por ejemplo, en 2.0122, 2.014, 2.11, 2.202, 2.203, 5.135 en relación con 2.062, etc.). Un estado de cosas, a su vez, es una combinación, relación o estructura de cosas u objetos (*Gegenständen*, *Sachen*, *Dingen*, 2.01; de estos términos, Wittgenstein utilizará con preferencia el primero).

Los objetos, que como ya vimos son los referentes de los nombres, son los elementos más simples de la realidad, de los que se componen las situaciones o estados de cosas. ¿En qué consisten los objetos? ¿Qué tipo de entidades son? En vano se buscarán ejemplos en el *Tractatus*. Es la misma dificultad que hemos encontrado a la hora de buscar ejemplos de proposiciones elementales. En una nota de 1915, Wittgenstein escribe con laconismo: «Nuestra dificultad era que hablábamos siempre de objetos simples y no sabíamos mencionar ni uno solo» (*Notebooks*, p. 68; es posible que el empleo del plural encierre una alusión a Russell).

Las propiedades que el *Tractatus* atribuye a los objetos clarifican la función que cumplen, pero no bastan para facilitarnos una representación de ellos. Se dice que son simples (2.02). Y es natural, puesto que corresponden a los elementos simples de las proposiciones, a los nombres. Si los objetos fueran compuestos no podrían ser nombrados, habrían de ser descritos, representados, y entonces serían sus partes componentes los constituyentes simples a los que se refirieran los nombres; esto es, bajo el supuesto de que es posible el análisis reductivo onto-lingüístico, supuesto que ya hemos visto operando en Russell, ha de haber objetos simples por definición (2.0201 ss.). Se dice, además, que los objetos son lo fijo, lo existente, por contraposición a su configuración, el estado de cosas, que es lo cambiante, lo variable (2.027-2.0272). Esta tesis es sumamente importante, ya que implica que la variabilidad de los acontecimientos del mundo consiste en la diversidad de las estructuras o relaciones que pueden darse entre los objetos, pero que por debajo de esta mutabilidad hay algo fijo e inmutable que son dichos objetos. Por eso afirma Wittgenstein que, por diferente que sea un mundo pensado respecto al mundo real, ha de tener algo en común con éste (2.022). ¿Qué? Simplemente una forma (*ibid.*). Podemos imaginar un mundo poblado por seres extraños, constituido de modo fantásticamente distinto, pero si podemos imaginarlo es porque tiene algo en común con el nuestro. Wittgenstein piensa que esta comunidad de todos los mundos posibles es una forma, una sustancia, constituida por los objetos (2.021, 2.023, 2.024). No me parece que el recurso a estos términos tenga en el *Tractatus* ningún sentido peculiar, y más bien creo que Wittgenstein los escogió por sus tradicionales connotaciones aristotélicas. Los objetos son la forma o sustancia de todo mundo posible porque son aquello que es necesario para que algo sea mundo. Un mundo es un determinado conjunto de relaciones entre los objetos. Relaciones distintas dan lugar a mundos diversos. Pero sean cuales fueren las relaciones hay algo inmutable y fijo que no difiere del mundo actual a cualquier mundo posible: los objetos. Por eso dirá Wittgenstein que la forma es la posibilidad de la estructura (2.033): pues la estructura es posible porque hay los objetos que la componen; o dicho de otra manera: los objetos contienen la posibilidad de todas las situaciones (2.014). La filosofía del *Tractatus*, como la de Russell, hunde sus raíces en la tradición occidental: se trata de salvar las apariencias de las cosas buscando lo necesario por debajo de lo contingente.

Wittgenstein no parece interesado en sacar especiales conclusiones con su terminología, y dar excesivo peso a sus palabras podría sumirnos en un pantano de irrelevancias. Así, puede uno sentirse tentado a elucubrar sobre el término «forma», en cuanto predicado de los objetos. Pero Wittgenstein se apresurará a conjugarlo con su opuesto, y dirá que «la sustancia es forma y contenido» (2.025). ¿Por qué llamarlo «sustancia»? Porque es lo que existe con independencia de los acontecimientos (2.024). Esto hay que entenderlo así: existen con independencia de cuáles sean en particular los acontecimientos en los que participan, pero no existen con independencia de cualquier acontecimiento en absoluto, y por sí mismos. Los objetos son independientes por cuanto pueden formar parte de todas las situaciones posibles, pero no son concebibles al margen de toda situación, de la misma manera que no tiene sentido concebir las palabras aisladas y al margen de las oraciones (2.0121-2.0123). Es esencial a los objetos poder formar parte de los estados de cosas, en el sentido de que es lógicamente necesario que los objetos aparezcan siempre relacionados entre sí; la propiedad que tienen los objetos de constituir situaciones o estados de cosas es, por ello, una propiedad que Wittgenstein llama, siguiendo la terminología de la época, «interna», esto es, no accidental, y que él considera como propiedad lógica o formal (2.011 ss., 2.01231, 4.122-4.123). El ámbito de todos los estados de cosas posibles constituye lo que, según vimos, el *Tractatus* denomina «espacio lógico» (1.13, 2.013, 2.11).

Hemos hablado hasta ahora del mundo, de los hechos, de los estados de cosas y de los objetos. Hay otro concepto más que interesa dilucidar. Las proposiciones elementales pueden ser verdaderas o falsas según representen estados de cosas existentes o inexistentes, pero sean lo uno o lo otro, y precisamente porque pueden serlo, son proposiciones con sentido, y esto significa que representan un estado de cosas que, sea existente o inexistente, es posible. El conjunto de los estados de cosas existentes constituye, según hemos visto, el mundo. Pues bien, esto más el conjunto de los estados de cosas inexistentes, pero posibles, es lo que Wittgenstein llama «realidad» (*Wirklichkeit*, 2.06). Puesto que los estados de cosas que existen, por existir, son *a fortiori* posibles, podemos decir que la realidad es el ámbito de lo posible, y que el mundo es una parte de lo anterior, la realidad realizada o actual. Esta terminología es coherente con las afirmaciones de Wittgenstein acerca de las representaciones y de las proposiciones, como cuando escribe que la proposición representa la realidad (4.01), que debe fijar la realidad positiva o negativamente (esto es, según represente una situación existente o inexistente, 4.023), o que la figura representa la realidad representando una posibilidad de existencia e inexistencia de estados de cosas (2.201). Las pruebas de este tipo pueden multiplicarse, no obstante lo cual muchos comentadores del *Tractatus* toman «mundo» y «realidad» como términos sinónimos. Solamente he advertido una afirmación de Wittgenstein que pueda justificar esta interpretación, la proposición 2.063, que dice: «la realidad total es el mundo». Creo que, salvo esta proposición, todo el resto de la doctrina del *Tractatus* es perfectamente coherente con

mi interpretación. De la proposición citada, lo único que se me ocurre pensar es que se trata de una pequeña incoherencia por parte de Wittgenstein. Téngase en cuenta que, cuando éste habla del *mundo*, afirma con toda claridad que es la totalidad de los estados de cosas existentes (2.04), y esto lo dice dos líneas antes de afirmar que la *realidad* son tanto los estados de cosas existentes como los inexistentes (2.06).

Tenemos, pues, que la estructura de la realidad, de acuerdo con la teoría del lenguaje que ya hemos estudiado, se analiza en el *Tractatus* por medio de las siguientes categorías:

Realidad: conjunto de todos los estados de cosas posibles (existentes o inexistentes). Corresponde al conjunto de todas las proposiciones elementales (verdaderas o falsas).

Mundo: conjunto de todos los estados de cosas existentes. Corresponde al conjunto de todas las proposiciones elementales verdaderas.

Estado de cosas (o situación): cualquier posible relación o configuración de elementos simples. Corresponde a la proposición elemental, que es una relación o configuración de nombres.

Hecho: conjunto de n estados de cosas existentes ($n \geq 1$).

Objetos (o cosas): elementos simples de los que se componen los estados de cosas. Corresponden a los nombres.

Objetos $\rightarrow$ realidad = posibilidad = estados de cosas	{	Existentes = hechos = mundo (<i>Tatsachen</i>) (<i>Welt</i>)
(<i>Gegenstände</i>) (<i>Wirklichkeit</i>) (<i>Möglichkeit</i>) (<i>Sachverhalte</i>)		Inexistentes

¿En qué coincide esta concepción de la realidad con la que hemos visto en el atomismo lógico de Russell, y en qué se distingue de ella? Russell había suministrado algunos ejemplos de proposiciones simples, y había sido mínimamente explícito respecto a su composición. Tales proposiciones constan, como ya sabemos, de nombres lógicamente propios, y de términos para propiedades y relaciones simples. Wittgenstein, en cambio, ha guardado silencio sobre este punto; incluso, como he subrayado más arriba, parece temer que el intento de dar ejemplos buscados *a priori* sólo conduzca al sinsentido. Su silencio, unido al hecho del todo evidente de que su doctrina no es más que una versión de la de su maestro, Russell, ha llevado a pensar a muchos que las proposiciones elementales de aquél sean como las proposiciones atómicas de éste, y que, en consecuencia, lo que Wittgenstein llama «nombres» sean indistintamente términos de particulares y términos de propiedades y relaciones. Si esta interpretación es correcta, entonces lo que en el *Tractatus* se llama «objetos» serán tanto entidades individuales, particulares en el sen-

tido de Russell, como propiedades y relaciones. Existe una breve anotación de 1915 que apoya esta interpretación; dice así: «Las relaciones y las propiedades, etc., son también *objetos*» (*Notebooks*, p. 61; a qué se refiera el «etc.» en esta cita lo ignoro).

En un cuidadoso trabajo, Irving Copi demostró hace tiempo («Objects Properties and Relations in the *Tractatus*», 1958), que la doctrina del *Tractatus* no puede ser ésa; por lo tanto, Wittgenstein debe de haber cambiado de opinión después de su anotación de 1915 (a menos que ésta tenga otra explicación; véase, por ejemplo, Griffin, *Wittgenstein's Logical Atomism*, p. 60). La argumentación de Copi ha sido criticada por Hintikka («Language-Games», 1976), pero a mí me sigue pareciendo del todo convincente. Dar aquí siquiera un resumen de ella sería excesivo para nuestros propósitos; en lo que sigue, argüiré por mi cuenta inspirándome en el artículo de Copi.

En primer lugar, se habrá advertido que Wittgenstein hace sobre los objetos afirmaciones muy parecidas a las que Russell hace acerca de los particulares. Así, Russell afirma —como referí en la sección 6.4— que los particulares son autosubsistentes, y los compara, desde este punto de vista, con las sustancias de nuestra tradición filosófica. De modo análogo, Wittgenstein se refiere a los objetos como lo que es fijo, y los denomina «sustancia» y «forma» del mundo. Es, sin duda, una primera indicación de que los objetos del *Tractatus* son entidades individuales como los particulares de Russell. No es, claro está, una prueba definitiva, pues bien podría ocurrir que Wittgenstein se hubiera apartado aquí de la posición de Russell. Una prueba definitiva, suponiendo que tal cosa sea posible respecto a un texto tan oscuro como lo es a veces el *Tractatus*, sólo puede encontrarse considerando con detenimiento lo que implican conjuntamente dos tesis básicas en él, el principio de la isomorfía entre el lenguaje y la realidad y la posibilidad de un análisis reductivo que conduzca a elementos simples.

Empecemos por las relaciones. Y admitamos la hipótesis de que también ellas, además de los particulares, sean objetos. Nos encontraremos entonces con que un estado de cosas puede estar formado por particulares y relaciones configurados entre sí de cierta manera, y esto equivale a decir que entre los particulares y las relaciones hay ulteriores relaciones, lo cual es ininteligible. Visto desde el lado del lenguaje, supone afirmar que los nombres que integran una proposición elemental incluyen nombres de relaciones; pero si esto es así, ¿para qué mantener que la proposición elemental es una configuración de nombres? Si las relaciones son también objetos y las proposiciones elementales incluyen nombres de relaciones, entonces el principio de representación isomórfica es inaplicable o bien hay que entenderlo metafóricamente, en contra del sentido aparente de las declaraciones del *Tractatus*. Pongamos un ejemplo. Sea el estado de cosas consistente en que el objeto *a* está encima del objeto *b*. Tal como hemos entendido la doctrina de Wittgenstein, una proposición que represente esa situación se reducirá, tras el oportuno análisis, a los nombres *a* y *b* dispuestos en determinada relación recíproca, tal vez así: $\frac{a}{b}$, o tal vez de esta

otra manera: a encima de b . Pero hay que tener en cuenta que, si la forma de la proposición elemental fuera esta última, la expresión «encima de» no sería un nombre de nada, sino tan sólo la forma de representación lingüística apropiada en castellano para ese hecho simple. En tal caso, habríamos de decir: la expresión «encima de» intercalada entre los nombres « a » y « b » representa que el objeto a está encima del objeto b . Y tal parece ser, en efecto, el sentido de una de las proposiciones más abstrusas del *Tractatus*: «No “El signo complejo ‘ aRb ’ dice que a está en la relación R con b ”, sino que ‘ a ’ está en cierta relación con ‘ b ’ dice que aRb » (3.1432). Si esto tiene algún sentido mínimamente claro es el de que « R » no es un nombre y, por tanto, el de que en una proposición analizada no pueden aparecer esos pseudonombres y, en consecuencia, que las relaciones no son objetos. Y adviértase que esto es muy distinto de lo que había dicho Russell; éste afirmaba simplemente que los hechos atómicos consisten en particulares más propiedades o relaciones simples, y que las proposiciones atómicas incluyen, por ello, nombres de particulares (pronombres demostrativos) y nombres de propiedades o relaciones simples. Russell hace, del principio de representación isomórfica, un uso diferente que Wittgenstein.

En cuanto a las propiedades, el *Tractatus* es más explícito. Hemos visto que los objetos constituyen la sustancia del mundo (2.021). Pues bien, en la proposición 2.0231 escribe Wittgenstein: «La sustancia del mundo sólo puede determinar una forma, y no propiedad material alguna. Pues las propiedades materiales son representadas sólo por las proposiciones, y se forman sólo por la configuración de los objetos.» Es decir: los objetos únicamente determinan la forma del mundo, o lo que es lo mismo (según vimos), las propiedades lógicas de lo real. Las propiedades materiales son el resultado de las relaciones o configuraciones de los objetos, y, en consecuencia, las relaciones no son nombradas, sino que se representan por medio de proposiciones. No puede estar más clara la diferencia entre objeto y propiedad material. Por eso añade Wittgenstein en la proposición siguiente: «Dicho sea de paso: los objetos carecen de color» (2.0232). Esto es: no sólo no es el color un objeto, sino que ni siquiera es una propiedad de objetos; es el resultado de una cierta configuración de objetos. El caso del color no parece que sea más que un ejemplo de propiedad material. De otra parte, puesto que las propiedades materiales o externas se contraponen en el *Tractatus* a las propiedades lógicas, formales o internas (cfr. 4.122 en relación con las proposiciones que se acaban de citar), parece claro que las primeras, las propiedades materiales, son las propiedades que podemos considerar empíricas o contingentes (como sugiere Copi). A estas consideraciones cabe añadir que, en un lugar en que Wittgenstein llama «objeto» a un color (4.123), se apresura a agregar entre paréntesis que esto constituye un uso vacilante de la palabra «objeto». Además, se puede mencionar también consideraciones de índole formal. Así, las afirmaciones, totalmente claras, de Wittgenstein, en el sentido de que, en una notación simbólica, los objetos están simbolizados por las variables, o por las constantes, de individuo (4.1211, 4.1272). Finalmente, cabe repetir el razona-

miento que se ha hecho a propósito de las relaciones. Si las propiedades fueran objetos, un estado de cosas podría consistir en una cierta relación entre particulares y propiedades, tesis que no me parece inteligible.

La conclusión es que lo que se denomina «objetos» en el *Tractatus* son entidades individuales como los particulares de Russell, esto es, entidades entre las cuales se dan relaciones. Lo peculiar de la doctrina de Wittgenstein consiste en que, para él, los hechos simples son siempre relaciones entre objetos; y por consiguiente, que lo que corrientemente llamamos «propiedades» (empíricas, no lógicas) son meramente el resultado de relaciones entre objetos, incluso aunque se trate de propiedades aparentemente simples como los colores; por lo tanto, las propiedades no son una categoría ontológica última, sino que se reducen a relaciones. Esto es, sin duda, del todo coherente con la concepción de Wittgenstein sobre lo real como conjunto de los mundos posibles, ya que, si las propiedades y las relaciones tuvieran asimismo la condición de objetos, y puesto que éstos son fijos e invariables para todos los mundos posibles, no habría respuesta para la pregunta: ¿qué es lo que varía de un mundo posible a otro? En la interpretación que he defendido, en cambio, la respuesta es clara: son fijas las entidades individuales, y podrían ser distintas sus relaciones, sus configuraciones y, en consecuencia, eso que llamamos «propiedades». Solamente con esta interpretación tiene sentido y aplicación el principio de representación isomórfica *.

En favor de esta interpretación hay, por último, una referencia muy clara en el escrito más representativo de la segunda filosofía de Wittgenstein. En las *Investigaciones filosóficas* (sección 46), después de reproducir un fragmento del *Teeteto* en el que Sócrates comenta que los elementos primarios carecen de toda determinación y no tienen nada más que el nombre, Wittgenstein afirma: «Estos elementos primarios eran tanto los individuos de Russell como los objetos del *Tractatus*.» Es comprensible, por ello, que se hayan comparado los objetos del *Tractatus* con las sustancias primeras de Aristóteles (así, Urmson, *El análisis filosófico*, cap. 5, secc. A), pero la semejanza no debe ocultar una importante diferencia: una propiedad, según el *Tractatus*, no se predica de un objeto, sino que se reduce a relaciones entre objetos.

* En un interesante trabajo, titulado «Use and Reference of Names», Hidé Ishiguro ha defendido la tesis de que los objetos del *Tractatus* no son entidades individuales carentes de propiedades, sino que son ejemplificaciones de propiedades no materiales, añadiendo que Wittgenstein no nos informa sobre qué tipo de propiedades serían éstas. Entrar aquí en los detalles de la intrincada y meritoria argumentación de Ishiguro sería demasiado largo para nuestros propósitos; sólo diré que no me ha convencido, porque me parece más oscura y abstrusa que las propias afirmaciones de Wittgenstein. La propia tesis defendida resulta, en mi opinión, menos inteligible aún que la interpretación de Copi que he desarrollado. Esta tiene en su favor no sólo razones internas a la doctrina de Wittgenstein que considero fácilmente comprensibles, sino además la virtud de situar el *Tractatus* en el contexto de la teoría de Russell, de donde en efecto procede.

Así pues, el mundo se compone de estados de cosas, que son relaciones entre objetos. Estos estados de cosas corresponden a los hechos atómicos de Russell, aunque, como hemos comprobado, no tienen la misma estructura. Pero Russell se vio obligado —como se recordará— a ampliar los tipos de hechos básicos a causa de la aplicación del principio de extensionalidad. Puesto que el principio es igualmente aceptado por Wittgenstein, ¿acepta éste las mismas consecuencias ontológicas?

Como ya insinué al tratar de Russell, Wittgenstein elude una ampliación de la ontología al modo de su maestro y reduce los hechos básicos a los estados de cosas. Su actitud frente a los hechos generales, negativos y de existencia deriva de su doctrina sobre las constantes lógicas y de su distinción entre lo que una proposición dice y lo que muestra. En la carta a Russell, que ya hemos citado varias veces, Wittgenstein considera la siguiente afirmación de este último: «Es necesario también que esté dada la proposición de que todas las proposiciones elementales están dadas»; y comenta: «Esto no es necesario porque es, incluso, imposible. ¡No hay tal proposición! Que todas las proposiciones elementales están dadas se *muestra* en que no hay ninguna que tenga un sentido elemental y que no esté dada» (*Notebooks*, p. 130). Trasladado esto al plano de los hechos significa que la generalidad no es un rasgo ontológico, que no es nada que pueda representarse, porque no es un hecho o un estado de cosas ulterior, que haya que añadir cuando tenemos todos los hechos o todos los estados de cosas, o todos los objetos. Dicho de otro modo: los cuantificadores, como cualesquiera otras constantes lógicas, no representan nada de la realidad ni hacen referencia a ella. Por eso dice Wittgenstein, un poco más arriba en su carta, que lo que se pretende expresar por medio de la *aparente* proposición «Hay dos cosas», se *muestra* en que hay dos nombres con significado distinto (esto es: que se refieren a objetos diferentes). ¿Qué quiere decir esto? Que el que haya tantos o cuantos objetos no es un estado de cosas y, por tanto no es algo que se pueda representar por medio del lenguaje, como tampoco se puede representar lingüísticamente el que haya objetos. Esto no es parte del mundo, de lo que ocurre, sino más bien el presupuesto para que haya mundo, para que algo acontezca, para que se den estados de cosas. Como tal presupuesto, el lenguaje no puede hablar acerca de ello, pero puede mostrarlo. Que hay objetos se muestra en que hay nombres. Y cuántos objetos haya se mostrará en el número de nombres que tengamos. Que hay dos cosas se reconocerá en cualquier proposición que exprese una relación entre dos nombres, por ejemplo, en cualquier proposición de la forma « aRb », siendo a y b dos nombres. O si empleamos variables, nuestra proposición tendrá la forma, en la notación de Wittgenstein, « $(\exists x, y) xRy$ » (4.1272). Ahora entendemos por qué llama Wittgenstein «aparente» a la proposición «Hay dos cosas»: porque esta expresión equivale a escribir « $(\exists x, y)$ » sin añadir nada más, es decir, sin predicar nada de los objetos denotados por las variables. Por lo mismo, tan poco sentido tiene decir que hay objetos, como que hay cien objetos, como hablar de la totalidad de los objetos (*ibid.*).

La idea, por consiguiente, es que lo que los cuantificadores indican lo deben mostrar las proposiciones analizadas, pero sin decirlo. Esta es, por cierto, la razón por la que, entre las proposiciones analizadas, no puede haber afirmaciones de identidad, ya que una afirmación de identidad, si no es tautológica, esto es, de la forma « $a = a$ », supone tener dos nombres para el mismo objeto, pues esto es lo que expresa una proposición de la forma « $a = b$ », pero en un lenguaje analizado cada objeto únicamente debe tener un nombre y cada nombre debe nombrar un objeto distinto (4.241 ss.). Dicho de manera más elegante: la identidad de objeto se expresará por medio de la identidad del nombre, no por medio del signo de identidad (5.53 ss.). El principio general es este: lo que las proposiciones muestran, no pueden decirlo las proposiciones. La cuantificación no es más que una operación de abreviatura que nos evita la repetición de los casos particulares. En lugar de repetir la atribución de mortalidad para cada uno de los seres humanos, decimos «Todos los hombres son mortales». En vez de ir pasando revista a diferentes cosas, para decir «O la nieve es blanca, o la hierba es blanca, o el azufre es blanco, o el yeso es blanco...», resumimos así: «Algunas cosas son blancas», o bien: «Existen cosas blancas». Pero ni el cuantificador universal ni el cuantificador particular expresan rasgo alguno de la realidad: no hay ni hechos generales ni hechos existenciales. Si existen objetos, y cuáles existan, se mostrará en los nombres que aparezcan en nuestras proposiciones simples. Que éstos sean todos los objetos, se mostrará en que no aparecen más nombres. Por eso dice Wittgenstein: «Si los objetos están dados, por lo mismo nos están dados todos los objetos. Si las proposiciones elementales están dadas, por lo mismo están dadas todas las proposiciones elementales» (5.524). Russell se había visto obligado a aceptar hechos generales y existenciales por asociar íntimamente los cuantificadores con las funciones veritativas. Wittgenstein mantendrá una actitud opuesta: «Yo separo el concepto *todo* de la función veritativa» (5.521). Y subrayará algo que Russell parece pasar por alto —como señalé en su momento—, y es la interdefinibilidad de los cuantificadores por medio de la negación, que hace redundante aceptar ambos cuantificadores como primitivos (5.441). En resumen: la cuantificación universal abrevia una conjunción de proposiciones elementales, y la cuantificación existencial abrevia una disyunción también de proposiciones elementales, pero no añaden nada nuevo a ambas funciones veritativas. Que esta tesis sea aplicable en la práctica, exige, naturalmente, que el alcance de la cuantificación sea un universo finito, pues de otro modo nunca tendríamos una conjunción, o una disyunción, completa. Por esto ha comentado Kenny que la teoría del *Tractatus* sobre la cuantificación parece requerir un «axioma de finitud» (Wittgenstein, cap. 5, p. 92 de la edición de Penguin).

Los hechos existenciales y los hechos generales quedan, de esta manera, excluidos de la ontología del *Tractatus*. Con mayor razón, los hechos negativos. La negación no es más que una operación veritativa, es decir, una operación que produce una función veritativa a partir de una proposición

elemental (5.3). Nada corresponde en la realidad al signo de negación, por cuanto lo mismo dicen una proposición y su negación (4.0621). «Sócrates vive» *dice* lo mismo que «Sócrates no vive»; el hecho que corresponda a la primera de estas proposiciones corresponde igualmente a la segunda. La negación no sirve para caracterizar el sentido de una proposición, y la prueba es que la doble negación se anula y nos deja con la proposición original: $\text{no-no-}p = p$ (*ibid.*). Esto confirma que no hay nada en la realidad que corresponda al «no» (5.44). ¿En qué consiste, entonces, la verdad de «Sócrates no vive», según el ejemplo de Russell? Esta proposición es verdadera, de acuerdo con las reglas veritativo-funcionales de la negación, porque «Sócrates vive» es falsa. Y esta última proposición es falsa porque representa un hecho inexistente; al principio del *Tractatus*, Wittgenstein había indicado que «también llamamos a la inexistencia de los estados de cosas, hecho negativo» (2.06). Dicho de otra manera: la ilusión russelliana de que a «Sócrates no vive» ha de corresponder un hecho negativo procede de que el hecho que corresponde a «Sócrates vive» es un hecho inexistente. Esto es lo que crea la apariencia de que es la negación la que pone a una proposición falsa de acuerdo con la realidad (5.512).

De la ontología ampliada de Russell quedan tan sólo por considerar los hechos de actitudes proposicionales, o hechos mentales, cuyo paradigma es la creencia. También esta clase de hechos queda excluida del ámbito de lo real por Wittgenstein. Al contrario de lo que pensaba Russell, las proposiciones como «*A* cree que *p*» no pueden construirse como si expresaran una relación entre una proposición, *p*, y un cierto objeto, en este caso, una persona, *A* (5.541). Wittgenstein enuncia su posición de forma escueta y más bien oscura: las proposiciones como «*A* cree que *p*», «*A* piensa *p*» y «*A* dice *p*», son de la forma «*p*' dice *p*», lo que significa que tales proposiciones no correlacionan un hecho con un objeto, sino que coordinan un hecho con otro hecho por medio de la coordinación de sus objetos (5.542). ¿Qué significa esto? Empecemos por el último caso: «*A* dice *p*». Cuando alguien pronuncia una proposición, ¿qué ocurre? Que emite, en una ocasión determinada, una serie de sonidos que constituyen una actualización de esa proposición considerada como tipo o entidad abstracta. Esto es: se da un signo acontecimiento que corresponde, por definición, al signo tipo que actualiza. Pero salvo la diferencia que hay entre el carácter abstracto del signo tipo y el carácter concreto del signo acontecimiento (recuérdense las secciones 2.2 y 3.1), nada hay que diferencie la proposición *p* de la emisión lingüística de *p* en una ocasión determinada por un hablante *A*. La pronunciación de *p* está, por eso, coordinada con la proposición *p*, y puesto que se trata de dos hechos, la pronunciación de la proposición y la proposición como tal, son dos hechos los que quedan coordinados (recuérdese que, para Wittgenstein, toda representación es un hecho, y por tanto la proposición y el pensamiento lo son asimismo). Vayamos ahora al otro caso: «*A* cree que *p*» y «*A* piensa *p*», afirmaciones entre las cuales sólo hay diferencias de matiz irrelevantes para nuestro problema. Lo que estas afirmaciones hacen, según Wittgenstein, es correlacio-

nar el pensamiento o creencia de *A* con la proposición *p*. ¿Pero en qué consiste el pensamiento o la creencia de *p*? En la propia proposición *p* pero sin palabras; en una representación o figura isomorfa de *p* mas carente de signos, puesto que, como vimos en su momento, la proposición es solamente el pensamiento exteriorizado por medio de signos. Por consiguiente, lo que se correlaciona en este caso es el hecho de la proposición tipo *p* con el hecho de la representación mental de *p*. En resumen: mientras que «*A* dice *p*» significa que los sonidos producidos por *A* corresponden (isomórficamente) a la proposición *p*, en cambio «*A* piensa *p*» (o «*A* cree que *p*») significa que la representación mental de *A* corresponde (isomórficamente) a la proposición *p*. Por tanto, esas afirmaciones no expresan una relación entre una proposición y un objeto, sino una correlación (y por ello, una relación isomórfica) entre dos hechos, el hecho de la proposición y el hecho de su pronunciación o el hecho de su representación mental, respectivamente. Lo cual supone, naturalmente, excluir la consideración del sujeto *A* como objeto simple; Wittgenstein no se encogerá ante esta consecuencia: «Esto muestra también que el alma, el sujeto, etc., tal como lo concibe la superficial psicología actual, es un absurdo. Ciertamente un alma compuesta ya no sería alma» (5.5421). A saber, lo que se rechaza es la consideración del sujeto como objeto simple (*cfr. Notebooks*, p. 118).

Los hechos quedan, pues, reducidos a estados de cosas, y se rechazan todas las razones que Russell había suministrado en favor de otros tipos de hechos, lo que da como resultado una ontología mucho más sobria y unitaria. Visto desde la perspectiva lingüística, significa que solamente tienen sentido aquellas proposiciones que puedan descomponerse en proposiciones elementales, o lo que es lo mismo, en configuraciones de nombres. No hace falta pensar mucho para reconocer que, con esta teoría del significado (llamémosla así, todavía), quedan situadas más allá del sentido proposiciones de muy diverso tipo. ¿Cuáles?

6.11 De lo que no puede hablarse

Wittgenstein considera en el *Tractatus* diferentes clases de lo que, desde el punto de vista del principio de representación isomórfica, hay que llamar «pseudoproposiciones», esto es, oraciones que carecen de sentido, que no dicen nada, y que, en definitiva, constituyen un intento de hablar de lo que no puede hablarse. Veámoslas.

1. En primer lugar, las pseudoproposiciones lógicas

Hemos visto que la forma lógica es lo que toda representación ha de tener en común con la realidad representada para poder representarla. Pues bien, las proposiciones, aun cuando pueden representar la realidad entera, esto es, la totalidad de los estados de cosas posibles, no pueden representar lo que han de tener en común con éstos, la forma lógica; la ra-

zón es que, para poder representar esta última, las proposiciones habrían de estar fuera de la lógica, y por consiguiente, fuera del mundo (4.12). ¿Pero cómo sabemos que las proposiciones participan, junto con los hechos, en la forma lógica, y por qué tenemos conocimiento de ella? Porque la forma lógica se refleja en las proposiciones; éstas la expresan, la exhiben, la muestran (4.121). Queda así introducido, de manera explícita, el concepto de mostrar, que Wittgenstein contrapone, aquí y en otros lugares, al concepto de decir: «lo que se puede mostrar no se puede decir» (4.1212); «lo que se expresa en el lenguaje, no podemos expresarlo por medio de él» (4.121). Las proposiciones cumplen, pues, dos funciones semánticas distintas: decir y mostrar. Y el principio de representación isomórfica o figurativa se aplica exclusivamente al decir.

Las proposiciones no pueden representar su forma lógica, sino que la muestran. Y, como ya vimos, la forma lógica incluye aquellas propiedades que podemos llamar, indistintamente, «formales», «lógicas», «estructurales» o «internas» (4.122). Que hay estas propiedades no puede afirmarse por medio de proposiciones con sentido, pero se muestra en ellas mismas (*ibidem*). Wittgenstein suministra un ejemplo de ello: que dos proposiciones sean contradictorias o que una se siga de la otra, lo mostrarán sus estructuras (4.1211), pero ni ellas, ni otras proposiciones cualesquiera, pueden decirlo (4.124, 4.125). Se advertirá que esta doctrina implica la imposibilidad de dar sentido a las proposiciones de la filosofía de la lógica, esto es, a aquellas proposiciones que tratan de las propiedades lógicas del lenguaje y, en su caso, del mundo.

La función de mostrar las propiedades lógicas del lenguaje da lugar a un tipo peculiar de verdades y falsedades, las verdades y falsedades lógicas. Una verdad lógica es una proposición compleja que es verdadera cualquiera que sea el valor veritativo de las proposiciones elementales componentes; Wittgenstein dirá que, en tal caso, las condiciones de verdad son tautológicas (4.46). Una falsedad lógica es una proposición compleja que es falsa sea cual fuera el valor de verdad de las proposiciones elementales que la componen; en este caso, las condiciones de verdad son contradictorias (*ibidem*). A las verdades lógicas, Wittgenstein les llamará «tautologías»; a las falsedades, «contradicciones» (*ibidem*).

Es oportuna, en este momento, una aclaración histórica. El término «tautología», que Wittgenstein parece haber sido el primero en usar con el sentido indicado, se ha conservado en la filosofía de la lógica hasta hoy, pero con un alcance más restringido. Actualmente es normal llamar «tautologías» sólo a las verdades lógicas de la parte más elemental de la lógica, de la lógica de proposiciones. La razón es que únicamente las verdades de la lógica de proposiciones son, según acuerdo unánime, verdades en virtud de los modos de composición veritativo-funcional, pero es disputable que esto sea cierto también de la lógica de predicados. Por esta razón, se considera que no todas las verdades lógicas son tautologías, sino que éstas constituyen tan sólo un subconjunto de aquéllas (puede verse, sobre esto, las secciones 10 y 16 de la *Lógica matemática* de Quine). Wittgenstein,

no obstante, llamaba «tautologías» también a las verdades de la lógica cuantificacional por cuanto, como se indicó anteriormente, mantenía que una proposición universal equivale a una conjunción de proposiciones elementales, y que una proposición cuantificada particularmente equivale a una disyunción de proposiciones elementales.

Pues bien, las tautologías y las contradicciones no dicen nada, no son representaciones isomorfas de la realidad, no representan ninguna situación posible (4.461, 4.462). Las tautologías carecen de condiciones de verdad, porque son compatibles con cualquier situación posible, y en consecuencia dejan abierta la totalidad del espacio lógico; las contradicciones, por el contrario, son incompatibles con toda situación posible, y cierran el espacio lógico sin dejar sitio a la realidad (4.461-4.463). Las tautologías y las contradicciones, dirá Wittgenstein con un pequeño juego de palabras, son carentes de sentido (*sinnlos*), pero no son sinsentidos (*unsinnig*) (4.461-4.4611). Esta pequeña distinción, que casi parece retórica, es muy significativa y suministra una pista clara para entender la concepción de Wittgenstein acerca de las proposiciones lógicas. Estas, sean tautologías o contradicciones, carecen de sentido puesto que carecen de condiciones de verdad y no cumplen con las exigencias del principio de isomorfía. Pero no por eso son absurdas o constituyen sinsentidos, y la razón, que Wittgenstein a continuación añade, es que tanto unas como otras pertenecen al simbolismo de la lógica (4.4611), y esto les da una función rigurosa, consistente en mostrar de modo patente las propiedades lógicas del lenguaje y del mundo (6.12). De igual manera que el «0» es parte del simbolismo aritmético, y en esta medida no es un signo absurdo (la comparación la añade Wittgenstein, 4.4611).

Lo anterior nos introduce en un tema metafísico íntimamente ligado a la lógica: la necesidad. A la luz de lo que hemos visto se comprenderá que Wittgenstein diga que la verdad de una tautología es cierta, la de una contradicción, imposible, y la de una proposición (se entiende: ni tautológica ni contradictoria), posible (4.464). Esto equivale a decir que fuera de la lógica todo es azar (6.3), y por tanto que no hay más necesidad que la necesidad lógica, ni más imposibilidad que la imposibilidad lógica (6.37, 6.375). No hay leyes científicas ni metafísicas, si por tal se entiende proposiciones necesariamente verdaderas. Pues lo único que hay en la ciencia son hipótesis (6.31, 6.341, 6.36311), y los principios metafísicos de la naturaleza, como el principio de causalidad, o el principio de razón suficiente, son simplemente intuiciones *a priori* sobre la forma posible de las proposiciones científicas (6.32, 6.34), y no sobre lo que éstas describen; versan sobre nuestro aparato conceptual, y no sobre el mundo (6.35 s.).

Las proposiciones de la lógica, tautologías y contradicciones, carecen, pues, de contenido, no representan situación alguna, y valen tan sólo por lo que muestran de las propiedades formales del lenguaje y del mundo. De aquí que la lógica sea anterior a la experiencia de que las cosas son de tal o cual manera, anterior a lo que sucede en el mundo, pero no anterior a que haya cosas o a que suceda algo; la lógica es anterior al cómo pero no

al qué (5.552). No exige que ocurra algo en particular, pero requiere que ocurra algo, que haya estados de cosas, relaciones entre objetos, y representaciones de éstas, como el lenguaje. La lógica no es una doctrina, puesto que no contiene afirmaciones sobre el mundo, sino un reflejo del mundo como en un espejo (5.511, 6.13). La lógica llena el mundo, y los límites del mundo son también sus límites (5.61). Lo que se expresa igualmente diciendo que la lógica es trascendental (6.13).

2. En segundo lugar, las pseudoproposiciones filosóficas

Lo que se acaba de señalar sobre la necesidad y los principios metafísicos constituye un primer ejemplo de las consecuencias que Wittgenstein está dispuesto a extraer de su doctrina sobre el sentido de las proposiciones. A decir verdad, estas implicaciones estaban ya declaradas en el prólogo del *Tractatus*: «Este libro trata de los problemas filosóficos, y muestra, según creo, que el planteamiento de estos problemas se debe al mal entendimiento de la lógica de nuestro lenguaje». La razón de este diagnóstico es la teoría de la proposición, que ya conocemos. En su escrito más antiguo, al menos de los que están publicados, ya en sus primeras líneas, Wittgenstein escribe: «La filosofía no suministra representaciones de la realidad» («Notas sobre la lógica», 1913). Por consiguiente, las proposiciones filosóficas no tienen sentido, son pseudoproposiciones, como dará a entender el *Tractatus*.

En efecto, cuando Wittgenstein se refiere aquí a la filosofía, lo hace después de haber desarrollado su doctrina de la proposición y contraponiendo la filosofía a la ciencia. El tema lo inicia una de las proposiciones más importantes y menos citadas del *Tractatus*; es también una de las más claras, la proposición 4.11, que dice: «La totalidad de las proposiciones verdaderas es la totalidad de las ciencias de la naturaleza». Y la proposición siguiente, añade: «La filosofía no es una de las ciencias de la naturaleza. (La palabra 'filosofía' debe referirse a algo que está por encima o por debajo de las ciencias de la naturaleza, pero no junto a ellas.)» (4.111).

Es claro, pues, que no hay proposiciones verdaderas que no sean científicas. Lo único que hay que advertir es que debe tomarse el término «ciencia de la naturaleza» con cierta amplitud, de manera que incluya, de una parte, también las ciencias humanas, en cuanto que describen determinados aspectos de lo que ocurre en el mundo; y de otra, asimismo, las observaciones cotidianas de lo que acontece, que en principio podrían parecernos demasiado triviales para formar parte del contenido de la ciencia, pero que, en una consideración más detenida, habrán de aparecer precisamente como las proposiciones que incorporan un conocimiento más inmediato y directo del mundo. Las ciencias de la naturaleza se contraponen, en el *Tractatus*, por un lado a la filosofía, y por el otro, a las ciencias formales, como la lógica y la matemática. La lógica, ya hemos visto que no es una doctrina, y que sus proposiciones carecen, en rigor, de sentido. En cuanto a la matemática, Wittgenstein piensa, acaso influido por el logi-

cismo de Russell, que no es más que un método lógico, y por ello, que sus proposiciones, en cuanto que son ecuaciones, son pseudoproposiciones, que no expresan ningún pensamiento, y que únicamente son útiles en su aplicación a procesos de inferencia entre proposiciones no matemáticas (6.2-6.211, 6.234).

Así formulado el problema, la filosofía, ciertamente, no es una ciencia natural. Su propósito no es representar lo que acontece. La verdad filosófica, como tal, aspira a estar más allá de la experiencia. Y esto es justamente lo que la hace cuestionable para Wittgenstein. No se trata de que las proposiciones filosóficas sean, en su mayoría, falsas, ni de que nos hallemos lejos de haber alcanzado verdades filosóficas. Se trata, más bien, de que la mayor parte de las proposiciones de los filósofos son sinsentidos (y nótese que Wittgenstein recurre aquí al término *unsinnig*, y no meramente al término, más neutro, más desprovisto de carga valorativa, *sinnlos*, que hemos visto utilizado para calificar las pseudoproposiciones de la lógica; 4.003). Las cuestiones filosóficas no son cuestiones que se pueda intentar responder; lo único que puede hacerse es establecer que son sinsentidos, originados en nuestro inal entendimiento de la lógica del lenguaje (*ibidem*). De aquí que la filosofía se convierta en una actividad; una actividad de clarificación. Hay que aclarar nuestros pensamientos, que de otra forma resultan confusos. La filosofía, por ello, no produce como resultado proposiciones, sino la clarificación de las proposiciones (4.112). ¿Cómo se lleva a cabo esta tarea clarificatoria? Poniendo límites a lo que se puede pensar, y por lo mismo, a lo que no puede pensarse. Más explícitamente: pensando lo que puede ser pensado hasta llegar a sus límites, que serán los que lo separen de lo que no puede ser pensado (4.114). Representando claramente lo que puede decirse, la filosofía se refiere, negativamente, por así decirlo, a lo indecible (4.115). En este sentido, la filosofía es crítica del lenguaje (4.0031). Pero, además, puesto que todas las proposiciones verdaderas son científicas, la delimitación del sentido viene a equivaler a una delimitación del ámbito de las ciencias de la naturaleza (4.113). De esta suerte, la filosofía se convierte en una especie de filosofía de la ciencia, si bien de magro contenido.

Naturalmente, hay que excluir del ámbito de la filosofía aquellas porciones residuales que correspondan a ciencias que se hallan en una etapa inicial de desarrollo. En estos años del *Tractatus*, la psicología se estaba constituyendo como ciencia, pero la psicología filosófica coleaba aún en muchas direcciones especulativas. Por eso, Wittgenstein se ocupa de aclarar que la psicología no está más relacionada con la filosofía que cualquier otra ciencia de la naturaleza (4.1121), y añade que la teoría del conocimiento es la filosofía de la psicología. Esto implica, claramente, que la teoría del conocimiento es rechazable como discurso con sentido, y que el estudio del conocimiento sólo puede ser objeto de la psicología científica. Por lo mismo, determinadas hipótesis sobre la naturaleza que, por su alcance general, puedan interesar a los filósofos o haber sido consideradas alguna

vez como filosóficas, tal como la teoría darwiniana, no tienen con la filosofía mayor relación que cualquier otra hipótesis científica (4.1122).

Se habrá observado que Wittgenstein no atribuye a las pseudoproposiciones filosóficas ninguna función que salve su utilidad o su validez, al contrario de lo que hace con las pseudoproposiciones de la lógica. Hay que exceptuar, no obstante, cierta clase particular de proposiciones filosóficas, las que tengan como propósito cumplir la tarea aclaratoria y delimitadora a que ha quedado reducida la filosofía. Tal es el caso de las propias proposiciones del *Tractatus*, pero volveremos sobre ello más adelante.

Nótese, asimismo, las siguientes implicaciones de lo que acabamos de estudiar. Vimos, en su momento, cómo los hechos atómicos para Russell consistían en datos sensibles, y en qué medida operaba aquí un supuesto epistemológico al margen de la teoría del lenguaje. Dado el contexto russelliano al que pertenece el *Tractatus*, cabría pensar, y muchos lo han dicho, que los estados de cosas de los que ahí se habla son también datos sensibles. Si se tiene en cuenta, sin embargo, que la totalidad de las proposiciones verdaderas es la totalidad de las proposiciones científicas, resulta difícil mantener lo anterior, pues la reducción de lo dado a datos sensibles implica una privatización del lenguaje ajena a la intersubjetividad del lenguaje científico. Fue esto, precisamente, lo que, pocos años después, condujo a Carnap a abandonar un intento parecido al de Russell en favor del cientifismo que caracterizó a los miembros del Círculo de Viena. Curiosamente, y si mi interpretación del *Tractatus* en este punto es correcta, Wittgenstein se encontraría mucho más cerca de los neopositivistas que de Russell en materia epistemológica. Y esto explica, por cierto, que aquéllos tuvieran tanta facilidad para extraer del *Tractatus* el principio de verificabilidad y fundar sobre aquél el empirismo lógico. Pero también sobre este punto retornaremos después.

3. La clase de pseudoproposiciones que vamos a considerar ahora no son una especie peculiar, sino tan sólo una aplicación de la tesis anterior que descalifica todas las aparentes proposiciones genuinamente filosóficas. Se trata de las que versan sobre el sujeto metafísico, sobre el sujeto trascendental, y en particular las que pretenden afirmar el solipsismo.

Es significativo que las afirmaciones de Wittgenstein sobre este tema sucedan a aquella proposición, que ya hemos comentado, en la que se dice que los límites del mundo son también los límites de la lógica (5.61). Esta proposición termina con estas palabras: «Lo que no podemos pensar, no podemos pensarlo; lo que no podemos pensar, tampoco podemos decirlo.» Esto, según agrega Wittgenstein a continuación, suministra la clave para decidir en qué medida sea verdad el solipsismo. La respuesta es: «lo que el solipsismo quiere decir es del todo correcto, sólo que no se puede decir, sino que se muestra a sí mismo» (5.62). ¿Cuál es la tesis solipsista, tal como la entiende Wittgenstein? Esta: el mundo es mi mundo. De aquí sacará el solipsista toda suerte de implicaciones sobre la imposibilidad de su comunicación con los demás o acerca del alcance de su conocimiento. La tesis,

según Wittgenstein, es correcta. Únicamente que no puede expresarse por medio del lenguaje, aunque sí cabe reconocerla en cuanto que se muestra en éste. ¿Por qué no se puede decir la tesis solipsista? Porque no representa ningún hecho, actual o posible, y por tanto no cumple con los requisitos del principio de representación isomórfica que ha de cumplir toda proposición para tener sentido. No es una tesis que describa hechos o estados de cosas, sino que es una afirmación acerca del mundo en su totalidad, y por tanto, en cierto sentido, debe estar más allá del mundo. ¿Por qué es correcta la tesis solipsista? Porque el mundo, por definición, lo encuentra cada cual en torno a sí, es circunstancia, en el estricto sentido orteguiano, por consiguiente, es el mundo de cada cual. Por eso añadirá Wittgenstein a continuación: «el mundo y la vida son lo mismo» (5.621). Esto, en rigor, es inexacto, como cualquier buen orteguiano podría demostrar, pero expresa bien hacia dónde discurre el pensamiento de Wittgenstein. ¿De qué manera se muestra a sí misma la tesis solipsista? A las preguntas anteriores hemos contestado nosotros en el espíritu del *Tractatus*. A esta última responde Wittgenstein expresamente: «Que el mundo es mi mundo se muestra en que los límites del lenguaje (del único lenguaje que yo entiendo) se refieren a los límites de mi mundo» (5.62).

La proposición anterior expresa también la razón última de que el solipsismo sea correcto. La tesis solipsista se muestra, aparece, en una característica formal, lógica (en el sentido del término que es propio del *Tractatus* y que ya hemos estudiado), del lenguaje, a saber, en que sus límites se refieren, remiten, a los límites de mi mundo. Hasta aquí, sin embargo, no hemos salido del solipsismo. Tan sólo hemos constatado que no es algo de lo que se pueda hablar, pero que es correcto por cuanto se muestra en cierta característica del lenguaje. ¿Es esto todo? Si así fuera, cabría acusar al autor del *Tractatus* de solipsista. Muchos lo han hecho, y entre otras razones, aparte de las que se acaban de esbozar, citaban el paréntesis de la proposición anterior. ¿Por qué? Porque el texto alemán es ahí ambiguo; dice: *der Sprache, die allein ich verstehe*, y esto se había traducido en muchas ocasiones, ya en la primera traducción inglesa, y por ello, también en la traducción castellana, como «del lenguaje que yo sólo entiendo». Con ello, Wittgenstein estaría aceptando que el lenguaje sólo es inteligible para cada sujeto, y su teoría caería, como la de Russell, en la privatización del lenguaje y en el solipsismo lingüístico. Hay, sin embargo, un testimonio fiable de que la significación que Wittgenstein daba a ese paréntesis corresponde a la traducción que he ofrecido más arriba: «del único lenguaje que yo entiendo», pues cuenta Lewy que Wittgenstein mismo hizo la oportuna corrección en un ejemplar de la primera traducción inglesa (Lewy, «A Note on the Text of the *Tractatus*», 1967; se encontrará una detenida discusión de este tema en García Suárez, *La lógica de la experiencia: Wittgenstein y el problema del lenguaje privado*, cap. II). Ahora bien, ¿cuál es el único lenguaje que yo entiendo? El mismo que entienden los demás: un lenguaje que comparte con la realidad la forma lógica y cuyas proposiciones significativas son representaciones isomórficas de hechos posibles. Por

consiguiente, un lenguaje intersubjetivo, y no privado. Y esto tiene como consecuencia que *mi* mundo, en cuanto representado por las proposiciones de mi lenguaje, coincidirá con el mundo de los demás por su estructura, ya que ésta es idéntica para todos los mundos posibles. ¿Cómo, si no, explicar que las proposiciones verdaderas sean las proposiciones científicas? La teoría de la representación del *Tractatus*, y la teoría de la realidad que acompaña a aquella, hacen imposible cualquier consideración del lenguaje y del mundo como algo privado y propio del sujeto; o más rigurosamente: muestran que tal consideración está vacía de contenido, pues es correcto afirmar que el mundo es mi mundo, pero la corrección de este aserto está relacionada con su falta de contenido.

Una cuestión es si el lenguaje es o no privado según Wittgenstein. He intentado explicar brevemente por qué no lo es. Otra cuestión, distinta, es si el lenguaje del *Tractatus*, esto es, las proposiciones de Wittgenstein en su obra, constituyen un lenguaje privado en la intención del autor. Es importante distinguir entre ambas cuestiones cuando se habla del problema del lenguaje privado en el *Tractatus*; las discusiones corrientes de este tema las mezclan y confunden con frecuencia. Pues bien, a la segunda cuestión mi respuesta es igualmente negativa. Mi argumento es que las proposiciones del *Tractatus* pretenden ser inteligibles para los demás, por cuanto no constituyen sino un ejemplo de proposiciones filosóficas aclaratorias, válidas, pues, por lo que muestran, y, como ya hemos visto, Wittgenstein acepta las pseudoproposiciones filosóficas cuando se limitan a ejecutar esa tarea delimitadora y elucidatoria. Considerar el lenguaje del *Tractatus* como un lenguaje privado me parece incompatible con la asignación al mismo de dicha función, que Wittgenstein hace al final de su obra, como comprobaremos en seguida.

Al margen del carácter lógico e intersubjetivo del lenguaje, hay otra línea de argumentos con la que Wittgenstein acaba por escapar del solipsismo, el cual, por cierto, debió de ser una fuerte tentación que, por razón de su personalidad, sufrió toda su vida (como parcialmente lo confirma Moore en «Conferencias de Wittgenstein en 1930-33», pp. 356-7). El solipsismo supone un yo, el yo propio, coordinado con el mundo. Pero ya consideramos antes, al hablar de las proposiciones de creencia, que el yo no es parte del mundo, no es un objeto simple, ni tampoco un peculiar hecho o estado de cosas. Wittgenstein es ahora todavía más explícito sobre este punto: «El sujeto pensante, el sujeto de representaciones, no existe» (5.631). Si escribiéramos un libro sobre el mundo tal y como lo hemos encontrado —continúa discuriendo Wittgenstein— habríamos de informar sobre nuestro cuerpo, sobre cuáles de sus partes están, y cuáles no, sujetas a la voluntad, etc.; «éste es un método de aislar el sujeto, o más bien, de mostrar que, en un importante sentido, no hay sujeto, pues solamente de él no podría tratar este libro» (*ibidem*). La argumentación de este párrafo está solamente insinuada, pero es clara: la descripción más completa del mundo no puede incluir el sujeto; por más que intente aislarlo nunca llegará a él; por mucho que hable sobre el cuerpo, sobre la

mente, sobre la voluntad, sobre las representaciones, los recuerdos y los sentimientos, siempre quedará detrás el sujeto al que atribuimos voluntad, pensamiento o sentimientos. Pues, ¿en qué lugar del mundo podemos localizar un sujeto metafísico? (5.633): «el sujeto no pertenece al mundo» (5.632).

Es curioso recordar que, por esta época, exactamente en 1914, Ortega había hecho una crítica presentando una concepción del sujeto trascendental como intimidad y ejecutividad que implica la imposibilidad de someterlo a ciertos manejos teóricos como la reducción husserliana, y probablemente también la imposibilidad de describir e, incluso, como en Wittgenstein, de hablar de él (Ortega, «Ensayo de estética a manera de prólogo»; *cfr.* Marías, Ortega. I. *Circunstancia y vocación*, secc. 72). Como para Wittgenstein, pero por razones diversas, tampoco para Ortega el sujeto es parte del mundo, y a esto se vincula la imposibilidad de hacer de él un objeto de representación. En una anotación de 1916, Wittgenstein escribió: «El yo no es un objeto. Yo me encuentro objetivamente frente a todo objeto, pero no frente al yo.» (*Notebooks*, p. 80). En el *Tractatus* lo comparará con el ojo, que, en cuanto órgano de la visión, no pertenece al campo visual, pero es condición necesaria para la existencia de éste (5.633 s.).

El yo del solipsista es condición para que haya mundo, pero no forma parte del mundo; es como un punto inextenso que coordina la realidad (5.64). Pero lo único que tenemos es la propia realidad; por eso concluye Wittgenstein: «el solipsismo, seguido estrictamente, coincide con el realismo puro» (*ibidem*). Este es el final de la supuesta aventura solipsista de Wittgenstein en el *Tractatus*. Puesto que el solipsismo exige un sujeto metafísico, y éste no es más que la condición para que haya realidad, es la realidad lo único que al fin podemos representarnos, y el solipsismo ha de dejar su sitio al realismo.

Del yo filosófico, o trascendental, no puede hablarse. Aquello de lo que puede hablarse, el hombre, el cuerpo, el alma de la que trata la psicología (esto es, la psique, la conciencia, etc.), no son el sujeto metafísico. Aquello es parte del mundo, es, en última instancia, reducible a estados de cosas, a relaciones. El sujeto filosófico no es parte del mundo, sino el límite del mundo (5.632, 5.641); y esto quiere decir: es supuesto y condición necesaria para que haya mundo («presupuesto» dice en los *Notebooks*, p. 79). Por eso es un sujeto trascendental. Hemos visto antes que Wittgenstein afirma de la lógica que es trascendental en cuanto que sus límites coinciden con los límites del mundo. Pues bien, el sujeto es trascendental por cuanto es límite del mundo. La terminología del *Tractatus* es aquí del todo coherente, excepto en que Wittgenstein debería haber hablado más bien de la realidad que del mundo, puesto que el límite en el que se sitúan tanto la lógica como el sujeto es, no tanto el límite entre lo que acontece y lo que no acontece, cuanto el límite entre lo que puede ocurrir y lo imposible. De esta manera, hay un sentido en el que puede haber un discurso filosófico, y no meramente psicológico, acerca del sujeto (5.641): el sentido en el cual proposiciones filosóficas de carácter elucidatorio, como las del

Tractatus, pueden mostrar de qué modo se muestra a sí mismo el sujeto en el lenguaje, y mostrar al tiempo que, por ello mismo, el solipsismo es correcto pero vacío.

4. Después de exponer su concepción de la lógica, y en particular su doctrina de las funciones veritativas, Wittgenstein hace unas consideraciones sobre el tema de la necesidad y sobre los principios metafísicos, especialmente el principio de causalidad. Lo hemos visto brevemente un poco más arriba, en el apartado 2 de esta sección. Pues bien, dicho tema le conduce a la cuestión de la relación entre el mundo y la voluntad, y con ello, al problema de las proposiciones éticas, estéticas y religiosas. Esta será la última parte del *Tractatus*, lo que podríamos llamar su axiología negativa.

De manera un tanto confundente, Wittgenstein empieza por hablar del valor de las proposiciones: «todas las proposiciones valen lo mismo» (6.4). Esto es: como descripciones de hechos posibles, todos los cuales son igualmente contingentes y entre los cuales no existe preeminencia alguna, no hay jerarquía ni diferencias de valor entre las proposiciones. ¿Qué proposiciones podrían parecernos más valiosas que las demás, y por tanto, superiores? De un lado, aquellas cuya verdad fuera necesaria. Pero ya hemos comprobado que las únicas proposiciones necesarias no son propiamente proposiciones, puesto que no dicen nada, sino pseudoproposiciones: son las tautologías. De otro lado, aquellas que declararan el sentido de los hechos, y que, por ello, estarían por encima de las meras descripciones. Pero, por definición, cualquier intento de expresar el sentido del mundo por medio del lenguaje debe infringir los requisitos del principio de isomorfía, pues o bien el sentido de los hechos es parte del mundo, en cuyo caso será un hecho más entre los hechos, y no se ve de qué modo pueda dar sentido a los demás hechos, o bien el sentido está fuera del mundo, y entonces el lenguaje no puede hablar de él. Esta última es la alternativa que Wittgenstein desarrolla: en el mundo todo es como es y ocurre como ocurre, por consiguiente, no hay en él ningún valor, porque si lo hubiera, sólo por esto no tendría valor (6.41; esto último es una forma paradójica de decir que considerar el valor como parte del mundo equivale a convertirlo en hecho y despojarlo de su condición de valor).

Una primera consecuencia de ello es que no puede haber proposiciones de ética (6.42 s.). ¿Por qué? Porque las proposiciones de ética no describen hecho alguno, sino que pretenden declarar, en un cierto aspecto, el sentido del mundo, al menos en lo que éste está conectado con la voluntad humana. Pero las proposiciones no pueden expresar nada que esté más alto (6.42) que el nivel de los hechos. ¿Significa esto que no existan en absoluto los valores éticos o que no tenga el mundo un sentido moral? Para entender todo lo que implica esta pregunta, así como lo que vamos a ver a continuación, conviene recordar que, para Wittgenstein, el mundo y la vida son lo mismo (5.621), y aunque esta última proposición está en el *Tractatus* bastante alejada de las que tratan sobre ética, es curioso notar

que una afirmación idéntica a esa se encuentra en los *Notebooks* (p. 77) junto con las que hablan de ética, y por cierto que aquí se cuida Wittgenstein de añadir que la vida de la que habla no es ni la vida fisiológica ni la vida psicológica. Pues bien, la respuesta a la pregunta anterior está ya indicada en este mismo lugar de los *Notebooks*: «la ética no trata del mundo; la ética debe ser una condición del mundo, como la lógica». En el *Tractatus* lo expresa de otra manera: «la ética es trascendental» (6.421). Como la lógica, como el sujeto metafísico, la ética es también una condición, un presupuesto, del mundo. Esto quiere decir, por lo que sabemos, que la ética está, también, en el límite del mundo. O dicho de otro modo: el valor moral no es parte del mundo, pero tampoco está más allá del mundo; es una condición necesaria para que haya mundo, para que haya vida. No hay mundo sin valores morales, como no hay mundo sin lógica o sin sujeto. Como la lógica, como el sujeto, la ética es igualmente trascendental.

Que no pueda haber mundo ni vida humana sin valores morales acaso sea discutible, pero es, curiosamente, algo en lo que las concepciones metafísicas de la ética coincidirían plenamente con Wittgenstein. Y algunas de ellas coincidirían, además, con la extensión que a seguidas hace Wittgenstein de su doctrina a la estética: «ética y estética son lo mismo» (6.421). Uno recuerda, casi involuntariamente, la doctrina escolástica según la cual lo bello es casi una especie de lo bueno. La identidad entre ética y estética aparece ya en el lugar indicado de los *Notebooks* (p. 77), pero además se encuentra en esta obra alguna pista por lo demás ausente en el *Tractatus*; así, dice: «La obra de arte es el objeto visto *sub specie aeternitatis*, y la vida buena es el mundo visto *sub specie aeternitatis*. Esta es la conexión entre el arte y la ética. La manera usual de considerar los objetos es verlos desde en medio de ellos; la consideración *sub specie aeternitatis*, desde fuera. De tal modo que tienen como trasfondo el mundo entero» (*Notebooks*, p. 83). El argumento, a primera vista, no es muy claro. Es patente la voluntad de salvar las cosas, los hechos, de su contingencia, digámoslo así. Usualmente estamos entre ellos; la ética y la estética representan un intento de sacarlos de su contexto mundano, espacio-temporal, contingente, y darles una validez eterna. La diferencia sería que la estética considera los objetos aisladamente (aquí el término «objeto» no parece tener el sentido técnico que recibe en la metafísica del *Tractatus*), mientras que la ética toma en consideración el mundo en su conjunto. Teniendo en cuenta que Wittgenstein, como hemos visto, ha equiparado el mundo a la vida, esto puede tal vez resultar más claro de lo que parece. La estética salva las cosas; la ética salva la vida humana. Al dar valor estético a una cosa transmutándola en obra de arte, al dar valor ético a la vida sometiénola a principios morales, salvamos una y otra colocándolas en una perspectiva de eternidad. La ética y la estética cumplen la misma función, y en esta medida, son idénticas. Ni el valor ético ni el valor estético están dentro del mundo, pero tampoco se hallan fuera de él; son condición

del mundo, y por ello, están en su límite. Esto significa que la ética y la estética son trascendentales.

Hay otro argumento por el cual puede arribarse a tal conclusión. Wittgenstein lo esboza sólo a propósito de la ética, y de forma más completa en los *Notebooks* (p. 79): «El bien y el mal aparecen únicamente con el sujeto. Y el sujeto no pertenece al mundo, sino que es un límite del mundo.» Es decir (y ampliándolo a la estética): los valores suponen el sujeto y aparecen sólo con él. Pero puesto que el sujeto es un límite del mundo (es trascendental), cuanto se refiere a los valores pertenece igualmente al límite del mundo (es trascendental). A la luz de esta consideración se entienden mejor las afirmaciones del *Tractatus* sobre las leyes éticas, sobre el premio y el castigo inoral, y sobre la voluntad (6.422-6.43). En resumen: no puede hablarse de la voluntad moral (tan sólo de la voluntad como objeto de la psicología científica), y la buena o mala voluntad no puede alterar los hechos, y por tanto: no puede alterar nada de lo que es posible expresar por medio del lenguaje, sino que, si puede modificar algo, será los límites del mundo (6.43). ¿Qué quiere decir esto? Que modificará el sentido que el mundo en su conjunto adquiera para el sujeto, de la misma forma que, para el hombre feliz, el mismo mundo es diferente que para el infeliz (*ibidem*).

De nada de lo que da sentido a la vida puede tratar el lenguaje. Ni la experiencia de la muerte, que no es propiamente una experiencia, puesto que no es parte de la vida (6.4311), ni de la vida eterna, que, en la medida en que se conciba como intemporal, no podemos representarnos, y la cual, en todo caso, es tan problemática como la vida presente (*ibidem* y 6.4312). Esta reflexión viene a dar en el tema religioso, o como Wittgenstein lo llama, «lo Místico», empezando por una concisa expresión de lo que podemos llamar «la teodicea negativa del *Tractatus*»: «Dios no se revela en el mundo» (6.432). Y por tanto el conocimiento del mundo no contribuye a otorgar sentido al mundo: «lo místico no consiste en cómo sea el mundo, sino en que es» (6.44). Y esto equivale a sentir el mundo como un todo limitado, que es tanto como contemplarlo *sub specie aeterni* (6.45; esta última expresión aproxima estrechamente lo místico a la ética, como era de esperar). Pero lo místico, que Wittgenstein parece tomar como epítome y compendio de lo que no puede expresarse, se muestra a sí mismo: «Hay ciertamente lo inexpressable. Esto se muestra a sí mismo, es lo místico» (6.522). ¿Dónde se muestra? Por lo que acabamos de ver, podemos pensar que se muestra en el sentimiento del mundo como un todo limitado. Así pues, en el lenguaje se muestra la estructura lógica del mundo y el sujeto trascendental; lo místico, en cambio, se muestra en un sentimiento. Aquí, el lenguaje no puede cumplir ni siquiera una tarea ostensiva. ¿Y la ética y la estética? Wittgenstein ha dicho que son trascendentales, pero no ha llegado a afirmar que se muestren aquí o allá. Cabe pensar que, coincidiendo con lo místico en ser formas de ver el mundo en una perspectiva de eternidad, se muestran igualmente en el sentimiento de finitud.

Sobre todo ello, pues, no se puede hablar. Mas esto no significa meramente que no puedan resolverse ciertos problemas o contestarse determinadas preguntas. Es que tampoco el problema o la pregunta tienen sentido en cuanto expresados lingüísticamente, porque «si puede en general hacerse una pregunta, también es posible responderla» (6.5). Por consiguiente, desde un punto de vista lingüístico, no hay problema (*ibidem*). La duda sólo puede existir cuando existe una pregunta, y ésta cuando hay una respuesta, y ésta, a su vez, cuando se puede decir algo (6.51). De aquí que el escepticismo, como actitud última, sea, no irrefutable, sino más bien un sinsentido, pues pretende dudar donde no es posible formular preguntas (*ibidem*). No hay más preguntas con sentido que las preguntas científicas; es cierto que cuando éstas han sido contestadas, los problemas vitales (el término es de Wittgenstein) están sin tocar, pero como estos problemas no constituyen propiamente preguntas, ya no hay más preguntas, y ésta es la respuesta (6.52). La solución del problema de la vida consiste en la disolución del problema (6.521). Wittgenstein piensa que ésta es acaso la razón por la que, quienes han encontrado el sentido de la vida, no han sido capaces de decir en qué consistía (*ibidem*).

Con esto, el *Tractatus* llega a su fin. Antes de poner la última palabra, Wittgenstein se cuidará de hacer una consideración sobre el método filosófico que sirve de marco para explicar qué función cumplen las pseudoproposiciones de su libro. El método correcto en filosofía —dice Wittgenstein— sería no decir nada excepto lo que se puede decir, a saber, las proposiciones de la ciencia natural, y por tanto, algo que no tiene nada que ver con la filosofía; y cuando alguien quisiera decir algo metafísico, demostrarle que no había asignado referencia a algunos de los signos de sus proposiciones. Este método sería insatisfactorio para el otro —concede Wittgenstein— pero sería el único estrictamente correcto (6.53).

En realidad, el método filosófico así concebido es aún más restringido de lo que Wittgenstein da a entender en otro lugar de su obra y de lo que en ella practica él mismo. El modo de enunciar la proposición anterior parece sugerir un ideal: ese sería el método estrictamente correcto, a saber, limitarse a demostrar que las proposiciones no científicas carecen de sentido a causa de la falta de referencia de algunos de sus signos. En tal medida, el método que antes ha expuesto en el *Tractatus*, no es el rigurosamente correcto, pero tiene la ventaja de ser más didáctico. Consiste, como se recordará, en mostrar, por medio del lenguaje, cuáles son los límites de éste y por qué no puede decirse lo que no se puede decir. Y esto es lo que hace el *Tractatus*, como explica Wittgenstein finalmente: sus proposiciones aclaran en cuanto que, quien lo entienda, las reconoce al fin como sinsentidos (6.54). Y aquí recurre Wittgenstein al término más fuerte, *unsinnig*. Las pseudoproposiciones del *Tractatus* no son simplemente carentes de sentido (*sinnlos*), como las proposiciones de la lógica; son auténticos sinsentidos, como las de la metafísica, pero se distinguen de éstas en que cumplen una tarea aclaratoria, sirven como los peldaños de una escalera, para subir por ella, y tras llegar arriba tirarla, porque no hay viaje de

vuelta. Quien entienda el *Tractatus*, «debe superar estas proposiciones, y ver entonces el mundo correctamente» (*ibidem*).

La última expresión del libro contiene una vuelta implícita al método riguroso de la filosofía, y suena, en consecuencia, casi como una crítica al propio *Tractatus*: «Sobre lo que no se puede hablar, se debe guardar silencio» (7). Como ya advertía Wittgenstein en el prólogo, todo el sentido del libro se puede resumir en esa afirmación más esta otra: lo que se puede decir, se puede decir claramente (4.116).

6.12 Balance del *Tractatus*

En esta sección pasaremos revista crítica a algunos de los principales puntos, supuestos y consecuencias del *Tractatus*, seleccionados, como es natural, en función de nuestros intereses y de la exposición anterior. Con ello, damos fin, no sólo al estudio de la doctrina del *Tractatus*, sino también al de toda una primera época de la filosofía analítica del lenguaje, caracterizada por este vano intento de hallar las condiciones, lógicas y semánticas, de un lenguaje perfecto. Dividiremos la presente sección en varios puntos.

1. El lenguaje.

El *Tractatus* habla poco del lenguaje; casi siempre habla de la proposición. Esto no tiene mayor alcance si tenemos en cuenta que «la totalidad de las proposiciones es el lenguaje» (4.001). Pero en todo caso, ¿de qué habla realmente: del lenguaje y de las proposiciones que todos usamos en nuestras lenguas, o de un lenguaje y de unas proposiciones pergeñados por Wittgenstein para servir como paradigma y patrón al lenguaje real? Vimos al comienzo de la sección 6.4 que es más bien esto último lo que pretendía Russell: exhibir en qué consiste un lenguaje lógicamente perfecto para que su idea nos sirva de contraste con el que detectar las imperfecciones lógicas de nuestras lenguas, y de este modo, evitar la caída en una errónea concepción del mundo. Puesto que el proyecto de Wittgenstein es, en otros aspectos, y como en su momento he subrayado, idéntico al de Russell, ¿hay que atribuirle, asimismo, la comunión en semejante idea de la filosofía del lenguaje? Muchos lo han mantenido; a mi juicio, con error. Pienso que, en este punto, Wittgenstein se aparta de Russell.

En primer lugar, no se encuentra en el *Tractatus* ninguna condena general del lenguaje ordinario como puede hallarse en Russell. Por el contrario, precisamente en el lugar en que Wittgenstein está escribiendo sobre la necesidad lógica de que haya proposiciones elementales, afirma: «Todas las proposiciones de nuestro lenguaje ordinario están de hecho, tal y como están, lógicamente bien ordenadas» (5.5563). Y por si fuera poco, añade: «Esa simpleza que aquí hemos de exponer, no es una imagen de la verdad, sino la verdad misma entera. (Nuestros problemas no son abstrac-

tos, sino acaso los más concretos que existen)» (*Ibidem*) *. ¿Cómo puede compaginarse semejante evaluación lógica del lenguaje ordinario con una doctrina de la proposición tan esotérica como la exigida por el principio de isomorfía? Distinguiendo entre la apariencia del lenguaje cotidiano y su estructura lógica subyacente. A primera vista no se percibe que las proposiciones sean concatenaciones de nombres; pero éste es el resultado al que se llega tras el análisis pertinente. Hay que traspasar la superficie multiforme, abigarrada y confusa de las expresiones lingüísticas usuales, para descubrir el orden lógico soterrado bajo ellas. Este orden lógico consiste en la relación de nombres que representa isomórficamente un estado de cosas; tal orden lógico corresponde, igualmente, al pensamiento contenido en la expresión lingüística. Lo que ocurre es que no siempre el pensamiento está patente en la proposición: «El lenguaje disfraza el pensamiento. Y tanto es así que, de la forma exterior del vestido, no se puede inferir la forma del pensamiento encubierto; porque la forma exterior del vestido ha sido diseñada para fines enteramente distintos que los de revelar la forma del cuerpo. Las convenciones tácitas para comprender el lenguaje ordinario son enormemente complicadas» (4.002). Dicho de otra forma: la gramática ordinaria no coincide con la gramática lógica, y más bien tiende a ocultarla, aunque puede ser incluso más compleja que ésta. «El lenguaje ordinario es parte del organismo humano y no menos complicado que éste. Es humanamente imposible captar directamente a partir de él la lógica del lenguaje» (4.002). Esto es lo que justifica la tarea del análisis: éste nos suministra el método para llegar a la lógica del lenguaje (y del mundo), que las expresiones lingüísticas ordinarias encubren y disimulan. Los hombres han creado diferentes lenguas en las que expresar todo hecho posible, pero esto no implica que sepan cuál es la referencia de cada palabra o signo simple, por lo mismo que muchos hombres saben hablar sin que sepan de qué manera se producen los sonidos (4.002).

Esta divergencia entre la gramática ordinaria y la gramática lógica es la culpable de varios de los defectos lógicos del lenguaje ordinario, como la ambigüedad, ya denunciada por Russell. Así, una misma palabra puede designar de modos distintos, como cuando se usa bien como nombre propio bien como predicado, lo que muestra que, aunque se trate del mismo signo, estamos ante dos símbolos distintos (4.323; recuérdese, si bien esto lo mencionamos en la sección 2.1 hablando de los signos, que para Wittgenstein el signo es lo que puede ser percibido de un símbolo, y el símbolo es el signo más el modo de designación, 3.32 s. y 3.326). Otro ejemplo,

* En una carta de 1922, comentando esta proposición, Wittgenstein escribe: «Con esto quería decir que las proposiciones de nuestro lenguaje ordinario no son, en modo alguno, menos correctas, o menos exactas, o más confusas, que las proposiciones escritas en, por ejemplo, el simbolismo de Russell, o en cualquier otro 'Begriffsschrift'. (Sólo que para nosotros es más fácil captar su forma lógica cuando están expresadas en un simbolismo apropiado.)» (*Letters to Ogden*, p. 50; agradezco a la profesora María Teresa Iglesias haberme llamado la atención sobre esta cita.) El contraste con la posición de Russell no puede resultar más evidente.

mencionado por Wittgenstein, es el de la palabra «es», que constituye tres símbolos diferentes según se use como cópula, como expresión de identidad o como indicación de existencia (3.323). Justamente para evitar estos defectos, y los errores que en ellos pueden alimentarse, recomienda Wittgenstein, en una vena semejante a la de Frege y Russell, utilizar un lenguaje de signos en el que cada signo sea no más que un símbolo, y en que símbolos distintos no se usen de modo aparentemente idéntico; esto es, «un lenguaje de signos que obedezca a la gramática lógica, a la sintaxis lógica» (3.325). Y añade entre paréntesis: la notación conceptual de Frege y de Russell son un lenguaje así, aunque no consigan eliminar todos los errores. Esta alusión, sin embargo, puede malentenderse. Podría parecer que el *Tractatus* trata de suministrar otro lenguaje de signos, otro cálculo lógico o escritura conceptual (esta última expresión es el título del primer libro de Frege) que mejore los de sus antecesores. Pero nada hay en el *Tractatus* que permita tomarlo así. No es, ciertamente, un lenguaje de signos, esto es, una notación, como lo son el *Begriffsschrift* de Frege o los *Principia Mathematica* de Russell y Whitehead. Es, más bien, una teoría filosófica sobre el lenguaje, y sobre la relación que liga a éste con la lógica, por un lado, y con la realidad, por otro. Que Wittgenstein reconozca la utilidad de las notaciones de Frege y Russell no quiere decir que su teoría de la proposición sea una teoría acerca de las expresiones de cualquiera de esos lenguajes formales. Justamente ésta es la diferencia, acaso sutil, que lo separa de Russell. Para éste, la teoría atomista del lenguaje es la doctrina del lenguaje lógicamente perfecto, un lenguaje que podría estar compuesto, como hemos visto, por una sintaxis como la de los *Principia* más un vocabulario de datos sensibles, y por ello fundamentalmente privado; para Wittgenstein, en cambio, la teoría atomista es, claramente, una concepción sobre el lenguaje ordinario, que intenta alcanzar, por debajo de su forma externa, la forma lógica que hace posible su poder representativo, su capacidad semántica. Es patente que se trata de una concepción trascendental del lenguaje. No cabe llamarla, en rigor, «teoría», por cuanto las proposiciones en las que se expresa carecen de sentido.

Esta importante particularidad de la posición de Wittgenstein no debe ocultar las muchas coincidencias que, por lo demás, tiene con Russell, e incluso con Frege. Por lo pronto, la reducción del discurso filosóficamente relevante al discurso descriptivo. Únicamente tienen sentido las oraciones que pueden ser empíricamente verdaderas o falsas; sólo ellas son propiamente proposiciones. Como afirmación sobre el lenguaje lógicamente perfecto, es unilateral y tan sólo resulta comprensible en cuanto vinculada a un cierto estadio del desarrollo de la lógica, un estadio en el que no se han elaborado todavía la lógica de los imperativos, de los juicios de deber o de las interrogaciones. Como afirmación acerca del lenguaje ordinario es totalmente arbitraria.

Su justificación teórica es la aplicación del principio de isomorfía, principio que ya Russell, para el lenguaje perfecto, había insinuado. ¿Pero sabemos cómo aplicar el principio a la relación entre el lenguaje y la reali-

dad? El principio se apoya, sin duda, en una aceptación de la función referencial como modo básico de significar, esto es, de relacionarse el lenguaje con la realidad extralingüística. La idea es que a cada elemento simple de la realidad le corresponderá un signo lingüístico simple. Y así, para explicar que el lenguaje con un conjunto finito de signos pueda expresar cualquier hecho posible, Wittgenstein tendrá que estipular que la realidad se compone de un conjunto finito de objetos, y que todo hecho posible está formado por objetos pertenecientes a este conjunto. Pero no sabemos ni en qué consisten dichos objetos ni cuáles son los signos simples de nuestro lenguaje. Por consiguiente, el principio de isomorfía no pasa de ser una exigencia de aplicación práctica imposible. Tanto el atomismo lingüístico como el atomismo ontológico (las dos caras del atomismo lógico) quedan sumidos en la oscuridad más impenetrable. Porque supongamos que el principio de isomorfía no hubiéra ido acompañado de la doctrina atomista; entonces podría haberse estipulado, por ejemplo, que los objetos sean todos aquellos agregados físicos que posean una determinada continuidad espaciotemporal mínima, y que los nombres que constituyen la proposición sean todas aquellas palabras que pueden referirse a dichos objetos. Aquí tendríamos una base para intentar aplicar el principio de isomorfía y comprobar si una aparente proposición posee o no sentido; con este método, las proposiciones científicas quedarían legitimadas y las metafísicas eliminadas, como en el *Tractatus*, pero sabríamos claramente por qué. Y esto es lo que nunca sabremos mientras nos atengamos a las afirmaciones de Wittgenstein. Esta deficiencia ontolingüística de la construcción wittgensteiniana va a quedar compensada, curiosamente, con la parte menos desarrollada de su doctrina, la epistemología.

Aparte de la falta de identificación separada de los objetos y de los nombres, que, en mi entender, hace imposible la aplicación del principio de isomorfía, hay que preguntarse si tal principio puede realmente explicar cómo tienen sentido las proposiciones descriptivas. Tomado en su literalidad, y llevado hasta el final, el principio establece que a cada posible relación entre objetos ha de corresponder una relación entre sus nombres. Pero me parece evidente que esto impone sobre el lenguaje exigencias por completo ajenas a él. Como vimos, Wittgenstein sugiere, a fin de comprender la esencia de la proposición, reparar en lo que significa la escritura jeroglífica (4.016). ¿Pero es que puede tomarse esta forma de escritura como representativa y característica de lo que es el lenguaje humano? A tal escritura acaso pueda aplicarse el principio de isomorfía; incluso para nuestra escritura alfabética podríamos tener convenciones tales como representar que un objeto se halla sobre otro escribiendo el nombre del primero encima del nombre del segundo, y otras reglas así. Pero con esto, no iríamos muy lejos. La capacidad expresiva (de todo tipo, y no sólo representativa) que posee el lenguaje, su creatividad semántica (recuérdese la sección 3.4), obedece, sin duda, a que las normas que regulan su relación con la realidad son de un tipo mucho más flexible que las propias de una relación de isomorfía. Lo característico del lenguaje no es expresar que un

objeto está sobre otro escribiendo una palabra encima de otra; lo peculiar del lenguaje consiste en expresar esa relación entre objetos no por medio de una relación entre palabras, sino por medio de palabras, alguna de las cuales, como la expresión «está encima de», corresponde a la relación que se quiere representar. Y se comprende que sea así puesto que la forma primaria del lenguaje humano es oral, y aquí la posibilidad de aplicar el principio de isomorfía es prácticamente nula. ¿Cómo podríamos representar isomórficamente en el lenguaje hablado que un objeto está sobre otro? ¿Haciendo una pausa de duración determinada entre los nombres de los objetos? ¿Pronunciando el nombre del primer objeto con tonalidad más alta que el del segundo? En mi opinión, la aplicación del principio de isomorfía a la teoría del significado procede de una indebida asimilación del lenguaje a las representaciones figurativas claramente pictográficas como los jeroglíficos, las pinturas, las fotografías, las imágenes ópticas, los diagramas, etc. Pero la capacidad semántica del lenguaje humano ciertamente excede a la mera capacidad de representación pictográfica. Aquí, Russell estaba en terreno más sólido; al limitar su doctrina a un lenguaje perfecto, se colocaba a cubierto de posibles contraejemplos tomados del lenguaje ordinario, aunque la consecuencia, que Russell aceptaba con presteza, fuera la de hacer de su lenguaje perfecto un instrumento perfectamente inútil para la comunicación.

2. Lo posible y lo verdadero.

¿De qué habla el lenguaje? De la realidad; o lo que es lo mismo, de lo posible. Las proposiciones tienen sentido en la medida en que representan isomórficamente posibles estados de cosas. ¿En qué consiste un estado de cosas desde el punto de vista de nuestro conocimiento? Algunos han dicho que a Wittgenstein no le interesa llegar a lo más básico en el conocimiento sino a lo más básico en la referencia. Ignoro cómo puedan separarse ambas cosas. Otros han pensado que, por razón del contexto histórico, deben tomarse los estados de cosas como datos de los sentidos, al estilo de Russell. Lo cierto es que las consideraciones de Wittgenstein sobre asuntos epistemológicos son escasas, tan escasas que hacen pensar en una deliberada voluntad de evitarlas. Esta voluntad sería coherente con su reducción de la filosofía a lógica más metafísica, y con su subordinación de la teoría del conocimiento a la psicología (4.1121). Pero creo que hay una vía indirecta por la que se pueden obtener consecuencias epistemológicas importantes, consecuencias que van más allá del fenomenalismo de Russell, y que los neopositivistas obtuvieron.

Esa vía transcurre a través de la teoría de la proposición, y por consiguiente pasa por el principio de isomorfía. Para poder ser verdadera o falsa, una proposición ha de tener sentido, y su sentido es el posible estado de cosas que representa. ¿Cuándo podemos decir que es posible un estado de cosas? Esta es la pregunta clave para la construcción ontolingüística del *Tractatus*. La respuesta es: cuando sabemos en qué circunstancias sería

verdadera la proposición que lo representa. Y éste es el principio de verificabilidad, eje de la epistemología neopositivista. Por eso Carnap ha hablado del «principio de verificabilidad de Wittgenstein» («Intellectual Autobiography», p. 45) y es común situar el origen de este principio en el *Tractatus*. ¿Parece exagerada esta interpretación? Medítese, entonces, qué otro significado puede darse a esta afirmación de Wittgenstein: «para poder decir: 'p' es verdadera (o falsa), debo haber determinado en qué circunstancias llamo a 'p' verdadera, y con ello determino el sentido de la proposición» (4.063). Lo cual es del todo coherente con lo que afirma en otro lugar: «lo que se puede describir, puede también suceder» (6.362). Que semejante germen del principio de verificabilidad tiene además el carácter cientifista que habría de recibir de los miembros del Círculo de Viena, corresponde bien a la tesis, ya subrayada anteriormente, y con frecuencia pasada por alto en las interpretaciones metafísicas del *Tractatus*, de que la totalidad de las proposiciones verdaderas es la totalidad de la ciencia natural (4.11). Con ello no quiero insinuar que Wittgenstein se hubiera propuesto formular un principio semejante al, llamado posteriormente, principio de verificabilidad, ni defender una epistemología cientifista. Lo que digo es que, dadas las afirmaciones del *Tractatus* que acabo de citar, y considerando la ausencia de criterios en él que nos permitan determinar en qué consisten los objetos y los nombres, es una fácil consecuencia, por lo demás compatible con el resto de su doctrina, la de que sólo pueden aceptarse como posibles estados de cosas aquellas situaciones que pueden ser contenido de una hipótesis científica, de tal modo que podamos prever en qué circunstancias las proposiciones que las describen serán verdaderas. Un corolario de esta consecuencia es que los objetos de los que habla el *Tractatus* serán aquellos elementos simples a los que pueda llegarse por medio del análisis científico, o sea, físico. Y los nombres serán aquellos términos con los que los designemos. Nótese que esto es perfectamente coherente con la posición de Wittgenstein en el sentido de que la determinación de cuáles sean los objetos, y por tanto, de cómo sean las proposiciones elementales, es una cuestión empírica en la que el filósofo no tiene competencia. Cuestión harto distinta es la de si la teoría física actual es compatible o no con este reduccionismo atomista. Si se acepta la tesis, defendida por diversos autores en los últimos años, de que no hay partículas físicas que sean propia y rigurosamente simples, sino que todo está compuesto de todo, y por consiguiente, que la estructura subatómica no es jerárquica, entonces la visión de Wittgenstein, como cualquier forma del atomismo filosófico, habrá que desecharla (véase D. Miller, «The Uniqueness of Atomic Facts in Wittgenstein's *Tractatus*», 1977).

La interpretación que estoy justificando será rechazada por muchos bajo la acusación de que hace del primer Wittgenstein una especie de cripto-fisicalista. Creo que esto es lo que realmente era el autor del *Tractatus*, y que éste es uno de los aspectos en los que, más claramente, se distingue de Russell. No de otro modo puede explicarse la radical negación que aquél hace de la filosofía. Pero además, es que las justificaciones positivas que he

indicado sobre la base de su teoría de la proposición son muy claras. Piénsese, por ejemplo, y por añadir aún otro argumento, cuándo carece de sentido una proposición para Wittgenstein: cuando alguno de sus signos carece de referencia (5.4733, 6.53). Ahora bien, ¿cómo podremos comprobar si los signos de una proposición poseen o no referencia a menos que la referencia venga dada en la experiencia? Quiénes se niegan a considerar a Wittgenstein empirista vacían de contenido su teoría del significado, haciéndola inútil (creo que es el caso de la profesora Anscombe en el capítulo 12 de su *Introducción al Tractatus*). La gran lección, acaso impremeditada, del *Tractatus*, y que está ya insinuada en Russell, es que una teoría del significado sin una teoría del conocimiento es vacua. Ciertamente que la nueva pregunta: «¿Cómo es posible el lenguaje?», se antepone a la pregunta tradicional en la filosofía moderna: «¿Cómo es posible el conocimiento?»; pero no la excluye. Una teoría del significado basada en la referencia exige ser completada con una teoría del conocimiento. Esto es lo que los neopositivistas vieron con extrema claridad, y en ello su doctrina constituye un desarrollo del *Tractatus* (y no una ruptura, como sugiere la profesora Anscombe, *op. cit.*, p. 152-154).

3. Decir y mostrar

De forma todavía más patente que para Russell, la lógica es común al lenguaje y al mundo, y funda la posibilidad de representar éste en aquél. La comunidad de forma lógica, en la que participan el lenguaje y el mundo, hace posible la correlación de sus estructuras, y permite tomar el principio de isomorfía como expresión del concepto de significado, o más exactamente, de sentido. La forma lógica es la posibilidad de que, tanto los objetos en el estado de cosas como los nombres en la proposición simple, se relacionen de alguna manera (2.151, 2.181 s.). La forma lógica es la posibilidad de que haya relaciones. Por eso, desde el punto de vista lógico, también el signo proposicional es un hecho, como lo es cualquier representación figurativa (2.141, 3.14), puesto que uno y otra consisten en una determinada relación de elementos (2.14).

Ahora bien, siendo esto así, parece que no tendría que haber inconveniente en representar estos hechos, hechos proposicionales, hechos representativos, por medio de otras proposiciones o de otras representaciones. Así acontece en la vida ordinaria: podemos pintar un paisaje, pero también podemos copiar la pintura de ese paisaje; podemos describir un edificio, pero podemos igualmente describir una pintura de ese edificio, e incluso podemos describir la descripción que alguien ha hecho de tal edificio. Es decir, no se ve que haya límite respecto al empleo metalingüístico del lenguaje o, en general, respecto al uso metafigurativo de las representaciones. Sin embargo, Wittgenstein no ha hecho ninguna utilización de esta idea, antes bien, ha negado que se pueda hablar con sentido acerca del lenguaje y de las relaciones entre el lenguaje y el mundo, condenando así al sinsentido las propias afirmaciones del *Tractatus*. Al final de su introducción

a la obra de Wittgenstein, Russell mostró ya sorpresa ante esta doctrina del *Tractatus* sobre la imposibilidad de hablar acerca del lenguaje, sugiriendo como solución la idea de una jerarquía de lenguajes, en la que cada uno de ellos hablara sobre la estructura del situado en el nivel inmediatamente inferior. ¿Qué razones tenía Wittgenstein para no hacer suya esta opinión?

Notemos, en primer lugar, que una jerarquía así no tiene límite. Si para comprender la estructura de un lenguaje L recurrimos a un lenguaje $L + 1$ en el que aquélla quede elucidada, nuestra explicación dependerá de que entendamos la estructura de $L + 1$, para lo cual habremos de recurrir a un nuevo lenguaje $L + 2$; pero la aplicación de este nuevo lenguaje quedará condicionada a la comprensión de su estructura, que no podremos obtener a menos que utilicemos un nuevo lenguaje $L + 3$; y así sucesivamente. Así entendida, la jerarquía de lenguajes no soluciona nuestro problema. Habiendo ostentado Wittgenstein la preferencia que hemos comprobado por la finitud en la interpretación de los cuantificadores, era de esperar que un recurso que, como la jerarquía de metalenguajes, introduce una serie infinita, no despertara su interés. Pero veamos, en segundo lugar, qué es lo que nos resuelve un metalenguaje. Supongamos que mantenemos la teoría isomórfica del significado. En tal caso, las proposiciones de nuestro metalenguaje tendrán sentido en la medida en que haya nombres, cada uno de los cuales se refiera a uno de los nombres del lenguaje objeto, y de tal modo que a cada relación de nombres en el lenguaje objeto corresponda una relación de metanombres en el metalenguaje. Pero con esto lo único que hemos conseguido es duplicar nuestro lenguaje, y nada conseguiremos añadiendo nuevos metalenguajes. Conservando la concepción isomórfica del sentido de las proposiciones, cada lenguaje de nuestra jerarquía quedará legitimado por la relación al lenguaje inmediatamente inferior, pero no hablará propiamente *acerca de* la estructura de éste ni *acerca de* la relación entre éste y el mundo. En definitiva, es lo que acontece en las metafiguras de la vida cotidiana. Si *Las meninas*; nuestra pintura tendrá sentido por relación al cuadro de Velázquez, pero no nos explicará nada sobre la estructura de éste ni sobre la relación que hay entre éste y la situación pintada por Velázquez.

La moraleja es clara: para hablar del lenguaje es inútil recurrir a metalenguajes a menos que, para éstos, se elabore una teoría del significado distinta del principio de isomorfía. La posición de Wittgenstein se comprende ahora mejor: puestos a prescindir del principio de isomorfía, más vale no entrar en las complejidades de una serie infinita de metalenguajes, y limitarse a proponer un concepto semántico diferente. Ese concepto es la idea de *mostrar*, que ya conocemos.

Lo malo es que el concepto de mostrar es totalmente oscuro e impreciso tal y como Wittgenstein lo utiliza. No es, por ello, extraño que sean comunes las críticas recibidas, ni que, quienes más desarrollaron las tesis del *Tractatus*, los neopositivistas, prescindieran de tal concepto enteramente. La realidad es que no sabemos en qué consiste mostrar ni cuáles son las

condiciones para mostrar algo con éxito. De un lado, según el *Tractatus*, las proposiciones con sentido *dicen* algo, verdadero o falso, en cuanto que representan un posible estado de cosas o situación. Pero de otro, Wittgenstein afirma que *muestran* su propia forma lógica, y puesto que ésta es común a la realidad, o sea, a toda situación posible, lo que muestran es, simultáneamente, la estructura lógica del lenguaje y de lo real (que incluye el mundo). El problema es que, cuando el filósofo quiere hablar de esta estructura lógica, no puede limitarse a pronunciar proposiciones con sentido, pues quien no haya percibido la forma lógica del lenguaje no cabe esperar que la capte por muchas proposiciones con sentido que se le repitan; pero todas aquellas proposiciones que el filósofo pretenda hacer *acerca de* la forma lógica carecerán, por definición, de sentido. Las proposiciones (con sentido) muestran, además, según vimos, el sujeto trascendental, como condición del lenguaje y de la realidad, y la vacía corrección de la actitud solipsista. Podemos decir que, en general, el lenguaje (con sentido) muestra todo aquello que es condición de sí mismo y de su función representativa. Pero además, el lenguaje sin sentido también muestra algo: ciertas proposiciones sin sentido muestran por qué las proposiciones sin sentido carecen de sentido, esto es, por qué no puede hablarse de lo que no puede hablarse. ¿Qué proposiciones son las que muestran esto? Las que cumplen una función aclaratoria, las únicas proposiciones filosóficas admisibles para Wittgenstein, de las que son ejemplo las propias proposiciones del *Tractatus*. Y por supuesto: en rigor se trata de pseudoproposiciones, puesto que no dicen nada. Finalmente, no son sólo las proposiciones las que muestran; hemos visto antes que los valores éticos, estéticos y religiosos no son tampoco tema del que se pueda hablar, sino que se muestran, según la interpretación que he presentado, en el sentimiento de finitud. Podemos resumir esto en el siguiente cuadro:

PROPOSICIONES	DECIR	SITUACIONES POSIBLES
	MOSTRAR	FORMA LÓGICA SUJETO TRASCENDENTAL
PSEUDOPROPOSICIONES (filosóficas aclaratorias)	MOSTRAR	FALTA DE SENTIDO
SENTIMIENTO DE FINITUD	MOSTRAR	VALORES (éticos, estéticos, religiosos)

Aquí queda patente la variada aplicación que del concepto de mostrar hace Wittgenstein y el carácter equívoco de este concepto, que en ningún lugar explica ni dilucida, frente a la precisa caracterización, aunque por ventura insatisfactoria, que el principio de isomorfía suministra para el

concepto de decir. Da la impresión de que, para todo aquello sobre lo que Wittgenstein quiere decir algo, y puesto que el principio de isomorfía no puede justificar sus afirmaciones ya que ninguna de ellas describe hechos, recurre al concepto de mostrar. En resumen: toda investigación trascendental es inexpresable, pero su contenido podrá, de alguna manera, mostrarse. Nótese que todo cuanto se muestra, la estructura lógica, el sujeto metafísico, el límite entre el sentido y el sinsentido, los valores, son condición del lenguaje y de la realidad y su tratamiento es explícitamente declarado trascendental por Wittgenstein. El trascendentalismo lingüístico del *Tractatus* es, por ello, una posición paralingüística, es una posición cuyo contenido no puede expresarse por medio del lenguaje sino tan sólo mostrarse. Resulta, por ello, tanto más lamentable que Wittgenstein no haya sido más claro y más explícito sobre los requisitos y condiciones del mostrar.

4. El solipsismo

Sobre la base de un presupuesto fenomenalista, Russell había acabado en el solipsismo. Empero, la cosa no era demasiado grave puesto que se trataba de un hipotético lenguaje lógicamente perfecto. El caso de Wittgenstein, como hemos visto, es muy distinto. Su posición epistemológica más bien es fisicalista que otra cosa, como, sin duda, lo manifiesta la disolución del solipsismo en el realismo que hemos visto sugerida en el *Tractatus*. Las consecuencias del solipsismo hubieran sido más graves para Wittgenstein en la medida en que éste no está formulando las condiciones de un lenguaje perfecto, sino mostrando la estructura del lenguaje real. Esto, unido a su presupuesto cientifista, lo libera de toda propensión a considerar los objetos como contenidos privados, y el intento de encontrar en el *Tractatus* rastros de una concepción privatista del lenguaje es inútil. Wittgenstein ha sacado todas las consecuencias posibles de la idea de que la lógica es común al lenguaje y al mundo. Puesto que la lógica pertenece al pensamiento como tal (3), y ya que la lógica exige que haya objetos, no hay lugar a considerar que los objetos puedan variar de unos sujetos a otros. El desinterés de Wittgenstein por la teoría del conocimiento expresa, tal vez, la ausencia de problematización que ve en ella. El cientifismo, en definitiva, es la sustitución de la teoría del conocimiento por la filosofía de la ciencia. Se comprende que el *Tractatus* tuviera tan gran influencia en la formación del neopositivismo, sobre todo si se piensa que, por esta época, Carnap estaba atravesando una etapa fenomenalista.

5. El sentido del *Tractatus*

Se ha comparado el empeño de Wittgenstein con la tarea de Descartes, y se ha dicho que, así como éste exploró las implicaciones filosóficas de la nueva ciencia creada por Copérnico y Galileo, Wittgenstein habría explorado en el *Tractatus* las consecuencias filosóficas de la nueva ciencia (for-

mal) creada por Frege y Russell, la lógica matemática. Hasta cierto punto, la comparación es aceptable, pero me parece confundente contraponer, en este aspecto, Wittgenstein a Frege y Russell; de una parte, estos mismos habían ya comenzado a formular la filosofía apropiada a la nueva lógica formal que habían creado; de otra, también Wittgenstein, aunque en grado menor, había participado en la construcción de la nueva ciencia. Más iluminadora es la comparación con Kant, que también se hace con frecuencia. Así como éste, contestando a la pregunta «¿Cómo es posible la ciencia?», explora los límites del conocimiento, Wittgenstein, respondiendo a la pregunta «¿Cómo es posible el lenguaje?», explora los límites del pensamiento. Estos límites son los que separan el discurso con sentido del discurso que carece de sentido. Dentro del pensamiento queda la ciencia, y más allá de él, la metafísica. En el puro límite, la lógica, la ética, la estética y la religión. Y todo esto nos lo enseña Wittgenstein por medio de un conjunto de afirmaciones que, según él mismo reconocerá, tampoco tienen sentido. Esta es la profunda, la desmesurada paradoja del *Tractatus*. ¿Pero se reduce la obra a este empeño didáctico? Se ha añadido que, del modo como Kant recurre a la razón práctica para superar los límites de la razón teórica, Wittgenstein recurre a lo que se puede mostrar para superar los límites de lo que se puede decir. La intención de Wittgenstein parece, sin embargo, haber querido ir más lejos, hasta el punto de otorgar a su libro un sentido moral, más sobre la base de lo que no dice que en razón de lo que afirma. Esto es, al menos, la descripción del *Tractatus* que, hacia 1919, Wittgenstein le hizo a su amigo Ficker. En una carta, le dice, en efecto: «el sentido del libro es un sentido ético»; y le da, como clave para entenderlo, la siguiente indicación que, según declara, tuvo una vez intención de incluir en el prólogo: «Mi obra consiste en dos partes, la que aquí presento y la que no he escrito. Y justamente esta segunda parte es la importante.» ¿Por qué? La respuesta la encontramos en lo que sigue: «Lo ético es delimitado por mi libro, por así decirlo, desde dentro; y estoy convencido de que, estrictamente, sólo así puede delimitarse.» En resumen: Wittgenstein traza los límites de lo ético a base de callar sobre ello allí donde tendría que haber hablado acerca de ello; tales límites quedan, pues, mostrados. Y éste es el sentido fundamental del *Tractatus*: mostrar lo poco que se consigue cuando se solucionan los problemas de los que trata (véase el fin del prólogo), pues lo realmente importante son los problemas éticos, como el sentido de la vida, y sobre ello nada puede decirse. (Las citas anteriores pertenecen a una carta cuyo original alemán y traducción inglesa se encontrarán en las páginas 143-144 de la recopilación de Engelmann, *Letters from Ludwig Wittgenstein, with a Memoir*.)

En fin, la búsqueda de un lenguaje lógicamente perfecto, que hemos visto iniciada en Frege y desarrollada por Russell más cerca de la lógica que del lenguaje ordinario, ha conducido a Wittgenstein, restableciendo el equilibrio entre aquella y éste, a postular la posibilidad de un análisis que nos suministre la estructura lógica subyacente al lenguaje ordinario. Ello probaría que el lenguaje lógicamente perfecto que perseguíamos, lo teníamos

ante nosotros sin darnos cuenta. Era el propio lenguaje ordinario, que bajo su multiforme apariencia ocultaba su escueta estructura. Pero todo esto supone que sabemos cómo llevar a cabo el análisis, suposición que, lamentablemente, y como ya he subrayado, no se cumple. En el fondo, Wittgenstein, como Frege y Russell, no ha salido de la lógica. Su salida tardará varios años en producirse, y será tan violenta que le sitúe en el extremo teórico opuesto.

Lecturas

Siendo éste el primer capítulo en el que se trata temáticamente de la filosofía analítica, sería conveniente hacer referencia, para empezar, a las obras generales que versan sobre aquélla, particularmente desde el punto de vista de la teoría del significado o de otros problemas en conexión con el lenguaje. De este tipo de libros, la bibliografía es escasa, y no sólo en castellano. Tenemos el libro de Urmson, *El análisis filosófico* (Ariel, Barcelona, 1978), que, según indica su subtítulo, se limita al período comprendido entre las dos guerras. De hecho, es, fundamentalmente, una buena exposición del atomismo lógico, con alusiones interesantes, aunque incompletas, al neopositivismo, y más breves aún a la filosofía del lenguaje ordinario. Hay que tener en cuenta que el original inglés es de 1956, y desde entonces se ha escrito mucho sobre estos temas. El titulado *Los orígenes de la filosofía analítica*, por Alston y otros autores (Tecnos, Madrid, 1976), es considerablemente más elemental y más simple. Trata de Moore (sobre quien yo hablaré en el próximo capítulo), de Russell, y de Wittgenstein en sus dos etapas. Consiste simplemente en la traducción de los artículos que a ellos se les dedica en *The Encyclopedia of Philosophy*, obra en ocho tomos que, bajo la dirección de Edwards, se publicó en 1967 (MacMillan, Londres). Estos trabajos pueden ser útiles para quienes deseen obtener, a un nivel sencillo, y con brevedad, una visión de conjunto sobre dichos autores. Carácter muy distinto posee la obra de Thomas Moro Simpson (argentino, pese a semejante nombre y apellido), *Formas lógicas, realidad y significado* (Editorial Universitaria de Buenos Aires, 1964; segunda edición corregida y aumentada, 1975). Este, que es el primer libro de filosofía analítica escrito originalmente en castellano, es un estudio monográfico sobre problemas de teoría del significado, y en él, como es natural, se subordina lo histórico a lo sistemático. No obstante, el lector encontrará allí con facilidad muchos de los puntos que hemos discutido en este capítulo, y otros que discutiremos en los sucesivos. Simpson ha realizado también, bajo el título *Semántica filosófica: problemas y discusiones* (Siglo XXI, Buenos Aires, 1973), una interesante recopilación de artículos que, aunque muy especializada, contiene, sobre todo en su primera parte, trabajos que hemos discutido en este capítulo o que habremos de discutir en los próximos. Carácter general tiene asimismo el libro de Christensen, *Sobre la*

naturaleza del significado (Lábor, Barcelona, 1968), aunque adolece en muchos puntos de falta de claridad. Los dos volúmenes que, con el título de *La concepción analítica de la filosofía*, ha recopilado Javier Muguerza (Alianza Universidad, Madrid, 1974) contienen muchos escritos analíticos que, por ser clásicos y centrales, son de cita obligada y de lectura inevitable.

Los principales trabajos de Frege que hemos estudiado en este capítulo están traducidos en *Estudios sobre semántica* (Ariel, Barcelona, 1971), así como en *Escritos lógico-semánticos* (Tecnos, Madrid, 1974); de ambas recopilaciones, me parece superior, y mejor traducida, la primera. El estudio de Thiel, *Sentido y referencia en la lógica de Gottlob Frege* (Tecnos, Madrid, 1972), considera, en su segunda parte, con mayor amplitud y detalle, los aspectos que hemos estudiado del pensamiento de Frege. Rompiendo la regla, hasta ahora seguida escrupulosamente, de no recomendar en estos apéndices de lecturas más que obras en castellano, me permitiré citar el libro más completo que conozco sobre los problemas del lenguaje en Frege, que es el de Dummett, *Frege. Philosophy of Language* (Duckworth, Londres, 1973), en el que tan sólo es de lamentar que la claridad no siempre acompañe a la morosidad de la argumentación.

De Russell, «La filosofía del atomismo lógico» está incluido en el volumen primero de la recopilación de Muguerza, ya citada, *La concepción analítica de la filosofía*; junto con otros escritos de Russell, algunos tan relevantes para nuestro tema como «Sobre la denotación», está también incluido en su libro *Lógica y conocimiento* (Taurus, Madrid, 1966). Es útil el pequeño estudio de Clack, *La filosofía del lenguaje de Bertrand Russell* (Fernando Torres, Valencia, 1976), aunque para algunos temas resulte excesivamente conciso.

El *Tractatus Logico-Philosophicus* de Wittgenstein fue publicado con traducción castellana por la editorial Revista de Occidente, Madrid, en 1957, en meritoria (aunque no del todo correcta) traducción de Tierno Galván, y ha sido editado por Alianza Editorial en 1987 con una nueva traducción de Jacobo Muñoz e Isidoro Reguera, la cual mejora mucho la anterior. Sobre Wittgenstein, comenzaré por citar obras generales, que, por consiguiente, no sólo comentan el *Tractatus*, sino también escritos posteriores de Wittgenstein que nosotros estudiaremos en el capítulo que viene. De tales obras, una de las mejores, y en mi juicio, la mejor de las traducidas al castellano, es el *Wittgenstein* de Kenny (Alianza Universidad, Madrid, 1974). Una de las características que más se ha destacado de este libro es el hincapié que hace en la continuidad del pensamiento de Wittgenstein y la atención que dedica a las etapas intermedias que ligán la doctrina del *Tractatus* con la madurez de su último pensamiento. No obstante, este procedimiento tiene, en mi entender, un inconveniente, y es que desdibuja un tanto el pensamiento maduro del último Wittgenstein, que sobre varios temas se halla sólo imperfectamente formulado en los escritos de etapas intermedias. Otro estudio general interesante y agudo, aunque

más breve y menos detallado, es el *Wittgenstein* de Pears (Grijalbo, Barcelona, 1973). Aún más breve, y al tiempo más sencillo y elemental, es *El concepto de filosofía en Wittgenstein*, de Fann (Tecnos, Madrid, 1975); quienes encuentren dificultoso seguir los libros de Kenny o Pears, harán bien en recurrir al de Fann, pues les facilitará el camino. Una obra de concepción semejante, pero considerablemente inferior, es *Wittgenstein y la filosofía contemporánea*, de Hartnack (Ariel, Barcelona, 1972); aunque por sus dimensiones, por su precio, y por ser la primera obra de este tipo publicada en castellano, fue durante un tiempo popular entre el público universitario, su lectura no me parece hoy aconsejable, pues el autor es, sobre muchos puntos, tan impreciso, e incluso confundente, que vale más abordar cualquiera de los libros citados antes. Los ensayos recopilados por Winch con el título *Estudios sobre las filosofías de Wittgenstein* (Editorial Universitaria de Buenos Aires, 1971) tienen interés, pero tratan aspectos muy especializados del pensamiento de Wittgenstein. Dentro del ámbito de estas obras generales citaré, por último, un libro no traducido, que, aunque es general por abarcar todo el desarrollo filosófico de Wittgenstein, es particular por su tema, la epistemología, o, en expresión del autor, la «metafísica de la experiencia»; esta formulación es intencionada, pues el libro explora la veta kantiana en la doctrina de Wittgenstein. Se trata de *Insight and Illusion*, de Hacker (Oxford University Press, 1972).

Sobre el primer Wittgenstein en particular no hay tanta bibliografía. Es recomendable, en primer lugar, la excelente *Introducción al Tractatus de Wittgenstein*, de Elizabeth Anscombe (El Ateneo, Buenos Aires, 1977), obra centrada en la filosofía de la lógica wittgensteiniana, pero muy rica en matices, y con todo el rigor que caracteriza a la autora. La revista *Teorema*, de Valencia, publicó en 1972 un número monográfico titulado *Sobre el Tractatus Logico-Philosophicus*, que incluye un temprano ensayo de Wittgenstein, «Notes on Logic» (1913), junto con estudios de especialistas españoles y extranjeros (estos últimos en la lengua original). Un libro de peculiares características, y de innegable atractivo, es *La Viena de Wittgenstein*, de Janik y Toulmin (Taurus, Madrid, 1974). Aquí se trata de situar al primer Wittgenstein sobre el trasfondo de la sociedad y de la cultura vienesas de principios de siglo, de las cuales, en definitiva, procede, y de extraer de ellas claves para la explicación del *Tractatus*. La pintura histórico-social que hacen los autores es compleja y brillante (si es exacta, no puedo juzgarlo yo), pero, en mi opinión, contribuye escasamente a hacer el *Tractatus* más inteligible; lo cierto es que, mal que pese a los autores, y por muy vienes que fuera Wittgenstein, las ideas contenidas en el *Tractatus* provienen fundamentalmente de Frege y Russell y sólo se entienden en este contexto teórico. El libro de Janik y Toulmin tiene, por otro lado, el inconveniente de cometer serios errores al explicar la doctrina de la representación en el capítulo 6. Para acabar, tampoco ahora resistiré la tentación de recomendar alguna obra no traducida. Me limitaré a una interpretación atomista del *Tractatus* que pudiéramos llamar «clási-

ca», como es la ofrecida por Griffin en *Wittgenstein's Logical Atomism* (Oxford University Press, 1964), y a una importante y completa recolección de artículos acerca del *Tractatus*, la realizada por Copi y Beard con el título de *Essays on Wittgenstein's Tractatus* (Routledge & Kegan Paul, Londres, 1966). También es recomendable el libro de H. O. Mounce, *Introducción al «Tractatus» de Wittgenstein*, Ed. Tecnos, Madrid, 1983.

Una obra de conjunto sobre la filosofía analítica del lenguaje, que tiene el valor de incluir algunos precedentes históricos así como desarrollos tan recientes como la teoría de Davidson, es el libro de Ian Hacking, *¿Por qué importa el lenguaje a la Filosofía?*, Ed. Sudamericana, Buenos Aires, 1979.

Una antología que incluye un buen número y variedad de artículos fundamentales en la filosofía analítica del lenguaje es la realizada por Luis Valdés con el título *La búsqueda del significado* (Tecnos, Madrid, 1991). Junto a ella resulta muy escasa, aunque todavía valiosa, la *Antología semántica* de Mario Bunge (Ediciones Nueva Visión, Buenos Aires, 1960), que durante muchos años hemos venido utilizando.

Quienes estén interesados en las relaciones que hay entre la obra de Wittgenstein y su compleja personalidad, pueden estudiar la primera siguiendo el hilo de la biografía del autor con obras como *Ludwig Wittgenstein*, de Baum (Alianza, Madrid, 1988), *Wittgenstein*, de Bartley III (Cátedra, Madrid, 1982), y *Wittgenstein. El joven Ludwig*, de McGuinness (Alianza, Madrid, 1991).

A lo anterior hay que añadir la publicación de artículos de Frege, *Escritos filosóficos*, en edición de Jesús Mosterín (editorial Crítica, Barcelona, 1996) así como el excelente libro de A. Kenny, *Frege. Una introducción* (ed. Tecnos, Madrid, 1997).

Capítulo 7

LOS ABUSOS DEL USC

Nacen las cosas cuando nacen las palabras; sin palabras no hay cosas, o si las hay, es como si no las hubiese, porque la cosa no existe por sí ni para otras cosas. (PÉREZ DE AYALA, *Belarmino y Apolonio*.)

7.1 Prolegómeno

Durante la primera guerra mundial, Wittgenstein, que luchaba en las filas austriacas, continuó trabajando en el *Tractatus*. Para 1918, la obra estaba terminada, y Wittgenstein, que había sido hecho prisionero por los italianos, hizo llegar una copia a Russell. Aunque la crítica reacción de Russell decepcionó a Wittgenstein, que en principio se negó a que la introducción de aquél se imprimiera con el *Tractatus*, éste se publicó finalmente como libro, acompañado de dicha introducción, en 1922. Pero para entonces Wittgenstein había abandonado del todo la filosofía. En coherencia con la tesis, allí mantenida, de que todos los problemas filosóficos son realmente pseudoproblemas, no dedicó a ellos más atención, al menos públicamente. Hasta 1926 trabajó de maestro de escuela en pequeñas aldeas de Austria y, posteriormente, de jardinero en un convento y como arquitecto en la construcción de una mansión para una hermana suya.

Vale la pena notar al paso que tan rigurosa actitud, cuya coherencia con su teoría acaba de indicarse y es comúnmente elogiada, no estaba exigida por el *Tractatus*. Sin traicionar sus convicciones, Wittgenstein podría haber dedicado sus esfuerzos a disipar y disolver problemas filosóficos concretos, aplicando así las tesis del *Tractatus*, al tiempo que proseguía cultivando disciplinas formales para las que ya había demostrado interés y genio, como la lógica y la matemática. En suma, la renuncia a la filosofía no tenía por qué implicar la renuncia a la teoría. Debió, pues, de haber algo más, lo que no es de extrañar dada la compleja personalidad de Wittgenstein. Sea como fuere, es el caso que, tras coronar una obra tan ardua

como el *Tractatus*, Wittgenstein se retiró de la vida intelectual pública durante los diez años siguientes.

En este tiempo, su obra fue adquiriendo progresiva influencia, y todavía más que en Inglaterra, en Alemania y Austria, hasta el punto de convertirse en uno de los pilares de lo que luego se iba a llamar «Círculo de Viena». El impacto del *Tractatus* en Carnap, en Schlick, en Feigl o en Waismann, fue muy fuerte; el primero de ellos narra que, en el Círculo, el *Tractatus* era leído en voz alta y discutido sentencia por sentencia («Intellectual Autobiography», p. 24). Esto no quiere decir que todas las tesis de Wittgenstein fueran incorporadas al ideario del Círculo de Viena; consecuentes con un positivismo más general y sistemático, sus miembros prescindieron de todo aquello que la sensibilidad metafísico-religiosa de Wittgenstein había introducido en el ámbito de lo que se muestra, y potenciaron, en cambio, el principio de verificabilidad que acompaña a la concepción isomórfica del lenguaje y el cientifismo que ya se apuntaba en el *Tractatus*. Partiendo de estos presupuestos se desarrolló toda una línea en la filosofía analítica del lenguaje que llega hasta hoy y ha sido extremadamente fecunda. Nosotros la examinaremos en el próximo capítulo.

Aquí vamos a tratar de otra línea, cuya vigencia alcanza igualmente hasta la actualidad, y que también procede del *Tractatus*, pero por modo contrario. Esta línea tiene su origen en una reacción contra el atomismo lógico, y es el propio Wittgenstein quien la lleva a cabo. Lo primero que hay que estudiar es, por tanto, la doctrina expuesta por éste en su segunda etapa filosófica. Llamaremos así a la etapa que se abre cuando Wittgenstein regresa a Cambridge en 1929 para reanudar su dedicación académica a la filosofía, y que se dilata ya hasta su muerte en 1951. Durante los años de su reclusión, Wittgenstein no dejó de recibir diversas invitaciones para su vuelta a la universidad, bien de amigos ingleses, como Ramsey, bien de miembros del Círculo de Viena, como los mencionados, con los cuales mantenía discusiones privadas. Según uno de sus biógrafos, Wittgenstein volvió a la filosofía porque tenía el sentimiento de que otra vez podía hacer una obra original (Von Wright, «Biographical Sketch», pp. 12-13). Es indudable que este sentimiento iba unido a una actitud crecientemente crítica del propio *Tractatus* y, en general, de todo el atomismo lógico.

Como vamos a ver, una de las tesis centrales que Wittgenstein había tomado de Russell, a saber: que la lógica suministra la estructura del lenguaje y de la realidad, fue sustituida por otra sumamente opuesta: que el lenguaje ordinario es mucho más rico que la lógica y que ésta no puede, por ello, darnos ninguna pista para entender aquél y menos aún para averiguar en qué consiste la realidad. Entre el lenguaje ordinario y la lógica es ahora el primero el que obtiene la primacía. Es por esta razón por la que algunos discípulos del segundo Wittgenstein han creído encontrar en Moore, filósofo de Cambridge contemporáneo de Russell y Wittgenstein, un parentesco, e incluso un precedente, respecto de la posición mantenida por este último en su segunda etapa filosófica.

Uno de los principales propagadores de esta imagen ha sido Malcolm, quien afirma, por ejemplo, que «la esencia de la técnica de Moore para refutar afirmaciones filosóficas consiste en señalar que tales afirmaciones van *contra el lenguaje ordinario*» («Moore and Ordinary Language», p. 349, subrayado en el original). Parecería que los filósofos tradicionales gustan de repudiar el lenguaje ordinario con sus afirmaciones, tan opuestas al sentido común. Frente a ellos, Moore habría reaccionado vindicando el lenguaje ordinario (*op. cit.*, p. 365). De aquí que el papel histórico de Moore en la filosofía haya consistido, al decir de Malcolm, en «haber sido tal vez el primer filósofo en darse cuenta de que cualquier enunciado filosófico que viole el lenguaje ordinario es falso, y el primero en defender de modo consistente el lenguaje ordinario de sus violadores filosóficos» (*op. cit.*, p. 368). De modo semejante, Lazerowitz ha descrito la defensa que hacía Moore del sentido común como una defensa del lenguaje del sentido común contra su modificación por los filósofos («Moore's Paradox», p. 393).

Es cierto que Moore se sentía profundamente sorprendido por la metafísica tradicional, y en particular por el neohegelianismo inglés de principio de siglo. Afirmaciones tales como «El tiempo no es real», «El mundo sólo existe mientras lo percibo», «No puedo tener certeza de que las demás personas piensen y sientan», y otras de esta guisa, le parecían sorprendentes, y contra ellas asumió, como tarea filosófica, la defensa de las creencias de sentido común. Sin duda que es posible interpretar esta tarca como una depuración del lenguaje filosófico, en el sentido de sustituir afirmaciones paradójicas, como las citadas, por las correspondientes afirmaciones de sentido común, pero esa no es fundamentalmente una tarea lingüística ni implica una teoría del lenguaje, a diferencia de lo que, como veremos, ocurre en el caso de Wittgenstein. No puede extrañar, por ello, que el propio Moore haya protestado contra esa interpretación de su método filosófico, contestando a Lazerowitz que, si todo lo que él (Moore) hacía era recomendar que no se usaran ciertas expresiones de manera distinta a como ordinariamente se usan, entonces había cometido una gran equivocación, pues estaba convencido de que lo que hacía no se reducía a eso («A Reply to My Critics», p. 675).

La imagen de Moore como filósofo del lenguaje ordinario tiene, por supuesto, cierta base en las apelaciones que hace, en ocasiones, al sentido ordinario de las palabras. Así, en un escrito tan típico como «Defensa del sentido común», acusa a determinados filósofos de utilizar el lenguaje de una manera tan peculiar que estarían dispuestos a llamar «verdadera» a una proposición que sea parcialmente falsa; y añade: «Deseo dejar claro que no estoy usando 'verdadero' en tal sentido; lo estoy usando de tal modo (y creo que éste es el uso ordinario) que si una proposición es parcialmente falsa, se sigue que no es verdadera, aunque pueda serlo parcialmente» («A Defence of Common Sense», p. 35). Es decir: frente a un uso filosófico peculiar de una palabra como «verdadero» habría un uso ordinario de cuyo lado se sitúa Moore. Lo propio ocurre cuando, a propósito de

otra argumentación, considera la expresión «hecho físico», y declara en su descargo: «se trata de una expresión de uso común, y creo que yo la estoy usando en su sentido ordinario» (*op. cit.*, p. 45). Por lo que respecta al lenguaje, la defensa que hace Moore del sentido común tiene como componente importante esta consideración: las expresiones poseen un significado «ordinario o popular», y esto es lo que, desgraciadamente, algunos filósofos son capaces de poner en duda (*op. cit.*, p. 36). De aquí que la defensa del sentido común, tal y como la lleva a cabo Moore, comporte el uso de las palabras en su sentido ordinario y la denuncia de aquellos usos extravagantes y peculiares que permiten a los filósofos formular sus paradojas o hacer afirmaciones incomprensibles. Desde el punto de vista metodológico, esta actitud está muy próxima a la del segundo Wittgenstein, como veremos. Pero desde el punto de vista teórico, la diferencia es clara: de un lado, Wittgenstein no está ocupado en una defensa de las creencias de sentido común; de otro, Moore carece de una teoría del lenguaje que justifique su proceder, teoría que, sin embargo, constituye el centro de la filosofía del segundo Wittgenstein. Por lo demás, las posibles influencias de uno de ellos sobre el otro no resultan claras, a pesar de que fueron colegas en Cambridge durante muchos años. «Defensa del sentido común» es de 1925, anterior, por consiguiente, al regreso de Wittgenstein. Moore asistió a las clases de Wittgenstein de 1930 a 1933, pero a la hora de decidir en qué medida este último le influyó, declara no saberlo, aunque reconoce que le hizo desconfiar de muchas cosas que, a no ser por él, se hubiera sentido inclinado a afirmar, y añade: «Me ha hecho pensar que lo que se requiere para la solución de los problemas filosóficos que me tienen perplejo, es un método muy diferente de cualquiera de los que yo he usado siempre, método que él utiliza con éxito, pero que yo nunca he sido capaz de entender con claridad suficiente para poder utilizarlo» («An Autobiography», p. 33). Estas palabras insinúan una distancia entre ambos mucho mayor de lo que da a entender la interpretación propuesta por algunos discípulos del segundo Wittgenstein, como los ya mencionados. Sólo cabe pensar que la tarea de Moore, contemplada a la luz de la enseñanza del segundo Wittgenstein, y por consiguiente destacando en aquélla la apelación al lenguaje ordinario, produce una impresión que, como la que les produjo a Malcolm y a Lazerowitz, no coincide con la idea que el propio Moore tenía. En conclusión, si hemos de atender a lo que Moore conscientemente intentaba hacer, no parece existir razón bastante para ver en él a un filósofo del lenguaje ordinario en el sentido del segundo Wittgenstein. Y puesto que el problema del lenguaje no se plantea como tema en la filosofía de Moore, queda justificado que no le dediquemos más atención.

He hablado en lo anterior del segundo Wittgenstein para referirme a la filosofía elaborada por él en su segunda etapa universitaria, esto es, a partir de su vuelta a Cambridge en 1929 y hasta su muerte. Muchos opondrían, sin embargo, que eso constituye una intolerable simplificación, y que el pensamiento de Wittgenstein es mucho más rico que todo eso. Aun cuando la división de su filosofía en dos etapas claramente diferenciadas,

y en algunos aspectos contrapuestas, la etapa lógico-atomista representada por el *Tractatus* y la etapa de primacía del lenguaje ordinario que representan las *Investigaciones filosóficas*, es usual en los estudios sobre Wittgenstein, algunos han prestado atención a los estadios intermedios que forman el puente entre ambas obras. Así, por ejemplo, y en relación con el tema de la crítica al solipsismo, se ha señalado, en esta segunda etapa, un primer estadio próximo al neopositivismo del Círculo de Viena entre 1929 y 1933; un segundo estadio, de 1933 a 1936, representado por el *Cuaderno Azul* principalmente, y un estadio final, correspondiente a las *Investigaciones* (Hacker, *Insight and Illusion*, cap. VII, secc. 1). Desde un punto de vista más general, Kenny, en su *Wittgenstein*, ha intentado hacer justicia al desarrollo de la segunda etapa de éste, dedicando capítulos distintos a las obras principales que representan esos estudios intermedios. Como ya indiqué en el apéndice de lecturas del capítulo anterior, semejante estudio tiende a desfigurar el hecho de que muchos de los temas centrales que Wittgenstein trata en esas obras están expuestos con mayor precisión en las *Investigaciones filosóficas*. Hay que tener en cuenta, además, que ninguna de esas obras, ni siquiera esta última, llegó a recibir de su autor la forma definitiva de libro, y que más bien constituyen conjuntos de notas o apuntes, pertenecientes a períodos determinados, y en diferente grado de elaboración. Desde este último punto de vista, parece que el manuscrito más elaborado es, con mucho, el de las *Investigaciones*, y por esta razón, para cualquiera de los temas que allí se tratan, debe darse primacía a esta obra sobre las demás. Todas ellas han sido publicadas póstumamente, pues Wittgenstein, después de publicado el *Tractatus*, no volvió a publicar nada, excepto un artículo de revista del que en seguida se sintió muy insatisfecho.

No hace falta añadir, por todo ello, que estaría fuera de lugar atender aquí al desarrollo del pensamiento de Wittgenstein en esta segunda etapa. Lo más que podemos hacer, de modo semejante a lo que se hizo en el capítulo anterior, es estudiar su filosofía del lenguaje en la obra más elaborada y representativa, en este caso las *Investigaciones filosóficas*, libro en el que trabajó aproximadamente desde 1935 a 1949. La exposición será completada ocasionalmente con referencias a otras obras, en especial a los llamados *Cuadernos Azul* y *Marrón*, que, si bien son meros apuntes de clase de los años 1933 a 1935, posteriormente revisados y corregidos, es el escrito que más se aproxima, por los temas y por su elaboración, a las *Investigaciones*.

7.2 Significado y uso en el segundo Wittgenstein

La mejor forma de entender las afirmaciones del segundo Wittgenstein acerca del lenguaje es compararlas con la doctrina del *Tractatus*, buscando siempre en aquéllas una crítica, aunque sea tácita, al atomismo lógico. De hecho, Wittgenstein lo reconoce así en el prólogo a las *Investigaciones*, donde afirma que sus nuevos pensamientos sólo quedarán correctamente

iluminados al ser contrapuestos a los antiguos y vistos contra el transfondo de ellos. Y hasta tal punto, que declara habersele ocurrido en una ocasión la conveniencia de publicar juntos el texto del *Tractatus* y la nueva obra (proyecto que posteriormente ha llevado a cabo la editorial alemana Suhrkamp como primer volumen de las *Obras* de Wittgenstein, incluyendo además con aquellos dos libros los *Cuadernos de 1914-1916*). Wittgenstein es igualmente lúcido y honesto respecto a las razones de esa comparación entre su vieja y su nueva teoría: se trata de que, desde que en 1929 volvió a trabajar en filosofía, se ha visto obligado a reconocer graves errores en su primer libro. Las *Investigaciones*, como vamos a ver, están en su mayor parte dedicadas a la denuncia y corrección de esos errores, muchos de los cuales, como era de esperar, se encuentran asimismo en Russell.

Su nueva obra consta de dos partes. La primera, que es la más revisada y pulida, está dividida en pequeños párrafos numerados, algunos tan cortos como un par de líneas; excepcionalmente, los hay de una página o página y media. Las referencias a esta parte las haré por el número del párrafo. Esta primera parte la dio Wittgenstein por terminada en 1945, fecha que, por cierto, lleva el prólogo que se acaba de citar. La segunda parte, en cuya revisión trabajó de 1947 a 1949, está claramente menos elaborada; se halla dividida en secciones, algunas de tan sólo media página, otras más largas, y la sección XI, con gran diferencia, alcanza a más de treinta páginas. Las citas de esta segunda parte las haré por el número de la sección en números romanos precedidos de la indicación «II», y cuando se trate de una sección larga, añadiré la página. El estilo de Wittgenstein en ésta, su más cuidada obra de la segunda época, consiste en notas sueltas, con frecuencia de estructura aporética, esto es, más sugiriendo dudas y dificultades que haciendo declaraciones claras; esto dificulta, a veces, entender cuál es la opinión precisa de Wittgenstein sobre el problema debatido, y obliga a examinar otras afirmaciones relacionadas. Tales notas carecen del cerrado sistematismo que posee el *Tractatus*, de modo que, al gusto wittgensteiniano por las declaraciones escuetas y breves, se añade ahora la ausencia de una línea argumental clara y continua. Wittgenstein confiesa en el prólogo su fracaso en el intento de sistematizar sus notas en un todo ordenado, y añade: «Esto estaba, por cierto, conectado con la naturaleza de la investigación. Pues ésta nos obliga a cruzar un dilatado campo del pensamiento en todas direcciones. Las consideraciones filosóficas de este libro son, en cierto modo, un conjunto de esbozos de paisajes originados en estos largos y enrevesados recorridos.» Wittgenstein es consciente de que el estilo de su obra la hace poco atractiva, enmarañada, repetitiva. Y acabará diciendo: «Por ello este libro es propiamente sólo un álbum.»

La idea básica de Wittgenstein sobre el lenguaje es ahora, en contra del *Tractatus*, que no hay nada común a todos los fenómenos lingüísticos en cuya virtud podemos englobarlos bajo el término «lenguaje» y, por consiguiente, que no ha lugar a una teoría sobre la forma general de las pro-

posiciones, como la desarrolla en el *Tractatus*. Lo que nos permite usar el término «lenguaje» para un amplio conjunto de fenómenos no es que éstos tengan algo en común, sino que están relacionados entre sí de muchas maneras distintas (65). ¿Justifican estas relaciones que llamemos «lenguaje» a todos esos fenómenos? Sí, porque son esas relaciones las que nos permiten pasar de un fenómeno a otro y reconocerlos, así como miembros de un mismo conjunto; pero los miembros de este conjunto lo son, no porque tengan en común una cierta propiedad, sino porque están entre sí relacionados unos con otros, aunque no necesariamente todos con todos. Para hablar de este tipo de conjuntos, como es el conjunto de los fenómenos lingüísticos, Wittgenstein utilizará un término corriente: «familia».

El ejemplo clásico, y conocido, al que Wittgenstein recurre para aclarar su concepción es el de los juegos. Su tratamiento del tema es bien representativo de su actual método. Consideremos —recomienda— eso que llamamos «juegos» (66): juegos de cartas, juegos de mesa, juegos de pelota, juegos de competición... ¿Qué tienen en común? No vale decir que algo habrán de tener en común porque en caso contrario no los llamaríamos «juegos»; lo que hay que hacer es mirar si lo tienen. Y mirando se advierte que no aparece nada que sea común a todos ellos. Lo que hay es una serie de parecidos y relaciones. Un grupo de juegos tiene en común con otros ciertas características, pero con un tercer grupo solamente coincide en algunas de ellas, mientras que surgen otras coincidencias nuevas; y así sucesivamente. ¿Son todos divertidos? Piénsese en el ajedrez. ¿Son todos competitivos? Piénsese en los solitarios. Las semejanzas aparecen y desaparecen según pasamos de unos juegos a otros. El resultado de este examen es: «Vemos una complicada red de semejanzas que se solapan y entrecruzan, unas son generales, otras particulares. Y no puedo caracterizarlas mejor que con la expresión «parecidos de familia»; pues así se solapan y entrecruzan, unas son generales, otras particulares. Y no puedo caracterizarlas mejor que con la expresión 'parecidos de familia'; pues así se solapan y entrecruzan las distintas semejanzas que hay entre los miembros de una familia: estatura, rasgos, color de los ojos, modo de andar, temperamento, etc. Y diré: los juegos constituyen una familia» (66-67).

La idea es, por consiguiente, que los miembros de una familia no se identifican por la posesión de una característica común, sino por su pertenencia a una determinada red de relaciones. Tal es el caso de los juegos, según piensa Wittgenstein. Por eso no puede darse una definición exacta de «juego»: el concepto de juego carece de límites estrictos (68-69). Mas esto no nos impide usar de él con éxito; no nos impide explicar a alguien a qué llamamos «juego», pues podemos dar ejemplos (71). La definición esencial no es el único medio de explicar un concepto, por lo mismo que un concepto de límites borrosos no deja por eso de ser un concepto. En este caso se encuentran, según Wittgenstein, el concepto de juego, el concepto de número (67) y el concepto de lenguaje, esto es, de fenómeno o de uso lingüístico. La comparación del lenguaje con los juegos o la idea de los parecidos de familia no son, por consiguiente, meras metáforas, sino

que son piezas centrales de la nueva teoría de Wittgenstein sobre el lenguaje, e importa no asumir a la ligera que lo que afirma sobre el tema es correcto.

¿Es, en efecto, indefinible el juego? El *Diccionario de la Real Academia Española* lo define así: «Ejercicio recreativo sometido a reglas, y en el cual se gana o se pierde.» Bien, ¿no puede afirmarse esto de todos los juegos? Con ello queda excluido el ejercicio profesional de los juegos, pero parece razonable que así sea. Aun cuando digamos de los deportistas profesionales que juegan, por ejemplo, al fútbol, ciertamente su ejercicio no constituye en propiedad un juego: no bajan al campo a recrearse o entretenerse, sino a ganarse la vida (aunque acaso se diviertan mucho jugando). Parece bastante claro que no llamamos nunca «juego» a una actividad que no sea recreativa, a menos que utilicemos el término en forma derivada o metafórica. Wittgenstein parece dudar de que el ajedrez o las tres en raya sean entretenidos o divertidos (*unterhaltend*, 66). Confieso que no entiendo esa duda. ¿Está todo juego sometido a reglas? Bien, si en todo juego se gana o se pierde, reglas habrá de haber, al menos las que determinen cuándo se gana y cuándo se pierde. Que ocurra esto último no supone, naturalmente, que haya competición. En los solitarios no hay competición; sin embargo, se gana o se pierde. Wittgenstein dice que esta característica falta cuando un niño tira la pelota contra la pared y la vuelve a coger (*loc. cit.*). Creo que se equivoca: cuando el niño la coge, gana, y si se le cae, pierde. Así es como el niño siente el juego, y ésta es su regla elemental: coger de nuevo la pelota cuando rebota de la pared. No se me ocurren juegos en los que no pueda señalarse la presencia de las características indicadas. Ni se me ocurren actividades que, poseyendo esos rasgos, no sean consideradas juegos.

Esto no significa que no pueda haber casos dudosos; pero la duda versará sobre la presencia de dichas características. Podemos dudar si cierta actividad es o no fundamentalmente recreativa, si tiene o no reglas. Tampoco se prejuzga cómo hayan de ser estas reglas. Pueden ser muy elementales, variables, e incluso es posible que se vayan modificando o inventando al tiempo que se juega, como Wittgenstein sugiere (83). La coincidencia de todos los juegos en las tres características mencionadas en la definición parece bastante general.

¿Es posible también, pese a las pretensiones de Wittgenstein, una definición del lenguaje? Definirlo lo hemos definido, en cuanto sistema de signos, en la sección 4.1, recurriendo a siete características o rasgos. Pero no es esto en lo que Wittgenstein piensa. Que todas las lenguas humanas coincidan en unas características muy generales que se hallan vinculadas a los caracteres biológicos que definen la especie humana, no es lo que Wittgenstein trata de refutar. El está pensando en una definición semántica del funcionamiento del lenguaje tal que reduce todos los posibles y variados usos lingüísticos a una única función. Es la definición de la función lingüística que se encuentra implícita en el *Tractatus*, y que podemos explicitar en los siguientes términos: el lenguaje es la totalidad

de las proposiciones; la proposición es el pensamiento expresado en sonidos; y el pensamiento es la representación lógica de los hechos posibles (*Tractatus*, 3, 3.1, 4, 4.001). La función lingüística queda así restringida a la función representativa o figurativa; la forma general de toda proposición es: así son los hechos; esto es lo único esencial para que una serie de signos formen una proposición (*Tractatus*, 4.5). Justamente contra esta reducción esencialista va dirigida toda la argumentación de Wittgenstein en su segunda época.

De aquí la comparación entre el lenguaje y los juegos. No hay una función lingüística única que defina al lenguaje como no hay —según cree Wittgenstein— ninguna característica única que defina al juego. Su posición es ahora pluralista: el lenguaje es, desde el punto de vista de su función, un conjunto de actividades o usos que forman una familia, tal y como ocurre con los juegos. Por ello, y a fin de evitar los errores y dificultades de la doctrina atomista, Wittgenstein recomienda sustituir la pregunta «¿Qué es el significado?» por esta otra: «¿Cómo se explica el significado?» (560, y *Cuaderno azul*, primera página). La respuesta es: enseñando a usar las expresiones. Y de ahí la conveniencia de sustituir la pregunta por el significado (*Bedeutung*) por una pregunta sobre el uso (*Gebrauch*) (561; se observará que en las *Investigaciones* Wittgenstein vuelve al sentido ordinario del término *Bedeutung*, abandonando el uso freguiano del término que hace en el *Tractatus*). Conveniencia que hace patente la siguiente propuesta: «Para una amplia clase de casos en los que utilizamos la palabra «significado», aunque no para todos los casos, se puede explicar dicha palabra así: el significado de una palabra es su uso en el lenguaje» (43, subrayado en el original).

Esta afirmación se ha citado frecuentemente de manera tergiversada, mencionando tan sólo su segunda parte: «el significado de una palabra es su uso en el lenguaje», y omitiendo la primera, que de modo explícito restringe el ámbito de su aplicación. La cuestión es: ¿en qué casos no procede la identificación del significado con el uso? La respuesta viene dada en la última parte del párrafo 43 que se acaba de citar, y reza así: «Y el significado de un nombre se explica, a veces, señalando a lo nombrado» (subrayado en el original). Podríamos, sin embargo, pensar: esto no nos impide decir que el significado de un nombre consiste en el uso que se hace de él para referirse al objeto o persona nombrados. Por tanto, también en este caso significado equivaldría a uso. Pienso que lo que Wittgenstein quiere dar a entender se capta, acaso, poniéndolo en relación con lo que acabamos de ver: en la mayor parte de los casos en los que hablamos del significado de las palabras éste puede explicarse hablando del uso que hacemos de ellas; pero cuando esas palabras son nombres propios, hay una manera más directa de explicar su significado, a saber: señalando al objeto designado por el nombre. La razón es que, en el caso de los nombres propios, el significado y la referencia coinciden, es decir, *Bedeutung* tiene al mismo tiempo su sentido ordinario y su sentido freguiano. En última instancia, no obstante, significado y uso coinciden también en este caso, pues

siempre podremos decir que el significado de un nombre propio es el uso que hacemos de él para referirnos a un objeto determinado individualizándolo entre los demás.

De esta analogía entre los juegos y los usos del lenguaje nace el gusto de Wittgenstein por un concepto al que recurre con frecuencia, aunque no siempre de modo uniforme, el concepto de juego de lenguaje o juego lingüístico (*Sprachspiel*). ¿Qué son los juegos lingüísticos? Yo lo pondría así: maneras particulares reales o imaginarias de usar el lenguaje, que tienden a mostrar cuáles son las reglas de un uso lingüístico. En general, puede decirse que son modelos simplificados de aspectos concretos del lenguaje. Cuando su argumentación lo requiere, Wittgenstein aconseja considerar el juego lingüístico de que se trate como si fuera un lenguaje primitivo total. Así, al comienzo de las *Investigaciones* considera unas palabras de San Agustín en las que se expresa una concepción del lenguaje claramente referencialista: las palabras funcionarían, fundamentalmente, como nombres de las cosas, y las oraciones serían combinaciones de nombres. Una concepción, como se ve, muy semejante a la del *Tractatus*. Entonces, Wittgenstein imagina un lenguaje que corresponda a esa concepción, y sugiere el siguiente: un albañil está construyendo un edificio con cuatro clases de piedras: bloques, pilares, losas y vigas; cuando necesita una piedra de una clase grita la palabra correspondiente, y su ayudante se la trae. Y añade Wittgenstein: «Concibamos esto como un lenguaje primitivo completo» (2). ¿Por qué? Porque al concebirlo así nos daremos cuenta de que tan simple sistema de comunicación no puede ser un lenguaje excepto en un sentido muy primitivo; por relación a nuestras lenguas actuales no constituye sino una pequeña porción de éstas. Por consiguiente, una concepción como la de San Agustín, como la del *Tractatus*, no pueden en manera alguna ser adecuadas a la totalidad del lenguaje humano. Sólo sirven para un lenguaje muy primitivo, porque son concepciones muy primitivas del lenguaje (1-3).

El modo de usar las palabras en el ejemplo del albañil podemos considerarlo —sugiere Wittgenstein— como uno de esos juegos por medio de los cuales aprenden los niños la lengua nativa. Y afirma a continuación: «Llamaré a estos juegos 'juegos lingüísticos', y hablaré a veces de un lenguaje primitivo como de un juego lingüístico. Y se podría llamar también juegos lingüísticos al proceso de nombrar las piedras y de repetir las palabras dichas por otro. (...) Asimismo, llamaré 'juego lingüístico' al todo formado por el lenguaje y las acciones con las que se halla entretelado» (7). Y en el *Cuaderno azul* dice de los juegos lingüísticos: «Son maneras de usar los signos, más simples que como los usamos en nuestro lenguaje cotidiano tan complicado. Los juegos lingüísticos son las formas del lenguaje con las que un niño comienza a hacer uso de las palabras. El estudio de los juegos lingüísticos es el estudio de las formas primitivas del lenguaje o lenguajes primitivos» (p. 17 del original y 44 de la trad. cast.). Por su parte, el *Cuaderno marrón* se inicia con una referencia a las palabras de San Agustín mencionadas antes y con el juego lingüístico del albañil. A partir

de éste, Wittgenstein va construyendo nuevos juegos a base de añadir pequeñas complicaciones sucesivas, como el uso de términos numerales o de nombrarse propios, y en un momento determinado se interrumpe para afirmar: «A sistemas de comunicación [como los anteriores] los llamaremos 'juegos lingüísticos'. Son más o menos semejantes a lo que llamamos juegos en el lenguaje ordinario. A los niños se les enseña su lengua nativa por medio de estos juegos, que en este caso poseen incluso el carácter entretenido de los juegos. No obstante, no estamos considerando los juegos lingüísticos que describimos como partes incompletas de un lenguaje, sino como lenguajes completos en sí mismos. Para mantener a la vista esta consideración, es con frecuencia útil imaginar que tan simple lenguaje constituye todo el sistema de comunicación de una tribu en un estadio primitivo de la sociedad» (p. 81 del original y 115-116 de la trad. cast.).

Las alusiones a la enseñanza infantil del lenguaje contienen, como se habrá notado, un eco de los años en que trabajó como maestro de escuela. Sin duda que esta experiencia debió de contribuir a la formación de su actitud crítica frente a la teoría atomista del lenguaje. Nótese asimismo el carácter más simple de los juegos lingüísticos en comparación con el lenguaje ordinario: justamente por eso sirven como instrumento de análisis para arrojar luz sobre el funcionamiento del lenguaje cotidiano. Y esto requiere separarlos suficientemente del resto de la compleja trama que compone el comportamiento lingüístico común; por esta razón conviene considerarlos como lenguajes totales, si bien, como es natural, de carácter muy primitivo. El juego de lenguaje, de otra parte, no son meramente las palabras o expresiones lingüísticas, sino el todo formado por éstas y las acciones con las que se hallan entretnejidas, según acabamos de ver. Esto no está explícito en los *Cuadernos*, pero sí en las *Investigaciones* (7). No obstante, es algo tan obvio para la definición de los juegos lingüísticos que no creo que pueda tomarse como una diferencia señalada entre ambas obras. De hecho, las diferencias que pueden encontrarse entre lo que Wittgenstein dice sobre los juegos de lenguaje en uno y otro libro, son, en mi opinión, diferencias de énfasis o de matiz, pero no diferencias sustanciales. Por cierto que esta caracterización del juego lingüístico como un todo formado a la vez por expresiones y acciones, ha sido interpretada por algunos de esta manera: como si Wittgenstein afirmara que el juego lingüístico es «el todo del lenguaje» más las acciones (así, Specht, *The Foundations of Wittgenstein's Late Philosophy*, III.2, pp. 42 y 45). Si esta interpretación fuera correcta, habría ahí ciertamente un nuevo concepto de juego lingüístico distinto del que hemos considerado. En esta extensión, el juego lingüístico por antonomasia, digámoslo así, sería el lenguaje entero, en cuanto conjunto de los juegos lingüísticos particulares. La afirmación de Wittgenstein en el párrafo 7 de las *Investigaciones* puede, no obstante, interpretarse como refiriéndose al todo formado por el lenguaje propio de un juego lingüístico junto con las acciones que son parte de tal juego (literalmente dice: «...das Ganze: der Sprache und der Tätigkeiten...»). Sea cual fuere la intención del autor en esta ambigua frase, lo cierto es que

el concepto de juego lingüístico usado por él sistemáticamente es el de modelo simplificado de un uso lingüístico, que hemos comentado anteriormente.

Con la noción de juego lingüístico están ligadas otras metáforas o comparaciones a las que Wittgenstein recurre frecuentemente. Así, por ejemplo, cuando afirma que las funciones de las palabras son tan distintas como las funciones de las herramientas que hay en un estuche (11), o como las manivelas en la cabina de una locomotora (12). Y no vale de nada suministrar una descripción tan general que resulte vacía, y decir, por ejemplo: todas las herramientas sirven para modificar algo (14), o bien todas las manivelas sirven para mover algo. Pues siempre habrá casos a los que nuestra descripción no alcance a menos que la forcemos de modo excesivo; en definitiva, ¿qué es lo que modifica un clavo, o una regla? (14). La idea básica es que el lenguaje es un instrumento, o mejor, un conjunto de instrumentos: las palabras, los conceptos, son instrumentos para jugar a una inmensa variedad de juegos lingüísticos (569, y *Cuaderno azul*, p. 67 del original y 101-102 de la trad. cast.). Lo que cuenta es el uso que hacemos de esos instrumentos, y para esto no basta fijarse únicamente en el instrumento, sino que hay que atender también a las acciones que acompañan a la pronunciación de las palabras (489), ya que hablar un lenguaje es parte de una actividad, de una forma de vida (19, 23). Lo fundamental aquí es que esas acciones nos van a revelar algo muy importante: que el uso de las palabras en el lenguaje, en los juegos lingüísticos, está sometido a reglas. Es la conexión regular entre los sonidos y las acciones lo que testimonia la existencia de un lenguaje (207). Son las reglas, por su parte, las que nos permiten hablar de corrección e incorrección en el uso del lenguaje, y las que asimismo nos permiten prever el comportamiento lingüístico de los demás. Aunque esto no significa que las reglas hayan de estar siempre perfectamente definidas ni que cubran todos los casos (82-85).

La doctrina anterior tiene, por lo pronto, claras consecuencias para una crítica devastadora de la teoría de la proposición que aparece en el *Tractatus*. Wittgenstein la afronta por derecho, y toma como uno de sus puntos más delicados la proposición 4.5 que antes hemos recordado: «La forma general de la proposición es: así son los hechos.» Este es —dice ahora, acaso con melancolía— el tipo de afirmación que uno se repite sin cesar, creyendo que está delimitando la naturaleza o esencia del objeto, la proposición; pero con ello, lo único que realmente hace es delimitar la forma a cuyo través contempla el objeto (114). Es decir: esta afirmación no expresa ningún descubrimiento acerca de la proposición, sino tan sólo el propósito de no llamar «proposición» más que a lo que tenga esa forma. En el fondo es tanto como decir que una proposición es todo aquello que puede ser verdadero o falso; o dicho de otro modo: que llamamos proposición a aquello a lo que, *en nuestro lenguaje*, le aplicamos el cálculo de las funciones veritativas (136, subrayado en el original). Parecería, entonces, que tenemos un concepto de verdad y falsedad, y que todo aquello que encaja con él es una proposición. Pero esto, dirá Wittgenstein, es una

mala imagen. Qué sea una proposición debe estar determinado, en un sentido, por las reglas de formación de oraciones en la lengua en cuestión, alemán, español, inglés, etc.; en otro sentido, por el uso del signo (o palabra) «proposición» en el juego lingüístico de que se trate. De este juego pueden formar parte los términos «verdadero» y «falso», y solamente de esta manera será el uso de estos términos parte del concepto de proposición (*loc. cit.*).

Lo que se quiere decir es esto: llamar «proposiciones» exclusivamente a las oraciones que pueden ser verdaderas o falsas, y aceptar como significativas sólo aquéllas es algo que puede estar justificado dentro de un determinado juego lingüístico como el cálculo veritativo-funcional, pero no es aceptable desde el punto de vista de los juegos en los que empleamos el lenguaje ordinario. En este último, lo que una proposición sea vendrá dado por el uso que hagamos del término «proposición», lo que sin duda habrá de remitirnos, al menos en parte, a las reglas gramaticales de formación de oraciones (recuérdese que, en este contexto, el término *Satz* que utiliza Wittgenstein puede significar tanto «oración» como «proposición»). Pero en el lenguaje ordinario no podemos pretender, sin falsearlo, que sólo son significativas las oraciones que pueden ser verdaderas o falsas. Esta mutilación del lenguaje que, como señalamos en el capítulo anterior, efectúa implícitamente el atomismo lógico, es ahora reconocida con toda justicia por el propio Wittgenstein.

La teoría isomórfica del lenguaje exigía hechos posibles que dieran sentido a las proposiciones. La cuestión clave, sobre la que llamé ya la atención en la última sección del capítulo anterior, es: ¿cuándo es posible un hecho? Allí aventuré una respuesta que, aun siendo, en mi parecer, perfectamente coherente con el *Tractatus*, no está clara y explícitamente formulada en éste: cuando sabemos en qué circunstancias es verdadera la proposición que lo representa. Esta salida epistemológica, como vimos, no parece vislumbrada por Wittgenstein de forma clara. Su posición permanecía encerrada en el lenguaje, como lo atestigua su afirmación: «lo que se puede describir, puede también suceder» (*Tractatus*, 6.362). O dicho de otro modo: es la proposición la que determina lo posible. Por eso Wittgenstein se pregunta ahora: «¿Depende enteramente de nuestra gramática a qué hemos de llamar (lógicamente) posible y a qué no, esto es, depende de lo que aquélla permita?» (520). La respuesta ha de ser negativa; la concepción isomórfica del lenguaje no puede, por sí sola, delimitar el ámbito de lo posible. La comparación del lenguaje con las representaciones figurativas es insuficiente: por lo pronto, no sabemos si lo estamos comparando con una representación histórica o con una pintura de género; la primera muestra algo que, supuestamente, ha ocurrido y, por tanto, puede ser falsa; ¿pero qué muestra la segunda? Wittgenstein sugiere esta respuesta: lo que muestra es a sí misma (522-523). Dicho en otras palabras: una representación histórica nos remite a un hecho, por comparación con el cual podemos decidir sobre su verdad o falsedad, pero una representación de género, como una obra de ficción, no remite a nada,

no podemos compararla con nada, y la concepción isomórfica le es, en rigor, inaplicable. Decir que la proposición es una representación o figura es irremediabilmente vago mientras no especifiquemos de qué clase de representaciones estamos hablando.

Una de las consecuencias de la teoría figurativa era que obligaba a hablar de hechos inexistentes, o si se prefiere, de la no existencia de hechos. Tal consecuencia es ahora ridiculizada por Wittgenstein: si digo que anoche no soñé, y esto es cierto, la proposición «Anoche soñé» será falsa. Pero para serlo ha de tener sentido. Y su sentido consiste en que expresa un hecho posible. ¿Significa esto que ha de haber algo así como el hueco que podría haber llenado un sueño? (448). Esta reflexión pretende mostrar cuán confusa es la idea de que la ausencia de un hecho debe contener la posibilidad de ese hecho.

El análisis atomista exigía, igualmente, la existencia de elementos últimos en la realidad que correspondieran a los signos más simples del lenguaje, a los nombres. Esto era producto de la concepción referencialista del significado. En el *Tractatus*, la identificación de estas entidades era considerada como tarea empírica más allá del alcance de la lógica. Wittgenstein se hace ahora cuestión de si la propia tarea, como tal, tiene buen sentido, y se pregunta: ¿cuáles son los constituyentes simples de una silla? ¿Los trozos de madera? ¿Los átomos, las moléculas? «Simple» quiere decir «no compuesto», y la cuestión es ¿en qué sentido de «no compuesto»? No hay un sentido absoluto de esta expresión. Lo que son partes simples para el carpintero no lo son para el físico o para el pintor. La pregunta acerca de las partes simples de un objeto solamente tiene sentido dentro de un determinado juego, ya que las palabras «simple» y «compuesto» las usamos en una infinidad de modos distintos (46-47).

Con esto cae por tierra todo el intento de reducir el lenguaje a nombres. La teoría referencialista construye el significado sobre la base de la relación entre el nombre y la cosa nombrada. ¿Pero en qué consiste esta relación? Esta relación no es absoluta; todo depende del juego lingüístico de que se trate. En un juego como el del albañil, la relación consiste en que, al oír el nombre, el ayudante trae una piedra del tipo que corresponda. En otros juegos, la relación puede ser diferente: puede consistir en que el nombre está escrito sobre el objeto, o en que alguien lo pronuncia cuando se señala a este último, etc. (37). El paroxismo de la concepción referencialista está en tomar como paradigma de los nombres los pronombres demostrativos, según hacía Russell. La posición de Wittgenstein sobre este punto es ahora rotunda: si no se quiere producir confusión, lo mejor es no decir que estas palabras nombran algo. (...) Pues tan extraña concepción proviene de una tendencia a sublimar la lógica de nuestro lenguaje, por así decirlo. «La respuesta apropiada es: llamamos 'nombre' a cosas distintas (...), a distintos tipos de usos de una palabra relacionados entre sí también de diferentes maneras, pero entre los cuales no se hallan los de la palabra 'esto'» (38). Wittgenstein piensa que, en esa concepción, el dar nombre se toma como un proceso oculto y misterioso que establece la conexión entre

la palabra y el objeto. Es la conexión que resulta cuando el filósofo se limita a contemplar el objeto repitiendo la palabra «esto», como a modo de bautismo, «extraño uso de esa palabra que ciertamente sólo aparece en la filosofía» (*ibidem*).

Wittgenstein ve la causa de esta doctrina en una confusión típica de la teoría referencialista, y es la que se da cuando se identifica el significado de un nombre con el referente, con el objeto nombrado. Entonces, si el referente desaparece, hay que concluir que el nombre pierde su significado. Pero es un hecho que hacemos uso de nombres, incluso de nombres propios, aun cuando no existan sus referentes. Luego tales nombres no serán auténticos nombres, y deben ser sustituidos, en el análisis, por otros términos. Tal era, como se recordará, el argumento de Russell al respecto. Para Wittgenstein, ahora, éste es un uso ilícito de la palabra «significado»: si el nombre «Miguel de Cervantes» hubiera perdido su significado cuando murió la persona, no podríamos hablar hoy de la muerte de Miguel de Cervantes. El nombre puede tener un uso y, por tanto, un significado, aunque haya desaparecido lo nombrado (39-40).

El defecto más profundo de una teoría referencialista es su primitivismo. Lo hemos visto a propósito de la cita de San Agustín que abre las *Investigaciones*: la concepción del significado como una relación de referencia es una idea primitiva del lenguaje, porque es la idea de un lenguaje primitivo (1-3). Desde el punto de vista de un lenguaje mínimamente complejo, en el que puedan hacerse descripciones de los objetos, nombrarlos no es más que la preparación para el juego lingüístico de describirlos, ni siquiera es un movimiento en este juego (49). No está, por consiguiente, justificado hacer de los nombres la base de una teoría del significado.

Toda la idea del análisis atomista pierde su sentido a efectos de la teoría del lenguaje. Porque los objetos pueden descomponerse en partes de diversas maneras, pero esto no implica que sus nombres hayan de descomponerse en forma análoga. Wittgenstein ridiculiza el análisis atomista intentando aplicarlo en la vida cotidiana: una escoba puede analizarse como compuesta de palo y escobilla, y un enunciado acerca de la primera puede sustituirse por un enunciado sobre dichas partes; ¿pero qué se gana con ello?, ¿es el segundo enunciado más claro que el primero?, ¿hay alguna razón para sustituir la orden «Tráeme la escoba» por «Tráeme el palo y la escobilla que está unida a él»? ¿tiene algún sentido afirmar que esta última orden subyace implícita en la primera? (60). La moraleja es ésta: desde el punto de vista del lenguaje cotidiano el análisis reductivo es inútil.

Para los atomistas, tal análisis estaba exigido, según vimos, por la lógica, en la que veían expresadas la estructura del lenguaje y del mundo. La lógica, como había subrayado Russell, era un lenguaje ideal. Wittgenstein protesta ahora contra lo que estas ideas sugieren: la lógica no es mejor, ni más perfecta, que el lenguaje ordinario. Es otra cosa, y resulta del todo confundente comparar éste con aquélla como un lenguaje imperfecto con un lenguaje ideal; lo más que puede decirse es que construimos len-

guajes ideales, pero éstos no representan ningún modelo al que haya de parecerse el lenguaje común (81).

Toda la crítica al análisis lógico-atomista descansa en la idea de que supone una falta de atención a la realidad del lenguaje ordinario que nos lleva a su falseamiento. Pero Wittgenstein nota que hay también una falta de imaginación que nos conduce a ver una ley, una necesidad, en el modo como, de hecho, usamos ciertas palabras (*Cuaderno azul*, p. 27 del original y 56 de la traducción castellana). Así, podemos preguntarnos: ¿qué significa negar? Y cabe recordar aquí la historia de los hechos negativos, que hemos seguido a través del atomismo lógico. Para evitar que el lenguaje embruje nuestro entendimiento, conviene ejercitar la imaginación. Imaginemos, sugiere Wittgenstein, un lenguaje que tuviera dos palabras distintas para negar, cuya única diferencia fuera que la doble negación con una de ellas produce una afirmación, mientras que el uso doble de la otra da como resultado una negación reforzada. ¿Tienen ambas palabras el mismo significado cuando se usan sin repetir? (556). Wittgenstein considera varias respuestas, dando a entender que no tiene por qué ser una de ellas mejor que la otra. Podemos decir que, puesto que en conjunto las dos negaciones tienen usos distintos, tienen asimismo distintos significados, aunque esta diferencia no tenga consecuencias cuando se usan sin repetir. Podemos señalar, por el contrario, que básicamente funcionan del mismo modo en los juegos lingüísticos, excepto por lo que respecta a esa pequeña peculiaridad que es cosa sin importancia debida a la costumbre, y subrayaremos, en este caso, que ambas palabras se enseñan de la misma manera, concluyendo, por tanto, que tienen el mismo significado. En fin, podemos mantener que las dos palabras expresan diferentes ideas, mostrando que llevan asociadas imágenes diversas: la primera negación cambia el sentido 180°, por así decirlo, y por ello, al duplicarse, vuelve al sentido original; la segunda es como un movimiento de cabeza, que al repetirse, se refuerza. Cualquiera de estas explicaciones da buenas razones en su favor. Podemos, por consiguiente, concluir con el mismo derecho, o bien que ambas negaciones poseen el mismo significado, o bien que tienen significado distinto.

Hasta aquí hemos visto la nueva concepción del significado que tiene Wittgenstein y el contexto de la crítica al atomismo lógico en el que aparece aquélla. Lo fundamental no es la relación de referencia entre las palabras y las cosas, sino los varios usos que hacemos del lenguaje. Entre estos usos no existe ninguna característica común, sino relaciones de índole diversa que forman como una red. No hay lugar, pues, para una definición del lenguaje. ¿Pero sería posible hacer una tipología de esos usos y reseñar al menos sus clases principales? Ni siquiera esto. Wittgenstein lo niega explícitamente y, en su lugar, se limita a dar ejemplos. Así, se pregunta: «¿Cuántas variedades de proposiciones hay? ¿Por ejemplo, aserción, pregunta y mandato?» (23). Y se replica: «Hay incontables variedades, innumerables maneras distintas de usar eso que llamamos 'signos', 'palabras' y 'oraciones'. Y esta multiplicidad no está fijada, y dada de una vez por

todas, sino que nuevos tipos de lenguaje, nuevos juegos lingüísticos, por así decirlo, nacen y otros envejecen y se olvidan.»

Lo primero que llama aquí la atención es el paso de una cuestión a otra aparentemente distinta. La pregunta versa sobre tipos de proposición, como son la aserción, la interrogación y el imperativo. Pero Wittgenstein contesta hablando de la variedad de los tipos de juegos lingüísticos. Este cambio de tema, que los comentaristas suelen subrayar, es, no obstante, del todo coherente con la posición de Wittgenstein. Hemos visto que, en contraste con el *Tractatus*, la unidad de análisis lingüístico no es ahora la proposición sino el uso lingüístico, y que éste queda reflejado en el modelo que es el juego de lenguaje. La reflexión de Wittgenstein, debidamente explicitada, viene a ser ésta: lo que interesa no es cuántos tipos de proposiciones hay, sino cuántas variedades de usos del lenguaje existen; y la respuesta es que éstas son innumerables, y que no pueden limitarse *a priori* porque están siempre en proceso de cambio. Wittgenstein no afirma, por consiguiente, que haya más o menos clases de proposiciones que las tres señaladas. Simplemente sustituye esa cuestión, más propia de un enfoque lógico, por la cuestión que ahora le preocupa: cuántas maneras distintas tenemos de usar las proposiciones, sean cuales fueren los tipos de éstas. Y su respuesta es: son innumerables y, además, no están dadas de una vez por todas. Dicho de otro modo: la clase de los usos lingüísticos es amplísima y está en perpetuo cambio. Wittgenstein suministra a continuación una larga lista de ejemplos, a saber: «dar órdenes y obedecerlas; describir un objeto, bien por su apariencia, bien dando sus medidas; construir un objeto a partir de una descripción o de un dibujo; informar sobre un acontecimiento; hacer suposiciones sobre ese acontecimiento; formular una hipótesis y comprobarla; representar los resultados de un experimento por medio de tablas y diagramas; inventar una historia, y leerla; hacer teatro; cantar jugando al corro; adivinar acertijos; hacer un chiste, y contarlo; resolver un problema práctico de aritmética; traducir de una lengua a otra; pedir, dar las gracias, maldecir, saludar, rezar» (23; se encontrará una lista análoga, pero más breve, en el *Cuaderno azul*, pp. 67-68 del original y 102 de la trad. cast.). Y para acabar, agrega: «Es interesante comparar la multiplicidad de las herramientas del lenguaje, y de los modos de usarlas, la multiplicidad de los tipos de palabras, y de los modos de usarlas, la multiplicidad de los tipos de palabras y de proposiciones, con lo que los lógicos han dicho sobre la estructura del lenguaje (y también el autor del *Tractatus Logico-Philosophicus*).» Estas últimas palabras reinciden en la suerte de confusión que ya se había insinuado al principio del párrafo: ¿en qué quedamos, se trata de tipos de proposición o de tipos de uso? Wittgenstein parece no distinguir entre lo uno y lo otro, pero los ejemplos de su lista, ciertamente, son ejemplos de usos, no de clases de proposiciones. Los ejemplos citados lo son de situaciones en las que hacemos algo utilizando el lenguaje; y nótese que aquí «utilizar el lenguaje» no significa tan sólo emitir locuciones sino también recibirlas: obedecer órdenes es un uso del lenguaje en la medida en que implica entender las órdenes, y lo mismo pue-

de decirse de la construcción de un objeto a partir de una descripción, o de la comprobación de una hipótesis. Todo esto son usos del lenguaje por cuanto requieren la comprensión del mismo. Pero es claro que los ejemplos de la lista no lo son de tipos de proposiciones, pues las mismas proposiciones pueden utilizarse al traducir, al hacer teatro, al hacer chistes, al narrar historias, etc.

El caso es que, desde el punto de vista de la teoría de la proposición, hay dos críticas diferentes que pueden hacerse al *Tractatus*. Una diría: es ilegítimo llamar «proposición» únicamente a las expresiones que pueden ser verdaderas o falsas; hay otros tipos de proposiciones en el lenguaje ordinario, tan legítimas como aquéllas y, por consiguiente, debemos elaborar una teoría del significado que valga para todas y no sólo para las primeras. Esta crítica ya la hemos visto formulada por Wittgenstein, por ejemplo, en el parágrafo 136, si no exactamente en estos términos, sí en términos parecidos. Para esta crítica es relevante, sin duda, la cuestión de determinar los tipos de proposiciones, cuestión que, sin embargo, Wittgenstein no aborda. La otra crítica rezaría así: hablar de proposiciones, como hacen los lógicos, es hacer injusticia al lenguaje ordinario incluso aunque se acepten otras clases de proposiciones además de las que pueden ser verdaderas o falsas; el lenguaje tal y como lo utilizamos posee una riqueza y una complejidad que desborda los límites de una teoría de las proposiciones, y que tan sólo se hace patente cuando, en lugar de considerar aquéllas, se atiende a los usos, esto es, a las formas de utilizar el lenguaje. Esta segunda crítica es más bien la que se desarrolla en el parágrafo 23 que acabo de comentar. Por eso Wittgenstein afirma en el parágrafo siguiente que quien no tenga a la vista la multiplicidad de los juegos lingüísticos se sentirá tentado de preguntarse cosas como ¿qué es una pregunta?, ¿es la afirmación de un estado de duda, o su descripción?, ¿es una petición? Y Wittgenstein recomienda aquí considerar cuántos tipos distintos de descripción tenemos (24). Lo que equivale a decir que no es adecuado plantear el problema de si las preguntas y las descripciones constituyen dos tipos de proposiciones o bien las primeras pueden reducirse a las segundas; lo correcto es atender a los usos que se hacen de unas y de otras, y entonces se comprobará que, por lo pronto, llamamos «descripciones» a tan diferentes usos del lenguaje, que aun cuando asimiláramos las preguntas a descripciones no habríamos, con ello, arrojado luz suficiente sobre el carácter de las primeras. Y esto implica que lo que interesa no son los tipos de proposiciones sino las clases de sus usos.

¿Cuáles son las diferencias más relevantes entre una clasificación de las proposiciones y una tipología de los usos? Importa obtener claridad sobre este punto si se quiere entender el desarrollo de la filosofía del lenguaje después de Wittgenstein. Puesto que una proposición según el *Tractatus* era toda oración que pueda ser verdadera o falsa, y ya que el contexto general en el que ahora nos movemos es el de la crítica a esa obra, los diferentes tipos de proposiciones corresponderán a los diversos modos de relación que hay entre las oraciones y la realidad extralingüística. Esto es,

la distinción entre tipos de proposiciones será una distinción semántica. Por esta vía, pues, la crítica al *Tractatus* se mantendría dentro del ámbito semántico, que es el ámbito propio de esta obra. Esa vía es justamente la que han seguido los neopositivistas y otros pensadores más o menos influidos por ellos. Algunos filósofos del lenguaje ordinario han estado también más próximos a ese enfoque que al que es típico del segundo Wittgenstein. El enfoque característico de éste consiste, como hemos visto, en atender a los usos que hacemos del lenguaje y, por consiguiente, a los propósitos de los hablantes, y a todas las demás circunstancias que rodean al comportamiento lingüístico, o dicho en palabras de Wittgenstein, en atender a la forma de vida en la que se hace uso de las palabras. Una clasificación de los usos será, por ello, una clasificación pragmática, y será, en consecuencia, más numerosa, abigarrada y multiforme que una clasificación puramente semántica como la primera. De aquí el contraste entre distinguir aserciones, mandatos y preguntas, por ejemplo, que serían tipos de oración, y distinguir entre contar chistes, traducir, narrar, hacer teatro, formular hipótesis, etc., etc., que constituyen tipos de uso.

En resumen, la nueva teoría de Wittgenstein sobre el lenguaje se basa en la idea de que lo importante no es una teoría de las proposiciones, sino una descripción de los usos lingüísticos. O lo que es lo mismo, aunque Wittgenstein no lo exprese así: propone sustituir la semántica por la pragmática, y por una pragmática al parecer empírica y particularmente vaga. Esto puede producir el sentimiento de que, con ello, la filosofía del lenguaje queda sin justificación y se disuelve en una especie de lexicografía aplicada. Si así fuera, desaparecería ese trascendentalismo lingüístico que había en el *Tractatus* y que, en la primera sección del capítulo anterior, señalé como característico de la filosofía analítica. La cuestión, en definitiva, es: ¿qué sentido filosófico tiene una descripción de los usos del lenguaje?

7.3 La crítica de los lenguajes privados

Hay un grupo de usos del lenguaje a los que Wittgenstein ha prestado particular atención: aquellos que tienen que ver con las experiencias internas, con los fenómenos mentales. Su propósito es mostrar que una teoría referencialista no puede dar cuenta del significado de las expresiones que se refieren a tales fenómenos, y, en esta medida, su tratamiento de este tema puede considerarse parte de su crítica general a ese tipo de teorías; en el curso de su desarrollo, sin embargo, Wittgenstein aprovechará la ocasión para arrojar nueva claridad sobre el concepto de lenguaje, desbordando aparentemente los límites de ese descriptivismo al que acabo de hacer mención.

Podemos partir de la siguiente pregunta: ¿cómo sabemos lo que significan las expresiones que hacen referencia a las experiencias internas? Por ejemplo: ¿cómo sabemos lo que significa la palabra «dolor»? Muchos, aplicando, aun inconscientemente, la teoría referencialista, contestarían:

cada cual lo sabe por su propio caso, por su propia experiencia del dolor; nadie puede sentir el dolor ajeno, luego no hay más que una posibilidad para llegar a conocer lo que significa el término «dolor»: sentir dolor. Para Wittgenstein, esto constituye una generalización irresponsable a partir del caso propio, y va a intentar mostrar que ni siquiera en el caso propio es posible dar significado a una palabra sobre la base exclusiva de conectarla con la experiencia interna (293). Lo hará por medio de una comparación. Supongamos que cada persona tuviera una caja conteniendo algo a lo que se denomina «escarabajo», y que nadie pudiera mirar en caja ajena. Cada cual dice que sabe lo que es un escarabajo mirando en su caja, aunque es perfectamente posible que cada uno tenga en su caja un objeto distinto al de los demás, e incluso que el objeto esté en perpetuo cambio. Y la pregunta de Wittgenstein es: ¿podría tener esa palabra un uso para esas personas? A lo cual responde: si lo tuviera, no sería el uso que consiste en la designación (*Bezeichnung*) de un objeto; el objeto no pertenece, en este caso, al juego lingüístico que se jugaría con la palabra «escarabajo». La razón parece clara: puesto que nadie sabe lo que tienen los demás en su caja, el uso de esa palabra no puede consistir en designar su contenido. Ni siquiera cabe pensar que «escarabajo» signifique algo así como «lo que hay dentro de cada caja», pues podría ocurrir que alguna caja estuviera vacía. Y concluye Wittgenstein: «Esto significa que si se construye la gramática de la expresión de la sensación según el modelo objeto y designación, el objeto queda fuera de nuestra consideración como irrelevante» (293).

El uso que se hace de una expresión es común, intersubjetivo, y ha de estar, por tanto, en conexión con objetos, fenómenos o manifestaciones que sean igualmente intersubjetivos, comunes. Si la palabra «escarabajo» tiene un uso para las personas del ejemplo, debe haber alguna conexión entre ella y la realidad intersubjetiva. La palabra no puede llegar a tener un uso si se conecta exclusivamente con algo que sea enteramente privado y exclusivo de cada cual. Trasladado al caso del término «dolor», y, en general, de las expresiones que se refieren a sensaciones y a experiencias internas o mentales, ello quiere decir que el significado de dichas expresiones no se puede reducir a una relación de referencia. O dicho de otra forma: que las teorías referencialistas, como la del *Tractatus*, no pueden explicar el significado de las palabras que se refieren a experiencias internas. La crítica al *Tractatus* sigue siendo aquí el trasfondo de los pensamientos de Wittgenstein. Importa, pues, tener bien claro que hay dos cosas que éste no está afirmando: primera, que no existan experiencias internas; segunda, que no se pueda hablar de ellas o expresarlas por medio del lenguaje. Wittgenstein no niega que cada cual siente su dolor y no el de los demás, como tampoco niega que, en su ejemplo, cada cual sólo tiene acceso a lo que hay en su propia caja. Wittgenstein tampoco dice que no se pueda hablar del dolor o expresarlo lingüísticamente, igual que no dice que no se pueda hablar de lo que hay en las cajas. Lo único que afirma es que las palabras con las que hablamos de esas experiencias, de lo

que cada uno tiene en su caja, no pueden considerarse como meramente designativas, esto es, al modo de los nombres propios de Russell o del *Tractatus*. ¿Por qué? La razón a la que ya he aludido más arriba, es porque «un proceso interno requiere criterios externos» (580).

¿Qué quiere decir esto último? Wittgenstein había distinguido en el *Cuaderno azul* entre criterios y síntomas (p. 24 s. del original y 53 s. de la trad. cast.). El criterio para afirmar la existencia de un fenómeno viene dado por la definición de este fenómeno. Wittgenstein ofrece este ejemplo: si la medicina define las anginas como una inflamación causada por cierto bacilo, la presencia de este bacilo en la sangre es el criterio para decir (y por tanto, para saber) que alguien tiene anginas. Un síntoma, en cambio, es un fenómeno que, por experiencia, hemos comprobado que coincide con el fenómeno que es criterio. Así, la inflamación de la garganta es un síntoma de anginas. La distinción entre el criterio y los síntomas de un fenómeno es, no obstante, arbitraria y vaga en la mayor parte de los casos, como Wittgenstein reconoce; esto sólo demuestra que el lenguaje ordinario carece de la exactitud propia de un cálculo o de un lenguaje artificial. Lo importante del criterio es que suministra la definición del término correspondiente. El criterio de las anginas suministra la definición del término «anginas». Pues bien, decir que un proceso interno requiere criterios externos es tanto como decir que los términos que se refieren a procesos internos han de ser definidos recurriendo a manifestaciones externas.

Volviendo a nuestra pregunta inicial, esta doctrina implica que sabemos lo que significa la palabra «dolor», no a causa de nuestra experiencia del dolor, sino en base a aquellas manifestaciones externas de dolor que constituyen el criterio para decir que alguien tiene un dolor, y que, por ello mismo, están mencionadas en la definición del dolor. A esta posición se la ha denominado, a veces, «conductismo lógico». Wittgenstein argüirá a su favor siguiendo, fundamentalmente, dos vías diferentes. De un lado, examinará el funcionamiento de conceptos tales como los de comprender, significar, imaginar, etc. De otro, mostrará que un lenguaje privado, esto es, un lenguaje cuyas palabras adquieran significado sin recurrir a objetos o fenómenos externos, es imposible. Veamos brevemente algunos ejemplos de los argumentos de Wittgenstein.

¿En qué consiste comprender? Wittgenstein recomienda: «¡No pensemos en la comprensión como un 'proceso mental'!»! Pues ésta es la forma de hablar que nos confunde. Preguntémonos en cambio: ¿en qué clase de casos, en qué tipo de circunstancias, decimos: «Ahora sé cómo seguir»? (154). La idea es que nuestra forma de hablar, nuestra terminología, nos confunde y nos plantea los problemas. Como he recordado antes, ya en el *Cuaderno azul* describía a quien se halla filosóficamente perplejo como alguien que ve una ley en el modo de usar las palabras (p. 27 del original y 56 de la trad. cast.), y en las *Investigaciones* describe la filosofía como una lucha contra el embrujamiento de nuestra inteligencia por el lenguaje (109). Es la expresión «proceso mental» la que nos crea el problema filo-

sófico. Lo que hay que hacer es atender a los juegos lingüísticos en los que utilizamos la palabra «comprender», y entonces comprobaremos que se trata, al menos en muchos casos, pero casos típicos, de circunstancias en las cuales podríamos haber dicho «¡Ya sé cómo seguir!» en lugar de «¡Ya lo comprendo!» Wittgenstein considera el caso de una serie numérica de la que se trata de averiguar la ley de sucesión. Y aquí, ciertamente, el criterio de que alguien ha comprendido la serie es que sabe continuarla (151-153). La comprensión no tiene por qué ser una especie de oculto proceso que ocurra paralelamente al proceso de escribir la serie numérica en cuestión. El problema filosófico de los procesos y estados mentales surge, precisamente, porque nos ponemos a hablar de estados y procesos sin decidir acerca de su naturaleza, con lo cual esta analogía con los procesos y estados físicos, que pretendía ser iluminadora, se torna fuente de problemas (308). Y parecería que se quisiera negar los procesos mentales, cuando lo que se quiere negar es, propiamente, «un proceso aún no comprendido en un medio todavía inexplorado» (*ibidem*). Wittgenstein, pues, no pretende negar que haya procesos mentales (154, 308). Lo que niega es que al llamarlos «procesos» estemos, por eso, más cerca de entenderlos; y lo que niega también es que, en ese sentido, comprender sea un proceso mental. Más bien, hay procesos mentales que acompañan a la comprensión (154). ¿Cuáles? Wittgenstein no es explícito al respecto, pero cuando da ejemplos de procesos mentales cita el aumento o la disminución de un dolor, o el oír una melodía o una expresión verbal (*ibidem*). Parece, pues, que son procesos mentales las sensaciones y percepciones, esto es, fenómenos en los que hay un estímulo externo orgánico, pero no lo es, en cambio, la comprensión. ¿Qué es la comprensión? Una capacidad de comportamiento que cumple con determinadas condiciones según cuál sea el objeto comprendido. En el ejemplo de la serie numérica, comprenderla es ser capaz de continuarla. Uno de los posibles criterios de comprender es ser capaz de proseguir.

Por lo mismo, no tiene sentido pensar que significar o querer decir algo por medio de una expresión lingüística sea hacer dos cosas al tiempo, emitir sonidos y desarrollar un cierto proceso mental que los acompañe. Utilizar el lenguaje correctamente, esto es, con la intención significativa adecuada a las palabras empleadas y a la ocasión en que se emplean, es simplemente dominar la técnica de emplear unos signos según ciertas reglas, pero no tiene por qué suponer la capacidad de emparejar el proceso externo de emitir sonidos con un supuesto proceso interno de darles significado; ¿pues cómo podríamos justificar que este proceso interno se realiza correctamente? (19-20). Pongamos que estoy sintiendo un dolor y, al propio tiempo, escuchando una melodía. Y digo: «Se va a terminar pronto». Puedo referirme al dolor o a la música: ¿pero en qué consiste la diferencia? (666). Wittgenstein tratará de mostrar, por medio de comparaciones, que no es posible buscar esa diferencia en una especie de ostensión o señalamiento, pues no hay nada que señalar cuando se trata del juego lingüístico con palabras como «querer decir» o «significar» (667-675). Una de sus

comparaciones es particularmente iluminadora: imaginemos que estoy hablando por teléfono y que le comunico a mi interlocutor «Esta mesa es demasiado alta», señalando al tiempo a la mesa que tengo al lado. ¿Qué función tiene el gesto de señalar? (670). Es claro que ninguna; si mi interlocutor no puede ver qué es lo que señalo, mi gesto no añade nada a mis palabras. Es decir: no se logra nada explicando el significar, el querer decir, como si esto consistiera en acompañar las palabras por un gesto mental de señalar. En el sentido en que un gesto corporal puede acompañar a mis palabras, un gesto mental no puede hacerlo (673). En suma: significar no es un proceso mental, no es una afección de la mente (675-676).

Pero la vía que Wittgenstein ha explorado más a fondo, a fin de elucidar en qué consiste el significado de las expresiones que se refieren a experiencias internas, es la vía que conduce a mostrar que no es posible un lenguaje privado. Wittgenstein llama «lenguaje privado» a aquel lenguaje «cuyas palabras han de referirse a lo que sólo puede conocer el hablante, a sus sensaciones inmediatas y privadas, de tal manera que nadie más pueda entender su lenguaje» (243). Según esta definición, un lenguaje es privado en cuanto sean privados sus referentes, esto es, en cuando sus palabras se refieran a algo que solamente pueda conocer el que lo usa, como sería el caso, a primera vista, de un lenguaje acerca de las experiencias internas del hablante. Pero en realidad, esto no basta. Podría ocurrir que el hablante se refiriera a sus sensaciones y sentimientos por medio de palabras que poseyeran significado en virtud de que son usadas para hablar de los objetos externos, y las cuales se usaran analógicamente para hablar de los estados mentales. Este lenguaje no sería privado, en el sentido de Wittgenstein. Para que un lenguaje sea privado no basta que sus referentes lo sean; han de ser privadas también sus reglas, y esto quiere decir que las expresiones de ese lenguaje habrán de poseer significado exclusivamente en virtud de una conexión directa entre la palabra y la experiencia interna designada por ella. No hay que olvidar que persiste en el trasfondo la crítica a la teoría referencialista como motivo. Estamos en un caso como el de la palabra «escarabajo», que ya vimos; en la medida en que se construya como meramente designativo, ese término es privado, pues se refiere a algo que sólo puede conocer el propio hablante, y adquiere su referencia únicamente sobre la base de una conexión directa con el objeto privado.

Una característica de las tendencias que Wittgenstein está criticando es la que consiste en subrayar con exceso la privaticidad de las experiencias internas. Así, cualquier pensador con tendencia al solipsismo propenderá a desconfiar de los sentidos externos, acentuando, en cambio, la inmediatez y privaticidad de sus experiencias internas, a las que otorgará primacía epistemológica, manteniendo que sólo el conocimiento de los estados mentales propios es seguro e infalible. Wittgenstein, que sin duda está aquí arguyendo en contra de una tendencia propia, tratará de disipar tan exageradas y vacuas pretensiones. Cuando alguien afirme que sólo él mismo

puede saber si realmente tiene o no dolores, y que los demás únicamente pueden suponerlo, Wittgenstein responderá: «si usamos la palabra 'saber' como se usa normalmente (¡y como deberíamos usarla si no!), entonces los demás saben con frecuencia si yo tengo o no dolores» (246). La cuestión es que no vale pretender que la palabra «saber» haya de recibir algún significado peculiar, propio de la filosofía, que nos permita decir que *yo sé*, con una certeza que los demás no pueden alcanzar, que tengo dolores. La afirmación «Sé que tengo un dolor» no tiene sentido excepto, quizá, como chiste, pues no puede significar otra cosa sino simplemente que tengo un dolor (*ibidem*). La falta de certeza queda aquí excluida por definición; usamos de tal modo palabras como «sensación», «dolor» o «intención», que el sujeto al que le atribuimos tales experiencias, si las tiene, no puede ignorarlo (247). Por eso, la afirmación «Las sensaciones son privadas» puede compararse con «Un solitario es un juego al que juega uno sólo» (248). La privaticidad de las experiencias internas se deduce de la propia definición de experiencia interna, y no tiene sentido intentar sacar de aquí consecuencias para una teoría del conocimiento. Solamente puede decirse que se conoce algo cuando sería posible desconocerlo; puesto que no tiene sentido afirmar de alguien que no sabe que tiene un dolor, tampoco tiene sentido decir que lo sabe. Y por ello mismo es correcto afirmar «Sé lo que estás pensando», pero es un error decir «Sé lo que estoy pensando» (*Investigaciones*, II, xi, pp. 221-222).

¿De qué manera usamos, entonces, las palabras para referirnos a nuestras vivencias? Ya hemos visto que no es posible construir esa relación como una mera designación sobre la base de una conexión directa u ostensión. ¿De qué otras alternativas disponemos? La más inmediata es la que explica el significado de tales palabras vinculándolas con lo que Wittgenstein llama las «expresiones naturales de las sensaciones» (256). Consideremos, por ejemplo, el dolor: son sus expresiones naturales el lloro, el quejido, el gesto (244). A fin de entender este punto mejor, imaginemos que los seres humanos no exteriorizáramos el dolor. ¿Podríamos enseñar a alguien lo que significa la expresión «dolor de muelas»? Bien, es de suponer que nuestras lenguas carecerían de expresiones así. Pero y si alguien inventara nombres para sus vivencias, ¿podría ser entendido por los demás? Estamos de nuevo, como se habrá notado, en el caso de las cajas y los «escarabajos», y la respuesta de Wittgenstein a ambas preguntas es, por consiguiente, negativa. Sin manifestaciones exteriores a las que recurrir, no hay posibilidad de entender los nombres de experiencias internas que otros pudieran usar. Resta tan sólo una cuestión más: aunque el hablante no pueda explicar a los demás lo que significan las expresiones que utiliza para designar sus vivencias, ¿podría acaso entenderlas él mismo en la soledad de su conciencia? Con esto hemos llegado al último reducto del solipsismo. El solipsista puede renunciar a comunicarse con los demás, pero se aferrará entonces al monólogo, al diálogo silencioso de la conciencia consigo misma, llevado en un lenguaje sólo comprensible para ella. Ahora bien, ¿qué quiere decir dar nombre a, o designar, una vivencia? Al pre-

guntarnos por su posibilidad en los términos anteriores, olvidamos que dar nombre a algo, en el lenguaje ordinario, requiere determinados presupuestos, entre ellos, fundamentalmente, un lenguaje en el que los nombres funcionen como tales. Pero no puede haber lenguaje si no hay una forma de vida, y, por tanto, manifestaciones externas que acompañen al uso de las palabras (257).

Una nueva comparación remachará definitivamente la argumentación en contra de la idea de un lenguaje privado. Supongamos que intento llevar un diario en el que quede constancia de la recurrencia de determinada sensación, y a estos efectos escribo el signo «S» cada día que tengo dicha sensación. ¿He dado nombre, con ello, a mi vivencia? ¿He dado significado, de esa manera, al nuevo signo? Utilizar un signo supone poder distinguir cuándo se usa correctamente y cuándo no, ¿pero cómo podríamos hacer esta distinción en semejante hipótesis? ¿De qué modo sería posible distinguir entre un uso correcto y un uso incorrecto del signo «S»? La realidad es que «en ese caso carecemos de un criterio de corrección, y podríamos decir que será correcto lo que me parezca correcto» (258). Pues, en efecto, cuando quiera que tenga una sensación, si me parece que ésa es la sensación que he decidido llamar «S», efectuaré la anotación correspondiente, ¿pero cómo sabré que ésa es la sensación «S», o lo que es lo mismo, qué garantía tendré de que estoy usando la expresión «S» correctamente y de acuerdo con mi primitiva intención? El signo «S» solamente puede tener significado en cuanto parte de algún juego lingüístico, y aquí no hay tal juego, porque no hay ninguna actividad con la que encaje el uso de ese signo (261). El habla es una actividad sometida a reglas, pero dado lo que significa la palabra «regla» no tiene sentido pensar que una sola persona y una sola vez en su vida puede seguir una regla (199). Actuar según una regla es una práctica, y no basta creer que se está cumpliendo una regla para que se la esté cumpliendo realmente; por ello no es posible seguir una regla de forma privada, pues no habría manera de distinguir, entonces, entre creer que se estaba siguiendo una regla y seguirla efectivamente (202). En suma: un lenguaje privado no es posible porque no se podría establecer diferencia entre la corrección y la incorrección en su uso; o dicho de otro modo: porque no existiría posibilidad de determinar si se estaba o no siguiendo reglas y cuáles.

El resultado de toda la argumentación precedente es que las expresiones que designan experiencias internas, vivencias, no constituyen un lenguaje privado ni reciben su significado de una mera conexión directa e inmediata entre la palabra y la vivencia. ¿Cómo se explica el significado de esas expresiones? De hecho, y brevemente, ya hemos respondido a esta pregunta un poco más arriba. Esas expresiones lingüísticas se hallan vinculadas a la expresión natural de las vivencias (256). ¿De qué manera? Wittgenstein, siguiendo la sugerencia que hemos visto en la sección anterior, se pregunta por el proceso de aprendizaje: ¿cómo aprendemos el significado de los nombres de las sensaciones? Y sugiere esta posibilidad: se conectan las

palabras con la expresión primitiva y natural de la sensación, y ulteriormente la sustituyen (244). Así, un niño se ha hecho daño y llora. El lloro es la expresión natural, primitiva, de lo que siente. Pero los adultos le hablan, y el niño aprende de ellos exclamaciones y frases que en lo sucesivo acompañarán a sus quejidos, e incluso los sustituirán. Y concluye Wittgenstein: «Los adultos le enseñan al niño un nuevo comportamiento expresivo del dolor» (*ibidem*). La idea, pues, parece ser que al decir «Me duele» no estoy designando mi vivencia sino expresándola, igual que podría expresarla por un gesto de dolor. En sentido parecido, había escrito en el *Cuaderno azul* que la diferencia entre «Tengo dolores» y «Fulano tiene dolores» corresponde a la diferencia entre quejarse y decir que alguien se queja (p. 68 del original y 103 de la traducción).

A esto se podría replicar que decir que alguien tiene dolores no es lo mismo que afirmar que se queja. Mas en defensa de Wittgenstein sería de justicia señalar que él no dice que sea lo mismo, sino que lo uno corresponde a lo otro, o más exactamente, que la primera de las diferencias citadas corresponde a la segunda. Sea como fuere, si la propuesta de Wittgenstein consiste en considerar la autoadscripción de una vivencia como expresión de esa vivencia, entonces hay que decir que es errónea por exagerada. La afirmación «Me duele aquí» no es siempre ni necesariamente una expresión de dolor; en muchos casos puede ser simplemente una información que damos, por ejemplo, al médico en la consulta o al amigo que nos visita. Que Wittgenstein intenta negar eso, lo parece por las citas mencionadas tanto como por lo siguiente: «Afirmer 'Tengo dolor' no es un enunciado acerca de una persona particular más de lo que lo es quejarse.» (*Cuaderno azul*, p. 67 del original y 101 de la trad. cast.) A la luz, sin embargo, de ciertas reflexiones en las últimas páginas en la segunda parte de las *Investigaciones*, más bien resulta que estaría dispuesto a admitir que afirmaciones como las citadas pueden constituir, en ciertos contextos, descripciones y no propiamente expresiones. Así, se pregunta sobre la frase «Tengo miedo»: «¿Qué es esto? ¿Un grito de temor? ¿O quieres comunicarme cómo te sientes? ¿O es una reflexión sobre tu estado actual?» Y después de examinar diversas formas en que podría usarse esa frase, dice: «Describir mi estado mental (por ejemplo, de miedo), eso lo hago en un contexto determinado.» Y más abajo: «Por cierto que no siempre decimos de alguien que se queja porque afirme que tiene dolores. Pues las palabras 'Tengo dolores' pueden ser una queja, y también pueden ser otra cosa.» (II, ix, páginas 187-189.)

Es indudable que lo que Wittgenstein dice sobre este punto es vago y se encuentra falto de precisiones. Que las expresiones mentales se aprenden en el contexto de comportamientos típicos no parece dudoso. Que muchas de ellas se usan normalmente, o al menos con frecuencia, como expresión de la vivencia a la que se refieren, es patente. Pero si también se usan como descripciones, o como designaciones, de dichas vivencias, queda por explicar cómo adquieren este significado y de qué manera se relaciona este uso con

su utilización expresiva. Si «Tengo un dolor» puede usarse, como admite Wittgenstein, para describir un estado mental, hay que explicar cómo ha llegado a tener este uso, este significado, o lo que tanto da, hay que explicar cómo hemos aprendido a usar esa expresión con esta función descriptiva. Hay, por ejemplo, la siguiente posibilidad: podríamos decir que, en su uso descriptivo, las expresiones mentales tienen como referencia las vivencias o estados mentales, y que su sentido deriva de las manifestaciones externas de dichas vivencias, esto es, de lo que Wittgenstein ha llamado sus «expresiones naturales». Así, la referencia de «dolor» es una vivencia de determinado tipo, y su sentido viene dado por las manifestaciones externas típicas de quien tiene un dolor. Como puede apreciarse, en esta hipótesis evitamos la situación de las cajas y de los «escarabajos», esto es, evitamos el solipsismo, pues si bien el referente es estrictamente privado (solamente yo puedo sentir mi dolor, por definición), el término no obtiene su significado por conexión inmediata y directa con el referente, sino por conexión directa con sus manifestaciones externas, y únicamente a través de estas últimas se conecta con la vivencia. Nótese que esto no quiere decir que «dolor» signifique, por ejemplo, «llanto». Wittgenstein, con razón, ha rechazado de forma explícita esta interpretación de su doctrina (244). Ello no ocurre en mi hipótesis puesto que el llanto, o cualquiera otra de las manifestaciones típicas del dolor, no son referentes del término «dolor»; el referente es la sensación, la vivencia, y las manifestaciones externas constituyen tan sólo las características de la realidad a través de las cuales se identifica el referente por medio de la expresión lingüística, en este caso, por medio del término «dolor». Las expresiones mentales tendrían, sin embargo, a diferencia de otras, una peculiaridad: cuando el sujeto las emplea para hablar de sí mismo en primera persona pueden funcionar como expresiones de sus estados mentales, vivencias, reemplazando a las manifestaciones naturales de éstos.

Lo anterior presupone que las experiencias internas tienen siempre manifestaciones naturales o primitivas, esto es, no lingüísticas. ¿Es así? No cabe aquí intentar resolver este punto, pero hay que tener en cuenta que, en ciertos casos, puede ser cuestión compleja la de determinar esas manifestaciones. ¿Cuál es la expresión natural de la angustia, de la esperanza, de la imaginación? No digo que no la tengan; lo que quiero decir es que, frente a estos ejemplos, el caso del dolor es demasiado sencillo. Wittgenstein nunca descendió a grandes detalles en lo que respecta a los aspectos constructivos de su doctrina en este tema. Pero no hay que pensar que no sea posible extenderla a todos los tipos y clases de vivencias o experiencias internas. Un estudio amplio y sistemático de los conceptos mentales dentro de los límites de un planteamiento como el de Wittgenstein fue el que hizo Ryle en *El concepto de lo mental*.

En mi opinión puede afirmarse que, desde el punto de vista de la teoría del conocimiento, toda la filosofía de Wittgenstein es un intento de presentar el solipsismo como algo imposible. Semejante esfuerzo está presente

tanto en el *Tractatus* como en las *Investigaciones*. En la primera de las dos obras, la doctrina de la representación figurativa hace imposible el solipsismo porque, en cuanto el solipsista intenta formular su postura, ha de recurrir a un lenguaje que solamente tiene sentido en la medida en que representa la realidad isomórficamente, con lo que la actitud solipsista queda contradicha por el propio uso de un lenguaje que, necesariamente, trasciende los límites del sujeto. En su segunda época, la necesaria conexión existente entre el lenguaje y la actividad extralingüística hace igualmente imposible la postura solipsista: las palabras que use el solipsista tan sólo pueden obtener su significado de manifestaciones externas, pero nunca exclusivamente de lo que halle en el recinto de su conciencia. Al utilizar los términos «necesaria» y «necesariamente», quiero poner de manifiesto que la conexión entre el lenguaje y la función figurativa, en la primera época, y la conexión entre aquél y las manifestaciones externas de las vivencias, en la segunda etapa, no son conexiones contingentes, sino, desde el punto de vista de la doctrina de Wittgenstein, conexiones que el lenguaje tiene por definición, o dicho con terminología tradicional, que le son esenciales. Con esto pretendo subrayar que Wittgenstein no ha abandonado totalmente, en su evolución posterior, ese trascendentalismo lingüístico que caracteriza al *Tractatus*. La crítica del concepto de lenguaje privado, como todo su tratamiento de las expresiones que se refieren a experiencias internas, involucra una teoría del significado que excede con mucho de la mera descripción de los usos lingüísticos. La crítica a los supuestos lenguajes privados no se apoya pura y simplemente en una descripción de los usos que hacemos del lenguaje. De aquí lo más que podríamos obtener es la prueba de que, de hecho, nunca se usa el lenguaje de modo privado en el habla cotidiana. Pero, como hemos comprobado, lo que Wittgenstein justifica es la imposibilidad de que el lenguaje llegue a usarse de esa forma privada, y su argumentación se reduce básicamente al siguiente enunciado: a un lenguaje privado no es posible aplicarle el concepto de significado ni el concepto de regla, conceptos sin los cuales no podemos explicar el concepto de lenguaje. Y esta argumentación ciertamente posee un carácter más trascendental que empírico y queda del todo próxima al tipo de argumentación propio del *Tractatus*. La crítica a la idea de un lenguaje privado es, además, y por supuesto, relevante para una crítica de la mayor parte de la filosofía a partir de Descartes, en la medida en que los sistemas filosóficos modernos participan en la idea de que es posible expresar lingüísticamente el contenido de la conciencia poniendo al mismo tiempo en duda la existencia del mundo exterior. Que esto no es posible parece deducirse fácilmente de la doctrina de Wittgenstein. Pero no da la impresión de que el amplio alcance de su posición crítica le interesara, o ni siquiera de que fuera consciente de ello. Aquí, como en todos los temas principales que toca en las *Investigaciones Filosóficas*, lo que parece obsesionarle son las limitaciones de la teoría referencialista, y por consiguiente, las insuficiencias del atomismo lógico en general, y del *Tractatus* en particular.

7.4 La filosofía como descripción de los usos lingüísticos

Este resto de trascendentalismo lingüístico que acabamos de ver implícito en la crítica a los lenguajes privados no se notaba apenas en la teoría del lenguaje como conjunto de usos, donde sólo podría advertirse débilmente en aspectos como el de la necesaria conexión entre el empleo del lenguaje y la forma de vida. En la formulación del método filosófico tal trascendentalismo es aún más difícil de encontrar.

Una de las afirmaciones que sería más fácil atribuir a una actitud trascendentalista es ésta: «Sentimos como si debiéramos comprender los fenómenos, pero nuestra investigación no se dirige a los fenómenos, sino, por así decirlo, a las 'posibilidades' de los fenómenos» (90). Este trascendentalismo, como en el *Tractatus*, es lingüístico: «Es decir, que lo que hacemos es recordar el tipo de enunciados que hacemos sobre los fenómenos. (...) Nuestra consideración es, por tanto, una consideración gramatical» (*Ibidem*). Y esta gramática, que podemos llamar «gramática filosófica», cumple una función de gran alcance, pues incorpora y agota todo cuanto tenga sentido considerar como tarea de la filosofía. Para empezar, investiga acerca de la función, de la estructura, en suma, de la esencia del lenguaje (92), pues «la esencia se expresa en la gramática» (371). No es, por supuesto, la gramática ordinaria o superficial, sino lo que en otro lugar llamará Wittgenstein «gramática profunda» (664). La esencia de la que ahí se habla no es, sin embargo, algo escondido, algo que hay que desenterrar y sacar a la luz (92); lo que pueda estar escondido, no nos interesa (126). ¿Qué quiere decir esto?

Quiere decir que este trascendentalismo gramatical se resuelve a la postre en una tarea puramente descriptiva, y aparentemente nimia: en aclarar aquellos equívocos y confusiones que conciernen al uso de las palabras, y que pueden obedecer, entre otras causas, a ciertas analogías existentes entre expresiones pertenecientes a diferentes regiones del lenguaje; una manera de llevar a cabo esta aclaración puede consistir en sustituir un tipo de expresión por otra (90). Por este camino, el programa filosófico de Wittgenstein cae, al menos tal y como él lo formula, en un mero descriptivismo de tan modesto alcance que aplicarle el nombre de «filosofía» parece irónico: «el trabajo del filósofo es un reunir recordatorios con una finalidad determinada» (127). ¿Recordatorios de qué? De cómo se usa el lenguaje en la vida cotidiana. Así, cuando el filósofo emplea términos como «saber», «ser», «objeto», «yo», «proposición», «nombre», e intenta aprehender la esencia de la cosa, hay que preguntarse: ¿se usa de hecho esa palabra de ese modo en el lenguaje en el que tiene su lugar de origen? (116). Y añade Wittgenstein: «Trasladamos las palabras desde su uso metafísico a su uso cotidiano.» ¿Recordatorios, con qué finalidad? Con la de deshacer el equívoco, la confusión, sobre la que descansa, en cada caso, el problema filosófico. De esta manera, «la filosofía simplemente coloca todo delante, y ni explica ni deduce nada» (126). Todo está a la vista, puesto que se trata de lo que todos hacemos a diario, los usos del lenguaje;

no hay una forma lógica, una forma de la proposición, que, como pensaba en el *Tractatus*, haya que sacar a la luz. Y por consiguiente, no hay nada que explicar: «debe desaparecer toda explicación, y sustituirla sólo la descripción» (109). En filosofía no hay conclusiones que sacar, pues lo que se enuncia es lo que todo el mundo admite (599); en consecuencia, tampoco hay nada que discutir en filosofía (128), y podría darse el nombre de «filosofía» a lo que es posible antes de todo descubrimiento y de todo invento (126). Dicho de otro modo: la filosofía no altera nada, puesto que no acrece nuestro conocimiento; la filosofía deja todo tal como está, y no puede modificar nuestro uso del lenguaje, ni tampoco suministrarle fundamento; en definitiva, únicamente puede describirlo (124).

Como acabamos de ver, esta descripción de los usos lingüísticos a la que queda reducida la filosofía tiene una justificación. ¿Cuál? La que responde al propósito de resolver los propios problemas filosóficos. Tales problemas no son, por supuesto, empíricos sino conceptuales, y se resuelven observando el funcionamiento del lenguaje; lo que se requiere no es nueva información, sino ordenación de lo que ya sabemos (109). El problema filosófico, para Wittgenstein, es un problema que se da en el lenguaje. En este punto, su posición es muy semejante a la que mantenía en el *Tractatus*. Aquí, los problemas filosóficos procedían del desconocimiento de la lógica de nuestro lenguaje.

Ahora, las proposiciones filosóficas expresan el resultado de una especie de calambre mental producido por una confusión respecto a las reglas que rigen el empleo del lenguaje. Pero, desgraciadamente, Wittgenstein no es ahora preciso en la caracterización del problema filosófico. Una de sus consideraciones más claras se encuentra ya en el *Cuaderno azul*: «El hombre que se halla filosóficamente perplejo ve una ley en el modo de usar una palabra, y al tratar de aplicar esta ley de modo consistente tropieza con casos en los que da con resultados paradójicos.» (Pág. 27 del original y 56 de la trad. cast.) Por lo demás, lo que abunda en este contexto son las metáforas: los problemas filosóficos surgen cuando el lenguaje está de vacaciones (38), cuando el lenguaje se mueve en el vacío, en lugar de funcionar (132); la filosofía rectamente entendida es una tarea (como lo era en el *Tractatus*), y esta tarea es una lucha contra el embrujamiento de nuestro entendimiento por el lenguaje (109 y *Cuaderno azul*, loc. cit.). Una de las metáforas que, por cierto, algunos de sus discípulos se tomaron más en serio, es la que presenta la tarea filosófica como una terapia: «el filósofo trata una cuestión como si fuera una enfermedad» (255), y por eso no hay un método filosófico, sino varios métodos, igual que hay diversas terapias (133). Quien tiene un problema filosófico se encuentra como perdido (123), y hay que enseñarle el camino como se ayuda a una mosca a salir de una botella (309).

Aparte de los problemas filosóficos que Wittgenstein trata con ocasión de hacer la crítica del atomismo lógico, y que ya hemos considerado en su momento comprobando cómo funciona su método, ha solido ser muy parco en la alusión a otros problemas. Una de las excepciones más notables es.

por lo que toca a los *Cuadernos azul y marrón*, el problema del tiempo, que podemos examinar brevemente como último ejemplo de su método de disolución de los problemas filosóficos.

Wittgenstein recuerda la preocupación de San Agustín por la definición del tiempo, y señala que la definición aclara con frecuencia la gramática de la palabra, pues es tal gramática la que suele causarnos la perplejidad, y, concretamente, la responsable de la contradicción que San Agustín formula: el tiempo no puede medirse, puesto que ni se puede medir el pasado, que ya no existe, ni se puede medir el futuro, que no existe todavía, ni es posible medir el presente, que carece de extensión (*Cuaderno azul*, p. 26 del original y 54-55 de la trad.). Esta contradicción, según Wittgenstein, se debe al conflicto entre dos usos distintos de la palabra «medir». Si el significado que damos a esta palabra es el que tiene cuando lo que medimos es la distancia entre dos marcas hechas en una cinta que se mueve, y de la que en cada momento sólo podemos ver una porción mínima, entonces no podremos explicar cómo sea posible medir el tiempo. El uso del término «medir» aplicado al tiempo presenta importantes diferencias, pues el tiempo no es como una cinta que se mueva; no puede compararse el pasado con la parte de la cinta que ya ha transcurrido ante nosotros y que hemos perdido de vista, ni el futuro con aquella porción que está por pasar, y menos aún el presente con el trocito que en cada momento tenemos a la vista. Lo que hay que recordar es que las reglas que rigen la palabra «medir» en el caso del tiempo son distintas de las que regulan su uso en el caso de una distancia. Dar una definición del tiempo no nos servirá, porque comprobaremos que, en algún sentido, es insuficiente, y entonces la sustituiremos por otra, que a la postre lo será asimismo, y así sucesivamente (p. 27, y 56 de la trad.).

En el *Cuaderno marrón* imagina una serie de juegos lingüísticos con los que podríamos enseñar a un niño a narrar acontecimientos pasados (páginas 104-109 del original y 141-147 de la trad.). Uno de ellos, por ejemplo, consiste en enseñarle a correlacionar dos series de ilustraciones. Una de ellas consiste en dibujos, fotos, o lo que quiera que sea, representando diversas posiciones del sol en relación con el paisaje en el que se desenvuelve la vida del niño. La otra contiene representaciones de las distintas actividades del niño a lo largo del día. El niño aprende a correlacionar ambas series, de manera que a cada figura de una actividad suya le haga corresponder una figura de la posición del sol, de acuerdo con la relación real que hay entre las actividades del niño y las diferentes posiciones del sol. Se supone que el niño es capaz de establecer la correcta relación entre las figuras en la medida en que los elementos del paisaje le ayudan a identificar las posiciones del sol; es decir, el niño sabe que el sol sale sobre aquella montaña, que se pone sobre aquellos árboles, etc. Para Wittgenstein, este tipo de correlaciones es lo que está implicado en la idea de medir el tiempo, y ocurre que, al considerar tales juegos, «no encontramos las ideas de pasado, presente y futuro, en su aspecto problemático y casi misterioso» (p. 107, y 144 de la trad.). Es el aspecto que se quiere indicar cuando uno se pregunta

adónde va el presente cuando se hace pasado. Pero este tipo de cuestiones nacen de metáforas desacertadas, como la que representa el tiempo como un río en el que flotarían las cosas; y entonces diremos que las cosas pasan, que el acontecimiento futuro está por llegar, y que el pretérito ya ha pasado. Al final, la propia metáfora mostrará su insuficiencia, pues acabaremos diciendo no sólo que las cosas pasan, sino también que el tiempo pasa, no sólo que partiré de viaje el sábado, sino que mi partida se aproxima, etcétera. Todo esto son, para Wittgenstein, ejemplos de una obsesión con el simbolismo; una obsesión que nos lleva construir analogías indebidas entre los diferentes usos de las palabras, y de aquí brota la pregunta filosófica.

Parece haber, pues, una constante en la forma en que Wittgenstein caracteriza los problemas filosóficos: se trataría de problemas que surgen de usar el lenguaje fuera de su contexto habitual, y la tarea debe consistir, por eso, en devolver las palabras a ese contexto cotidiano. Al igual que en el *Tractatus*, los problemas filosóficos no son problemas a resolver sino problemas a disolver. Ahora bien, mientras que en el *Tractatus* esa condena de la filosofía es del todo coherente con la teoría del lenguaje que allí se propone, con la doctrina de la representación figurativa, semejante coherencia no resulta tan clara en su segunda época. Pues ahora el lenguaje se analiza como conjunto de sus usos, y ya hemos visto declarado que no hay, en principio, límites a la variedad de tales usos (23). Cabe, entonces, preguntar: ¿por qué rechazar la utilización metafísica del lenguaje? ¿No puede haber un uso filosófico del lenguaje como hay un uso religioso, un uso lógico, un uso científico natural, o un uso poético? ¿Por qué no va a ser la filosofía un juego lingüístico más como tantos otros? La respuesta de Wittgenstein a estas preguntas probablemente discurriera en torno a una característica del lenguaje que hemos visto señalada por él, y que, como he indicado al comienzo de esta sección, constituye un débil resto de trascendentalismo en su nueva teoría del lenguaje: se trata de la conexión necesaria entre el uso del lenguaje y el resto de las actividades que componen una forma de vida. Puede, en verdad, hablarse de una forma de vida religiosa, en la que una serie de actividades y actitudes suministran un contexto extralingüístico en el que cobra sentido el juego lingüístico religioso; así como la actividad del científico en su trato con las cosas da sentido a su peculiar uso del lenguaje cuando formula sus teorías; el lógico, por su parte, opera con el lenguaje para descubrir en él ciertas relaciones formales, las relaciones de inferencia. ¿Y el poeta? También, acaso, el medio en el que se desenvuelve su actividad es el lenguaje, sólo que lo que en él busca es la belleza de la expresión, lo insólito del significado, la sonoridad de las palabras. ¿Pero de qué trata el filósofo? ¿Cuál es la actividad que da sentido a su uso del lenguaje? ¿Es el propio lenguaje el medio en el que opera? Pero en última instancia, el filósofo, incluso el atomista lógico, pretende hablar de la realidad. Sin embargo, no hay ningún trato con ésta que parezca ser peculiar de la filosofía, no hay actividad alguna que suministre un contexto para el pretendido juego filosófico con el lenguaje.

Pienso que tales argumentos podrían justificar la posición de Wittgenstein. ¿Son aceptables? En mi opinión apoyan con éxito una caracterización que es aplicable a buena parte de los problemas filosóficos, pero que ni da cuenta de la filosofía en su conjunto ni excluye la posibilidad de continuar haciendo filosofía. Wittgenstein no parece darse cuenta de que si el uso filosófico del lenguaje carece de un particular contexto de actividades y comportamientos es porque trata de todo, porque todo puede ser tema de la filosofía si se considera más allá de cualquier contexto particular, más allá de cualquier actividad determinada. Por esto precisamente, y pese a su condena de la filosofía, el *Tractatus* y las *Investigaciones* son dos obras filosóficas de las más importantes de nuestro siglo. Que ciertos problemas filosóficos nacen de un uso peculiar e ilegítimo de las palabras es por su parte una tesis filosófica; cómo es posible que el lenguaje pueda llegar a crearnos esas confusiones mentales es a su vez un problema filosófico. Y una descripción terapéutica de los usos ordinarios del lenguaje, que nos rememore las reglas que implícitamente aplicamos en nuestro comportamiento lingüístico, puede evitar nuestra persistencia en esas confusiones, pero no impedirá que continuemos planteándonos problemas filosóficos, porque los problemas filosóficos surgen desde todos los puntos, se insinúan en todas las direcciones. Lo difícil no es plantearse un problema filosófico, sino no planteárselo. Cada vez que nos preguntamos por las relaciones entre las diferentes ciencias, por la relación entre las ciencias y la práctica cotidiana, por las condiciones y límites del conocimiento, por la validez última de nuestras normas y valoraciones de todo tipo, estamos planteando problemas filosóficos. Por ello, la pregunta por el tiempo seguirá siendo una pregunta filosófica siempre que responda a la insatisfacción causada por una definición determinada, o cuando quiera que aspire a una integración de enfoques diversos, por ejemplo, de la consideración física del tiempo y de su vivencia en la experiencia interior. Que en la respuesta hay que eludir la metáfora confundente o la analogía descaminada es obvio. Pero con ello no se elimina al problema filosófico; simplemente se lo depura.

La injusticia —llamémosla así— que Wittgenstein comete con el uso filosófico del lenguaje, afecta también a su crítica del atomismo lógico. Hemos visto que uno de los puntos criticados es el de la distinción entre lo simple y lo compuesto. Wittgenstein intenta mostrar que esta distinción es siempre relativa a un contexto determinado, a unos propósitos particulares, a un cierto procedimiento de análisis, y que nada de esto está presente en el *Tractatus* cuando se exige que las proposiciones se compongan de nombres o que los hechos se hallen compuestos por objetos. No sabemos —viene a decir Wittgenstein— ni para qué propósitos ni con qué procedimientos pretendemos descomponer la realidad y el lenguaje. Y recurre en esta crítica, según vimos, a la ignominiosa comparación con el análisis de una escoba: ¿de qué se compone una escoba?; ¿de palo y escobilla?; ¿y se gana algo sugiriendo que el término «escoba» equivale realmente a «palo

con una escobilla fijada a un extremo?» Pero Wittgenstein es aquí tremendamente injusto con el autor del *Tractatus*, y éste podría responderle: el análisis lógico-atomista no está pensado para cosas como escobas; es un análisis lógico-metafísico que pretende dar cuenta de la estructura de lo real e ir más allá de los fenómenos. Es cierto que lo que es simple para el carpintero puede no serlo para el pintor o para el físico, pero la cuestión interesante es: ¿hay algún tipo de simplicidad específica?, ¿hay una simplicidad metafísica? Y si no la hay, ¿existe alguna simplicidad que, por su alcance más general, sea particularmente relevante para el conocimiento de la realidad? Parece, en efecto, que, si se rechaza el análisis metafísico, ocupará su lugar algún otro tipo de análisis que, por su alcance y por su solidez metodológica, satisfaga nuestras aspiraciones de conocimiento. Por ejemplo, el análisis científico; y entonces, aceptaremos como constitutivos simples de la realidad aquellos de los que habla la física atómica. Tal sería la posición de un neopositivista. Pero no vale decir que lo simple es relativo y que varía según el punto de vista en el que uno se sitúa; porque no todos los puntos de vista son igualmente relevantes para el problema debatido. Para la cuestión de si hay elementos simples de los que se componga la realidad, y de cuáles sean, es relevante el punto de vista del físico, pero no lo es en absoluto el del pintor ni el del carpintero. El relativismo de Wittgenstein, explicable como reacción contra el *Tractatus*, no es más aceptable que el absolutismo del atomismo lógico, y por desgracia carece de la grandeza lógico-metafísica de dicha obra.

Desde la teoría de la proposición, del *Tractatus*, a la teoría de los usos, en las *Investigaciones Filosóficas*, hay un proceso que podemos llamar, tomando el término de otros, de «destrascendentalización» (*cfr.* Rorty, «Epistemological Behaviorism and the De-Trascendentalization of Analytic Philosophy»). Los usos, infinitamente multiformes y variados, siempre cambiantes, sustituyen a la proposición como representación isomórfica; las condiciones necesarias que hacen posible al lenguaje se difuminan en una vaga conexión con la forma de vida, con las actividades extralingüísticas, la cual deja fuera únicamente los lenguajes privados. Podemos decir que la semántica trascendental tiende a ser sustituida por la pragmática empírica (y digo «tiende» para salvar los restos de trascendentalismo que he señalado en Wittgenstein). Ahora bien, la descripción de los usos deriva su interés filosófico de un peculiar diagnóstico acerca del problema filosófico, pero si este diagnóstico no es válido más que para una clase de esos problemas, otros métodos pueden adquirir, o recobrar, más importancia, reduciendo al tiempo la del procedimiento descriptivo defendido por Wittgenstein. En particular, podría pasar a primer término una teoría semántica de los tipos de proposiciones, que hemos visto explícitamente pospuesta por Wittgenstein. Más aún: podría adquirir relevancia una nueva teoría sistemática del lenguaje en la que la variedad de los usos lingüísticos quedara ordenada de algún modo con criterios claros, resultando de esta forma más manejable. Es lo que ocurre, por ejemplo, en la teoría de los actos de habla.

7.5. La herencia de Wittgenstein

Los discípulos y continuadores de Wittgenstein han solido ser agrupados según se vinculen a una u otra de las dos grandes universidades inglesas, Oxford y Cambridge. Tal clasificación, sin duda, no pasa de ser aproximada, y por fuerza ha de dejar fuera a muchos pensadores ajenos a la filosofía inglesa, pero posee cierto sentido. De un lado, la filosofía vigente en Oxford después de la guerra ha estado centrada en el lenguaje ordinario de una manera que ciertamente tiene una estrecha relación con la doctrina de Wittgenstein y acusa su influencia; ha sido denominada, por ello, «filosofía del lenguaje ordinario» o «escuela de Oxford», y son sus representantes algunos de los filósofos ingleses más importantes entre los años cuarenta y sesenta. A la llamada «escuela de Cambridge» pertenecen, en cambio, algunos pensadores que se caracterizan por llevar adelante el programa terapéutico de Wittgenstein, aplicándose de forma continuada y temática a la disolución de problemas filosóficos diversos.

El representante más conocido de esta última tendencia es Wisdom. Su método consiste en considerar los problemas filosóficos desde todos los puntos de vista posibles, a fin de mostrar que se trata de perplejidades producidas por afirmaciones paradójicas. Iluminado desde múltiples ángulos, el problema aparece situado en un contexto en el que acaba perdiendo la problematicidad. De aquí el paralelismo entre las afirmaciones del filósofo y las afirmaciones del psicoanalista («Filosofía, metafísica y psicoanálisis»). Pero a pesar de que engendren perplejidad, los enunciados paradójicos deben tomarse en serio, pues sólo al considerar la paradoja literalmente se concede al detalle concreto aquella atención que puede permitir romper los viejos hábitos mentales y reconocer en qué medida y de qué modo una vieja clasificación desfigura la realidad (*op. cit.*, p. 273 del original y 445 de la trad. cast.). La paradoja puede servir para captar aspectos de la realidad sobre los que antes no se había reparado. Como en el psicoanálisis, así en la filosofía una visión más completa de las raíces y conexiones del problema puede contribuir a curarlo. El tratamiento de un problema filosófico es como el tratamiento de una enfermedad, había dicho Wittgenstein; y Wisdom precisa: como el de una enfermedad mental, pues «cuando consideramos las dudas obstinadas del metafísico, tales como «¿Puede alguien saber lo que está bien y lo que está mal?», «¿Puede alguien saber lo que otros piensan o sienten?», al punto nos recuerdan las dudas crónicas del neurótico y del psicópata, tales como «¿He cometido el pecado que no se perdona?», «¿No están todos realmente en contra de mí?» (*op. cit.*, páginas 281-282 del original, y 453 de la trad. cast.). Como en Wittgenstein, los argumentos de Wisdom se desarrollan con frecuencia en forma de diálogo, pero aquí las argumentaciones tienden, muy claramente, a poner de manifiesto todas las razones posibles a favor y en contra de una posición, de modo tal que, a la postre, ninguna quede victoriosa, viéndose que todas ellas son en parte aceptables y en parte inaceptables, y que cada una de las posiciones en torno a un problema tiende a destacar o subrayar un aspecto

de la realidad que las demás pasan por alto. Wisdom ha mantenido, por ello, que el método propio y característico de la filosofía consiste, no en un razonamiento deductivo, ni en uno inductivo, sino en un tipo de razonamiento comparativo, analógico, según el cual se subsumen casos distintos bajo la misma categoría en base a sus semejanzas relativas («Gods», p. 157 de *Philosophy and Psychoanalysis*; recuérdense los parecidos de familia en Wittgenstein).

En escritos de cuidado estilo literario, Wisdom se limita a reelaborar y aplicar el diagnóstico wittgensteiniano sobre el problema filosófico y su terapia adecuada. No hay en aquél ningún interés particular por un estudio del lenguaje ni aportación alguna destacable a su estudio. La preocupación por el lenguaje sí es, en cambio, un motivo central en los filósofos ligados a la escuela de Oxford. El primero que ha de ser citado aquí es Ryle. Muchas de las insinuaciones que Wittgenstein desliza o sugiere se encuentran formuladas en él de manera explícita y con mayor precisión. Así, he propuesto, en la sección anterior, una posible razón por la que Wittgenstein rechazaba el uso filosófico del lenguaje, a saber: que no hay ninguna actividad peculiar que suministre un contexto para el mismo. Pues bien, en una vena levemente más comprensiva que la de Wittgenstein, Ryle ha señalado que una de las razones por las que los argumentos filosóficos parecen, en ocasiones, discusiones entre sordos, y por las que unos filósofos parecen emplear términos del todo heterogéneos con los de otros, es precisamente porque no existe actividad alguna que constituya el dominio de la habilidad filosófica; los expertos que usan términos técnicos de derecho, de química o de fontanería, aprenden a emplear esos términos, en parte siguiendo las instrucciones oficiales propias de su oficio, pero en mayor parte por el ejercicio directo de las técnicas especiales propias de su profesión. Pero los términos filosóficos no son de este tipo, pues no hay un campo peculiar del conocimiento ni una aptitud especial que sean propios del filósofo («El lenguaje común», pp. 123-124 del original y 53 de la trad. cast.).

De aquí que el filósofo consciente de su condición acabe por recurrir al lenguaje ordinario o común, en contraposición tanto a los lenguajes técnicos de las diferentes profesiones y disciplinas como a los lenguajes formalizados propios de los cálculos lógicos. Recogiendo también aquí un motivo claramente wittgensteiniano, Ryle se cuida de indicar las insuficiencias de la lógica formal para la resolución de los problemas filosóficos. El sueño del formalizador, como Ryle lo llama (*loc. cit.*, p. 125), consiste en creer que los poderes lógicos de las expresiones ordinarias pueden reducirse sin pérdida a los de las expresiones de un cálculo lógico. Mas el lenguaje ordinario tiene su propia lógica, una lógica informal, que es justamente la relevante para el planteamiento y la discusión de los problemas filosóficos (*Dilemmas*, cap. VIII), y que excede con mucho de las abstractas y tipificadas relaciones presentes en la lógica formal. El conjunto de las relaciones lógicas que cada proposición tiene con las demás, lo que Ryle llama los «poderes lógicos de las proposiciones» («Argumentos filosóficos», p. 331 del original), no se dejan trasladar salvo en pequeña medida a un lenguaje

formal. Por ello, desde el punto de vista de la lógica, entendida con esta amplitud y de tal modo que es, en su mayor parte, lógica informal, queda justificado el estudio del lenguaje ordinario.

Ryle ha cuidado también, asimismo, de clarificar en qué sentido tiene interés filosófico el uso del lenguaje, y tal como él lo entiende, la conclusión es que se trata de algo muy tradicional en filosofía. A la pregunta: ¿Tiene algo que ver la filosofía con el uso del lenguaje?, Ryle responde: esto equivale a preguntar si las discusiones conceptuales son discusiones filosóficas. Y es claro que las discusiones sobre conceptos, por ejemplo, como los de causa, número, voluntad, etc., siempre han sido consideradas discusiones filosóficas, y lo siguen siendo. Es decir, no hay para Ryle ninguna novedad radical en esta atención al uso del lenguaje, y la razón es que la pregunta por el uso de una expresión equivale, para él, a la vieja pregunta por un concepto. Si nos preguntamos por el uso del término «causa», lo que estamos cuestionándonos es para qué sirve, qué función o tarea cumple, y esto es precisamente lo que, por ejemplo, Hume se planteaba sobre el concepto de causa. Ni la pregunta por el uso ni la pregunta por el concepto son preguntas acerca de una palabra de una lengua particular. En la medida en que «causa» pueda traducirse, en los contextos que nos interesan, por *cause* o por *Ursache*, nuestra investigación afectará indistintamente a cualquiera de ellas, puesto que versa sobre las tareas, sobre la función, que esas palabras cumplen, y en la medida y grado en que sean intertraducibles coincidirán en su uso. No otra cosa hay en la investigación tradicional de los conceptos, pero la nueva terminología tiene la ventaja de que evita los problemas ontológicos que solían plantear estos últimos, por ejemplo, respecto a cómo existen o dónde están, así como evita también los similares problemas que surgen cuando se sustituye los conceptos por los significados (por otra parte, el interés de hablar del uso más que del significado ya ha quedado bien documentado con la discusión de Wittgenstein). Por último, hablar sobre el uso tiene la ventaja de que permite distinguir el uso correcto del uso incorrecto, y de esta manera deshacer más fácilmente ciertas confusiones filosóficas («El lenguaje común», pp. 112-114 y 126 del original, y 42-44 y 55 de la trad. caste.). Entendido el uso lingüístico de esta manera, digamos conceptual, no corresponde exactamente a la extrema amplitud con la que Wittgenstein habla de «uso» en las *Investigaciones filosóficas*, ni encaja con algunos de los ejemplos que suministra en el párrafo 23 de esta obra (hacer teatro, adivinar acertijos, contar chistes...). El concepto de uso que tiene Ryle es más preciso y más restringido, y al tiempo que contribuye a acercar la filosofía lingüística a la filosofía tradicional, le resta a aquella gran parte de la novedad y del radicalismo que posee Wittgenstein. Ryle, por su parte, ha subrayado su interpretación lógico-informal del concepto de uso, distinguiendo entre éste y lo que, en la traducción castellana, se ha vertido como «usanza» (en el original, distinguiendo entre *use* y *usage*; «El lenguaje común», p. 115 y ss. del original y 46 y ss. de la trad.). La distinción se refiere a aquellos aspectos del uso que son del todo irrelevantes para la determina-

ción de eso que Ryle ha denominado «los poderes lógicos de las proposiciones», esto es, aquellas relaciones de implicación, consistencia, inconsistencia, apoyo inductivo, etc., que una proposición tiene con otras. Y es patente que, para la determinación de estas características, es irrelevante si la expresión en cuestión se usa mucho o poco, si es propia del lenguaje culto o del habla popular, si está de moda o resulta anticuada, etc. Estos son los aspectos que Ryle atribuye al *usage* de las expresiones, dejando para el *use* lo que se refiere a la tarea o función que realizan las expresiones. De hecho, en castellano hay que decir que ambos son aspectos del uso de las expresiones, pero el aspecto conceptual, que a Ryle le interesa, es el que corresponde a lo que podríamos llamar «la semántica de la expresión», mientras que las características sociológicas que excluye de su consideración forman parte más bien de un tratamiento pragmático y sociolingüístico del lenguaje. La necesidad en que Ryle se ha visto de distinguir ambos aspectos, aunque no lo haga en estos términos, muestra por contraste la enorme vaguedad del concepto de uso, algo que, por cierto, ya hemos visto ilustrado a propósito de Wittgenstein (y lo que digo se aplica igualmente en inglés, pues la distinción de Ryle entre *use* y *usage* ciertamente no se da en el lenguaje ordinario con nitidez, y tal como él la utiliza más bien resulta un recurso técnico para distinguir estos dos aspectos del tema).

Lo anterior sugiere ya una importante cuestión en el estudio del lenguaje, y es la de cómo se relacionan el uso de las palabras entendido como empleo o funcionamiento con su uso entendido como costumbre o hábito social. Ryle nota, con razón, que, cuando se trata de actividades intersubjetivas, aprender el uso de ciertos instrumentos supone conocer los usos sociales que regulan dichas actividades. Pero en el caso del lenguaje la relación es compleja y sutil. Ryle ha admitido que puede hablarse de un mal *use*, pero no de un mal *usage* (*loc. cit.*); ahora bien, lo importante es que un *use* incorrecto suficientemente prolongado e imitado puede convertirse en un nuevo *usage* que a su vez legitime como correcto dicho *use*. O dicho en castellano: que el empleo correcto de las palabras obedece a prácticas sociales, y que una utilización incorrecta puede, a la larga, originar nuevas prácticas que vengan a justificar como correcta dicha utilización. Podría objetarse que esta consideración introduce un punto de vista diacrónico o histórico que es ajeno, e irrelevante, para una posición que, como la de Ryle, es en definitiva, y aunque en un sentido amplio, lógica, y, por consiguiente, hace abstracción de los procesos de modificación en el lenguaje. Aun aceptando la objeción, hay que añadir que el tema de la relación entre ambos aspectos del uso es el propio Ryle quien lo menciona. Nótese, por cierto, que son estos aspectos del uso lingüístico en cuanto práctica social los que se acentúan cuando se estudia el uso del lenguaje en una perspectiva como la de la filosofía orteguiana (por ejemplo, en Ortega, *El hombre y la gente*, cap. XI, o en Marías, «La realidad histórica y social del uso lingüístico»).

El mejor ejemplo del método de Ryle es su obra sobre *El concepto de lo mental*. Se trata aquí de estudiar los poderes lógicos de un cierto tipo de proposiciones, a saber: las que versan sobre los procesos llamados «mentales», o, dicho con otra expresión que el autor emplea en la introducción de su libro, de rectificar la geografía lógica de los conceptos mentales. Teniendo en cuenta que, para él, un concepto es simplemente el modo de usar una expresión, hay que concluir que su proyecto consiste en rectificar, al menos en parte, nuestro uso de las expresiones que se refieren a fenómenos mentales. Sus argumentos van dirigidos de modo particular contra todas aquellas teorías que categorizan los fenómenos mentales como constituyentes de un mundo interno semejante y paralelo al mundo físico. Tales teorías serían culpables de aceptar lo que llama Ryle «el dogma del fantasma en la máquina» o mito de Descartes, según el cual lo que convierte nuestro comportamiento externo en propiamente humano es el conjunto de acontecimientos paralelos que supuestamente tienen lugar en su interior; un cuerpo sería humano porque está animado por un alma. Ryle consideraba este tipo de teorías como un desgraciado producto de la filosofía cartesiana, e intentaba mostrar en su obra que semejante forma de hablar acerca de la mente conduce a diversas falacias y confusiones y que, en suma, constituye un ejemplo de lo que denominaba «error categorial», esto es, el error de asignar a un tipo de realidades una categoría diferente de la que les corresponde. Su posición era que a los fenómenos mentales no puede aplicarse las categorías propias del análisis físico, y que por ello no tiene sentido contraponer la mente y la materia, que lo mental no es un mundo paralelo al mundo físico, sino una clase particular de características del comportamiento humano. Lo interesante es que, en el curso de sus argumentaciones, Ryle, de acuerdo con el propósito rectificatorio que he mencionado, hacía recomendaciones diversas sobre cómo usar o no usar las palabras; así, por ejemplo, recomendaba prescindir de la expresión «en la mente», a fin de evitar la falsa localización de los procesos mentales (cap. II, secc. 5). Pero en otras ocasiones, en cambio, recurría al uso ordinario para descalificar el uso de los filósofos por divergente con aquél; y así, indicaba que los términos «voluntario» e «involuntario» se aplican comúnmente a las acciones que no deben realizarse, mientras que es típico de los filósofos extender el uso de esos términos a los actos meritorios, dando así lugar a numerosos pseudoproblemas acerca de la libertad de la voluntad (cap. III, secc. 3). Sean cuales fueren los méritos específicos de la investigación de Ryle sobre la lógica de los conceptos mentales, hay que decir que su manera de recurrir al lenguaje ordinario distaba mucho de ser clara y homogénea. Pues si bien la divergencia del uso filosófico respecto al primero puede constituir un testimonio adicional de que tal uso corresponde a una doctrina filosófica equivocada, es patente que la reconstrucción de los conceptos mentales emprendida por Ryle implicaba una reforma de ciertos usos ordinarios del lenguaje, tanto más cuanto que, por mucho que el dogma dualista que intentaba desacreditar tenga como uno de sus representantes modernos más eximios a Descartes, se trata evi-

dentemente de un dogma, no sólo constante en la filosofía occidental desde los griegos, sino muy anterior a ésta y procedente de lejanas instancias míticas y religiosas; y por ello se trata de un dogma profundamente arraigado en el lenguaje ordinario y presente en numerosas locuciones.

El contraste entre lógica formal y lógica informal o del lenguaje ordinario aparece bien marcada asimismo en Strawson, perteneciente también al grupo de Oxford. En el curso de una obra sobre el carácter de la teoría lógica, Strawson ha comparado repetidamente la lógica formal, caracterizada por reglas exactas, precisas y rígidas, con la lógica del discurso común, carente de semejante exactitud, pero notablemente más compleja y multiforme, razón por la que no puede esperarse de la primera que recoja o sistematice del todo las reglas y distinciones características del lenguaje ordinario (*Introducción a la teoría lógica*, cap. 2, seccs. 5 y 16, y cap. 8, seccs. 7 y 8). Por lo mismo, el discurso común nos pondrá en presencia de relaciones no reducibles a la deducibilidad y de incoherencias no asimilables a la contradicción lógica. Una de esas relaciones es precisamente la de presuposición, desarrollada por Strawson como parte de su crítica al tratamiento logicista del lenguaje.

Según Strawson, un enunciado *A* presupone un enunciado *B* cuando la verdad de *B* es una condición necesaria tanto para la verdad de *A* como para su falsedad (*op. cit.*, cap. 6, secc. 7). Si la verdad de *B* fuera una condición necesaria exclusivamente para la verdad de *A*, tendríamos que decir que de *A* se deduce *B*, y estaríamos en presencia de la relación de deducibilidad que los lógicos formalizan; es la relación que hay entre «Todos los hombres son mortales» (*A*) y «El hombre que escribió *Fragmentos de apocalipsis* es mortal» (*B*). Pero hay casos en que un enunciado es condición no sólo de la verdad, sino también de la falsedad de otro. Así, según el ejemplo que vimos en las secciones 6.6 y 6.7, la verdad del enunciado

- (1) Hay una entidad única que es el segundo satélite natural de la Tierra,

es condición tanto para la verdad como para la falsedad de

- (2) El segundo satélite natural de la tierra se encuentra a 800.000 kilómetros de ésta

ya que no tiene sentido intentar averiguar si esta afirmación es verdadera o falsa a menos que la primera sea verdadera. Strawson dirá que (2) presupone (1), y defenderá que esta manera de enfocar el problema de la relación entre ambos enunciados hace innecesaria la teoría de las descripciones de Russell, la cual, por contraste, ha de resultar extremadamente artificiosa. Como se recordará, la teoría de Russell exige, en el caso de que la oración (1) sea falsa, que (2) también lo sea, ya que esta última equivale, según él (secc. 6.6), a

- (3) Hay una entidad única que es el segundo satélite natural de la Tierra y se encuentra a 800.000 kilómetros de ésta

Strawson, en cambio, dirá que, si (1) es un enunciado falso, entonces no tiene sentido plantear la cuestión de si (2) es verdadero o falso, pues cuestión semejante sólo puede plantearse sobre el supuesto de que (1) es cierta, y añadirá que analizar (2) como equivalente a (3) no es más que una fuente de problemas típicos de quienes intentan reducir el lenguaje ordinario a los estrechos moldes de la lógica formal, y que viene al fin a dar en el absurdo de considerar como únicos nombres, en sentido lógico, los pronombres demostrativos.

La crítica de Russell por Strawson se halla en el artículo de éste «Sobre la referencia» («On Referring», 1950). Lo que nos interesa ahora no es tanto la crítica como tal (de hecho, fue implícitamente recogida en la sección 6.7) cuanto el espíritu con el que está hecha y la diferente actitud frente al lenguaje que ella ejemplifica. Y ese espíritu es, claramente, el de una vindicación del lenguaje ordinario frente a las estrechas exigencias de la lógica. Hay que notar, por cierto, que en su artículo no aparece todavía el término «presuposición» que Strawson empleará posteriormente y que, de forma más elaborada, utilizan hoy filósofos y lingüistas para el análisis de ciertos aspectos del lenguaje; Strawson recurre en su artículo al término «implícitar» (*imply*), cuidando de añadir que se trata de «un sentido muy especial y extraño» de dicho término, y ciertamente distinto del que tiene cuando se trata de la implicación lógica (*op. cit.*, secc. III). El término «presuponer» (*presuppose*), que introduce en su *Introducción a la teoría lógica*, constituye, desde luego, una mejor elección, pues la propia distinción terminológica contribuye a manifestar que se trata de una relación peculiar al discurso ordinario y ajena al ámbito de la lógica formal. Esta diferencia de nivel —llamémosla así— se relaciona con otra distinción que es típica de Strawson y bien característica de un filósofo del lenguaje ordinario: la distinción entre una oración (*sentence*), el uso de una oración (*use*) y la proferencia o emisión de una oración (*utterance*), y otro tanto por lo que toca a las expresiones que son parte de una oración («Sobre la referencia», secc. II). Así, una oración o una expresión puede ser proferida en diversas ocasiones o por diferentes personas, y tendremos entonces otras tantas proferencias de la expresión u oración de que se trate; pero distintas personas, o la misma en distintas ocasiones, pueden referirse con una expresión al mismo objeto, en cuyo caso habrá que decir que hacen de ella el mismo uso, y que hacen un uso diferente sólo cuando se refieran a distinto objeto o entidad. Pues bien, la idea de Strawson es que la verdad y la falsedad han de predicarse del uso de las oraciones, pero no de las oraciones mismas, así como la referencia ha de predicarse del uso de las expresiones y no de las expresiones mismas. Veamos un ejemplo: la oración «El presidente del gobierno español nació en Cebreros» no es por sí verdadera ni falsa ni tiene sentido pretender que lo sea, igual que la expresión «El presidente del gobierno español» no se refiere por sí sola a nadie ni a nada. Esta expresión adquiere referencia al ser usada para decir algo sobre alguien, así como la oración anterior adquirirá un valor de verdad al ser usada en una ocasión determinada. El uso de tal oración en 1975

hubiera constituido un enunciado falso, y su uso en 1977, por el contrario, hubiera constituido un enunciado verdadero, pues la expresión que en dicha oración hace de sujeto habría sido usada para referirse a una persona distinta en cada ocasión. Lo que es verdadero o falso no es la oración, sino la aserción o enunciado (*statement*) hechos al usarla, o, dicho de otro modo, la proposición expresada por la oración en ese uso. Por ello, si la expresión que hace de sujeto no tiene referencia en una ocasión determinada, la oración de la que forma parte no podrá tener, a su vez, valor de verdad. Así, la expresión «El presidente de la República española» no podía usarse literalmente en 1980 para referirse a nadie, por lo que la oración «El presidente de la República española es un médico» tampoco podía usarse en ese tiempo para hacer un enunciado, y si alguien la hubiera usado, su aserción hubiera carecido, en rigor, de valor veritativo.

En resumen: la verdad y la falsedad no son propiedades de oraciones, sino de usos de oraciones; referirse no es algo que haga una expresión, sino algo que hace un hablante cuando usa esa expresión. En última instancia, estas consideraciones son parte de una crítica al referencialismo de la doctrina de Russell sobre el lenguaje, críticas muy semejantes a las que hemos visto dirigidas por Wittgenstein también contra el *Tractatus*. Según Strawson, Russell confundió significar con referirse, pero significar es «al menos en un importante sentido, una función de la oración o de la expresión», y dar el significado de una expresión es, en este sentido, «dar directrices generales para su uso en la referencia o mención de objetos o personas particulares», así como dar el significado de una oración es «dar directrices generales para su uso al hacer afirmaciones verdaderas o falsas» («Sobre la referencia», secc. II). La semejanza entre esta forma de hablar del significado y las consideraciones del segundo Wittgenstein es patente. La cuestión es: ¿se gana algo realmente hablando de «uso» y no de «oración»? Es claro que una oración, entendida como entidad abstracta, como eso que hemos llamado en la sección 2.2 «signo tipo», no es ni verdadera ni falsa, pero no es menos claro que cuando los filósofos, incluido Russell, hablan de oraciones verdaderas o falsas se refieren a oraciones concretas, en cuanto utilizadas en una ocasión determinada para hacer afirmaciones, esto es, en cuanto acontecimientos o signos concretos. Recurrir, como hace Strawson, a ejemplos del tipo de «El presidente del gobierno español» contribuye a hacer sus distinciones más recomendables, pues no puede dudarse que la referencia de esta expresión cambiará según el momento en que se use, y distinguir entre la expresión y su uso permite dar razón cómodamente de esos cambios de referencia y, por tanto, también de los cambios de valor veritativo en las oraciones en que tal expresión aparezca como sujeto. Pero todo esto puede conseguirse igualmente distinguiendo entre la expresión tipo, que como tal carecerá de referencia, y la expresión utilizada, que la tendrá o no según cuándo se use. Empeñarse, entonces, en que no es una expresión la que se refiere sino un hablante por medio de una expresión, no pasaría de ser empeño bizantino, pues una expresión utilizada es una expresión en cuanto pronunciada, escrita o entendida por un hablante.

Una argumentación en todo similar justificaría, por su parte, la atribución de valores veritativos a las oraciones en cuanto utilizadas. La justificación para estas distinciones, aun cuando no desaparece, resulta más cuestionable si, en lugar de ejemplos como los considerados, pensamos en expresiones como «El presidente del gobierno español durante 1975», cuya referencia, evidentemente, no depende del uso que hagamos de ella, ni puede cambiar de un tiempo a otro. Igualmente, la oración «El presidente del gobierno español durante 1975 es socialista» es falsa para quienquiera que la use y cuandoquiera que se use, y no se ve qué ventaja puede haber aquí en recurrir al uso. Pues es claro que la referencia de la expresión que hace de sujeto en esta oración viene determinada por su significado y es independiente del contexto de su uso (siempre que éste sea literal, y excluyendo, por tanto, el uso literario, poético, figurado, etc.). En suma: ¿es mejor decir, como Frege, que las expresiones tienen sentido y referencia, o afirmar, como Strawson, que tienen sólo sentido (o significado) y que la referencia es propia de su uso? Creo que ambos enfoques son compatibles con tal que atribuyamos referencia a las expresiones en cuanto utilizadas y que no olvidemos que una expresión tipo no es más que la abstracción de una expresión.

La cuestión de a qué atribuyamos sentido, referencia y valores de verdad es, sin embargo, cuestión cuya solución no prejuzga que aceptemos o no la teoría russelliana de las descripciones. Aún hemos de resolver qué hacer cuando el sujeto de una oración carece de referencia, por ejemplo, si alguien afirma en 1980 «El presidente de la República española es un filósofo». ¿Vamos a considerar esta oración así utilizada como falsa, o simplemente como carente de valor veritativo? Por las razones que ya se indicaron en la sección 6.7, considerarla como falsa y analizarla de acuerdo con la teoría de las descripciones me parece demasiado artificioso. Aquí hay que reconocer, con Strawson, que, *desde el punto de vista del lenguaje ordinario*, resulta más natural considerarla como una oración que, en cuanto usada en esa fecha, no es ni verdadera ni falsa. ¿Impide esta decisión un tratamiento lógico de esa oración? Bueno, no se ve muy bien qué interés puede tener el tratamiento lógico de oraciones declarativas que carezcan de condiciones veritativas; pero supongamos que alguien está muy interesado en averiguar qué proposiciones pueden deducirse de esa oración o cuáles son incompatibles con ella: en tal caso podemos asumir por hipótesis que la oración fuera verdadera y proceder en consecuencia. O bien podemos formalizar esa oración en un sistema lógico trivalente y asignarle el tercer valor de verdad. El tratamiento lógico de una oración siempre depende de los procesos de inferencia que queramos considerar, y éste es un punto de vista lo bastante restringido como para que no haya de afectar a una consideración lingüística que atienda al modo de funcionar la oración en el discurso ordinario. Por esta razón, hay que coincidir con Strawson cuando da fin a su artículo con estas palabras: «Ni las reglas aristotélicas ni las de Russell proporcionan la lógica exacta de ninguna de las expresiones del lenguaje ordinario, pues el lenguaje ordinario carece de una

lógica exacta.» Pero, por eso mismo, tampoco puede refutarse desde el lenguaje ordinario una doctrina lógica como la teoría de las descripciones. Todo lo más, podrá mostrarse que aplicada al lenguaje común resulta confusiva o artificiosa, y que se pueden conseguir los mismos o mejores resultados por medios más económicos. (Strawson no parece consciente de que su crítica a Russell debe moverse dentro de estos márgenes de relatividad, y menos aún lo es Russell en la respuesta que le da en el capítulo final de *La evolución de mi pensamiento filosófico*, pero la relatividad de esta disputa está sugerida en el bello artículo de Lemmon, «Sentences, Statements and Propositions».)

Se habrá notado que en lo anterior nos hemos topado con otro uso del término «uso». Tal y como Strawson emplea el término, se hace un uso de una expresión para referirse a un objeto o persona particular. Strawson se cuida de anotar esta peculiaridad, y escribe en una nota (al comienzo de la sección II de su artículo): «Este uso (*usage*) de 'uso' (*use*) es, desde luego, distinto de: (a) el uso (*usage*) corriente en el que 'uso' (*use*) (de una palabra, frase u oración particular) es igual (aproximadamente) a 'reglas para usar', que a su vez es igual (aproximadamente) a 'significado'; y de (b) mi propio uso (*usage*) en la frase 'uso referencial individualizador (*uniquely referring use*) de las expresiones', en el cual 'uso' es igual (aproximadamente) a 'manera de usar'.» (El lector reconocerá que el parecido que este párrafo tiene con un trabalenguas no es culpa mía.) En la acepción (a), «uso» tiene un sentido que me parece muy próximo al que hemos visto que le da Wittgenstein. En la acepción (b), que es más restringida, se aproxima al carácter conceptual que «uso» tiene en Ryle. Tenemos, pues, una tercera acepción propia de Strawson, y es una acepción importante puesto que en ella el uso aparece como aquello que da referencia a una expresión y que convierte una oración en un enunciado, esto es, en algo capaz de ser verdadero o falso. Que las tres acepciones de «uso» tienen mucho en común salta a la vista, y que las diferencias pueden ser sólo de matiz no hace falta subrayarlo. ¿Por dónde pasa la línea que distingue las reglas para usar una expresión de la manera de usarla? De otra parte, ¿no ha reconocido Strawson que las reglas de uso suministran directrices generales para el uso referencial de las expresiones y para el uso de las oraciones que consiste en hacer enunciados verdaderos o falsos? Aquí podemos empezar a preguntarnos si un concepto como el de uso, con tan variadas manifestaciones, que, por otra parte, quedan sumidas en esta vaguedad, constituye realmente para el estudio del lenguaje un instrumento más útil que los viejos conceptos de sentido y referencia.

7.6 Cómo hacer cosas con palabras

Esa atención al lenguaje común, de la que acabamos de ver algunos ejemplos, adquiere en Austin tal primacía que parece en ocasiones independizarse de cualquier otro propósito filosófico ulterior. Aunque muerto

prematuramente, sus artículos y conferencias han ejercido tan grande influencia en la filosofía inglesa, y en particular dentro del círculo de Oxford, que en buena medida las características de esta escuela coinciden con las de la doctrina de Austin.

Sus declaraciones metodológicas no parecen apartarse de las que hemos visto en el segundo Wittgenstein. Así, reconoce que las palabras son herramientas, que debemos estar preparados contra las trampas que nos tiende el lenguaje, y que el conjunto de las expresiones que usamos incorpora todas las distinciones y conexiones que los hombres, a lo largo de muchas generaciones, han creído que valía la pena hacer («Alegato en pro de las excusas», pp. 129-130 del original y 174 de la trad. cast.). Esto justifica suficientemente el interés por el lenguaje ordinario. Pero Austin añade además que no hay que confundir las palabras con las cosas, y que es necesario comparar aquéllas con éstas a fin de percatarnos de la arbitrariedad y de la inadecuación de nuestras expresiones, de manera que podamos contemplar el mundo sin anteojeras (*ibidem*). El propósito último es, por consiguiente, obtener una visión correcta de la realidad, y puesto que esta visión, si ha de ser intersubjetiva y comunicable, ha de expresarse en palabras, lo primero es conocer bien lo que significan las palabras y lo que por medio de ellas podemos decir. Esto no implica que la tarea filosófica acabe aquí; los significados de las expresiones en su uso común derivan, fundamentalmente, de las necesidades prácticas de la vida cotidiana, y no cabe esperar que satisfagan todas nuestras necesidades de expresión teórica. No hay que olvidar que «la superstición, el error y la fantasía de todo tipo se hallan incorporados al lenguaje ordinario, y a veces incluso sobreviven al paso del tiempo» (*loc. cit.*, p. 133 del original y 177 de la trad. cast.). Y añade: «Ciertamente, el lenguaje común no es la última palabra: en principio, puede ser siempre completado, mejorado y sustituido; recuérdese tan sólo que es la primera palabra» (*loc. cit.*). Es patente que, en este punto, Austin va más allá de Wittgenstein y es más claro que Ryle.

En la obra publicada de Austin, esa «primera palabra», esto es, el estudio del lenguaje ordinario, ocupa casi todo su tiempo y su atención. Teniendo en cuenta su temprana muerte, y que casi todas sus obras estaban sin acabar y sin pulimentar, sería injusto pensar que fue infiel a su método o que no fue capaz de sacarle más rendimiento. Pero lo cierto es que Austin ha pasado a la historia de la filosofía del lenguaje por su ocupación sistemática, minuciosa y artesanal con el lenguaje, ocupación para la que sugirió el nombre de «fenomenología lingüística» (*op. cit.*, p. 130). Su idea era que, cuando se trata de estudiar un concepto filosófico, el primer paso es considerar las diversas expresiones comunes que tienen relación con tal concepto y observar de qué modo se usan y cómo se vinculan entre sí. Por esta razón recomendaba buscar un ámbito del lenguaje abundante en giros y expresiones, y lo más ajeno posible a las teorías filosóficas, que de otra forma podrían haberlo contaminado con su jerga técnica (*ibidem*). De aquí que, en el artículo que vengo citando, escogiera el campo lingüístico de las excusas como objeto de estudio relevante para una elucidación de los con-

ceptos de libertad, responsabilidad y acción, y por ello, en general, para toda la ética. Como método de trabajo recomendaba Austin formar grupos de investigación que se dividieran el trabajo y recopilaran conjuntamente las expresiones y usos del campo elegido; el grupo tenía además una importante ventaja a la que Austin era muy sensible: permitía contrastar intersubjetivamente las afirmaciones de cada uno acerca de los significados y usos de las palabras, y suministraba, por tanto, cierta garantía sobre la exactitud de las afirmaciones. No parece que Austin tuviera una confianza decisiva y última en las llamadas «intuiciones del hablante nativo», a las que otros pensadores, como Searle (*Actos de habla*, cap. 1, secc. 3), parecen constituir en piedra angular de la metodología. Austin, al menos, usaba con frecuencia del diccionario y así lo recomienda a los demás. Tampoco tenía inconveniente en recurrir a usos lingüísticos técnicos, y no meramente comunes, cuando pudiera ser de ayuda; por lo que toca a las excusas, recomienda tener en cuenta el discurso jurídico y la psicología científica («Alegato en pro de las excusas», pp. 134-137 del original y 178-180 de la traducción castellana).

La principal aportación de Austin a la teoría del significado es, sin duda, su análisis de los actos de habla (*speech acts*). La idea originaria es que una buena parte de las expresiones que usamos las usamos para hacer algo por medio de ellas, para realizar algún acto (se entiende: distinto del mero acto de decir algo). Es peculiar de estas expresiones que en ellas aparezca un verbo en primera persona. Así, decimos «Sí, acepto» para realizar el acto de contraer matrimonio, y en respuesta a la pregunta «¿Aceptas como esposo (o esposa) a..., etc.?»; o también: «Prometo que...» para hacer una promesa; o «Te bautizo... con el nombre de ...» para bautizar; o «Fallo que...» (o «Fallamos que...») para emitir un fallo judicial; e, igualmente, «Te aconsejo que...» para dar un consejo, o «Te pido perdón» para pedir perdón, etc. A este tipo de expresiones, en cuanto así utilizadas, las llamaba Austin «preferencias realizativas» (*performative utterances*); lo que en ellas interesa destacar no es la expresión como tal, sino su emisión o pronunciación para realizar el acto de que se trate: bautizar, prometer, aconsejar, etc. Naturalmente que si pensamos en las expresiones correspondientes en cuanto usadas de esta manera, podemos sin confusión hablar de «expresiones realizativas». Al llamar la atención sobre estas expresiones, Austin intentaba hacer una aportación al estudio de los usos lingüísticos. En una clara alusión a Wittgenstein, aunque sin citarlo, Austin declara irónicamente que los filósofos, en cuanto pueden enumerar diecisiete usos del lenguaje, empiezan a hablar de los *infinitos* usos del lenguaje, pero incluso aunque fueran diez mil —añade— con tiempo se podrían enumerar todos; no basta con invocar un nuevo uso del lenguaje cada vez que deseamos escapar a un problema filosófico: lo que necesitamos es «un marco en el que discutir estos usos del lenguaje» («Performative Utterances», p. 221).

El concepto de expresión realizativa pretende contribuir a suministrar ese marco. De hecho, sin embargo, no lo consigue enteramente. La razón

es que Austin contrapone ese tipo de expresiones a las que llama «constatativas» (*constative*), esto es, a las que pueden ser verdaderas o falsas, a los enunciados («Performative-Constative», p. 22). En un aspecto, la contraposición es válida, desde luego. Tiene sentido decir que la afirmación «Prometo que...» no describe nada, y que no puede ser ni verdadera ni falsa; no dice que las cosas sean de tal o cual manera ni que suceda esto o lo otro; simplemente, esa afirmación, en las condiciones usuales (esto es: supuesta la sinceridad del que habla, excluyendo que se pronuncie en el curso de una representación teatral, etc.), constituye la realización del acto de prometer. Una afirmación, en cambio, como «El promete que...» describe algo que ocurre, es propiamente un enunciado o proferencia constativa (hace constar cómo son las cosas), y puede ser verdadera o falsa. Pero en otro aspecto, la contraposición no es aceptable. Pues supongamos que en lugar de decir «El promete que...», digo: «Te informo de que él promete que...». Parece que lo coherente sería afirmar aquí que esta última es también una expresión realizativa, y que ejecuta el acto de informar, aunque su contenido sea un enunciado y ese enunciado pueda ser verdadero o falso. Pero si esa última expresión es realizativa, entonces también lo es la primera, puesto que yo puedo realizar el acto de informar diciendo

(1) Te informo de que él promete que vendrá

o bien diciendo simplemente

(2) El promete que vendrá

Con lo que resulta que (2), a pesar de ser un enunciado (y como tal verdadero o falso), es *al mismo tiempo* una oración realizativa. Por consiguiente, las oraciones constativas o enunciados son también realizativas y la contraposición originaria es incorrecta. A las oraciones realizativas como (1), las llamaba Austin «explícitas», y a las que son como (2) las denominaba «primarias», «primitivas» o «implícitas» («Performative Utterances», p. 231, y *Palabras y acciones*, conferencias III y IV). Los ejemplos pueden multiplicarse; puedo dar un consejo diciendo

(3) Te aconsejo que lo pienses bien

o afirmando meramente

(4) Piénsalo bien

Puedo hacer una promesa empleando las palabras

(5) Prometo que vendré

o limitándome a decir con laconismo

(6) Vendré

si por otras circunstancias está claro que se trata de una promesa. Las oraciones (3) y (5), como (1), son explícitamente realizativas, mientras que (4) y (6), como (2), lo son de manera implícita o primaria.

De este modo, la teoría de los actos de habla, entendidos como los tipos de actos que realizamos, implícita o explícitamente, por medio del lenguaje, le suministra una interpretación más rigurosa para la doctrina de los usos lingüísticos, si bien a condición de extender aquélla a todas las expresiones y de desarrollarla de modo sistemático. Esto es lo que hizo Austin en una serie de conferencias póstumamente publicadas bajo el título *How to Do Things with Words* (la expresividad de este título: «Cómo hacer cosas con palabras», se pierde un tanto en el título de la traducción castellana: *Palabras y acciones*, aunque éste es un título que, en inglés, usó a veces el propio Austin para sus conferencias). Más de la mitad del libro está dedicada a explorar las posibles diferencias entre las expresiones realizativas y las constativas, pero ante la imposibilidad de encontrar criterios definitivos que apoyen la contraposición, Austin replantea el problema de forma más radical preguntándose: ¿en qué sentidos puede afirmarse que decir algo es hacer algo? La respuesta a esta pregunta le llevará a introducir unas distinciones que han sido de vasta influencia no sólo en la filosofía del lenguaje sino también en la teoría lingüística posteriores.

Al acto de decir algo, en el sentido usual de esta expresión, Austin lo llama «acto locucionario» (*locutionary act*), y en él distingue tres aspectos: primero, el acto fonético, que consiste en pronunciar ciertos sonidos; segundo, el acto fático (*phatic*), que consiste en la pronunciación de unos sonidos en cuanto pertenecientes a un léxico y regulados por una gramática, y, por último, el acto rético (*rhetic*), que consiste en pronunciar esos sonidos (expresiones, palabras) con un sentido y con una referencia más o menos determinados (*Palabras y acciones*, conferencia VIII). Aquí lo primero que llama la atención es que se llame «acto» a cada uno de los aspectos de un acto locucionario; parecería que realizar un acto locucionario es hacer al tiempo tres actos, fonético, fático y rético. Es claro, sin embargo, que no es esto lo que Austin quiere decir: se trata, simplemente, de distinguir tres aspectos sucesivamente más abstractos en el acto de decir algo. Cuando tenemos este acto completo, tenemos la pronunciación de unas expresiones lingüísticas con sentido y referencia, o, dicho de otra forma, tenemos una oración en uso (y esta acepción de «uso» me parece que coincide con la que es peculiar de Strawson, y que ya conocemos). Si abstraemos el sentido y la referencia (aspecto rético), nos quedamos con unos sonidos gramaticales, pero sin significado. Es el caso de alguien que pronuncia una expresión cuyo significado ignora. Si abstraemos finalmente la gramaticalidad (aspecto fático), lo que resta son simplemente sonidos, que o bien no constituyen palabras de una lengua, o bien, si lo son, no se hallan gramaticalmente estructuradas. Tendremos entonces un fenómeno meramente fonético, y como tal aún no lingüístico, en el sentido del lenguaje verbal.

Mas con esto tan sólo hemos hablado del acto *de* decir algo; nada se ha mencionado todavía de lo que se hace *al* decir algo, esto es, de lo que antes se ha llamado «acto de habla». A esto lo denominaré ahora Austin «acto ilocucionario» (*illocutionary act*), y es claro que se realiza al mismo tiempo que el acto locucionario. ¿Cómo distinguir entre uno y otro? Curiosamente, la distinción involucra, de nuevo, el concepto de uso: «Para determinar qué acto ilocucionario se realiza, debemos determinar de qué manera estamos usando la locución» (*loc. cit.*, p. 98 del original y 142 de la trad. cast.). A continuación, y como ejemplos de maneras de usar las expresiones, da la siguiente lista: preguntar o responder; dar una información, dar seguridad, hacer una advertencia; anunciar un veredicto o una intención; dictar sentencia; hacer una cita, o una apelación, o una crítica; identificar o dar una descripción. Al lector le habrá recordado inmediatamente la lista de juegos lingüísticos que ofrece Wittgenstein en la sección 23 de las *Investigaciones filosóficas*, pues tiene con ella un gran parecido. Austin se muestra consciente de la vaguedad enorme que aqueja a la expresión «manera de usar», y que hace que pueda aplicarse con justicia no sólo a los actos ilocucionarios, sino también a los locucionarios e incluso a un tercer tipo de actos que serán mencionados en seguida. En todo caso, resaltar la ambigüedad del término «uso» como la del término «significado» es parte de la doctrina de Austin (*loc. cit.*, p. 100 del original y 145 de la trad. cast.). Ahora tenemos que un primer sentido del uso de una oración es el que consiste en usarla para decir algo sobre algún objeto o sobre alguien: es el acto *de* decir algo o acto locucionario. Un segundo sentido del uso de una oración es el que se da cuando la usamos para responder a una pregunta, para advertir, para criticar, para describir, para aconsejar..., esto es, cuando realizamos un cierto acto *al* decir algo: es el acto ilocucionario. Una vez más debe evitarse la confusión de pensar que se trata de dos actos distintos. Cuando aconsejo a alguien «Piénsalo bien», no hago dos cosas: darle un consejo y decirle que lo piense bien; realizo en realidad un solo acto: el de aconsejarle que lo piense bien. Pero podemos abstraer en este acto las palabras usadas, como podemos abstraer, por otro lado, el tipo de acto para el que han sido usadas. Tanto más cuanto que esas mismas palabras podrían haberse utilizado para realizar otro tipo de acto, por ejemplo, para formular un ruego. De hecho, Austin sugiere una terminología menos confundente: sugiere hablar más bien de «fuerzas ilocucionarias» que de «actos»; diríamos entonces que las palabras «Piénsalo bien» tenían en tal ocasión la fuerza de un consejo, mientras que en tal otra poseían, por el contrario, la fuerza de un ruego (*loc. cit.*, p. 99 del original y 144 de la trad. cast.). Esto permite contrastar la fuerza con el significado. El significado, en cuanto integrado por el sentido más la referencia, pertenece al aspecto locucionario del acto verbal; la fuerza, en cambio, pertenece a su aspecto ilocucionario.

Austin distingue todavía un tercer tipo de acto que puede realizarse por medio de las palabras, el que llama «acto perlocucionario» (*perlocutionary act*). Consiste en los efectos que el acto verbal produzca «en los senti-

mientos, pensamientos o acciones del auditorio, del hablante o de otras personas» (*loc. cit.*, p. 101 del original y 145 de la trad. cast.). Podemos ahora comparar los tres tipos de actos al modo de Austin: si le comunico a alguien «Me dijo 'Piénsalo bien'», estaré haciendo referencia a un acto locucionario; si, por el contrario, digo «Me aconsejó que lo pensara bien», estaré mencionando un acto ilocucionario; si, finalmente, afirmo «Me convenció de que lo pensara bien», entonces hablaré de un acto perlocucionario. A diferencia del acto *de* decir algo, y del acto realizado al decir algo, el acto perlocucionario es el acto realizado *por*, o a consecuencia de, decir algo. Así, son actos perlocucionarios convencer, decepcionar, desahogarse, impresionar, y toda la amplia variedad de efectos que puede tener el uso del lenguaje sea en el auditorio sea en el hablante. Nótese que también el acto perlocucionario puede incluirse bajo un concepto tan amplio como el de uso: podemos afirmar, sin duda, que el lenguaje se usa para convencer, para impresionar, para desahogarse...; pero esto es de índole muy diferente a decir que el lenguaje se usa para advertir, para prometer, para mandar... ¿Cómo expresar esta diferencia de manera rigurosa? Austin sugiere que el uso ilocucionario es, a diferencia del uso perlocucionario, convencional, en el sentido de que puede hacerse explícito por medio de una fórmula realizativa como las que ya conocemos. De este modo se puede decir: «Te advierto que...», «Prometo que...», «Te mando que...», etc.; por el contrario, ello no es posible con el uso perlocucionario, esto es, no se puede decir: «Te convengo de que...», «Te impresiono de que...», «Me desahogo de que...», «Te decepciono de que...», etc. En otras palabras: el acto perlocucionario no se realiza de modo convencional (*loc. cit.*, p. 103 del original y 147-148 de la trad. cast.). La triple distinción de Austin, por consiguiente, introduce orden y rigor en el maltratado concepto de uso lingüístico, aunque naturalmente no agote todos los sentidos o acepciones en que puede tomarse. Sin duda, muchos de los ejemplos de la lista de Wittgenstein (*Investigaciones*, secc. 23) escapan a las categorías de Austin, pero eso es algo de lo que éste es plenamente consciente (*loc. cit.*, p. 104 del original y 148 de la trad.): el uso del lenguaje para hacer chistes (que Wittgenstein cita) y el uso poético del lenguaje (que Wittgenstein no cita) son usos lingüísticos en un sentido que es ajeno al de las categorías de Austin. Puede haber, por otro lado, casos que, aun quedando cubiertos por su clasificación, sea dudoso dónde incluirlos, si en el aspecto ilocucionario o en el perlocucionario: Austin cita, como ejemplos, insinuar y expresar emociones. Es importante tener en cuenta que el acto perlocucionario, puesto que consiste meramente en las consecuencias de lo que se dice, no está determinado por las convenciones lingüísticas ni tiene por qué corresponder a la intención del hablante; así, mis palabras pueden, en contra de mi intención, atemorizar a mi auditorio, y esto por causa de ciertas características del contexto extralingüístico que pueden ser ajenas a mis palabras y a mi intención.

Es patente que la nueva clasificación suministra un instrumento de análisis mucho más riguroso que la inicial distinción entre expresiones rea-

lizativas y expresiones constativas. Ahora tenemos que todas las expresiones, en cuanto usadas o proferidas, son realizativas, bien de modo explícito, bien de modo implícito. Y que la verdad y la falsedad, que al principio se atribuían exclusivamente a las expresiones constativas, hay que atribuir las ahora a las expresiones cuando éstas realizan determinados actos ilocucionarios. Así, la oración

(7) Hay un toro en el prado

puede ciertamente ser verdadera o falsa en una ocasión determinada, ya se use para informar, para advertir, para formular una conjetura, para expresar acuerdo con lo que otra persona ha dicho, etc. Ello significa que si sustituimos esa oración implícitamente realizativa que es (7) por una que lo sea de forma explícita, como

- (8) Informo de que hay un toro en el prado
- (9) Sospecho que hay un toro en el prado
- (10) Advierto que hay un toro en el prado

etcétera, tendremos que decir que lo que es verdadero o falso no son estas oraciones en cuanto tales, sino su contenido locucionario, aquello que se sospecha, se advierte o de lo que se informa. Las oraciones realizativas explícitas, como tales, no son ni verdaderas ni falsas, sino, en términos de Austin, felices o infelices, según que se den todos los requisitos necesarios para que efectivamente realicen, en la ocasión de que se trate, el acto ilocucionario que expresan.

Por todo lo anterior, una clasificación general de los actos de habla se convierte en una clasificación de las fuerzas ilocucionarias o actos ilocucionarios, y para confeccionarla puede tomarse como guía una clasificación de los verbos realizativos, esto es, de los verbos que sirven para introducir una oración realizativa explícita. Austin sugiere las cinco clases siguientes (*op. cit.*, conferencia XII).

1. **Veredictivos** (*verdictives*). Típicamente, estos actos consisten en la emisión de un veredicto, no sólo como sentencia, sino también como estimación o evaluación. Austin considera de este tipo los siguientes verbos: absolver, condenar, valorar, describir, diagnosticar, estimar, medir, caracterizar...

2. **Ejercitativos** (*exercitives*). Consisten en el ejercicio del poder, del derecho o de la influencia, expresando una decisión en favor o en contra de una manera de actuar. Austin da estos ejemplos: nombrar (para un cargo), destituir, mandar u ordenar, avisar, rogar, aconsejar, vetar, promulgar, perdonar, recomendar...

3. **Compromisorios** (*commissives*). Consisten en comprometerse con una forma de acción. Por ejemplo: prometer, proponerse, hacer voto de,

jurar, tener la intención de, dar palabra de, consentir, apoyar, oponerse, declarar la intención de... Aquí nota Austin que existen indudables diferencias entre las diversas formas de compromiso y las meras declaraciones de intención o propósito; lo que entre unas y otras hay de común es la anticipación realizada por el hablante de lo que será su comportamiento en el futuro, y esto es lo que le da unidad al presente grupo de actos de habla.

4. Comportativos (*behabitives*). Se trata, según Austin, de una categoría muy variada, que tiene que ver con el comportamiento social, y que incluye la noción de reacción a la conducta ajena y a los avatares de la vida de los demás; están, por ello, muy ligados a la expresión de los sentimientos. Entre los ejemplos de Austin se encuentran: disculparse, dar las gracias, dar la enhorabuena, dar el pésame, felicitar, criticar, quejarse, dar la bienvenida, desear que lo pase bien, maldecir, bendecir, brindar, desafiar...

5. Expositivos (*expositives*). Consisten en expresar de qué modo encajan nuestras expresiones en una conversación o en una discusión. Austin se muestra sumamente intranquilo respecto a este grupo, pensando que sus ejemplos podrían ser incluidos también en una u otra de las categorías anteriores, y declara no estar en claro él mismo sobre este punto. De sus ejemplos cabe señalar: afirmar, negar, enunciar, subrayar, responder, describir, informar, suponer, aceptar, retirar, objetar, corregir, preguntar, argüir, deducir, definir, explicar, referirse, querer decir...; y coloca una interrogación junto a estos otros: saber, creer y dudar.

En esta exposición he elegido, de los ejemplos de Austin, aquellos que me parecen más claros y, en todo caso, los que tienen una mayor relevancia filosófica. El lector debe tener en cuenta que Austin suministra muchos ejemplos más, en su mayoría considerablemente más oscuros y debatibles que los que hemos visto. La clasificación es, sin duda, muy fácil de criticar, pero dado que el propio Austin no la presenta más que como aproximación y con toda clase de cautelas, la crítica tiene un interés relativo. Puede afirmarse, en primer lugar, que buena parte de los expositivos podrían incluirse en otras categorías, y especialmente entre los veredictivos (de hecho, «describir» aparece en ambos grupos), pero acabamos de ver que Austin es consciente de esta insuficiencia. En segundo lugar, no se advierte que la clasificación obedezca a ningún criterio general, y la caracterización de los grupos resulta en exceso vaga y heterogénea, pero tampoco se muestra Austin ajeno a esta objeción. En tercer lugar, y principalmente, hay que preguntarse si la consideración de los verbos realizativos es realmente la mejor guía para una clasificación de los actos de habla. En mi opinión, Austin está aquí todavía demasiado sujeto a su concepto original de expresión realizativa, lo que curiosamente contrasta con el hecho de que varios de los ejemplos que ofrece difícilmente podríamos aceptarlos como verbos realizativos. Considérese, por ejemplo, «querer decir». Podemos comenzar una afirmación así: «Quiero decir que...»; pero querer decir no

es un tipo de actos de habla, sino más bien una condición necesaria para la realización de cualquier acto de habla. Tampoco se ven razones suficientes para aceptar que saber, creer, dudar, suponer, proponerse, y tener una intención sean tipos de actos de habla; lo único que tendría sentido considerar como acto de habla es la expresión lingüística del saber, de la creencia, de la duda, de la suposición, de la determinación o de la intención. Aquí, Austin parece haber incurrido en el error de creer que todo verbo que pueda iniciar en primera persona del presente de indicativo una oración, es un verbo realizativo. Y es claro que no. Es cierto que coloca una interrogación junto a «creer», «dudar» y «saber», pero no parece dudar respecto a «tener la intención de» o «proponerse», y es patente que esto no constituyen actos de habla, actos que se realicen por medio del lenguaje, aun cuando la declaración de una intención o de un propósito sí lo sea. Quiero decir lo siguiente: la expresión «Tengo el propósito de venir» no es ilocucionaria, aunque sí lo es la expresión «Declaro mi propósito de venir»; el acto de habla es declarar la intención, no el tenerla; sin embargo, Austin acepta ambos como actos de habla en el grupo de los compromisorios.

7.7 Actos de habla

Las dificultades precedentes, y otras más que por sutiles o por dudosas no he mencionado, han conducido a Searle a proponer una clasificación alternativa de los actos ilocucionarios («Una taxonomía de los actos ilocucionarios»). Esta propuesta forma parte, no obstante, de una teoría general de los actos de habla que se aparta, en aspectos importantes, de la posición de Austin. Lo primero que rechaza Searle es la distinción entre el aspecto locucionario y el aspecto ilocucionario, distinción que también a otros les ha parecido difícilmente defendible o explicable (véase Strawson, «Austin and 'Locutionary Meaning'», y Hare, «Austin's Distinction between Locutionary and Illocutionary Acts»). La razón principal de Searle es que, puesto que puede distinguirse entre actos ilocucionarios genéricos y específicos, los actos locucionarios pueden quedar reducidos al tipo más general de acto ilocucionario, a saber: decir. *Decir* sería también, por tanto, un acto ilocucionario, el más genérico, y no habría posibilidad de identificar el acto locucionario como realidad separable del acto ilocucionario (Searle, «Austin on Locutionary and Illocutionary Acts», secc. II). Otra razón adicional, que Searle también señala (*loc. cit.*, secc. I), es la de que en las oraciones explícitamente realizativas el significado determina ya el acto ilocucionario que se realiza, y sería, por consiguiente, imposible abstraer separadamente el acto locucionario y el acto ilocucionario; así, en la oración «Te ruego que me escuches», su propio significado determinaría que se trata de un ruego, y al abstraer su aspecto locucionario (y en particular, su aspecto rético), estaríamos abstrayendo al mismo tiempo su fuerza ilocucionaria.

En mi opinión, estas razones, que naturalmente Searle desenvuelve de forma más prolija en su artículo, no son suficientes para rechazar la distinción de Austin. Ciertamente su concepto de acto locucionario adolece de todas las imprecisiones y vaguedades que aquejan al concepto de significado, e incluso al concepto de gramaticalidad, puesto que ambos van incluidos en el aspecto locucionario. Es cierto, también, que en el caso de los realizativos explícitos la fuerza ilocucionaria es parte del significado y, por tanto, que no puede separarse de éste. Pero para todas las demás oraciones que no son explícitamente realizativas la distinción es aplicable y posee un sentido suficientemente claro. Lo que ya no me parece aceptable es extender el concepto de acto ilocucionario de tal modo que podamos afirmar que *decir* es un acto de este tipo. El concepto de acto ilocucionario pretende suministrar un instrumento para distinguir y clasificar los diversos tipos de actos que podemos ejecutar cuando decimos algo, y de nada sirve a este propósito considerar el decir mismo como acto ilocucionario. Decir es lo que hay de común en todos los actos de habla, y el presupuesto para que pueda realizarse un acto de habla particular, sea el que fuere; por ello, en decir es en lo que consiste propiamente el acto locucionario, y en este punto la posición de Austin está bien fundada. Congruentemente con su negación de esta distinción, Searle reformula los conceptos básicos de la teoría, aceptando como tipos de actos de habla los siguientes (*Actos de habla*, secc. 2.1). En primer lugar, los actos que consisten en proferir palabras u oraciones; segundo, los actos proposicionales, como los llama Searle, que consisten en la predicación y en la referencia; y tercero, los actos ilocucionarios, en el sentido que ya conocemos. Como en el caso de Austin, se trata aquí también no de actos separados, sino de aspectos que pueden abstraerse en el acto de habla total y efectivo. La teoría de Searle, desarrollada con muchos matices en su libro, tiene un interés propio y dificultades asimismo características, en especial por lo que toca a lo que llama «actos proposicionales», pero no nos es posible demorarnos más en este tema. Examinaremos únicamente su clasificación de los actos ilocucionarios.

En un esfuerzo por ofrecer una clasificación razonada y sistemática, Searle ha mencionado una serie de doce criterios con los que se puede diferenciar unos actos ilocucionarios de otros («Una taxonomía de los actos ilocucionarios», secc. II). Me detendré en los que tienen mayores consecuencias y citaré otros de pasada. En primer lugar, el propósito del acto ilocucionario; dicho propósito es común a todos los actos de un cierto tipo; así, las órdenes y los ruegos tendrían en común servir al propósito de intentar que el oyente haga algo, a diferencia, por ejemplo, de una descripción, cuyo propósito sería representar cómo es algo. En segundo término, la dirección de ajuste entre las palabras y el mundo; este criterio deja abiertas tres alternativas: o bien las palabras pretenden ajustarse al mundo, como ocurre con las descripciones, o bien pretenden que sea el mundo el que se ajuste a ellas, como es el caso de las peticiones, o bien no hay relación de ajuste, como acontece en el caso de los saludos. El ter-

cer criterio es lo que denomina «condición de sinceridad», esto es, lo que se refiere al estado psicológico que el acto ilocucionario, si es sincero, manifiesta. Así, las descripciones, enunciados, explicaciones, etc., manifiestan las creencias del hablante; las promesas, juramentos y amenazas manifiestan sus intenciones; las órdenes, peticiones y ruegos, sus deseos, etc. Los tres criterios referidos son, para Searle, los más importante. Brevemente cita, además, otros cuantos. Por ejemplo, las diferencias de fuerza en la ilocución (así, la diferencia entre una promesa y una mera declaración de intención); las diferencias de posición jerárquica entre el hablante y el oyente, cuando son relevantes (por ejemplo, para distinguir si se trata de una orden o de una petición); las diferencias del modo en que las palabras se relacionan con los intereses de hablante o del oyente (y que distinguen, por ejemplo, una felicitación de un pésame); la diferencia entre que el verbo ilocucionario tenga un uso realizativo o que no lo tenga («mandar» lo tiene: «Te mando que...», pero «amenazar» no, pues no se amenaza diciendo «Te amenazo de que...»); y otras clases más de diferencias que no es necesario detallar. Vale la pena, por cierto, reparar en la última de las citadas, pues la forma en que Austin ha presentado su clasificación de las fuerzas ilocucionarias parece sugerir la idea, ciertamente errónea, de que a toda fuerza ilocucionaria le corresponde un verbo realizativo.

Sobre la base de los criterios precedentes, y en particular de los tres primeros, Searle propone la siguiente clasificación de actos ilocucionarios (*loc. cit.*, secc. IV):

1. Representativos. Su propósito es comprometer al hablante con que algo es de tal o cual manera; su dirección de ajuste es de las palabras hacia el mundo, como se muestra en que los actos de este tipo sean valorables como verdaderos o falsos; el estado psicológico al que corresponden, cuando se ejecutan sinceramente, es la creencia. Esta clase incluye, según Searle, muchos de los actos que Austin clasificaba bien como veredictivos, bien como expositivos.

2. Directivos. Cumplen el propósito de intentar (en diverso grado) que el oyente haga algo; el ajuste va desde el mundo a las palabras; el estado psicológico correspondiente es el deseo. Se incluyen aquí: ordenar, pedir, preguntar, rogar, permitir y aconsejar.

3. Compromisorios. El propósito de éstos es comprometer al hablante (también en diferentes grados) con un comportamiento futuro; el ajuste se produce, como en el caso anterior, desde el mundo hacia el lenguaje; la condición de sinceridad es que el hablante tenga la intención de obrar como dice.

4. Expresivos. Sirven al propósito de expresar el estado psicológico del hablante, y por ello pertenecen a este grupo: agradecer, disculparse, dar el pésame, felicitar, etc.; según Searle, carecen de dirección de ajuste,

puesto que aquí no se trata ni de hacer corresponder las palabras a la realidad ni de hacer que esta última corresponda al lenguaje, ya que, si el acto es sincero, siempre corresponderá al estado psicológico que exprese; la condición de sinceridad es variable, puesto que son diversos los estados psicológicos que pueden expresarse por medio de un acto ilocucionario.

5. Declaraciones. Su propósito es producir una modificación en una situación determinada, creando una situación nueva. Pertenecen a este grupo: nombrar (a alguien para un cargo), dimitir, cesar (a alguien en un cargo), declarar la guerra, declarar marido y mujer, bautizar, etc. Aquí, piensa Searle, la dirección de ajuste es recíproca, puesto que, al ejecutar el acto con éxito, se crea la nueva situación y, por consiguiente, el lenguaje y la realidad se ajustan mutuamente. Lo que falta, en cambio, es la condición de sinceridad, pues el estado psicológico del hablante es irrelevante en este tipo de actos y sus palabras no pretenden, por ello, manifestarlo. Se habrá notado que esta categoría incluye muchos de los ejemplos más típicos de actos de habla, precisamente los que Austin solía citar cuando inició su estudio de las expresiones realizativas. También se habrá reparado en que los actos de este tipo declarativo suelen darse en el marco de instituciones sociales determinadas, especialmente jurídicas y políticas.

A las declaraciones, en el sentido anterior, agrega Searle un grupo mixto que denomina «declaraciones representativas». Se trata de aquellos actos representativos ejecutados por una autoridad en cuanto tal, y por tanto con carácter de declaración, esto es, de actos en los que la representación de los hechos tiene carácter definitivo por realizarse en el marco de determinadas instituciones. Así, el juez declara probados tales hechos, o el árbitro declara que el balón salió fuera o que el jugador se encontraba dentro del área. En estos casos, la declaración no crea la situación, sino que la reconoce de modo oficial y definitivo; el balón no sale fuera porque lo diga el árbitro, pero, si así lo afirma, todo procede como si efectivamente hubiera salido fuera, aun cuando no haya ocurrido de esta manera. En estos casos, como es natural, sí hay condición de sinceridad, y es la creencia, pues el acto presupone que la autoridad cree lo que dice.

Lo anterior es propiamente una clasificación de los actos ilocucionarios, pero no de los verbos ilocucionarios. Este es un punto importante, y que contribuye al rigor y a la exactitud de la taxonomía de Searle. Ello significa que hay verbos ilocucionarios que no designan tipos de actos ilocucionarios; así, por ejemplo, «insistir» y «sugerir». Puedo decir: «Insisto en que vengas»; «Insisto en que es tarde»; «Insisto en que te agradezco lo que has hecho»; «Insisto en que voy a venir». Tenemos sucesivamente un acto directivo, un acto representativo, un acto expresivo y un acto compromisorio. ¿Qué es lo que añade, entonces, ese «insisto»? Según Searle, marca el grado de intensidad con el que se presenta, en el acto, el propósito ilocucionario. Compárense los ejemplos anteriores con «Sugiero que vengas», «Sugiero que es tarde», etc. (excepto en el tercer caso; no se dice «Sugiero que te agradezco...»). Otros verbos ilocucionarios

pueden, a su vez, señalar otros aspectos del acto distintos del propósito; así, los verbos «responder», «anunciar», «considerar», «advertir»... Por no aludir al propósito ilocucionario, estos verbos no constituyen tampoco tipos diferenciados de actos ilocucionarios.

Las virtudes de la clasificación de Searle son tan patentes que no necesitan recomendación, en particular cuando se la compara con la de Austin. No significa esto, sin embargo, que todo en ella esté claro. En primer lugar, hay dos ejemplos del grupo de los actos directivos sobre los que vale la pena llamar la atención: «preguntar» y «permitir». Sobre el carácter semántico de las preguntas se ha discutido mucho en los últimos años; ya en la sección 7.2 vimos esta cuestión suscitada por Wittgenstein, y como se recordará éste resuelve la cuestión remitiéndola a la variedad de los usos que podemos hacer de una pregunta (*Investigaciones*, 24). Ahora tenemos que, según Searle, preguntar es una especie de acto ilocucionario perteneciente a la clase de los directivos. Esto implica que se trata de un uso del lenguaje con el que se pretende que el oyente haga algo, a saber: responder. Toda pregunta, por lo tanto, equivaldrá a una oración del tipo de «Dime si...» o «Dime qué...». Básicamente, esta concepción me parece correcta. Habría que dar cuenta, luego, del uso retórico de las preguntas así como de su utilización en el monólogo; pero no creo que esto presente dificultades, pues tales usos del lenguaje no pueden explicarse, en general, más que con referencia al uso intersubjetivo y comunicativo, que es el primario. Por lo que respecta al permiso, el asunto me parece más debatible. No entiendo cómo pueda pretenderse que permitir es intentar que el oyente haga algo; el lector recordará, acaso, que esta misma posición la vimos ejemplificada por Schaff, en la sección 2.3, al interpretar el significado de la luz verde de un semáforo. Pero permitirle a alguien que haga algo no es decirle que lo haga, es remover una prohibición existente y, por consiguiente, se trata de un acto que pertenece al grupo de las declaraciones, en el sentido de Searle, ya que crea la situación de permiso. Me refiero con esto, claro está, al uso personal del término «permitir». Cuando se usa impersonalmente, como en la frase «Está permitido...», el acto realizado pertenece, en cambio, a la clase de los representativos, puesto que meramente consiste en manifestar que no existe prohibición al respecto. En ningún caso veo que se puedan construir los permisos como actos directivos.

Un punto más importante es el de la dirección de ajuste en los actos expresivos y en las declaraciones (grupos 4 y 5). Los primeros, según Searle, carecen de dirección de ajuste, pues quien expresa su estado psicológico no pretende que sus palabras se ajusten a la realidad, ni tampoco, por supuesto, que sea la realidad la que se ajuste a sus palabras. Yo diría, sin embargo, que esto es así porque el acto sólo es posible si se da ese ajuste. Expresar agradecimiento, condolencia, felicitación, etc., sólo es posible (cuando se es sincero) si existe el estado psicológico que las palabras pretenden transmitir; por consiguiente, el ajuste es inmediato. Yo no tengo que preocuparme de que mis palabras se ajusten a mi estado psico-

lógico; si hablo sinceramente, si mis palabras expresan mi actitud, entonces ya están ajustadas por principio. Como en los actos representativos, la dirección de ajuste es desde el lenguaje hacia la realidad (el estado psicológico), pero, a diferencia de los actos representativos, el ajuste se da inmediatamente, pues no hay acto expresivo sin ese ajuste (en cambio, los actos representativos pueden darse sin ajuste; una descripción no es menos descripción porque sea falsa). Algo semejante diría yo acerca de las declaraciones. De éstas, Searle piensa que tienen doble dirección de ajuste, ya que, al quedar creada la situación por las palabras, lenguaje y realidad se ajustan recíprocamente. Pero tampoco en este caso veo justificada la posición de Searle, pues parecería que no hay ninguna prioridad temporal entre la declaración y la situación que ésta crea. Sin embargo, es claro que primero son las palabras, y sólo como resultado de éstas, en ciertas condiciones, queda creada la nueva situación; luego es la realidad la que se ajusta al lenguaje. Únicamente que, a diferencia de lo que ocurre en los actos directivos y en los compromisarios, el ajuste se produce de modo inmediato y sin concurso posterior de ninguno de los participantes en el acto de habla. Se verá claro del todo con un ejemplo. Si aconsejo a una pareja de amigos «Casaos», la realidad quedará ajustada a mis palabras en el caso de que sigan mi consejo y se casen; si le comunico a alguien «Me voy a casar», la realidad se ajustará a mi declaración de intención si efectivamente cumplo mi propósito; pero si, como ministro de una religión, o como juez, a una pareja que ha venido a casarse, y que por lo demás cumple con los requisitos exigidos, le digo «Os declaro marido y mujer», la realidad queda ajustada a mis palabras por el mero hecho de haberlas pronunciado en las condiciones exigidas, y sin necesidad de ningún comportamiento ulterior, mío o de los otros participantes. En este sentido digo que la realidad queda ajustada al lenguaje de forma inmediata. Pero son las palabras las que crean la situación, y por tanto aquéllas preceden a ésta, y en consecuencia es la realidad la que se ajusta al lenguaje.

Lo anterior no son sino objeciones de detalle. La cuestión más importante es la que se refiere a los criterios de la clasificación. La taxonomía de Searle se basa en los tres criterios que ya hemos visto: el propósito ilocucionario, la dirección del ajuste entre el lenguaje y la realidad, y la condición de sinceridad. Consideremos, en primer lugar, este último criterio. Hay casos en los que parece estar claro: que, ejecutados sinceramente, los actos representativos correspondan a una creencia, o que los compromisarios correspondan a una intención, no parece disputable. Pero ciertamente no se trata de un criterio muy decisivo, puesto que en el caso de los actos expresivos la condición de sinceridad va referida a una variedad indeterminada de estados psicológicos posibles (¿también la creencia y la intención?; Searle no lo aclara), y en las declaraciones dicha condición no existe. En cuanto a los actos directivos, mientras que parece razonable afirmar que el estado psicológico relevante para la mayor parte de ellos sea el deseo o el querer, no veo que tal cosa no pueda mantenerse ni para los consejos ni para los permisos. Sobre estos últimos ya he hablado

antes, para proponer su inclusión entre las declaraciones sacándolos del grupo de los directivos donde Searle los coloca, y la razón era que no se trata de actos por los que se pretenda que alguien haga algo; de aquí se desprende *a fortiori* que no implican ningún deseo por parte del hablante. Por lo que toca a los consejos, aun cuando aquí se trate de que el oyente haga algo, que por cierto se supone que va en su propio interés, en modo alguno puede mantenerse que un consejo sincero haya de expresar el deseo de quien lo da; mi consejo no es peor ni menos sincero por el hecho de que me sea indiferente lo que haga el destinatario e incluso de que yo no desee que haga lo que le aconsejo que haga. El consejo se da en beneficio de su destinatario, y un consejo sincero, precisamente por serlo, puede muy bien ir en contra de los deseos del que lo da.

El segundo criterio es la dirección de ajuste entre las palabras y el mundo. Puesto que es un criterio que recurre a la relación entre el lenguaje y la realidad se trata de un criterio propiamente semántico. Si éste fuera el principal criterio para clasificar los actos de habla, tendríamos que sospechar que lo que estábamos clasificando eran clases de expresiones más que tipos de actos. Nos queda, no obstante, el primero de los criterios expuestos por Searle, que es, en su entender, el más importante, el del propósito ilocucionario del acto. Los diferentes propósitos determinan, en la forma que hemos visto, los distintos tipos de actos ilocucionarios; cada clase de actos se caracteriza por el propósito típico al que responde. La cuestión es: el propósito ¿de quién? La respuesta obvia es: del hablante, naturalmente. Pero se habrá observado que Searle habla siempre del propósito del acto: del propósito de una descripción, de una orden, de una promesa... ¿Qué sentido tiene esto? Pienso que lo que le interesa a Searle es, de todos aquellos propósitos posibles que puede tener alguien al realizar un acto de habla, aquel que es común a todos los actos de una misma clase y que puede servir para caracterizarla. Aquí, Searle tiene buen cuidado de excluir aquellos propósitos que consisten en la producción de determinados efectos, pues tales propósitos tendrían carácter perlocucionario, y no servirían, por ello, para caracterizar actos ilocucionarios. Esta es la razón por la que no atribuye Searle a los actos directivos el propósito de conseguir que alguien haga algo, sino únicamente el propósito de *intentar* que alguien haga algo; por lo mismo, el propósito de los actos representativos no es convencer al oyente de cómo son las cosas, sino simplemente exponer la opinión del hablante al respecto. Pese a estas cautelas, sin embargo, el criterio no es aplicable de modo inequívoco y coherente. Ello obedece, a mi entender, a que se trata de un criterio más psicológico que lingüístico. Considérese, por ejemplo, el caso de los actos directivos. Parece que tiene buen sentido afirmar que su propósito es intentar que el hablante haga algo; pero es al mismo tiempo patente que tal propósito se lo podemos atribuir igualmente a actos de otro tipo. Para aconsejar a alguien que estudie filosofía (un consejo un tanto perverso, pensarán algunos), puedo decirle simplemente: «Estudia filosofía» (acto directivo); pero puedo también decirle: «Todos los que estudian filosofía se sienten

felices» (acto representativo, por más que su verdad resulte dudosa); o bien: «Si estudias filosofía, te prestaré gran ayuda en tu estudio» (acto compromisorio). Podría, incluso, henchir los pulmones, sonreír de oreja a oreja, y exclamar: «¡Qué feliz me siento yo de haber estudiado filosofía!» (acto expresivo). De cualquiera de estos actos podemos decir que tiene como propósito intentar que mi oyente estudie filosofía.

Naturalmente que, puestos a defender a Searle, se podría señalar que sólo en el caso del acto directivo tal intento es propiamente un propósito ilocucionario, puesto que, en los demás actos ejemplificados, ese propósito tiene carácter ulterior y mediato. Así, podríamos subrayar que, en el acto representativo, el intento de que mi oyente estudie filosofía es un propósito mediato y secundario respecto al propósito inmediato y propio de ese acto que es, a saber, representar a los estudiosos de la filosofía como seres felices. E igualmente podríamos declarar que, en el acto compromisorio, mi propósito inmediato es comprometerme a prestar ayuda a mi oyente en el caso de que estudie filosofía, y sólo de forma mediata e indirecta es mi propósito intentar que estudie filosofía. Etcétera...

Pero la cuestión es: ¿por qué puede atribuirse un propósito propio, directo e inmediato al acto de habla, distinguiéndolo de otros propósitos mediatos o indirectos que el acto pueda cumplir? O dicho en otros términos que el propio Searle ha empleado (*Actos de habla*, cap. III): ¿por qué pueden distinguirse entre condiciones que son esenciales para los actos ilocucionarios y otras que no lo son? La única respuesta que me parece rigurosa es ésta: el propósito inmediato, propio y tipificador de un acto de habla es el que puede atribuirse al hablante en función de lo que significan las palabras empleadas en el acto y del modo de realizarse el ajuste entre ellas y la realidad. Por esta razón se justifica que el acto de decirle a alguien «Si estudias filosofía te ayudaré en tu estudio» tenga como propósito directo comprometer al que habla, aun cuando pueda tener también el propósito indirecto de intentar que el oyente estudie filosofía. De lo que significan las palabras se infiere que expresan la intención del hablante y que, por tanto, el ajuste entre ellas y el mundo se producirá cuando, dándose la condición indicada, el hablante actúe de conformidad con la intención expresada. En última instancia, dado que la relación entre el acto ilocucionario y el estado psicológico del hablante no constituye un criterio claro ni aplicable en todos los casos, y ya que el propósito ilocucionario tampoco parece poseer contornos más precisos, lo que nos queda es el significado gramatical de las expresiones utilizadas en el acto de habla y la forma de relacionarse éstas con el mundo. Esto invita a preguntarse: ¿no estaremos clasificando más bien tipos de expresiones que clases de actos? En todo caso, las cinco clases de actos ilocucionarios aceptadas por Searle no parecen ser otra cosa que cinco tipos muy generales de funciones que cabe atribuir al lenguaje, a saber: representar cómo son las cosas y cuáles son los hechos; decir a alguien que haga algo; comprometer al hablante a hacer algo; expresar un estado psicológico; crear una situación.

Se observará que, de estas funciones, hay tres que aparecen recogidas en las tipologías usuales: la función representativa, la función expresiva y la función directiva. Así es precisamente la antigua, y bien conocida, distinción de Karl Bühler entre las funciones semánticas del signo lingüístico (*Teoría del lenguaje*, secc. 2.2). El signo es símbolo en tanto que cumple una función de representación de los hechos y las cosas, es síntoma en cuanto posee la función de expresar algo del emisor, y es señal en la medida en que su función es apelar al oyente para dirigir su conducta. Las expresiones lingüísticas, en tanto que usadas, cumplen simultáneamente las tres funciones, aunque predomina una u otra sobre las demás según los casos. En este punto, la teoría de los actos ilocucionarios no difiere, pues, como Searle reconoce, es perfectamente normal realizar varios de esos actos al mismo tiempo («Actos de habla indirectos», primera pág.; por otra parte, éste es el supuesto del ejemplo que acabamos de ver sobre el consejo de estudiar filosofía). La clasificación de Searle añade la función compromisoria y la función creativa. Podría intentarse la reducción de ambas a las categorías anteriores; por ejemplo, reduciendo la función compromisoria a una forma de la función directiva (el hablante dirigiría su propia conducta), y la función creativa a una especie de la función representativa (la expresión lingüística que crea una situación se convertiría en una representación inmediata y directa de la situación creada). No obstante, los cinco tipos de Searle ponen al descubierto particularidades importantes del funcionamiento del lenguaje que la clasificación de Bühler, más simplista, tiene que describir con menos recursos.

La idea de que las funciones son concurrentes, aunque predomine una sobre otra, se encuentra asimismo en la clasificación de Jakobson, que incluye también las tres funciones de Bühler, a las que llama, respectivamente, función referencial, función emotiva y función conativa. A ellas añade otra más: la función fática, ya señalada por Malinowski, que consiste en utilizar el lenguaje para mantener la comunicación, por medio de frases hechas y fórmulas ritualizadas; la función metalingüística, con la que el lenguaje se torna sobre sí mismo; y la función poética, que es la que adquiere el lenguaje cuando se lo usa desde el punto de vista del estilo, de la sonoridad y, en suma, de la estética («Lingüística y poética», páginas 353 ss.). La función metalingüística puede incluirse en la categoría de los actos representativos, puesto que decir cómo son las cosas incluye decir cómo es el lenguaje. En cambio, las funciones fática y poética sin duda constituyen maneras de usar el lenguaje de las que no da cuenta la clasificación de los actos ilocucionarios que estamos viendo, y que vale la pena tener en cuenta en una teoría general de los actos de habla.

El concepto de función del lenguaje no suele quedar demasiado claro en tipologías como las que acabo de recordar, y en esta medida podría considerarse problemático un acercamiento de ese concepto al de tipo de acto ilocucionario. Pero si los tipos de actos ilocucionarios no pretenden ser otra cosa que modos generales de usar el lenguaje, entonces no se ve qué diferencias sustanciales puede haber en principio entre clasificar aquellos

o clasificar las funciones del lenguaje, excepto que, a través del concepto de acto ilocucionario, se ha conseguido una precisión en este tema que las consideraciones usuales sobre las funciones del lenguaje no habían alcanzado. Desde el punto de vista de este momento que estamos estudiando ahora en la evolución de la teoría del significado, lo que podemos buscar, tanto en una clasificación de los actos ilocucionarios como en una clasificación de las funciones del lenguaje, es una restricción rigurosa que introduzca orden y precisión en la descripción de los usos lingüísticos, la cual había dejado Wittgenstein sumida en una considerable oscuridad. Que el propósito último de la clasificación de Searle es precisamente ése, queda declarado por él mismo en los siguientes términos: «La conclusión más importante que hay que sacar de esta discusión es la siguiente. No hay, como Wittgenstein (en una posible interpretación) y otros muchos han pretendido, un número infinito e indefinido de juegos lingüísticos o usos del lenguaje. Más bien, la ilusión de usos ilimitados del lenguaje se engendra a causa de una enorme falta de claridad acerca de lo que constituyen criterios para delimitar un juego o uso del lenguaje de otro.» («Una taxonomía de los actos ilocucionarios», última pág.) Adoptando criterios como los de Searle, tendríamos que todos los usos, pese a su aparente variedad, se reducen a los cinco tipos generales citados: representar la realidad, decir que se haga algo, comprometer al que habla, expresar un estado psicológico, y crear una situación. A estos usos podemos añadir algunos otros, acaso marginales, como el uso fático y el uso estético, que hemos visto recogidos por Jakobson. Ya antes comprobamos que, en Austin, la distinción entre los aspectos locucionario, ilocucionario y perlocucionario, y su división de las fuerzas ilocucionarias en cinco clases, obedecía también al propósito de introducir orden y limitación en la idea de uso lingüístico. La infinita diversidad de los juegos lingüísticos, que para Wittgenstein parecía desafiar todo intento teórico de definición y clasificación, queda así reducida a los estrechos márgenes de unas pocas categorías. (Tanto la teoría de Austin como la de Searle, así como la crítica de aquél por este último, la inclusión de la teoría de los actos de habla en la lingüística transformatoria, que aquí no podemos detenernos a considerar, y otros muchos puntos de vista y aportaciones a este tema, están tratados en una interesante obra de conjunto como es la de Holdcroft, *Words and Deeds. Problems in the Theory of Speech Acts*; por lo demás, la teoría de los actos de habla es uno de los temas que más bibliografía han originado durante los últimos diez años, así en la filosofía del lenguaje como en la lingüística, aunque mi impresión es que la abundancia de la bibliografía se debe más bien a la imprecisión de la teoría que a su rendimiento.)

7.8 Tipos de discurso

Varios de los criterios que hemos visto utilizados en las dos secciones precedentes para distinguir entre clases de fuerzas o actos ilocucionarios

vamente los imperativos y las oraciones en las que se usan expresiones deónticas, esto es, expresiones como «debes», «es obligatorio», «te prohíbo», etc. (con excepciones como la que se acaba de indicar en el paréntesis anterior). Quedan fuera del discurso directivo los juicios de valor en sentido propio, o sea, las oraciones en las que típicamente se usan términos como «bueno», «malo», «correcto», «bello», etc. En cambio, Hare incluía, como hemos visto, los juicios valorativos dentro del discurso prescriptivo; hay que aclarar que Hare considera juicios de valor tanto a las oraciones con términos valorativos («bueno», etc.) como las que tienen términos deónticos («debes», etc.), aunque reconoce que, para ciertos propósitos, deben distinguirse ambas variedades (*Freedom and Reason*, 2.8, nota). Para Ross, la diferencia fundamental se halla en que los juicios propiamente valorativos, si bien expresan actitudes a favor y en contra de las cosas, de las personas y de las acciones, no presentan ninguna forma definida de comportamiento, como lo hacen los juicios normativos o deónticos y los imperativos (*Lógica de las normas*, secc. 9). Ross reconoce, asimismo, la existencia de otro tipo peculiar de discurso, ajeno a las categorías anteriores, el que forman las exclamaciones, cuyo significado es típicamente emotivo (*op. cit.*, secc. 19).

En Ross, por tanto, el criterio fundamental es semántico, si bien éste, a su vez, determina diferencias sintácticas y pragmáticas. Y es semántico porque recurre al significado; aquí, el significado parece estar tomado como sentido en la acepción freguiana: el significado de una oración indicativa es un tema o asunto presentado como real, mientras que el significado de una oración directiva es la idea de una acción presentada como algo a realizar.

Desde un punto de vista diverso, y varios años antes, Morris había intentado una clasificación de los tipos de discurso que, utilizando criterios parcialmente coincidentes con los anteriores, excede con mucho por los detalles sobre lo que hemos visto. La propuesta de Morris se produce en el marco de su teoría conductista de la semiótica, de la que ya di noticia en el capítulo 2. Morris recurre a un doble criterio de clasificación, el modo de significar de los signos complejos (expresiones, oraciones) y el propósito de su uso. Según esto, los signos pueden significar de los siguientes modos distintos (*Signos, lenguaje y conducta*, cap. 3): de modo designativo, cuando el signo significa características del entorno del organismo que interpreta el signo; de modo apreciativo, cuando el signo significa para su intérprete una categoría preferencial de algo, respecto a lo cual le dispone para reaccionar en favor o en contra; de modo prescriptivo, cuando el signo significa para el intérprete la necesidad de un proceso de conducta; de modo formativo, cuando el signo significa propiedades de las combinaciones de signos, y es, por consiguiente, por lo que respecta al lenguaje, un modo metalingüístico de significar. Morris reconoce un quinto modo de significar, si bien éste carece de relevancia para la determinación de los tipos de discurso; es el modo identificativo, en el cual un signo significa la colocación de algo en el espacio y en el tiempo. Como se habrá notado,

los modos de significar dependen de cómo opere el signo para el intérprete.

El intérprete puede ser, naturalmente, tanto el propio productor o emisor del signo como el oyente. Cuando se adopta el punto de vista del productor y se atiende al propósito con el que ha producido sus signos, pueden distinguirse cuatro usos primarios de un signo (*op. cit.*, cap. 4): el uso informativo, que no requiere más aclaración; el uso valorativo, que consiste en otorgar a algo estado preferencial; el uso incitativo, cuando se persigue provocar una reacción, de sí mismo o de otro; y el uso sistemático, con el cual se busca organizar la conducta dentro de un todo. Como se ve, existe una patente analogía entre cada uno de estos usos y los cuatro modos principales de significar que ya conocemos, y Morris no deja de señalarla (cap. 4, secc. 2). El modo designativo parece especialmente indicado para informar; el modo apreciativo resulta en particular apropiado para valorar; el modo prescriptivo sería el adecuado para incitar; y, finalmente, aunque sólo sea por exclusión, el modo formativo correspondería al propósito sistemático. Sin duda, es lo referente a estas dos últimas categorías lo que puede resultar más difícil de entender, pero lo consideraremos con ejemplos dentro de un momento. La idea de Morris es que, no obstante esa correspondencia, cada modo de significar puede combinarse con cada uno de los usos, proporcionando cada combinación un tipo de discurso que vendrá definido, según eso, por el modo de significar de los signos y por el uso dado a éstos. En total habrá, pues, dieciséis tipos de discurso. Cada uno de estos tipos los ejemplifica Morris de acuerdo con la tabla siguiente (cap. 5).

Uso:	Informativo	Valorativo	Incitativo	Sistemático
<i>Modo:</i>				
Designativo	Científico	De ficción	Legal	Cosmológico
Apreciativo	Mítico	Poético	Moral	Crítico
Prescriptivo	Tecnológico	Político	Religioso	De propaganda
Formativo	Lógico-matemático	Retórico	Gramatical	Metafísico

Antes de apresurarse a la crítica debe tenerse en cuenta que los ejemplos suministrados por Morris no son sino eso: ejemplos. Como él mismo manifiesta (*op. cit.*, cap. 5, secc. 1), el discurso científico es sólo un caso de discurso designativo-informativo, pero no tiene por qué ser único; y lo mismo para los demás tipos. Dicho esto, las cuestiones principales que plantea esa tabla pueden dividirse en dos órdenes: de un lado, si las dieciséis categorías que resultan de la combinación de uso y modo están bien definidas, son suficientemente unívocas y carecen de solapamientos

importantes; de otro, si bastan para caracterizar los ejemplos citados, u otros que pudieran darse.

Por lo que respecta a la primera cuestión, comencemos por examinar los criterios empleados por Morris. A primera vista, se combina un criterio semántico, la manera de significar, con un criterio pragmático, el propósito del uso. Sólo esto debería otorgar a la clasificación de Morris una superioridad sobre las que hemos considerado antes, las cuales eran más imprecisas respecto a los criterios a aplicar y al mismo tiempo más simples en cuanto a las categorías resultantes. Ahora bien, si examinamos con detención los modos de significar, advertiremos que los tres primeros son del todo parejos con los tres tipos de significado que determinan, en Ross, los correspondientes tipos de discurso. El modo designativo de Morris corresponde al significado indicativo de Ross, el modo apreciativo corresponde al discurso valorativo, y el modo prescriptivo al significado directivo (para no mal interpretar este paralelismo hay que recordar que el libro de Morris es más de veinte años anterior al de Ross, y debe añadirse, además, como curiosidad, que no hay una sola mención de Morris en el libro de aquél). Resta el modo formativo, en el cual se significan las propiedades de las combinaciones de signos tal y como éstas se muestran en signos como la negación, la disyunción, los términos de cuantificación, los numerales, etc. Si nos limitáramos a los modos de significar, no obtendríamos, pues, más que los tres tipos principales de discurso antes discutidos, con la adición de un nuevo tipo, el discurso metalingüístico de carácter lógico-matemático. ¿Qué añaden las cuatro clases de usos introducidas por Morris? En principio, cada uso parece la proyección pragmática de cada uno de los modos de significar: el uso típico de un signo designativo sería informar, el de un signo apreciativo consistiría en valorar, y el de un signo prescriptivo se resumiría en incitar. Menos claro resulta en el último caso, pues ciertamente no se ve por qué el uso típico de un signo formativo habría de ser sistematizar la conducta, pero la relación en los tres primeros casos es patente. La peculiaridad de la clasificación de Morris, y su aparente riqueza y complejidad deriva de entrecruzar los modos con los usos. Y es justamente aquí donde empiezan las dificultades.

¿En qué consiste, por ejemplo, un tipo de discurso en el que las palabras significan designativamente y se usan para valorar? Resulta difícil imaginarlo. Morris sugiere como ejemplo el discurso de ficción, el discurso literario narrativo, razonando que por medio de él se pretende inducir en el lector actitudes preferenciales frente a su contenido, personajes, o hechos, circunstancias, etc. Que esto sea, sin embargo, una característica de las narraciones literarias en cuanto tales es del todo dudoso, aunque, desde luego, haya cierto género de narraciones que se caracterizan de esa forma; pero en estas últimas, su condición les viene precisamente de ciertas consideraciones que en ellas hace el autor y que justamente constituyen expresiones con significado apreciativo y prescriptivo, y no meramente designativo. Cuando el propósito es incitar a una forma de conducta y el significado de las palabras es designativo, tenemos un tipo de discurso

como es, por ejemplo, el de la ley. ¿Qué es lo que designan aquí las palabras? Según Morris, los castigos que la comunidad reserva para quienes actúen o dejen de actuar de ciertas maneras. En este punto, la posición de Morris puede aceptarse por lo que toca a las leyes penales. Al menos en España, el *Código Penal* no afirma, por ejemplo, que no se deba matar a otro, sino que el que matare a otro será condenado, como homicida, a la pena de reclusión menor (artículo 407). Parece claro, en efecto, que aquí las palabras simplemente designan, y que su propósito principal es el de incitar a los ciudadanos a que no maten. Pero no se ve cómo podría decirse lo propio a propósito de otras leyes. Así, cuando el *Código Civil* estipula que son españoles los hijos de padre español (art. 17), ¿a qué incita? O cuando establece que un medianero no puede sin consentimiento del otro abrir ventana en la pared medianera (art. 580), ¿qué es lo que designa? El último tipo de discurso designativo, aquel cuyo propósito es sistemático, lo ilustra Morris con el ejemplo del discurso cosmológico, entendiendo por tal el discurso filosófico sobre el mundo. Según esto, la ciencia coincidiría con la cosmología filosófica en significar de modo designativo, pero se distinguiría en que la primera tiene como propósito informar mientras que la última pretende organizar la conducta en un todo sistemático. Esta distinción, sin embargo, oculta las relaciones históricas entre la cosmología filosófica y la ciencia, y el hecho palmario de que la primera pierde todo su sentido a partir del momento en que la ciencia consigue elaborar una explicación amplia y comprensiva del proceso del cosmos. Parecería, según lo que sugiere la clasificación de Morris, que ambas son compatibles puesto que responden a propósitos diversos, pero aparte de que la diferencia decisiva entre ellas no sea la que deriva de sus diferentes propósitos, tampoco queda demostrado que la cosmología responda al propósito que Morris le asigna.

Dudas análogas pueden repetirse en torno a casi todos los demás ejemplos de la tabla. Que las expresiones poéticas tengan un significado apreciativo y sirvan a un propósito valorativo no puede por menos de resultar arbitrario. Morris cita en su favor una poesía de Whitman que ciertamente responde a esa caracterización apreciativo-valorativa, ¡pero cuántas poesías podrían recordarse que carecen de expresiones apreciativas y de propósito valorativo! ¿Cómo justificar que las expresiones morales tienen significado apreciativo, y en cambio las religiosas lo tienen descriptivo, y por qué no más bien al contrario? ¿Cómo puede olvidarse toda la porción de discurso designativo que toda religión incluye, y que suministra al creyente una cierta representación del universo? Nótese el anacronismo de considerar el discurso gramatical como incitativo por su propósito, lo que supone que la gramática no trata de cómo es una lengua, sino de cómo debe emplearse. En fin, y como absurdo superlativo, el lector no habrá pasado por alto que el discurso metafísico queda caracterizado en la tabla como compuesto de expresiones cuyo significado es formativo y, por consiguiente, que tratan de las propiedades de las combinaciones de signos. El razonamiento de Morris en este punto se apoya en la idea tradicional de

que las afirmaciones metafísicas son necesariamente verdaderas, y por tanto se hallan más allá de las posibilidades de una confirmación empírica. Si esto es así —piensa Morris— entonces las expresiones metafísicas no pueden poseer significado designativo, y habrá que equipararlas, en este aspecto, a las expresiones lógico-matemáticas, que también son necesariamente verdaderas. Tal argumento, sin embargo, no convence, pues aunque las afirmaciones metafísicas fueran necesariamente verdaderas, no lo serían a causa de las reglas que regulan las combinaciones de signos, como acontece en la lógica y en la aritmética. Por más que la metafísica se distinga de la ciencia, no se ve cómo pueda negarse que sus expresiones significan de modo designativo, y más bien, si acaso, cabría distinguir entre ambas por el propósito sistematizador de la metafísica, que debería ocupar, en consecuencia, la misma casilla de la tabla que ocupa la cosmología.

¿Qué decir del discurso filosófico? La razón de que no aparezca, como tal, recogido en la tabla de Morris, es porque, según él (cap. 8, secc. 6), está formado por los cuatro tipos de discurso cuyo uso es sistematizador. El discurso filosófico lo integran, por tanto, fundamentalmente, la cosmología, la metafísica, la crítica y la propaganda. Por lo que respecta a estos dos últimos ingredientes, no está de más alguna aclaración. El discurso crítico es parte de la filosofía en cuanto se aplica a los valores de todo tipo, y aparece así en la ética y en la estética. La propaganda lo es en cuanto que aspira a propagar métodos de investigación y pensamiento. Estos cuatro ingredientes de la filosofía coinciden en una característica: como formas de discurso tienen los cuatro un uso cuyo propósito es sistematizar y organizar la conducta en un todo completo. La filosofía es, para Morris, una síntesis simbólica, totalizadora y crítica (*op. cit.*, págs. 258-259).

Es de justicia señalar que Morris mismo no se sentía satisfecho con su clasificación, reconociendo que la complejidad de estos fenómenos desborda el intento de encajarlos en su esquema semiótico (*op. cit.*, cap. 6, secc. 12, p. 205). Desde luego que las rápidas críticas que, casi a vuela pluma, le hemos hecho, no hacen justicia al detalle de sus razones y de sus argumentos en favor de las ilustraciones de su tabla. No obstante, nada convencer ni éstos ni aquéllas, y a pesar de las numerosas páginas que Morris les dedica, sus consideraciones son parciales, caprichosas e incluso peregrinas. La primera conclusión es que los ejemplos suministrados no pueden caracterizarse suficientemente sólo en términos de los cuatro modos y de los cuatro usos. Otras características, más pragmáticas que semánticas sin duda, deben tenerse en cuenta si se quiere distinguir con rigor entre las peculiaridades que el empleo del lenguaje tiene respectivamente en la política, en las leyes, en la poesía, en la religión, etc. Sin olvidar que de algunos de estos ámbitos del lenguaje pueden formar parte expresiones con diferente modo de significar y con diverso propósito (piénsese, por ejemplo, en la narración literaria, en la religión o en los textos legales).

Lo anterior ha servido, en todo caso, para dejar patente la ambigüedad del término «discurso». Se habla del discurso científico, del discurso metafísico, del discurso religioso... Y se habla también del discurso declarativo,

del discurso prescriptivo, del discurso valorativo... Y acontece, además, que dentro de un discurso en la primera acepción pueden encontrarse expresiones que pertenecen a diversos discursos en la segunda acepción, y también que expresiones pertenecientes a un tipo de discurso en la segunda acepción aparecen en diferentes tipos de discurso en la primera acepción. Por ejemplo: es del todo obvio que dentro del discurso religioso se encuentran tanto expresiones declarativas acerca de la realidad supuestamente sobrenatural, como expresiones valorativas de los hechos y acciones junto a expresiones prescriptivas de formas de conducta; e inversamente, es claro que las expresiones valorativas son parte indispensable tanto del discurso religioso como del discurso estético y del político.

La distinción entre tipos de discurso en la primera de estas dos acepciones es, en mi entender, cuestión pragmática fundamentalmente. Son las diferentes actividades humanas y las distintas formas de cultura las que han originado esos variados tipos de discurso y sus características peculiares. En cambio, el criterio más claro, más seguro y más riguroso para distinguir entre tipos de discurso en el segundo sentido sólo puede ser de carácter semántico, a saber: la forma de relacionarse la expresión lingüística con la realidad. Atendiendo a este criterio, propongo distinguir los siguientes tipos de discurso.

1. *Discurso declarativo*.—Es el constituido por aquellas oraciones que, debido a lo que significan, pueden corresponder o no a lo que ocurre, a cómo son las cosas, y pueden, en consecuencia, ser *verdaderas* o *falsas*. Estas oraciones son típicas de las ciencias, así como de aquellos procesos de comunicación que responden a un propósito informativo. Son igualmente típicas de aquellos usos del lenguaje con los que se pretende representar situaciones posibles, como acontece en la ficción literaria.

2. *Discurso prescriptivo*.—Está formado por aquellas oraciones que por lo que significan, pueden ser *cumplidas* o *incumplidas*, según que se lleve o no a cabo lo que la oración dice. Incluye tanto las oraciones imperativas en general como las llamadas normativas o deónticas, esto es, aquellas que califican una forma de conducta como obligatoria o debida. Téngase en cuenta que estamos hablando aquí, como en los demás casos, de oraciones usadas, y no meramente mencionadas. Cuando yo informo a alguien de cuáles son sus obligaciones según determinado código de normas, no estoy propiamente usando las expresiones normativas correspondientes sino mencionándolas en el contexto lingüístico de mis afirmaciones que son, como es patente, declarativas. Son también prescriptivas ciertas oraciones por medio de las que se realizan usualmente actos compromisorios, puesto que comprometerse no es sino obligarse uno a sí mismo. Cuando expresa compromiso una oración como «Lo haré» es prescriptiva. Las interrogaciones las considero, asimismo, oraciones prescriptivas puesto que constituyen peticiones de respuesta.

3. *Discurso expresivo*.—Consiste en oraciones que expresan, a causa de su significado, el estado psicológico del hablante o sus actitudes, y que por ello pueden ser *sinceras* o *insinceras*. Incluyo aquí tanto las exclamaciones de todo tipo como las expresiones de deseo, dolor, amor, creencia, y demás estados y sensaciones. No hay que olvidar que muchas de estas oraciones se pueden interpretar, asimismo, como declarativas de hechos internos al sujeto; así, «Me siento angustiado» se puede tomar como una expresión del estado del sujeto o como una declaración acerca del mismo (recuérdese la discusión de Wittgenstein en torno a este tipo de oraciones, que vimos en la sección 7.3). También incluyo aquí, aunque parezca extraño, las oraciones valorativas. La razón es que su significado consiste en expresar actitudes favorables o desfavorables. El hecho de que estas actitudes se justifiquen en virtud de determinadas propiedades del objeto valorado ha conducido en muchas ocasiones a tomar los valores como propiedades de los objetos, convirtiendo las oraciones valorativas en oraciones declarativas y haciendo de la valoración una cuestión de verdad o falsedad. Pero los valores no son propiedades sino relaciones. El valor es la relación que algo (persona, hecho o cosa) tiene, en virtud de sus propiedades, con determinadas exigencias o necesidades del sujeto. Cuando resultan satisfechas, expresamos nuestra actitud diciendo del objeto que es bueno, bello, útil, provechoso, etc., según de qué tipo sean esas necesidades o deseos. En caso contrario, decimos que es malo, feo, inútil, etc. Puesto que usualmente esas exigencias o necesidades son compartidas por la mayoría de los miembros de la comunidad las valoraciones tienden a hacerse sociales, adquiriendo apariencia de objetividad y ocultando que simplemente son manifestaciones de nuestra actitud preferencial. Así, en una situación de sólidas vigencias sociales de carácter moral, decir de alguien que es bueno (en sentido moral) tendrá como consecuencia que el oyente inmediatamente le atribuya una forma de comportarse muy específica y bien determinada; pero esto no significa que el término «bueno» describa esa forma de conducta; más bien, el significado de esta palabra consiste en expresar que esa manera de comportamiento satisface las exigencias morales del que habla, las cuales, en este ejemplo, coinciden con las de los demás hablantes en general.

Ciertamente, otros tipos de discurso pueden dar lugar a actos de habla sinceros o insinceros. Uno es sincero al pronunciar una oración declarativa si cree lo que afirma, y es insincero en caso contrario; es sincero cuando pronuncia un imperativo si desea que su prescripción se cumpla, y no lo es si no lo desea; es sincero al prometer si tiene intención de hacer lo que promete y no lo es si carece de tal intención. Esto ocurre porque el uso de estas oraciones presupone cierta actitud en el hablante: creencia, deseo o intención. Pero el significado de estas oraciones *no consiste* en expresar dicha actitud. Por eso aquí es el acto de habla el que es sincero o insincero. Por el contrario, las oraciones expresivas tienen como significado expresar la actitud del hablante, y por ello la sinceridad o insinceridad no es solamente una característica del acto verbal (que también lo es), sino además

de la propia oración; la sinceridad es, precisamente, la manera de relacionarse la oración con la realidad, es su valor semántico. No debe olvidarse que lo que estamos clasificando son tipos de oraciones y no de actos.

4. *Discurso realizativo*.—Lo constituyen aquellas oraciones que, en virtud de lo que significan, enuncian el propio acto de habla que por medio de ellas se realiza. Incluyen, en primer lugar, lo que anteriormente hemos llamado «oraciones realizativas explícitas», esto es, oraciones de la forma «Te mando que...», «Informo de que...», «Prometo que...», etc. Y en segundo lugar, las oraciones que crean situaciones. Son de este tipo las fórmulas que, en el contexto de las instituciones jurídicas, se emplean para producir aquellos efectos que caracterizan una nueva situación. Por ejemplo, la fórmula con la que el Rey de España sanciona las leyes: «Yo vengo en sancionar la siguiente Ley...». Así como las siguientes expresiones ya citadas anteriormente: «Por la presente nombro a... (alguien para un cargo)», «Fallamos que debemos condenar y condenamos...», etc. De ambos tipos de expresiones podemos decir que son *válidas* o *inválidas*. Estos términos son especialmente adecuados para el segundo tipo de los mencionados, esto es, para las fórmulas jurídicas, pero pueden aplicarse, faltando otro término mejor, a las expresiones del primer tipo. Una oración de la forma «Prometo que...» es *válida* cuando con ella se realiza el acto de prometer, esto es, cuando se cumplen todas las condiciones para que efectivamente sea una promesa. Así, la oración sería *inválida* si, por ejemplo, el hablante promete algo que no está en su capacidad hacer. Una oración de la forma «Te mando que...» será *válida* cuando constituya efectivamente un mandato, y no lo será en caso contrario, por ejemplo, porque el hablante carezca de la autoridad necesaria para mandar al oyente. Es patente que la oración es *válida* o *inválida* según lo sea el acto realizado con ella. En el primer caso, es *válida* la oración si es *válida* la promesa; en el segundo, lo es si es *válido* el mandato. Esto ocurre porque, significando estas oraciones el propio acto realizado con ellas, la relación entre ellas y la realidad consiste justamente en que el acto sea efectivamente realizado. Si el acto es *inválido*, la oración que lo expresa también lo es. Austin, como se recordará, evaluaba las oraciones realizativas como *felices* o *infelices*; mis términos conservan la intención de Austin, pero me parecen más claros y tienen en su favor la legitimidad de su uso jurídico, si bien he de reconocer que resultan forzados al aplicarlos a ciertos actos verbales como los ruegos o los consejos.

Las principales diferencias entre esta clasificación y la de Searle derivan de que éste clasifica actos de habla (y más específicamente, actos ilocucionarios), mientras que ahora hemos clasificado tipos de oraciones. Desde este punto de vista tenemos que distinguir entre una oración como «¡Súmate a la huelga!», que expresa un comportamiento a realizar, y puede ser cumplida o incumplida, y una oración como «Te mando que te sumes a la huelga», que expresa el propio acto que realiza y puede ser *válida* o *inválida*. Naturalmente, la primera de ambas oraciones puede usarse tam-

bién para dar un mandato, el cual será a su vez válido o inválido. Es claro que el imperativo «¡Súmate a la huelga!» está contenido en la oración realizativa «Te mando que te sumes a la huelga», pero ello no es diferente de lo que ocurre en el discurso declarativo indirecto. Cuando digo «Está lloviendo» expreso un hecho y la oración puede ser verdadera o falsa; cuando digo «Informo de que está lloviendo» expreso mi propio acto verbal y no tiene sentido decir de esta oración que sea verdadera o falsa, aunque esto pueda decirse de la oración declarativa subordinada que contiene.

Por lo demás, y siempre que nos limitemos a oraciones que no sean explícitamente realizativas, existe una clara correspondencia entre algunos de los tipos de discurso recogidos y algunos de los tipos de actos ilocucionarios distinguidos por Searle. Así, es patente la correspondencia entre los actos representativos y el discurso declarativo, entre los actos directivos y el discurso prescriptivo, entre los actos expresivos y el discurso del mismo nombre. Como resumen y recordatorio, recojo en la tabla que sigue la clasificación de las funciones del lenguaje de Jakobson, la de los actos ilocucionarios de Searle y la de los tipos de discurso que he propuesto.

<i>Funciones del lenguaje</i>	<i>Actos ilocucionarios</i>	<i>Tipos de discurso</i>
Referencial	Representativos	Declarativo
Conativa	Directivos	Prescriptivo
Emotiva	Expresivos	Expresivo
Metalingüística	Compromisorios	Realizativo
Poética	Declaraciones	
Fática		

La clasificación de los actos ilocucionarios clasifica los actos que realizamos al usar la palabra; la clasificación de los tipos de discurso clasifica los tipos de oraciones por los que realizamos esos actos. La clasificación de las funciones del lenguaje, a su vez, está más cercana de la primera, pero toma en consideración criterios más variados y resulta, por ello, más confusa (es obvio que las tres segundas funciones no pertenecen al mismo nivel taxonómico que las tres primeras). Los tipos de discurso, tal y como ahí los hemos entendido, corresponden a los modos más generales de significar gramaticalmente, y pertenecen, por consiguiente, al nivel del análisis locucionario.

7.9 Una teoría pragmática del significado

Hemos visto en las secciones anteriores diversas maneras de desarrollar la teoría del lenguaje sobre la base de atender al uso que hacemos de él.

De aquellas, la elaboración más completa es la del concepto de acto ilocucionario y de sus clases, que inicia Austin y perfecciona Searle, y que suministra una interpretación parcial, pero más precisa y rigurosa, para las vagas declaraciones de Wittgenstein acerca de los usos lingüísticos. Aunque en tales desarrollos no están del todo ausentes las consideraciones semánticas, y ni aun las sintácticas, es patente que dan el mayor peso a la perspectiva pragmática, de acuerdo con el giro que toma la filosofía del lenguaje en el segundo Wittgenstein. Pues bien, una teoría propiamente pragmática del significado es la elaborada por Grice durante los años en que se desarrolla la teoría de los actos de habla, aunque de modo independiente. Si bien su primera presentación data de mediados los años cincuenta, tan sólo en los últimos tiempos ha empezado a tener influencia y ha sido objeto de discusión frecuente.

El programa de Grice consiste en explicar, recurriendo a las intenciones del hablante, lo que éste quiere decir en una ocasión determinada por medio de sus palabras, para explicar, sobre esta base, tanto el significado permanente de las expresiones como el significado aplicado que éstas poseen al ser usadas. De esta manera, la intención comunicativa del hablante queda colocada como piedra angular en la teoría del significado.

La definición general y primaria de Grice es la siguiente («Meaning», 1957, p. 46): afirmar de un hablante que quiso decir algo por medio de sus palabras equivale a decir que el hablante pretendió que la pronunciación de sus palabras produjera cierto efecto en su auditorio precisamente por medio del reconocimiento de su intención. Formulado de modo más técnico («Utterer's Meaning and Intentions», 1969, p. 151): «Un hablante, *H*, quiso decir algo al proferir una expresión, *x*», es verdad si y sólo si, para un auditorio, *A*, *H* profirió *x* con la intención de que:

- (1) *A* produjera una respuesta particular *r*
- (2) *A* reconociera dicha intención de *H*, y
- (3) *A* produjera su respuesta *r* sobre la base de dicho reconocimiento.

De esta formulación, todavía simple y elemental, hay varias cosas que decir. La primera, que Grice no la limita a la comunicación lingüística, sino que la extiende a todos los casos de lo que llama «significado no natural», concepto que corresponde muy aproximadamente al de los signos culturales o convencionales, de los que hablamos en la sección 2.3, aunque Grice niega que la correspondencia sea total. Ciertamente los ejemplos que da de significado natural son casos de esos signos que hemos llamado «vestigios» o «síntomas»; así es significado natural el que tienen ciertas manchas cutáneas cuando decimos de ellas que significan sarampión, o también el que se da cuando decimos que el presupuesto nacional de este año significa que tendremos un año duro. Como se ve, en el primer caso el efecto significa la causa, mientras que en el segundo, la causa significa el efecto. Como ejemplos de significado no natural, y aparte de las expresiones lingüísticas, Grice menciona el caso en el que tres timbrazos en el

autobús significan que va lleno. Las razones de Grice para negar que su distinción entre significado natural y no natural corresponda con exactitud a la distinción entre signos naturales y convencionales son oscuras, e incluyen consideraciones muy peculiares, tales como que las palabras no son signos o que, en el ejemplo citado, el presupuesto no es signo de lo que significa, esto es, que no es signo de que tendremos un año duro («Meaning», p. 41). Puesto que nosotros hemos aceptado por definición, en el capítulo 2, que todo cuanto tiene significado es signo, todos los ejemplos de Grice han de ser, para nosotros, ejemplos de signos, y ya que Grice no justifica este punto con más detalle, podemos desentendernos de otras formas de significado no natural y limitarnos al propio de las expresiones lingüísticas.

En segundo lugar, hay que decir que al hablar de que el auditorio produzca su respuesta *sobre la base* del reconocimiento de la intención del hablante, lo que Grice quiere decir es que tal reconocimiento sea, al menos, parte de su razón para producir su respuesta, y no que constituya meramente la causa de la misma.

Tenemos, pues, que querer decir algo por medio de una expresión consiste en proferir la expresión con una triple intención: (1) con la intención de que el oyente produzca una respuesta; (2) con la intención de que el oyente se dé cuenta de la intención (1); y (3) con la intención de que el oyente produzca su respuesta por razón (al menos, en parte) de haberse dado cuenta de la intención (1). ¿En qué consiste esa respuesta del oyente que el hablante tiene intención de originar con sus palabras?

Grice ha considerado dos casos básicos, el de las expresiones declarativas y el de las imperativas, sugiriendo que, en el primero, la intención del hablante es inducir una creencia en el oyente, mientras que, en el segundo, consistiría en que el oyente haga algo («Meaning», p. 45). Esto, por lo pronto, está abierto a dos objeciones, la primera importante, la segunda, no tanto. Por un lado, es claro que, en muchas ocasiones, lo que pretende el hablante no es convencer al oyente de que crea algo, sino simplemente comunicarle lo que cree él. Por otro, la diferencia entre inducir una creencia y conseguir que el oyente haga algo es una diferencia cualitativa que obliga a un tratamiento distinto para cada uno de estos dos tipos de efecto. De aquí que, en trabajos posteriores, Grice haya reformulado este punto del siguiente modo. La respuesta o efecto que el hablante pretende conseguir con sus palabras consiste, si se trata de una oración declarativa, en que el oyente piense que el hablante cree algo, y si se trata de una oración imperativa, en que el oyente tenga la intención de hacer algo. De esta forma, la respuesta del oyente será primariamente, en ambos casos, una actitud proposicional: o bien atribuir al hablante una creencia, o bien tener la intención de hacer algo («Utterer's Meaning Sentence-Meaning and Word-Meaning», 1968, p. 59; y en forma algo distinta, considerablemente más complicada, y examinando diversos contraejemplos, «Utterer's Meaning and Intentions», 1969, secc. IV). Tenemos, en definitiva, como respuesta, dos estados psicológicos de carácter intencional. Pero

¿y en el caso de que se trate de oraciones de otro de los tipos que hemos aceptado al final de la sección anterior, por ejemplo, expresivas o realizativas? No creo que resulte difícil extender el análisis de Grice a estos casos. Si se trata de oraciones expresivas de las actitudes o estados del hablante, su intención comunicativa primaria será conseguir que el oyente le atribuya la actitud o estado en cuestión; si es una oración realizativa, su intención será que el oyente reconozca el acto realizado con la proferencia de tal oración.

Hasta ahora hemos hablado de lo que el hablante quiere decir al proferir una expresión en una ocasión determinada, y al reproducir la definición de Grice más arriba he añadido que era todavía simple. En efecto, con ocasión de discutir diferentes objeciones y críticas, Grice ha ido refinando su definición hasta llegar a notables extremos de complejidad («Utterer's Meaning and Intentions», secc. III). Sin entrar en detalles que no son de este lugar, y tan sólo a modo de ilustración, reproduciré una de las definiciones más elaboradas, simplificándola levemente a fin de que se entienda mejor.

Un hablante, *H*, quiso decir algo al proferir una expresión, *x*, si y sólo si, para un auditorio *A*, *H* profirió *x* con la intención de que:

- (1) *A* reconociera que *x* tenía ciertas características *c*
- (2) *A* reconociera la intención (1) de *H*
- (3) *A* reconociera cierta correlación *w* entre las características *c* de la expresión *x* y cierto tipo de respuesta *r*
- (4) *A* reconociera la intención (3) de *H*
- (5) *A* reconociera, sobre la base del cumplimiento de las intenciones (1) y (3) de *H*, la ulterior intención de *H* de conseguir que *A* produjera la respuesta *r*
- (6) *A* produjera la respuesta *r* sobre la base del cumplimiento de la intención (5) de *H*, y
- (7) *A* reconociera la intención (6) de *H*.

¿De dónde procede la complejidad de esta nueva definición frente a la que ya conocíamos? De que aquí se ha introducido un nuevo e importante elemento: las características de la expresión utilizada y la relación que poseen con cierto tipo de respuestas. Dicho de otra forma: la nueva definición reconoce que, para explicar lo que el hablante quiere decir por medio de sus palabras, hay que tener en cuenta ciertas características de estas últimas así como la forma en que tales características están correlacionadas con ciertos tipos de respuesta, y que todo ello debe caer bajo la intención comunicativa del hablante y ser reconocido como tal intención por el oyente.

Esto por lo que se refiere a lo que un hablante quiere decir al usar una expresión en una ocasión determinada. Ahora hay que explicar sobre este fundamento, otros problemas relativos al significado de las expresiones. En primer lugar, el significado permanente o intemporal de las expresiones tipo, esto es, aquello de lo que hablamos cuando decimos «La ex-

presión x significa...». Aparentemente, lo que aquí está en cuestión es lo que significa la expresión como tal, y no lo que el hablante quiera decir por medio de ella. Pero para Grice se trata, claro está, de explicar lo primero en función de esto último. ¿De qué manera? Tal y como se recoge en la siguiente definición:

Para un hablante H , la expresión-tipo x significa algo, si y sólo si H tiene en su repertorio el procedimiento siguiente: proferir un ejemplar de la expresión x cuando tiene la intención de que el oyente produzca determinada respuesta («Utterer's Meaning, Sentence-Meaning and Word-Meaning», p. 61, def. 2; he simplificado la formulación a fin de que se entienda mejor).

Lo que esta definición quiere decir es lo siguiente. La expresión-tipo de «Está lloviendo» significa algo para mí en la medida en que yo cuento con el procedimiento siguiente: pronunciarla siempre que tengo la intención de que mi auditorio produzca la respuesta que consiste en atribuirme la creencia de que está lloviendo. Naturalmente que al proferir o pronunciar una expresión lo que hago es producir un ejemplar de dicha expresión, puesto que la expresión tipo no es más que una abstracción (recuérdese la distinción entre signo tipo y signo acontecimiento en la sección 2.2).

La definición está pensada, en principio, para lo que Grice llama «expresiones tipo no estructuradas» (*unstructured utterance-types*). Grice considera el ejemplo de un gesto. Puesto que antes hemos decidido limitarnos a las expresiones lingüísticas, no consideraremos otros casos, pero hay que hacer constar que la expresión que he dado como ejemplo: «Está lloviendo», requiere en rigor una definición más compleja, pues se trata de una expresión estructurada y completa. Por supuesto que Grice no tiene ninguna dificultad en complicar su definición lo suficiente para que cubra tanto el caso de las expresiones estructuradas y completas (oraciones) como el de aquéllas que poseen estructura pero son incompletas (frases o sintagmas). Unas y otras se caracterizan porque su significado deriva del significado de los elementos léxicos que las componen; por ello, Grice recurre aquí a la noción de procedimiento resultante, que define así: un procedimiento para una expresión tipo x es resultante si está determinado por el conocimiento de procedimientos para las expresiones tipo que son elementos de x , y tales que esos procedimientos lo son también para cualquier secuencia de expresiones tipo que ejemplifique una ordenación sintáctica particular (*op. cit.*, pp. 63-64). La definición de significado de una expresión tipo sin estructura puede, ahora, extenderse a las expresiones tipo estructuradas, sean oraciones o frases, sin más que especificar que el procedimiento de proferir la expresión para conseguir la oportuna respuesta del oyente sea un procedimiento resultante.

Lo anterior se refiere exclusivamente a lo que significan las expresiones tipo, con estructura o sin estructura, para un hablante determinado, o, dicho de modo más técnico, en un idiolecto. Es patente que resta por explicar lo que podríamos considerar más importante: el significado de una expresión tipo en un lenguaje, esto es, para una comunidad de ha-

blantes. Por lo que toca a las expresiones sin estructura, Grice suministra la definición siguiente.

Para una comunidad *C*, la expresión-tipo *x* significa algo si y sólo si algunos (¿muchos?) miembros de *C* tienen en su repertorio el procedimiento de proferir un ejemplar de *x* cuando tienen la intención de que el oyente produzca determinada respuesta, estando condicionada la conservación de tal procedimiento por el supuesto de que, por lo menos, algunos otros miembros de *C* tengan o hayan tenido dicho procedimiento en su repertorio (*op. cit.*, p. 62, def. 3).

La definición es, sin duda, imprecisa. La duda entre si han de ser algunos o muchos quienes coincidan en sus procedimientos comunicativos resulta ilamativa, porque hace pensar que la identificación de una comunidad lingüística sea cuestión de número o de proporción. Da, además, la impresión de que poseer un lenguaje común fuera cosa de ponerse de acuerdo en la utilización de procedimientos comunicativos, lo que, si bien puede admitirse para lenguajes artificiales, no tiene sentido en el caso de las lenguas naturales. Con su definición, Grice pretende —según declara— captar la idea de conformidad, así como la de uso correcto o incorrecto de las expresiones, pero su definición me parece quedarse todavía muy lejos de una mínima claridad. El punto central es la noción de procedimiento, pero su intento de definir en qué consiste tener un procedimiento en el repertorio propio acaba también en el fracaso, según él mismo reconoce. Como aproximación avanza la idea de que ello equivale a estar dispuesto a hacer algo, pero advierte en seguida la insuficiencia de esta explicación, pues uno puede no estar de ningún modo dispuesto a usar una expresión aun cuando dicha expresión signifique algo para él. En vista de lo cual sugiere Grice que lo que necesitamos es más bien la idea de estar equipado para usar una expresión, pero renuncia a precisar más esta idea (*op. cit.*, p. 62). Uno no puede por menos de pensar que, formulado el problema en estos términos, se ha llegado, por otro camino, al debatido problema de la adquisición del lenguaje (que discutimos en el capítulo 5), y que el problema ha pasado al ámbito de la psicolingüística.

Hemos visto, pues, lo que hace relación al significado permanente o intemporal de las expresiones tipo tanto para un hablante individual como para un grupo. Resta por examinar lo que Grice llama el significado intemporal *aplicado* de esas expresiones. Es a lo que nos referimos cuando decimos «La expresión *x* significa en este caso...». La definición de Grice, un tanto complicada, se deriva en realidad fácilmente de las anteriores. Es como sigue:

Cuando el hablante *H* profirió la expresión *x*, ésta significaba algo si y sólo si *H* tenía la intención de que su auditorio reconociera lo que *H* quería decir por medio de *x*, sobre la base del conocimiento (o presunción) que el auditorio tiene de que al repertorio de *H* pertenece el procedimiento de proferir *x* cuando *H* tiene la intención de que su auditorio produzca determinada respuesta (*op. cit.*, p. 63; he refundido en mi formulación las definiciones 4 y 4' de Grice).

De esta manera, el significado permanente aplicado se sitúa como eslabón intermedio entre el significado permanente y el significado ocasional de las expresiones. Ni que decir tiene que, como Grice reconoce (p. 57), el significado aplicado y el significado ocasional pueden coincidir. La definición de significado aplicado que acabamos de ver está pensada para expresiones sin estructura. Su extensión a expresiones estructuradas, completas e incompletas, requiere simplemente, como en el caso del significado intemporal o permanente, la introducción de la idea de procedimiento resultante, de la manera que ya conocemos.

Por complejo que lo anterior pueda parecer, deja fuera multitud de detalles y aspectos técnicos de la teoría de Grice, así como de aplicaciones y desarrollos marginales. De éstos, y aun sin entrar en una exposición detallada, no quiero dejar de mencionar el esbozo ofrecido por Grice para un tratamiento del significado de las palabras, al menos para dos categorías sintácticas fundamentales, el sintagma nominal y el sintagma adjetival. La teoría de Grice consiste en asignar a cada uno de estos dos tipos de frases o palabras un tipo de correlación; la correlación que llama referencial, al sintagma nominal, siendo lo correlacionado con la frase un objeto particular; y la correlación que llama denotacional, al sintagma adjetival, siendo aquí lo correlacionado cualquiera de los miembros de una clase de objetos. Así, en la oración «El perro de Rodríguez es peludo», el sintagma nominal está correlacionado con un objeto particular que es, a saber: cierto perro, y el sintagma verbal (en este caso, predicativo o adjetival) está correlacionado con uno de los miembros del conjunto de los objetos peludos. Sobre esta base, y para explicar el significado completo de la oración, bastará recurrir a la noción de procedimiento resultante que ya conocemos. La cuestión es: ¿en qué consiste la correlación mencionada? Lo más destacable es que Grice no encuentra modo de definirla si no es recurriendo a la intención por parte del hablante de establecerla. De aquí que pueda afirmar, en contra de gran parte de la filosofía actual del lenguaje, que la intencionalidad parece estar incrustada en los fundamentos de la teoría del lenguaje.

A pesar de su técnica enunciativa, la teoría de Grice constituye claramente una última consecuencia del nuevo giro que la obra del segundo Wittgenstein imprimió a la teoría del lenguaje. Si lo que importa es nuestro uso del lenguaje, resulta obvio que haya que preguntarse por la intención con la que lo usamos. El problema es: ¿basta esto para explicar incluso el significado intemporal (o gramatical) de las expresiones? Searle, por ejemplo, ha criticado que Grice no complete su teoría con una consideración de las convenciones que regulan el uso del lenguaje. Sin embargo, hemos podido apreciar que, en una definición elaborada, Grice acepta que la expresión ha de estar correlacionada, en virtud de sus características, con el tipo de respuesta que el hablante pretende obtener del oyente, y que el hablante ha de tener la intención de que el oyente reconozca dicha correlación. En este aspecto, la crítica de Searle (*Actos de habla*, 2.6) no es del todo válida, aunque hay que decir en descargo de este último que,

debido a la fecha de su libro, Searle tal vez sólo pudo tener en cuenta el primero de los artículos de Grice, pero no los trabajos posteriores.

Pero incluso sobre la definición refinada podemos todavía preguntarnos: ¿en qué consiste la correlación entre ciertas características de las expresiones lingüísticas y ciertos tipos de respuesta? O dicho con un ejemplo: ¿por qué cuando digo «Está lloviendo» el oyente me atribuye la creencia de que está lloviendo? ¿Cuáles son las características de esa oración que están correlacionadas con dicha actitud del oyente? En un fragmento de las *Investigaciones filosóficas* que Searle recuerda (*loc. cit.*, p. 45), Wittgenstein escribe: «Haz el siguiente experimento: di 'Hace frío aquí', queriendo decir 'Hace calor aquí'. ¿Puedes hacerlo?» (secc. 510). En mi opinión, la respuesta es que podemos hacerlo si cambiamos el significado de la palabra «frío»; pero para ello no basta la intención del hablante. Hace falta algo más. ¿Qué? Esto es lo que la doctrina de Grice deja sumido en la oscuridad: lo que se requiere para que las palabras tengan tal o cual significado en un lenguaje determinado. Ya vimos que su definición de significado permanente para una comunidad es particularmente insatisfactoria.

Searle ha acusado también a Grice de que confunde el aspecto ilocucionario y el perlocucionario por definir el significado en términos de los efectos que el hablante persigue (*loc. cit.*). Tal crítica, sin embargo, no parece justificada, pues no se advierte en ninguna afirmación de Grice semejante confusión. Además, en rigor, Grice no define el significado en términos de efectos, sino en términos de la intención (del hablante) de producir esos efectos. Esto significa que, si el efecto pretendido no se produce, si, por ejemplo, mi auditorio no me atribuye la creencia en que está lloviendo, no por eso mi afirmación «Está lloviendo» carecerá de significado, con tal que haya sido proferida con la serie de intenciones que Grice exige.

Lo que sí resulta llamativo es esta compleja serie de intenciones entrelazadas entre sí que la doctrina de Grice invoca. Teniendo en cuenta que la adición de sucesivas intenciones obedece, en principio, al propósito de dar cuenta de sucesivos contraejemplos, se ha sugerido que podría darse un regreso infinito que permitiera siempre una nueva clase de intención cuando recurrir a ella convenga para salvar la teoría (Platts, *Ways of Meaning*, III.3, p. 87). Aunque Grice se ha mostrado sensible a esta crítica («Utterer's Meaning and Intentions», pp. 156 y ss.), la adición de nuevas intenciones, más allá de las que aparecen recogidas en la definición elaborada que conocemos, solamente estaría exigida por ejemplos tan complejos y rebuscados que no resultan verosímiles, y menos aún interesantes. En tal medida, parece justificado desentenderse de la posibilidad de un regreso infinito.

Las definiciones de Grice, como toda teoría del lenguaje que se centre en el uso, atienden primordialmente a lo que el hablante quiere decir por medio de sus palabras, y derivadamente —según hemos visto— y en función de ello, al significado de las oraciones como tales. Pero vimos en la sección 3.4 que las oraciones de una lengua son potencialmente infinitas,

y esto implica que la mayoría de ellas nunca serán usadas. Se ha dicho, por esta razón, que tal teoría no sirve para dar cuenta del significado de las oraciones en los lenguajes naturales (Platts, *Ways of Meaning*, 111.3, p. 89). La objeción, sin embargo, creo que podría ser contestada desde el punto de vista de las definiciones más elaboradas. Hemos visto, en definitiva, que Grice distingue entre las expresiones no estructuradas y las expresiones con estructura, y que al hilo de estas últimas llega a ofrecer una definición para el significado de expresiones simples, tomando como ejemplos los casos de sintagmas nominales y adjetivales. Teniendo en cuenta que el conjunto de los lexemas de una lengua es siempre finito, no parece que Grice pueda, en principio, tener dificultad para explicar, sobre la base anterior, el significado de cualquier oración, usada o no, de una lengua. Las nociones de procedimiento en el repertorio de un hablante, y de procedimiento resultante, están pensadas justamente para resolver este problema. Téngase en cuenta que Grice no sólo ha explicado lo que el hablante quiere decir, sino también lo que, intemporalmente, significan sus palabras. Lo que caracteriza su posición es que ha explicado esto último en función de lo primero, y no, como era sólo en la teoría lingüística y en buena parte de la filosofía del lenguaje, lo primero en función de lo último. Si la explicación de Grice funciona, funciona también para las oraciones no usadas, y la crítica nunca podrá prosperar por esta vía.

Pero aquí, como ya sabemos, la noción central es la noción de procedimiento. Hasta donde llega mi conocimiento, no hay en los escritos de Grice precisión mayor que la, ya mencionada, de que tener un procedimiento es hallarse equipado para usar la expresión de que se trate. El lector habrá notado aquí una curiosa y sospechosa semejanza con el concepto chomskiano de competencia. El hablante usa el lenguaje como lo usa porque está equipado para usarlo así, porque en eso consiste su competencia de hablante. Resulta difícil evitar la impresión de que la teoría de Grice conduce directamente a un mentalismo tan oscuro como el de Chomsky. Y no podía ser menos habiendo puesto como piedra angular de la teoría del significado el concepto de intención, concepto oscuro si los hay. Porque, en última instancia, ¿cómo reconoce el oyente las complejas y entrelazadas intenciones comunicativas del hablante? Grice nunca nos ha ilustrado al respecto. En realidad, la respuesta más obvia parece ésta: por lo que significan sus palabras. Más arriba nos hemos preguntado: ¿por qué cuando digo «Está lloviendo» mi auditorio me atribuye la creencia de que está lloviendo? No se ve qué otra respuesta pueda haber si no es ésta: por lo que significa esa oración, esto es, por lo que significan juntas las palabras «está» y «lloviendo». Y en esto justamente consiste la correlación existente entre las características de la oración y la actitud que por medio de ella el hablante pretende inducir en el oyente. Esas características, sobre las que Grice nos había dejado ignorantes, no pueden ser otra cosa que eso que llamamos el significado gramatical de la oración.

Pero, como hemos estudiado en las secciones anteriores, hay algo que no es el significado gramatical, y que es también parte de la intención

comunicativa del hablante: la fuerza ilocucionaria. Es cierto que la fuerza depende, al menos en parte, del significado de la oración: no puedo ordenar a alguien que cierre la puerta diciendo «La puerta está cerrada», ni informarle de que la puerta está cerrada diciéndole «Cierra la puerta». Pero es claro que para la determinación de la fuerza específica de una oración es fundamental, aparte de otras características del contexto extralingüístico, la intención del hablante, pues es ésta la que hace que sus palabras constituyen una advertencia o una simple información, un consejo o un ruego. (No hace falta advertir que la fuerza se confunde con el significado gramatical cuando se trata de una oración explícitamente realizativa.) De aquí que se haya separado, dentro de la teoría del significado, la teoría del sentido y la teoría de la fuerza, para matizar que, mientras que la teoría de Grice no es en absoluto aplicable a la primera, en cambio resulta fecunda aplicada a la segunda (Platts, *op. cit.*, p. 93). Searle mismo, que como hemos comprobado tiene sus objeciones a la doctrina de Grice, ha llegado incluso a reformularla de manera que resulte particularmente apta para una teoría del lenguaje basada en la fuerza ilocucionaria (*Actos de habla*, 2.6, *ad finem*).

Lo que mantengo es, pues, que no basta el uso que hacemos de las palabras, la intención con la que las pronunciamos, para explicar su significado. Más bien, las usamos como las usamos, y podemos poner en su utilización las intenciones que ponemos porque significan lo que significan. Esto no implica negar que el significado de lo que decimos depende en parte de nuestra intención; en un sentido amplio, el término «significado», aplicado a una expresión lingüística en cuanto usada por alguien en una ocasión determinada, incluye tanto el significado gramatical de la expresión como lo que añade el contexto extralingüístico y la intención del hablante. Pero lo primero y primario es el significado gramatical. En la sección siguiente estudiaremos ciertos aspectos en los que el contexto coopera al significado total de la oración, y consideraremos otros requisitos de la comunicación lingüística que son parte importante de una teoría pragmática del lenguaje.

7.10 La implicación pragmática y la implicación contextual

La atención concedida por Grice a las intenciones comunicativas del hablante es asimismo responsable de una categoría introducida por él en el análisis de las relaciones entre oraciones: la categoría de implicatura conversacional. Para entender mejor el hueco que esta noción viene a llenar, conviene que recordemos cuáles son las formas usuales de relacionarse entre sí las proposiciones a efectos de su tratamiento lógico y semántico.

Cuando tenemos dos oraciones p y q , tales que la oración compleja compuesta de ambas en ese orden es falsa solamente en el caso de que p sea verdadera y q falsa, la composición de p y q constituye una función

y esto implica que la mayoría de ellas nunca serán usadas. Se ha dicho, por esta razón, que tal teoría no sirve para dar cuenta del significado de las oraciones en los lenguajes naturales (Platts, *Ways of Meaning*, III.3, p. 89). La objeción, sin embargo, creo que podría ser contestada desde el punto de vista de las definiciones más elaboradas. Hemos visto, en definitiva, que Grice distingue entre las expresiones no estructuradas y las expresiones con estructura, y que al hilo de estas últimas llega a ofrecer una definición para el significado de expresiones simples, tomando como ejemplos los casos de sintagmas nominales y adjetivales. Teniendo en cuenta que el conjunto de los lexemas de una lengua es siempre finito, no parece que Grice pueda, en principio, tener dificultad para explicar, sobre la base anterior, el significado de cualquier oración, usada o no, de una lengua. Las nociones de procedimiento en el repertorio de un hablante, y de procedimiento resultante, están pensadas justamente para resolver este problema. Téngase en cuenta que Grice no sólo ha explicado lo que el hablante quiere decir, sino también lo que, intemporalmente, significan sus palabras. Lo que caracteriza su posición es que ha explicado esto último en función de lo primero, y no, como era sólo en la teoría lingüística y en buena parte de la filosofía del lenguaje, lo primero en función de lo último. Si la explicación de Grice funciona, funciona también para las oraciones no usadas, y la crítica nunca podrá prosperar por esta vía.

Pero aquí, como ya sabemos, la noción central es la noción de procedimiento. Hasta donde llega mi conocimiento, no hay en los escritos de Grice precisión mayor que la, ya mencionada, de que tener un procedimiento es hallarse equipado para usar la expresión de que se trate. El lector habrá notado aquí una curiosa y sospechosa semejanza con el concepto chomskiano de competencia. El hablante usa el lenguaje como lo usa porque está equipado para usarlo así, porque en eso consiste su competencia de hablante. Resulta difícil evitar la impresión de que la teoría de Grice conduce directamente a un mentalismo tan oscuro como el de Chomsky. Y no podía ser menos habiendo puesto como piedra angular de la teoría del significado el concepto de intención, concepto oscuro si los hay. Porque, en última instancia, ¿cómo reconoce el oyente las complejas y entrelazadas intenciones comunicativas del hablante? Grice nunca nos ha ilustrado al respecto. En realidad, la respuesta más obvia parece ésta: por lo que significan sus palabras. Más arriba nos hemos preguntado: ¿por qué cuando digo «Está lloviendo» mi auditorio me atribuye la creencia de que está lloviendo? No se ve qué otra respuesta pueda haber si no es ésta: por lo que significa esa oración, esto es, por lo que significan juntas las palabras «está» y «lloviendo». Y en esto justamente consiste la correlación existente entre las características de la oración y la actitud que por medio de ella el hablante pretende inducir en el oyente. Esas características, sobre las que Grice nos había dejado ignorantes, no pueden ser otra cosa que eso que llamamos el significado gramatical de la oración.

Pero, como hemos estudiado en las secciones anteriores, hay algo que no es el significado gramatical, y que es también parte de la intención

comunicativa del hablante: la fuerza ilocucionaria. Es cierto que la fuerza depende, al menos en parte, del significado de la oración: no puedo ordenar a alguien que cierre la puerta diciendo «La puerta está cerrada», ni informarle de que la puerta está cerrada diciéndole «Cierra la puerta». Pero es claro que para la determinación de la fuerza específica de una oración es fundamental, aparte de otras características del contexto extralingüístico, la intención del hablante, pues es ésta la que hace que sus palabras constituyan una advertencia o una simple información, un consejo o un ruego. (No hace falta advertir que la fuerza se confunde con el significado gramatical cuando se trata de una oración explícitamente realizativa.) De aquí que se haya separado, dentro de la teoría del significado, la teoría del sentido y la teoría de la fuerza, para matizar que, mientras que la teoría de Grice no es en absoluto aplicable a la primera, en cambio resulta fecunda aplicada a la segunda (Platts, *op. cit.*, p. 93). Searle mismo, que como hemos comprobado tiene sus objeciones a la doctrina de Grice, ha llegado incluso a reformularla de manera que resulte particularmente apta para una teoría del lenguaje basada en la fuerza ilocucionaria (*Actos de habla*, 2.6, *ad finem*).

Lo que mantengo es, pues, que no basta el uso que hacemos de las palabras, la intención con la que las pronunciamos, para explicar su significado. Más bien, las usamos como las usamos, y podemos poner en su utilización las intenciones que ponemos porque significan lo que significan. Esto no implica negar que el significado de lo que decimos depende en parte de nuestra intención; en un sentido amplio, el término «significado», aplicado a una expresión lingüística en cuanto usada por alguien en una ocasión determinada, incluye tanto el significado gramatical de la expresión como lo que añade el contexto extralingüístico y la intención del hablante. Pero lo primero y primario es el significado gramatical. En la sección siguiente estudiaremos ciertos aspectos en los que el contexto coopera al significado total de la oración, y consideraremos otros requisitos de la comunicación lingüística que son parte importante de una teoría pragmática del lenguaje.

7.10 La implicación pragmática y la implicación contextual

La atención concedida por Grice a las intenciones comunicativas del hablante es asimismo responsable de una categoría introducida por él en el análisis de las relaciones entre oraciones: la categoría de implicatura conversacional. Para entender mejor el hueco que esta noción viene a llenar, conviene que recordemos cuáles son las formas usuales de relacionarse entre sí las proposiciones a efectos de su tratamiento lógico y semántico.

Cuando tenemos dos oraciones p y q , tales que la oración compleja compuesta de ambas en ese orden es falsa solamente en el caso de que p sea verdadera y q falsa, la composición de p y q constituye una función

veritativa diádica que la lógica de proposiciones reconoce bajo el nombre de «condicional veritativo-funcional», y que se lee usualmente: *si p entonces q*. Es frecuente, desde Russell y Whitehead, llamar a este condicional «implicación material», y leerlo así: *p implica materialmente q*. Quine ha objetado, con toda razón, que esto no es exacto, pues mientras que el condicional veritativo-funcional es una forma de unir oraciones para formar una oración más compleja, la implicación es una relación entre oraciones, y como tal sólo puede afirmarse utilizando nombres de esas oraciones, por ejemplo así: «*p*» *implica materialmente* «*q*» (*Lógica matemática*, secc. 5). Puesto que en estas definiciones que estamos recordando no hay peligro, para nosotros, de confusión ni error en este punto, podemos pasar por alto la recomendación de Quine, que, desde luego, es aplicable tanto en el presente caso como en los que vamos a repasar sucesivamente. En conclusión: podemos decir que *p* implica materialmente *q* siempre que la composición veritativo-funcional de *p* y *q* (en este orden) es un condicional verdadero.

Cuando dos oraciones *p* y *q* son tales que su condicional es lógicamente verdadero, decimos que *p implica lógicamente q*. Un condicional es lógicamente verdadero cuando no deja de ser verdadero para cualquier sustitución uniforme de sus términos no lógicos. Así, es lógicamente verdadero el condicional:

- (1) Si todos los filósofos son neuróticos, entonces todos los filósofos analíticos son neuróticos

pues el condicional sigue siendo verdadero de cualquier forma en que se sustituyan sus términos no lógicos, siempre que la sustitución sea uniforme; por ejemplo, así:

- (2) Si todas las mujeres son adorables, entonces todas las mujeres rubias son adorables

Con ello se da a entender que, en este ejemplo, son términos lógicos «todos» y «si... entonces». De conformidad con lo anterior, podremos afirmar que la oración «Todos los filósofos son neuróticos» *implica lógicamente* la oración «Todos los filósofos analíticos son neuróticos»; y de modo análogo para el ejemplo (2).

Cuando dos oraciones *p* y *q* son tales que la afirmación «*p* y no *q*» es contradictoria, se dice, en inglés, que *p entails q*. Esta expresión se ha traducido a veces como «*p* entraña *q*», y por lo mismo, el sustantivo *entailment* se ha traducido como «entrañamiento». Yo, por mi parte, ante la dificultad de encontrar un término exacto e idiomático al mismo tiempo, le he dado la vuelta a la relación y he traducido en otras ocasiones «*p entails q*» como «*q* se deduce de *p*» o «*q* es deducible de *p*». Así definida, es obvio que esta relación va estrechamente unida a la de implicación lógica, aunque hay una larga discusión sobre si coinciden exactamente o no.

Provisionalmente, y para nuestros efectos, podemos aceptar que siempre que p implica lógicamente q , q se deduce de p , pero no viceversa; esto es, que hay casos en los que q se deduce de p , pero no está lógicamente implicada por ella. Esto ocurre cuando la deducibilidad se base en el significado de términos no lógicos que aparecen en p y q . Así, decir de alguien que es un adulto y negar que ha pasado la pubertad es contradictorio, en virtud de lo que significan los términos (no lógicos) «adulto» y «pubertad». En este sentido, de la afirmación de que alguien es adulto se deduce que ha pasado la pubertad, aunque lo primero no implica lógicamente lo segundo. Podríamos tal vez decir que lo implica analíticamente, y reservar en consecuencia la expresión «implicación analítica» para estos casos, aunque sin olvidar que, en rigor, toda implicación lógica es una subespecie de la implicación analítica. Así como al concepto de implicación lógica corresponde el concepto de verdad lógica, al concepto de implicación analítica corresponde el de verdad analítica. En el capítulo próximo estudiaremos una importante discusión de este punto por parte de Quine.

Cuando dos oraciones p y q son tales que si q es falsa, p no es verdadera ni falsa, se dice que p *presupone* q . Como se recordará, esta relación fue introducida por Strawson como medio de dar cuenta de las oraciones cuyo sujeto carece de referencia, sin tener que recurrir a una explicación tan artificiosa como la teoría de las descripciones definidas. Así, la oración «El asesino de Aristóteles era de Siracusa» presupone la oración «Aristóteles murió asesinado», ya que la verdad de esta última oración es condición necesaria tanto para la verdad como para la falsedad de la primera, y siendo esa última oración falsa (como creemos que es), la primera no puede decirse que sea ni lo uno ni lo otro.

Con esto hemos revisado la implicación material, la implicación lógica, la implicación analítica y la presuposición. A estas relaciones se añade la implicatura conversacional de Grice, que tiene, por cierto, semejanzas con lo que, por los mismos años, otros venían llamando implicación pragmática o implicación contextual.

Así, Nowell-Smith, en su *Ética* (1954, cap. 6, secc. 2) habla de la implicación contextual definiéndola del siguiente modo: un enunciado (*statement*) p *implica contextualmente* otro enunciado q si cualquiera que conozca las convenciones normales del lenguaje está autorizado a inferir q a partir de p en el contexto en que ambos aparecen. Y ofrece las tres reglas siguientes sobre este concepto: primera, hacer un enunciado implica que el que lo hace cree lo que dice; segunda, todo hablante implica contextualmente que tiene lo que, en su opinión, son buenas razones en favor de su enunciado; y tercera, todo cuanto uno dice puede asumirse que es relevante para los intereses de su auditorio. Estas tres reglas pueden, y suelen, quebrantarse en el habla cotidiana, pero parece que deben ser convenciones que, en principio, rijan el uso normal del lenguaje.

La imprecisión de la formulación anterior es algo que salta a la vista. En primer lugar, no está claro si las convenciones normales a las que se hace referencia son convenciones de un lenguaje determinado o convencio-

nes propias de todo lenguaje por el hecho de serlo. A juzgar por las tres reglas, habría que concluir que se trata de las convenciones del lenguaje en general. Siendo esto así, es del todo llamativo que se llame «contextual» a una implicación que no depende para nada del contexto. Pues, en efecto, las reglas citadas no toman en cuenta el contexto, el cual sólo sería relevante para excluir la aplicación de alguna de ellas. Por ejemplo, si se ve claro por el contexto que estoy hablando en broma, nadie tendría por qué asumir que creo lo que digo o que tengo buenas razones para decirlo. Finalmente, no está claro qué es lo que produce la implicación, pues si bien se atribuye ésta, en la definición, al enunciado, esto es, a la afirmación que se hace al usar la oración, la regla primera más bien hace que la implicación derive del acto de hacer el enunciado, mientras que la regla segunda claramente atribuye la implicación al propio hablante.

Todos estos puntos oscuros han sido elucidados por Nowell-Smith en un trabajo posterior («Contextual Implication and Ethical Theory», 1962). Aquí, la implicación contextual ya no se llama así, sino implicación pragmática, que evidentemente es más exacto; las reglas de las que esta implicación se deriva son reglas necesarias para que todo lenguaje cumpla sus propósitos, y el sujeto de la implicación es el propio hablante, pues es éste quien, al hacer un enunciado, implica pragmáticamente que cree lo que dice, que posee sus buenas razones para creerlo así y que juzga de interés para su auditorio el decirlo.

Frente a la implicación pragmática, así definida, la implicación contextual resulta ser una subespecie de aquélla. La implicación contextual es también propiamente pragmática, puesto que es una forma de implicación que involucra al hablante; se trata de lo que el hablante puede dar a entender por medio de sus palabras a causa de haberlas proferido en un determinado contexto extralingüístico. Un excelente ejemplo es el suministrado por Grice («The Causal Theory of Perception», 1961): en un tribunal de exámenes uno de los miembros, comentando el ejercicio escrito de un alumno, se limita a decir con laconismo: «Tiene buena letra y no comete faltas de ortografía.» En ese contexto, el hablante da a entender, sin duda, que el contenido del examen es tan pobre que no vale la pena hablar de ello. Lo característico del caso presente es que la implicación no se produce a no ser que se utilicen esas palabras en tal contexto o en uno semejante. La oración, por sí sola y en razón de lo que significa, carece de esa implicación, y la adquiere tan sólo al ser usada en un tipo de contexto muy determinado. De aquí que la implicación pueda ser cancelada —en término de Grice— si el hablante añade algún comentario en este sentido, por ejemplo: «Pero no quiero decir con ello que el ejercicio sea malo.»

Grice da otro ejemplo más, de carácter muy distinto, como caso también de este modo de implicación. Se trata del caso en el que alguien afirma «Mi mujer está en la alcoba o en el comedor», implicando por ello que ignora en cuál de ambas habitaciones se encuentra. Es patente que esta forma de implicación puede igualmente ser cancelada añadiendo las pa-

labras oportunas, pero, en contra de lo que sugiere Grice, es muy diferente a la anterior, lo bastante para que pueda ponerse en duda la asimilación de ambas. En el caso precedente, en el que irónicamente se descalificaba un examen escrito, la implicación se debía sin duda al contexto particular en el que se usaban las palabras. Era una forma de implicación contextual. En el último caso, sin embargo, el contexto es irrelevante para justificar que se dé la implicación. En cualquier situación normal en que alguien haga la afirmación citada, se producirá la implicación y el oyente espontáneamente atribuirá al hablante ignorancia respecto a cuál es la habitación donde se halla la persona referida. En mi opinión, éste es un caso más bien de implicación pragmática en el sentido que hemos visto. Son más bien las convenciones generales del lenguaje las que justifican esta implicación, pues es una convención general del uso del lenguaje que el hablante diga lo que cree y, por consiguiente, que solamente utilice una expresión de duda desconocimiento cuando éste es su estado mental.

Grice compara con los casos anteriores un tercer ejemplo digno de mención, el de quien afirma «Es pobre pero honrada», implicando, según parece, que existe algún contraste entre la pobreza y la honradez, acaso porque resulte especialmente meritorio ser honrado cuando se es pobre. Los comentarios de Grice en torno a este ejemplo me parecen particularmente confusos y no vale la pena discutirlos en detalle, pero mencionaré que parece inclinado a aceptar que esta implicación sea cancelable; esto es, defiende que sería inteligible afirmar: «Es pobre pero honrada, aunque no quiero decir que haya nada llamativo en que sea ambas cosas a la vez.» En mi opinión, una afirmación como esta última encierra un uso erróneo del término «pero», y solamente es inteligible en la medida en que pensemos que el hablante, por descuido o por ignorancia, ha empleado ese término para decir lo que podría haber expresado simplemente con la conjunción «y». La única razón por la que he citado este tercer caso de los que Grice comenta es porque pienso que se trata de otro ejemplo de lo que anteriormente hemos considerado implicación analítica, ya que la implicación viene dada exclusivamente por lo que significan nuestras palabras, en este ejemplo y, en particular, la palabra «pero».

Los tres ejemplos de Grice que hemos considerado pueden, pues, categorizarse, respectivamente, como casos de lo que hemos venido llamando implicación contextual, implicación pragmática e implicación analítica. Posteriormente, Grice ha desarrollado y sistematizado su tratamiento de este tema en la siguiente forma («Logic and Conversation», 1975). Distingue ahora entre lo que llama, respectivamente, implicatura (*implicature*) convencional e implicatura conversacional. La primera es una relación que obedece exclusivamente a los rasgos convencionales de las expresiones usadas. El último de los tres ejemplos anteriores, el que involucraba la conjunción «pero», sería un caso de implicatura convencional. La implicatura conversacional, a su vez, puede ser de dos tipos, particularizada y generalizada. La primera se debe a los rasgos propios del contexto extralingüístico particular en el que se usan las palabras. El primero de los tres ejem-

plos comentados, en el que se condenaba implícitamente un examen encomiando su buena letra y ortografía, constituye un caso de este tipo. Finalmente, la implicatura conversacional generalizada se produce por el uso normal de ciertas expresiones, y a menos que se den determinadas circunstancias que la excluyan. Tal vez sea ejemplo de ésta el caso en el que la afirmación de que alguien está en tal sitio o en tal otro implica que el hablante no conoce en cuál de ambos sitios está. Grice da ahora (*op. cit.*, p. 56) este ejemplo: la afirmación acerca de alguien, de que va a ver a una mujer esta tarde, implica conversacionalmente de forma generalizada que esa mujer no es su esposa, ni su madre, ni ningún pariente cercano. Según esto, parece que las categorías de Grice corresponden a los tres conceptos que ya habíamos visto. Su implicatura convencional es lo que habíamos considerado implicación analítica, y tiene, por consiguiente, carácter semántico (aunque Grice, desde su teoría pragmática del significado, que ya conocemos, podría defender el carácter igualmente pragmático de ese tipo de implicación). La implicatura conversacional generalizada equivaldrá a la implicación pragmática general, y la particularizada corresponderá a una implicación pragmática relativa al contexto extralingüístico, o sea, lo que hemos llamado implicación contextual.

Según Grice, las implicaturas conversacionales están esencialmente conectadas con ciertos rasgos generales del discurso que derivan de que la comunicación lingüística es una forma de conducta cooperativa, que sirve a un propósito común a los hablantes (*op. cit.*, p. 45). Esto es lo que recoge el *principio de cooperación*, como lo llama Grice, y que formula así: haz que tu contribución a la conversación sea la requerida, en el momento en que tiene lugar, por el propósito o dirección aceptados del intercambio lingüístico en el que tomas parte (*loc. cit.*). Este principio se particulariza en cuatro tipos de máximas, que son los siguientes:

1. Por lo que se refiere a la *cantidad* de información a suministrar: haz que tu contribución sea tan informativa como se requiera para los propósitos normales de la comunicación, pero no más de lo que se requiera.

2. Por lo que hace a la *calidad* de la comunicación: intenta que tu contribución sea verdadera. Que se descompone en no decir lo que uno cree que es falso y en no decir aquello para lo que se carece de pruebas adecuadas.

3. Por lo que toca a la *relación* entre las palabras y el tema de la comunicación: sé relevante.

4. Y, finalmente, por lo que respecta al modo de la *comunicación*: sé claro. Lo cual se descompone en evitar expresiones oscuras, evitar la ambigüedad, ser breve y ser ordenado.

¿Cuál es la condición de estas máximas, así como del principio general de cooperación, del cual derivan? Para Grice (*op. cit.*, pp. 47 y ss.), se

trata, por lo pronto, del supuesto, que todos aceptamos en principio, de que los hablantes proceden, en general y mientras no se demuestre lo contrario, de acuerdo con tales máximas. Y la base de esta suposición es, por una parte, el simple hecho de que los hablantes se suelen comportar así, y tienen el hábito de hacerlo. Pero no es sólo eso. Es que, además, es razonable seguir dichos principios si la conducta lingüística ha de ser racional y teleológica, y sin ellos no parece que pueda concebirse el uso del lenguaje en general. Estaríamos en presencia, por tanto, de una especie de principios *a priori* de la racionalidad lingüística comunicativa. Ello no deja de ser llamativo si se repara en el aire tan híbrido que tienen dichas máximas, pues mientras que algunas, en efecto, suenan a mandamientos de un decálogo («no dirás lo que es falso»), otras más bien parecen máximas de calendario («sé breve y ordenado»). La importancia relativa de unas y otras es, desde luego, muy distinta. Mientras que el principio de decir lo que uno cree cierto y relevante es sin duda un principio sin cuyo cumplimiento general no parece que se pueda explicar el uso normal del lenguaje, en cambio el principio de ser breve, claro y ordenado no se ve que pueda constituir sino una aspiración de justificación muy relativa. Supongamos que una comunidad fuera tan consciente de la estética del habla, y de la retórica, que prefiriera, en general, usar expresiones largas, y aun oscuras y desordenadas, con tal de que fueran bellas, novedosas y elegantemente construidas. No encuentro razones para considerar su actitud más irracional que la contraria. Esto no es más que una hipótesis que me parece verosímil, pero, ya en el terreno de los hechos, hay que añadir que parece haber una comunidad lingüística que no cumple con algunos de los principios de Grice. Si es cierto lo que cuenta E. O. Keenan («On the Universality of Conversational Implicatures», cit. en Gazdar, *Pragmatics*, p. 54), los hablantes del malgache, esto es, la generalidad de los habitantes de la antigua Madagascar, hoy República Malgache, procuran que sus expresiones sean lo menos informativas posible, empleando disyunciones en lugar de aserciones simples y categóricas, recurriendo a formas verbales sin sujeto y prefiriendo pronombres indefinidos o nombres comunes a nombres propios y descripciones definidas. ¿Es irracional la conducta de estos hablantes? ¿O más bien deberíamos pensar que la racionalidad lingüística definida por los principios de Grice es relativa a una cultura y a una forma de vida?

De conformidad con lo anterior, Grice caracteriza la noción de implicatura conversacional así: al decir que *p*, un hablante implica conversacionalmente que *q*, con tal que, primero, haya que presumir que está observando las máximas conversacionales mencionadas, o, al menos, el principio de cooperación; segundo, que para que su proferencia de *p* sea consistente con la presunción anterior, se requiera el supuesto de que el hablante se da cuenta de que *q*; y tercero, que el hablante piense que está al alcance de la competencia del oyente darse cuenta de que se requiere el supuesto que se acaba de mencionar (*loc. cit.*, pp. 49 y ss.).

Las máximas de Grice, así como sus tipos de implicatura, dan origen a gran variedad de cuestiones y han sido objeto de discusión detallada a propósito de ejemplos concretos (v. Gazdar, *Pragmatics*, cap. 3, y Kasher, *Conversational Maxims and Rationality*). Por lo que a nosotros respecta, baste señalar la conexión que esas máximas tienen con lo que antes hemos llamado implicación pragmática. Esta forma de implicación deriva, como vimos, de aquellas reglas que podríamos considerar necesarias para que todo lenguaje cumpla su propósito al ser usado. Y esto es lo que expresan los principios de Grice, luego parece que, básicamente al menos, la noción es la misma. Como se habrá apreciado, las reglas de implicación pragmática enunciadas por Nowell-Smith, a saber: que el hablante crea en la verdad de lo que dice, tenga buenas razones para creerlo y lo juzgue de interés para su auditorio, se encuentran recogidas en los principios 2 y 3 de Grice, que son precisamente los que resultan menos discutibles. ¿Por qué hablar de implicación pragmática mejor que de implicatura conversacional? Por un lado, el término «conversacional» me parece indebidamente restrictivo, y parecería limitar el alcance de este tipo de relación entre oraciones al caso de ciertas formas de comunicación oral, siendo así que estas implicaciones, y los principios en que se apoyan, se aplican igualmente a todo uso del lenguaje, sea o no una conversación y sea o no oral. En cuanto al término «implicatura», la elección de este neologismo puede justificarse en base a las diferencias que separan esta relación de la implicación semántica y, en especial, de su caso paradigmático, la implicación lógica. Pero teniendo en cuenta que también usamos el término «implicación» para un caso tan alejado de la implicación semántica como el condicional veritativo-funcional (implicación material), no veo grave inconveniente en usarlo para esta relación pragmática y sí, en cambio, la ventaja de mantener la uniformidad en la terminología.

Teníamos, pues, definidas con cierta claridad y exactitud, suficientes al menos para nuestros propósitos, la implicación material o condicional veritativo-funcional, y tres tipos de implicación semántica, la presuposición, la implicación analítica y la implicación lógica. Hemos de considerar a estos tres tipos de implicación como semánticos, puesto que la implicación deriva de lo que significan los términos y expresiones de la oración implicante, bien sean los términos y expresiones de carácter no lógico (en la presuposición y en la implicación analítica), bien sean los que tienen carácter lógico (en la implicación lógica). Lo que está en juego en estas relaciones es la verdad o falsedad de las oraciones relacionadas. Así, si p presupone q , la falsedad de q hace que p no sea ni verdadera ni falsa; si, en cambio, p implica analítica o lógicamente q , la falsedad de q hace a p falsa.

A estas relaciones tenemos ahora que añadir dos relaciones, no semánticas, sino pragmáticas: la implicación pragmática general y la implicación contextual (a las que pueden reducirse, por las razones que hemos visto, las implicaturas conversacionales de Grice). Se trata de relaciones en las que, como era de esperar siendo pragmáticas, no es la verdad o falsedad de las oraciones lo que está en juego. Si la oración «Está lloviendo ahora

ahí fuera», en cuanto pronunciada por mí en una ocasión determinada, implica pragmáticamente la oración «Creo que está lloviendo ahora ahí fuera», entonces la falsedad de esta última oración, en el caso, por ejemplo, de que yo creyera en ese momento que hace sol, no tiene consecuencia alguna para la verdad o la falsedad de la primera; en definitiva, que esté o no lloviendo no depende de lo que yo crea. La razón última es que no es propiamente una oración-tipo *p* la que implica pragmáticamente otra oración *q*, sino que es el uso o proferencia de *p* que hace un hablante el que implica pragmáticamente que el hablante se encuentra en tal o cual estado psicológico o actitud mental. Teniendo en cuenta esta última precisión, podemos intentar una definición aproximada en estos términos: *p* implica pragmáticamente *q* si y sólo si la afirmación «*p* y no *q*» es contraria a los propósitos generales del uso del lenguaje. Es el caso de quien afirmara «Está lloviendo, aunque creo que no lo está», «Márchate, pero no deseo que te marches», etc.

Los propósitos generales del uso del lenguaje que justifican la implicación pragmática pueden ser aquellos en los que coinciden todos los lenguajes por el hecho de serlo, podríamos decir: los propios del lenguaje en cuanto tal, o bien los propios de una lengua en particular e incluso de algún subgrupo lingüístico. Quiero insinuar con ello que, por muy peculiares características que pueda tener el uso de una lengua, como en el supuesto del malgache, tales características darán lugar, por su parte, a peculiares implicaciones pragmáticas.

La implicación contextual tiene igualmente carácter pragmático, pues resulta producida por el uso que se hace del lenguaje y no meramente por lo que significan los términos y expresiones. Por lo mismo, no afecta tampoco al valor de verdad de las oraciones relacionadas. Como en el caso anterior, no es propiamente una oración-tipo *p* la que implica contextualmente *q*, sino el uso de *p* en una ocasión determinada el que implica contextualmente *q*. ¿Cómo definir esta relación? Aun cuando se trate de una relación pragmática, la definición de la implicación pragmática general no es estrictamente aplicable. Lo que está aquí en juego no son los propósitos generales del lenguaje que justifican el que se digan las cosas de cierto modo o se emplee la lengua de determinada manera. Lo que está aquí en juego es la posibilidad de dar a entender, por medio de unas palabras, algo totalmente distinto de lo que significan esas palabras, y de hacerlo en la medida en que el contexto extralingüístico completa lo que ellas significan por sí solas. Se trata, en definitiva, de lo que llamamos el doble sentido, como se aprecia bien en el ejemplo de Grice acerca del examinador que descalifica irónicamente un examen haciendo un elogio de su ortografía. Podríamos intentar una definición del modo siguiente: *p* implica contextualmente *q* en un cierto tipo de contexto *c* si y sólo si es usual atribuir al hablante la intención de dar a entender *q* cuando profiere *p* en ese tipo de contexto *c*. Esta definición, como la precedente para la implicación pragmática general, adolece de enorme vaguedad, especialmente en cuanto a términos como «usual», pero recoge en todo caso el sentido de nuestras

conclusiones, por provisionales que éstas puedan ser. No está de más recordar que la definición rigurosa de implicatura conversacional es muy problemática (como puede verse en Gazdar, *Pragmatics*, pp. 41-43).

Con esto podemos poner fin a nuestra breve exploración de la teoría pragmática, un tratamiento de los problemas del significado que, a través de Austin, se relaciona, en última instancia, con el giro impuesto por el segundo Wittgenstein a la filosofía del lenguaje. Hemos estudiado los problemas principales planteados por este último, algunas aportaciones de sus herederos más distinguidos, el desarrollo de la teoría de los actos de habla en Austin y Searle y, finalmente, la teoría pragmática de la comunicación lingüística en Grice. Junto a descubrimientos e iluminaciones de indudable valor, hemos tropezado también con dificultades y confusiones. Que puedan éstas resolverse dentro de este enfoque, o que se requiera completarlo, enriquecerlo o sustituirlo, es algo sobre lo que aún no podemos decidir. Debemos, antes, considerar otra línea de investigación que discurre paralela y coetánea con la que hemos explorado en este capítulo y que, como ella, procede también originariamente del atomismo lógico, y en particular del *Tractatus* de Wittgenstein.

Lecturas

Hay que empezar por hacer una remisión a las obras generales de filosofía analítica recomendadas para el capítulo anterior, pues en todas ellas se trata algunos de los temas considerados en el presente capítulo. La remisión ha de incluir, naturalmente, los libros de conjunto sobre Wittgenstein que allí se recomendaron, pues en todos ellos se estudia su segunda etapa filosófica, e incluso, por lo que respecta al libro de Kenny, también sus etapas intermedias. Con esto podemos pasar a la bibliografía que se refiere, en particular, al segundo Wittgenstein.

La obra fundamental, más elaborada y definitiva, del segundo Wittgenstein, las *Philosophische Untersuchungen* (*Investigaciones filosóficas*) ha sido publicada con traducción al castellano por la editorial Crítica en Barcelona en 1988. Hay, desde luego, una traducción inglesa de toda confianza que acompañaba a la edición original del texto alemán (Blackwell, Oxford, 1953), y que ha sido publicada después separadamente (*Philosophical Investigations*, Blackwell, Oxford, 1963).

Los *Cuadernos azul y marrón* (Tecnos, Madrid, 1968), aunque no tienen el interés de las *Investigaciones*, presentan con bastante aproximación todos los temas característicos de estas últimas.

Entre las obras dedicadas a la segunda filosofía de Wittgenstein, hay que empezar por citar *Lenguaje, filosofía y conocimiento*, de José Luis Blasco (Ariel, Barcelona, 1973), que hace muy accesible la introducción a la teoría del lenguaje del segundo Wittgenstein, la cual examina en conexión tanto con sus precedentes lógico-atomistas como con sus continuadores en la filosofía del lenguaje ordinario. Carácter más monográfico tiene

la excelente investigación de Alfonso García Suárez, *La lógica de la experiencia: Wittgenstein y el problema del lenguaje privado* (Tecnos, Madrid, 1976), que, aunque contiene un interesante capítulo sobre el solipismo en el *Tractatus*, está dedicado al problema que enuncia en su título, del que ofrece un tratamiento riguroso y exhaustivo. Hablando de este tema, recordaré que el libro de Hacker, *Insight and Illusion*, tiene un interesante tratamiento de estos problemas en el segundo Wittgenstein. Y puesto que la bibliografía castellana específica sobre el segundo Wittgenstein es tan escasa, me permitiré añadir un par de títulos no traducidos. De un lado, una obra clásica de conjunto, *The Foundations of Wittgenstein's Late Philosophy*, de Specht (Manchester University Press, 1969; el original es alemán y de 1963, pero asumo que el lector tendrá más facilidad para leerlo en inglés). De otra parte, una excelente y completa recopilación de artículos sobre las *Investigaciones filosóficas*, la realizada por Pitcher bajo el título *Wittgenstein: The Philosophical Investigations* (Doubleday, Nueva York, 1966).

Sobre la filosofía del lenguaje ordinario, aparte de lo que pueda encontrarse en los libros* y recopilaciones generales ya mencionados, y en particular en la recopilación de Muguerza, *La concepción analítica de la filosofía*, debe verse la de Chapell, titulada *El lenguaje común* (Tecnos, Madrid, 1971), que contiene ensayos muy representativos. También hay trabajos de interés en la recopilación de Parkinson, *La teoría del significado* (Fondo de Cultura Económica, México).

La obra fundamental de Austin, que es también la más importante dentro del análisis postwittgensteiniano del lenguaje ordinario, está traducida con el título de *Palabras y acciones* (Paidós, Buenos Aires, 1971; 2.ª edición titulada *Cómo hacer cosas con palabras*). Se han traducido también sus artículos, algunos muy importantes, bajo el título *Ensayos filosóficos* (Alianza, Madrid, 1975). Finalmente, está traducida la obra de Searle que sistematiza, en su versión propia, la teoría de los actos ilocucionarios: *Actos de habla* (Cátedra, Madrid, 1980).

Sobre actos y fuerzas ilocutivas se encontrarán artículos muy representativos en la compilación de Luis Valdés, *Significado y acción* (Ediciones Rubio, Valencia, 1983), además de lo que hay en su otra antología ya mencionada, *La búsqueda del significado* (Tecnos, Madrid, 1991).

Un agudo y polémico análisis de las tesis centrales del segundo Wittgenstein sobre el concepto de lenguaje privado, en relación con seguir una regla, es el que ofrece Saul Kripke en *Wittgenstein: reglas y lenguaje privado* (Universidad Autónoma de México, 1989). La interpretación de Kripke ha originado una singular discusión. Sobre el tema puede verse J. Hierro, «Significado y reglas lingüísticas», en su colección de artículos *Significado y verdad. Ensayos de semántica filosófica* (Alianza, Madrid, 1990).

Capítulo 8

DESDE UN PUNTO DE VISTA LOGICO

Quienes insisten en la claridad y en la lógica fracasan con frecuencia en hacerse entender. (HENRY MILLER, *Sexus*.)

8.1 La teoría verificacionista del significado en Carnap

Como ya se indicó al paso en la sección 7.1, la influencia más fuerte e inmediata del *Tractatus* se dejó sentir en Austria y Alemania, hasta el punto de que fue uno de los elementos que contribuyó a aglutinar a un grupo de filósofos, lógicos y científicos que había comenzado a reunirse a principios de los años veinte, en torno a Schlick, a la sazón catedrático en la Universidad de Viena. Hacia 1929, el grupo se constituyó oficialmente como Círculo de Viena. Como vimos, Carnap refiere que gran parte del *Tractatus* se discutió allí sentencia por sentencia («Intellectual Autobiography», p. 24), y es además sabido que Schlick conoció a Wittgenstein en 1927, y que éste tuvo varias reuniones con aquél, con Carnap y con Waismann (Carnap, *op. cit.*, pp. 25 y ss.). Esto no significa que aceptaran todas las doctrinas del *Tractatus*. Aceptaron unas y rechazaron otras, pero lo que es cierto es que no puede explicarse la filosofía del Círculo de Viena sin el *Tractatus*. Carnap, en particular, menciona al primer Wittgenstein como el pensador que tal vez le influyó en mayor grado, aparte de Frege y Russell (*loc. cit.*). En esta medida, Wittgenstein se halla también en el origen de una línea de investigación en la filosofía del lenguaje que posee características propias claramente contrapuestas a las de la línea que hemos examinado en el capítulo anterior. No hace falta decir que de la filosofía del Círculo de Viena, del neopositivismo o empirismo lógico, no vamos a hacer aquí un examen de conjunto. Simplemente examinaremos los problemas del lenguaje, en especial la teoría del significado, y los temas más estrechamente vinculados con ella. Y lo haremos a propósito del autor que más atención les ha dedicado, Carnap, para pasar

posteriormente a autores más recientes que cabe considerar herederos de esta línea.

Siete años después de la publicación del *Tractatus*, en 1928, Carnap publica su primer libro importante, *La estructura lógica del mundo*. El objeto del libro era efectuar una reconstrucción racional de todos los conceptos del conocimiento sobre la base de conceptos que se refieren a lo inmediatamente dado. La expresión «reconstrucción racional» pretende indicar que se trataba de definir y explicar los conceptos epistemológicos en los términos indicados, y no de describir la génesis psicológica real de los mismos. De esta forma, se intentaba llevar a término el viejo programa positivista de reducir todos los conceptos a algo básicamente dado en la experiencia. Como base de su sistema, Carnap escogió una base fenomenalista: lo inmediatamente dado es una corriente continua de vivencias elementales, que son unidades complejas formadas por pensamientos, percepciones, sentimientos, etc. Carnap reconoce que la elección de una base fenomenalista se debió a la influencia de los positivistas alemanes de fines del XIX, Mach y Avenarius, y, sobre todo, a la influencia de Russell («Intellectual Autobiography», p. 18). Pero ha insistido también en que las razones de la elección eran puramente metodológicas, y en que no implicaban ninguna tesis ontológica respecto a qué entidades componen la realidad (*loc. cit.*). El sistema de la constitución de conceptos, como Carnap llama a su método, utiliza de esa corriente de vivencias una determinada propiedad estructural: el recuerdo de la semejanza entre vivencias distintas. Así, la relación «recordado como semejante a» suministra la base sobre la que constituir, o construir, todos los conceptos epistemológicos y sus relaciones lógicas. Se advertirán aquí dos características decisivas. Por una parte, Carnap apoya todo su sistema de reconstrucción en una propiedad relacional, esto es, en un modo de relacionar entre sí las vivencias elementales. Lo que le interesa no son tanto los elementos del conocimiento cuanto las relaciones básicas entre ellos. Carnap atribuye este interés suyo a la influencia de los *Principia Mathematica*, y en particular al tratamiento que aquí se daba a la lógica de las relaciones (prólogo a la segunda edición de *Der Logische Aufbau der Welt*, e «Intellectual Autobiography», p. 16; debe recordarse que la mayor novedad conceptual de la lógica simbólica frente a la lógica aristotélica fue precisamente la lógica de relaciones). Por otra parte, esa propiedad estructural o relacional se considera tal y como cada cual puede advertirla en su propio recuerdo. Epistemológicamente, el sistema de Carnap es solipsista. Está el sujeto con sus propias vivencias, en última instancia sus recuerdos, y sobre esa base hay que reconstruir todos los demás conceptos, hasta los que se refieren a las mentes y a las vivencias ajenas. Este solipsismo, ciertamente, recuerda al de Russell, aunque sea de otro género; en definitiva, el solipsismo es siempre la amenaza que se cierne sobre el fenomenalismo. Pero el sistema de Carnap es muy distinto del de Russell. Los elementos no son simplemente datos de los sentidos, sino conjuntos más complejos y organizados, vivencias. Ello se debe a que Carnap sufrió aquí la influencia, ausente en Russell, de la

psicología de la *Gestalt*, que explicaba los contenidos de la percepción como totalidades organizadas.

Es patente que lo anterior constituye un tema específico de teoría del conocimiento, lo que nos exime de entrar en más detalles; la referencia a esa primera obra de Carnap nos permitirá, en todo caso, entender mejor el origen de su pensamiento y sus influencias primeras (se encontrará una excelente exposición del fenomenalismo carnapiano en el capítulo III del libro de Ulises Moulines, *La estructura del mundo sensible*). Aun cuando el planteamiento de estos problemas por Carnap parecía muy prometedor, muy poco tiempo después de la publicación de su libro él mismo abandonó ese enfoque. Como resultado de las discusiones habidas con otros miembros del Círculo de Viena, y en particular con Neurath, Carnap acabó por reconocer que lo fundamental es el lenguaje en el que se exprese el conocimiento de lo dado, esto es, que lo fundamental son las proposiciones básicas o, como ellos las llamaban, protocolares (*Protokollsätze*). Ahora bien, si estas proposiciones consistieran en expresiones de vivencias, el lenguaje básico o protocolar resultaría ser individual y no podría explicarse bien cómo se relaciona con el lenguaje científico, en el que quedan recogidos todos nuestros conocimientos. Es cierto que el solipsismo que acompañaba a su construcción fenomenalista era, al menos en la intención de Carnap, un solipsismo meramente metodológico, pero no metafísico; no obstante, las dificultades de dar cuenta del lenguaje científico sobre esta base eran obvias, y la conveniencia de explicar la intersubjetividad del conocimiento sobre una base lingüística que fuera igualmente intersubjetiva era algo que, defendido con ardor por Neurath, Carnap no pudo por menos de reconocer. Así se impuso en el Círculo la idea de que lo fundamental no es lo dado en la experiencia personal, sino lo científicamente confirmable, y se tomaron como proposiciones protocolares aquellas que podían considerarse básicas desde el punto de vista de la ciencia empírica más exacta y más desarrollada, la física. A esta posición se la llamó «físicalismo», y las proposiciones protocolares, así entendidas, consistían, en consecuencia, en descripciones cuantitativas de regiones espaciotemporales. Si añadimos a esto la tesis de que todas las ciencias empíricas, en la medida en que lo son, participan del mismo carácter epistemológico (tesis de la unidad de la ciencia), tendremos una visión mínimamente completa, aunque escueta y resumida, de las tesis más generales del Círculo de Viena que suministran el contexto filosófico en el que Carnap va a desarrollar su pensamiento acerca del lenguaje.

Algunos de los conceptos básicos de la teoría del lenguaje que Carnap está elaborando por esta época, así como la relación que guardan con las tesis epistemológicas que hemos visto, aparecen ya en un famoso artículo de 1932, esto es, publicado cuatro años después del libro que acabamos de comentar. El artículo, loado y denostado por igual, defendía la tesis que claramente expresaba su título: «Superación de la metafísica por medio del análisis lógico del lenguaje.» Carnap intentaba justificar aquí que la metafísica se compone de enunciados carentes de significado, los cuales,

por consiguiente, son en rigor pseudoenunciados, enunciados sólo en apariencia. Tales pseudoenunciados lo son de dos clases: o bien por contener alguna palabra sin significado, o bien porque, aun estando formados por palabras todas ellas significativas, éstas estén compuestas entre sí con infracción de la sintaxis.

Examinemos primero el caso de las palabras. ¿Cuándo es significativa una palabra? Carnap suministra varias formulaciones alternativas de lo que considera la condición suficiente y necesaria para que un término posea significado, a saber (*op. cit.*, secc. 2):

1. Que se conozcan los criterios empíricos que regulan el uso o aplicación del término.
2. Que se haya estipulado de qué proposiciones protocolares se deducen aquellas proposiciones elementales en las que el término aparece.
3. Que estén fijadas las condiciones de verdad de las proposiciones elementales en las que el término aparece.
4. Que se conozca el método para verificar dichas proposiciones elementales.

Las cuatro formulaciones vienen —según Carnap— a parar en lo mismo. Y se habrá observado ya en qué grado involucran criterios epistemológicos. Por lo que respecta a la primera formulación, exige simplemente que sepamos en qué casos y en qué condiciones, relativas a nuestra experiencia, emplearíamos la palabra de que se trate. Carnap sugiere un ejemplo: si inventamos un término y lo aplicamos única y exclusivamente a objetos cuadrangulares, sin que se tome en cuenta ninguna otra característica del objeto excepto su carácter cuadrangular, nuestro nuevo término será sinónimo de «cuadrangular», y tal será su significado, sin que quepa pretender que tiene algún otro significado ulterior no reducible al del término «cuadrangular». El criterio de aplicación es el que suministra el significado del término.

Ahora se entenderá por qué la segunda formulación puede tomarse como alternativa de la primera. Si sabemos en qué casos utilizamos un término, sabremos entonces cuáles son las proposiciones básicas de las que se deducen aquellas proposiciones elementales en las que aparece dicho término, pues las primeras, las proposiciones básicas o protocolares, serán descripciones, en términos de propiedades observables, de aquellos casos o situaciones en las que aplicamos el término en cuestión (cfr. «Proposiciones protocolares», de Neurath). Y viceversa, si tenemos esas proposiciones protocolares, tendremos en ellas descripciones observacionales de las situaciones en las que se aplica el término. Esto nos lleva a la tercera formulación, pues lo anterior equivale a fijar las condiciones de verdad de las proposiciones elementales en las que aparezca el término en cuestión; tales condiciones de verdad serán aquellas descritas por las proposiciones

protocolares relevantes. Finalmente, es claro que quien conozca, a propósito de un término, lo que enuncia cualquiera de las tres formulaciones anteriores, conoce el método para verificar las proposiciones elementales en que aparezca aquél, pues conocer esto último no es sino conocer las condiciones de aplicación del término.

Según Carnap, la mayor parte de los términos metafísicos típicos no cumple con esa condición: «idea», «absoluto», «infinito», «cosa en sí», «esencia», «yo», «principio», «Dios» (en su acepción filosófica)... Así, por ejemplo, en el caso del concepto metafísico de principio —que Carnap considera con más detenimiento—, las condiciones observables que son características en un uso empírico de ese término se desechan como irrelevantes, pero no se establecen nuevas condiciones de aplicación para la palabra «principio» en su uso filosófico. En rigor, no se le ha dado un nuevo significado, puesto que no se le ha asignado un nuevo método de verificación a su contenido (*op. cit.*, secc. 3).

El otro caso en el que un enunciado puede carecer de significado es cuando sus palabras, aunque sean significativas, están compuestas infringiendo la sintaxis, con lo que no constituyen una oración significativa o enunciado. Pero la sintaxis a la que se refiere Carnap es la sintaxis lógica, no la sintaxis gramatical. La razón es que esta última permite ciertos pseudoenunciados, que, sin embargo, la primera identifica y excluye. Carnap pone como ejemplo de tales pseudoenunciados éste: «César es un número primo.» Y aunque enunciados semejantes llamen fácilmente la atención y no haya dudas sobre su falta de significación, muchos pseudoenunciados metafísicos, en cambio, pueden pasar desapercibidos y ser tomados como significativos (*op. cit.*, secc. 4). Como es bien sabido, Carnap cita aquí, a modo de ejemplo, varias afirmaciones características de Heidegger, tomadas de su opúsculo *¿Qué es metafísica?*, tales como: «La nada es anterior al no y a la negación... La angustia revela la nada... La nada misma anonada...» (secc. 5). La sintaxis lógica resulta aquí quebrantada al sustantivar un adverbio de negación y privación, asignándole, por analogía con otros sustantivos, una supuesta referencia, la cual no sólo ignoramos en qué consiste, sino que además contradice el significado literal del adverbio que hemos sustantivado. La cuestión es: ¿qué tiene ello que ver con la sintaxis?

En realidad, como se habrá apreciado, el problema es semántico y epistemológico, como lo era en el caso de los términos carentes de significado que acabamos de ver. La razón por la que la sintaxis gramatical no impide la formación de una oración como «César es un número primo», es porque no son propiamente reglas sintácticas, sino semánticas, las que resultan quebrantadas en esa oración (recuérdese lo que estudiamos en la sección 4.4, apartado c, sobre el ejemplo de Chomsky «Las verdes ideas duermen furiosamente»). Se trata, en suma, de una regla semántica que impedirá aplicar un predicado matemático a una persona humana; esto es, una regla de selección, según el nombre que reciben esta clase de reglas en la teoría semántica chomskiana, como vimos en su momento (*loc. cit.*).

Por lo que respecta a los ejemplos de Heidegger comentados por Carnap, la cuestión es aún, si cabe, más claramente semántica, pues lo que aquí se halla en discusión es cuál pueda ser la referencia de la aparente expresión nominal «la nada». Y esta cuestión va ligada, como en el caso general de los términos no significativos que ya conocemos, a una cuestión patentemente epistemológica: la cuestión de cuáles son los criterios empíricos de aplicación de la expresión «la nada», o lo que tanto vale, la de cuál es el método para verificar las afirmaciones acerca de la nada. Y desde este punto de vista la crítica de Carnap es del todo sólida y consistente. La razón por la que Carnap pone su análisis en la cuenta de la sintaxis lógica es que, por esta época, está convencido de que todo análisis riguroso del significado ha de ser un análisis lógico-sintáctico. Así, escribe unos años después del artículo que estamos comentando: «Una de las tareas principales del análisis lógico de un enunciado determinado es descubrir el método de verificación de dicho enunciado. La cuestión es: ¿qué razones puede haber para afirmar tal enunciado?, o ¿cómo se puede estar seguro de su verdad o falsedad? Los filósofos denominan problema epistemológico a esta cuestión. La epistemología o teoría filosófica del conocimiento no es más que una parte especial del análisis lógico mezclado normalmente con algunas cuestiones psicológicas relativas al proceso de conocer» (*Filosofía y sintaxis lógica*, 1935, I.1).

La idea de Carnap es que el significado de un enunciado está en el método de su verificación («Superación de la metafísica...», secc. 6). A menos que se trate, claro está, de enunciados analíticos, de lo que Wittgenstein había llamado, en el *Tractatus*, tautologías y contradicciones, y que Carnap, siguiendo a aquél, califica de enunciados verdaderos o falsos en virtud solamente de su forma. Tales enunciados no dicen nada de la realidad, y esto ciertamente ya lo había subrayado Wittgenstein en el *Tractatus*, según hemos visto. Por consiguiente, todo enunciado que pretenda decir algo acerca de la realidad únicamente será significativo en la medida en que se posea algún método para comprobar si es verdadero o falso. Puesto que todo método de verificación ha de acabar en oraciones protocolares, y puesto que la exigencia de tal método es característica de los enunciados científicos, la consecuencia es que todo enunciado significativo que no sea analítico es un enunciado científico. De esta manera adquiere contornos más precisos la vaga afirmación de Wittgenstein de que lo único que puede afirmarse son las proposiciones científicas (*Tractatus*, 4.11 y 6.53). Y la crítica a la metafísica recibe una formulación de carácter más epistemológico que la de Wittgenstein, pero en todo acorde con ésta (*Tractatus*, 4.003 y 6.53). El análisis lógico declara carente de significado todo supuesto conocimiento que pretenda ir más allá de la experiencia, según Carnap (*loc. cit.*). Cae bajo esta condena todo pretendido conocimiento de esencias, de valores, de principios metafísicos. ¿Qué queda, entonces, para la filosofía? En una vena del todo próxima a la de Wittgenstein, Carnap dirá: lo que queda para la filosofía no son enunciados, no es una teoría, sino solamente un método, el método del análisis lógico

(*loc. cit.*, p. 77 del original). En su aspecto negativo sirve para denunciar todo tipo de pseudoenunciados. En su aspecto positivo sirve para clarificar los conceptos y colocar los fundamentos de las ciencias, tanto empíricas como formales.

A diferencia de Wittgenstein, sin embargo, Carnap les reconoce una función a los pseudoenunciados metafísicos: expresar una actitud general ante la vida (*op. cit.*, secc. 7). Sólo que tal expresión, que adquiere en la poesía, en la pintura, en la composición musical, una manifestación adecuada, en la metafísica recurre a un medio inadecuado, pues se trata de un medio teórico: enunciados, aparentemente significativos, ligados entre sí por relaciones aparentemente lógicas. El medio de expresión crea, tanto en el metafísico como en su lector, la ilusión de que se está haciendo afirmaciones verdaderas o falsas, consistentes o inconsistentes; pero, en realidad, no se está haciendo afirmaciones: simplemente se está expresando una actitud, como el artista. La argumentación y la polémica están, por ello, tan fuera de lugar en la metafísica como en la poesía o en la música. En palabras de Carnap: «El metafísico es un músico carente de habilidad musical» (*loc. cit.*).

Aunque en su artículo no está explícita, en una nota añadida en 1957 Carnap recuerda que, posteriormente, se introduciría una distinción entre significado cognitivo y significado expresivo o emotivo, la cual es aplicable aquí. Ahora podremos decir que los pseudoenunciados metafísicos carecen de significado cognitivo, no expresan conocimiento de ningún tipo, pero que poseen significado emotivo, como lo poseen las expresiones de sentimientos y deseos. En términos de la clasificación de los tipos de discurso propuesta al final de la sección 7.8, habremos de decir que las afirmaciones metafísicas pertenecen al discurso expresivo. Es claro, sin embargo, que la posición de Carnap peca de extremosa, y que distorsiona la efectiva función histórica de la metafísica (asunto distinto es su posibilidad actual). Que la metafísica pueda expresar una actitud ante la vida no es dudoso, y que en muchos casos lo hace puede comprobarse fácilmente; pero hay casos, también, en los que resultaría muy forzado reducir la obra metafísica a esa función expresiva. De otro lado, una obra científica puede también expresar una actitud ante la vida sin perder por ello su carácter científico. La argumentación no está fuera de lugar en la metafísica; antes bien, es central en ella como en cualquier obra teórica, y es acaso lo único que puede hacerla todavía respetable. El discurso metafísico es un discurso teórico de la máxima generalidad y abstracción, y el metafísico no es un artista inhábil, sino un teórico que no se conforma con las limitaciones de la experiencia. ¿Puede llamarse conocimiento al contenido de una teoría metafísica? Esta es otra cuestión. Creo que la aportación positiva y duradera de Carnap, como de los demás miembros del Círculo de Viena, ha sido el estudio riguroso de las condiciones del conocimiento empírico, y en particular del conocimiento científico, todo lo cual ha creado las bases para una teoría del conocimiento y para una filosofía de la ciencia sólidas y rigurosas. En este sentido, las consideraciones que brevemente hemos

visto han contribuido a formular conceptualmente con claridad las razones epistemológicas de la separación que entre la ciencia y la metafísica progresivamente viene aconteciendo a lo largo de la época moderna. Consideraciones que son, no obstante, y como he señalado, erróneas respecto al papel que la metafísica ha representado en relación con la ciencia en la historia del pensamiento occidental. Sobre este punto, la consideración de Russell, que tomaba las doctrinas metafísicas como hipótesis precientíficas de alcance muy general, es más justa y exacta (*Los problemas de la filosofía*, caps. 14 y 15). Lo que hay que concederle a Carnap, y en esto coincide con Russell, es que no hay propiamente conocimiento hasta que esas hipótesis no adquieren algún método de verificación, o, lo que es lo mismo, alguna relación lógicamente aceptable con enunciados intersubjetivos o de observación, esto es, con proposiciones protocolares.

La desmesura de Carnap se halla en haber tomado, como criterio de significado, lo que constituye tan sólo un criterio para determinar si una oración posee o no contenido cognoscitivo de carácter empírico. Este es un criterio que, si acaso, servirá para delimitar la ciencia de la metafísica; pero un criterio de demarcación es cosa distinta de un criterio de significado, como ha subrayado Popper («La demarcación entre la ciencia y la metafísica»). Hay que distinguir, en todo caso, las características propias de la metafísica, en cuanto disciplina o producto cultural, de los rasgos que pueda tener alguna de sus manifestaciones. Reconocer el carácter teórico no científico de la metafísica no impide aceptar que determinadas doctrinas sean un conjunto de absurdos galimatías, o que otras sean vacías y triviales. Al escoger como ejemplo las afirmaciones de Heidegger sobre la nada, Carnap, aun cuando fuera sin pretenderlo, se facilitó extremadamente su labor crítica. Pues es obvio que tales afirmaciones no tienen, por razones estrictamente gramaticales de carácter semántico, significado literal (otra cosa es que dichas afirmaciones puedan interpretarse como una manera metafórica de expresar el miedo a la muerte y como el intento de construir sobre él una concepción de la realidad).

De otro lado, el acierto de Carnap fue ligar la teoría del significado y la teoría del conocimiento de forma tal que evitaba el solipsismo en el que paraba Russell, salvando en cambio la intersubjetividad del significado y vinculando explícitamente la teoría del significado a la filosofía de la ciencia en la forma que el *Tractatus* meramente apuntaba. Lo que había de aceptable en esta actitud persistiría, como vamos a ver, en autores posteriores.

Aparte de las dificultades que se acaban de mencionar, la propia formulación del principio de verificabilidad encontró en seguida numerosas dificultades de aplicación a los propios enunciados científicos. Para empezar, se subrayó —y esto es razonable— que el principio no requiere, para que una proposición sea significativa, que podamos comprobar actualmente su verdad o su falsedad, sino únicamente que sepamos de qué manera podríamos hacer esa comprobación; lo que el principio exige no es la verificabilidad actual, sino la verificabilidad en principio (Ayer, *Lenguaje, ver-*

dad y lógica, cap. 1; Hempel, «Problemas y cambios en el criterio empirista de significado», secc. 2). En segundo lugar, hubo que apresurarse a rechazar la idea, que algunos habían mantenido en los primeros días del Círculo de Viena, de que el principio exige de toda proposición significativa que sea completamente verificable, esto es, que haya algún conjunto finito de proposiciones básicas, observacionales o protocolares, del que se deduzca aquélla. La razón principal es que, así entendido, el principio excluiría las leyes científicas, pues excluiría toda proposición universal basada en datos parciales (Hempel, *loc. cit.*, Popper, *La lógica de la investigación científica*, secc. 6). Precisamente para evitar este inconveniente, y buscando una formulación más congruente con el método hipotético-deductivo, Popper (*loc. cit.*) sugirió un principio de falsabilidad, que requiere, para que una proposición sea científica, que sea en principio posible refutarla completamente recurriendo a la experiencia; o dicho de otra forma: que la negación de esa proposición se deduzca de un conjunto finito de proposiciones básicas. (Se notará que el principio de falsabilidad simplemente caracteriza a una proposición como científica, pero no se ofrece como condición necesaria para que sea significativa; la razón es que se trata, para Popper, de un criterio de demarcación de la ciencia, y no de un criterio de significado. La cuestión de formular un criterio de significado cognitivo era, para él, un pseudoproblema, como puede verse en su artículo «La demarcación entre la ciencia y la metafísica», donde critica extensamente a Carnap.) Este criterio posee, sin embargo, un inconveniente semejante al anterior, aunque de signo inverso: rechaza como no científicas las proposiciones particulares, esto es, proposiciones como «Algunos seres vivos son monocelulares», ya que la negación de una proposición particular es una proposición universal (en el caso de nuestro ejemplo: «Ningún ser vivo es monocelular»), y acabamos de ver que de datos parciales, como son siempre el contenido de las observaciones, no puede deducirse una proposición universal.

Por esta razón, Ayer prefirió prescindir de requisitos tan exigentes como los anteriores y propuso un principio de verificabilidad débil que se limitaba a exigir, para que una proposición tuviera contenido fáctico, que hubiera observaciones que fueran relevantes para la determinación de su verdad o falsedad. Dicho con más precisión, el principio exigía que pudieran deducirse algunas proposiciones de experiencia a partir de la proposición en cuestión junto con otras premisas, sin que fuera posible deducir tales proposiciones de experiencia exclusivamente a partir de las citadas premisas (*Lenguaje, verdad y lógica*, cap. 1, p. 39 del original). Este principio está evidentemente mal formulado, puesto que permite asignar significado fáctico a cualquier proposición como, por ejemplo, «La angustia revela la nada»; basta, para ello, adoptar como premisa adicional una proposición condicional como «Si la angustia revela la nada, entonces ahora es de día», pues de esta premisa en conjunción con la proposición anterior se deduce, por *modus ponens*, la proposición observacional «Ahora es de día», la cual no se deduce de la premisa por sí sola. Ayer mismo

reconoció la corrección de esta crítica, y en la introducción a la edición revisada de su libro (que es diez años posterior) reformuló su criterio del modo siguiente. Ahora afirma que un enunciado es directamente verificable si es un enunciado de observación o de él, junto con otros enunciados de observación, se deduce al menos un enunciado de observación que no sea deducible exclusivamente de esos otros enunciados. Y dice de un enunciado que es indirectamente verificable si de él, junto con otras premisas, se deduce al menos un enunciado directamente verificable que no sea deducible de esas premisas por sí solas, y con la condición de que estas premisas incluyan exclusivamente enunciados que sean o bien analíticos, o bien directamente verificables, o bien tales que se puedan determinar de manera independiente como indirectamente verificables (p. 13 del original). Aunque todavía se le han puesto peros a esta formulación (Hempel, *loc. cit.*) su progreso respecto a las anteriores es innegable, pues elimina la posibilidad de inferencias como la que se acaba de mencionar. También para Ayer, más próximo a Carnap que a Popper, el principio de verificabilidad es un criterio de significado fáctico o empírico.

Carnap, por su parte, había seguido también una tendencia liberalizadora por lo que al principio de verificabilidad respecta. Ya en 1936, en su trabajo «Testability and Meaning», sustituyó el concepto de verificación por el de confirmación, aceptando en lo sucesivo que para que una proposición posea significado empírico basta que sea, en algún grado, confirmable por medio de observaciones adecuadas. La noción de grado de confirmación fue la base sobre la que Carnap realizó, posteriormente, sus importantes aportaciones a la teoría lógica de la probabilidad. Sobre el fundamento de las nuevas ideas de Carnap, Hempel propuso, a su vez, un nuevo criterio de significado cognoscitivo que llama, también, la atención por su liberalidad («Problemas y cambios en el criterio empirista de significado», secc. 3). De acuerdo con este criterio, una oración tiene significado cognitivo si y sólo si es traducible a un lenguaje empirista. Qué lenguaje pueda considerarse empirista es cuestión ulterior, y dependerá de sus reglas de construcción, así como de su vocabulario. Hempel sugería, como un posible lenguaje de este tipo, aquel cuyas reglas sintácticas fueran las de cualquier sistema lógico contemporáneo (por ejemplo, los *Principia Mathematica*), y cuyo vocabulario pudiera reducirse a las constantes lógicas y a ciertos predicados de observación. Estos últimos constituirán el vocabulario empírico básico de tal lenguaje. Aquí, la mayor dificultad se encuentra en suministrar las reglas que permitan esa traducción de la proposición dudosa al lenguaje empirista. Hempel lo reconoció así, eliminando el requisito de la traducción, y fijando, en su lugar, ciertas condiciones que habría de cumplir el vocabulario de una proposición para poder atribuir a ésta significado cognitivo. Tales condiciones se resumen en que los términos de ese vocabulario sean o bien términos lógicos, o bien términos con significación empírica, siendo de esta última clase los predicados de observación y aquellos términos conectados con ellos por medio de definiciones o por otros medios lógicamente satisfactorios («The Concept of

Cognitive Significance: A Reconsideration» y «Empiricist Criteria of Cognitive Significance: Problems and Changes»). El problema ahora era caracterizar esas conexiones. A estos efectos, una de las primeras conclusiones que se impuso es la de que el problema hay que plantearlo con relación al conjunto de una teoría científica: «Si la significación cognitiva puede atribuirse a algo, es a sistemas teóricos completos formulados en un lenguaje con una estructura bien determinada. La marca decisiva de la significación cognitiva en un sistema de ese tipo parece ser la existencia de una interpretación para el mismo en términos de entidades observables.» (Hempel, *Empiricist Criteria of Cognitive Significance...*, secc. 4). A ello se añadió, como un paso más en este proceso liberalizador, la idea de que el significado cognitivo de una teoría es una cuestión de grado, y que más bien que distinguir entre teorías significativas y no significativas, lo razonable es evaluar las teorías en función de diversas características, tales como su claridad y precisión, su poder explicativo y predictivo, su simplicidad formal, y el grado en el que han sido confirmadas por la experiencia (Hempel, *loc. cit.*). Con estos criterios, las teorías metafísicas, ciertamente, no alcanzan valores muy altos. Queda por demostrar que ése sea el punto de vista adecuado para la evaluación de éstas. En todo caso, el problema ahora era dar cuenta de la relación entre los términos teóricos y los términos de observación en las teorías científicas. La empresa de trazar una distinción más general entre oraciones con significado cognitivo y oraciones carentes del mismo había perdido interés por su propia falta de viabilidad.

Una palabra ahora sobre el propio principio de verificabilidad en cuanto criterio de significado. Voces apresuradas se alzaron al momento con lo que parecía la crítica definitiva a principio tan irrespetuoso con nuestra venerable tradición metafísica: el principio de verificabilidad es significativo y, sin embargo, no es verificable. Pero crítica tan fácil es igualmente fácil de deshacer. Porque el principio de verificabilidad no es una proposición acerca de la realidad, y en esta medida no tiene sentido pretender que sea verificable. ¿Qué tipo de proposición es entonces? Recordemos que, en principio, tan sólo se reconocían dos clases de significado, cognitivo y emotivo, por lo que si no le asignamos este último (lo que sería absurdo), hemos de asignarle significado cognitivo, y esto sólo será posible si lo tomamos como proposición analíticamente verdadera. La dificultad estaba en tomar el principio como verdadero, pues ¿qué reglas lógicas podrían dar cuenta de su pretendida verdad analítica? De aquí que se eligiera presentarlo como una definición del concepto de significado fáctico, añadiendo que esta definición no pretende ser arbitraria sino justificada, y que quien no la acepte quedará obligado a suministrar otra definición más adecuada (Ayer, *Lenguaje, verdad y lógica*, introducción a la edición revisada, p. 16 del original). Hempel, por su parte, lo presenta como una propuesta lingüística que trata de explicar la noción de enunciado asertórico, y que es adecuada en estos dos aspectos: porque analiza el concepto común de significado cognitivo, y porque ofrece al mismo tiempo una «recons-

trucción racional» de éste que elimina ambigüedades y permite un uso teórico general y coherente de dicho concepto. («Problemas y cambios...», secc. 6.) En ambas formulaciones una cosa está clara: según ellas, el principio de verificabilidad no es ni verdadero ni falso, en la medida en que una definición o propuesta terminológica, por adecuada que sea y por justificada que esté, no es ni verdadera ni falsa. Acusar al principio de verificabilidad de no ser verificable equivale, por consiguiente, a no haber entendido en qué consiste. Ahora bien, queda una cuestión pendiente: si el principio no es ni verdadero ni falso, entonces carece de significado cognitivo. Y puesto que sería absurdo sugerir que tenga significado emotivo, ¿qué clase de significado tiene entonces? Esto muestra, una vez más, la insuficiencia de la teoría del significado que se expresa en el principio de verificabilidad.

8.2 El modo material y el modo formal

Hemos visto por qué razones las supuestas realidades estrictamente filosóficas, aquellas que son objeto de las afirmaciones metafísicas, sólo dan lugar a pseudoenunciados y han de ser expulsadas del ámbito de la teoría y del conocimiento. Hay, no obstante, realidades que, aún siendo objeto de las ciencias, son también tema de estudio y consideración por la filosofía: el hombre, el lenguaje, la sociedad, la historia, el tiempo, el espacio, la naturaleza... ¿Qué es lo propio del tratamiento que la filosofía da a estos objetos? ¿Hay un punto de vista específico de la filosofía que le permita a ésta decir sobre esos objetos algo diferente a lo que dicen las respectivas ciencias? La reducción de todo conocimiento al conocimiento científico, que caracteriza a los neopositivistas, es claramente incompatible con una respuesta afirmativa. Carnap hará aquí, sin embargo, una consideración que salva el valor de la filosofía; deja para la filosofía un ámbito más amplio y más positivo que el que le reconocía el *Tractatus*. El tratamiento filosófico de esos temas es un tratamiento lógico; esto quiere decir que la filosofía trata, no del hombre, del lenguaje, de la sociedad, de la historia..., sino de las características lógicas de las proposiciones científicas acerca de esos objetos. La filosofía se convierte así en filosofía de las ciencias, y más concretamente, en lógica de las ciencias: «la lógica de la ciencia sustituye a la intrincada maraña de problemas que se conoce como filosofía» (*Logische Syntax der Sprache*, secc. 72). Tal y como Carnap la entiende ahora, la lógica de la ciencia constituye un estudio meramente sintáctico de las proposiciones científicas, estudio que tiene como su parte general una teoría de la sintaxis lógica del lenguaje. Pero esta sintaxis tiene un contenido más amplio que lo usual: trata no sólo de las reglas que determinan cuándo una proposición está bien formada (reglas de formación), sino también de las reglas que especifican qué proposiciones pueden obtenerse a partir de otras (reglas de transformación).

Como puede apreciarse, la posición de Carnap deriva directamente, también aquí, del *Tractatus*, a saber, de aquellas proposiciones en que Wittgenstein reduce la filosofía a crítica del lenguaje y a delimitación del ámbito de las proposiciones científicas (4.0031, 4.113). Pero Wittgenstein nunca llegó a adoptar una actitud como la de Carnap, pues, como ya sabemos, pensaba que cualquier proposición acerca de otra proposición es un sinsentido (incluyendo, por tanto, las propias proposiciones del *Tractatus*). Esto es lo primero que Carnap va a rechazar de Wittgenstein: que no se pueda formular con sentido la teoría del lenguaje, y más todavía tratándose simplemente de una teoría sintáctica, esto es, de una teoría acerca de la forma de las proposiciones. La sintaxis lógica puede formularse; tanto la sintaxis pura, que es como una construcción, y por tanto meramente analítica, como la sintaxis descriptiva, esto es, aplicada a un lenguaje concreto. Y si esto es así, entonces también la filosofía, en cuanto sintaxis lógica de la ciencia, puede formularse en proposiciones correctas y significativas. Para ello basta, en definitiva, suministrar un conjunto de reglas que nos permitan construir la teoría formal del lenguaje. Para demostrar que ello puede hacerse, Carnap se limitará a mostrar su propia obra sobre la sintaxis lógica del lenguaje. A la luz de ésta indicará, además, un criterio que permite distinguir, mucho más claramente que en el *Tractatus*, entre las proposiciones filosóficas aceptables, esto es, las de la lógica de la ciencia, y las que son rechazables, a saber, las proposiciones metafísicas. La diferencia se halla en que las primeras, pero no las últimas, son convertibles en proposiciones sintácticas, o lo que es lo mismo, son traducibles a lo que llama Carnap el modo formal del discurso (*op. cit.*, secc. 73). Veamos en qué consiste esto.

Según lo que implícitamente hemos considerado, las proposiciones sintácticas se contraponen a las proposiciones de objeto; éstas tratan sobre la realidad, aquéllas, del lenguaje. Entre unas y otras se deslizan, sin embargo, las que Carnap denomina pseudoproposiciones de objeto, proposiciones que hablan aparentemente de objetos, pero que en realidad se refieren al lenguaje, y de modo más específico, a las formas de designación de los objetos. Así, son oraciones de objeto, por ejemplo, las siguientes (*op. cit.*, secc. 74 ss.):

- (1) El 5 es un número primo
- (2) Babilonia fue una gran ciudad

Son, en cambio, oraciones sintácticas, éstas:

- (3) La palabra «cinco» es un término de número, no de cosa
- (4) En la conferencia de ayer aparecía la palabra «Babilonia» (o una expresión sinónima de ésta)

A las cuales corresponderían las siguientes pseudoproposiciones de objeto:

- (5) El 5 es un número, no una cosa
- (6) La conferencia de ayer trató de Babilonia.

Como se ve, (1) y (2) hablan de ciertos objetos para predicar de ellos determinadas propiedades: de un objeto matemático, como es el número 5, se afirma que es primo; de un objeto físico, construido por el hombre, como es una ciudad, se predica grandeza. Lo característico de (5) y (6) es que parecen hablar asimismo de esos objetos, para decir del número 5 que es un número y para afirmar sobre la ciudad de Babilonia que de ella trató la conferencia de ayer. Para Carnap, esto es sólo una apariencia; en rigor, (6) no dice nada sobre Babilonia, ni (5) sobre el número 5: no es una propiedad de Babilonia que se hable de ella en una conferencia, ni una propiedad del 5 ser un número. Lo que (5) y (6) quieren decir realmente está dicho de forma clara y explícita en (3) y (4), respectivamente, a saber, que «cinco» es una palabra que designa un número, y que en la citada conferencia se ha usado la palabra «Babilonia» o alguna expresión sinónima de ésta. Las proposiciones sintácticas, como (3) y (4), pertenecen a lo que llama Carnap el modo formal del discurso; las proposiciones cuasisintácticas, o pseudoproposiciones de objeto, pertenecen al modo material del discurso (*loc. cit.*).

Esta distinción permite dejar al descubierto, según Carnap, el auténtico carácter de los problemas filosóficos. En la medida y grado en que son problemas verdaderos, y por tanto rigurosamente formulables en un lenguaje, se descubren como problemas acerca de la forma de nuestras expresiones. Toda proposición filosófica con sentido, y por consiguiente, no metafísica, es una proposición cuasisintáctica que, como tal, puede y debe convertirse en una proposición sintáctica, o lo que es lo mismo, puede y debe pasarse del modo material al modo formal. No se trata de rechazar el modo material como tal, sino de mostrar que es engañoso y que puede equivocarnos, haciéndonos creer que estamos hablando de objetos cuando nuestro discurso tan sólo posee sentido si se entiende como metalingüístico. El uso del modo material nos induce a tomar las proposiciones filosóficas como absolutas, haciéndonos olvidar que son relativas al lenguaje (*op. cit.*, secc. 78). Carnap ofrece diversos ejemplos, de los que seleccionaré algunos como muestra. De los siguientes pares de proposiciones filosóficas, la primera está en el modo material y la segunda en el modo formal, siendo esta última, por tanto, la formulación más explícita y rigurosa de la primera:

- (7) Los números son clases de clases de cosas
- (8) Las expresiones numéricas son expresiones de clase del segundo nivel
- (9) Los números son objetos primitivos
- (10) Las expresiones numéricas son expresiones del nivel cero
- (11) Las relaciones son datos primitivos
- (12) Los predicados poliádicos son símbolos no definidos

(13) Las relaciones no son primitivas, sino que dependen de las propiedades

(14) Los predicados poliádicos se definen sobre la base de los predicados monádicos

La mayor parte de las disputas filosóficas —piensa Carnap— proceden del uso del modo material y se desarrollan en ese caldo de cultivo. La adecuada traducción de las proposiciones en disputa al modo formal mostraría los caracteres reales de lo que está en discusión y evitaría polémicas estériles. Así, la discusión sobre la naturaleza de los números se convertiría en una disputa sobre cómo construir las expresiones numéricas; la cuestión acerca de si las relaciones son primitivas o derivan de las propiedades aparecería como el problema de si tomar o no los predicados poliádicos como símbolos no definidos; etc. Véanse estos otros ejemplos:

(15) Todo color está en un lugar

(16) Toda expresión de color va siempre acompañada, en una oración, por una designación de lugar

(17) El tiempo es continuo

(18) Las expresiones de números reales se usan como coordenadas temporales

La traducción al modo formal se puede aplicar también, con toda facilidad, al *Tractatus*. Entre los ejemplos ofrecidos por Carnap, están éstos:

(19) El mundo es la totalidad de los hechos, no de las cosas (1.1)

(20) La ciencia es un sistema de oraciones, no de nombres

(21) Si conozco un objeto, conozco también todas sus posibilidades de formar parte de hechos (2.0123)

(22) Si está dado el género de un símbolo, están dadas también todas sus posibilidades de formar parte de oraciones

La idea no es que el modo material sea, como tal, rechazable, sino que, al usarlo, debe tenerse en cuenta que se trata de una forma de transposición, de una manera de hablar no literal, puesto que se formulan como afirmaciones sobre objetos lo que realmente son afirmaciones sobre el lenguaje. Si una oración está en el modo material, entonces puede traducirse al modo formal, y esto prueba que se trata de una oración metalingüística, y por tanto que no constituye una proposición de la ciencia natural, pero que, sin embargo, es una proposición con sentido. Por eso puede afirmar Carnap: «La traducibilidad al modo formal del discurso constituye la piedra de toque para todas las proposiciones filosóficas» (*op. cit.*, secc. 81). Pues, en efecto, tan sólo las proposiciones filosóficas que sean así traducibles son aceptables, porque sólo ellas pueden pertenecer a la lógica de la ciencia, sólo ellas versarán sobre la sintaxis lógica del lenguaje científico; las demás proposiciones filosóficas serán metafísicas, y por consiguiente, pseudo-proposiciones.

La crítica de Carnap a las proposiciones metafísicas, que vimos en la sección anterior, junto con su reducción de las demás proposiciones filosóficas a proposiciones lógico-sintácticas del modo formal, a la manera que acabamos de estudiar, representa uno de los esfuerzos más rigurosos realizados en nuestro siglo para acomodar la filosofía a una concepción científica del mundo y del conocimiento. La primacía de la ciencia en el ámbito del conocimiento y descripción de la realidad, únicamente dejaría libre a la filosofía el ámbito metalingüístico en el que se produce la aplicación de la lógica a la ciencia. Pero ya en la sección precedente vimos que la tarea crítica encuentra más dificultades de las aparentes a primera vista, y que sólo en ciertos casos puede llevarse a cabo con éxito fácilmente. Por lo que respecta a la triple distinción entre lenguaje de objetos, discurso material y discurso formal, toca ahora considerar brevemente sus dificultades.

Que las proposiciones del modo material que Carnap da como ejemplo pueden, en general, ponerse en el modo formal a la manera que hemos visto, no parece discutible. Podría diferirse, si acaso, de la traducción particular que propone para algunos casos particulares. Así, en un ejemplo que antes no he citado, ofrece la siguiente proposición material («Filosofía y sintaxis lógica», II.7):

(23) La conferencia trató de la metafísica

y da para ella la siguiente traducción al modo formal:

(24) La conferencia contenía la palabra «metafísica»

Ahora bien, es patente que un texto o discurso puede tratar de metafísica sin contener esta palabra ni ninguna expresión sinónima con ella. Para versar sobre una disciplina determinada, un texto no requiere emplear un nombre de la misma. Basta conocer el tipo de cuestiones tratado y la manera de tratarlas para determinar que una conferencia versó sobre metafísica, sobre sociología, sobre lógica, o sobre lo que fuera. Ello no significa que no podamos hacer una traducción de (23) al modo formal; podríamos, acaso, afirmar:

(24') La conferencia contenía numerosos términos metafísicos

o bien:

(24'') La conferencia constaba fundamentalmente de pseudoproposiciones metafísicas

etcétera. Otra cosa es si se trata de un tema u objeto particular, como en el caso de Babilonia. Aquí tal vez haya que aceptar, que para tratar de esa ciudad, una conferencia, texto o discurso debe contener algún nombre o designación de la misma.

Cuestión más importante es la siguiente. No son sólo las proposiciones cuasisintácticas del modo material las que pueden pasarse al modo formal, sino también las proposiciones de objeto. No se ve nada incorrecto en traducir:

(2) Babilonia fue una gran ciudad

a:

(25) «Babilonia» es un término que designa una gran ciudad

Ni tampoco en traducir:

(1) El 5 es un número primo

a:

(26) La palabra «cinco» es un término que designa un número primo

Con esto damos un paso más hacia el centro de lo que presupone el análisis de Carnap. De lo que se trata con él es de eliminar pseudoproblemas filosóficos sustituyendo el discurso sobre entidades de dudosa condición ontológica por el discurso acerca del lenguaje en el que hablamos de ellas. Y la pregunta es: ¿qué se gana con esto? ¿Quedan eliminadas realmente esas cuestiones porque las reformulemos como cuestiones metalingüísticas? Podríamos sentirnos inclinados a una respuesta afirmativa si estas últimas fueran cuestiones solubles por sí solas y al margen de la correspondiente cuestión ontológica. Pero, desgraciadamente, al menos en buena parte de los casos, no es así. Porque usamos el lenguaje para hablar de lo extralingüístico, de la realidad, sea existente, ficticia, inventada, construida, empírica, especulada, física, matemática, o del tipo que sea, y no tenemos otras expresiones, empleamos las expresiones que empleamos, y en los contextos que Carnap considera, no por sí mismas ni como fines en sí mismas. Por ejemplo: parece que si decidimos construir las expresiones numéricas como expresiones de clase del segundo nivel será porque concebimos los números como clases de clases, y no viceversa. ¿Qué ganamos con plantear nuestro problema en términos del modo formal? Si nuestras expresiones de color van siempre acompañadas de designaciones espaciales es porque no concebimos el color sino como una manifestación de la realidad física, pero no viceversa. Si el conocimiento se expresa en oraciones y no en nombres aislados es porque la realidad se nos da en forma de conexiones o hechos, y no en forma de cosas separadas. Con esto no pretendo negar que la traducción al modo formal pueda contribuir, en numerosas ocasiones, a dar claridad sobre el tema en discusión, y que sirva para revelar, en muchos casos, el carácter metalingüístico del problema debatido. Pero lo que no puede aceptarse es que la traducibilidad al modo formal constituya el rasgo característico de las pseudoproposiciones de objeto, pues, en principio, también las auténticas proposiciones de objeto son susceptibles de tal tra-

ducción. Qué aparentes proposiciones de objeto hayamos de rechazar es cuestión que tendremos que decidir en función de nuestra concepción de la realidad y del conocimiento. En principio, que una afirmación puede formularse en el modo formal no es de ninguna manera indicio de que no pueda aceptarse como afirmación sobre objetos. Puede muy bien ocurrir que ambas formulaciones sean correctas. Por eso el *Tractatus* enuncia con frecuencia las mismas afirmaciones de ambos modos: como afirmaciones sobre la realidad y como afirmaciones sobre el lenguaje. Lo que sí puede mostrar la traducción de una proposición al modo formal es la posible raíz lingüística que pueda tener un determinado problema filosófico. Así visto, el método de Carnap no diferiría mucho del empleado por el segundo Wittgenstein para disolver los problemas filosóficos.

Por último, Carnap ha presentado el modo formal como un tipo de discurso que versa sobre la forma del lenguaje, y compuesto por proposiciones sintácticas. Pero, como ya se habrá advertido, varios de los ejemplos anteriores sólo muy forzosamente pueden caracterizarse así. Cuandoquiera que se alude a la designación de las palabras, a su significado o a su sinonimia con otra, como ocurre claramente por lo menos en los ejemplos (3), (4), (16) y (20) de los que he citado (y también en (25) y (26), pero éstos no son de Carnap), estamos ante una caracterización semántica, puesto que se involucra la relación entre la palabra y la realidad. Para tal caracterización la sintaxis lógica es del todo insuficiente. Esto es algo que el propio Carnap no tardaría en reconocer.

8.3 De la sintaxis lógica a la semántica formal: el concepto semántico de verdad

En *Sintaxis lógica del lenguaje* (1934), cuya última parte acabamos de comentar, Carnap había presentado una teoría formal del lenguaje en el sentido más riguroso. Esta teoría es formal porque en ella no se hace alusión al significado de los símbolos (palabras) ni al sentido de sus concatenaciones (oraciones), sino única y exclusivamente a las clases de aquéllos y al orden en que son admisibles para constituir una secuencia bien formada. Las reglas que rigen este último aspecto son las reglas de formación; las que regulan qué secuencias de símbolos pueden derivarse a partir de otras, son las reglas de transformación. Una regla de formación para un lenguaje determinado puede ser ésta: toda secuencia formada por un símbolo de predicado seguido por uno o más símbolos de individuo es correcta, esto es, es una oración de ese lenguaje. Una regla de transformación, también para un determinado lenguaje, puede ser así: de una oración de la forma p junto con otra de la forma $\text{si } p \text{ entonces } q$, puede derivarse q (es la vieja regla de inferencia lógica llamada *modus ponens*). La idea que Carnap explicita y defiende en su obra como más novedosa es la de que las características lógicas de las oraciones dependen exclusivamente de su forma, esto es, de su estructura sintáctica, o lo que tanto da, de las reglas de for-

mación y de transformación propias del lenguaje al que pertenecen las oraciones en cuestión, reglas que pueden formularse sin aludir para nada al significado de las oraciones ni de las palabras que las componen. Según esto, basta la sintaxis para caracterizar a una oración como analíticamente verdadera o falsa, para decidir si dos oraciones son entre sí compatibles o contradictorias, o si una se deduce de otra, etc. (*Logische Syntax der Sprache*, secc. 1; «Filosofía y sintaxis lógica», II.5).

En su libro, y a causa —según dice— de las deficiencias de los lenguajes naturales, Carnap se limita a desarrollar la sintaxis lógica de un par de lenguajes simbólicos artificiales contruidos a tal efecto, formulando en una parte ulterior la teoría general. Sin embargo, apenas publicado aquel, Carnap había llegado ya a la conclusión de que la definición sintáctica de los conceptos lógicos mencionados es insuficiente, y de que el estudio sintáctico del lenguaje, o estudio acerca de la forma de las expresiones, ha de ser completado con un estudio semántico, en el que consideremos la relación entre las expresiones y la realidad. Este estudio incluirá, por ello, conceptos que en la sintaxis estaban ausentes, en especial los de significado y verdad. El cambio de orientación, que no excluye la sintaxis sino que la completa con la semántica, se debió a la influencia de los lógicos polacos, y en particular de Tarski, como el propio Carnap reconoce («Intellectual Autobiography», p. 60).

Tarski había publicado, en 1933, en polaco, su importante estudio sobre «El concepto de verdad en los lenguajes formalizados», que fue publicado en alemán en 1936 (de hecho, ese título, que es el usual, corresponde a la versión alemana; en polaco, el título decía «... en los lenguajes de las ciencias deductivas»). En su trabajo, Tarski formulaba las condiciones que ha de reunir una definición satisfactoria del concepto de verdad, y suministraba tal definición, y lo hacía, por las razones que ahora veremos, para los lenguajes formalizados, tomando como ejemplo el cálculo de clases. Su concepción semántica de la verdad estaba formulada en términos lógicos tan rigurosos que Carnap no necesitó más para aceptar algo que ya venía rumiando en los años anteriores: que los conceptos semánticos pueden también expresarse en el lenguaje riguroso de la lógica, y que por consiguiente puede hablarse de manera formal acerca de la relación entre el lenguaje y el mundo. Aunque los nuevos conceptos encontraron gran resistencia en principio, resultaron ser extremadamente fecundos, y de ellos procede la importante obra semántica de Carnap. Consideremos primero, brevemente, en qué consistía la aportación de Tarski.

Las condiciones establecidas por éste para una definición de la verdad exigían que tal definición fuera materialmente adecuada y formalmente correcta. Veamos primero lo que se refiere a la primera condición. El concepto de verdad que Tarski pretendía definir de forma rigurosa era el concepto clásico de la verdad entendida como conformidad de la proposición con la realidad, esto es, la verdad como correspondencia o acuerdo entre nuestras afirmaciones y los hechos; lo que Tarski va a definir es la expresión «oración verdadera». Es el concepto de verdad que está presente

en la vieja sentencia aristotélica: «Decir de lo que es que no es o de lo que no es que es, es falso; decir de lo que es que es, o de lo que no es que no es, es verdadero» (*Metafísica*, 1011 b, 26 ss.; citado por Tarski en su estudio, secc. 1, nota 2, así como en «La concepción semántica de la verdad», I.3; este último es un artículo de 1944 en el que presenta una versión informal y abreviada de su trabajo original, contestando a algunas críticas). Pues bien, para una definición de este concepto, Tarski propone como condición de adecuación material la siguiente: una definición sólo será adecuada si de ella se siguen todas las equivalencias de la siguiente forma:

(T) X es verdadera si y sólo si p

¿Qué quiere decir esto? Consideremos el siguiente ejemplo de Tarski: «La nieve es blanca». Dado el concepto de verdad que aspiramos a definir, ¿cuándo podemos decir que esta oración es verdadera? La respuesta es obvia: si, y solamente si, la nieve es blanca. Si designamos esa oración por medio de la letra X , podemos escribir:

(1) X es verdadera si y sólo si la nieve es blanca

y si queremos ser más explícitos:

(2) «La nieve es blanca» es verdadera si y sólo si la nieve es blanca

Esta oración es una equivalencia de la forma de (T), es un ejemplo que responde a ese esquema. De una definición de verdad deben, pues, seguirse todos los ejemplos que tienen la forma del esquema (T), o dicho de otro modo, que cumplen con la condición (T) de adecuación material. Lo dicho implica, naturalmente, que, a pesar de lo que muchas veces se afirma mal-entendiendo a Tarski, (T) no expresa la definición semántica de la verdad, sino únicamente una de las condiciones que ésta ha de cumplir. Tomando (T), y sustituyendo « p » por una oración declarativa y « X » por un nombre de esa oración, obtenemos una equivalencia que es, en palabras de Tarski, «una definición parcial de la verdad, que explica en qué consiste la verdad de esta oración individual». Y añade: «La definición general debe ser, en cierto sentido, una conjunción lógica de todas estas definiciones parciales.» («La concepción semántica de la verdad», I.4).

Podemos pasar ahora a la otra condición estipulada por Tarski, la que exige una definición de la verdad que sea formalmente correcta. Lo primero que esta condición exige es que la definición se formule en un lenguaje de nivel más alto que aquel para el cual se formula. Puesto que la definición se da para oraciones, éstas pertenecerán a un lenguaje L ; pues bien, la definición a su vez debe pertenecer, no a L , sino a un nuevo lenguaje que trate de L , esto es, a un metalenguaje de L , que podemos llamar $L + 1$. ¿Cuál es la razón de esta exigencia? Evitar paradojas como la del mentiroso, que al decir «Todo lo que yo digo es falso» da lugar a la siguiente situación:

si lo que dice es verdad, entonces es falso, y si lo que dice es falso, entonces es verdad. Evitar esta paradoja, que haría imposible establecer un concepto rigurosamente lógico de verdad, requiere evitar un lenguaje en el que una oración pueda afirmar la verdad o la falsedad de sí misma (a estos lenguajes los llama Tarski «semánticamente cerrados»). De aquí que la definición de verdad para un lenguaje haya de enunciarse en su metalenguaje. ¿Qué diferencias relevantes habrá entre ambos? Puesto que la definición de la verdad viene a equivaler al conjunto de equivalentes de la forma (T), el examen de ésta nos puede suministrar la respuesta, ya que el esquema (T) está, en efecto, formulado metalingüísticamente. Y examinando dicho esquema observamos lo siguiente: para una oración p del lenguaje objeto, (T) contiene, en primer lugar, a p misma; en segundo lugar, contiene un nombre o designación de p , que es X (podría serlo igualmente « p », esto es, el resultado de colocar p entre comillas, o cualquier otro recurso que pudiéramos idear para referirnos a p); contiene, en tercer lugar, el predicado «es verdadera»; y contiene, finalmente, el functor lógico bicondicional representado por la expresión «si y sólo si». En resumen, el metalenguaje debe contener como mínimo: el predicado «es verdadera», constantes lógicas proposicionales, todas las oraciones asertóricas del lenguaje objeto (esto es, todas las oraciones del lenguaje objeto que pueden ser verdaderas o falsas), y los nombres de estas oraciones. Por lo que respecta a las oraciones del lenguaje objeto, la exigencia de que estén contenidas en el metalenguaje puede sustituirse —según Tarski— por esta otra: que sean traducibles a oraciones del metalenguaje («La concepción semántica de la verdad», I.9). Así interpretado el esquema (T), p no sería en él una oración del lenguaje objeto, sino una oración del propio metalenguaje que traduciría una cierta oración del lenguaje objeto. Por consiguiente, para que una definición de verdad sea formalmente correcta se requiere que esté formulada en un metalenguaje, y que éste cuente, al menos, con los recursos indicados. A esto sólo hay que añadir que tanto la estructura del lenguaje objeto como la de su metalenguaje tienen que ser formalmente especificables, ya que necesitamos saber con todo rigor cuáles son las oraciones correctas y bien formadas de uno y otro. Pero esto no lo sabremos si no conocemos cuáles son exactamente las reglas de formación de ambos lenguajes. Esto explica que Tarski construya su definición para lenguajes formalizados.

Con esto tenemos las dos condiciones que ha de cumplir una definición aceptable de la verdad. Ahora hay que construir la definición. Para ello, Tarski recurrirá al concepto semántico de satisfacción (*op. cit.*, I.11; «El concepto de verdad en los lenguajes formalizados», secc. 3). El concepto de satisfacción, a diferencia del concepto de verdad, es aplicable a aquellas expresiones lógicas con las que se forman las oraciones, esto es, a las funciones proposicionales. Así, la oración «Todos los hombres son mortales», que formalizamos en el cálculo de predicados como $\Lambda x (Fx \rightarrow Gx)$, se compone cuantificando universalmente el condicional formado por las funciones proposicionales Fx y Gx ; esto es, equivale a afirmar: vale para

todo x que si x tiene la propiedad F , entonces tiene la propiedad G . De Fx y Gx no puede decirse que sean verdaderas o falsas, sino que son satisfechas o no satisfechas. Siendo « F » el predicado «es hombre», la función Fx es satisfecha cuando la variable « x » es sustituida por nombres de seres humanos, como «Sócrates», «Tarski», «Napoleón», etc., pero no cuando es sustituida por nombres de objetos de otra clase, como animales, edificios, entidades matemáticas, etc. Cuando las funciones contienen una sola variable individual libre, como en el ejemplo que acabamos de considerar, son satisfechas por objetos singulares. Las funciones con varias variables individuales, o sea, las relaciones, son satisfechas por n -tuplas ordenadas de objetos. Así, la función proposicional Rxy , donde « R » es la relación «es padre de», es satisfecha por cualquier par ordenado de objetos tales que el primero sea padre del segundo, como en el caso $\langle \text{Abraham, Isaac} \rangle$. Y análogamente para funciones proposicionales más complejas. A las funciones proposicionales se les denomina también oraciones abiertas, esto es, no cuantificadas. Puesto que no hay límite para el número de variables que pueden componer una función proposicional, nos encontraríamos con sucesivas nociones de satisfacción según el número de variables de la función. Así, para funciones con dos variables, el concepto de satisfacción relacionaría la función con pares ordenados; para funciones con tres variables, la satisfacción sería en términos de tríadas y así sucesivamente. Con lo que tendríamos infinitas nociones de satisfacción. A fin de conseguir una noción general de satisfacción, Tarski define ésta como una relación entre funciones proposicionales y secuencias o sucesiones de infinitos objetos, conviniendo que cualquier secuencia infinita satisface una función proposicional con n variables individuales libres en caso de que la serie ordenada de los n primeros objetos de esa secuencia satisfaga la función dicha, con independencia de cuáles sean los demás objetos de la secuencia. Por ejemplo: sea la función proposicional $Rvxyz$, significando « v equidista de x , y y z ». Pues bien, tal función es satisfecha por cualquier secuencia infinita cuyos cuatro primeros objetos sean tales que el primero equidiste de los tres siguientes; el resto de los objetos de la secuencia es irrelevante a estos efectos.

Esto por lo que se refiere a una función proposicional positiva. Si consideramos ahora la negación de una función tal, diremos que la negación de una función proposicional Z es satisfecha por todas aquellas secuencias que no satisfacen Z . Si consideramos la conjunción de dos funciones proposicionales Z y W , diremos que esta conjunción es satisfecha por todas aquellas secuencias que satisfacen Z y W . Si lo que tenemos, en cambio, es la disyunción incluyente de Z y W , diremos que tal disyunción es satisfecha por todas aquellas secuencias que satisfacen Z , por todas las que satisfacen W , y naturalmente, por todas las que satisfacen Z y W al tiempo. Y así sucesivamente. Como habrá podido apreciarse, a pesar del curso informal y mínimamente técnico de esta explicación, el concepto de satisfacción es definido por Tarski de modo recursivo, partiendo del caso más simple y

definiendo a partir de él la satisfacción para casos sucesivamente más complejos.

Podemos pasar ahora de las funciones proposicionales, u oraciones abiertas, a las oraciones cerradas, u oraciones en sentido estricto, esto es, aquellas que carecen de variables libres. Acabamos de ver que, para la satisfacción de una función proposicional con n variables individuales, solamente son relevantes los n primeros miembros de cualquier secuencia de infinitos objetos. Puesto que una oración (cerrada) carece de variables individuales libres, puede afirmarse que para su satisfacción no es relevante ninguno de los objetos de una secuencia. De aquí que, si la oración es verdadera, entonces es satisfecha por todas las secuencias de objetos; y si es falsa, entonces no es satisfecha por ninguna. Y ésta es la definición semántica de la verdad que propone Tarski:

Una oración es verdadera si y sólo si es satisfecha por toda secuencia de objetos, o lo que es lo mismo, por todos los objetos en general; y una oración es falsa si y sólo si no es satisfecha por ninguna secuencia de objetos, o lo que tanto da, por ningún objeto («La concepción semántica de la verdad...», I.11; «El concepto de verdad en los lenguajes formalizados», secc. 3, definición 23).

A primera vista, la definición resulta esotérica. ¿Cómo puede decirse que una oración verdadera es satisfecha por todos los objetos y que una oración falsa no lo es por ninguno? Para entenderlo, hay que tener a la vista que las oraciones se explican lógicamente como el resultado de cuantificar las funciones proposicionales ligando las variables libres por medio de cuantificadores, o bien como el resultado de sustituir las variables libres por nombres. Considérese ahora una función proposicional tan simple como Fx , donde « F » sea el predicado «es hombre». Esta función, como hemos visto, es satisfecha por toda secuencia de objetos cuyo primer objeto sea un ser humano. Por ejemplo, cualquier secuencia cuyo primer objeto sea Cervantes satisface esa función. Cerremos ahora esa oración abierta ligando la variable « x » con el cuantificador particular, así:

$\forall x Fx$. Esto lo leeremos como:

(3) Al menos para un x vale que x es hombre

que aproximadamente equivale a la afirmación «Hay hombres». Puesto que lo que aquí decimos es que en algún caso, al menos en uno, la función Fx es satisfecha, la oración (3) es a su vez satisfecha no sólo por las secuencias que satisfacen Fx , sino también por las que no satisfacen Fx , ya que toda secuencia que no satisfaga Fx es compatible con $\forall x Fx$. Pensemos ahora en el caso de una oración falsa: «Hay unicornios», esto es:

(4) Al menos para un x vale que x es un unicornio

que podemos formalizar como $\forall x Hx$, donde « H » significa «es un unicornio». Que esta oración es falsa significa que no hay ningún caso en el que

la función proposicional Hx sea satisfecha, o dicho de otra forma, que no hay ningún objeto que la satisfaga, y por consiguiente, que no hay ninguna secuencia infinita de objetos que la satisfaga. Luego *a fortiori*, tampoco habrá secuencia alguna que satisfaga la oración (4). Luego una oración falsa es una oración que no es satisfecha por ninguna secuencia.

Lo dicho para las oraciones cerradas con el cuantificador particular puede aplicarse igualmente a las oraciones generales en razón de la interdefinibilidad de los cuantificadores, que hace que cualquier oración general puede convertirse en una oración particular y viceversa. Considérese la oración general «Todos los hombres son mortales», que leeríamos lógicamente como:

(5) Para todo x vale que, si x es un hombre, entonces x es mortal

y que formalizaríamos como $\Lambda x (Fx \rightarrow Gx)$. Por definición de los cuantificadores, esto equivale a $\neg \vee x (Fx \wedge \neg Gx)$, o sea, a la oración:

(6) No vale ni siquiera para un x que x sea hombre y x no sea mortal

Y la verdad de esta oración consiste en que, no habiendo ningún objeto que satisfaga la función Fx (« x es hombre») y no satisfaga a la vez la función Gx (« x es mortal»), no hay tampoco ninguna secuencia que satisfaga la oración:

(7) Vale al menos para un x que x es hombre y x no es mortal

Y puesto que (7) es la negación de (5) y de (6), no habiendo ninguna secuencia que satisfaga (7), hay que concluir que todas las secuencias satisfacen las oraciones (5) y (6), las cuales son por eso verdaderas.

El procedimiento puede igualmente aplicarse a la otra manera de transformar las funciones proposicionales en oraciones, a saber, sustituir las variables individuales por nombres. Hagámoslo en el simple caso de la función Fx , « x es hombre», y transformémosla, por ejemplo, en:

(8) Cervantes es un hombre

La función Fx , dijimos antes, es satisfecha por cualquier secuencia infinita de objetos que comience con un ser humano. Pues bien, la oración (8) es satisfecha por todas las secuencias, pues lo es, en primer lugar, por cualquier secuencia que comience precisamente con Cervantes, ya que todas las secuencias que comiencen con Cervantes satisfacen Fx ; y lo será luego por todas las demás secuencias en cuanto que éstas sólo difieren de aquéllas en comenzar con un objeto distinto de Cervantes y son por tanto irrelevantes para el caso. Consideremos ahora que la oración fuera falsa. Por ejemplo:

(9) Cervantes es un matemático

No hay ninguna secuencia de objetos que dé comienzo con Cervantes y que satisfaga la función « x es un matemático», y puesto que las demás secuencias son irrelevantes a estos efectos, hay que concluir que ninguna secuencia satisface (9). (No hace falta añadir que las letras predicativas F , G , H , R , usadas en los ejemplos anteriores, han funcionado como constantes, puesto que representaban, en cada caso, predicados determinados.)

Si el lector echa de menos en todas estas explicaciones el conocido ejemplo de Tarski que mencionamos al principio, «La nieve es blanca», puede aplicarle el análisis anterior teniendo en cuenta que, por ser el término «nieve» un término de masa, como «agua», «aire», etc., conviene transformar el ejemplo en «Para todo x , si x es un fragmento suficientemente grande de nieve, entonces x es blanco» (si el fragmento fuera muy pequeño, no se vería propiamente blanco sino más bien incoloro o grisáceo).

La idea fundamental de Tarski al definir la verdad en términos de satisfacción es la siguiente: o bien una oración es satisfecha por todas las secuencias, o bien no es satisfecha por ninguna. En el primer caso es verdadera; en el segundo caso, es falsa.

El tratamiento tarskiano del concepto de verdad presupone, como se habrá notado, que toda oración es o verdadera o falsa, y que si una oración es verdadera, entonces su negación es falsa, y viceversa. Se asumen, pues, los principios de bivalencia y de tercero excluido. Hay que subrayar también, y esto es muy notable, que la definición semántica de la verdad es tan aplicable a las verdades empíricas como a las verdades lógicas o analíticas, pues si todas las secuencias de objetos satisfacen la oración «Todos los hombres son mortales», con mayor razón satisfarán una oración como «Todos los hombres son hombres».

La explicación precedente se ha dado en términos lo más intuitivos posible, y los ejemplos han sido tomados del lenguaje común. Ello podría hacer pensar que la definición de Tarski es aplicable a los lenguajes naturales. Sin embargo, y como ya indiqué al principio, la definición está pensada tan sólo para lenguajes formalizados, como muestra con toda claridad el propio título del estudio de Tarski. De hecho, ya la primera sección de su trabajo, que dedica a la verdad en el lenguaje ordinario, lo lleva a conclusiones del todo negativas. Para Tarski, un lenguaje natural tiene el inconveniente de ser al propio tiempo su propio metalenguaje, con lo que se incumple la exigencia de que el metalenguaje en el que se defina la noción de verdad sea más rico que el lenguaje objeto para el cual se defina, y se originan paradojas semánticas como las del mentiroso que imposibilitan la definición de verdad. De otra parte, un lenguaje natural no es formalmente especificable, pues no hay reglas exactas que determinen qué oraciones son correctas, y están bien formadas, a lo que contribuye además el hecho de que los lenguajes naturales se encuentren en perpetua evolución. En suma: que una definición de verdad para un lenguaje natural no podría ser formalmente correcta. De aquí que Tarski desista del intento y dé fin a la primera sección de su trabajo concluyendo: «Si las considera-

ciones anteriores son exactas, entonces la posibilidad misma de usar la expresión «proposición verdadera» de forma coherente y de acuerdo tanto con los principios de la lógica como con el espíritu del lenguaje común, así como la posibilidad de construir una definición correcta de esa expresión, parece muy cuestionable.» La opinión de Tarski resulta hoy, sin embargo, excesivamente pesimista. Davidson, como veremos en su momento, intenta construir la teoría del significado en los lenguajes naturales sobre la base de una teoría de la verdad de corte tarskiano. Por otro lado, es indudable que la lingüística transformatoria ha dado un gran paso hacia la definición de «oración correcta» para un lenguaje natural, y si bien, como vimos en el capítulo 4, el paradigma chomskiano está poblado de alternativas y la situación es aún confusa, en particular por lo que respecta a las reglas de tipo semántico, es patente que no se puede mantener hoy frente al lenguaje natural una actitud tan desconfiada y negativa como la que se podía tener en los años treinta. Pero volveremos sobre este tema más adelante. (Se encontrará una hábil defensa de los lenguajes naturales frente a las críticas de Tarski, a propósito de la paradoja del mentiroso, en el trabajo de Francisco Gracia, «La paradoja del mentiroso en los lenguajes naturales».)

La posible significación de la definición de Tarski para resolver las discusiones filosóficas tradicionales acerca del concepto de verdad es cuestión ulterior que pertenece más bien a la teoría del conocimiento y que nos alejaría en exceso del tema del lenguaje, que primariamente nos ocupa. Como ya señalé, Tarski pensaba que su concepto semántico de la verdad hace justicia a una teoría que, como la aristotélica, entiende la verdad como correspondencia o adecuación entre nuestras palabras y la realidad, pero desde luego él, por su parte, no pretendía en modo alguno mediar en ninguna disputa filosófica, y explícitamente declara no creer que haya nada que corresponda a la expresión «el problema filosófico de la verdad»; en todo caso, habría problemas diversos, filosóficos y no filosóficos, en torno a la noción de verdad («La concepción semántica de la verdad», II.18). Es más, declara asimismo que el concepto semántico de la verdad es epistemológicamente neutral, y compatible, por ello, tanto con el idealismo como con el realismo (*ibidem*). Siendo así la posición de Tarski, huelga añadir que la evaluación filosófica de su definición de la verdad ha oscilado entre dos polos que podemos ver ejemplificados en Black y en Popper, respectivamente. Para el primero, la definición tarskiana, precisamente a causa de su neutralidad, carece de relevancia filosófica (Black, «The Semantic Definition of Truth», secc. 9). El segundo, por el contrario, piensa que, por haber suministrado una formulación lógica del concepto de correspondencia, la definición de Tarski ha resuelto viejas disputas filosóficas sobre la verdad y viene a apoyar el realismo metafísico (Popper, «Comentarios filosóficos sobre la teoría de la verdad de Tarski», secc. I.). (Se encontrará una breve y asequible exposición de la teoría tarskiana en Haack, *Philosophy of Logics*, cap. 7, y, en términos algo más técnicos, en Quine, *Filosofía de la lógica*, cap. 3.)

Por lo que nosotros habíamos abordado el concepto semántico de verdad no era, sin embargo, por sus posibles méritos filosóficos, sino porque constituye el primer tratamiento rigurosamente lógico de un concepto semántico, y porque fue el ejemplo que Carnap tuvo a la vista para el desarrollo de su nuevo programa. La contribución de Carnap a la semántica formal se materializó principalmente en tres libros que aparecieron con cortos intervalos durante los años cuarenta: *Introduction to Semantics* (1942), *Formalization of Logic* (1943), y *Meaning and Necessity* (1947). Lo primero que hay que tener en cuenta es que la nueva atención a los conceptos semánticos no pretende reemplazar al estudio de la sintaxis sino completarla (*Introduction to Semantics*, secc. 39). Las modificaciones más importantes conciernen a la distinción entre las constantes lógicas y los signos descriptivos (símbolos de individuos y de propiedades, por ejemplo), y afectan por consiguiente a la distinción entre verdad empírica o de hecho y verdad lógica. Tales distinciones han de trazarse —piensa ahora Carnap— primariamente en la semántica, si bien podrán ser formalizadas, esto es, representadas por medio de conceptos sintácticos en un cálculo construido al efecto (*loc. cit.*). Por lo que se refiere a la traducción de oraciones filosóficas al modo formal o sintáctico, que hemos considerado en la sección anterior, Carnap reconoce que, para todas aquellas que tienen que ver con los conceptos de designación y de significado, parece más natural traducirlas a un modo semántico, esto es, a un metalenguaje que trate de propiedades semánticas y no simplemente formales; en este caso estarían los ejemplos que mencionamos al final de la sección precedente. Por lo demás, Carnap sigue manteniendo que, por las razones que ya conocemos, el modo material es peligroso. Por último, y de acuerdo con este giro, la tesis de que la filosofía se reduce a sintaxis lógica del lenguaje, en general, y del lenguaje científico, en particular, queda ampliada para admitir también la semántica como parte de la lógica de la ciencia. La filosofía consiste, en suma, en el análisis semiótico del lenguaje en general, y del lenguaje científico en particular (*loc. cit.*).

Puesto que de las tres obras que Carnap dedica a la semántica, la última es la más elaborada y la que ha ejercido mayor influencia, estudiaremos a continuación algunos de los conceptos semánticos más característicos con referencia a esta obra.

8.4 Extensión e intensión: ontología y semántica

En *Meaning and Necessity* Carnap presenta un método de análisis del significado que denomina «método de la extensión y de la intensión». Este método se contrapone a aquellos otros métodos que, como el de Frege, toman las expresiones lingüísticas primariamente como nombres, y construyen fundamentalmente el significado como la relación de designación o referencia. El método de Carnap pretende evitar este unilateral enfoque distinguiendo dos operaciones que podemos realizar con las expre-

siones. La primera es analizar la expresión semánticamente (o, en un sentido amplio del término, lógicamente) para entender su significado. La segunda consiste en investigar la situación a la que se refiere la expresión, a fin de determinar su verdad; esta última es, por tanto, una operación empírica. La primera operación nos da la intensión de la expresión; la segunda nos suministra su extensión. La segunda operación, siendo una operación de aplicación, presupone la primera. Si no sabemos lo que significa una oración, no podemos saber en qué casos es verdadera. Por eso afirma Carnap que una expresión tiene primariamente una intensión, y secundariamente una extensión (secc. 45). Hablar de la intensión de una expresión es nada más que una manera técnica, y suponemos que por ello más rigurosa, de hablar de su significado. Hablar de la extensión equivale, simplemente, a hablar de la aplicación de la expresión, y por tanto, si se trata de oraciones, de su verdad.

Para formular su método, Carnap elige un lenguaje objeto compuesto por las constantes lógicas usuales, variables individuales, y, como constantes descriptivas, constantes de individuo y de predicado, estas últimas en número finito. Este lenguaje, que podemos llamar *L*, y que será ampliado en lugares ulteriores de su libro para propósitos que ahora no nos interesan, contiene reglas específicas de cuatro clases fundamentales (la razón de escoger como lenguaje objeto un lenguaje simbólico o artificial es precisamente contar con reglas determinadas y rigurosas). En primer lugar, reglas de formación que determinen qué series de signos constituyen oraciones sintácticamente correctas o bien formadas. En segundo lugar, reglas de designación, esto es, reglas que determinen qué objeto individual designa cada constante individual, o qué propiedad o relación designa cada constante predicativa. En tercer lugar, reglas de verdad, y en cuarto lugar, reglas de ámbito (*range*). Veamos ahora las reglas de verdad. Estas son definidas recursivamente partiendo de la verdad de la oración atómica, y explicitando reglas sucesivas para oraciones compuestas con cada una de las conectivas. La regla de verdad para el caso más simple, el de la oración atómica, es así: «una oración atómica de *L* que consista en un predicado seguido por una constante individual es verdadera si y sólo si el objeto al que se refiere la constante posee la propiedad a la que se refiere el predicado» (*Meaning and Necessity*, secc. 1, 1-3). Como se ve, esta regla deriva de una definición tarskiana de la verdad (que Carnap había recogido explícitamente en la sección 7 de *Introduction to Semantics*). La regla de verdad para cada clase de oración compuesta corresponde, naturalmente, a la tabla de verdad de la conectiva con la que se efectúe la composición: conjunción, condicional, etc. Se habrá apreciado que las reglas de verdad presuponen las reglas de designación, pues no podemos determinar la verdad o falsedad de una oración sin saber lo que designan sus términos. Así, por ejemplo, si la constante individual *a* designa a Cervantes, y la constante predicativa *F* designa la propiedad de ser un escritor, entonces la oración *Fa* dice que Cervantes es un escritor (entendiendo el «es» intemporalmente), y es verdad puesto que, en efecto, Cervantes es un escritor. Se apreciará asimismo

que las oraciones de nuestro lenguaje L son expresiones que pueden ser verdaderas o falsas, y por consiguiente que el método de Carnap está pensado para este tipo de oraciones, esto es, para oraciones declarativas. O dicho de otra manera que resultará familiar: que el significado que Carnap pretende analizar es el significado que en la sección 8.1 hemos llamado significado cognitivo. A todas las expresiones, sean oraciones o partes de oraciones, que poseen este tipo de significado, y que por tanto pueden ser objeto del análisis que estamos considerando, las llama Carnap «designadores». Esto no quiere decir que haya que entender estas expresiones como nombres, pues, según ya hemos visto, se trata de evitar una teoría referencialista del significado. Los designadores incluyen, según esto, las oraciones, los predicados y los nombres o descripciones de objetos individuales.

Veamos ahora las reglas que habíamos mencionado en cuarto lugar, las reglas de ámbito. Para ello hay que introducir el concepto de descripción de estado: una descripción de estado en L es un conjunto D de oraciones de L , tal que D contiene, para toda oración atómica de L , o bien esta oración o bien su negación, pero no ambas, y ninguna otra oración (secc. 2, p. 9). Es patente que una descripción de estado suministra una descripción completa de un posible estado de nuestro universo del discurso, esto es, del universo formado por aquellas entidades individuales, propiedades y relaciones designadas por las constantes descriptivas de nuestro lenguaje objeto L . Por ello sugiere Carnap que las descripciones de estado corresponden a los posibles estados de cosas del *Tractatus* o a lo que Leibniz había llamado mundos posibles, concepto este último que la teoría semántica más reciente ha vuelto a poner en circulación, según veremos. Ahora podemos dar reglas semánticas que determinen, para cada oración de L , si la oración vale o no en una descripción de estado determinada, o sea, si la oración sería o no verdadera en el caso de que la descripción de estado lo fuera. Tales reglas son sumamente sencillas, pues se trata de reglas como éstas: una oración atómica vale en una descripción de estado si y sólo si pertenece a ella; la negación de una oración vale en una descripción de estado si y sólo si la oración no vale en ella; la disyunción (incluyente) de dos oraciones vale en una descripción de estado si y sólo si una de las dos oraciones, o las dos, valen en ella; etc. Así, por ejemplo, si el universo de nuestro discurso contiene a Cervantes y las propiedades de ser un escritor y de ser un matemático, la oración atómica «Cervantes es un escritor» vale en cualquier descripción de estado en la que sea verdadera, o lo que es lo mismo, en cualquier situación posible en la que sea verdad que Cervantes es un escritor (y esta situación posible es, por cierto, la situación histórica real). En esta misma descripción de estado vale también la disyunción «Cervantes es un escritor o es un matemático». En la descripción de estado correspondiente a la situación histórica real no vale, en cambio la oración «Cervantes es un matemático», si bien es claro que esta oración atómica vale en la descripción de estado correspondiente a una posible situación en la que Cervantes hubiera sido un matemático, pues

en tal situación esta oración hubiera sido verdadera. A la clase de todas las descripciones de estado en las que vale una cierta oración X se le denomina ámbito de X . Y las reglas de ámbito son reglas como las que acabamos de considerar, esto es, reglas que determinan en qué descripciones de estado vale cada oración de L . Para entenderlo mejor, considérese el siguiente ejemplo. Sea un universo constituido por tres personas a , b y c , que están siendo sometidas a un experimento de neurofisiología, y cuyas únicas propiedades relevantes a estos efectos son estar dormido y estar despierto. Sea L el lenguaje en el que estamos describiendo las situaciones de este experimento, y sea, por consiguiente, una oración atómica de L la oración:

(1) La persona a está dormida

¿Cuál es el ámbito de esta oración? La respuesta es: el conjunto de las descripciones de estado en las que vale (1), o dicho de otra forma, el conjunto de todas las situaciones en las que (1) es verdadera. Y este conjunto está formado por las siguientes oraciones:

(2) La persona a está dormida, y b y c también lo están

(3) La persona a está dormida, pero b y c están despiertas

(4) La persona a está dormida, pero b está despierta y c dormida

(5) La persona a está dormida, y b también lo está, pero c está despierta

Las oraciones (2) a (5) constituyen todas las descripciones de estado de L en las que vale la oración (1); por consiguiente, esas oraciones son el ámbito de (1). (Desde el punto de vista lógico, las oraciones (2) a (5) son simples conjunciones de tres oraciones atómicas; en el ejemplo he preferido, sin embargo, enunciarlas de modo más idiomático; el lector debe tener en cuenta que términos como «pero» y «también» son lógicamente irrelevantes.)

¿En qué descripciones de estado no vale (1)? En las que corresponden a todas aquellas situaciones en que a está despierto; por ejemplo en esta descripción:

(6) La persona a está despierta, b está dormida y c está despierta etcétera. El análisis es igualmente aplicable a oraciones compuestas. Considérese, por ejemplo, la disyunción (incluyente):

(7) O a está despierto o b está dormido (o ambas cosas)

Esta oración (7) vale en toda descripción de estado en la que valga alguna de sus partes (o las dos). Vale, por tanto, entre otras, en las descripciones (2), (5) y (6) de las que hemos visto.

Las reglas de ámbito, junto con las reglas de designación, suministran una interpretación para toda oración del lenguaje objeto L , pues, en pala-

bras de Carnap: «conocer el significado de una oración es saber en cuáles, de todas las situaciones posibles, sería verdadera, y en cuáles no, como ha señalado Wittgenstein» (sect. 2, p. 10). Este breve párrafo suministra, desde un punto de vista formál, un correlato a lo que hemos visto formulado, desde un punto de vista epistemológico, en el principio de verificabilidad. Y es de subrayar la invocación que ahí se hace de Wittgenstein, pues muestra, una vez más, que la tesis de que el significado consiste en las condiciones de verdad la había tomado Carnap del *Tractatus*.

¿Qué relación tienen estos conceptos con la definición semántica de la verdad? La siguiente. De entre todas las descripciones de estado de un lenguaje *L*, hay una, y sólo una, que corresponde a la situación real, y es por tanto, la descripción verdadera. Está compuesta por todas las oraciones atómicas verdaderas y por las negaciones de las falsas. Por lo tanto, todas sus oraciones son verdaderas. Podemos afirmar, en consecuencia, que una oración es verdadera si y sólo si vale en aquella descripción de estado que es verdadera. Naturalmente, esto no es una definición de la verdad, puesto que el concepto de descripción de estado verdadera presupone el concepto de oración atómica verdadera, que a su vez puede definirse en términos de satisfacción al modo tarskiano que ya conocemos.

Una de las aportaciones más características de Carnap en *Meaning and Necessity* es que suministra, sobre la base anterior, una definición rigurosa para lo que la tradición filosófica venía llamando verdad lógica, necesaria o analítica. Puesto que una de las definiciones usuales de este concepto es la que hace depender la verdad analítica exclusivamente del significado de las palabras, y ya que el significado está determinado, según el método anterior, por las reglas de ámbito y de designación, que conjuntamente constituyen las reglas semánticas de nuestro lenguaje *L*, cabe pensar que una oración será analíticamente verdadera cuando su verdad dependa tan sólo de dichas reglas semánticas. Esto es lo que establece Carnap en la siguiente condición, que por recordar a la convención (T) de Tarski podemos llamar «convención (L)»:

- (L) Una oración *X* es *L*-verdadera en un lenguaje *L* si y sólo si *X* es verdadera en *L* de tal manera que su verdad se pueda establecer exclusivamente sobre la base de las reglas semánticas de *L*, sin ninguna referencia a los hechos extralingüísticos.

La expresión «*L*-verdadera» es el término técnico que sustituye a expresiones como «lógicamente verdadera», «analíticamente verdadera» o «necesariamente verdadera». Y no debe olvidarse que nuestro lenguaje *L* es un lenguaje simbólico, construido por nosotros, y, por tanto, artificial (de hecho, Carnap lo llama «sistema semántico», y lo representa con la letra «*S*», pero yo he preferido mantener a la vista el carácter «cuasi-lingüístico» de estos sistemas semánticos utilizando más bien la letra «*L*»).

¿Cómo podemos definir la verdad-*L* sobre el fundamento de la convención anterior? Puesto que una verdad necesaria es, de acuerdo con Leibniz,

una verdad que vale en cualquier mundo posible, y ya que al concepto de mundo posible corresponde en nuestra teoría el concepto de descripción de estado, Carnap propone la definición siguiente:

Una oración X es L -verdadera en $L =_{df}$ X vale en toda descripción de estado de L (secc. 2, 2-2).

La definición es adecuada puesto que, si una oración vale en toda descripción de estado, entonces es que vale con independencia de los hechos extralingüísticos y exclusivamente en función de las reglas semánticas. Por consiguiente, la definición cumple con la convención (L). Así, por ejemplo, puede mostrarse fácilmente por qué la oración « a está dormido o no lo está» es una verdad necesaria, pues en unas descripciones de estado valía « a está dormido», y en otras valía « a no está dormido», luego la disyunción de ambas oraciones valdrá en todas las descripciones de estado, y esto lo sabemos solamente en virtud de las reglas de designación y de ámbito, sin necesidad de averiguar si a está, de hecho, dormido o no.

A partir de la definición de verdad- L podemos definir otros conceptos- L como los siguientes:

X es L -falsa en $L =_{df}$ $\neg X$ es L -verdadera en L

Añadiendo a esto la definición de verdad- L se sigue que una oración es L -falsa si y sólo si no vale en ninguna descripción de estado. Este concepto pretende explicar formalmente el concepto filosófico de falsedad necesaria o lógica, esto es, de contradicción.

X L -implica Y en $L =_{df}$ el condicional $X \rightarrow Y$ es L -verdadero en L

Lo que significa que una oración X L -implica otra oración Y si y sólo si Y vale en toda descripción de estado en la que valga X . Esta definición corresponde al concepto de implicación lógica o implicación analítica (*entailment*). Se recordará que nosotros hemos distinguido ambos conceptos anteriormente en la sección 7.10; volvemos sobre este punto dentro de un poco.

X es L -equivalente a Y en $L =_{df}$ el bicondicional $X \leftrightarrow Y$ es L -verdadero en L

El resultado es que dos oraciones son L -equivalentes si y sólo si ambas valen en las mismas descripciones de estado. Y finalmente:

X es L -determinada en $L =_{df}$ X es o bien L -verdadera o bien L -falsa en L

Los conceptos- L así definidos tienen una interesante relación con los conceptos modales. Pues si una oración L -falsa es una oración que no vale en ninguna descripción de estado, es, entonces, una oración cuya verdad

no es posible. Por lo mismo, si una oración L -implica otra, esto equivale a decir que no es posible que haya una descripción de estado en la que valga la primera y no la segunda.

Por contraposición al concepto de determinación- L , Carnap define un concepto que puede corresponder al de oración contingente o fáctica, esto es, no necesaria.

X es L -indeterminada o fáctica en $L =_{df}$ X no es L -determinada

Esto significa que una oración es contingente si y sólo si hay por lo menos una descripción de estado en la que vale y una en la que no vale. Sobre esta base pueden a su vez definirse una serie de lo que llama Carnap conceptos- F , esto es, conceptos relativos a la verdad contingente, fáctica, empírica o sintética, según las varias denominaciones que a lo largo de la historia de la filosofía se han usado a este respecto.

X es F -verdadera en $L =_{df}$ X es verdadera en L pero no L -verdadera
 X es F -falsa en $L =_{df}$ $\neg X$ es F -verdadera en L

O lo que es lo mismo: una oración es F -falsa si y solo si es falsa pero no L -falsa.

X F -implica Y en $L =_{df}$ $X \rightarrow Y$ es F -verdadero en L
 X es F -equivalente a Y en $L =_{df}$ $X \leftrightarrow Y$ es F -verdadero en L

Una característica ulterior de este método de análisis semántico consiste en extender el concepto de equivalencia a todos los designadores, con la sola condición de que los designadores unidos por el signo de equivalencia (o bicondicional) sean del mismo tipo. Hemos usado hasta ahora ese signo, como es usual, entre oraciones o símbolos de oraciones; Carnap lo extenderá también a símbolos de individuo y a predicados de cualquier grado (monádicos, diádicos, etc.). En este sentido, dos expresiones de individuo son equivalentes si y sólo si designan la misma entidad individual; y dos expresiones predicativas son equivalentes si y sólo si valen para los mismos objetos (si se trata de predicados monádicos) o para las mismas series de objetos (si el predicado es poliádico). La equivalencia puede ser, naturalmente, de carácter L o de carácter F . Ambas equivalencias se definen de manera simple para todos los designadores en general, incluyendo las oraciones, como sigue (secc. 3, 3-5). Sean D y D' dos designadores cualesquiera del mismo tipo:

D es equivalente a D' en $L =_{df}$ la oración $D \leftrightarrow D'$ es verdadera en L
 D es L -equivalente a D' en $L =_{df}$ la oración $D \leftrightarrow D'$ es L -verdadera en L
 D es F -equivalente a D' en $L =_{df}$ la oración $D \leftrightarrow D'$ es F -verdadera en L

Esta ampliación del concepto de equivalencia permite ahora definir la identidad de extensión y de intensión para los designadores del modo siguiente (secc. 5, 5-1 y ss.):

Dos designadores tienen la misma extensión en $L =_{df}$ son equivalentes en L

Dos designadores tienen la misma intensión en $L =_{df}$ son L -equivalentes en L

Esto confirma lo que se anticipó al comienzo de esta sección: la extensión depende de cómo sea la realidad, la intensión depende de cómo sea el lenguaje. Pero con esto solamente sabemos en qué casos dos designadores tienen la misma extensión o la misma intensión. Y podemos preguntarnos: ¿en qué consisten una y otra? Comencemos por los predicados. La respuesta debe venir sugerida por la siguiente consideración: la extensión será algo que tengan en común las expresiones predicativas equivalentes; la intensión será algo en lo que coincidan las expresiones predicativas L -equivalentes. Carnap sugiere estas definiciones (secc. 4, 4-14 y ss.):

La extensión de una expresión predicativa es la clase de los objetos individuales a los que se aplica; su intensión es la propiedad que expresa.

Por ejemplo: la extensión del predicado «humano» o «es hombre» es la clase de los seres humanos; su intensión es la propiedad que hace que algo sea un ser humano. Dos predicados pueden poseer la misma extensión sin coincidir en su intensión; así «humano» y «bípedo implume», ya que, suponemos, ambas expresiones, aunque se aplican exactamente a los mismos objetos, no expresan la misma propiedad. Sí tendrían, en cambio, la misma intensión las expresiones «humano» y «animal racional», pues, suponemos, expresan la misma propiedad. Este cauteloso «suponemos» que he repetido se debe a lo siguiente: se notará que la afirmación de que dos predicados P y Q tienen la misma intensión significa que la oración « a tiene la propiedad P si y sólo si tiene también la propiedad Q » es una oración L -verdadera, esto es, que vale en toda descripción de estado y, por consiguiente, que es verdadera en función de las reglas semánticas de nuestro lenguaje L . Y esto presupone que algunas de esas reglas definen «humano» como «animal racional» pero no como «bípedo implume». La cuestión es: ¿por qué optar por una definición y no por otra? Volveremos sobre ello.

Las definiciones anteriores de «extensión» e «intensión» están formuladas para predicados monádicos, o de propiedades. Pueden ser reformuladas de manera apropiada para predicados poliádicos, o de relaciones, haciendo alusión, no a objetos, sino a pares de objetos (si es una relación diádica), a tríadas de objetos (si lo es triádica), etc.

Para expresiones de individuo podemos pensar en definiciones semejantes. Parece razonable pensar aquí que dos expresiones de este tipo tienen la misma extensión cuando designan el mismo objeto individual. Pero ¿cuándo diremos que poseen la misma intensión? Si para dos predicados tener la misma intensión es expresar la misma propiedad, lo que sin duda no acaba

de estar del todo claro, la solución de Carnap es aún más oscura en el caso de los designadores individuales, pues dirá que dos expresiones de individuo son *L*-equivalentes cuando expresan el mismo concepto individual. De aquí las siguientes definiciones (secc. 9):

La extensión de una expresión individual es el objeto individual al que se refiere; su intensión es el concepto individual que expresa.

La idea es que dos expresiones de individuo como «el manco de Lepanto» y «el autor de *El Quijote*» tienen la misma extensión pero diferente intensión, pues se refieren al mismo objeto o individuo pero expresan dos conceptos individuales distintos. ¿Cuándo expresarán dos designadores individuales el mismo concepto individual? Suponemos que cuando digan lo mismo acerca del objeto. Así, tendrán la misma intensión las expresiones «el autor de *El Quijote*» y «el escritor que escribió *El Quijote*», o «el manco de Lepanto» y «el que quedó inútil de una mano en Lepanto». Pues, en efecto, parece claro que la afirmación «Alguien es el autor de *El Quijote* si y sólo si es el escritor que escribió *El Quijote*» es una oración *L*-verdadera, esto es, verdadera en función de las reglas semánticas exclusivamente, y con independencia de los hechos extralingüísticos. Depende, en cambio, de estos últimos que la expresión «el autor de *El Quijote*» y «el manco de Lepanto» se refieran al mismo individuo.

Examinemos, por último, el caso de las oraciones. ¿Cuál será su extensión? De acuerdo con la línea de nuestro razonamiento tendremos que contestar: algo que tengan en común las oraciones equivalentes. Y lo que las oraciones equivalentes tienen en común es su valor de verdad. Pues dos oraciones son equivalentes cuando ambas son al tiempo verdaderas o falsas. Por una vía distinta a la de Frege, llegamos aquí a un resultado muy parecido. Frege había dicho, extrañamente, que las oraciones son nombres de su valor de verdad. Carnap dice que tienen como extensión su valor de verdad. Y su justificación es más convincente, recurriendo, como acabamos de ver, al concepto de equivalencia. En cuanto a la intensión, habremos de preguntarnos en qué coinciden dos oraciones que sean *L*-equivalentes, esto es, equivalentes por razón de las reglas semánticas únicamente. Y la respuesta es: en lo que significan, en la proposición que expresan o, como había dicho Frege, en su sentido. En resumen (secc. 6, 6-1 y ss.):

La extensión de una oración es su valor de verdad; su intensión es la proposición que expresa.

Según Carnap se encarga de aclarar, el término «proposición», en este uso, como el término «propiedad» anteriormente, no se refieren a entidades mentales subjetivas como ideas o pensamientos, sino a «algo objetivo que puede o no estar ejemplificado en la naturaleza» (secc. 6, p. 27; es el sentido objetivo en el que Frege utilizaba el término «pensamiento»). Para que una entidad de cierta clase sea una proposición debe cumplir estas dos condiciones: que a toda oración de nuestro lenguaje *L*, las reglas de éste le asignen exactamente una entidad de esa clase; y que a dos oraciones de *L* se les asigne la misma entidad si y sólo si dichas oraciones son *L*-equivalentes en *L* (pp. 31-32).

Con esto hemos obtenido, al menos, una idea de cómo procede el método de Carnap y de sus conceptos básicos. Innumerables conceptos y desarrollos deben quedar fuera inevitablemente, y entre ellos algunos muy relevantes para la teoría posterior del lenguaje. Por ejemplo, el concepto de estructura intensional, que sirve para suministrar una versión técnica del concepto de sinonimia. Cuando dos oraciones están construidas de la misma manera y a partir de designadores con la misma intensión en una y en otra, podemos decir que poseen la misma estructura intensional, o, lo que es lo mismo, que son intensionalmente isomórficas (en el habla cotidiana diríamos que son sinónimas; secc. 14 de *Meaning and Necessity*). Así, en el ejemplo de Carnap, las expresiones « $2 + 5$ » y «II suma V», en un lenguaje L en el que contáramos con estos diferentes signos, tienen la misma estructura intensional, pues están formadas del mismo modo a partir de signos con la misma intensión. Y nótese que la identidad de estructura intensional no es lo mismo que la equivalencia- L , pues cualquiera de esas expresiones es L -equivalente a «7», pero no es intensionalmente isomórfica con ella.

Los conceptos de extensión e intensión suministran un recurso más claro y menos problemático que los conceptos de referencia y sentido manejados a la manera de Frege, y, en conjunto, el método de Carnap contribuye decisivamente a liberarnos del persistente prejuicio de que significar es fundamental y básicamente nombrar. Dentro del análisis de Carnap hay que destacar su formulación rigurosa del concepto de verdad lógica, que constituye un progreso notable en comparación con las formulaciones de Wittgenstein en el *Tractatus*; y también hay que subrayar asimismo su nítida distinción entre verdad lógica y verdad empírica, distinción que quedaba difuminada bajo el concepto tarskiano de verdad. Sin embargo, en la sección 7.10, al hilo de distinguir entre la implicación lógica y la implicación analítica, distinguimos también entre verdad lógica y verdad analítica, y esta distinción no se halla recogida en la teoría carnapiana que hasta aquí hemos estudiado. Pero si una verdad lógica en nuestro lenguaje L es una oración verdadera por razón solamente de las reglas semánticas de L , entonces parece razonable distinguir entre las reglas que fijan el significado de las constantes lógicas y las reglas que estipulan el significado de las constantes descriptivas. Ambos tipos de reglas darán lugar a dos tipos de verdad- L , como Carnap la llama. Del primer tipo será una oración como:

(8) Rodríguez está despierto o no lo está

Esto es un ejemplo de lo que en 7.10 hemos considerado como una oración lógicamente verdadera. Del segundo tipo será, en cambio:

(9) Si Rodríguez está despierto, entonces no está dormido

Esto es lo que, en dicho lugar, hemos llamado una oración analíticamente verdadera, aunque, en rigor, el concepto de verdad analítica es el género que incluye ambos tipos de oraciones verdaderas, (8) y (9).

Haciéndose eco de esta clase de consideraciones, que habían sido formuladas por Quine, Carnap recogió la distinción en un trabajo posterior del modo siguiente («Meaning Postulates», 1952). Las reglas de ámbito ciertamente muestran que (8) vale en toda descripción de estado de nuestro lenguaje y, por tanto, que es *L*-verdadera, esto es, necesariamente verdadera. Pero esas reglas semánticas no son suficientes para mostrar que también (9) lo es, pues si en nuestro lenguaje tenemos los predicados «está dormido» y «está despierto», y la expresión individual «Rodríguez», las reglas de ámbito no pueden, por sí solas, hacer de (9) una oración necesariamente verdadera, o sea, una oración que valga en todas las descripciones de estado. Para esto hemos de añadir una nueva regla que constituye lo que llama Carnap un postulado de significado, ya que establece las relaciones que hay entre lo que significa «estar dormido» y lo que significa «estar despierto», y en virtud de las cuales podemos afirmar que nadie puede estar a la vez despierto y dormido. Si representamos «está despierto» como «*F*» y «está dormido» como «*G*», nuestro postulado de significado tendrá esta forma:

$$(10) \quad \Lambda x (Fx \rightarrow \neg Gx)$$

o alternativamente:

$$(11) \quad \Lambda x (Gx \rightarrow \neg Fx)$$

Nótese que estos postulados no establecen lo que significan «*F*» y «*G*». Tan sólo dicen que, dado lo que significan, lo cual habrá sido establecido por las reglas de designación, ninguno de los objetos del universo de nuestro discurso puede tener al tiempo ambas propiedades. Y esto es lo único imprescindible para determinar que (9) es necesariamente verdadera, que vale en toda descripción de estado de nuestro lenguaje, pues, en efecto, nuestro postulado de significado *L*-implica la oración (9).

Una vez que tenemos las reglas de designación, y que sabemos, por consiguiente, lo que significan los predicados de nuestro lenguaje, ¿cómo sabemos qué relaciones hay entre ellos y, por lo tanto, qué postulados de significado hemos de establecer? Si nuestro lenguaje objeto fuera el lenguaje cotidiano, tendríamos que intentar determinar, por los medios empíricos oportunos, cuáles son esas relaciones. Esto es, precisamente, lo que pretende la teoría semántica actual, como vimos brevemente acerca de Katz en la sección 4.4. En algunos casos, como el de los predicados «despierto» y «dormido», que hemos usado en el ejemplo, puede ser fácil; en otros, puede ser considerablemente más difícil. Si, por el contrario, nuestro lenguaje, como en el caso de Carnap, es un lenguaje artificial, entonces depende de nosotros decidir qué postulados de significado aceptamos, en función de los propósitos e intereses con los que hayamos construido nuestro lenguaje objeto.

Con esto hemos tocado una cuestión disputada: la aplicabilidad que tenga este método a los lenguajes naturales. Sobre esta cuestión, que habremos de tratar en las próximas secciones, se ha pronunciado Carnap en otro trabajo que viene a completar de modo importante lo que estamos estudiando («Significado y sinonimia en los lenguajes naturales», 1955). Supongamos un lingüista enfrentado a un lenguaje que no entiende ni conoce. ¿Podría determinar por algún medio la extensión y la intensión de las expresiones de ese lenguaje? Parece razonable pensar que, observando el comportamiento de los hablantes, comenzará por averiguar a qué objetos aplican éstos los diferentes predicados (asumiendo, para simplificar, que se trata de predicados de observación). De esta manera, fijará la extensión parcial, dentro de unos ciertos límites espaciotemporales, de cada predicado que tome en consideración. Del mismo modo fijará la clase de los objetos a los que los hablantes rehúsan aplicar dichos predicados; e incluso, si es el caso, fijará también la clase de los objetos a los que los hablantes dudan si aplicarlos o no. Sobre esta base, nuestro lingüista puede formular hipótesis acerca de la aplicación de esos predicados a objetos no considerados, pero relevantemente semejantes a los conocidos (y puede haber, por supuesto, las dudas y errores propios de todo procedimiento inductivo). ¿Cómo se podría ahora pasar a fijar la intensión de esos predicados?

Para Carnap, la determinación de la intensión es, exactamente igual que la fijación de la extensión, objeto de una hipótesis empírica contrastable por el recurso al comportamiento lingüístico de los hablantes. La peculiaridad de la intensión consiste en que, para determinarla, debemos tener en cuenta no sólo los casos reales, sino también todos los casos lógicamente posibles. La razón es ésta: para determinar la intensión de un predicado nos interesa averiguar qué variaciones podría sufrir un objeto determinado en sus propiedades sin que ello nos impida atribuirle el predicado en cuestión. Ello nos obliga tomar en cuenta no sólo los objetos reales a los que se aplica el predicado, sino también aquellos objetos posibles a los que habría que aplicarlo, entendiendo por posibles los «lógicamente posibles», esto es, todo cuanto podamos imaginar, incluso contrario a las leyes actuales de la naturaleza. Esto es necesario, además, por cuanto muchas expresiones del lenguaje ordinario carecen de extensión, ya que no se aplican a ningún objeto, como, por ejemplo, «unicornio», «centauro», «duende», etc., y puesto que tienen significado, y las distinguimos entre sí semánticamente, hay que concluir que tienen una intensión, para determinar la cual de nada nos serviría limitarnos a observar los casos de su aplicación. Nótese, pues, que la determinación de la intensión de las expresiones no presupone la existencia de objetos determinados.

La intensión de un predicado es, en definitiva, la condición general que tiene que cumplir un objeto para que le apliquemos tal predicado. Esto significa que se prueba cuál sea la intensión de un predicado para un hablante por la disposición de éste a aplicarlo a aquellos objetos que poseen determinada característica o propiedad; el concepto de intensión es un

concepto disposicional. En términos de él podemos, ahora, definir dos expresiones como sinónimas si y sólo si poseen la misma intensión, y podemos decir que una oración es analítica si y solo si su intensión comprende todos los casos posibles (*op. cit.*, secc. 5). Así, pues, los conceptos de extensión e intensión no presentan al parecer ninguna dificultad en su empleo para la descripción de los lenguajes naturales.

La aceptación de entidades intensionales, como las propiedades, las proposiciones y los conceptos individuales, ha sido debatida y rechazada por muchos; en particular, y en los términos que estudiaremos en las secciones próximas, por Quine. Carnap, quien, como hemos visto, hace de ellas un recurso indispensable en el análisis semántico, las ha justificado en general del modo que vamos a considerar, defendiendo su perfecta compatibilidad con una actitud empirista («Empirismo, semántica y ontología», 1950). La acusación usual a la que responde aquí Carnap es la de platonismo: ¿conduce a una ontología platónica la admisión de proposiciones, propiedades y conceptos, además de objetos físicos? ¿Dónde está el límite que hemos de poner a los tipos de entidades que podemos aceptar en nuestra ontología? El procedimiento de Carnap para tratar esta cuestión recuerda en extremo al método de traducir del modo material al modo formal.

Cuandoquiera que deseamos hablar en nuestro lenguaje sobre un nuevo tipo de entidades tenemos que introducir en él las expresiones adecuadas; dicho en términos de Carnap, tenemos que construir un marco lingüístico adecuado para hablar de tales entidades (*op. cit.*, secc. 2). Esto nos permite distinguir dos órdenes de cuestiones. De un lado, cuestiones que se refieren a la existencia o no de ciertas entidades del nuevo tipo: son cuestiones internas a nuestro nuevo marco lingüístico. De otra parte, cuestiones relativas a la existencia o no del nuevo tipo de entidades como un todo: son cuestiones externas al nuevo marco lingüístico. Las primeras, las cuestiones internas, se plantean y resuelven usando las nuevas expresiones introducidas, esto es, sirviéndose del marco lingüístico creado al efecto. Así, aceptado un marco lingüístico que trata sobre objetos físicos, como es buena parte del lenguaje ordinario, podemos hacernos preguntas como: «¿Han existido realmente los unicornios?», «¿Existió don Quijote?», «¿Hay seres vivos en Marte?»... Todo esto constituye preguntas internas respecto a ese marco. A diferencia de éstas, una cuestión externa sería una cuestión acerca de la existencia de los objetos físicos en general. Y ésta es, para Carnap, una cuestión típicamente filosófica, una cuestión que ni el hombre de la calle ni el científico se plantean. Porque la pregunta «¿Existen realmente los objetos físicos?» es una pregunta que no puede hacerse dentro del marco lingüístico propio de los objetos físicos, pues en realidad cuestiona el propio marco. Es una pregunta mal planteada. En realidad, equivale a la cuestión siguiente: «¿Vamos a aceptar un lenguaje que hable de objetos físicos?» Es, más bien, una cuestión sobre el marco lingüístico que hemos de elegir. En palabras de Carnap: «Aceptar el mundo de las cosas no significa sino aceptar una cierta forma de lenguaje, o, en

otras palabras, aceptar reglas para formar enunciados y para comprobarlos, aceptarlos o rechazarlos» (*loc. cit.*, p. 208). Esto no quiere decir que la decisión de usar o no ese lenguaje sea arbitraria. La forma científica de nuestro trato epistemológico con la realidad, a su vez justificada por la eficacia en el dominio de la misma, acaso ofrezca la razón suficiente de que ese lenguaje de objetos físicos se haya incorporado al lenguaje cotidiano de modo casi universal e incuestionado. Pero esto —según Carnap— no es prueba de que existan tales objetos, sino tan sólo de que nos conviene hablar de ellos.

Lo propio puede decirse de otros marcos lingüísticos. La filosófica pregunta: «¿Existen los números?», tiene sentido si se entiende como: «¿Vamos a aceptar un lenguaje con términos de números?» Esto es, se trata de una cuestión externa al marco lingüístico numérico, a diferencia de una pregunta propiamente interna, como: «¿Hay algún número primo mayor que 100?», la cual presupone la utilización de ese marco lingüístico. Igualmente, la cuestión «¿Hay proposiciones necesarias?» es una cuestión interna a un lenguaje en el que se hable de proposiciones; pero la pregunta «¿Hay proposiciones?» es de carácter externo, y debe entenderse más bien como «¿Conviene emplear un lenguaje en el que se hable de proposiciones?» Otro tanto habría que decir sobre cuestiones como «¿Son reales las propiedades?», «¿Existen el espacio y el tiempo?», etc. La idea de Carnap es que las cuestiones externas, formuladas como cuestiones ontológicas y no como cuestiones metalingüísticas, son pseudocuestiones metafísicas, y como tales carentes de significado cognitivo. Como puede apreciarse, reaparece en este contexto semántico la clásica tesis carnapiana que ya conocemos sobre los problemas filosóficos: o son problemas acerca del lenguaje o no son problemas.

Desde este punto de vista, la introducción de nuevas entidades en el universo de nuestro discurso requiere dos movimientos fundamentales (*op. cit.*, secc. 3). En primer lugar, la introducción de un término general para la clase de estas entidades, lo que nos permitirá decir si una entidad determinada es o no de esa clase; por ejemplo, el término «propiedad» si se trata de introducir propiedades, el término «número» si queremos introducir entidades numéricas, etc. En segundo lugar, la introducción de un nuevo tipo de variables, cuyos valores sean las nuevas entidades introducidas (Carnap reconoce aquí que ha sido Quine el primero en subrayar la conexión entre aceptar entidades determinadas y asignar ciertos valores a las variables del lenguaje; la tesis clásica de Quine, que consideraremos más adelante, se enuncia así: ser es ser el valor de una variable). Podríamos pensar que la introducción de nuevas entidades es algo que hay que decidir, por las razones apropiadas, con independencia del lenguaje, y que tan sólo cuando se ha tomado una decisión positiva tiene sentido introducir las expresiones que se refieran a esas entidades. La posición de Carnap es la opuesta: «la introducción de nuevas expresiones no requiere justificación alguna porque no implica ninguna afirmación de realidad» (*loc. cit.*, p. 214). Con ello, Carnap aspira a eludir la acusación

de platonismo antes mencionada, pues si aceptar ciertos tipos de expresión no supone comprometerse con la realidad o con la existencia de ciertas entidades, entonces la semántica carnapiana no tiene por qué implicar la amplia ontología que parece implicar.

Nótese que Carnap se ve obligado, para ello, a un difícil movimiento. Por una parte, niega que, al aceptar un marco lingüístico, uno se vea comprometido a afirmar la existencia o realidad de ciertas entidades; pero, por otra, tiene que mantener que esas entidades son los valores de las nuevas variables, o, dicho de otra forma, que son los *designata* de las expresiones pertenecientes al nuevo marco lingüístico. Es decir, separa la cuestión semántica y la cuestión ontológica como independientes: un marco lingüístico nos compromete con la aceptación de ciertas entidades como *designata* de ciertas expresiones, y esta es una cuestión semántica; pero no nos compromete a aceptar que tales entidades existan o sean reales, lo cual constituye una cuestión ontológica. Y téngase en cuenta, para no malentender a Carnap, que decir que tal entidad es designada por la expresión *X* no es decir que *X* sea el nombre de la entidad; es, simplemente, decir que la entidad forma parte de la extensión de *X* en virtud de la intensión de ésta. El movimiento de Carnap es, pues, virtualmente idéntico al que vimos en la sección 8.2. Tienen sentido las afirmaciones:

- (12) Existen números primos mayores que 10
- (13) Existen seres vivos monocelulares

pero no, en cambio:

- (14) Los números existen
- (15) Existen seres

Sin embargo, parece que de (12) se deduce rigurosamente (14), y de (13), (15). Si existen números primos, entonces existen números; si hay seres vivos monocelulares, entonces hay seres. ¿Cómo puede justificarse que (12) y (13) sean aceptables, pero no (14) y (15)? Todo lo más, podremos reconocer que las dos primeras afirmaciones son interesantes y que, siendo verdaderas, acrecientan nuestro conocimiento; mientras que las dos últimas son triviales y no nos enseñan nada. La distinción radical entre semántica y ontología, tal y como Carnap la propone, no parece fácil de mantener. Por otra parte, parece una distinción relativa. ¿Cómo caracterizar los límites de un marco lingüístico y separarlo de otro? Si aceptamos un lenguaje que hable de objetos físicos, la pregunta sobre si existen los unicornios será interna, y la que versa sobre si, en general, hay objetos físicos será una cuestión externa, y como tal espuria. Ahora bien, supongamos que lo que aceptamos es un marco lingüístico en el que podemos hablar indistintamente sobre objetos físicos, sobrenaturales, imaginarios, presentidos, etc. Aquí, la pregunta sobre si existen objetos físicos es interna y, por consiguiente, perfectamente legítima. En la medida en

que nuestras expresiones se apliquen igualmente a objetos físicos y a objetos imaginarios, podremos continuar usando nuestro lenguaje aunque mantengamos que no existen objetos físicos. Considérese este otro caso. Con relación a un lenguaje de objetos físicos, en general, la cuestión de si hay centauros es interna. Pero si alguien quisiera aceptar la existencia de centauros y evitar una penosa discusión al respecto, podría limitarse a decir que él aceptaba un lenguaje en el que puede hablarse de centauros, y que cualquier intento de argumentar en contra implica hacer pseudo-afirmaciones metafísicas. Naturalmente, este sujeto tendría que justificar las ventajas de aceptar ese lenguaje. ¿Pero podemos asegurar, por anticipado, que no las tuviera? Podrían ser ventajas de orden estético, o psicológico... Ventajas que podría señalar en abundancia, por ejemplo, el defensor de un lenguaje de entidades sobrenaturales.

Pienso que la separación de Carnap entre semántica y ontología es, además de difícilmente sostenible por sí, innecesaria para justificar las categorías de su análisis semántico. Porque lo que apoya que aceptemos o no un cierto marco lingüístico es, en gran parte, justamente el tipo de entidades con las que nos comprometemos, y no vale pretender que uno es el compromiso semántico, en el que las aceptamos como *designata*, y otro el compromiso ontológico, en el que afirmamos que existen. Porque si las aceptamos como *designata*, entonces, en algún sentido, existen. Lo que podremos hacer es distinguir diferentes sentidos, modos o formas de la existencia. Afirmar que existen los números no es particularmente llamativo si uno se apresura a añadir que existen en cuanto construcciones matemáticas que obedecen a unas reglas, y que nos permiten unas determinadas operaciones con los objetos físicos. La existencia de éstos poco tendrá que ver, entonces, con la de aquéllos. Por lo mismo, afirmar que existen las propiedades, por ejemplo, el color verde, no implica afirmar que existen en el mismo sentido y de la misma manera que los objetos verdes. Mantener que hay proposiciones, no significa defender que su existencia sea como la de los hechos. Etcétera. Esto tampoco supone que haya que aceptar cualquier tipo de entidades. Es el éxito en nuestro trato con la realidad el que legitimará qué clases de entidades vamos a aceptar, o, lo que es lo mismo, qué teorías vamos a emplear. Y en virtud de este criterio podremos, acaso, aceptar objetos físicos, propiedades, proposiciones y números, y rechazar, en cambio, círculos cuadrados, esencias inmutables, demonios y dioses. Pero en el orden de la justificación racional (no en el orden de la explicación psicogenética) únicamente tiene sentido que aceptemos un lenguaje cuando hemos decidido qué entidades nos parecen admisibles.

8.5 La teoría conductista del significado en Quine y la crítica al concepto de analiticidad

Al año siguiente de la publicación del trabajo de Carnap que acabamos de estudiar, Quine publicó uno de los más conocidos y celebrados ensayos

filosóficos de estos años: «Dos dogmas del empirismo» (1951). Contenía una crítica muy severa de algunos puntos fundamentales de la teoría carnapiana del lenguaje que hemos revisado en las secciones precedentes. En especial, se criticaban dos aspectos: la distinción radical entre verdades analíticas y verdades fácticas o sintéticas, y la reducción de todo enunciado cognitivamente significativo a una construcción lógica de elementos simples directamente conectados con la experiencia inmediata.

La crítica a este último punto iba dirigida en particular contra la teoría verificacionista del significado. Lo menos que puede decirse es que, por esta época, hacer tal crítica era un poco como alancear muertos. Ya hemos visto, en la sección 8.1, que al comienzo de los años cincuenta tanto Carnap como Hempel y Ayer eran plenamente conscientes de las dificultades del principio de verificabilidad y del problema de relacionar adecuadamente los términos observacionales y los términos teóricos en una teoría científica. Curiosamente, Quine dedica la mayor parte de su crítica, no a estos problemas, sino a las insuficiencias del fenomenalismo que Carnap había mantenido en su primera época (¡antes de 1929!). Lo más sugestivo en la crítica de Quine es su tesis, que posteriormente se convertiría en una de las más características, de que nuestros enunciados sobre el mundo externo no se confirman individual y separadamente, sino en conjunto, formando un todo, una teoría (*op. cit.*, secc. 5). Es la tesis que se ha denominado «holismo» (del griego ὅλος, «todo», «entero»), y que en la sección 8.1 hemos visto recogida por Hempel.

Más interés tenía la crítica al otro aspecto de la posición empirista. Quine distingue, dentro de las verdades analíticas, las dos clases que ya mencionamos al paso en la sección 7.10. De una parte, las verdades propiamente lógicas, esto es, aquellos enunciados que son verdaderos en razón de los términos lógicos que intervienen en ellos. Es característico de estos enunciados que continúan siendo verdaderos comoquiera que se sustituyan sus términos no lógicos, siempre que esta sustitución sea uniforme. Podemos decir que son verdaderos por razón de su estructura o forma lógica (*op. cit.*, secc. 1; *Lógica matemática*, introducción; *Filosofía de la lógica*, cap. 4). Por ejemplo, el enunciado:

- (1) Nadie que no esté casado está casado

sigue siendo verdadero de cualquier forma en que se sustituya el predicado «estar casado»; por ejemplo, en:

- (2) Nadie que no sea escritor es escritor
(3) Nadie que no sea famoso es famoso

e incluso:

- (4) Nada que no sea un ser vivo es un ser vivo
(5) Nada que no sea fácilmente comprensible es fácilmente comprensible

etcétera (el cambio de «nadie» por «nada» simplemente indica que cambia el tipo semántico de predicado, pero no afecta a la forma lógica de estos enunciados). En suma, la estructura o forma lógica de estos enunciados es:

- (6) Para toda entidad, si ésta no tiene una cierta propiedad F , entonces no tiene dicha propiedad F

En símbolos:

$$(7) \quad \Lambda x (\neg Fx \rightarrow \neg Fx)$$

Se comprenderá que para distinguir, de esta manera, entre verdades lógicas y verdades que no lo son, tenemos que haber determinado previamente cuáles son los términos, elementos o constantes lógicas. Afortunadamente, sobre este punto, los lógicos suelen estar de acuerdo, al menos en los aspectos más básicos. Como muestra (7), los elementos lógicos en los ejemplos anteriores son los que podemos representar por medio de las expresiones «todo», «no» y «si... entonces».

Pero hay otro tipo de verdades analíticas: aquellas que dependen de lo que significan los términos no lógicos y, en particular, los predicados. Son verdades que, según Quine, pueden convertirse en verdades lógicas recurriendo a los términos sinónimos apropiados. Así, el enunciado:

- (8) Nadie que sea soltero está casado

es verdadero en virtud, aparentemente, de lo que significan los predicados «ser soltero» y «estar casado», y no sólo en virtud de su forma lógica (aunque, naturalmente, también ésta es relevante). Sustituyendo la expresión «ser soltero» por la expresión sinónima «no estar casado», obtenemos la verdad lógica (1). Nótese que, en rigor, ambas expresiones no son sinónimas, pues «no estar casado» incluiría también a los viudos y a los divorciados; podríamos decir que «no estar casado» es *parte* del significado de «ser soltero», y esto basta para que sea posible la conversión de (8) en (1). Precisamente algo semejante había dicho Kant en este respecto, a saber: que un juicio es analítico cuando el predicado está contenido en el concepto del sujeto como parte de él (*Crítica de la razón pura*, introducción, secc. IV); dicho en términos actuales: cuando el significado del predicado es parte de lo que significa el sujeto. Y esto es lo que parece ocurrir en (8).

He dicho «parece» porque aquí es justamente donde Quine sitúa el centro de sus dificultades. Aceptar este tipo de verdades analíticas requiere hablar de los significados de los predicados, y admitir que lo que uno significa puede estar contenido en lo que significa otro, o coincidir con ello. En suma, presuponemos la noción de sinonimia, y con ella, la de significado en el sentido intensional que ya conocemos. ¿Cómo podríamos justificar que dos expresiones significan lo mismo, o que lo que una significa es

parte de lo que significa la otra? Alguien podría decir que esto se justifica por definición, pero Quine señala, con toda razón, que las definiciones de las palabras del lenguaje cotidiano las hacen los lexicógrafos con la pretensión de recoger el uso que realmente se hace de esas palabras en la comunidad lingüística y, por consiguiente, que, a la postre, no es la definición la que funda la sinonimia, sino al revés. La sinonimia estaría fundada, más bien, en la conducta lingüística de los hablantes («Dos dogmas del empirismo», secc. 2). Una segunda posibilidad es identificar la sinonimia con la sustituibilidad *salva veritate*. Dos expresiones serían sinónimas, según esto, siempre que fueran sustituibles la una por la otra sin que variara el valor de verdad de la oración en la que se realizara la sustitución. Así, tenemos que cualquier oración verdadera (o falsa) con la que digamos algo acerca de los solteros, seguirá siendo verdadera (o falsa) aunque sustituyamos el término «solteros» por la expresión «quienes no han estado casados». Ahora bien, aquí la queja de Quine, asimismo razonable, es que la sustituibilidad *salva veritate* tan sólo nos asegura que ambas expresiones se aplican a los mismos objetos, esto es, que son extensionalmente equivalentes, pero no que signifiquen lo mismo, que sean intensionalmente equivalentes, en suma, que expresen la misma propiedad (*op. cit.*, secc. 3). Hay, sin duda, infinidad de casos en los que expresiones que claramente tienen distinto significado se aplican a los mismos objetos, como «animal con corazón» y «animal con riñones», y no por eso pretenderíamos que la oración:

(9) Todo animal que tiene corazón, tiene riñones

es analítica, pues parece claro que más bien se trata de una verdad empírica, fáctica.

¿Qué decir de la propuesta de Carnap de recurrir a postulados de significado para determinar la clase de estos enunciados analíticos? La primera objeción es que el método de Carnap está restringido a un lenguaje artificial *L*, y que aun cuando determinemos la clase de los enunciados analíticos en ese lenguaje *L*, con ello no hemos explicado lo que significa la expresión «enunciado analítico»; sabemos cuáles son los enunciados analíticos en *L*, pero no sabemos en qué consiste su analiticidad (*op. cit.*, secc. 4). A esto podría responderse que sabemos que un enunciado analítico es un enunciado que vale en toda descripción de estado de *L* porque está *L-implicado* por los postulados de significado de *L*, lo que sin duda suministra una explicación para el concepto de «enunciado analítico en *L*». Pero ahora viene la segunda objeción. ¿Cuál es el significado de la expresión «postulado de significado» o «regla semántica»? Para Quine, estas expresiones son tan oscuras como la de «enunciado analítico», y en última instancia tendremos que un postulado de significado se define como aquella regla que determina qué enunciados son analíticos, con lo cual no habremos aclarado nada. De aquí la conclusión de Quine: «Las reglas semánticas que determinan los enunciados analíticos de un lenguaje artificial

poseen interés tan sólo en la medida en que ya comprendamos la noción de analiticidad, pero no ayudan nada a conseguir tal comprensión» (*loc. cit.*, p. 36). Tendríamos primero que saber cuáles son las características mentales o de comportamiento relevantes para la analiticidad, y únicamente después serían útiles los postulados de significado. Mientras continuemos ignorando qué factores determinan la analiticidad, la distinción radical entre verdades analíticas y verdades sintéticas seguirá siendo un «dogma metafísico» (p. 37).

El único lugar en el que, por lo que hemos visto, Quine estaría dispuesto a buscar esos factores, es en la conducta lingüística de los hablantes. Aquí se encontraría, al parecer, el fundamento de la sinonimia y, por tanto, de las verdades analíticas. De acuerdo con su aspiración a un pragmatismo riguroso, coherente y sistemático, Quine ha explorado esa vía posteriormente con cierto detenimiento en *Palabra y objeto* (1960), una de las obras más importantes de los últimos veinte años en filosofía del lenguaje. Todo el carácter del libro está recogido en las palabras que abren su prólogo: «El lenguaje es un arte social. Al adquirirlo hemos de depender enteramente de indicios intersubjetivos respecto a qué decir y a cuándo decirlo. De aquí que no esté justificado comparar los significados lingüísticos a menos que sea en términos de las disposiciones de las personas a responder abiertamente a estimulaciones socialmente observables.»

En esta teoría del significado, que podemos considerar como social y pragmática, tiene un puesto central el concepto de estimulación. Este concepto nos sirve, por lo pronto, para definir un concepto empírico de significado como es el concepto de significado estimulativo, según lo llama Quine (*Palabra y objeto*, secc. 8). El significado estimulativo afirmativo de una oración declarativa *O* para un hablante *H* será la clase de las estimulaciones que provocarían el asentimiento de *H* si se le preguntara «¿*O*?». El significado estimulativo negativo será, por su parte, la clase de las estimulaciones que provocarían el disentimiento del hablante en dichas condiciones. Ambos, el positivo y el negativo, componen el concepto de significado estimulativo. Las estimulaciones mencionadas son, naturalmente, los diferentes tipos de irradiación que excitan los sentidos externos; cada estimulación, por otra parte, está constituida por aquellas irradiaciones que alcancen hasta una determinada duración establecida de modo convencional según nos interese en cada caso. Como puede apreciarse, el significado estimulativo se basa en las disposiciones de los hablantes a asentir o disentir de las oraciones cuando se encuentran sometidos a estimulaciones diversas.

Lo anterior puede aplicarse al problema de la sinonimia, comprobando qué resultados proporciona, pero antes conviene distinguir varios tipos de oración a los que habremos de aludir. Cuando una oración exige asentimiento o disentimiento sólo si se pregunta tras una estimulación apropiada, se trata de una oración ocasional. Por ejemplo: «Me duele», «Tiene la cara sucia»... Por el contrario, una oración es permanente si el asentimiento a la misma, o su disentimiento, puede repetirse cada vez que se

pregunta aun cuando no haya en ese momento estimulación adecuada alguna. Por ejemplo: «Ha llegado el correo», «Han florecido los almendros»... Las oraciones permanentes tienden a convertirse en ocasionales cuando disminuye el intervalo temporal entre las posibles estimulaciones sucesivas, ya que la distinción entre ambos tipos de oración es, como el significado estimulativo, relativa a la duración que se tome como módulo temporal de la estimulación (*op. cit.*, secc. 9). Del carácter de la distinción entre oraciones ocasionales y permanentes se desprende que el significado estimulativo, puesto que depende de las estimulaciones, es más importante para las primeras que para las segundas.

En vista de que las oraciones ocasionales son las más inmediatamente ligadas a la estimulación, podríamos buscar aquí un fundamento empírico para la sinonimia, y pensar que dos oraciones ocasionales serán sinónimas siempre que posean el mismo significado estimulativo, esto es, cuando las estimulaciones que provocarían el asentimiento a una oración sean las mismas que provocarían el asentimiento a la otra, y cuando las estimulaciones que darían lugar al disentimiento de una de las oraciones sean, por su parte, las mismas que darían lugar al disentimiento de la otra. Pero aun en caso tan simple, Quine ve dificultades. Supóngase que se trata de traducir entre dos lenguajes que nunca han tenido relaciones entre sí: es lo que llama Quine «traducción radical» (secc. 7). Supongamos que una de esas lenguas es el castellano, siendo la otra una lengua desconocida propia de una comunidad indígena encontrada en el curso de una exploración. Un explorador hispanohablante intenta traducir y halla, como su primer caso práctico, el siguiente: cuando se divisa un conejo, los indígenas, mirándolo, dicen «Gavagai». Podríamos aventurar la hipótesis de que esta expresión es sinónima de «Conejo», en la medida en que las mismas estimulaciones que impulsarían al explorador a asentir a «¿Conejo?», harán asentir también a los nativos a «¿Gavagai?»; y en la medida en que las mismas estimulaciones que harían al explorador contestar negativamente a la pregunta, provocarían la réplica negativa de los nativos a la pregunta correspondiente (podemos asumir que la contestación afirmativa o negativa de los nativos se realice por medio de gestos cuyo significado esté suficientemente claro). Diríamos, entonces, que ambas expresiones, «Conejo» y «Gavagai», poseen el mismo significado estimulativo. La dificultad para Quine, sin embargo, está en que la disposición de los nativos a asentir a la pregunta «¿Gavagai?» puede obedecer, en ciertos casos, a una información adicional que se suma a la estimulación; por ejemplo: pueden contestar afirmativamente no porque hayan percibido claramente un conejo, sino porque sepan que aquí hay una madriguera de conejos, o porque hayan percibido determinados insectos que se encuentran siempre en los lugares donde hay conejos, etc. Podríamos mantener que estos aspectos no forman parte del significado de «Gavagai», y que solamente a la presencia de conejos se refiere esta expresión; pero a ello objetará Quine: ¿por qué decir esto y no, en cambio, que «Gavagai» significa también la presencia de aquellos elementos que acompañan siempre a los

conejos?; no hay prueba definitiva en favor de lo uno o de lo otro; no hay criterio evidente que permita separar la información adicional del significado estricto (secc. 9). Como puede apreciarse, esto constituye una aplicación de la tesis, ya defendida en «Dos dogmas del empirismo», de que la distinción entre cuestiones de hecho y cuestiones de significado es relativa y de grado.

Por consiguiente, el significado estimulativo no alcanza a cumplir las funciones que solemos atribuir al concepto filosófico de significado, y la identidad de significado estimulativo es algo que no se da ni siquiera en casos tan simples como el de la traducción de «Gavagai» por «Conejo». Pero no todas las expresiones son igualmente afectadas por la información adicional. Las que designan, por ejemplo, colores, lo son menos, según Quine (secc. 10). Así, una expresión tal como «Rojo», tomada como oración ocasional, esto es, afirmada acerca de un objeto que está a la vista, o acaba de estarlo, resulta menos afectada por informaciones adicionales que «Conejo»; por esta razón, en el caso de términos como los de colores —admite Quine— su significado estimulativo nos acerca considerablemente a lo que es sólo esperar del concepto de sinonimia. En otros casos, sin embargo, es precisamente la información previa la que nos permite asentir a, o disentir de, una oración ocasional; por ejemplo, respecto de la expresión «Soltero», pues es lo que sabemos de una persona, y no ninguna de sus características que podemos percibir cuando la vemos, lo que nos permite decir que es, o no, soltera. El significado estimulativo de «Soltero» es del todo insuficiente para determinar el asentimiento o el disentimiento en una ocasión determinada. Aquellas oraciones ocasionales cuyo significado estimulativo no varía por efecto de la información adicional, las considera Quine oraciones observacionales (*loc. cit.*). En ellas, lo que, al margen de la teoría de Quine, llamaríamos su significado, llega casi a coincidir con el significado estimulativo de Quine. Es el caso de las expresiones de colores, que acabamos de considerar. Justamente en las oraciones observacionales es donde la noción de significado estimulativo puede sustituir a la noción vaga, preteórica, de significado, añadiéndole el rigor que le falta. Estas oraciones de observación, así entendidas, coinciden, por su carácter, con las proposiciones básicas en filosofía de la ciencia, o con las proposiciones protocolares de los neopositivistas. Pero se distinguen de las proposiciones atómicas de Russell en que aquéllas versan sobre cosas, en el sentido ordinario, y no exclusivamente sobre datos sensibles como estas últimas.

Por ahora, el único caso que se aproxima un tanto a nuestro concepto vago y cotidiano de sinonimia es el caso de aquellas oraciones de observación que poseen el mismo significado estimulativo, como, por ejemplo, en «Rojo», «Colorado» y «Encarnado», tomando estas expresiones como oraciones. Sin embargo, si nos limitamos a un solo hablante, encontraremos que incluso las oraciones que no son observacionales pueden dar lugar a sinonimia, con tal que tengan para dicho hablante el mismo significado estimulativo, ya que aquí no aparecen las dificultades que, en el caso de

varios hablantes, derivan de la posibilidad de que éstos posean información adicional diferente unos de otros. Así, para un solo hablante, «Soltero» y «Varón que no ha estado casado» tendrán el mismo significado estimulativo, ya que asentirá a ambas expresiones exactamente ante las mismas estimulaciones, esto es, ante las mismas personas. No hace falta subrayar que esto puede no acontecer en el caso de dos hablantes, a saber: cuando las personas solteras que uno conoce sean distintas de las que conoce el otro, pues en este caso el significado estimulativo de «Soltero» será diferente para uno y otro. Esto puede contribuir a confirmar la idea, ya apuntada, de que, para expresiones de este tipo, el significado estimulativo queda muy lejos de lo que, tanto en el habla cotidiana como en teorías semánticas más confiadas que la de Quine, llamaríamos simplemente *el significado* de la expresión (por ejemplo, de «Soltero»). Es claro que la sinonimia interna al idiolecto de un hablante es relevante para el concepto de analiticidad, ya que es una sinonimia que, como en las expresiones «Soltero» y «Varón que no ha estado casado», se basa no en que ambas expresiones tengan el mismo significado, sino simplemente en el hecho de que la información adicional que posee el hablante tiene como consecuencia que ambas expresiones tengan para él idéntico significado estimulativo (secc. 11).

Hasta ahora hemos considerado oraciones preferentemente ocasionales, pues son las ligadas más directamente a nuestras estimulaciones, y también con preferencia formadas por una sola palabra o expresión. Naturalmente, la forma completa de esas oraciones sería «Eso es un conejo», «Esto es rojo», etc. Pero explicar el significado de estas oraciones a quien no conozca la lengua es considerablemente más complicado a causa de esos otros términos: «eso», «es», etc. Quine ha tomado el caso más simple por lo que respecta a la conexión entre el lenguaje y el comportamiento. La cuestión ahora es: ¿podríamos extender nuestro análisis a las mismas palabras o expresiones tomadas esta vez no como oraciones, sino como términos? Supongamos que «Conejo» y «Gavagai» son, como oraciones, estimulativamente sinónimas. Pues bien, ello no garantiza que, como términos, «conejo» y «gavagai» sean, no ya sinónimos, sino ni siquiera coextensivos, aplicables a los mismos objetos. La razón, según Quine (secc. 12), es que, en rigor, no sabemos de qué se predica exactamente «gavagai». Puede ocurrir que se predique de animales, como el término «conejo». Pero puede acontecer también que se predique de estadios temporales de conejos; o bien de todas y cada una de las partes de un conejo en cuanto unidas formando un todo. En todos estos casos, «gavagai» sería, como «conejo», un término general. Pero puede darse el caso, incluso, de que fuera un término singular, que nombrara, por ejemplo, aquella porción del mundo constituida por todos los conejos, y de la que cada conejo singular sería tan sólo una parte; o bien que nombrara un universal, la propiedad que hace de algo un conejo, la conejidad. En todos los casos mencionados, el significado estimulativo de la oración «Gavagai» podría coincidir con el de la oración «Conejo», a pesar de las diferencias que

habría entre el término «gavagai» y el término «conejo». Lo cual prueba que la diferencia entre objeto concreto (un conejo) y objeto abstracto (la conejidad), así como la diferencia entre término general («conejo») y término singular («conejidad»), son independientes del significado estimulativo. Pues, como dice Quine (*loc. cit.*, p. 52 del original y pp. 65-66 de la trad. cast.): «Señalemos un conejo y habremos señalado un estadio temporal de un conejo, las partes constitutivas de un conejo, la fusión de todos los conejos, y la manifestación de la conejidad.» La estimulación no basta para resolver entre estas alternativas. Habría que recurrir a cuestionar al nativo con preguntas más complicadas, como «¿Es éste el mismo gavagai que aquél?», «¿Hay aquí un gavagai o dos?», etc. Pero esto requiere un dominio de la lengua nativa muy superior al propio del caso que estamos considerando. La conclusión es que, mientras que las oraciones ocasionales y el significado estimulativo suministran un puente entre dos lenguas, la clasificación de los términos y la determinación de su referencia pertenecen al esquema conceptual propio de cada una y pueden divergir ampliamente de una a otra.

Tornemos ahora al caso de la sinonimia estimulativa para un solo hablante. ¿Nos garantiza la sinonimia estimulativa de las oraciones «Soltero» y «Varón que no ha estado casado» que también sean estimulativamente sinónimos los términos correspondientes, «soltero» y «varón que no ha estado casado»? Una vez más, la ventaja de esta situación es que no plantea las dificultades de la traducción radical. Para aceptar la sinonimia estimulativa de esas expresiones, tomadas como términos, basta que las correspondientes oraciones ocasionales sean estimulativamente sinónimas y que el hablante esté dispuesto a asentir a la oración permanente «Todo soltero es un varón que no ha estado casado, y viceversa». No se plantea aquí la cuestión de si «soltero» se refiere a personas, partes de personas, estadios de personas, etc. La cláusula estipula, en definitiva, que ambas expresiones se aplican a los mismos objetos, sean lo que fueren. De acuerdo con ello, la definición de sinonimia estimulativa para términos generales rezará así: los términos (o predicados) «*F*» y «*G*» son estimulativamente sinónimos para un hablante *H* en un tiempo *t* si y sólo si «*F*» y «*G*» tienen, en cuanto oraciones ocasionales, el mismo significado estimulativo para *H* en *t*, y *H* está dispuesto en *t* a asentir a la cláusula «Todo *F* es *G* y viceversa». Puesto que esta cláusula establece la condición de sinonimia estimulativa de «*F*» y «*G*», podemos considerarla como una oración estimulativamente analítica, y definir este concepto así: una oración es estimulativamente analítica para un hablante si éste está dispuesto a asentir a ella a continuación de toda estimulación (dentro del módulo temporal establecido). La definición de sinonimia estimulativa para términos generales puede reformularse para dos términos singulares «*a*» y «*b*», simplemente requiriendo que el hablante esté dispuesto a asentir a la cláusula «*a* es idéntico a *b*», la cual será asimismo una oración estimulativamente analítica. Así, la sinonimia estimulativa de «Azorín» y «José Martínez Ruiz», tomando estas expresiones como términos, derivará de la

analiticidad estimulativa de la oración «Azorín es (idéntico a) José Martínez Ruiz». Nótese que las definiciones anteriores únicamente son aplicables a lenguas a las que puedan traducirse expresiones como «todo», «es» e «idéntico a» (secc. 12 de *Palabra y objeto*).

Con esto, tenemos los conceptos de sinonimia estimulativa (para términos y para oraciones ocasionales) y de analiticidad estimulativa. Al parecer, es lo más cerca que, desde el punto de vista del comportamiento lingüístico, podemos llegar a los disputables conceptos de analiticidad y sinonimia. ¿Qué decir sobre aquéllos? En primer lugar, están restringidos al idiolecto de un hablante, lo que, en principio, los haría inutilizables para el análisis del lenguaje de una comunidad. Este defecto tiene, sin embargo, un sencillo arreglo que Quine permite. Podemos socializar el concepto de sinonimia estimulativa de oraciones ocasionales estipulando que dos oraciones ocasionales serán socialmente sinónimas (desde el punto de vista de la estimulación) siempre que sean estimulativamente sinónimas para casi todos los hablantes de la comunidad (secc. 11, fin; nótese la cautela de la formulación). De modo semejante, dos términos estimulativamente sinónimos lo serán socialmente siempre que lo sean para casi todos los hablantes (secc. 12). Y una oración permanente, que sea estimulativamente analítica, lo será socialmente siempre que lo sea para casi todos los hablantes (secc. 14). De esta manera, podemos aceptar que una oración es estimulativamente analítica para una comunidad lingüística siempre que lo sea para casi la totalidad de sus miembros (cómo pueda determinarse esto último es otra cuestión). Otro tanto diremos si se trata de decidir que dos oraciones, o dos términos, son estimulativamente sinónimos para la comunidad.

Pero en segundo lugar, y esto es más importante, el concepto de sinonimia estimulativa (individual o social) entre oraciones o términos lo único que nos proporciona son oraciones o términos que, *de hecho*, son coextensivos, esto es, aplicables a los mismos objetos o estimulaciones. No sabemos en qué medida dependa esto de los significados de las expresiones en cuestión. Y lo mismo para el concepto de analiticidad estimulativa. Sabemos que una oración estimulativamente analítica es una oración a la que estamos dispuestos a asentir a continuación de cualquier estimulación. Pero ignoramos por qué estamos así dispuestos: ¿simplemente por lo que significan nuestras palabras, o también a causa de lo que sabemos acerca de la realidad? De acuerdo con esos conceptos, tendríamos que tan sinónimas son las expresiones «soltero» y «varón que nunca ha estado casado» como «moneda dorada» y «moneda de una peseta» (se entiende, para monedas de curso legal en España durante los últimos decenios); y que tan analítica es la oración «Ningún soltero está casado» como «Toda moneda dorada de curso legal vale una peseta». Y, sin embargo, entre ambos casos hay una diferencia, diferencia que aspiramos a recoger cuando distinguimos entre enunciados analíticos y enunciados sintéticos. ¿En qué consiste? Según Quine (secc. 12), la diferencia podría derivar de cómo aprendemos unas y otras expresiones, pues aprendemos «soltero» por medio de aso-

ciaciones de palabras, pero aprendemos «moneda de peseta» conectando la expresión con determinados objetos. Es decir, las expresiones que tenemos a considerar sinónimas serían aquellas que aprendemos por su interconexión con otras expresiones, mientras que las que aceptamos como meramente coextensivas serían las que aprendemos por su asociación con estimulaciones.

Más recientemente, en *Las raíces de la referencia* (1973; secc. 21), Quine ha acotado, dentro de las oraciones estimulativamente analíticas, una subclase de las mismas que se aproxima, aún más, a las oraciones tradicionalmente aceptadas como analíticas. Se trata de aquellas oraciones que todos los hablantes de la comunidad aprenden a entender al mismo tiempo que aprenden que son verdaderas. Así, se aprende lo que significa la oración «Un perro es un animal» al tiempo que se aprende a asentir a ella (el ejemplo es de Quine). En estos casos, la analiticidad estimulativa está conectada con el proceso de aprendizaje del lenguaje de tal manera que entender la oración y aceptarla como verdadera son procesos simultáneos. Naturalmente, éste sería también el caso de «Ningún soltero está casado». Esta aproximación al concepto de analiticidad la hace depender de la uniformidad social en el aprendizaje de ciertas palabras o expresiones. Pero esto no supone que se pueda trazar una división radical entre las verdades analíticas y las verdades sintéticas. Más bien lo que hay es una serie gradual, en la que hay que situar, en primer lugar, aquellas oraciones que todos los hablantes aprenden al tiempo que conocen su verdad, y que son las que podemos llamar «analíticas». En segundo lugar, aquellas que una mayoría, aunque no todos, aprenden de esa manera. Luego, las que aprenden así algunos; después, las que tan sólo unos pocos llegan a aprender de semejante modo, y, finalmente, las que nadie ha aprendido en la forma dicha. En resumen: hay oraciones más o menos analíticas o, mirándolo desde el punto de vista inverso, más o menos sintéticas.

8.6 La indeterminación de la traducción radical

Ya hemos visto en qué consiste la traducción radical: en traducir entre dos lenguajes que no habían tenido previamente relaciones culturales recíprocas. Quine dedicará buena parte del capítulo segundo de *Palabra y objeto* a defender la tesis de que semejante traducción no puede estar del todo determinada por la conducta lingüística de los hablantes. Dicho de otra forma: puede haber sistemas distintos, incompatibles entre sí, de efectuar tal traducción, y compatibles en cambio con la totalidad de las disposiciones lingüísticas de los hablantes (secc. 7). Naturalmente, la divergencia de traducción entre un sistema y otro será tanto menor cuanto más inmediata se halle la oración traducida a la estimulación no verbal. Salta a la vista que el principio de la indeterminación de la traducción radical va a suministrar a Quine un importante punto de apoyo para su teoría del lenguaje.

Según vimos, Quine describe, como ejemplo más simple, el caso de que, al divisarse un conejo en las inmediaciones, el nativo, dirigiendo su mirada en esa dirección, dice «Gavagai». El traductor, que hemos supuesto de habla española, anota «Conejo». Y comprueba ulteriormente que en todas aquellas ocasiones en las que él asentiría si se le preguntara «¿Conejo?», los nativos asienten cuando él les pregunta «¿Gavagai?» (hemos supuesto asimismo que el asentimiento o la negación lo realizan por medio de gestos cuyo sentido está claro). Una primera hipótesis sería, entonces, la de que «Gavagai» y «Conejo» poseen el mismo significado estimulativo, son estimulativamente sinónimas, puesto que son expresiones a las que los respectivos hablantes asienten y de las que disienten tras las mismas estimulaciones. Y vimos también que aun tan simple hipótesis presenta dificultades, pues el asentimiento del nativo podría obedecer, no a la estimulación relevante, sino a información adicional que él tiene, pero que no conoce el visitante; por ejemplo, información sobre la presencia usual de conejos junto a determinadas plantas. Esto hará que, en tales ocasiones, el nativo asienta a «¿Gavagai?» cuando el visitante daría una respuesta negativa a «¿Conejo?». A pesar de estas diferencias de significado estimulativo entre ambas expresiones, sin embargo, y asumiendo que no se produzcan más que en una pequeña proporción de los casos, el visitante pensará razonablemente que su traducción está justificada y seguirá adelante con ella, poniéndola de nuevo a prueba cada vez que se presente una clara ocasión (secc. 9). Para ciertas expresiones, como las que designan colores, la influencia de informaciones adicionales puede ser aún menor que para expresiones como «Gavagai»; puede ser nimia. Como ya sabemos, Quine considera como oraciones observacionales aquellas oraciones ocasionales cuyo significado estimulativo no varía a consecuencia de la información adicional (secc. 10). Podemos afirmar, pues, que, aun con la incertidumbre usual en los procedimientos inductivos, las oraciones observacionales son traducibles y escapan en principio al principio de indeterminación.

Al pasar de las oraciones a los términos, la situación empeora. Incluso si aceptamos que «Gavagai» y «Conejo» son oraciones estimulativamente sinónimas, y en esta medida traducibles, ello no nos asegura que, tomadas esas expresiones como términos, sean no ya sinónimas, sino ni siquiera coextensivas. Según comprobamos en la sección 8.5, «gavagai» podría designar conejos, estadios temporales de conejos, la conejidad, etc., sin que ello afectara a su significado estimulativo al ser usada como oración ocasional. Por ello, la traducción de términos queda bajo el principio de indeterminación, a diferencia de lo que acontece con las oraciones observacionales.

Por lo que toca a las oraciones ocasionales no observacionales, la posición de Quine es que no pueden traducirse con seguridad, si bien el traductor podrá determinar que dos oraciones son estimulativamente sinónimas para los nativos cuando compruebe que las estimulaciones que los provocan a asentir a una provocan también el asentimiento a la otra, y que las

estimulaciones que los impulsan a disentir de una los impulsan asimismo a disentir de la otra. El traductor será igualmente capaz de identificar las oraciones que son estimulativamente analíticas para los nativos, puesto que puede comprobar que hay oraciones a las que los nativos asienten a continuación de cualquier estimulación, y otras a las que replican negativamente también tras cualquier estimulación (secc. 15).

Las únicas oraciones que, hasta ahora, se prestan directamente a la traducción son las oraciones observacionales. Pero hay una clase muy distinta de oraciones que, según Quine, escapa también al principio de indeterminación de la traducción: son las que consisten en funciones veritativas (secc. 13). Recurriendo al criterio pragmático que estamos utilizando, el del asentimiento y el disentiimiento (aunque Quine lo llama «semántico»), podemos comprobar si una oración compleja suficientemente corta del lenguaje nativo expresa una función veritativa. Así, si el resultado de añadir una cierta expresión a una oración es que el nativo deja de asentir para pasar a disentir, o viceversa, podemos concluir que estamos ante la negación. Si la consecuencia de añadir una expresión a dos oraciones suficientemente cortas es que el nativo asiente a la nueva oración compleja sólo si antes asentía a cada una de sus partes, disintiendo en otro caso, podemos pensar razonablemente que la expresión representa la conjunción. Si al añadir una expresión a dos oraciones encontramos que el nativo tan sólo disiente en el caso de que esté dispuesto a disentir por separado de las dos oraciones componentes, asintiendo en caso contrario, cabe confiar en que hemos hallado la expresión de la disyunción incluyente. Y de modo semejante para otras funciones veritativas.

Posteriormente, Quine ha señalado algunas insuficiencias en el método anterior. Por ejemplo, puede ocurrir que el nativo no esté dispuesto a disentir de la oración compleja (supuestamente disyuntiva) y, en cambio, disienta de cada una de sus partes por separado. Esto puede ocurrir, en general, con afirmaciones sobre acontecimientos futuros; así, el nativo puede asentir a la disyunción «Lloverá o se perderá la cosecha», sin estar dispuesto a asentir por separado a alguna de las oraciones «Lloverá» y «Se perderá la cosecha» («Existencia y cuantificación», pp. 103-104 del original). En cuanto a la conjunción, el nativo puede, en cierto casos, disentir de ella, sin disentir de ninguno de los componentes por separado (*Las raíces de la referencia*, secc. 20). A pesar de estos casos marginales, es patente que las funciones veritativas pueden traducirse con un alto grado de seguridad. La alusión de Quine a que las oraciones componentes sean cortas se debe a que, si son largas, el nativo puede confundirse más fácilmente haciendo inútil la comprobación del traductor; se trata, pues, de una mera cautela. Y no hace falta añadir que, al traducir las expresiones nativas de las funciones veritativas al lenguaje del traductor, es inexcusable tomar precauciones, ya que, como es bien sabido, palabras como «no», «y» u «o» no coinciden exactamente con las funciones veritativas mencionadas, por más que las representen con suficiente aproximación.

¿Puede extenderse el método anterior a otras partes de la lógica? Quine piensa que no. Únicamente las relaciones veritativofuncionales pueden reconocerse en una lengua ajena recurriendo al comportamiento lingüístico de sus hablantes. En la lógica de predicados nos encontramos con oraciones cuantificadas del tipo de «Todo F es G » o «Algún F es G », y con afirmaciones de identidad como « $a = b$ ». Pero para traducir este tipo de oraciones se requiere determinar a qué objetos se aplican los predicados « F » y « G », o a qué se refieren los nombres « a » y « b ». Y esto es algo que no está determinado por el significado estimulativo de las expresiones. En tal medida, la cuantificación y la identidad son intraducibles; no constituyen patrimonio común a los hablantes de lenguas radicalmente diversas, sino que son parte de los mecanismos de referencia internos a cada una.

Aquí acaba todo cuanto nuestro intrépido traductor puede averiguar acerca del lenguaje nativo sobre la base del comportamiento lingüístico de los hablantes, y en particular provocando su asentimiento o disentimiento de las expresiones tras diferentes tipos de estimulación. ¿Cómo puede ir más allá de estos límites? Cabe imaginar que el traductor irá dividiendo las expresiones que escuche en partes convenientemente breves, y que por su repetición pueda considerar como palabras, estableciendo, con carácter hipotético, las equivalencias entre éstas y las diversas palabras y expresiones de su lengua, del castellano. Esto es lo que constituye las hipótesis analíticas del traductor, como las llama Quine. Naturalmente, deben ser acordes con los resultados de su traducción directa. Esto es, deben estar conformes con la traducción de oraciones observacionales y de funciones veritativas, y deben ser tales que las oraciones identificadas como estimulativamente analíticas para los nativos, queden traducidas por oraciones que sean estimulativamente analíticas para los hablantes del castellano, y las oraciones que sean entre sí estimulativamente sinónimas para los primeros, resulten traducidas por oraciones que lo sean igualmente para los últimos. Las hipótesis analíticas se caracterizan porque van más lejos de cuanto puede obtenerse exclusivamente a partir de las disposiciones de los nativos para el comportamiento lingüístico, autorizando traducciones para las que no existe justificación concluyente alguna (secc. 15).

Podría pensarse que el traductor siempre tiene la posibilidad de tomar el camino más largo, a saber: hacerse bilingüe, aprendiendo la lengua nativa a la manera de un niño. Quine concede que esto le permitiría traducir todo tipo de oraciones ocasionales, y no sólo las de observación, ya que estará ahora en posesión de toda la información adicional necesaria. Pero ésta es la única ampliación que experimentará su capacidad de traducir. En todo lo demás permanecerá igual, tendrá que recurrir a hipótesis analíticas, con la sola diferencia de que, en la medida en que haya llegado a dominar el lenguaje nativo, podrá actuar él mismo como informador nativo y como traductor, y no necesitará andar cuestionando a los demás hablantes ni provocando su asentimiento o disentimiento. El punto importante es éste: por más que domine la lengua nativa no encontrará

criterios que le permitan traducirla con seguridad excepto dentro de los límites señalados. En realidad, el propio supuesto de su bilingüismo carece de verosimilitud, pues no podrá en ningún caso aprender la nueva lengua como un niño; en su aprendizaje estará continuamente apoyándose en hipótesis analíticas, aun inconscientemente.

Esos criterios rigurosos de traducción no se encuentran en las hipótesis analíticas, pues es posible formular diversos sistemas de hipótesis incompatibles entre sí e igualmente compatibles con la conducta y con las disposiciones lingüísticas de los hablantes nativos. Por ejemplo, «gavagai», en cuanto término, puede traducirse como «conejo» bajo un cierto sistema de hipótesis, y como «parte de conejo» según un sistema distinto. La diferencia entre ambos sistemas podría consistir meramente en que determinada expresión nativa, utilizada para preguntar a los nativos sobre este punto, se tradujera en un sistema como «es lo mismo que» y en el otro como «es parte de lo mismo que». La posición de Quine es que, en casos como éste, no hay prueba independiente alguna de que una de las traducciones sea más exacta que la otra (secc. 15). En realidad, la idea de que una de las traducciones sea más correcta que la otra no tiene buen sentido: «La cuestión no es que no podamos estar seguros de si la hipótesis analítica es correcta, sino que ni siquiera hay, al contrario de lo que ocurría en el caso de 'Gavagai', un asunto objetivo sobre el que decir algo correcto o incorrecto» (secc. 16, principio).

¿Adónde lleva todo esto? No se trata, por lo pronto, de pretender que toda traducción es imposible. Podemos traducir, y traducimos de hecho con suficiente seguridad, del inglés, del ruso o del chino al castellano. Lo que Quine subraya es que la traducción se funda, en estos casos, en el parentesco de las lenguas entre sí y en las relaciones culturales existentes entre sus hablantes. El problema se plantea cuando se trata de lenguas entre las que no ha habido relaciones previas: éste es el caso de la traducción radical. Aun aquí, ya hemos visto que, para las oraciones observacionales y las funciones veritativas, es posible justificar la traducción. ¿Y para las demás expresiones? Quine no pretende que, llegado el caso, no tenga interés intentar su traducción, ni que la traducción que se proponga no sea suficiente para todos los efectos prácticos. Pero en el grado en que tal traducción no puede estar objetivamente mejor justificada que otras, no cabe pretender que la traducción consista en equiparar expresiones de lenguas distintas que tengan el mismo significado. El concepto de identidad de significado pierde todo su sentido y, con ello, también los conceptos de sinonimia, traducción, analiticidad y proposición (entendida como lo que significa una oración). El principio de la indeterminación de la traducción radical va dirigido, en suma, contra el mentalismo en la teoría del significado, contra la tesis de que toda oración, junto con todas sus oraciones sinónimas y con todas sus traducciones posibles a otras lenguas, expresan una misma e inalterable idea, un mismo y común significado, que tendría en su mente cualquier hablante de cualquier lengua siempre que pronunciara cualquiera de las oraciones del conjunto así definido. Esto es,

rechaza la tesis de que el lenguaje da expresión material a proposiciones o significados que existen previa e independientemente, tesis que ha encontrado gran apoyo en algunas de las afirmaciones de Chomsky, y que puede verse formulada en el siguiente párrafo de Katz: «Más o menos, la comunicación lingüística consiste en la producción de un fenómeno acústico externo, públicamente observable, cuya estructura fonética y sintáctica codifica los pensamientos o ideas interiores y privados de un hablante, y en la descodificación de la estructura fonética y sintáctica, presente en ese fenómeno físico, que realizan otros hablantes en la forma de una experiencia interior y privada de los mismos pensamientos e ideas» (*Filosofía del lenguaje*, cap. 4, principio). Estos pensamientos o ideas, que según Katz pueden repetirse en las mentes de hablantes diversos, serán los que den a las expresiones su significado. Y contra esta teoría del significado aspira Quine a ponernos en guardia.

8.7 La regimentación lógica del lenguaje y el criterio de compromiso óntico

Como vimos en el capítulo 6, la moderna teoría del significado nació, con Frege, bajo el signo de la queja contra el lenguaje ordinario: quejas fundamentalmente acerca de la ambigüedad y de la vaguedad de muchas de sus palabras y expresiones, y precisamente de aquellas que poseen importancia lógica mayor; y quejas también sobre la falta de referencia y sobre la indeterminación del sentido de muchos nombres y predicados («Sobre sentido y referencia», p. 70). Para paliar esta deficiencia lógica del lenguaje natural, Frege concibió un lenguaje lógico al que llamó «conceptografía» (*Begriffsschrift*, título de la propia obra donde lo expone). Este intento lograría una primera madurez en el sistema lógico de Whitehead y Russell, *Principia Mathematica*, en el que Russell veía, como ya sabemos, la estructura de un lenguaje lógicamente perfecto con el que comparar, para su descrédito filosófico, el lenguaje natural. El *Tractatus* de Wittgenstein contiene, a estos efectos, un cambio de posición sutil, pero significativo: a pesar de la ambigüedad que en tan gran medida afecta al lenguaje común, éste se halla lógicamente bien estructurado y la forma lógica hay que buscarla directamente en él, si bien una notación simbólica, a la manera de Frege y Russell, es útil precisamente para evitar aquella ambigüedad (*Tractatus*, 3.325 y 5.5563). En su segunda etapa filosófica, Wittgenstein, según hemos visto, se apartó notablemente de estos caminos para investigar más bien las peculiaridades del lenguaje natural al margen de cualquier exigencia lógica. Algunos de sus herederos y continuadores destacarán la riqueza de funciones lingüísticas y de relaciones semánticas que la lógica no puede recoger. Ahora más bien es la lógica la que aparece como deficiente. Esto se aprecia con nitidez en pensadores como Strawson y Austin. Carnap, por su parte, había continuado en este tema la línea de Russell, cuya influencia se advierte tanto en su primera

época, la de la sintaxis lógica, como en su época semántica, aquí enriquecida con la extensión de los procedimientos formales a los conceptos semánticos que había iniciado Tarski, según estudiamos anteriormente. Aun asumiendo y defendiendo la posibilidad de aplicar sus categorías al lenguaje ordinario, Carnap solamente investigó determinados lenguajes artificiales o sistemas semánticos, de alcance reducido y contruidos a propósito para ejemplificar su teoría.

Quine se distingue de Carnap por volver la atención al lenguaje ordinario. Su reelaboración de conceptos semánticos como los de significado y sinonimia en términos de respuestas a estimulaciones está pensada para el lenguaje natural en su uso común, y también sobre él versa el principio de indeterminación de la traducción radical. Pero a diferencia de lo que es propio del segundo Wittgenstein y de su escuela, y prosiguiendo en esto la línea de Frege, Russell y Carnap, Quine va a defender la posibilidad y la conveniencia de recurrir a la lógica a fin de introducir en el lenguaje natural orden y concierto, o, como dice él, de «regimentarlo» (cap. 5 de *Palabra y objeto*; la versión castellana traduce como «regulación» el expresivo término original *regimentation*). ¿Cuál es el propósito de esta regimentación? Uno que afecta a la relación entre el lenguaje y la realidad y, por consiguiente, a nuestro conocimiento: clarificar nuestro esquema conceptual, entender cómo cumple el lenguaje su función referencial y, sobre todo, simplificar nuestra teoría del lenguaje hasta los mayores extremos, suministrando un conjunto de «formas canónicas» a las que *traducir* (esperamos que justificadamente) las variadas expresiones del lenguaje natural (*op. cit.*, secc. 3). Quine ha subrayado aquí algo digno de mención: puesto que la notación lógica se explica por medio del lenguaje y tiene en él traducción regular, si no totalmente exacta, parafrasear una expresión ordinaria en términos de la lógica es virtualmente lo mismo que dar su paráfrasis en otros términos del propio lenguaje natural. Tal paráfrasis no pretende, por supuesto, encontrar expresiones sinónimas; de las consideraciones de Quine sobre la sinonimia, que ya conocemos, se desprende con facilidad que este concepto es ajeno a su programa de regimentación lógica del lenguaje. Para lo que pretende conseguir, el programa no requiere el descubrimiento o la justificación de sinonimias. Y lo que pretende conseguir es la reducción de la gran variedad de expresiones lingüísticas a unas pocas formas básicas, lo que equivale a buscar las categorías básicas del pensamiento o los rasgos más generales de la realidad (secc. 33, fin). Por su aspiración, Quine entronca directamente con Russell y el primer Wittgenstein.

El resultado de esta regimentación, sin embargo, no es nada llamativo. Ni siquiera encontraremos las pequeñas sorpresas metafísicas que nos deparaba el análisis de Russell o el del *Tractatus*, lo que sin duda casa con la tesis de Quine de que no hay una filosofía primera previa a la ciencia («Reply to Chomsky», en *Words and Objections*, p. 303). Lo único que resulta es que un lenguaje natural como el inglés puede reducirse en su totalidad a los elementos sustanciales de la porción más básica de la lógica,

la lógica de predicados de primer orden. ¿De qué manera? Veámoslo muy en esbozo, prescindiendo de multitud de detalles, dificultades y desarrollos técnicos que Quine trata en su momento. Asumiremos, naturalmente, que el castellano es traducible al inglés, y en consecuencia aplicaremos al primero las conclusiones de Quine sobre el último.

En primer lugar, los términos singulares indefinidos: expresiones como «alguno», «cada uno», «ninguno», «cualquiera», etc. No hay que decir que todos ellos pueden sustituirse por expresiones que comiencen con «todo» o «algún», según los casos, en concurso además con la negación cuando proceda. Las expresiones que obtengamos pueden ser, desde luego, menos idiomáticas, menos elegantes o más largas que las originales, pero la regimentación lógica del lenguaje no pretende ciertamente servir a la estética ni al estilo literario ni a la retórica. Sirva esta aclaración para lo sucesivo. Tenemos así, las expresiones «todo» y «algún», que son las que traducen los operadores lógicos llamados «cuantificadores», base de la lógica de predicados. Todavía puede conseguirse una simplificación ulterior, pues, como es sabido, cualquiera de los cuantificadores puede definirse en términos del otro más la negación (el lector inexperto en lógica puede consultar sobre este punto en cualquier manual). A esos términos sólo hay que añadir una expresión como «tal que», ya que las oraciones que obtengamos serán del tipo de «Todo objeto es tal que...», y de «Algún objeto es tal que...» (secc. 34).

Consideremos, en segundo lugar, los términos singulares definidos, como las descripciones definidas, que ya estudiamos en Frege y Russell, y los demostrativos singulares. El atomismo lógico de Russell aceptaba como términos singulares de un lenguaje lógicamente perfecto, o nombres lógicamente propios, expresiones como los pronombres demostrativos del lenguaje natural, palabras del tipo de «esto», «eso», «aquello». Y una de las críticas que se le hicieron, como se recordará, fue que un demostrativo puede siempre sustituirse por una descripción como «el objeto al que estoy apuntando». Pues bien, tal es la posición de Quine; los demostrativos singulares quedan reducidos a descripciones acompañadas con gestos de ostensión: «Esta mesa» será parafraseada como «La mesa aquí delante» (apuntando hacia ella), etc. En cuanto a las descripciones definidas, su forma canónica será: «El objeto tal que...»; por ejemplo, para la expresión «El satélite natural de la Tierra», su forma canónica será «El objeto tal que es satélite natural de la Tierra». Tenemos además términos singulares de clase, términos como «la humanidad», que designa la clase de los seres humanos. Su forma canónica será «La clase de los objetos tales que son seres humanos»; que a su vez puede reducirse a una descripción definida de la manera siguiente: «El objeto tal que todo objeto es miembro de él si y sólo si es un ser humano» (secc. 34). La oscuridad idiomática que pueda aquejar a estas formulaciones en el lenguaje cotidiano desaparece, por supuesto, en el simbolismo lógico, pues aquí contamos con variables diversas para evitar la ambigüedad; así, la expresión que se acaba de mencionar, utilizando variables diría: «El x tal que, para todo y , y es miembro

de x si y sólo si y es un ser humano» (prescindiendo de la traducción lógica de las demás palabras).

Otra clase de términos singulares son los nombres propios. Ya a propósito de Russell tuvimos ocasión de considerar los peculiares problemas que plantean, en especial cuando se trata de nombres sin referencia y aparecen en ciertos contextos lingüísticos. Así, la proposición «Pegaso existe» es falsa, mientras que su negación «Pegaso no existe» es verdadera. Pero, asumiendo que nos referimos al caballo alado del que habla la mitología, hay que reconocer que el término «Pegaso» carece de referencia en la realidad, y si por consiguiente no estamos hablando de ningún objeto o entidad existente, ¿cómo es que podemos hacer afirmaciones verdaderas o falsas sobre él? La propuesta de Quine es considerar los nombres propios, en estos casos, como términos generales con función predicativa dentro de una oración cuantificada particularmente (secc. 37). Lo que esto quiere decir es simplemente que la oración:

(1) Pegaso existe

tiene como forma canónica:

(2) Algún objeto es Pegaso

o lo que es lo mismo:

(3) Para algún x , x es Pegaso

tomando el «es» como «es idéntico a». Traduciendo a símbolos todo ello, excepto el término «Pegaso», tendremos:

(4) $\forall x x = \text{Pegaso}$

Como se recordará, para Russell la eliminación de los nombres propios ordinarios tenía como primer paso su sustitución por descripciones definidas, que ulteriormente desaparecían en el contexto de oraciones cuantificadas (secc. 6.6). Ello requiere que podamos hacer del supuesto objeto o entidad del que hablamos una descripción que sustituya al nombre. El método de Quine evita este requisito, que en algunas ocasiones puede ser dificultoso de cumplir, y abrevia el proceso reduciendo los nombres propios directamente a predicados de identidad. La idea es que las oraciones a las que llegamos en este análisis son de la forma:

(5) $\forall x Fx$

donde el predicado « F » es de la forma « $= a$ »; así, en (4), el predicado es « $= \text{Pegaso}$ », y (4) es falsa justamente porque no hay ninguna entidad que satisfaga el predicado « $= \text{Pegaso}$ », o dicho en términos de Quine, porque el término general en que hemos convertido el nombre «Pegaso»

no es verdadero de ninguna entidad. Cuando una oración de este tipo es verdadera, es porque el término general correspondiente es verdadero de un objeto, pero como ese término procede de un nombre propio, ocurrirá que es verdadero exclusivamente de un sólo objeto. Es lo que acontece en el caso de la oración:

(6) $\forall x \ x = \text{Borges}$

Nótese que aquí, como en el análisis de Russell, y de acuerdo con la tesis kantiana, «existe» desaparece como predicado. El inconveniente principal es que se recurre a un operador tan oscuro como el de identidad; ya en el *Tractatus*, Wittgenstein escribió: «La identidad del objeto la expreso por la identidad del signo, y no por medio de un signo de identidad, y la diversidad de objetos por la diversidad de signos» (5.53). Pues ¿qué quiere decir que algo es idéntico a Borges o que nada es idéntico a Pegaso? Simplemente que el término singular «Borges» posee referencia en el mundo real, y que «Pegaso» carece de ella. Por tanto, se trataría, en última instancia, de afirmaciones metalingüísticas, disfrazadas como pseudoproposiciones de objeto. Tal vez esto rompa la armonía de la regimentación al modo de Quine, pero resulta más claro.

La propuesta de Quine es de aplicación aún más artificiosa en el caso de los términos singulares que aparecen en posiciones no referenciales, esto es, en posiciones que no presuponen que ellos posean referencia. Por ejemplo, la oración:

(7) Picasso está dibujando a Pegaso

se convertiría, según Quine, en:

(8) Picasso está haciendo un dibujo del que imagina que se parece a Pegaso

que parcialmente traducida a símbolos, tendría esta forma:

(9) $\forall y \ (\text{Picasso está haciendo } y \wedge \text{Picasso imagina } x \ [x \text{ se parece a Pegaso}] \text{ de } y)$

La cuestión es que, si el nombre «Pegaso» carece de referencia, entonces no sabemos en qué consiste que alguien se imagine de un dibujo que se parece a Pegaso; del retrato de Dalí hecho por Picasso podemos imaginar que se parezca a Dalí, en el sentido de que podemos imaginar que haya cierta correspondencia entre el retrato y el personaje real, lo que, mientras Dalí no muera, siempre podremos comprobar. ¿Pero cómo podemos imaginar una correspondencia entre un dibujo de Pegaso y Pegaso, si éste no existe ni ha existido realmente? La paráfrasis que Quine recomienda nos obliga a atribuir a los pintores de temas mitológicos imaginaciones un tanto extrañas. Su propuesta tiene, no obstante, sus ventajas: por una

parte, elimina la ambigüedad del término «es» entre tomarlo como cópula y tomarlo expresando identidad, pues en la notación canónica será siempre cópula entre sujeto y predicado, ya que no habrá diferencia lógica entre decir «Ese es un pintor» y «Ese es Picasso». De otra parte, elimina los vacíos de valor veritativo cuando se trata de nombres sin referencia; así, la afirmación «Pegaso vuela», a la que, en principio, podríamos negar un valor de verdad, queda regimentada como «Algún objeto es Pegaso y ese objeto vuela», que es patentemente falsa en virtud de la falsedad de su primera parte. (Esta eliminación de vacíos veritativos es lo que pretendía la teoría de las descripciones de Russell, una vez dado el paso previo de transformar los nombres propios en descripciones; por lo que respecta a la propuesta de Quine que acabamos de discutir, además de las seccs. 37 y 38 de *Palabra y objeto*, puede verse su *Lógica matemática*, secc. 27, y su *Filosofía de la lógica*, cap. 2).

Los nombres propios, pues, quedan analizados como términos generales en posición predicativa. Restan, por último, las descripciones definidas. Aquí la posición de Quine es virtualmente idéntica a la de Russell: las descripciones definidas o singulares desaparecen sustituidas por predicados y cuantificadores según proceda en cada caso. Así, «El caballo alado robado por Belerofonte vuela» será analizada como «Algún objeto es tal que solamente él es un caballo alado robado por Belerofonte, y ese objeto vuela». Puesto que, según vimos antes, los demostrativos y los nombres abstractos de clase se traducen a descripciones para su regimentación lógica, también ellos quedarán finalmente disueltos en predicados y cuantificadores de la manera que se acaba de indicar. Lo propio se aplica a los nombres de propiedades y relaciones («la blancura», «la paternidad», etc.), que pueden igualmente transformarse en descripciones de la forma «el objeto tal que...».

Al final, los únicos términos singulares conservados en la notación canónica son las variables, oscuramente representadas en nuestros anteriores ejemplos por las palabras «objeto», «entidad», «él», «ése», etc., lo que, según Quine, podría considerarse que da testimonio de «la primacía del pronombre» (secc. 38, fin). Además de las variables, tenemos términos generales en posición predicativa, que junto con aquéllas componen las oraciones atómicas, así: Fx , Fxy , $Fxyz$, etc. Las demás oraciones se construirán a partir de éstas por medio de funciones veritativas y cuantificadores, añadiendo acaso algún otro recurso artificial que pueda hacernos falta para fines específicos (como los operadores introducidos por Quine al final de la sección 38, y eliminados luego en la 44, que no tiene interés reproducir aquí).

Nuestro lenguaje canónico nos permite, por cierto, definir un concepto que parece sustituir con ventaja a conceptos rechazados antes por Quine como el concepto vago de sinonimia o el concepto carnapiiano de isomorfía intensional. El nuevo concepto que así se insinúa es el de sinonimia estructural, como lo llama Quine (secc. 42), y se define así: las oraciones formuladas en la notación canónica son sinónimas entre sí si y sólo si una

de ellas puede transformarse en otra por medio de transformaciones pertenecientes a la lógica de predicados y de las funciones veritativas, junto con la sustitución recíproca de términos generales que sean estimulativamente sinónimos. Esto significa que una oración de la forma de $\bigwedge x (Fx \rightarrow Gx)$ será estructuralmente sinónimo de otra de la forma $\neg \bigvee x (Hx \wedge \neg Ix)$ en el caso de que sean estimulativamente sinónimos los predicados «F» y «H», y de que lo sean asimismo los predicados «G» e «I». Esto nos permitiría tomar como sinónimas aquellas oraciones eternas de nuestra notación canónica que cumplieran con esos requisitos, y podríamos afirmar entonces, con justificación, que tales oraciones eternas expresan la misma proposición. Llama Quine «oración eterna» a toda oración permanente cuyo valor de verdad permanece fijo a lo largo del tiempo y para sucesivos hablantes (secc. 40). Suelen ser de esta clase las proposiciones teóricas en las ciencias, así como ciertas oraciones informativas cuando los datos de persona, tiempo y lugar involucrados en ellas están inequívocamente determinados y no varían según el contexto en el que la oración se usa.

Pues bien, ni siquiera la sinonimia estructural, así definida para oraciones eternas regimentadas, puede aceptarse como criterio que nos permita decir que dos oraciones expresan la misma proposición. Aparte de otras objeciones menores que el propio Quine examina, está la siguiente: nuestro concepto tan sólo es aplicable al lenguaje canónico que hemos adoptado, pero no a las oraciones eternas de lenguajes ajenos. En particular, depende de la sinonimia estimulativa de los términos generales, y ya vimos en su momento que la sinonimia estimulativa de términos (a diferencia de la de oraciones) es totalmente interna a una lengua y escapa a toda posibilidad de traducción radical. Esto significa que la sinonimia estructural que hemos definido no nos permite salir de nuestro lenguaje (ya regimentado), que no nos permite justificar la traducción de oraciones eternas, y por consiguiente, que no suministra tampoco justificación adecuada para admitir proposiciones como lo significado por oraciones (secc. 42). Pues, en definitiva, el concepto de proposición pretende expresar aquello que es común a una oración y a todas sus posibles traducciones a otros lenguajes. Y el caso es que nuestro concepto de sinonimia estructural no nos suministra criterio alguno que nos permita establecer tal sinonimia entre oraciones eternas de lenguajes diversos recurriendo a las disposiciones de los hablantes para el comportamiento verbal. La traducción de oraciones eternas está irremisiblemente afectada por el principio de indeterminación, y en consecuencia, el concepto de proposición condenado a desaparecer (secc. 43).

Con él, desaparecen también otros conceptos intensionales (del término «intensión», en el sentido que le da Carnap, y que ya conocemos) como los de propiedad y relación, como lo prueba el estrecho paralelismo que puede trazarse entre éstos y las proposiciones. Pues así como las proposiciones pueden considerarse como los significados de oraciones eternas cerradas (esto es, con todas sus variables ligadas por algún cuantificador), las propiedades y relaciones pueden tomarse como los significados de ora-

ciones eternas abiertas, esto es, oraciones eternas con variables libres. Y las mismas dificultades tendremos —según Quine— para decidir cuándo estamos ante la misma propiedad o la misma relación, que tenemos para resolver cuándo nos encontramos ante la misma proposición (secc. 43). La propuesta de Quine es aquí sustituir las propiedades por clases, ya que estas últimas sí poseen criterios controlables de identidad; en efecto, dos clases son idénticas cuando tienen los mismos miembros. Si la propiedad «rojo» es o no idéntica a la propiedad «encarnado» tiene, pues, una formulación alternativa preferible que es ésta: la de si a la clase de los objetos rojos pertenecen o no las mismas entidades que a la clase de los objetos encarnados. Y de modo parecido se argumentará respecto a las relaciones; sólo que éstas se sustituirán por clases de pares ordenados. Y en lugar de argüir sobre si la relación «tío carnal de» es o no idéntica a la relación «hermano del padre o de la madre de» (suponiendo que alguien quisiera argüir sobre ello), nos preguntaremos más bien si los mismos pares ordenados de personas están incluidos bajo una y otra relación, o sea, si cada par de personas tal que la primera es tío de la segunda es también tal que la primera es un hermano del padre o de la madre de la segunda. Como se ve, el carácter general de la propuesta de Quine es sustituir los conceptos intensionales por conceptos extensionales.

Tenemos, por consiguiente, que, desde el punto de vista lógico, el lenguaje puede reducirse a las siguientes construcciones básicas: la predicación, la cuantificación y las funciones veritativas. Queda por determinar, de un lado, una cuestión de vocabulario: cuál sea el léxico de términos generales que van a oficiar de predicados. Aquí entrarán, no sólo adjetivos y nombres comunes, sino también verbos, adverbios e incluso, como ya sabemos, nombres propios. De otra parte, resta la cuestión de conectar el lenguaje con la realidad, esto es, de determinar de qué objetos va a hablar nuestro lenguaje, o dicho de otro modo, de cuáles serán los valores que adquieran las variables de las que predicamos términos generales y que sometemos a cuantificación. Esto es propiamente una cuestión ontológica, y el modo de plantearla que es característico de Quine aparece ya tempranamente en una fórmula que se ha citado con frecuencia: *ser es ser el valor de una variable* («Sobre lo que hay», 1948). Como todas las fórmulas simples y sentenciosas, es fácil de malentender y se presta al abuso.

Por lo pronto, la fórmula no pretende expresar, pese a su apariencia, un criterio para determinar lo que hay, lo que existe, sino más bien lo que en un lenguaje determinado se dice que hay, o tanto da, lo que una teoría afirma que existe. Esto es, nos ofrece un criterio para determinar los compromisos ónticos de una teoría (*op. cit.*, pp. 15 y 19). La idea es que, para delucidar qué tipos de entidades acepta como existentes una teoría, el criterio más seguro es traducirla a nuestra notación canónica, pues ésta nos mostrará, a través de la cuantificación, cuanto deseamos saber al respecto (*Palabra y objeto*, secc. 49). Los objetos o entidades aceptados por una teoría o lenguaje serán precisamente aquellos que compongan el universo de valores que constituyan el dominio de las variables sometidas a cuanti-

ficación en la teoría o lenguaje en cuestión (la indiferencia entre hablar de «lenguaje» o de «teoría» obedece a la imposibilidad, que hemos visto defendida por Quine, de distinguir radicalmente entre las cuestiones de significado y cuestiones de hecho). Dicho de otra forma: los objetos reconocidos por una teoría serán aquellos de los que ésta haga afirmaciones del tipo de «Todo objeto es tal que...» o «Algún objeto es tal que...». En esta medida, la notación canónica suministra un medio en el que discutir con más claridad las cuestiones ontológicas, pues es un medio, en definitiva, objetivo y neutral; pero, naturalmente, un medio tan sólo accesible a quien esté dispuesto a aceptar la traducción de sus afirmaciones a la lógica de predicados de primer orden.

Con lo anterior solamente hemos tocado la cuestión de lo que una teoría dice que hay. Queda otra cuestión: ¿qué es lo que debemos reconocer como existente? El criterio referido no responde a esta pregunta, pues responderla equivale a fijar el universo de valores de nuestras variables; o lo que es lo mismo: equivale a elegir una determinada ontología. ¿Cuál elige Quine? Al final de «Sobre lo que hay» había dejado esta cuestión abierta, aconsejando «tolerancia y espíritu experimental» a la hora de comparar el fenomenalismo con el fisicalismo o la ciencia natural con la matemática platonizante. Si hemos de juzgar por sus alusiones en *Palabra y objeto* (secc. 48, por ejemplo), y por declaraciones posteriores (*Aspectos de la filosofía de W. V. Quine*, p. 153), su posición es fisicalista, pues defiende una ontología básicamente integrada por objetos físicos, a los que se reducen, por ejemplo, todas las entidades mentales (*Palabra y objeto*, secc. 54; *Las raíces de la referencia*, secc. 9). Pero esta calificación hay que matizarla, pues no significa que Quine excluya absolutamente cualquier otro tipo de objetos, como los objetos abstractos. Es cierto que de éstos rechaza las proposiciones, los significados, las propiedades y las relaciones; pero no rechaza, en cambio, las clases (*Palabra y objeto*, seccs. 48 y 55) ni los números (*op. cit.*, secc. 50), y por esta razón ha protestado enérgicamente contra las frecuentes acusaciones que se le hacen de nominalismo (secc. 49, nota 5), subrayando las dificultades que se interponen en el camino de todo programa radicalmente nominalista (secc. 55). Esta ontología es coherente con la reducción que efectúa Quine de las disposiciones a mecanismos físicos (*Las raíces de la referencia*, seccs. 3 y 4), con la conversión de la epistemología en una parte de la psicología empírica («La epistemología naturalizada») y, en suma, con su negativa a admitir la filosofía como algo previo a la ciencia o más firme que ésta.

Lo importante aquí es insistir en la relatividad de toda ontología. La ontología es relativa a un marco conceptual de referencia tal y como éste viene dado por el lenguaje que le sirve de trasfondo y desde el cual formulamos nuestras preguntas ontológicas («La relatividad ontológica»), del mismo modo que sólo tiene sentido preguntarse por la posición y la velocidad de un móvil con relación a un sistema de coordenadas. Por supuesto que las cuestiones de referencia acerca del lenguaje de trasfondo requerirán un lenguaje ulterior que suministre el marco de referencia, y así sucesiva-

mente. ¿No hay límite a este progreso? Teóricamente, no, pero en la práctica —según Quine— sí lo hay: ponemos término a ese regreso de una teoría a otra, de un lenguaje a otro, cuando llegamos al lenguaje natural, a nuestra lengua nativa. Ahí detenemos nuestro peregrinaje, pues éste es el marco en el que, en definitiva, todos coincidimos; se entiende: todos los que hablamos la misma lengua (*loc. cit.*, p. 49 del original).

Hablar de ontología es hablar de la referencia de las palabras, y por consiguiente exige tratar tanto de la realidad como del lenguaje. En esta vinculación entre la ontología y la semántica, de la que ya hemos encontrado ejemplos en las páginas anteriores, resuena claramente el eco de Carnap. Al acabar *Palabra y objeto* (secc. 56), Quine se ocupará de señalar con él ciertas coincidencias y divergencias que son muy significativas. Carnap había mantenido, según vimos, que toda cuestión filosófica genuina es una cuestión lingüística; Quine responderá: ¿y por qué no decir lo propio de toda cuestión teórica? En última instancia toda cuestión teórica versa también sobre el lenguaje, pues en el curso del desarrollo de las teorías se va produciendo lo que llama Quine «ascenso semántico», un proceso que, por cierto, y como él mismo reconoce, es como el paso del modo material al modo formal, en la expresión de Carnap. Consiste en pasar de hablar de objetos a hacer afirmaciones sobre las palabras por medio de las cuales nos referimos a ellos. Pero mientras que Carnap veía aquí lo característico de la filosofía genuina, Quine percibe este proceso en todo ámbito teórico. Como ya indicamos en su momento (secc. 8.2), lo cierto es que toda proposición de objeto puede ponerse en el modo formal. ¿Cuáles son las ventajas de hacerlo? Según Quine, situar la discusión en un campo en el que sea más fácil alcanzar un acuerdo. El ascenso semántico traslada nuestras divergencias desde los objetos a las palabras, y cuando se trata de objetos abstractos y de carácter teórico, el cambio puede contribuir de modo decisivo al acuerdo. De aquí que resulte especialmente recomendable en filosofía, aunque también es usual recurrir a él en las ciencias, como lo recuerdan los casos aducidos por Quine (*loc. cit.*). La formulación metalingüística de los problemas no es ya, al contrario de lo que ocurría en Carnap, la característica metodológica de la filosofía frente a la ciencia. Y no puede serlo, puesto que la distinción entre verdades fácticas y verdades analíticas ha quedado disuelta en un continuo en el que tan sólo es posible hacer diferencias de grado. Lo que caracteriza al filósofo no es un punto de vista propio sino simplemente la amplitud de sus categorías, que le conducen a tratar de los compromisos ónticos de todas las teorías.

8.8 Crítica a Quine y defensa de la analiticidad

Igualmente que nos ocurrió con otros autores estudiados anteriormente, el resumen precedente de la filosofía del lenguaje de Quine ha tenido

que omitir, no sólo detalles, sino incluso temas enteros, algunos de indiscutible interés. Hemos ignorado, así, todo el estudio de la ontogénesis de la referencia, que se contiene en el capítulo 3 de *Palabra y objeto*, y que considera el desarrollo, en el individuo, de los varios modos de referirse por medio del lenguaje al mundo, dentro de los límites concretos de la lengua inglesa. Tampoco hemos examinado determinados problemas de la referencia como la llamada «opacidad referencial», esto es, la imposibilidad, en ciertos contextos lingüísticos, de sustituir entre sí términos con la misma referencia sin que cambie el valor de verdad del enunciado (*op. cit.*, secciones 30 y ss.). Pero no había tiempo para tratar todo, y a fin de conservar la proporción en el espacio dedicado a Quine ha sido imprescindible, igual que en otros casos, limitarnos a los aspectos de alcance más general. Tampoco se ha hecho alusión al problema de la adquisición del lenguaje, sobre el que Quine tiene interesantes sugerencias y ha intervenido en una sustanciosa polémica con Chomsky; pero en este punto puedo, al menos, remitir al lector a otro lugar donde he tratado el tema con detenimiento (*La teoría de las ideas innatas en Chomsky*, caps. 5 y 6).

El pensamiento de Quine sobre el lenguaje recoge, en su conjunto, la herencia de Carnap en algunos de sus aspectos más positivos, y se elabora en gran parte como respuesta polémica a ciertas tesis de aquél. La teoría de Quine, sumamente compleja en general, muy concisa y detallada en algunos puntos, demasiado vaga en otros, con frecuencia malentendida, probablemente a causa del peculiar estilo literario del autor, no es fácil de sintetizar, y menos aún de evaluar. Sin duda, uno de sus principales puntos de apoyo es la negación de la distinción radical entre verdades analíticas y verdades sintéticas. El resto de esta sección la dedicaremos, por ello, a una discusión, e intento de refutación, de esta tesis. Junto a ella, o más bien entrelazada con ella, está la tesis de la indeterminación de la traducción radical. Esta última es una tesis tan laboriosamente elaborada por Quine que no resulta fácil de atacar. La cuestión es qué prueba. Desde luego, no pretende probar que, para lenguas emparentadas filológicamente o relacionadas culturalmente, no sea posible traducir de una a otra de modo aceptable a efectos prácticos. Lo que pretende mostrar es que una traducción única y segura, especialmente cuando las lenguas carecen de relaciones previas conocidas, es imposible excepto para oraciones de observación, y en su mayor parte, para las funciones veritativas. Y esto a su vez pretende dejar manifiesto que los significados, tomados como entidades mentales comunes a las oraciones sinónimas, carecen de criterios empíricos que permitan utilizarlos para justificar la traducción y, por consiguiente, que el concepto de significado, así entendido, es inútil en la teoría semántica (el significado estimulativo no es más que un sustitutivo conductista, como el propio Quine reconoce). A esto podríamos oponer que los criterios empíricos que hacen utilizable el concepto de significado son precisamente esa comunidad de pautas de comportamiento en que consisten las relaciones culturales y que acompañan a las relaciones históricas entre las diversas lenguas. Las oraciones ocasionales que no son de observación, las oraciones

permanentes, e incluso las oraciones eternas, como Quine las llama respectivamente, podrían coincidir en significado de una lengua a otra en la medida en que los usuarios de ambas lenguas poseyeran instituciones culturales y formas de vida suficientemente parecidas. Cuanto más lo fueran, más fácil sería mostrar la sinonimia de sus oraciones respectivas, y mayor sería el número de casos en que podríamos mostrarlo. Naturalmente, esto no afecta a la traducción propiamente radical; que en el caso extremo imaginado por Quine no existan criterios empíricos que nos permitan justificar una traducción mejor que otra, parece indiscutible. Lo que queda en cuestión es el valor que hayamos de atribuir a este tipo de casos: tal vez sean útiles en la crítica a una concepción platónica, o incluso meramente mentalista, del significado, pero no parecen relevantes para una concepción del significado que podríamos llamar «culturalista» o pragmática.

La traducción del lenguaje a la lógica de predicados tiene también, por su parte, limitaciones muy precisas. Obedece al propósito de arrojar claridad sobre el aparato referencial del lenguaje, mostrando cómo reducir todos los tipos principales de expresión a los escasos recursos de la porción más básica, más sencilla y más ampliamente aplicable, de la lógica. Que tal regimentación sea posible es, sin duda, importante por lo que respecta a la utilización del lenguaje en la formulación de teorías. Pero no se atribuya a Quine ningún intento de ofrecer una teoría sobre la estructura del lenguaje natural. Por esta razón, el que a veces resulte necesario forzar las expresiones ordinarias para hacerlas entrar en la horma lógica no es algo que a Quine le preocupe. Ya vimos un ejemplo de esta situación en su tratamiento de los nombres propios ordinarios, y de qué manera su propuesta se apartaba considerablemente de la forma ordinaria de usarlos, llegando a parafrasear algunas oraciones con artificiosidad extremada. La idea era convertir todos los términos singulares en predicados, dejando en posición de sujeto, y por consiguiente con función referencial, únicamente las variables. De aquí el criterio de compromiso óntico: ser es ser el valor de una variable.

Pero hay aquí un problema: ¿cómo saber cuál es el dominio de individuos que constituyen los valores posibles para las variables de nuestro lenguaje? No parece que haya otro camino que atender a los predicados (así, Blasco en «Compromiso óntico y relatividad ontológica»). Si afirmamos «Hay seres vivos unicelulares», su expresión canónica semiformalizada será «Para algún x , x es ser vivo y unicelular». Nuestra pregunta es: ¿a qué individuos puede referirse « x »? No se ve otra manera de averiguarlo que buscando seres vivos e inquiriendo si son o no unicelulares. Esto es, que los valores de x en esa oración son ¡seres vivos! Por consiguiente, son los predicados los que determinan el dominio de las variables, y en su caso, de aquellos términos del lenguaje natural que ofician de variables, como «ello», «algo», «alguno», etc. La consecuencia es que podemos reformular nuestro criterio de compromiso óntico así: en lugar de «ser es ser el valor de una variable», diremos «ser es satisfacer un predicado» o

«ser es ser denotado por un término general» (fórmulas expresamente admitidas por Quine en su respuesta a Blasco, p. 165 de *Aspectos de la filosofía de W. V. Quine*). ¿Qué razones había para construir el criterio sobre las variables? Simplemente razones lógicas de carácter histórico, a saber, «el papel central de la variable en la lógica neoclásica»; y añade Quine: «Desde este punto de vista la noción de término general o de predicado es derivativa. Defino el término general o el predicado como lo que se aplica a una variable para formar una sentencia abierta.» (*Loc. cit.* pp. 165-166.) En suma: lo que hay, para un determinado lenguaje o teoría, es simplemente todo aquello de lo que, en ese lenguaje o teoría, se predica algo. No hay otra vía hacia los objetos excepto sus propiedades y relaciones; o lo que es lo mismo: un lenguaje, una teoría, no habla propiamente de objetos, puesto que no contiene criterios para especificar los valores de sus variables; en rigor, tan sólo habla de clases de objetos, a saber, aquellas clases que vienen definidas por cada predicado. Esto último se debe a que, como hemos visto, Quine sustituye las propiedades y relaciones por clases; cuestión aparte es si, desde el punto de vista ontológico y epistemológico, puede operarse con clases sin reconocer propiedades y relaciones.

Pasemos, por fin, a la analiticidad. ¿Es posible escapar a las objeciones de Quine? Para defender el concepto de verdad analítica, muchos dirían que estas verdades son aquellas que estamos dispuestos a mantener en cualquier circunstancia, ante cualquier eventualidad. Quine contesta: si esto quiere decir «tras cualquier estimulación», no hemos salido de la analiticidad estimulativa (*Palabra y objeto*, secc. 14). Y con ello, según hemos visto, no podemos distinguir la analiticidad de la coextensionalidad. Incluso aun socializada del modo que ya conocemos, la analiticidad estimulativa no constituye una reconstrucción conductista del concepto de analiticidad, sino, todo lo más, como el propio Quine reconoce, su «sucedáneo conductista». Porque lo que clásicamente pretende explicar el concepto de analiticidad es por qué estamos dispuestos a mantener ciertas verdades frente a cualquier cambio en los hechos. Es curioso que Quine mismo llega en algunas ocasiones a tocar en este punto del problema; así cuando comenta que, ante la negación de una verdad analítica por otro hablante, nuestra reacción se parece más a la reacción que tendríamos ante una expresión extranjera que no hubiéramos comprendido (*loc. cit.*): tenderemos, entonces, a considerar analíticas aquellas oraciones cuya negación por nuestro interlocutor nos haría preguntarnos qué es lo que éste quiere decir, o nos haría dudar de que nos entendemos. ¿Pero por qué es esto así? A pesar de la resistencia de Quine, no se advierte respuesta más clara que ésta: porque lo que una oración analítica manifiesta es una regla lingüística, y por consiguiente negar aquella equivale a salirse fuera de las reglas del juego —del juego lingüístico, podríamos ahora decir recordando a Wittgenstein—; equivale a no hablar el lenguaje en el que la oración se formula.

Que los factores que determinan la analiticidad de las oraciones, y la sinonimia de las expresiones, están en el comportamiento de los hablantes,

es algo en lo que no es posible disentir de Quine. Es el comportamiento de los hablantes el que determina cuanto se refiere al lenguaje natural, puesto que usar un lenguaje no es sino una forma de comportarse. Pero el comportamiento lingüístico es un comportamiento sujeto a reglas, y las reglas que ese comportamiento pone de manifiesto y a las que, muchas veces incluso de manera consciente, pretende sujetarse, pueden en principio formularse de forma explícita y mínimamente rigurosa. Al invocar el comportamiento no se debe pasar por alto que las explicaciones acerca de la conducta lingüística son, a su vez, una forma de conducta lingüística; y nuestra conducta lingüística incluye la distinción entre lo que atribuimos a las reglas de nuestro lenguaje y lo que asignamos a cómo es el mundo. Frente a Carnap, Quine lleva la razón en esto: poseer un concepto de analiticidad para un lenguaje artificial o sistema semántico no implica que ese concepto sea aplicable al lenguaje natural. Pero Quine parece dar por supuesto que no es posible suministrar, para una lengua natural, reglas semánticas o postulados de significado al modo de los que da Carnap para sus sistemas semánticos. Y, sin embargo, esto es lo que hace un lingüista cuando estudia la semántica de una lengua.

Así, por ejemplo, dentro de la lingüística chomskiana, Katz, en los términos propios de su teoría semántica, que brevemente estudiamos en la sección 4.4, ha ofrecido una definición de la analiticidad del siguiente corte: una oración *O* es analítica en una lectura (o sentido) *L* si y sólo si todo indicador semántico no complejo del predicado de *O* aparece también en el sujeto de *O*; y para todo indicador semántico complejo del predicado de *O*, entonces hay un elemento de éste que aparece en el sujeto de *O*. (*Filosofía del lenguaje*, cap. 5; *Teoría semántica*, 2.6 y 4.5, y para la discusión de Quine, 6.2). Se notará aquí una formulación lingüística de la idea de Kant, recordada antes, de que un juicio analítico es aquel cuyo predicado está incluido en el concepto del sujeto. La definición de Katz incluye la cautela de referirse sólo a un sentido determinado, o en su terminología, «lectura», de una oración, pues es patente que una oración ambigua puede ser analítica en un sentido pero no en otro. En rigor, podríamos considerar cada lectura de una oración como una oración distinta de las demás. La idea de Katz es que, al descomponer el significado del sujeto y del predicado en indicadores semánticos, al modo que vimos en la sección 4.4, podemos encontrarnos con que los indicadores simples que constituyen el significado del predicado, forman parte también del significado del sujeto. En tal caso, la oración es analítica. Por lo que toca a nuestros ejemplos anteriores, ello implica que los componentes semánticos del significado de «animal» están incluidos en el significado de «perro», y que los que pertenecen al significado de «no casado» están asimismo en «soltero». En cambio, la no analiticidad de la oración:

(10) Toda moneda dorada de curso legal es una moneda de una peseta

quedará mostrada por el análisis semántico al hacer explícito que hay componentes del significado de la expresión «moneda de una peseta» que no forman parte del significado de la expresión «moneda dorada de curso legal». Lo propio puede añadirse para el caso de la oración:

(11) Todo animal que tiene corazón es un animal que tiene riñones

Lo anterior, por cierto, no debe sugerir que únicamente las oraciones formadas por sujeto y predicado puedan ser analíticas. Como Katz recuerda (*Teoría semántica*, 4.5), son igualmente analíticas oraciones como:

(12) Se recuerda lo que no se olvida

(13) Se vende a quien compra

etcétera.

Acomodar estos casos en la definición anterior requiere simplemente reelaborarla de forma un poco más compleja, como hace Katz en el lugar indicado. Aquí pasaremos por alto esta complejidad.

Katz ha formulado en términos de la moderna teoría semántica una vieja definición de la analiticidad, esbozando un procedimiento de decisión que nos permita averiguar cuándo es analítica una oración. Pienso que la propuesta de Katz es un paso adelante que nos pone en camino de conseguir un criterio a la vez riguroso y operativo para identificar las oraciones analíticas en el lenguaje ordinario. El problema es que no está claro aún cómo se ha de realizar el análisis semántico; es cuestión que depende del progreso en la teoría semántica y que va ligada a su desarrollo. Es perfectamente obvio que el concepto de analiticidad está vinculado a otros conceptos semánticos como los de significado, sentido y sinonimia, y que la crítica de Quine afecta por igual a todos estos conceptos y obliga a prescindir de ellos, excepto por lo que hace a sus reformulaciones en términos estimulativos, que ya conocemos. Ahora bien, los nuevos conceptos de significado estimulativo y analiticidad estimulativa no nos sacan del ámbito de la extensionalidad (ni lo pretenden, por supuesto), y sin salir de este ámbito no se ve cómo pueda darse cuenta cabal de aspectos tan usuales en el comportamiento lingüístico como la distinción entre no creer lo que nos dicen y no entenderlo. Si alguien nos comunicara que había encontrado un animal con pulmones y sin corazón, podríamos dudarle mientras no nos lo enseñara; si, en cambio, afirmara haber hallado un soltero casado, ciertamente no esperaríamos que pudiera mostrarnos tan extraño ejemplar. Frente a lo que nos dicen, distinguimos siempre, en general de forma coherente y con bastante seguridad, entre entenderlo y creerlo (como han subrayado Strawson y Grice en «In Defense of a Dogma»). Que no hay solteros casados no es simplemente una creencia en la que todos los hablantes del castellano coincidamos. Expresa, aunque formulada incorrectamente, la regla del castellano que regula el uso del término «soltero».

La defensa del concepto de analiticidad es considerablemente más difícil con respecto al lenguaje ordinario, que es el que a nosotros nos interesa en particular. Aquí, el estado todavía primitivo de la teoría lingüística y las dificultades que aún se han de vencer para alcanzar una teoría semántica rigurosa, repercuten considerablemente en nuestro intento de deslindar con precisión las oraciones analíticas de las demás. La tarea parece, en cambio, bastante más sencilla con relación a lenguajes particulares como los de las teorías científicas axiomatizadas, o sea, teorías tales como la geometría de Euclides, la aritmética o la mecánica clásica. En estos casos se puede estipular lo siguiente: serán proposiciones analíticamente verdaderas en la teoría sus axiomas y sus teoremas, y serán analíticamente falsas las negaciones de unos y otros; mientras que serán proposiciones sintéticas las definiciones de la teoría y las proposiciones indeterminadas en ésta (esto es, aquellas proposiciones tales que ni ellas ni sus negaciones son teoremas; esta propuesta es de Ulises Moulines en «Lo analítico y lo sintético: dualismo admisible»). La noción de que rechazar una proposición analítica implica rechazar un lenguaje posee en el terreno de las teorías axiomáticas perfiles más claros, pues es patente que rechazar un axioma o un teorema equivale a rechazar la teoría. Pero la idea de que una proposición analítica es verdadera exclusivamente a causa de lo que significan sus palabras sólo resulta aplicable en el caso de los axiomas y teoremas de un sistema formal, como los de la aritmética o la geometría (y con mayor razón en la lógica, aunque aquí las verdades analíticas se reducen a las verdades lógicas). No así en el caso de una teoría empírica, pues es obvio que aquí los axiomas no pueden aceptarse como verdaderos con independencia de la realidad extralingüística. Considérese, por ejemplo, el segundo axioma de la mecánica de Newton: «El cambio en el movimiento de un cuerpo es proporcional a la fuerza motriz aplicada y actúa en la línea recta según la cual se aplica ésta.» ¿Puede considerarse que la verdad de esta afirmación sea independiente de la realidad como lo es la verdad de «Ningún soltero está casado»? Si dijéramos tal cosa, no veo cómo podríamos distinguir entre teorías formales y teorías empíricas, entre sistemas lógico-matemáticos y sistemas físicos. Cosa distinta ocurre con los axiomas de una teoría matemática. Considérese el primero de los postulados de Euclides: «Por dos puntos distintos pasa una única recta»; o el segundo de los axiomas de Peano para la aritmética: «El sucesor de cualquier número es un número.» Aquí sí tiene sentido ver, en estas verdades, la mera formulación de las reglas que constituyen el lenguaje de estas teorías, y la semejanza con «Ningún soltero está casado» es aceptable. Para las teorías empíricas, como la mecánica de Newton, la propuesta de Moulines casa mal con alguna de las más clásicas caracterizaciones del concepto de analiticidad.

Muchos, aun defendiendo la distinción entre enunciados analíticos y sintéticos, la han matizado de diferentes formas que permitan acomodar las quejas de Quine. Así, Putnam («The Analytic and the Synthetic») piensa que, aunque hay que reconocer ciertos enunciados como analíticos y

otros como sintéticos, Quine tiene en el fondo razón en su crítica, pues se trata de una distinción trivial y carente de consecuencias filosóficas, que más que ayudar en la resolución de problemas filosóficos puede conducir a serios errores. Bunge, por su parte («Análisis de la analiticidad»), relativiza la distinción a cada tipo de contexto y de teoría, señalando diversas clases de enunciados analíticos que van desde las definiciones a los teoremas.

Aquí prescindiremos de las verdades lógicas, cuyo carácter analítico no ha sido puesto en duda por Quine. Prescindiremos también de los axiomas y teoremas de las teorías formales, que, claramente en el caso de la aritmética, y de modo más disputable en el caso de la geometría, tienen también carácter analítico. Por lo que respecta a los axiomas y teoremas de las teorías empíricas, atribuirles carácter analítico requiere, como hemos visto, prescindir de la idea de que la verdad analítica es independiente de la realidad, con lo que tendríamos un concepto de analiticidad excesivamente amplio. Pero no entraremos en más detalles sobre este problema, pues nos obligaría a abordar una serie de cuestiones propias de la filosofía de la ciencia. En cuanto a las definiciones de una teoría, tomadas como propuestas, es claro que no son ni verdaderas ni falsas, y por tanto que no son propiamente analíticas. A continuación, nos limitaremos, fundamentalmente, a las proposiciones analíticas en el lenguaje ordinario, aquellas que son verdaderas por lo que significan sus términos no lógicos, y que algunos han llamado «tautonomías» (así, Bunge, *loc. cit.*), como nuestro maltratado ejemplo «Ningún soltero está casado».

Al defender la distinción entre proposiciones analíticas y sintéticas para el lenguaje ordinario, lo que pretendo es subrayar la diferencia que tienen entre sí las dos oraciones que forman cada uno de los siguientes pares:

- (14) Ningún soltero está casado
- (15) Ningún soltero se libra de la neurosis
- (16) Todo dolor es desagradable
- (17) Todo dolor es transmitido por fibras nerviosas de tipo A-delta y C
- (18) Los mamíferos poseen mamas
- (19) Los mamíferos tienen pelo

Las oraciones (14), (16) y (18) son analíticas porque su verdad deriva únicamente de lo que significan sus palabras, está determinada tan sólo por las reglas semánticas del castellano; de quien las negara diríamos que no usa las palabras «soltero», «dolor» o «mamífero» en su sentido aceptado. Esto no podríamos afirmarlo, en cambio, de quien negara (15), (17) y (19), pues la verdad de estas oraciones depende, no sólo de lo que significan sus palabras, sino también de cómo es la realidad; un mejor conocimiento de ésta podría, en principio, llevarnos a concluir la falsedad de esas afirmaciones; ningún aumento en nuestros conocimientos podría, en cambio, hacernos pensar en la falsedad de las tres primeras, (14), (16) y (18). Siendo esto así, ¿cómo habremos de caracterizar la clase de las proposiciones ana-

líticas? Existen varias posibilidades, algunas de las cuales han ido apareciendo implícitamente en las páginas anteriores.

Una primera caracterización, de tipo semántico, estipularía que son analíticamente verdaderas aquellas proposiciones que están aseguradas contra la falsedad, esto es, que son necesariamente verdaderas; y lo contrario para las que son analíticamente falsas. Lo malo de esta caracterización es que acentúa la necesidad de la verdad o de la falsedad de la proposición analítica, con lo que parecería que se trata de una verdad o de una falsedad muy importantes, cuando en realidad son una verdad y una falsedad totalmente triviales en cuanto que independientes de los hechos extralingüísticos.

Otra caracterización, también de orden semántico, consistiría en calificar como analíticas aquellas proposiciones que son verdaderas o falsas en todos los mundos posibles, retomando la vieja idea de Leibniz (así, Lewis, *Convention*. V.3). El problema es que la noción de mundo posible es extremadamente oscura, y no se ve bien cómo definirla excepto recurriendo a las verdades analíticas: «Los mundos posibles son modelos... para un lenguaje suficientemente rico; pero, por supuesto, no todos los modelos para ese lenguaje ¡sólo aquellos que satisfacen sus verdades analíticas!» (Lewis, *op. cit.*, p. 207).

Podríamos mantener, en tercer lugar, que son analíticas aquellas proposiciones que son independientes de nuestro conocimiento de los hechos extralingüísticos. Naturalmente que para reconocer una proposición como analíticamente verdadera o falsa tenemos que conocer ciertos hechos lingüísticos, aquellos que consisten en que las palabras relevantes tienen tal o cual significado. Pero bastará saber esto, bastará conocer, en suma, las reglas lingüísticas relevantes, para declarar la oración verdadera o falsa, sin que se requiera saber nada más sobre la realidad. Sería ésta una caracterización epistemológica. En mi opinión es adecuada, y expresa una condición necesaria y suficiente para calificar una proposición como analítica. No quiero decir con ello que sea del todo transparente. Está, en primer lugar, el problema de que todo conocimiento de reglas lingüísticas incluye en ciertos aspectos conocimiento de la realidad extralingüística. ¿Cómo podría entender lo que significa «soltero» quien desconociera la institución del matrimonio, o cómo podría saber lo que quiere decir «dolor» quien nunca hubiera percibido una manifestación de desagrado? Y en segundo término, hay que contar con que el desarrollo y modificación de nuestro conocimiento de la realidad puede llevarnos a modificar el significado de las palabras, con lo que ciertas proposiciones que antes eran analíticas dejarán de serlo, y otras que no lo eran pueden llegar a serlo. Esto nos conduce a una cuarta caracterización.

Es la idea, que hemos visto antes recogida, de que son analíticas aquellas proposiciones que son indispensables para usar un lenguaje. Podemos considerar esta caracterización como pragmática. Deja, sin embargo, abierta esta cuestión: ¿por qué son ciertas proposiciones indispensables para utilizar un lenguaje? Pienso que en la respuesta a esta pregunta se encierra lo que es decisivo para nuestro problema. Y la respuesta que propongo es

muy sencilla: porque esas proposiciones expresan aquellas convenciones de significado que en cada momento es necesario mantener para poder manifestar los cambios en nuestros conocimientos y comunicarlos a los demás. Lo que quiero decir es que si las palabras no tuvieran un significado suficientemente estable no podrían ser usadas para la comunicación por una comunidad de hablantes. Pues bien, las proposiciones analíticas (del lenguaje ordinario) expresan esas convenciones de significado, únicamente que lo hacen de manera confundente. En efecto: la forma clara, explícita y literal de expresar una regla de significado es por medio de un enunciado metalingüístico, tal como:

(20) La palabra «soltero» significa en castellano «varón que nunca ha estado casado»

(21) La palabra «dolor» significa en castellano «sensación molesta y aflictiva»

(22) La palabra «mamífero» significa en castellano «animal cuya hembra posee mamas»

Y lo que mantengo es que las correspondientes oraciones analíticas, como nuestros ejemplos (14), (16) y (18), expresan exactamente lo mismo que estas oraciones metalingüísticas, pero como no están formuladas metalingüísticamente parecen versar sobre la realidad extralingüística, y justamente esto es lo que nos confunde. Nótese que la formulación metalingüística no es posible en los ejemplos (15), (17) y (19), pues daría este resultado:

(23) La palabra «soltero» significa en castellano «varón neurótico»

(24) La palabra «dolor» significa en castellano «sensación transmitida por fibras nerviosas de tipo A-delta y C»

(25) La palabra «mamífero» significa en castellano «animal con pelo»

Oraciones, todas ellas, cuya falsedad puede probarse empíricamente recurriendo al empleo práctico y real del castellano en la actualidad. Esto prueba que las correspondientes oraciones (15), (17) y (19) son sintéticas.

Ahora bien, si (23), (24) y (25) son empíricamente falsas, de (20), (21) y (22) hay que decir que son empíricamente verdaderas, y por consiguiente, que son oraciones sintéticas, no analíticas. Y se comprende que así sea, ya que se trata de afirmaciones acerca de cuáles son determinadas reglas lingüísticas del castellano. Las definiciones (20) a (22) no forman parte de un lenguaje artificial que yo acabe de inventar y proponga aquí. Si así fuera, constituirían propuestas, y en tal caso no serían propiamente ni verdaderas ni falsas (y esto es lo que más arriba hemos considerado sobre las definiciones pertenecientes a teorías). Las definiciones que estamos estudiando, en cuanto pertenecientes al lenguaje ordinario, pretenden expresar reglas de significado que forman parte del sistema semántico de una lengua particular ya en uso, y por esta razón hay que tomarlas como afirmaciones

empíricas, sintéticas. (No hay que olvidar que muchas de estas reglas pueden regir sólo en ciertas comunidades nacionales o en ciertos subgrupos sociales de hablantes, pero ello no afecta a la presente argumentación ni —según creo— tampoco a los ejemplos que estamos considerando.)

Volviendo a las oraciones analíticas, como (14), (16) y (18), lo que pretendo sugerir es lo siguiente: es característico de las oraciones analíticas en el lenguaje ordinario el expresar una regla de significado como si se tratara de una afirmación sobre la realidad extralingüística. Aplicando por nuestra cuenta la expresión de Carnap, podríamos decir que se trata de pseudoproposiciones de objeto: no nos dan información alguna sobre los objetos, pues lo que afirman sobre ellos se deriva exclusivamente de lo que significan las palabras que los designan.

Esto por lo que respecta a la caracterización de las proposiciones analíticas. ¿Cómo puede decidirse si una oración determinada es o no de esta clase? El desarrollo de un procedimiento efectivo de decisión para estos propósitos va ligado al progreso de la teoría semántica, y por lo que hemos visto (sección 4.4) a ésta le resta aún mucho trecho para alcanzar un estatuto de solidez y rigor. Mientras no conozcamos claramente la manera de determinar cuáles son los significados de una palabra, no podremos decidir con cierta seguridad sobre si la correspondiente oración en que aparece esa palabra es o no analítica. Pero esto no implica que no tengamos un concepto claro de lo que es una oración analítica. Saber en qué consiste una oración analítica y poseer un procedimiento para decidir si una determinada oración lo es o no, son cuestiones separadas y distintas. Ante un cierto enunciado podemos dudar si es o no analítico en la medida y grado en que dudemos sobre cuál es la regla de significado aplicable al caso. Creo que los ejemplos que he utilizado anteriormente son especialmente claros, pero sin duda no todos los casos son así. Hay proposiciones más claramente analíticas que otras; porque la claridad es aquí relativa al estado de nuestra teoría semántica, así como a la mayor o menor rapidez en el cambio de significado de las palabras. Pero no hay proposiciones más o menos analíticas; la analiticidad no es una cuestión de grado, como no lo es tampoco el significado.

Qué proposiciones sean analíticas es relativo a un lenguaje en cuanto interpretado por una teoría semántica. Pero una teoría semántica para un lenguaje natural es una teoría empírica que pretende dar cuenta de cómo puede analizarse el significado en ese lenguaje de acuerdo con el uso que los hablantes hacen de él. Esto es considerablemente más difícil que suministrar postulados de significado para un lenguaje artificial al modo de Carnap. Pero, como ya vimos, fue este mismo quien sugirió cómo buscar las reglas de significado en los lenguajes naturales, y nadie ha demostrado que no sea posible hacerlo. Estas reglas o convenciones de significado describirán los límites semánticos del lenguaje en cuestión. Y tomando fuera de su contexto un término típico del primer Wittgenstein, podemos añadir: mientras que las reglas de significado describen esos límites, las proposiciones analíticas simplemente los *muestran*.

(*Nota bene*: en esta sección he empleado «oración» como «oración con cierto significado», esto es, semánticamente unívoca; «proposición» como «lo que una oración significa»; y «enunciado» como «oración en cuanto utilizada»; creo que, por esta razón, mi uso indistinto de los tres términos no afecta al problema debatido ni debe producir confusión. Y por supuesto: Quine rechazaría el término «proposición» así entendido.)

8.9 Significado y verdad

Como vimos en su momento, Tarski pensaba que no es posible formular una definición rigurosa del concepto de verdad para un lenguaje natural. La razón era que los lenguajes naturales no son formalmente especificables, pues carecerían de reglas exactas de formación, y que en ellos se confunden lenguaje objeto y metalenguaje. La actitud de Tarski —señalé— resulta hoy demasiado pesimista, y algunos, por el contrario, piensan que es posible aplicar a las lenguas naturales su noción de verdad y, más todavía, que ésta es la base y el fundamento para construir una teoría del significado. El defensor más destacado de esta tesis, en los últimos años, ha sido Davidson.

La mayor parte de los filósofos del lenguaje a partir de Frege, y muchos lingüistas actuales, como Katz (recuérdese la sección 4.4), han defendido que la teoría del significado debe explicar de qué manera depende el significado de una oración de los significados de las palabras y expresiones que intervienen en ella. El problema es que, cuando tenemos los significados de las expresiones más simples de una oración, nos queda aún por explicar cómo contribuyen éstos al significado total de la oración (esto es lo que aspiraban a explicar las reglas de proyección en la teoría de Katz, según vimos). La objeción básica de Davidson a esta clase de teorías es que las explicaciones que suministran son vacuas, pues no aclaran cómo resolver sobre el significado de la oración, sino que se quedan en el estadio inicial, en la asignación del significado a los nombres y predicados («Truth and Meaning», 1967). Como en Quine, los significados quedan también ahora en entredicho; pero no porque sean entidades oscuras o porque no sepamos cómo decidir si dos expresiones tienen o no el mismo significado, sino por una razón más simple y más elemental: porque de hecho carecen de aplicación a las oraciones como tales (*op. cit.*, p. 453).

La sugerencia de Davidson es adoptar una concepción en cierto modo holista del significado; esto es, tomar el significado de cada oración como algo unitario y propio de ella, que tiene en la medida en que pertenece a un determinado lenguaje. Dicho de otro modo: una teoría del significado debe ser tal que de ella puedan deducirse todas las oraciones de la forma «*X* significa *s*», para el lenguaje de que se trate (siendo «*X*» cualquier oración de dicho lenguaje y «*s*» su significado). Pero puesto que lo que está oscuro es a qué pueda referirse «*s*», o sea, en qué consista el significado de una oración, Davidson sugiere a continuación sustituir «*s*», supuesto nombre de un significado, por algo más claro y más manejable, por una

oración que nos dé el significado de «X». Recurriendo a la variable de oraciones usual en estos casos, p , nuestra afirmación anterior tendrá esta forma: «X significa que p ». ¿Qué quiere decir aquí la expresión «significa que»? Este es justamente el centro de la cuestión. La idea de Davidson es que la mejor manera de elucidar esa expresión en tal contexto es sustituirla por el predicado «es verdadera», añadiendo el juntor «si y sólo si». Tras semejantes modificaciones, la afirmación en cuestión tendrá esta forma: «X es verdadera si y sólo si p ». Y como se habrá notado, ésta es justamente la condición que debe cumplir toda definición de la verdad para ser materialmente adecuada, según Tarski. La condición general de significatividad de la que habíamos partido ha quedado transformada en la convención (T) de Tarski (recuérdese la sección 8.3). Si sustituimos « p » por una oración y «X» por la descripción estructural de esa oración, la convención (T) expone en qué consiste que dicha oración tenga significado. Dicho en palabras de Davidson: «Una teoría del significado para un lenguaje L muestra 'de qué manera dependen los significados de las oraciones de los significados de las palabras' si contiene una definición (recursiva) de verdad-en- L » (*op. cit.*, p. 455).

En la medida en que la definición de Tarski suministra condiciones necesarias y suficientes para la verdad de cada oración, dar dichas condiciones es *una manera* (*sic* en Davidson) de dar el significado de cada oración. Así entendida, la teoría del significado es una teoría empírica, pues da lugar a consecuencias que pueden confirmarse empíricamente, ya que puede siempre comprobarse de este modo si las condiciones de verdad de una oración determinada son las que la teoría especifica para ella. Semejante comprobación es, desde luego, trivial si el lenguaje cuyo significado se está elucidando es parte del lenguaje que usa y entiende el que está aplicando la teoría. No lo es, en cambio, si se trata de averiguar el significado de las oraciones en un lenguaje ajeno; para esta empresa, Davidson invoca un principio de caridad: asumir que los hablantes de la lengua investigada son, en general, tan coherentes como pueda serlo el investigador, e intentar que el acuerdo entre las condiciones de verdad de las oraciones ajenas y las de las oraciones del investigador sea lo mayor posible. De no seguir estos principios, podría ocurrir que se hiciera una interpretación del lenguaje ajeno injustificadamente divergente respecto al lenguaje propio. En todo caso, si la desconexión previa entre ambos lenguajes es completa, toparemos con la indeterminación de la traducción radical puesta de manifiesto por Quine, y que Davidson asimismo reconoce.

La propuesta de Davidson implica, por consiguiente, que es posible aplicar la noción tarskiana de verdad a los lenguajes naturales. Que en éstos puedan aparecer paradojas semánticas como la del mentiroso no le parece razón suficiente para desistir de la empresa (*op. cit.*, p. 459); que en ellos falten reglas precisas que establezcan con rigor la gramaticalidad de las expresiones y que permitan su tratamiento formal, es hoy menos defendible que en tiempos de Tarski, en razón del estimable progreso que ha tenido en los últimos decenios la gramática formal y la lingüística choms-

kiana (*ibid.*). Pero esta fundamentación del significado sobre la verdad tiene problemas propios. Acaso el que más de inmediato salte a la vista sea la exigencia de que podamos predicar la verdad de todo tipo de oraciones, sean propiamente declarativas, o bien valorativas, prescriptivas, expresivas o realizativas. Tal pretensión contraría la clasificación de los tipos de discurso propuesta anteriormente (7.8); allí pretendimos hacer una clasificación semántica de los tipos de oraciones atendiendo a la relación entre éstas y la realidad, y atribuimos la verdad y la falsedad exclusivamente a las oraciones declarativas. Davidson no ve aquí problema sustancial alguno; el problema consistiría en descubrir el significado de palabras determinadas como «bueno», «deber», etc. Pero lo único que impone su teoría es que las oraciones en las que intervienen dichas palabras sean analizadas semánticamente en términos de sus condiciones de verdad; es decir, la teoría simplemente implica equivalencias del tipo de «'Fulano es bueno' es verdadera si y sólo si Fulano es bueno», «'Se debe nacionalizar la banca' es verdadera si y sólo si se debe nacionalizar la banca», etc. Lo que aquí no está claro es cómo podemos comprobar empíricamente si se dan o no las condiciones de verdad propias de tales oraciones. Pienso que hay aquí un problema complejo, pero no insistiré en él por el momento. En cambio, donde sí reconoce Davidson este problema es en las oraciones imperativas; no parece que tendría sentido predicar de ellas condiciones de verdad, e incluso ni que pueda formularse para ellas la convención (T), ya que daría lugar a oraciones aparentemente mal formadas, del tipo de: «'Vuelva usted mañana' es verdadera si y sólo si vuelva usted mañana.»

Otra cuestión es que muchas oraciones, por contener términos deícticos o referencialmente ambiguos, pueden variar en su valor de verdad según quién y cuándo las use. Para estos casos, la recomendación de Davidson es simple y atinada: la verdad y la falsedad pueden relativizarse a una oración en cuanto usada por un cierto hablante en un tiempo determinado. En fin, otras varias dificultades de diversa entidad han quedado indicadas por el propio Davidson (*op. cit.*, p. 465; todo lo anterior puede encontrarse resumido también en su artículo «Semantics for Natural Languages»).

La desconfianza frente al concepto de significado (tomado como sentido, no como referencia), que hemos visto ejemplificada en Quine, recibe, pues, de Davidson nueva justificación. Era de esperar, por ello, que Quine acogiera con toda simpatía la teoría de Davidson, encontrando entre la de éste y la suya numerosas coincidencias («Reply to Davidson», en *Words and Objections*, pp. 333-335). La idea básica de Davidson, a saber: que una vez que se han dado las condiciones de verdad de las oraciones, quedan con ello dados sus significados, encaja del todo con diversas tesis de Quine que ya conocemos. Por lo pronto, toma como unidad semántica básica la oración, lo que sin duda está en el espíritu de la posición de Quine, y aspira a clarificar la función semántica de la oración dentro del sistema lingüístico. En segundo lugar —piensa Quine— suministra una justificación adicional para la regimentación lógica del lenguaje en térmi-

nos veritativo-funcionales, pues tal regimentación lleva a cabo los procedimientos recursivos que exige la definición tarskiana. Finalmente, y por citar sólo puntos que hemos tratado con anterioridad, la teoría de Davidson respeta la indeterminación de la traducción radical, puesto que la verdad, dicho en palabras de Quine, «es ella misma inmanente al esquema conceptual» (*loc. cit.*, p. 334).

La teoría de Davidson, por consiguiente, propone sustituir la expresión intensional «significa que» por la expresión extensional «es verdadera si y sólo si». Se advertirá que semejante propuesta está en el espíritu del *Tractatus*, del cual constituye una generalización aprovechando la definición tarskiana. Ahora la cuestión fundamental es: ¿nos da la nueva expresión todo cuanto aspirábamos a obtener de la expresión sustituida? De aquí sale justamente la objeción más importante que se le ha formulado a Davidson: que una afirmación extensional de la forma:

(T) *X* es verdadera si y sólo si *p*.

dice menos de lo que dice una afirmación intensional de la forma:

(S) *X* significa que *p*

La idea de la objeción es que un enunciado de la forma (T) para una oración determinada tiene menos fuerza interpretativa que un enunciado de la forma (S) para dicha oración (Foster, «Meaning and Truth Theory», p. 10). La razón básica es que para que un enunciado tenga efectivamente la forma (T), se requiere que «*p*» sea la traducción de la oración designada por «*X*». Con esta condición, lo que un enunciado de tipo (T) afirma es que, siendo el mundo como es, la oración designada por «*X*» es verdadera si y sólo si *p*. Pero esto no es, en rigor, dar las condiciones de verdad de la oración designada por «*X*»; para hacer esto último tenemos que decir en qué situaciones, de todas las posibles, sería verdadera dicha oración. Sólo esto nos ofrece las condiciones que son necesarias y suficientes para la verdad de la oración cuya descripción estructural viene dada por «*X*». Tan sólo esto, por consiguiente, nos podría dar el significado de la oración en cuestión (y esto parece ser lo que Wittgenstein quiere decir en el *Tractatus*). En conclusión: que los enunciados de la forma (T) son insuficientes para suministrar el significado de las oraciones de un lenguaje, pues no proporcionan, propiamente, las condiciones de verdad de las mismas. Si se quiere formular la teoría del significado sobre la teoría de la verdad, se requiere un instrumento más poderoso que la convención (T) de Tarski. De aquí que Foster (*op. cit.*, pp. 31-32) proponga reformular la teoría del significado en los siguientes términos: una teoría del significado para un lenguaje *L* consiste en la construcción de una variedad apropiada de mundos que agoten las posibles situaciones (o estados de cosas, dicho con expresión del *Tractatus*) permitidas por nuestro punto de vista filosófico; más un conjunto finito de axiomas de los que se deduzca, para cada

oración *O* de *L*, la reformulación canónica relevante obtenida a partir del esquema siguiente: «Para todo *m*, *X* es verdadera-en *m* si y sólo si, si se diera *m*, ocurriría que *p*», siendo «*X*» la descripción estructural de *O*, «*p*» la traducción canónica de *O* al metalenguaje, y «*m*» cualquiera de los mundos posibles contruidos. Como puede apreciarse, aquí no se predica la verdad en general, sino tan sólo la verdad para mundos posibles.

La teoría, así reformulada, pretende dar cuenta de lo que hay que saber para poder interpretar correctamente las expresiones de un hablante. Davidson ha reconocido que, en su primera versión, su teoría era insuficiente a este respecto, y ha dedicado sus esfuerzos en los últimos años a enmendarla de forma adecuada («Reply to Foster», pp. 33 y ss.). El centro de su interés está ahora en la noción de traducción a la que se acaba de hacer referencia. Por trabajar con lenguajes formales, y estando primordialmente interesado en el concepto de verdad, Tarski podía dar por supuesto el concepto de traducción. Pero esto no es posible en la teoría de la interpretación radical, como se conoce a la teoría de Davidson, en la que más bien hay que asumir el concepto de verdad y justificar el de traducción. Esto no implica, según Davidson, abandonar la convención (T), sino leerla de una nueva manera (*loc. cit.*, p. 35). La teoría de Davidson aspira a satisfacer la convención (T), añadiendo a la vez restricciones empíricas que justifiquen la traducción de la oración designada por «*X*» por la oración representada como «*p*». Si la teoría propugnada por Davidson consigue esto, podremos entonces aceptar que, al mismo tiempo, da cuenta del saber que caracteriza al intérprete de un lenguaje, o, lo que es lo mismo, de lo que es suficiente saber para entender un lenguaje. Y una teoría que pase tal prueba empírica es —piensa Davidson— una teoría que puede ser proyectada a cualesquiera situaciones contrafácticas o mundos posibles, y esto ha de resultarle claro a cualquiera que conozca el carácter de la teoría y sepa cómo puede confirmarse por la experiencia (*loc. cit.*, p. 36).

La teoría pretende, pues, explicar eso que intuitivamente y de forma oscura llamamos «significado», tomándolo en la acepción de «sentido», según viene utilizándose este término desde Frege y el *Tractatus*. Pero al recurrir al concepto de verdad, la teoría obliga a pasar por alto ciertos matices semánticos como los que distinguen a «y» de «pero», o los que caracterizan a las diferentes fuerzas ilocucionarias (cfr. Foster, *op. cit.*, nota 9). Ello no quiere decir que no haya forma alguna de acomodar estas características del significado en una teoría del tipo de la teoría de Davidson. En realidad, son dificultades propias de una teoría que no está más que esbozada, y cuya aplicación concreta está por hacer. Algunos discípulos de Davidson, conscientes de la importancia de unir la teoría veritativa del significado con la teoría de las fuerzas ilocucionarias, han reformulado la primera de tal manera que dé cabida a esta última. Así, si consideramos el enunciado de partida «En el lenguaje *L*, la oración *X* significa que *p*», Platts, por ejemplo, propone definirlo del modo siguiente (*Ways of Meaning*, p. 67): hay, para *L*, una teoría *T* de la verdad tal que, primero, es un teorema de *T* que *X* es verdadera si y sólo si *p*, y segundo, las afirma-

ciones de esta teoría se combinan con una teoría aceptable de la fuerza ilocucionaria, así como con el comportamiento observado, lingüístico y no lingüístico, de forma que permita la adscripción de actitudes proposicionales aceptables a los hablantes de *L*. Se reconocerá que esta formulación es excesivamente vaga, pero vale como síntoma de una voluntad de integración, motivada, acaso, por las propias limitaciones de toda teoría que intente construir el concepto de significado sobre el concepto de verdad. Limitaciones contra las cuales se viene debatiendo la teoría del significado desde el *Tractatus* en la dirección que estamos estudiando en este capítulo.

8.10 El significado como función

Al examinar ciertas objeciones a Davidson, hemos visto cómo se echaba mano del concepto de mundo posible a fin de formular las condiciones de verdad de una oración. Un mundo posible no es, en el presente contexto, más que una situación, bien real, bien imaginaria o contrafáctica, esto es, contraria en algún aspecto a los hechos, a lo realmente ocurrido. La noción está directamente relacionada con el concepto de posibilidad tal y como Wittgenstein lo utiliza en el *Tractatus* (2.012 y ss.). En esta sección trataremos de una teoría reciente en la que se apela al concepto de mundo posible a fin de explicar el concepto de significado, así como otras nociones semánticas. El caso que consideraremos es el de Lewis. Por lo que respecta a la función del concepto de verdad en la teoría semántica se halla bien próximo a Davidson. Para Lewis, lo primero que hay que saber acerca del significado de una oración es en qué condiciones sería verdadera; y añade: «La semántica que carece de un tratamiento de las condiciones de verdad no es semántica» («General Semantics», 1970, p. 18.)

Lewis empieza por construir una gramática categorial, que es una gramática de estructura sintagmática libre de contexto con las siguientes características. Se parte de unas pocas categorías básicas, y las demás categorías, o categorías derivadas, se definen a base de las primeras. En rigor, bastan dos categorías, por ejemplo, nombre propio y oración, como es sabido desde que Ajdukiewicz construyó este tipo de gramáticas en 1935. Supongamos que aceptamos como categorías básicas las de oración (*O*), nombre propio (*N*) y nombre común (*C*). Podemos representar entonces el sintagma verbal como *O/N*, que quiere decir: la categoría que, combinada con un nombre propio (*N*) da como resultado una oración (*O*); un adjetivo será lo que, combinado con un nombre común (*C*), da otro nombre común (*C*), o sea, *C/C*; y así sucesivamente, aunque de forma más compleja, para otras categorías (*op. cit.*, secc. II).

Siguiendo las directrices que ya había marcado Montague con anterioridad, Lewis considera que el significado de una oración será aquello que determina las condiciones en las que la oración es verdadera o falsa; esto es, lo que determina el valor de verdad de la oración en diferentes situaciones, lugares, tiempos, y para diversos hablantes. El significado de un

nombre propio será lo que determina qué es lo que aquél nombra, si es que nombra algo, en las diferentes situaciones, tiempos, etc. (y naturalmente, lo nombrado puede no existir realmente, pero sí en alguna situación posible, como ocurre con los entes de ficción). De modo parejo, el significado de un nombre común será aquello que determina a qué cosas (reales o posibles) se aplica éste en las diferentes situaciones, tiempos, etc. Con ello tenemos una primera aproximación al significado de nuestras tres categorías básicas. Si, al modo de Carnap, llamamos «extensión» de una oración a su valor veritativo; «extensión» de un nombre propio al objeto nombrado por él, y «extensión» de un nombre común a la clase de los objetos a los que éste se aplica, tendremos que el significado es, en principio, aquello que, para cada una de esas categorías, determina su extensión en cada caso. Pero no la determina por sí solo. La verdad o la falsedad de una oración, por ejemplo, no depende exclusivamente de su significado, sino además de la situación en la que se emplee, de quien la pronuncie y en qué tiempo y lugar, etc. Más bien tendremos que decir, por ello, que el significado determina de qué manera depende la extensión de esos otros factores. El significado, según esto, equivale a una función que, teniendo como argumento el conjunto de esos factores, tiene como valor la extensión. Llamando «índice» a cada conjunto de tales factores, los significados, así concebidos, serán funciones de índices a extensiones. Para denominar a los significados en esta teoría, Lewis propone resucitar el viejo término carnapiano de «intensión».

Aunque el parentesco entre la teoría de Lewis y la de Carnap salta a la vista, el concepto de intensión no coincide en ambos totalmente. Para Carnap, las intensiones tenían como argumentos descripciones de estado, que podemos tomar como mundos posibles (recuérdese la sección 8.4). En Lewis, las intensiones toman como argumentos conjuntos de diversas características relevantes para la determinación de la extensión (conjuntos a los que acabamos de llamar «índices»). Estos índices incluyen como uno de sus elementos un mundo posible o descripción de estado. Un mundo posible es una posible totalidad de los hechos determinada en todos sus aspectos, pero únicamente algunos de éstos serán relevantes para determinar la extensión. A la vista de las diferentes clases de términos deícticos que poseen las lenguas, tales aspectos incluirán el tiempo, el lugar, el hablante, el oyente, etc. De acuerdo con la extensión propia de cada categoría, diremos que la intensión de una oración es una función de índices a valores de verdad; que la intensión de un nombre propio es una función de índices a objetos individuales; y que la intensión de un nombre común es una función de índices a conjuntos de objetos (*op. cit.*, secc. III). Las intensiones así entendidas no equivalen del todo a los significados intuitivamente considerados, pues pueden darse —según Lewis— diferencias de significado en el discurso ordinario que no vayan acompañadas por diferencias de intensión. Lewis sugiere el caso de las tautologías: la intensión sería la misma para todas ellas, ya que es la función constante que, para todo índice como argumento, tiene como valor

la verdad. ¿Diríamos, entonces, que todas las tautologías tienen el mismo significado? Lewis piensa que sería absurdo afirmarlo así. Pero téngase en cuenta que las tautologías son en rigor proposiciones compuestas y que, por tanto, siempre podríamos descomponerlas en las proposiciones simples correspondientes. Bien es cierto que este recurso no es aplicable si tomamos en consideración verdades lógicas no estrictamente tautológicas, como las verdades cuantificacionales, y las de la lógica de la identidad (por ejemplo, «Todo objeto es idéntico a sí mismo»). Pero cabe dudar, fundamentalmente, que en el discurso ordinario reconozcamos significado a enunciados de ese tipo. En definitiva, no creo que haya aquí ningún problema grave para la teoría de Lewis.

Otra cuestión a resolver es la que se plantea cuando, para un determinado índice, la oración carece de valor veritativo o el nombre propio no nombra nada. Es el caso de las oraciones acerca de entidades inexistentes, que Russell declaró falsas en su teoría de las descripciones, y que Strawson prefirió declarar carentes de valor veritativo. Lewis opta por considerar que, en tales casos, la intensión es una función no definida. Así, para un índice que contenga como mundo posible nuestro mundo real, el nombre «Pegaso» tendrá una intensión indefinida, e igualmente por lo que toca a una oración como «El presidente de la República española en 1980 era médico» (pero naturalmente hay mundos posibles en los que esta oración será verdadera, y otros en los que será falsa, así como mundos posibles en los que «Pegaso» nombrará algún objeto). Por supuesto que no diríamos que esta oración carece de significado, y esto muestra que, para casos como el presente, la teoría de Lewis deja el significado y la intensión muy apartados uno de otro. No obstante, acaso la situación tampoco sea grave si pensamos que el significado de esa oración tal vez pueda entenderse como lo que nos permite imaginar en qué condiciones, esto es, en qué mundos posibles, esa oración sería verdadera o falsa. Para la determinación de su significado sólo serán, entonces, relevantes aquellos índices que contengan mundos en los que la oración sea verdadera o falsa, esto es, mundos en los que España fuera una República en 1980.

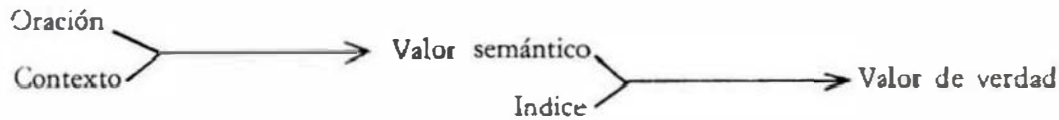
Queda, por fin, la dificultad de las oraciones no declarativas que, en principio, y aparentemente, no poseen valores de verdad. La propuesta de Lewis, a fin de reducirlas a su teoría, es considerarlas como paráfrasis de las correspondientes oraciones realizativas. El caso más típico, y más extremo, es el de las oraciones imperativas; por ejemplo: «Escúchame». Para Lewis, sería una forma de expresar lo mismo que dice la oración: «Te ruego que me escuches»; o tal vez: «Te mando que me escuches». Según esto, parece que los imperativos serían ambiguos, según qué acto ilocucionario se realizara por medio de ellos en cada ocasión. La ambigüedad quedaría eliminada al expresarlos en la forma realizativa. O tal vez podríamos pensar que hay una forma realizativa común a todos ellos, únicamente que carecemos en nuestras lenguas de una palabra que la exprese; podríamos recurrir a alguna palabra que realizara esa función, por ejemplo: «imperar», y reformular así el imperativo citado: «Te impero

que me escuches», que serviría tanto si se trata de un ruego como si se trata de un mandato. En todo caso, la idea de Lewis es que las oraciones realizativas deben ser consideradas como declarativas y, por consiguiente, atribuirles valores de verdad. Según esto, la diferencia entre oraciones declarativas y no declarativas sería puramente sintáctica, y se daría en el nivel de la estructura superficial, pero no en el de la estructura profunda (*op. cit.*, secc. VIII).

La propuesta de Lewis presenta, sin embargo, inconvenientes. ¿Cuándo es verdadera y cuándo es falsa una oración realizativa? Si aplicamos la convención (T) al ejemplo anterior, obtendremos el enunciado siguiente: «La oración 'Te ruego que me escuches' es verdadera si y sólo si te ruego que me escuches». Esto quiere decir que una oración realizativa es verdadera cuando efectivamente se realiza el acto que la oración expresa. O, lo que es lo mismo: que una oración realizativa, pronunciada en serio, sinceramente, y en las condiciones normales del discurso ordinario, es siempre verdadera. ¿Cuándo será falsa? Cuando el hablante hable en broma, o esté representando una obra de teatro, o practicando la elocución, o haciéndose pasar por otra persona, etc. Esto da lugar a una curiosa asimetría entre las oraciones declarativas y las no declarativas. Pues considérese este ejemplo de oración declarativa: «Está lloviendo». Esta oración es verdadera si y sólo si está lloviendo en el lugar y en el momento a los que haga referencia. Supongamos que alguien la pronuncia equivocadamente cuando no está lloviendo. La oración es, en tal caso, falsa. Pero toda oración declarativa puede también ser parafraseada por medio de una oración realizativa correspondiente. Supongamos entonces que el hablante hubiera hecho su afirmación en esta forma: «Te informo de que está lloviendo». Tendremos entonces que decir que esta oración, así usada, es verdadera, aun cuando no lo sea la oración «Está lloviendo» contenida en ella. En resumen: las oraciones declarativas tendrán dos clases de condiciones de verdad, según que se formulen en su forma realizativa o en su forma no realizativa, mientras que las oraciones no declarativas solamente tendrán una clase de condiciones de verdad, las de su enunciación en forma realizativa. Semejante asimetría me parece antiintuitiva y falta de justificación. Por ello consideré más ajustado al uso del lenguaje ordinario distinguir, en la sección 7.8, diferentes tipos de discurso sobre la base de diversos valores semánticos, reservando la verdad y la falsedad para las oraciones declarativas.

Debido a la dificultad que entraña especificar, para cada uso de una oración, todos los factores relevantes que deben entrar en los índices, Lewis ha reformulado su teoría posteriormente de diversas maneras (*Index, Context and Content*, 1981). En la que me parece preferible, toma como argumento inicial la oración y el contexto en el que se usa; la gramática será la función que, para este par de factores como argumento, tome como valor lo que llama Lewis «valor semántico», esto es, un contenido proposicional; añadiéndole a este último un índice (en el sentido que ya

conocemos), ambos, índice y proposición, serán a su vez el argumento de una función cuyos valores son valores de verdad. En esquema:



De esta manera se separan el contexto de uso de la oración y el conjunto de factores (índice) que hace a la oración verdadera o falsa. Como se ve, el precio es una mayor complejidad en la teoría (este tipo de complejidad había sido anticipado por David Kaplan y por Stalnaker, cuyas sugerencias recoge Lewis). Parece que la nueva formulación permite expresar con más claridad la diferencia entre una oración que es verdadera en todo contexto para un determinado índice o mundo posible, y una oración que es verdadera en un contexto determinado para todo mundo o índice posible. Así, presumiblemente, la oración «El agua hierve a 100° al nivel del mar» es verdadera en unos mundos pero no en otros; ahora bien, en los mundos en los que es verdadera, lo es en todos los contextos. Por el contrario, la oración «Yo soy el hijo de Napoleón» es verdadera tan sólo en el contexto en que es pronunciada por alguien que sea efectivamente hijo de Napoleón (por ejemplo, el rey de Roma), pero será verdadera en todo mundo posible en el que sea así pronunciada, ya que ser hijo de los padres propios es una condición que no se puede dejar de tener en cualquier situación posible para seguir siendo uno mismo (sobre este punto, véase la sección 8.11).

Lo anterior está, como se habrá apreciado, en la línea Tarski-Carnap, y considera el lenguaje como sistema abstracto, posiblemente común a todas las lenguas humanas. Lewis ha distinguido cuidadosamente entre este enfoque y una consideración psicosociológica del lenguaje, abordando en otros lugares esta segunda perspectiva. Es curioso que también aquí el concepto de verdad tiene un lugar central. Como resumen, y sin entrar en detalles, baste señalar que, para Lewis, la convención básica que rige el uso de un lenguaje en una comunidad es una convención de veracidad (*Convention*, V.4). Es la veracidad uniforme y general de los hablantes la que permite la comunicación con éxito por medio de un lenguaje común. Esto no significa que nadie en absoluto mienta. Quiere decir simplemente que, en general, casi todos los hablantes procuran casi siempre decir lo que creen que es verdad, porque asumen que los demás así lo hacen, y porque todos saben que los demás cuentan con que cada uno así lo hará, y esto es lo que prefieren a fin de coordinar sus acciones lo mejor posible. Aquí puede plantearse de nuevo el caso de los imperativos: ¿cómo se les aplica la convención de veracidad? La respuesta de Lewis es totalmente insatisfactoria: la aplica el oyente con su comportamiento intentando cumplir el imperativo y, en esta medida, convertirlo en verdadero. Reaparece aquí la asimetría entre las oraciones declarativas y las imperativas, pues mien-

tras para las primeras la convención obliga al hablante, para las segundas obliga al oyente. Una asimetría aún más llamativa que la que hemos encontrado hace un momento, pues muestra que la convención de veracidad sirve para explicar el uso que hace el hablante de las oraciones declarativas, pero no el que hace de las imperativas; en este último caso lo único que explica es el comportamiento de aquel a quien se dirige el imperativo. En todo caso, lo único que me importaba subrayar es la función central que el concepto de verdad puede tener incluso para un enfoque psicosociológico que, como tal, se ocupe de la convención fundamental que rige el uso del lenguaje.

El concepto de mundo posible forma parte también de la teoría de Montague sobre el lenguaje. Su carácter altamente técnico y la complejidad de los instrumentos lógicos que utiliza la sitúan en un nivel muy por encima de cuanto hemos estudiado en la presente obra, por lo que no creo que sea recomendable ni tan siquiera intentar ofrecer un esbozo de ella (el lector interesado puede consultar el artículo de Daniel Quesada, «Lógica y gramática en Richard Montague»). La posición básica de Montague es que no hay diferencia teórica sustancial entre los lenguajes naturales y los lenguajes artificiales (esta declaración abre dos de sus trabajos más característicos, «Gramática universal» y «English as a Formal Language», ambos de 1970). Sobre esta base, y con vistas a llegar a una teoría de la verdad para el lenguaje natural, Montague desarrolla un metalenguaje lógico extremadamente complicado en el que recoger todos los aspectos, tanto sintácticos como semánticos y pragmáticos, del lenguaje (o fragmento de lenguaje) estudiado. El considerable mérito de Montague se encuentra no sólo en el complejo aparato formal que pone en funcionamiento, sino más aún en haberlo aplicado a fragmentos concretos del inglés, intentando mostrar cómo funciona su teoría y a qué resultados puede llegar. En qué medida estos resultados ayuden efectivamente a entender la estructura y el funcionamiento de un lenguaje natural presumo que es todavía disputable. Aunque Montague ofrece su teoría como una alternativa a la lingüística transformatoria, hay lingüistas que piensan que sus conclusiones tienen poco o nada que ver con el lenguaje natural y que tan sólo sirven para lenguajes artificiales contruidos a propósito. Frente a todos los autores que hasta aquí hemos considerado, Montague ciertamente representa el intento más ambicioso de aplicar técnicas formales al estudio de las lenguas.

8.11 Una teoría de los nombres propios

La necesidad de estipular condiciones de verdad para las oraciones no sólo en relación al mundo real, sino también a situaciones posibles contrafácticas, recurriendo, por tanto, al concepto de mundo posible, ha sido igualmente destacada por Kripke. En el desarrollo de esta idea, Kripke ha formulado una teoría sobre los nombres propios que, oponiéndose a la

posición de Frege y Russell que ya conocemos, y que ha dominado en la filosofía analítica, vuelve a la posición de Stuart Mill que estudiamos en la sección 6.2, incidiendo además en otros conceptos como el de verdad necesaria.

Frege pensaba que el sentido de un nombre propio viene dado por cualquier descripción definida verdadera acerca del objeto al que se refiere el nombre. Según esto, los nombres propios, así llamados en el lenguaje ordinario, tienen sentido, si bien tal sentido puede diferir de unos hablantes a otros. Russell, más aún, sostenía que los llamados nombres propios no son más que maneras de abreviar descripciones, y que un auténtico nombre propio en sentido lógico tendría que ser una palabra totalmente carente de sentido, una palabra que se limitara a denotar el objeto sin decir nada acerca de él, una palabra, en suma, del tipo de los pronombres demostrativos. A diferencia de ambos, Mill había mantenido que los nombres propios del lenguaje ordinario, aun cuando muchas veces posean sentido, o, como él lo llama, «connotación», no la requieren para funcionar efectivamente como nombres propios, y que para serlo les basta tener denotación o referencia. A esta concepción vuelve Kripke por las razones que vamos a ver.

A los nombres propios y a las descripciones definidas llama Kripke «designadores» («Naming and Necessity», 1972, p. 254). Entre éstos distingue los designadores rígidos de los designadores no rígidos o accidentales: un designador rígido es aquel que designa o denota el mismo objeto en cualquier mundo posible; en otro caso, se trata de un designador accidental (*op. cit.*, p. 269). Con esto queda introducido el concepto de mundo posible. ¿De qué se trata? Ya lo hemos visto brevemente y al paso en las páginas precedentes. Un mundo posible no es más que un posible estado o situación del mundo, que puede ser naturalmente el mundo real, y además cualquier otra situación contrafáctica imaginable, esto es, contraria a los hechos reales en algún aspecto. Por consiguiente, un mundo posible no es algo que descubramos, sino algo que estipulamos; en palabras de Kripke: «un mundo posible está dado por las condiciones descriptivas que asociamos con él» (*op. cit.*, p. 267). Siempre que decimos: «Supongamos que...», construimos un mundo posible, al menos con tal que lo que suponemos sea efectivamente posible. «Supongamos que las elecciones francesas de 1981 las hubiera ganado Giscard...», «Supongamos que me tocara el premio mayor de la lotería en el próximo sorteo...», constituyen construcciones de mundos posibles. Por supuesto, no aceptaremos como mundo posible toda estipulación: «Supongamos que hubiera cuadrados redondos...», «Supongamos que el agua mojara y no mojara al mismo tiempo...», no suministran condiciones para un mundo posible; de hecho, el término «suponer» es incorrecto en semejantes afirmaciones. Pues bien, Kripke mantiene que los nombres propios son designadores rígidos, esto es, que un nombre propio se refiere al mismo objeto en cualquier mundo posible. Es precisamente porque podemos hablar con sentido de lo que le podría haber ocurrido a Giscard, a Aristóteles, al Partenón, a España...,

por lo que cada uno de estos nombres designa un objeto que es el mismo en todo mundo posible. Se entiende: en todo mundo posible en el que exista el objeto. No se pretende que en todo mundo posible haya de existir Giscard, Aristóteles, el Partenón o España. La madre de Aristóteles podría haber muerto antes de darlo a luz, los griegos podrían haber decidido no construir el Partenón, los movimientos de la corteza terrestre podrían haber sido tales que no se formara la Península Ibérica, o los padres de Giscard podrían no haber tenido hijos. Pero en todos aquellos estados del mundo en los que cualquiera de estos objetos exista, debe tratarse del mismo objeto, pues si no, ¿cómo sabemos que estamos hablando de lo que le podría haber ocurrido a Giscard, al Partenón, etc.?

Tenemos, pues, que los nombres propios son designadores rígidos. Pero no sólo ellos. También pueden funcionar como designadores rígidos los pronombres demostrativos y personales. Si digo «Este es demasiado grande» refiriéndome a un cierto objeto, «éste» denotará ese objeto en cualquier mundo en el que el objeto exista siempre que me refiera a él por medio de tal pronombre; pues yo podría también decir: «Ojalá éste hubiera sido un poco más pequeño». Por lo mismo, «yo», en cuanto empleado por mí, se refiere a mí en cualquier mundo posible; por eso puedo hacer afirmaciones tales como: «Yo podría haberme dedicado a la abogacía». El hecho de que la referencia de estos pronombres varíe según quién, dónde y cuándo los usa, no impide que funcionen como designadores rígidos. Incluso, y para objetos indeterminados, también las variables libres de un cálculo funcionan como designadores rígidos (*op. cit.*, nota 16, p. 345). En contraste, las descripciones definidas son, *en principio*, designadores accidentales, puesto que una descripción se refiere a un objeto dando de él ciertas características que lo identifican, pero cabe suponer que, en otros mundos posibles, tal objeto podría carecer de esas características sin dejar de ser ese objeto. Así, «el preceptor de Alejandro Magno» no se refiere rigidamente a Aristóteles, ya que podemos imaginar una situación en la que Aristóteles no hubiera sido el maestro de Alejandro Magno. Esto prueba que las descripciones no dan el sentido de un nombre propio y, por consiguiente, que no son parte de su significado; ninguna de las descripciones verdaderas que podamos hacer de un objeto es parte del significado de su nombre, ni cada una de ellas por separado ni la disyunción de todas ellas, pues, en principio, toda descripción recurrirá a características que el objeto podría no haber tenido sin dejar de ser tal objeto. El sentido de «Aristóteles» no viene dado ni siquiera por la descripción «la persona llamada 'Aristóteles'», pues Aristóteles, el filósofo griego que fue discípulo de Platón y preceptor de Alejandro Magno, además de no haber sido nada de esto, podría igualmente no haberse llamado «Aristóteles». Es decir: que el nombre «Aristóteles» designa rigidamente el mismo objeto en todos los mundos posibles en los que ese objeto exista, incluso aun cuando en algunos de ellos no se llame «Aristóteles».

He dicho hace un momento que «en principio» toda descripción recurrirá a características que el objeto podría no haber tenido sin dejar de

ser tal objeto. La inserción de la frase «en principio» se debe a lo siguiente. Parece que todo objeto debe tener algunas características constantes en todos los mundos posibles si efectivamente se trata del mismo objeto en todos ellos. Esas características serán entonces esenciales para el objeto, y así lo acepta Kripke (*op. cit.*, pp. 272 y 276). Si resulta que en nuestra descripción del objeto recurrimos a características esenciales, entonces nuestra descripción será un designador rígido, pues denotará el mismo objeto en todos los mundos posibles. Es lo que acontece, por ejemplo, con la descripción «la razón de la circunferencia al diámetro», que designa en todo mundo posible el mismo objeto, a saber: el número real que designa también el nombre « π ». Por consiguiente, no todas las descripciones definidas son designadores accidentales; algunas designan rígidamente, a saber: cuando recurren a propiedades esenciales. No obstante, ni siquiera en estos casos podemos tomar la descripción como el sentido de un nombre. Pues supongamos que describiéramos a Aristóteles por medio de alguna propiedad esencial. Teniendo en cuenta que llamarse «Aristóteles» no es una propiedad esencial, nuestra descripción seguiría designando a Aristóteles incluso en un mundo en el que no se llamara «Aristóteles» y, por lo tanto, no hay ninguna justificación para ver en tal descripción el sentido del nombre «Aristóteles».

Podemos suponer que las propiedades esenciales son necesarias para identificar a un objeto en todos los mundos posibles en los que exista. En rigor, según Kripke, no es así. Aunque admitamos que un objeto tiene propiedades esenciales, sin las cuales no sería el que es, no necesitamos recurrir a ellas para identificarlo a través de mundos posibles. La identificación es posible justamente porque podemos usar un designador rígido para referirnos a él. Porque tenemos a Aristóteles identificado en el mundo real, y porque tenemos un nombre para referirnos a él (aunque cualquier otro tipo de designador rígido cumpliría la misma función), podemos hablar de él en cualquier posible situación contrafáctica especulando sobre lo que pudo o no haberle ocurrido. El uso de un designador rígido no requiere la identificación transmundana del objeto, sino que la hace posible. Lo que sí se requiere para que sea posible la designación rígida de un objeto en mundos posibles es: primero, que el objeto exista en esos mundos, pues podría no existir en algunos de ellos (por ejemplo, podríamos imaginar múltiples situaciones posibles en las que Aristóteles no hubiera existido), y segundo, que utilicemos el lenguaje para hablar de esos mundos posibles con el mismo sentido y referencia con que lo usamos para hablar del mundo real. Es patente que podríamos imaginar un mundo en el que no existiera el nombre «Aristóteles», pero para estipular si en ese mundo existe o no Aristóteles o qué le haya de acontecer en él, tenemos que usar el nombre «Aristóteles» como lo usamos actualmente. Dicho de otra forma: podemos suponer un mundo posible en el que el castellano fuera muy distinto de como es actualmente, pero para hablar actualmente de esa posible situación tenemos que hacerlo en el lenguaje tal y como es actualmente. Que «Aristóteles» designa el mismo objeto en todos los mundos

posibles en que tal persona exista no significa que en todos esos mundos posibles haya de existir la palabra «Aristóteles» (*op. cit.*, p. 289).

Si no es necesario identificar el objeto por sus propiedades a fin de referirnos a él, ¿cómo podemos justificar que efectivamente nos estamos refiriendo a él? En el mejor de los casos, la justificación puede venir de que hayamos tenido ocasión de conectar el nombre con el objeto por un procedimiento ostensivo. En defecto de esto, puede ocurrir que sepamos sobre el objeto lo suficiente para identificarlo. Son, como se habrá notado, los dos tipos de conocimiento que destacaba Russell, conocimiento por familiaridad y conocimiento por descripción. Pero el caso que acabamos de plantearnos es aquel en el que no se da ninguno de ambos supuestos, es decir, es el peor de los casos. Pues bien, aquí, según Kripke, semejante justificación viene dada por el proceso de la comunicación lingüística. Yo puedo usar un nombre para referirme a aquello que tal nombre designa aun cuando no sepa del objeto designado lo suficiente para identificarlo; mi justificación será entonces que he oído usar ese nombre para denotar tal objeto, y que mi intención al usarlo es asimismo referirme a dicho objeto (pp. 299 y ss.). Es la cadena comunicativa la que va pasando el nombre de hablante a hablante. Por ejemplo, si me piden que cite un poeta griego del siglo xx, puedo dar el nombre de «Cavafis», con la intención de referirme a aquel que tal nombre designa, aun cuando yo no sepa sobre él nada que pueda identificarlo (pues saber que se trata de un poeta griego contemporáneo sin duda no es suficiente para identificarlo). Puedo incluso hacer afirmaciones contrafácticas sobre él: «Supongamos que Cavafis nunca hubiera escrito poesía...». Lo que cuenta es mi intención de referirme a aquel que designa el nombre «Cavafis» quienquiera que sea; naturalmente, yo asumo que el nombre se refiere a alguien puesto que lo he oído, o lo he leído, usado de forma coherente y en contextos normales.

De la doctrina anterior extrae Kripke consecuencias para el concepto de verdad necesaria. Por lo pronto, respecto de los enunciados de identidad. Si un enunciado de identidad se formula empleando designadores rígidos, de ser verdadero lo será necesariamente; pero si emplea designadores accidentales, entonces será contingentemente verdadero. Así, el enunciado «Azorín es José Martínez Ruiz» es necesariamente verdadero, puesto que ambos nombres designan el mismo individuo en todos los mundos posibles y, por consiguiente, no hay ninguna situación en la que ese enunciado pudiera ser falso. Alguien podría objetar: por el contrario, ese enunciado sería falso para un mundo en el que Martínez Ruiz nunca hubiera decidido adoptar el seudónimo «Azorín». Pero esto es erróneo, ya que, incluso para ese posible mundo, nuestro enunciado sigue siendo verdadero, por cuanto constituye un supuesto de la teoría que las palabras se usan con la referencia que poseen en nuestro mundo actual. En cambio, si afirmamos «Azorín es el autor de *Castilla*», o «El autor de *Doña Inés* es el escritor de la generación del 98 nacido en Monóvar», éstos son enunciados de identidad contingentemente verdaderos, pues Azorín podría no haber nacido en Mo-

nóvar, no haber sido escritor o, de serlo, no haber escrito ninguna de esas obras.

Naturalmente, y por lo que hemos visto antes, también tendremos un enunciado necesariamente verdadero si en lugar de nombres propios empleamos descripciones definidas que recurran a propiedades esenciales, ya que tales descripciones son igualmente designadores rígidos. La dificultad de dar ejemplos con esas descripciones proviene de la oscuridad que reina en torno a la cuestión de cuáles sean las propiedades esenciales. Kripke apenas toca este punto. Para personas, sugiere que una propiedad esencial es el origen. Cabe imaginar que a una persona le hubieran acontecido desde su nacimiento sucesos totalmente distintos de los que realmente le han acaecido, pero ¿cabe imaginar que una persona hubiera nacido de padres diferentes de los reales? Aquí, piensa Kripke (p. 314), estaríamos más bien ante otra persona distinta. En realidad, no basta determinar los padres para determinar la individualidad; en definitiva, todos los hermanos tienen los mismos padres en común, lo que no les impide ser distintos entre sí. Más exactamente, lo que se requeriría es la determinación del espermatozoo y del óvulo de los que se formó la persona en cuestión. El proceder de la unión de tal espermatozoo y de tal óvulo es una propiedad esencial de toda persona, una propiedad que no puede dejar de tener en ningún mundo posible. Para objetos materiales artificiales, como una mesa, Kripke señala como propiedades esenciales, primero, el fragmento de materia con el que se ha fabricado, por ejemplo, tal bloque de madera, o tales troncos; y segundo, la forma de objeto que se le ha dado, ya que es patente que, con esa misma materia, podríamos haber fabricado, por ejemplo, un sillón, una cama, etc. (pp. 314 y nota 57 en p. 351). Por lo demás, esto son tan sólo ejemplos, pues Kripke no se detiene en el tema. Los numerosos e intrincados problemas que habrían de surgir a la menor exploración del tema saltan a la vista.

Curiosamente, Kripke, que a propósito de los nombres propios, y por encima de toda la tradición analítica, vuelve a Mill, difiere de éste en cuanto a los nombres comunes, o al menos en cuanto a ciertos tipos de nombres comunes. Pues los nombres comunes o generales de clases naturales, como las especies animales, «gato», «tigre», etc.; los términos de masa, como «agua», «oro», etc., y ciertos nombres de fenómenos naturales, como «calor», «luz», «sonido», etc., son, para Kripke, designadores rígidos y, por consiguiente, se hallan muy próximos a los nombres propios. La idea es que la referencia de estos nombres se fija por medio de su definición, y que b que ellos designan lo designan en todo mundo posible en el que exista el fenómeno o la clase de objetos en cuestión. Así, si definimos el calor como el movimiento de las moléculas, la palabra «calor» designará ese fenómeno físico en cualquier mundo en el que tal fenómeno se dé; otras características que pueda poseer, tales como ser sentido por los seres humanos de esta manera o de la otra, podrían variar de un mundo a otro, pero no aquello que hace que llamemos a cierto fenómeno «calor». O consideremos el caso de «gato»; con este término designaremos rígida-

mente, esto es, en toda situación posible, aquel animal que tenga lo que es esencial para ser un gato, sea ello lo que fuere. Por supuesto, tal propiedad esencial puede no ser ninguna de las que usualmente asociamos con los gatos y encontramos recogidas en los diccionarios. Por ejemplo: ¿por qué no pensar que los gatos dejaran de ser domésticos? En consecuencia, las afirmaciones «el calor es el movimiento de las moléculas», «la luz es una corriente de fotones», «el agua es un compuesto de dos átomos de hidrógeno por cada átomo de oxígeno», etc., son afirmaciones que, si son verdaderas, lo son necesariamente, pues no pueden dejar de ser verdaderas en ninguna situación posible en las que exista el calor, la luz, el agua, etc. (pp. 315 y ss.). La posición de Kripke asume, claro está, lo siguiente: que si en un mundo posible hubiera una sustancia con todas las propiedades externas del agua, pero con otra composición atómica, esto no sería agua; o que si lo que en un posible mundo nos permitiera la visión fuera algo diferente de la impresión de los fotones, en tal mundo no existiría la luz, sino otro fenómeno que, en todo caso, produciría los mismos efectos. Mill había dicho que los nombres generales poseen connotación, que connotan propiedades de los objetos a los que se aplican. En base a su aplicación del concepto de mundo posible, Kripke objeta: no hay tal, al menos para nombres como los citados; lo único que connota el término «oro» es la propiedad de ser oro; lo único que connota la palabra «calor» es la propiedad de ser calor; lo único que connota el nombre «gato» es la propiedad de ser gato. Etcétera.

Como se habrá observado, las definiciones a las que Kripke recurre en sus ejemplos son definiciones científicas. Asume, en efecto, que son éstas las que nos dan la esencia del fenómeno o de la clase de objetos definidos. Otras propiedades que podamos atribuirles, y que no se deduzcan de la definición, hay que tomarlas como contingentes y, por lo tanto, como propiedades que tales fenómenos u objetos podrían no tener en otros posibles mundos. Ahora bien, si las definiciones científicas expresan la esencia de lo definido y constituyen, por tanto, en caso de ser verdaderas, verdades necesarias, hay que concluir que las verdades descubiertas por la ciencia son, en general, verdades necesarias, y no contingentes como es usual aceptar (pp. 320 y ss.). Esto no significa que las cosas no pudieran ser diversas de como son. Desde luego que podrían no existir los gatos y, acaso, tampoco la luz. Pero una vez que hemos averiguado qué tiene que tener esencialmente algo para ser un gato, o qué tiene que haber esencialmente para que haya luz, no cabe imaginar un mundo posible en el que los gatos o la luz fueran de otra manera. Podría haber objetos muy parecidos a nuestros gatos, o un fenómeno que produjera los mismos efectos que la luz, pero no serían propiamente gatos ni luz.

Lo anterior no implica que tales verdades se conozcan *a priori*. Está claro que tan sólo se alcanzan a fuerza de investigación empírica y son, en consecuencia, verdades *a posteriori*. Pero son al mismo tiempo verdades necesarias, ya que expresan lo que necesariamente ha de tener algo para ser esa clase de objeto, o ese tipo de fenómeno (pp. 330 y ss.). Precisamente uno

de los puntos que Kripke ha querido dejar bien claros es que «*a priori*» y «*a posteriori*» son términos de la teoría del conocimiento, mientras que «*necesario*» y «*contingente*» son términos metafísicos, y que cualquier relación que se pretenda establecer entre ellos hay que justificarla filosóficamente (pp. 260 y ss.). Tenemos, pues, que hay verdades necesarias y *a posteriori*, a saber: las de las ciencias. Pero también hay, según Kripke, verdades contingentes y *a priori*. Por ejemplo: la afirmación de que el metro patrón, esto es, la barra de platino que se conserva en la Oficina Internacional de Pesas y Medidas de Sèvres, mide un metro, es obvio que se conoce *a priori*, puesto que se apela a la longitud de esa barra para fijar la referencia de la expresión «un metro». Pero dicha afirmación es contingente, puesto que es contingente que tal barra tenga la longitud que efectivamente tiene, y cabe suponer que su longitud hubiera podido ser un poco mayor o un poco menor de lo que es realmente (p. 275). En la medida en que tener esa longitud no es una propiedad esencial de la barra, la afirmación de que la barra en cuestión mide un metro es contingente, por más que su verdad nos sea conocida por definición y, por tanto, *a priori*.

La teoría de Kripke que acabamos de considerar en resumen, ha sido objeto de constante debate en los últimos años. Si bien, a diferencia de otras doctrinas que hemos venido examinando en este capítulo, no contiene una tesis general acerca del lenguaje o del significado, muestra, de modo a la vez riguroso y elegante, las consecuencias que, dentro de la teoría del significado, puede tener la introducción de conceptos modales, como los conceptos de necesidad y posibilidad, y, elaborando con atención particular un tema muy concreto como es el de la teoría de los nombres propios, pone de manifiesto la interconexión existente entre la semántica, la metafísica y la teoría del conocimiento, alcanzando sus consecuencias hasta la filosofía de la ciencia. Puesto que la teoría analítica del significado la iniciamos estudiando la teoría de los nombres de Stuart Mill, no habrá resultado impertinente que la cerremos (aunque sea provisionalmente) con una nueva teoría de los nombres que, en los aspectos que hemos visto, enlaza con la de aquél.

Los problemas que encierra la teoría de Kripke son complejos y numerosos, y sin duda es obligado esperar nuevos desarrollos antes de dar una opinión definitiva. En mi juicio, hay una cuestión fundamental que está todavía por aclarar. Ya la mencioné al tratar el tema de la posibilidad en el *Tractatus*, y es ésta: ¿qué es posible? Parece que debemos excluir en principio del ámbito de lo posible todo cuanto encierre contradicción, como que una figura sea a la vez redonda y cuadrada, o que el agua moje y no moje. ¿Pero es posible simplemente todo cuanto no encierre contradicción? ¿Es posible un mundo en el que sean las ondas sonoras las que nos permitan ver? (Esta posibilidad es reconocida por Kripke en la p. 324.) ¿Es posible un mundo en el que el agua no moje? ¿Es posible un mundo en el que un animal tenga toda la apariencia externa de un tigre pero tenga la estructura interna de un reptil? (Kripke reconoce esta posibilidad incluso para el mundo real, p. 317.) ¿Es posible un mundo en el que *El Quijote*

fuera escrito por un niño de cinco años? Aplicado al caso de las personas: ¿es que no hay límites para lo que podemos imaginar que podría haber hecho, o haberle ocurrido a Aristóteles? Dadas su herencia genética y los estímulos recibidos del medio, ¿es igualmente contingente que se dedicara a la filosofía y que se hiciera bandolero? Dada la historia de una persona hasta un momento determinado de su vida, ¿es igualmente contingente todo cuanto a partir de ese momento hace y le ocurre? Dicho en términos más generales: ¿no hay condiciones sobre las descripciones que podemos asociar a un mundo posible? Por otra parte, un mundo posible es una situación o estado total del mundo. Si suponemos un mundo posible en el que Aristóteles no se dedicó a la filosofía, ¿qué más hemos de suponer para que este mundo sea coherente? Es de esperar que, en muchos aspectos, las vidas tanto de Aristóteles como de otras personas habrían sido distintas y, por supuesto, la historia ulterior del pensamiento occidental. ¿Podemos realmente suponer un mundo así? Quiero decir: ¿podemos imaginarlo? Y si no podemos imaginarlo, ¿tiene sentido hablar de él? Son cuestiones todas ellas que, en el mejor de los casos, producen perplejidad.

Otra dificultad que también me parece muy importante es la que se refiere a las definiciones de los nombres de cosas y fenómenos naturales. He señalado, a propósito, que las definiciones que Kripke parece aceptar son las que pueden encontrarse en las ciencias. Esto le permite caracterizar las verdades científicas como verdades necesarias. Pero supongamos que aceptáramos las definiciones fenoménicas más bien que las científicas. Quiero decir: que definiéramos el calor como la sensación que se tiene en tales y cuales circunstancias, por ejemplo, estando desnudo cuando la atmósfera tiene tal grado de humedad y cuando el termómetro marca a partir de tal temperatura ambiente, siendo tal y cual la temperatura del cuerpo. En tal caso, sería perfectamente contingente que el calor es el movimiento de las moléculas, pues cabe pensar que esa sensación podría haber sido producida por otro fenómeno que no fuera éste. Tendríamos un caso análogo si en lugar de definir la luz como una corriente de fotones la definiéramos como aquello que impresiona nuestro sentido de la vista permitiéndonos la visión de los objetos. En tal caso, sería contingente que lo que impresiona nuestro sentido de la vista es una corriente de fotones, pues podría haber sido otro fenómeno físico distinto. La diferencia que hay entre definir ciertos conceptos de una forma o de otra resulta aún más patente cuando, al final de su trabajo, Kripke, tratando el tema de la identidad entre la mente y el cuerpo, considera el caso del dolor. Pues supongamos que es cierto, como parece hoy día, que el dolor es el resultado de la estimulación de las fibras nerviosas de tipo A-delta y C. Si definimos el dolor de esta forma (a saber: «Es el resultado de la estimulación...»), tendremos que tal definición expresa una verdad necesaria, puesto que, en todo mundo posible en el que haya seres con este tipo de fibras nerviosas y reciban estímulos a través de ellas, tendremos que decir que hay dolor. Y esto incluso aunque ocurriera (¿sería posible?) que la estimulación de esas fibras, debido acaso a ciertas particularidades de su

organización cerebral, les produjera a esos seres lo que nosotros denominamos en nuestro mundo «placer». Como esta consecuencia es absurda, Kripke tiene que reconocer que el caso del dolor es diferente, pues para que haya dolor tiene que haber sensación de dolor, ya que producir sensación de dolor es una propiedad esencial de la estimulación de las fibras A-delta y C (en cambio, producir sensación de calor sería una propiedad puramente accidental del movimiento de las moléculas; pp. 339 y ss.). Al margen de que tal distinción entre el dolor y el calor no acaba de estar del todo justificada, lo que no resulta inteligible es por qué no podría haber sido el dolor producido por la estimulación de otro tipo de fibras nerviosas, y por qué no podemos imaginar un mundo en el que existiera la sensación de dolor pero se produjera por mecanismos diferentes de los que la producen en el mundo actual. En suma, no veo qué razones nos impiden afirmar que el enunciado «El dolor es el resultado de la estimulación de fibras nerviosas de tipo A-delta y C» es una verdad contingente.

Téngase en cuenta que esta cuestión de las definiciones es, como vimos en la sección 8.8, directamente relevante para el problema de las verdades analíticas. A estos efectos lo que interesa, como ya sabemos, es la definición léxica, que puede o no coincidir con la definición científica según los casos. De aquella definición depende que califiquemos un enunciado como verdadero en razón de lo que significan sus palabras, o sea, analítico, o como verdadero en razón de la realidad extralingüística y, por tanto, sintético. Desde este punto de vista sugerí en dicha sección que la afirmación de que el dolor es una sensación desagradable es analíticamente verdadera, mientras que el enunciado de que el dolor es producido por la estimulación de tales o cuales fibras nerviosas es una afirmación sintética. Aunque Kripke elude explícitamente entrar en el problema de las verdades analíticas (nota 63, p. 352), su teoría tiene también consecuencias para este tema, ya que propone considerar analíticas aquellas verdades que son a la vez necesarias y *a priori*, en el sentido de su teoría (*ibidem* y nota 21 en p. 346). De acuerdo con la caracterización de las verdades analíticas que propuse en la sección 8.8, tanto su necesidad como su conocimiento *a priori* son triviales, en el sentido de que ambas características vienen dadas simplemente por el lenguaje en el que se formulan las verdades analíticas.

Todo ello, desde luego, no afecta a la tesis de Kripke sobre los nombres propios, con la cual, por otra parte, estoy conforme. Pero acaso la tesis no requiera ni siquiera la introducción del concepto de mundo posible. Pues para justificar que los nombres propios carecen de sentido o connotación, y que son totalmente independientes de descripciones identificatorias, basta mostrar, y creo que de ello es fácil dar ejemplos cotidianos, que los nombres propios se usan en muchos casos sin saber apenas nada del objeto al que se refieren, y desde luego sin saber lo suficiente para identificarlo. Recuérdese el ejemplo acerca del poeta Cavafis.

Lecturas

Los temas y autores tratados en este capítulo han tenido, en comparación con los de los dos capítulos precedentes, menor eco en la bibliografía en castellano. Podrá encontrarse referencias de interés en el libro de Simpson, *Formas lógicas, realidad y significado* (Editorial Universitaria de Buenos Aires, 1964), y están traducidos algunos importantes artículos en su antología *Semántica filosófica: problemas y discusiones* (Siglo XXI, Buenos Aires, 1973). También la *Antología semántica* recopilada por Mario Bunge (Nueva Visión, Buenos Aires, 1960) contiene trabajos muy relevantes para este capítulo, algunos de los cuales mencionaré más abajo.

Para una visión global del neopositivismo puede consultarse con provecho el *Examen del positivismo lógico*, de Weinberg (Aguilar, Madrid, 1959), y *El Circulo de Viena*, de Kraft (Taurus, Madrid, 1966). El libro clásico de Ayer, *Lenguaje, verdad y lógica* (Editorial Universitaria de Buenos Aires, 1965; hay una edición posterior en Martínez Roca, Barcelona), continúa siendo una excelente introducción a los problemas centrales del neopositivismo. La conocida recopilación de Ayer, *El positivismo lógico* (Fondo de Cultura Económica, México, 1965), contiene una muestra muy representativa de los artículos más influyentes.

De Carnap están traducidas las conferencias tituladas *Filosofía y sintaxis lógica*, recogidas en el primer volumen de la recopilación de Muguerza ya citada, *La concepción analítica de la filosofía* (Alianza, Madrid, 1974); una traducción anterior de dichas conferencias fue publicada en 1963 por la Universidad Nacional Autónoma de México. El celebrado y denostado artículo de Carnap «Superación de la metafísica por medio del análisis lógico del lenguaje» se encuentra traducido en la citada recopilación de Ayer, *El positivismo lógico*. Su posterior trabajo «Empirismo, semántica y ontología» se halla en el volumen segundo de la recopilación de Muguerza. Otro trabajo importante de la misma época, «Significado y sinonimia en los lenguajes naturales», se encontrará en la *Antología semántica* de Bunge. Desgraciadamente, ninguna de las grandes obras de Carnap sobre el lenguaje se encuentra traducida al castellano.

También en la *Antología semántica* de Bunge está recogida la traducción del artículo donde Tarski hace más fácil y accesible su concepción semántica de la verdad: «La concepción semántica de la verdad y los fundamentos de la semántica». Su trabajo original, «El concepto de verdad en los lenguajes formalizados», que es sumamente técnico, se encontrará en traducción inglesa en la recopilación de sus trabajos *Logic, Semantics, Metamathematics* (Oxford University Press, 1956); y en traducción francesa, en la recopilación, *Logique, sémantique, métamathématique* (Armand Colin, París, 1972), que es una recopilación más amplia que la inglesa.

Para Quine, en cambio, contamos con traducción de casi todos sus escritos. Su obra fundamental acerca del lenguaje, *Palabra y objeto*, está publicada por Labor (Barcelona, 1968). Su primera y más conocida recopilación de artículos, *Desde un punto de vista lógico*, está publicada en

Ariel (Barcelona, 1963), y contiene su famoso artículo «Dos dogmas del empirismo». Otra recopilación traducida es *La relatividad ontológica y otros ensayos* (Tecnos, Madrid, 1974); casi todos los trabajos aquí incluidos son posteriores a *Palabra y objeto* y tienen con este libro conexión directa. El último libro de Quine sobre problemas del lenguaje lleva por título *Las raíces de la referencia* (Revista de Occidente, Madrid, 1977) y continúa y amplía algunos de los temas de *Palabra y objeto*. La revista *Teorema* ha dedicado a Quine un interesante número monográfico con el título *Aspectos de la filosofía de W. V. Quine* (Valencia, 1976), donde a los comentarios y críticas de los colaboradores siguen las respuestas de Quine.

El resto es silencio. (*Hamlet*.)

¿No depende todo de nuestra interpretación del silencio que nos rodea? (LAWRENCE DURRELL, *Justine*.)

En las últimas secciones hemos visto sugerencias, más o menos explícitas, en favor de una teoría del significado que integre la teoría de los actos verbales junto con la teoría de la verdad, tomando por tanto en consideración tanto las fuerzas ilocucionarias y las intenciones comunicativas del hablante como las condiciones semánticas que determinan la relación entre el lenguaje y el mundo. Una teoría así recogería lo mejor de los dos Wittgenstein, y combinaría, en lo posible, las dos líneas de investigación sobre el lenguaje en las que ha estado dividida la filosofía analítica a partir del *Tractatus*. En estas páginas finales quiero desarrollar un poco este motivo, con la esperanza de que ello sirva también para mostrar cómo pueden encajarse entre sí algunas de las piezas que hemos estudiado anteriormente, y cuya conexión recíproca puede haber sido olvidada o puede permanecer aún oscura.

Consideremos una oración cualquiera en cuanto usada en una ocasión determinada por un hablante. ¿Qué es lo que hay que tener en cuenta para poder decir que hemos entendido lo que significa esa oración? Pienso que no se puede pasar por alto ninguno de los siguientes aspectos.

1. *El significado gramatical*.—Es lo que la oración significa con independencia de quién y cuándo la utilice; lo que la oración expresa para cualquier hablante de la lengua en virtud de las reglas sintácticas y semánticas de dicha lengua. Es lo que permite usar la oración en una variedad de ocasiones distintas, y por consiguiente viene dado por la gramática. En muchas ocasiones basta para determinar a qué se refiere el hablante, con lo que la oración cumple su función referencial sin otra ayuda que la de lo que significan sus propias palabras. Así, basta entender lo que significa la

oración «El agua hierve a cien grados centígrados al nivel del mar» para saber a qué se refiere. Se refiere a una cierta sustancia natural, de la que afirma cierta propiedad. Pero no siempre ocurre así; de aquí que otras veces haya que tomar en consideración otro aspecto.

2. *La referencia contextual.*—Tiene lugar cuando es necesario atender al contexto en el que se emplea una oración para determinar a qué se refiere. La función referencial requiere entonces la participación del contexto; a qué se refiera el hablante dependerá, no sólo del significado de la oración, sino también de quién, cuándo, dónde, cómo y de qué esté hablando. La oración «Por allí lo veo venir», se refiere a alguien que habla, y que es quien ve aproximarse algo, y se refiere igualmente a algo que va hacia el hablante y a un lugar por el que va. Pero no basta entender lo que la oración significa por sí sola para saber a qué se refiere cuando es utilizada. Hay que saber además quién habla y de qué habla.

Sobre esta cuestión conviene varias precisiones. Aunque por brevedad decimos, y en lo anterior está dicho, «Esa oración se refiere a...», en rigor es el usuario del lenguaje, el intérprete del signo, quien propiamente se refiere, esto es, quien actualiza la función referencial. Utilizo en general el término «referencia» para esta función, distinguiéndola del referente, esto es, de lo denotado en el cumplimiento o ejercicio de dicha función. Y de todo lo anterior se desprende que no tomo como referente de una oración su valor de verdad, al contrario de Frege, sino aquellas entidades a las que se refieran ciertas partes o expresiones de la oración, como los nombres propios, las descripciones, y tal vez también los predicados (no entraré ahora en más detalles sobre este punto). Dicho de otra forma: que una oración no denota un referente peculiar, distinto de lo que denoten sus partes.

3. *La fuerza ilocutiva implícita.*—Tomemos el ejemplo del párrafo anterior: «Por allí lo veo venir». Por mucho que entendamos el sentido de esta oración, y por bien que sepamos quién habla y quién y por dónde viene, no puede decirse que sepamos exactamente lo que significa en cuanto usada en una ocasión particular si no sabemos que fue pronunciada para avisar al oyente de un peligro. O por tornar a un ejemplo más sencillo y más repetido, aun cuando entienda el sentido de la oración «Súmate a la huelga», y sepa que va dirigida a mí y conozca a qué huelga se refiere, no habré captado del todo su significado, en cuanto usada en una determinada ocasión, si ignoro si se trata de una orden, de un consejo o de una súplica. Naturalmente, este tercer aspecto tan sólo se añade al significado gramatical si la fuerza ilocutiva (o ilocucionaria) está implícita; si está explícita, la oración tendrá forma realizativa y la fuerza vendrá ya dada por el sentido de la oración. Es lo que acontece en: «Te aviso que por allí lo veo venir», y en «Te aconsejo que te sumes a la huelga». Que se trata, respectivamente, de un aviso y de un consejo es algo que ya sabe quien entienda el sentido de estas oraciones.

4. *La implicación contextual.*—Si una oración, en cuanto empleada en cierta ocasión, implica por razón del contexto otra afirmación distinta, no podremos decir que hemos entendido totalmente su significado, en cuanto así empleada, a menos que hayamos captado la afirmación contextualmente implicada. Recordando el fácil ejemplo de Grice, si en una junta calificadora de exámenes uno de sus miembros afirma de un ejercicio escrito: «Tiene buena letra y no comete faltas de ortografía», sin añadir más, no habremos entendido lo que esta oración significa, en cuanto usada en esta ocasión, a menos que captemos la valoración negativa que, implícitamente, pretende expresar.

No cabe incluir aquí, claro está, las demás formas de implicación que tratamos en la sección 7.10. La implicación pragmática general se refería a las convenciones generales que regulan el uso de cualquier lenguaje en la comunicación. Por consiguiente, tal implicación no afecta a lo que signifique una oración particular en una ocasión determinada. Así, por ejemplo, la convención de veracidad, que hemos visto subrayada por Nowell-Smith, por Grice y por Lewis (secciones 7.10 y 8.10), esto es, la implicación de que todo el que profiere una oración declarativa cree, en la mayor parte de los casos, que lo que dice es verdad, debe ser, si la teoría es correcta, una condición necesaria para toda comunicación lingüística en cualquier lenguaje, y por tanto sin consecuencias para determinar el significado particular de una oración usada. En cuanto a las tres formas de implicación semántica que se citó en la sección 7.10, la implicación analítica, la implicación lógica y la presuposición, tampoco pueden añadir nada nuevo a la determinación del significado, ya que, siendo implicaciones que dependen de ciertas palabras o expresiones que aparecen en la oración son, en definitiva, parte de su significado.

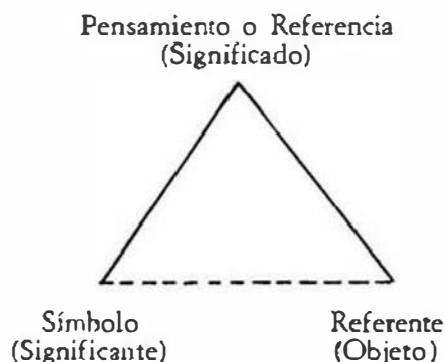
5. Por último, *la fuerza expresiva o connotación subjetiva.*—Este es un aspecto que apenas ha sido mencionado en los capítulos precedentes, y en general ha quedado fuera de las teorías estudiadas. Tiene relación con lo que ya Frege llamaba la «coloración del sentido». El estudio de este aspecto, que es sin duda el más subjetivo e idiosincrásico de cuantos pueden distinguirse en el significado de una oración, tiene gran importancia cuando se trata de investigar las características individuales de un hablante, como acontece en psicología y en psiquiatría. Por esta razón me parece legítimo afirmar que, ante una oración usada en determinado contexto, no sabemos totalmente lo que significa hasta que no sabemos, además de lo que toca a los cuatro aspectos anteriores, lo que esa oración expresa implícitamente acerca de la personalidad del hablante, por ejemplo, angustia, ambición, sentimiento de superioridad, complejo de culpa, etc. Con esto descendemos ya a lo puramente individual y nos alejamos del sistema de la lengua y de toda abstracción. Así se explica que las teorías del significado pasen sobre este aspecto como sobre ascuas. Pero no puede negarse que forma parte también de lo que significa una oración en cuanto empleada en un cierto contexto, aunque, desgraciadamente, ello tenga la siguiente consecuencia

ineludible: que de buena parte de las oraciones que escuchamos, e incluso que pronunciamos, en nuestros intercambios lingüísticos, nunca llegamos a conocer su significado totalmente. Pues no conocemos lo suficiente de nuestro interlocutor, y con frecuencia ni de nosotros mismos, para captar lo que las palabras dan implícitamente a entender acerca del hablante. Al intento de averiguarlo es a lo que algunos llaman «hermenéutica del lenguaje» (a este intento responde el libro de Castilla del Pino, *Introducción a la hermenéutica del lenguaje*, si bien difiero de su planteamiento en muchos puntos). La fuerza expresiva a la que aquí me refiero no debe confundirse con la expresividad paralingüística que una oración tiene en virtud de los gestos, tono de voz y otras características corporales que acompañen a su proferencia.

He venido hablando de una oración en cuanto utilizada en un contexto u ocasión determinados. Pues, en efecto, la unidad de comunicación es la oración usada, y ella es la base empírica sobre la que construir toda teoría del significado. A una oración, en cuanto usada, es a lo que algunos llaman «proferencia», traduciendo así el término inglés *utterance*, que es corriente emplear a estos propósitos. Corresponde a lo que, en la sección 2.2, hablando de los signos en general, llamábamos «signo acontecimiento», que otros llaman «ejemplar» o «caso», en contraposición al signo tipo. No haría falta añadir que una oración puede ser, al menos en su estructura superficial, sumamente simple. Recuérdense los ejemplos de Quine (secc. 8.5), en los que «Conejo» y «Rojo» funcionaban como oraciones. Precisamente por tratarse de oraciones usadas, el significado puede presentar la complejidad que hemos visto. En cada uno de sus aspectos juegan, en diversa medida, tres factores: la gramática, el contexto y la intención comunicativa. Juntos configuran lo que podemos llamar «modelo tridimensional del habla». Dentro de la gramática se puede distinguir entre la gramática particular de cada lengua, la gramática universal (de la que tratamos en el capítulo 4), y la gramática lógica (o sistema formal de las relaciones más generales tanto de carácter sintáctico como semántico y pragmático). En el contexto hay que considerar, no sólo grupos reducidos, como el que daba lugar a la implicación contextual del ejemplo de Grice, sino también toda una clase social, toda una comunidad histórica, e incluso toda una época cultural, según el tipo de implicaciones que se quiera descubrir. El estudio de las vigencias sociales a través del uso del lenguaje, de lo que en él se dice tanto como de lo que no se dice porque se da por sabido, estudio que propugnaba Ortega y que Marías ha realizado a propósito de algunas épocas de la historia de España, constituye, desde nuestro punto de vista, el estudio de un tipo peculiar, e importante, de implicaciones contextuales. En cuanto a la intención comunicativa, es ésta la que explica que el hablante escoja unas u otras palabras para emitir en una ocasión determinada, pero hacer de tal intención la base y el fundamento de una teoría del significado porta en sí el peligro de olvidar que la intención comunicativa está limitada por el contexto en el que se habla y, más aún, por las reglas de la lengua que se emplea. De hecho, lo que pretendemos comunicar es,

salvo casos excepcionales o geniales, lo que nuestra lengua nos permite comunicar. En gran medida, somos hablados por nuestra lengua. De aquí que una teoría como la de Grice, que acentúa más la intención del hablante que las reglas de la lengua, haya de afrontar particulares dificultades a la hora de explicar el significado gramatical, y a la postre dé la impresión de haber colocado la carreta delante de los bueyes (secc. 7.9).

Según esto, ¿cómo puede definirse el significado gramatical? Toda teoría del significado debe asentarse, como ya hemos dicho, sobre la base de la oración en uso, y debe explicitar aquello que hace que una oración pueda usarse en diferentes contextos y ser entendida por el oyente de acuerdo con la intención comunicativa del hablante. Esto es: debe dar cuenta de aquello que regula, y por tanto limita, la intención comunicativa del hablante de acuerdo con el contexto o situación en que se habla. Una definición del significado gramatical debe explicar, por consiguiente, cómo se relacionan las palabras con las diversas situaciones en las que se emplean, y de qué manera remiten a los diferentes aspectos de la realidad y de las cosas. Pienso que esto es lo que claramente ofrece una concepción del significado como función, una concepción como la esbozada por Carnap y que hemos visto desarrollada por Lewis (secc. 8.10). Esta es la que, con algunas variantes, resumiré en lo que sigue. Es cierto que considerar el significado gramatical como función puede parecer excesivamente técnico, pero por lo pronto evita las dificultades de teorías más primitivas. Desde luego está libre de los inconvenientes de una teoría isomorfista, que se ve obligada a postular una relación biunívoca entre los elementos más simples del lenguaje y los elementos más simples del mundo, asumiendo la posibilidad de un análisis que nadie ha sido capaz de llevar a la práctica, y que además tan sólo sería aplicable a una porción del lenguaje, a las oraciones declarativas. Y queda al margen también de los problemas de esas teorías que han dominado durante años, en particular dentro de la lingüística, y que construían el significado como una especie de imagen o entidad mental, que a veces llamaban «concepto», y a través de la cual conectaba el significante con la realidad. Esta concepción procede, en última instancia, de Saussure, que consideraba el significado como concepto, y ha tenido su versión más conocida en el triángulo de Ogden y Richards, popularizado luego por Ullman, que representa la relación de significación así:



La figura quiere dar a entender que el significante no se relaciona directamente con el objeto al que se refiere, sino que tal relación es indirecta, pues tiene lugar a través del pensamiento, que es en lo que consiste el significado. Este se convierte así en una entidad mental y en consecuencia permite todo género de manipulaciones arbitrarias y de afirmaciones incontrastables, como las que antaño permitía el concepto de lenguaje mental, del cual es su directo heredero. (La mayor parte de los libros de semántica lingüística desarrollan, en una u otra forma, este tipo de teoría; puede verse, en particular, Ogden y Richards, *El significado del significado*, capítulo I; Ullmann, *Semántica*, 3.I, y Baldinger, *Teoría semántica*, primera parte. En cambio, Leech, en su *Semántica*, cap. 2, elude estos lugares comunes y ofrece una concepción más compleja y elaborada, aunque en mi opinión excesivamente amplia.)

Consideremos, para empezar, las dos categorías básicas de orden inferior a la oración que suelen tomarse en este tipo de análisis, el nombre común y el nombre propio. Diremos que el significado gramatical de un nombre común es una función que, teniendo como argumento un contexto de uso, tiene como valor un conjunto de objetos, aquellos a los que el nombre se aplica. En principio, podemos considerar que el contexto incluye un mundo posible, si bien con todas las precauciones que se desprenden de las consideraciones que sobre este concepto hicimos en la sección 8.11. El contexto incluirá, además, un cierto hablante, un auditorio (que en el monólogo coincidirá con el propio hablante), un lugar y tiempo de su acto verbal y una situación en la que éste se produzca. Así entendido, un contexto corresponde aproximadamente a lo que llama Lewis «índice», si bien este concepto es más complejo y depende, en parte, de su concepción metafísica de los mundos posibles (*cfr.* «General Semantics», p. 24 s.). Muchos sustantivos son multívocos, como ocurre con el castellano «gato», que con otro propósito examinamos en la sección 4.4. En estos casos, por analogía con lo que estipulamos acerca de los signos en la sección 2.4, tendríamos que decir que la función significativa del nombre adquiere valores múltiples, puesto que el nombre designa varios conjuntos diferentes de objetos, y por tanto que no es función en el sentido lógico del término. Hay, no obstante, otra alternativa, que ahora puede ser preferible. Consiste en asignar a los términos ambiguos varios significados, como de hecho suele hacerse en el discurso ordinario; diremos, entonces, que «gato» tiene, al menos, siete significados, y cada uno de ellos será una función distinta que tomará como valor uno de los siete conjuntos de objetos que esa palabra puede designar. Para esta opción, el significado será una función en sentido lógico. Nótese que esta caracterización del nombre común no prejuzga cómo haya de definirse cada nombre ni si es o no un designador rígido, y es por tanto compatible con la tesis de Kripke sobre los nombres comunes de clases naturales tanto como con la de Mill.

Al introducir el concepto de mundo posible en nuestro análisis hay que tener en cuenta que estamos tratando del significado gramatical en una lengua real, y como vimos en la sección 8.11, hablar de un mundo

posible requiere mantener constante la significación de nuestras palabras. Por ejemplo: cabe imaginar un mundo posible en el que se hablara un castellano tal que la palabra «gato» tuviera un significado diferente de los que tiene actualmente; pero puesto que estamos describiendo ese mundo en el castellano actual, al hablar de gatos en ese mundo tendremos que emplear la palabra «gato» con su significado actual. Naturalmente, no se pretende negar que a lo largo del tiempo se producen cambios semánticos, y que se modifica el significado de las palabras; pero al intentar la caracterización semántica de una lengua, se asume siempre un corte sincrónico para el que la caracterización aspira a ser válida. En resumen: el significado de un nombre común es lo que hace que el nombre se aplique a un cierto conjunto de objetos; es la función que, para un contexto como argumento, adquiere como valor un conjunto de objetos. La referencia del nombre consiste en remitir a tal conjunto de objetos; el referente del nombre es cada uno de los miembros de dicho conjunto. Su sentido es la descripción asociada a ese conjunto, o dicho de otra forma, la definición del objeto que decide su pertenencia al conjunto.

Este conjunto puede serlo de objetos inexistentes. Es lo que acontece con el nombre «sirena» o con el nombre «centauro». ¿Por qué tienen significado estos nombres? Porque remiten a objetos de determinadas características. ¿Pero cómo sabemos que es así si tales objetos no existen? Una respuesta es: porque podrían existir; porque hay mundos posibles en los que existen centauros y sirenas. Tal respuesta me parece insatisfactoria y para ella estaban pensadas las cautelas que insinué antes. ¿Son posibles los centauros y las sirenas? Yo diría que no; pero desde luego lo que no encuentro son razones para decir que sí. Tal vez ni siquiera la pregunta tenga buen sentido. Pero lo que está fuera de duda es que los nombres respectivos sí tienen significado. ¿Por qué? Porque los objetos a los que remiten pueden ser descritos. Hay un procedimiento de construcción para tales objetos, que nos permite incluso representarlos gráficamente, en pinturas o esculturas. La función significativa de un nombre común consiste en remitir a conjuntos de objetos que llevan asociada alguna descripción unívoca, aunque no existan. Un objeto no se define por su existencia, sino por su descripción, por el procedimiento de su construcción. Por supuesto, no todo nombre común que pueda formarse sintácticamente designa un objeto: la expresión «círculo cuadrado» no designa nada y por tanto carece de significado, puesto que no hay descripción posible ni método de construcción para un círculo cuadrado.

Pasemos a los nombres propios. Los así llamados en el lenguaje ordinario agotan su significado en la función puramente referencial. Su significado es, por tanto, aquella función que, para un contexto como argumento, toma como valor el objeto individual al que se refiere el nombre. Un nombre propio, para serlo, no requiere que el objeto lleve asociada ninguna descripción; o dicho de otra manera: un nombre propio no necesita tener sentido. De hecho, sin embargo, algunos nombres propios, como los de personas, suelen tener, en algunas lenguas, algún sentido más o menos

complejo. Entre nosotros, la mayor parte de tales nombres tienen como sentido masculinidad o feminidad, ya que se aplican a varones o a hembras pero no a ambos. Aunque, como es sabido, hay nombres que no cumplen siquiera con esta condición mínima; por ejemplo, «Trinidad». Es patente que los nombres propios son casi todos multívocos, pues se aplican indistintamente a diversos objetos. En congruencia con lo anterior, diremos que un nombre multívoco posee varios significados, tantos como objetos haya que ostenten ese nombre, puesto que en este caso el significado se reduce a la función referencial.

En cuanto a las descripciones definidas, también su significado es una función que toma como valor un objeto individual, pero en virtud de determinadas características del objeto que la descripción recoge. Dicho de otra forma: el significado de la descripción definida depende del significado de las palabras que la componen, y sólo en virtud del significado de éstas cumple aquélla su función referencial. Esto justifica las diversas propuestas que se han hecho para reducir las descripciones definidas a otros tipos de expresión, y que vimos en Russell y Quine (seccs. 6.6 y 8.7). La falta de referente real en los nombres propios y en las descripciones definidas puede resolverse como en los nombres comunes. En cuanto al nombre propio sin referente real, su función significativa tomará como valor el objeto caracterizado por aquellas descripciones que estipulen las características del objeto. Es el caso de los nombres de ficción como «Don Quijote». En el caso de las descripciones definidas, dichas características vienen ya dadas por la propia descripción, y constituyen su sentido, que deriva, en última instancia, del significado de las palabras que la componen. Así en «El Presidente de la República Española en 1980».

Puesto que las demás categorías gramaticales pueden definirse en función de las categorías de nombre propio, de nombre común y de oración, como vimos en Lewis (secc. 8.10), podemos definir el significado de las primeras sobre la base de lo que hemos estipulado para estas últimas (por lo que toca a las oraciones, véase un poco más abajo). Por ejemplo: si un adjetivo es lo que añadido a un nombre da como resultado otro nombre, el significado de un adjetivo será una función que, teniendo como argumento el significado de un nombre, tenga como valor el significado del nombre compuesto resultante. Así, el significado de «pródigo» será la función que, teniendo como argumento el significado de, por ejemplo, «hijo», adquiera como valor el significado de «hijo pródigo»; a su vez, el significado de «hijo pródigo» será aquella función que, para cada contexto como argumento, adquiera como valor el conjunto de los hijos pródigos. Si aplicamos esta propuesta a un nombre propio, tenemos la siguiente diferencia. El significado de la expresión «el pródigo Onassis» será una función que tendrá como valor el individuo Onassis. Pero éste es también el valor que adquiere la función significativa del nombre «Onassis», luego no hay diferencia en la función referencial del nombre, esté o no adjetivado. La diferencia está en que el nombre («Onassis») carece de sentido, y por tanto no lleva asociada con su referente ninguna descripción, mientras que

el nombre adjetivado («el pródigo Onassis») incorpora un sentido, el que viene dado por el adjetivo, y por consiguiente asocia una descripción (no definida, desde luego) con dicho referente.

No entraré aquí en más detalles respecto a otras categorías derivadas. En principio, no hay por qué esperar grandes dificultades en su análisis; en todo caso, mayor complejidad. Recuérdese que, a pesar de diferencias gramaticales superficiales, muchas de estas categorías pueden unificarse. Así, vimos en la sección 4.6 que algunos lingüistas han propuesto la reducción de nombres, adjetivos y verbos a una sola categoría, y hemos visto subrayado por Quine, en la sección 8.7, que adjetivos, verbos y adverbios, junto con los nombres comunes, forman parte por igual de los predicados en un lenguaje lógicamente regimentado. Lo propio puede decirse de las oraciones. En cuanto a los pronombres, funcionan en general como nombres o como descripciones definidas. Carácter distinto tienen las conjunciones, cuyo significado consiste en remitir al propio discurso, conectando entre sí las oraciones para formar oraciones compuestas; aquí pueden ser relevantes las consideraciones de Morris sobre el modo formativo de significar (secc. 7.8).

Vayamos ahora a las oraciones, y de ellas, a las oraciones simples. Hemos encontrado en Davidson y en Lewis la aspiración a dar razón del significado de las oraciones en términos de sus condiciones de verdad (seccs. 8.9 y 8.10). La distinción semántica entre tipos de discurso que propuse en la sección 7.8 contraría esa pretensión, pero elude las dificultades que, para las oraciones imperativas, realizativas y expresivas, tienen que afrontar Davidson y Lewis, y que indiqué en su momento. Por tal razón, si queremos conservar un marco de análisis como el ofrecido por ellos, me parece preferible distinguir tipos de condiciones semánticas, esto es, formas diversas de relacionarse las oraciones con la realidad. Diremos, entonces, que el significado de una oración declarativa es una función que, teniendo como argumento un contexto, tiene como valor la verdad o la falsedad. En el caso de las oraciones eternas, como «El agua hierve a 100 grados centígrados al nivel del mar», su valor veritativo es el mismo para todo contexto en un lapso indefinido de tiempo (al menos mientras se mantengan las leyes actuales de la naturaleza; aquí entraría de nuevo en discusión la idea de un mundo posible en el que esa oración fuera falsa). En esa medida, su verdad depende exclusivamente del sentido de la oración, el cual estipula justamente todas las condiciones necesarias y suficientes para determinar su valor veritativo. Si la oración contiene elementos deícticos, palabras cuya referencia depende del contexto, y es por tanto una oración de las que Quine ha llamado «ocasionales», su valor veritativo variará según el contexto. Así, el significado de la oración «Por allí lo veo venir» es la función que, teniendo como argumento una determinada persona como hablante, un cierto objeto (o persona) como referente de «lo», y un cierto lugar como referente de «por allí», adquiere como valor la verdad o la falsedad, según sea el caso. En las oraciones acerca de entidades de ficción, su significado tendrá como argumento una espe-

cificación del contexto en el que la oración puede adquirir un valor de verdad. Esto permite distinguir la verdad literaria de «Don Quijote arremetió contra los molinos de viento», de la falsedad (relativa a ese mismo contexto) de «Don Quijote huyó ante los molinos de viento». Si la oración es sobre una entidad inexistente y carece de un contexto especial, podemos declararla falsa; tal es el caso de cualquier afirmación que queramos hacer sobre el supuesto Presidente de la República Española en 1980. Es decir: seguimos la solución de Russell, pero sin exigir que dichas oraciones hayan de ser parafraseadas como oraciones cuantificadas; las razones para declararlas falsas son, por nuestra parte, estrictamente semánticas: toda oración cuyo sujeto carezca de referencia es falsa, puesto que una condición necesaria para su verdad es la referencia del sujeto. Esta razón es básicamente la de Russell, pero la cuestión de cómo haya de formularse lógicamente una oración es cuestión diferente, que corresponde a la tarea de traducir el lenguaje ordinario al lenguaje de la lógica.

Nada de lo anterior me parece aplicable, sin retorcer y desfigurar el discurso ordinario, a oraciones no declarativas. Empecemos por las oraciones que en 7.8 hemos llamado «prescriptivas». Diré que el significado de una oración prescriptiva es una función que, teniendo como argumento un contexto determinado, adquiere como valor el de ser cumplida o incumplida, según el curso que siga la realidad. Es el caso, como vimos, de las oraciones imperativas, de las interrogaciones y de las oraciones con términos deónticos, como el verbo «deber». Lo dicho a propósito de las oraciones declarativas vaíe aquí de nuevo: una oración imperativa puede ser cumplida o incumplida en el contexto peculiar de una ficción literaria; y una oración imperativa dirigida a una persona inexistente es siempre incumplida.

Por lo que respecta a las oraciones expresivas, su significado será una función que, para un contexto como argumento, toma el valor de sincera o insincera. «Ojalá se llegue a un desarme mundial» tiene como significado una función que, para cada argumento, adquiere el valor de sincera o insincera, pues lo que expresa es la relación entre un posible estado de cosas y la actitud del hablante.

Finalmente, si la oración es realizativa, su significado será una función cuyo valor sea la validez o la invalidez. «Te aconsejo que no te cases», o «Le nombro Director General del Medio Ambiente», tienen como significado una función cuyo argumento incluirá un cierto hablante y un cierto destinatario, y cuyo valor será la validez o la invalidez, precisamente en virtud de quiénes sean el hablante y el destinatario. Si aconsejo que no se case a alguien que ya está casado, mi consejo es inválido; y si nombro a alguien Director General mi nombramiento es igualmente inválido, pues carezco de facultades jurídicas para realizar tal nombramiento.

La presente propuesta tiene menos elegancia que una teoría general de las condiciones de verdad, ya que obliga a distinguir diferentes clases de oraciones según los valores semánticos que les sean apropiados. No deja de ser paradójico que una teoría unitaria del significado haya de re-

nunciar a unificar los valores semánticos de las oraciones. La unificación puede intentarse, desde luego. Una vía para conseguirla es reformular todas las oraciones en su forma realizativa, como hace Lewis, según vimos. Pero tiene el grave inconveniente, ya señalado, de que convierte a todas las oraciones declarativas sinceramente pronunciadas en verdaderas, ya que quien equivocadamente afirme «Declaro que la Luna no es un satélite de la Tierra» estará haciendo una afirmación verdadera, aunque su contenido sea falso. No pudiendo afrontar esta consecuencia, Lewis prefiere aceptar la asimetría entre oraciones declarativas y oraciones no declarativas, exigiendo la forma realizativa como forma canónica para estas últimas, pero no para las primeras («General Semantics», p. 59). Y semejante asimetría me parece un precio más alto, por lo injustificada y artificiosa, que la pluralidad de valores semánticos. Otra vía para intentar la unificación es reducir los valores semánticos de la siguiente manera. Una oración prescriptiva sería verdadera si cumplida, y falsa en caso de no ser cumplida. Una oración expresiva sería verdadera cuando fuera sincera, y falsa si no lo fuera. Y una oración realizativa sería verdadera en caso de ser válida, y falsa no siéndolo. Tal vez esta vía resulte atractiva para muchos, pero en mi opinión modifica, ampliándolo, el sentido usual de los términos «verdadero» y «falso»; aunque acaso podría insinuarse que la modificación no es mayor que la que tiene lugar cuando llamamos «verdadera» a una tautología. En todo caso presentaría una dificultad, ya señalada anteriormente, al intentar aplicar la convención (T) de Tarski a algunas de estas oraciones, pues obligaría a fórmulas tales como: «'Escúchame' es verdadera si y sólo si escúchame», «'¡Viva el Betis!' es verdadera si y sólo si viva el Betis», etc. Podría oponerse, no obstante, que la cuestión se reduce a encontrar una traducción adecuada de las oraciones imperativas y expresivas al metalenguaje en el que formulamos la condición (T), y podría recomendarse una traducción tal como: «'Escúchame' es verdadera si y sólo si me escuchas», o «'Viva el Betis!' es verdadera si y sólo si deseo que el Betis progrese». Sin embargo, esta contrapropuesta tiene el inconveniente de que asigna las mismas condiciones de verdad a una oración imperativa como «¡Escúchame!» y a una oración declarativa como «Me escuchas», con lo que borra toda distinción semántica entre ellas.

Si aceptamos alguna de estas vías, o de otras semejantes, tendremos la satisfacción de poder resumir el concepto de significado de una oración en una breve fórmula como ésta: el significado gramatical de una oración es aquella función que, teniendo como argumento un contexto determinado, tiene como valor la verdad o la falsedad. Si, por el contrario, mantenemos la distinción semántica entre cuatro tipos de discurso, nos veremos obligados a definir separadamente el significado para los cuatro tipos correspondientes de oración. La única definición general que podremos enunciar entonces, será de este corte: el significado gramatical de una oración es aquella función que, teniendo como argumento un contexto determinado, adquiere como valor alguno de los valores semánticos pertenecientes a los siguientes

pares: verdad-falsedad, cumplimiento-incumplimiento, sinceridad-insinceridad y validez-invalidéz.

Pero esta cuestión, como otras mencionadas antes, y como tantas más que ni siquiera ha habido lugar para citar, está todavía abierta.

APENDICE

¡Todo lo importante del mundo se resume en palabras, abren o cierran, atan o libran! (TORRENTE BALLESTER, *La isla de los jacintos cortados*.)

Aunque dentro de la tradición filosófica marxista los problemas del lenguaje no han recibido nunca la atención que, como hemos visto, les dedicó la filosofía analítica, hay un tema central en aquella tradición que por su interés general debe tenerse en cuenta en una teoría del significado: es el tema de las relaciones entre la ideología y el lenguaje. En lo que sigue esbozaré este tema, intentando mostrar de qué manera enlaza con la teoría del significado que se ha apuntado en las páginas precedentes.

Si bien en Marx y Engels no hay ciertamente una teoría del lenguaje, sí puede encontrarse consideraciones de interés sobre las relaciones entre el pensamiento, el lenguaje y la vida real (por ejemplo, en *La ideología alemana*; para un examen detallado de estas aportaciones puede verse *Language et marxisme*, de Houdebine, caps. II y III). Desgraciadamente estas sugerencias no fueron recogidas en la tradición posterior y quedaron sin fructificar en una teoría más amplia que pudiera tener interés (sobre las razones de ello, Houdebine, *op. cit.*, IV.I). Con el nacimiento de la lingüística en la obra de Saussure, y a partir de su divulgación en la Unión Soviética, comienzan a configurarse las condiciones para que se despierte el interés por un planteamiento de los problemas lingüísticos en la perspectiva del pensamiento dialéctico. La primera forma explícita de tal interés será, por desgracia, la teoría del lingüista ruso Nicolás Marr acerca de las relaciones entre el lenguaje y la superestructura ideológica.

Según Marr, el lenguaje tendría carácter clasista y pertenecería a la superestructura ideológica de la sociedad, siendo instrumento de dominación al servicio de la clase dominante. Esta posición aparecía, por cierto, en el contexto de una curiosa teoría acerca del origen del lenguaje («Ueber die Entstehung der Sprache», 1925). Según esta teoría, todas las lenguas

tendrían un origen común en el lenguaje de gestos manuales propio de las comunidades primitivas. De aquí se habría pasado a un lenguaje sonoro partiendo de los sonidos que los brujos y hechiceros utilizaban para la celebración de sus ritos. De esta manera, el lenguaje vocal aparecía originariamente vinculado a la clase dominante, la de los hechiceros, únicos que conocían el significado de sus voces, y podía así ser empleado como un instrumento para ejercer y asegurar su dominio sobre la tribu. Todas las lenguas, ligadas en su evolución al desarrollo histórico de los modos de producción, habrían retenido este carácter de dominación clasista.

No hace falta decir que la doctrina de Marr no recibió nunca el menor asomo de comprobación empírica, pese a lo cual gozó de amplia difusión e incluso hizo escuela en la Unión Soviética de los años veinte. La polémica fue notable, y condujo al final a afirmaciones desmesuradas como la de que cada clase social tiene su propio lenguaje. Aunque, según se ha escrito, la teoría de Marr convenía por razones políticas a la Unión Soviética (Marcellesi y Gardin, *Introducción a la sociolingüística*, II.1), su influencia fue bruscamente cortada por una famosa intervención explícita de Stalin (sobre la doctrina de Marr y su contexto sociopolítico puede verse, además de la obra citada de Marcellesi y Gardin, cap. II, el número 46 de la revista *Langages*, 1977, enteramente dedicada al tema).

Fue bastante tiempo después, en 1950, cuando apareció el escrito de Stalin que desautorizaba la doctrina de Marr, llegando a negar su carácter marxista («Sobre el marxismo en la lingüística», original publicado en *Pravda*). La posición de Stalin comienza por distinguir nítidamente entre lenguaje y superestructura, señalando que, mientras que en Rusia, tras la revolución, la infraestructura capitalista ha sido sustituida por una infraestructura socialista, cambiando de modo acorde la superestructura ideológica y jurídico-política, la lengua rusa, en cambio, continúa siendo sustancialmente la misma. ¿Qué ha cambiado en ella? Según Stalin, sólo ciertos aspectos semánticos y léxicos. Han aparecido nuevas palabras y expresiones vinculadas con la nueva forma socialista de producción y con las nuevas instituciones, al tiempo que otras, conectadas al modelo social anterior, han desaparecido o han cambiado de sentido. Pero ni el fondo léxico esencial ni la estructura gramatical han sufrido modificación seria. La lengua no es producto de una infraestructura determinada ni de una clase particular, sino que es medio de comunicación entre todas las clases sociales, a las que sirve indiferentemente; la lengua rusa sirve, por ello, a la cultura socialista tanto como sirvió antes a la cultura burguesa.

La argumentación de Stalin es hasta aquí de un sentido común que casi resulta burgués (lo que da lugar a una crítica sarcástica en Houdebine, *Langage et marxisme*, IV.III). A ello hay que añadir aún otro argumento, de mayor alcance teórico, pero muy impreciso. Según Stalin, mientras que la superestructura se halla ligada a la producción tan sólo de manera indirecta, a través de la infraestructura económica que constituye la base, el lenguaje, en cambio, está directamente conectado con la actividad productiva, pues lo está, en última instancia, con todas las actividades del hombre.

De aquí que el lenguaje refleje de modo inmediato las transformaciones que tienen lugar en la producción, enriqueciendo su fondo léxico y perfeccionando su estructura gramatical. El lenguaje, por consiguiente, no coincide con la superestructura, sino que desborda sus límites, y no tiene tampoco carácter de clase. Esto no significa que las clases no dejen su impronta en el lenguaje. Hay lo que llama Stalin «dialectos» o «jergas» de clase, caracterizados por el empleo de expresiones y términos propios de una clase social, y que sirven a la pretensión de utilizar el lenguaje para la defensa de los intereses de clase; pero no constituyen más que una pequeña porción de una lengua, y naturalmente ninguno de tales dialectos constituye un lenguaje.

Frente a las exageraciones de Marr, Stalin se ha movido al otro extremo. No parece reconocer relaciones importantes entre la ideología y el lenguaje ni articulación alguna entre el uso del lenguaje y la división en clases. La posibilidad de una teoría sobre estos aspectos del lenguaje queda prácticamente negada. Se entiende que un crítico haya llegado a decir que «nunca se ha estado tan lejos de Marx como en este texto de Stalin» (Houdebine, *op. cit.*, p. 159).

Y, sin embargo, muchos años antes, en los propios años de florecimiento de la doctrina de Marr, existía un círculo de estudiosos rusos que habían alcanzado claridad mucho mayor sobre estos problemas, dejando incluso publicaciones más agudas que otras muy posteriores; me refiero en particular a la obra *Marxismo y filosofía del lenguaje*, publicada originalmente en Leningrado, en 1929 (el título de la traducción castellana es *El signo ideológico y la filosofía del lenguaje*; teniendo en cuenta que la traducción apareció en Buenos Aires en 1976, el lector puede figurarse las razones del cambio de título; yo lo citaré por el título literal en castellano, aunque sigo la traducción inglesa, realizada directamente sobre el ruso). El libro apareció bajo el nombre de Valentin Voloshinov, y así han aparecido varias de sus recientes traducciones (incluida la castellana, aunque no la francesa), pero, al menos en su mayor parte, era la obra del maestro y cabeza del grupo, Michail Bachtin, inspirador de esta escuela y conocido por sus análisis estilísticos de diversos escritores clásicos. Por diferentes motivos que van desde lo político a lo temperamental, Bachtin dejó que el libro saliera bajo el nombre de su discípulo y con algunas adiciones y modificaciones de la pluma de éste. Por cierto que Voloshinov desapareció en los años treinta, aunque Bachtin sobrevivió al estalinismo y falleció no hace mucho, en 1975. En lo sucesivo nos referiremos a ambos como autores del libro.

Lo primero que se encuentra subrayado en esta obra es la íntima relación existente entre la teoría de la ideología y la semiótica. Todo lo ideológico posee significado, en cuanto que representa o remite a otra cosa distinta de sí mismo; esto es: todo lo ideológico es signo; «sin signos, no hay ideología» (parte primera, cap. 1). Para Bachtin y Voloshinov, el ámbito de la ideología y el ámbito de los signos son equivalentes: todo signo es susceptible de una valoración ideológica, y todo lo ideológico po-

see valor semiótico. La ideología no se encuentra en la conciencia como ámbito independiente: pretender lo contrario es psicologismo, y en suma idealismo. La ideología está en los signos, por medio de los cuales se desarrolla la comunicación humana y en los cuales se constituye la conciencia: «la conciencia únicamente llega a ser conciencia una vez que está llena de contenido ideológico (semiótico), y por consiguiente sólo en el proceso de interacción social» (*ibidem*).

Pero hay muchas clases de signos: obras de arte, símbolos religiosos y políticos, etc. De todos ellos, el signo por excelencia es la palabra, y esto hace de ella también el fenómeno ideológico por excelencia. Bachtin y Voloshinov indican varias características de la palabra, que justifican esa condición. En primer lugar, la palabra es el signo más puro, pues no es más que signo, toda su realidad se reduce a signo. En segundo lugar, es neutral ante cualquier campo ideológico, ya que sirve para cualquier función ideológica: científica, estética, moral... Tercero, la palabra está ligada a un campo ideológico genérico, pero muy importante, que es la conversación, la comunicación cotidiana entre los hablantes. Cuarto, la palabra es el medio en el que se da la conciencia, es el material semiótico de la vida interior. Y en quinto y último lugar, y a consecuencia de lo anterior, la palabra acompaña a todo acto de creación ideológica y a todo acto de comprensión o interpretación, aunque sólo sea como lenguaje interior. En estas últimas características, que hay que subrayar especialmente, late acaso la influencia de Vygotsky, que había destacado la importancia de la función egocéntrica en el uso y desarrollo del lenguaje (recuérdese la sección 5.4).

Destacada, de esta forma, la significación ideológica del lenguaje, cabe preguntar: ¿de qué manera conectan la ideología y el lenguaje? La idea de Bachtin y Voloshinov (primera parte, cap. 2), es que el contenido de todo signo lingüístico tiene una determinada acentuación o carga valorativa socialmente adquirida, en virtud de la cual el signo funciona ideológicamente. Esta acentuación es múltiple, de acuerdo con las diferentes clases sociales; cada clase social utiliza el lenguaje con una peculiar carga valorativa. Puesto que la clase social y la comunidad semiótica no coinciden, ya que esta última incluye varias clases, hay que distinguir el lenguaje común a éstas de aquello que, en el uso de este lenguaje, las distingue, a saber: la carga o acento valorativo que caracteriza a cada una de ellas. Estas diferentes cargas valorativas se entrecruzan en el lenguaje, y manifiestan así la lucha de clases, a través de la cual la existencia social queda refractada en el lenguaje. ¿Por qué refractada y no meramente reflejada? Porque el lenguaje no es un medio neutral en cuyo uso cada clase social tenga la misma capacidad y autonomía, sino que «la clase dominante se esfuerza por impartir un carácter supraclasista y eterno al signo ideológico, haciéndolo uniacentado, y extinguiendo o reprimiendo la lucha entre distintos juicios de valor sociales que tiene lugar en él» (p. 23 de la trad. inglesa).

De este modo quedan trazadas las líneas para un tratamiento de las relaciones entre la base y la superestructura a través del medio lingüístico. Para entender adecuadamente lo anterior debe recordarse el sentido en el

que se emplean determinados términos técnicos. En el presente contexto, una sociedad se caracteriza por un modo de producción típico, al que pueden acompañar restos de modos de producción residuales pertenecientes a formaciones sociales históricamente anteriores. Así, las sociedades actuales en Occidente se caracterizan por el modo de producción capitalista, al que acompañan de hecho restos del modo de producción agrícola, artesanal, etcétera. Un modo de producción consta de una infraestructura económica o base y de una superestructura institucional e ideológica. La base incluye unas fuerzas productivas (fuerza de trabajo, técnicas y medios de producción, etc.) y unas relaciones de producción, en las que se articulan aquellas con el resto de la sociedad. La superestructura incluye aquellas instituciones políticas y jurídicas que tienen a su cargo mantener la estabilidad de las relaciones sociales junto con la ideología, esto es, el conjunto de las actitudes, ideas y doctrinas que sirven al mantenimiento de las actuales relaciones sociales y de producción. En este modelo de análisis sociológico, las clases sociales se caracterizan por su función en las relaciones de dominación que se dan en toda sociedad. De aquí que cada modo de producción se caracterice por dos clases antagónicas entre sí, la clase dominante y la clase dominada. La primera tiende a mantener estables las relaciones sociales impidiendo los cambios estructurales, para lo que recurre a todos los factores que integran la superestructura, desde el derecho a la religión. La clase dominada tendrá como interés, acaso no consciente pero sí real, cobrar conciencia de su situación para ponerle fin, y en esto consiste la lucha de clases. Puesto que en toda sociedad coexisten junto al modo de producción característico otros procedentes de épocas anteriores, se explica que además de las dos clases propias de aquél puedan detectarse otras clases sociales residuales pertenecientes a esos modos de producción igualmente residuales. Este es el análisis más clásico, y no hace falta añadir que, en su aplicación a ejemplos concretos, requiere multitud de precisiones, y que, desde el punto de vista teórico, cuenta también con muy diversas reformulaciones. El análisis de las clases sociales, por ejemplo, se ha desarrollado extraordinariamente en los últimos decenios. (No es necesario subrayar que el concepto de clase social tiene un sentido muy distinto en la sociología ajena a la tradición marxiana: para todo lo anterior puede verse Fioravanti, *El concepto de modo de producción*; Dahrendorf, *Las clases sociales y su conflicto en la sociedad industrial*; Parkin, *Orden político y desigualdades de clase*; y Poulantzas, *Las clases sociales en el capitalismo actual*.)

Al incluir la ideología dentro de la superestructura, en el breve recordatorio anterior, hemos tomado la ideología como aquellas doctrinas y formas culturales que sirven a la clase dominante para asegurar su dominio, y en cuanto así sirven. Es lo que otras veces se ha considerado como pensamiento deformado por los intereses de clase dominante. Si bien este sentido es el más típico en el concepto marxista de ideología, se encuentra también, incluso en Marx y Engels, un sentido más neutral del término «ideología», cuando éste designa cualquier doctrina, idea u opinión meramente en virtud de su vinculación a los intereses de una clase social, sea cual fuere.

Pues bien, como se habrá observado, en este último sentido neutral toman Bachtin y Voloshinov la ideología cuando estudian su relación con el lenguaje. Si por el contrario tomamos la ideología en el sentido valorativo, como visión deformada de las cosas, podremos afirmar que el uso ideológico del lenguaje consiste justamente en dificultar que la clase dominada pueda expresar por medio de él adecuadamente sus propios intereses. De esta manera se refractan, se deforman, las verdaderas condiciones de la comunicación social, pues se impide que el lenguaje exprese libremente y por igual el conflicto de clases. Si el análisis es correcto, eso es lo que intenta la clase dominante en su manipulación del lenguaje. En esta medida, la clase dominada se encontrará en una situación lingüísticamente alienada.

Que el planteamiento de Bachtin y Voloshinov es incomparablemente más lúcido que las afirmaciones de Marr y de Stalin, es patente. Que debidamente completado y prolongado podría haber sido muy fértil salta a la vista, y sólo hace más lamentable que su libro no haya circulado nunca, y que no haya sido conocido en Occidente hasta los años setenta. Nada de cuanto conozco, escrito después de ellos, es más preciso ni está mejor fundado: ni las vagas consideraciones de Rossi-Landi sobre la alienación lingüística (*Linguistics and Economics*, 7.4), ni su identificación dialéctica entre ideología y lenguaje, que me parece arbitraria y vacía (*Ideología*, 2.8), ni su forzado intento de explicar la comunicación lingüística como una forma de la producción e intercambio de bienes, aplicándole categorías económicas (*Linguistics and Economics*, passim, y «El lenguaje como trabajo y como mercado»).

El problema fundamental, en mi opinión, era encontrar un lugar teórico, un plano adecuado, en el que situar este orden de problemas. Que ese lugar no es la gramática, no es el sistema de la lengua, es obvio, y esto es lo que el escrito de Stalin se limitaba a señalar. Que ese lugar tampoco está en el habla individual es igualmente claro, pues no es el individuo el que crea la ideología, ni ésta es un fenómeno individual, sino al contrario, es justamente una consecuencia de la ubicación del individuo en una clase social. Todo dualismo del tipo de la distinción entre lengua y habla tenía que hacer imposible el planteamiento de estos problemas (y de aquí la crítica de Ponzio a Saussure en *Producción lingüística e ideología social*, III.1). Sin llegar a encontrar las categorías adecuadas, Bachtin y Voloshinov planteaban el problema con perspicacia. Para acabar de formularlo correctamente basta, en mi opinión, recurrir a aquellos aspectos del uso del lenguaje que, siendo compatibles con el sistema de la lengua, pero sin estar exigidos por ella, no derivan tampoco del individuo como tal ni son particulares. Tales aspectos son los que entran bajo el concepto de norma de Coseriu, que hemos estudiado en otro lugar (secc. 3.2). Como allí señalé, las peculiaridades que posea el uso del lenguaje en cuanto sirve a la defensa de los intereses de clase puede categorizarse como una norma clasista en el uso del lenguaje. ¿Qué peculiaridades son ésas? Peculiaridades pragmáticas, puesto que pertenecen al uso del lenguaje, y por consiguiente peculiaridades que no se mostrarían apenas en el sistema gramatical. Entre estas

peculiaridades estarán el empleo más frecuente de ciertos términos, por ejemplo, los de carácter espiritualista, los que se refieren a la autonomía del individuo, los de carácter estético, etc. Pero lo notable es que el funcionamiento ideológico del lenguaje no se advierte tanto en la literalidad del propio discurso, en las oraciones por sí mismas, y por tanto en las puras palabras, cuanto en lo que éstas dan por sabido, y en lo que pretenden contextualmente dar a entender, en especial, aquellos juicios de valor contextualmente implicados por lo que se dice, así como aquellos otros que implícitamente se pretende evocar en el oyente. Ya hemos visto mencionados los juicios de valor en la doctrina de Bachtin y Voloshinov sobre la acentuación social del signo lingüístico; también constituyen la categoría prominente en los análisis empíricos más recientes, como los de Mattelart (por ejemplo, en «El marco del análisis ideológico», 1970), y los de Verón (así, en «Ideología y comunicación de masas: la semantización de la violencia política», 1971) y los de Dorfman (en *Para leer a Pato Donald*, 1972, junto con Mattelart).

Eliseo Verón, por ejemplo, al explicitar las coordenadas teóricas de su trabajo, se ha ocupado de señalar que la ideología «no es un tipo particular de mensajes, o una clase de discursos sociales, sino uno de los muchos niveles de organización de los mensajes, desde el punto de vista de sus propiedades semánticas» (*op. cit.*, p. 141); la ideología aparece así como «un nivel de significación que puede estar presente en cualquier tipo de mensajes, aun en el discurso científico» (*ibidem*). Este nivel de significación sólo se descubre al descomponer los mensajes, ya que la ideología propiamente «no se comunica, sino que se metacomunica», o si se prefiere, que «opera por connotación y no por denotación»; puede afirmarse, por ello, que «la lectura ideológica de la comunicación social consiste en descubrir la organización implícita o no manifiesta de los mensajes» (*ibidem*). La función conativa o prescriptiva de toda comunicación lingüística (que consideramos en la sección 7.7) se realiza por medio de la comunicación ideológica en cuanto ésta tiende a reforzar determinadas formas de comportamiento social, y esta función prescriptiva se cumple además de cualquier otra función, por ejemplo, informativa, que el mensaje puede tener por razón de sus características semánticas (*op. cit.*, p. 142).

La interpretación ideológica del uso del lenguaje constituye, pues, un momento fundamental en su análisis pragmático, y requiere atender primordialmente al contexto en el que se produce ese uso. Más bien por lo que en un contexto determinado se calla que por lo que se dice, y más bien por lo que se asume que por lo que se pone en cuestión, el discurso cumple una función ideológica, que no puede, por ello, calibrarse adecuadamente más que teniendo en cuenta el contexto extralingüístico. Considerando que, según vimos en el epílogo, no hay límite a la amplitud del contexto, éste puede ser en ocasiones únicamente un grupo reducido, como el de las amas de casa, los obreros agrícolas o los profesores de universidad; pero puede ser, en otras, toda una sociedad o todo un grupo de sociedades. De aquí que el tipo de relación entre oraciones donde debe buscarse el condici-

miento ideológico sea, en mi opinión, la relación de implicación contextual. El significado ideológico de una oración proviene, la mayor parte de las veces, del contexto en el que se pronuncia. La implicación contextual, como vimos en la sección 7.10, no es más que un tipo de implicación pragmática. Y como allí estudiamos, hay una implicación pragmática de carácter general que incluye aquellos supuestos más generales sobre los que se apoya el uso del lenguaje en la comunicación. Según esto, el uso del lenguaje implica, entre otras cosas, que el hablante intente ser veraz y decir lo que sea relevante para la comunicación en la que participa. Pues bien, si las relaciones de dominación son una condición que aparece en toda comunidad humana, a esos supuestos pragmáticos que subyacen al uso del lenguaje y que han investigado algunos filósofos analíticos, cabe añadir otro más: el uso del lenguaje implica pragmáticamente que lo que se dice tiene, en general, un significado ideológico, esto es, que en alguna medida está al servicio de intereses de clase. Pienso que ésta fue la intuición profunda de Bachtin y Voloshinov.

BIBLIOGRAFIA

- Acero, J.; Bustos, E., y Quesada, D.: *Introducción a la filosofía del lenguaje*, Cátedra, Madrid, 1982.
- Alarcos Llorach, E.: *Fonología española*, Gredos, Madrid, 1950 (4.ª ed. aumentada, 1965).
- Alcina, J., y Blecua, J. M.: *Gramática española*, Ariel, Barcelona, 1975.
- Alston, W. P.: *Philosophy of Language*, Prentice-Hall, Englewood-Cliffs, 1964 (*Filosofía del lenguaje*, Alianza, Madrid, 1974).
- Alston, W., y otros: *Los orígenes de la filosofía analítica: Moore, Russell, Wittgenstein*, Tecnos, Madrid, 1976.
- Andrés, T. de: *El nominalismo de Guillermo de Ockham como filosofía del lenguaje*, Gredos, Madrid, 1969.
- Anscombe, E.: *An Introduction to Wittgenstein's Tractatus*, Hutchinson, Londres, 1959 (*Introducción al Tractatus de Wittgenstein*, El Ateneo, Buenos Aires, 1977).
- Aristóteles: *Metafísica*, edición trilingüe de V. García Yebra, Gredos, Madrid, 1970.
- : *Sobre la interpretación*, Cuadernos Teorema, Valencia, 1977.
- Austin, J. L.: «A Plea for Excuses», 1956, incluido en *Philosophical Papers*.
- : «Performative Utterances», 1956, incluido en *Philosophical Papers*.
- : *Philosophical Papers*, Oxford University Press, 1961, por donde cito (hay 2.ª edición ampliada de 1970, traducida como *Ensayos filosóficos*, Revista de Occidente, Madrid, 1975).
- : «Performatif-Constatif», en *La Philosophie Analytique*, Cahiers de Royaumont, 4, Les Editions, de Minuit, París, 1962 (trad. inglesa: «Performative-Constative», en la recopilación de la Ch. Caton, *Philosophy and Ordinary Language*, University of Illinois Press, 1963).
- : *How to Do Things with Words*, Oxford University Press, 1962 (*Palabras y acciones*, Paidós, Buenos Aires, 1971).
- Ayer, A. J.: *Language, Truth and Logic*, Gollancz, Londres, 1936; 2.ª ed. revisada 1946 (*Lenguaje, verdad y lógica*, Editorial Universitaria de Buenos Aires, 1965; una edición posterior en Martínez Roca, Barcelona).
- (recop.): *Logical Positivism*, The Free Press, Nueva York, 1959 (*El positivismo lógico*, Fondo de Cultura Económica, México, 1965).
- Ayer, A. J.; Gellner, E., y Kuznetsov, I. V.: *Filosofía y Ciencia*, Cuadernos Teorema, Valencia, 1975.
- Bach, E.: *Syntactic Theory*, Holt, Rinehart and Winston, Nueva York, 1974 (*Teoría sintáctica*, Anagrama, Barcelona, 1976).

- Bachtin, M., y Voloshinov, V.: *El signo ideológico y la filosofía del lenguaje*, Nueva Visión, Buenos Aires, 1976 (traducción del título original: *Marxismo y filosofía del lenguaje*, publicado en ruso en Leningrado en 1929; trad. inglesa: *Marxism and the Philosophy of Language*, Seminar Press, Nueva York, 1973).
- Baker, G. P., y Hacker, P. M.: *Wittgenstein. Understanding and Meaning*, Blackwell, Oxford, 1980.
- Baldinger, K.: *Teoría semántica*, Ediciones Alcalá, Madrid, 1970.
- Bayés, R. (recop.): *¿Chomsky o Skinner? La génesis del lenguaje*, Fontanella, Barcelona, 1977.
- Berlin, B., y Kay, P.: *Basic Color Terms*, University of California Press, Los Angeles, 1969.
- Bever, T. G.: «The Cognitive Basis for Linguistic Structures», en J. R. Hayes (recop.), *Cognition and the Development of Language*, Wiley, Nueva York, 1970.
- Black, M.: «The Semantic Definition of Truth», *Analysis*, 1948; reimpresso en su obra *Language and Philosophy*, Cornell University Press, Ithaca, 1949.
- : «La relatividad lingüística: las opiniones de B. L. Whorf», en *Modelos y metáforas*, Tecnos, Madrid, 1966 (original: *Models and Metaphors*, Cornell University Press, Ithaca, 1962).
- : *The Labyrinth of Language*, Praeger, 1968 (*El laberinto del lenguaje*, Monte Avila, Caracas, 1969).
- : «Comment», en Borger y Cioffi (recops.), *Explanation in the Behavioural Sciences*, Cambridge University Press, 1970.
- Blasco, J. L.: *Lenguaje, filosofía y conocimiento*, Ariel, Barcelona, 1973.
- : «Compromiso óntico y relatividad ontológica», en el volumen colectivo *Aspectos de la filosofía de Quine*, Teorema, Valencia, 1976.
- Block, L. (recop.): *Perspectives on the Philosophy of Wittgenstein*, Blackwell, Oxford, 1981.
- Brown, R.: «Development of the First Language in the Human Species», *American Psychologist*, 1973.
- : *A First Language. The Early Stages*, Allen & Unwin, Londres, 1973.
- Bruner, J. S.; Olver, R., y Greenfield, S.: *Studies in Cognitive Growth*, Wiley, Nueva York, 1966.
- Bueno, G.: *El papel de la filosofía en el conjunto del saber*, Ciencia Nueva, Madrid, 1970.
- Bühler, K.: *Sprachtheorie*, Fischer, Jena, 1934 (*Teoría del lenguaje*, Revista de Occidente, Madrid, 1950).
- Bunge, M.: «Análisis de la analiticidad», en su recopilación citada a continuación.
- (recop.): *Antología semántica*, Nueva Visión, Buenos Aires, 1960.
- Busnel, R. G., y Classe, A.: *Whistled Languages*, Springer, Berlín, 1976.
- Campbell, R., y Wales, R.: «The Study of Language Acquisition», en J. Lyons (recop.), *New Horizons of Linguistics*, Penguin, Harmondsworth, 1970 (*Nuevos horizontes de la lingüística*, Alianza, Madrid).
- Carnap, R.: *Der Logische Aufbau der Welt*, Weltkreis Verlag, Berlín, 1928.
- : «Überwindung der Metaphysik durch Logische Analyse der Sprache», *Erkenntnis*, 1932; trad. inglesa en la recopilación de Ayer, *Logical Positivism* («Superación de la metafísica por medio del análisis lógico del lenguaje», en la traducción de dicha recopilación).
- : *Logische Syntax der Sprache*, Springer, Viena, 1934 (trad. inglesa: *The Logical Syntax of Language*, Kegan Paul, Londres, 1937).
- : *Philosophy and Logical Syntax*, Kegan Paul, Londres, 1935 («Filosofía y sintaxis lógica», en la recopilación de Muguerza, *La concepción analítica de la filosofía*, vol. 1).
- : «Testability and Meaning», *Philosophy of Science*, 1936-37; abreviado, en la recopilación de M. Feigl y M. Brodbeck, *Readings in the Philosophy of Science*, Appleton-Century-Crofts, Nueva York, 1953).
- : *Introduction to Semantics*, Harvard University Press, 1942.

- : *Formalization of Logic*, Harvard University Press, 1943.
- : *Meaning and Necessity*, The University of Chicago Press, 1947.
- : «Empiricism, Semantics and Ontology», *Revue Internationale de Philosophie*, 1950; reimpreso en la 2.ª edición de *Meaning and Necessity* («Empirismo, semántica y ontología», en la recopilación de Muguerza, *La concepción analítica de la filosofía*, vol. 2).
- : «Meaning Postulates», *Philosophical Studies*, 1952; reimpreso en la 2.ª edición de *Meaning and Necessity*.
- : «Meaning and Synonymy in Natural Languages», *Philosophical Studies*, 1955; reimpreso en la 2.ª edición de *Meaning and Necessity* («Significado y sinonimia en los lenguajes naturales», en la recopilación de Bunge, *Antología semántica*).
- : «Intellectual Autobiography», en la recopilación de P. A. Schilpp, *The Philosophy of Rudolf Carnap*, Open Court, La Salle, 1963.
- Casares, J.: *Diccionario ideológico de la lengua española*, Gustavo Gili, Barcelona, 2.ª edición puesta al día, 1971.
- Cassirer, E.: *An essay on Man*, New Haven, 1944 (*Antropología filosófica*, Fondo de Cultura Económica, México, 1945).
- Castilla del Pino, C.: *Introducción a la hermenéutica del lenguaje*, Península, Barcelona, 1972.
- Cerezo, P.: *Metafilosofía*, Labor, Barcelona, en preparación.
- Clack, R. J.: *Bertrand Russell's Philosophy of Language*, Martinus Nijhoff, La Haya, 1969 (*La filosofía del lenguaje de Bertrand Russell*, Fernando Torres, Valencia, 1976).
- Contreras, H. (recop.): *Los fundamentos de la gramática transformacional*, Siglo XXI, México, 1971.
- Cooper, D. E.: *Knowledge of Language*, Prism Press, Londres, 1975.
- Copi, I. M.: «Objects, Properties and Relations in the *Tractatus*», *Mind*, 1958; incluido en la recopilación de Copi y Beard que se cita a continuación.
- Copi, I. M., y Beard, R. W. (recops.): *Essays on Wittgenstein's Tractatus*, Routledge and Kegan, Paul, Londres, 1966.
- Corominas, J.: *Diccionario crítico etimológico de la lengua castellana*, Gredos, Madrid, 4 vols., 1954-57.
- Coseriu, E.: *Sistema, norma y habla*, Montevideo, 1952; y en *Teoría del lenguaje y lingüística general*, Gredos, Madrid, 1967.
- Chappell, V. C. (recop.): *Ordinary Language*, Prentice-Hall, Englewood-Cliffs, 1963 (*El lenguaje común*, Tecnos, Madrid, 1971).
- Chomsky, N.: *Syntactic Structures*, Mouton, La Haya, 1957. (*Estructuras sintácticas*, Siglo XXI, México, 1975).
- : «A Review of B. F. Skinner's *Verbal Behavior*», *Language*, 1959 (traducción castellana en la recopilación de R. Bayés).
- : *Aspects of the Theory of Syntax*, M. I. T. Press, Cambridge, Mass., 1965 (*Aspectos de la teoría de la sintaxis*, Aguilar, Madrid, 1970).
- : «The Formal Nature of Language», en E. Lenneberg, *Biological Foundations of Language*, Wiley, Nueva York, 1967. («La naturaleza formal del lenguaje», en la traducción del libro de Lenneberg, y también en Gracia (recop.), *Presentación del lenguaje*.)
- : «Recent Contributions to the Theory of Innate Ideas», *Synthese*, 1967 (traducción castellana en *Teorema*, 1973).
- : «Language and the Mind», *Psychology Today*, 1968 («El lenguaje y la mente», en Contreras (recop.), *Los fundamentos de la gramática transformacional*).
- : *Language and Mind*, Harcourt Brace, Nueva York, 1968 (*El lenguaje y el entendimiento*, Seix Barral, Barcelona, 1971).
- : «Linguistics and Philosophy», en S. Hook (recop.), *Language and Philosophy*, Nueva York University Press, 1969.
- : «Deep Structure, Surface Structure and Semantic Interpretation», en *Studies in General and Oriental Linguistics*, recop. por R. Jakobson y S. Kawamoto, TEC Co., Tokio, 1970 (incluido en la obra de Chomsky, *Studies on Semantics in Generative*

- Grammar, Mouton, La Haya, 1972; trad. cast. en Sánchez de Zavala, *Semántica y sintaxis en la lingüística transformatoria*, vol. I).
- : «Knowledge of Language», en K. Gunderson (recop.), *Language, Mind and Knowledge*, University of Minnesota Press, Minneapolis, 1975.
- : *Reflections on Language*, Pantheon Books, Nueva York, 1975 (*Reflexiones sobre el lenguaje*, Ariel, Barcelona, 1979).
- : *Dialogues avec... por Mitsou Ronat*, Flammarion, París, 1977. (*Conversaciones con...*, Granica, Barcelona, 1978.)
- : *Rules and Representations*, Columbia U. Press, Nueva York, 1980.
- Chomsky, N., y Miller, G. A.: «Introduction to the formal analysis of natural languages», en R. D. Luce, R. Bush y E. Galanter (recops.), *Handbook of Mathematical Psychology*, II, Wiley, Nueva York, 1963. (*El análisis formal de los lenguajes naturales*, Alberto Corazón Editor, Madrid, 1972).
- Christensen, N. E.: *On the Nature of Meanings. A Philosophical Analysis*, Munksgaard, Copenhagen, 1962. (*Sobre la naturaleza del significado*, Labor, Barcelona, 1968.)
- Church, A.: *Introduction to Mathematical Logic*, I, Princeton University Press, 1956.
- Dahrendorf, R.: *Sozialen Klassen und Klassenkonflikt in der industriellen Gesellschaft*, Enke, Stuttgart, 1957. (*Las clases sociales y su conflicto en la sociedad industrial*, Rialp, Madrid, 1974, 3.ª ed.)
- Danto, A.: *What Philosophy Is*, Harper & Row, Londres, 1968. (*Qué es filosofía*, Alianza, Madrid, 1976.)
- Davidson, D.: «Truth and Meaning», *Synthese*, 1967; incluido en la recopilación de J. Rosenberg y Ch. Travis, *Readings in the Philosophy of Language*, Prentice-Hall, Englewood-Cliffs, 1971, por donde he citado.
- : «Semantics for Natural Languages», en el volumen colectivo *Linguaggi nella società e nella tecnica*, Edizioni di Comunità, Milán, 1970.
- : «In Defense of Convention T», en la recopilación de H. Leblanc, *Truth, Syntax and Modality*, North Holland, Amsterdam, 1973.
- : «Radical Interpretation», *Dialectica*, 1973.
- : «Belief and the Basis of Meaning», *Synthese*, 1974.
- : «Reply to Foster», en la recopilación de G. Evans y J. McDowell, *Truth and Meaning. Essays in Semantics*, Oxford University Press, 1976.
- : «Moods and Performances», en A. Margalit (recop.), *Meaning and Use*, Reidel, Dordrecht, 1979.
- : *Inquiries into Truth and Interpretation*, Clarendon Press, Oxford, 1984, incluyendo los artículos anteriores.
- Davidson, D., y Harman, G. (recops.): *Semantics of Natural Language*, Reidel, Dordrecht, 1972.
- Davidson, D., y Hintikka, J. (recops.): *Words and Objections. Essays on the Work of W. V. Quine*, Reidel, Dordrecht, 1969.
- Davis, M.: *Computability and Unsolvability*, McGraw-Hill, Nueva York, 1958.
- Demonte, V.: *La subordinación sustantiva*, Cátedra, Madrid, 1977.
- Derwing, B. L.: *Transformational Grammar as a Theory of Language Acquisition*, Cambridge University Press, 1973.
- Diccionario de la lengua española*, Real Academia Española, Espasa-Calpe, Madrid, 1970, 19.ª ed.
- D'Introno, F.: *Sintaxis transformacional del español*, Cátedra, Madrid, 1979.
- Dorfman, A., y Mattelart, A.: *Para leer al Pato Donald*, Siglo XXI, Buenos Aires, 1972.
- Dubois, J., y otros: *Dictionnaire de la linguistique*, Larousse, París, 1973. (*Diccionario de lingüística*, Alianza Madrid, 1979.)
- Ducrot, O., y Todorov, T.: *Dictionnaire encyclopédique des sciences du langage*, Du - Seuil, París, 1972. (*Diccionario enciclopédico de las ciencias del lenguaje*, Siglo XXI, Buenos Aires, 1974.)
- Dummett, M.: *Frege. Philosophy of Language*, Duckworth, Londres, 1973.
- Eco, U.: *Segno*, ISEDI, Milán, 1973. (*Signo*, Labor, Barcelona, 1976.)

- : *Trattato di Semiotica Generale*, Bompiani, Milán, 1975. (*Tratado de semiótica general*, Lumen, Barcelona, 1977.)
- Edwards, P. (director): *The Encyclopedia of Philosophy*, 8 vols., Macmillan, Londres, 1967.
- Ernout, A., y Meillet, A.: *Dictionnaire étymologique de la langue latine*, Klincksieck, París, 1951, 3.ª ed.
- Evans, G., y McDowell, J. (recops.): *Truth and Meaning. Essays in Semantics*, Oxford University Press, 1976.
- Fann, K. T.: *Wittgenstein's Conception of Philosophy*, Blackwell, Oxford, 1969. (*El concepto de filosofía en Wittgenstein*, Tecnos, Madrid, 1975.)
- Ferrater Mora, J.: *Indagaciones sobre el lenguaje*, Alianza, Madrid, 1970.
- : «Pinturas y modelos», en *Filosofía y ciencia en el pensamiento español contemporáneo*, Tecnos, Madrid, 1973; incluido también en la obra del autor *Las palabras y los hombres*, Península, Barcelona, 1972.
- : *Diccionario de filosofía*, Alianza, Madrid, 1979, 6.ª ed.
- Fioravanti, E.: *El concepto de modo de producción*, Península, Barcelona, 1972.
- Fishman, J.: «A Systematization of the Whorfian Hypothesis», *Behavioral Science*, 1960.
- Foster, J. A.: «Meaning and Truth Theory», en la recopilación de G. Evans y J. McDowell, *Truth and Meaning. Essays in Semantics*, Oxford University Press, 1976.
- Fouts, R.: «Acquisition and Testing of Gestural Signs in Four Young Chimpanzees», *Science*, 1973. («La adquisición y comprobación del uso de los signos gestuales en cuatro chimpancés jóvenes», en Sánchez de Zavala (recop.), *Sobre el lenguaje de los antropoides*.)
- Frege, G.: *Begriffsschrift*, Halle, 1879. (*Conceptografía*, Universidad Nacional Autónoma de México, 1972.)
- : «Funktion und Begriff», 1891, incluido en *Kleine Schriften*. («Función y concepto», en *Estudios sobre semántica*.)
- : «Ueber Sinn und Bedeutung», 1892, incluido en *Kleine Schriften*. («Sobre sentido y referencia», en *Estudios sobre semántica*.)
- : «Ueber Begriff und Gegenstand», 1892, incluido en *Kleine Schriften*. («Sobre concepto y objeto», en *Estudios sobre semántica*.)
- : «Ausführungen über Sinn und Bedeutung», 1892-1895, incluido en *Nachgelassene Schriften*. («Consideraciones sobre sentido y referencia», en *Estudios sobre semántica*.)
- : «Der Gedanke», 1918-1919, incluido en *Kleine Schriften*.
- : *Translations from the Philosophical Writings of Gottlob Frege*, recop. por P. Geach y M. Black, Blackwell, Oxford, 1952.
- : *Kleine Schriften*, recop. por I. Angelelli, Georg Olms, Hildesheim, 1967.
- : *Nachgelassene Schriften*, Hamburgo, 1969.
- : *Estudios sobre semántica*, Ariel, Barcelona, 1971.
- : *Escritos lógico-semánticos*, Tecnos, Madrid, 1974.
- Frisch, K. von: *Aus dem Leben der Bienen*, Springer, Berlín, 1969 (ed. revis.). (*La vida de las abejas*, Labor, Barcelona, 1976.)
- Gadamer, H. G.: *Wahrheit und Methode*, Mohr, Tubinga, 1960. (*Verdad y método*, Sígueme, Salamanca, 1977.)
- Galmiche, M.: *Sémantique Générative*, Larousse, París, 1975. (*La semántica generativa*, Gredos, Madrid.)
- Gallego, A.: «Fetomonas», *Boletín Informativo de la Fundación Juan March*, abril, 1978.
- García Suárez, A.: *La lógica de la experiencia. Wittgenstein y el problema del lenguaje privado*, Tecnos, Madrid, 1976.
- Gardner, R. A., y Gardner, B. T.: «Teaching Sign Language to a Chimpanzee», *Science*, 1969. («Cómo enseñar el lenguaje de los sordomudos a un chimpancé», en Sánchez de Zavala (recop.), *Sobre el lenguaje de los antropoides*.)
- : «Early Signs of Language in Child and Chimpanzee», *Science*, 1975. («Algunos signos tempranos de lenguaje en el niño y el chimpancé», en Sánchez de Zavala (recop.), *Sobre el lenguaje de los antropoides*.)

- Gazdar, G.: *Pragmatics*, Academic Press, Nueva York, 1979.
- Geach, P.: *Mental acts*, Routledge, Londres, 1957.
- Goodman, N.: «The Epistemological Argument», *Synthèse*, 1967 (trad. castellana en *Teorema*, 1973).
- Gracia, F.: (recop.): *Presentación del lenguaje*, Taurus, Madrid, 1972.
- : «La paradoja del mentiroso en los lenguajes naturales», en la recopilación de F. Gracia, J. Muguerza y V. Sánchez de Zavala, *Teoría y sociedad*, Ariel, Barcelona, 1970.
- Greenberg, J. H.: «Some Universals of Grammar with Particular Reference to the Order of Meaningful Elements», en la recopilación del mismo autor, *Universals of Language*, MIT Press, Cambridge, Mass., 1963.
- : *Language Universals*, Mouton, La Haya, 1966.
- (recop.): *Universals of Human Language*, Stanford University Press, 1978, cuatro volúmenes.
- Greene, J.: *Psycholinguistics. Chomsky and Psychology*, Penguin, Harmondsworth, 1972.
- Grice, H. P.: «Meaning», *The Philosophical Review*, 1957; incluido en la recopilación de P. F. Strawson, *Philosophical Logic*, Oxford University Press, 1967, por donde cito.
- : «The Causal Theory of Perception», *Aristotelian Society Supplementary Volume*, 1961; incluido en la recopilación de G. J. Warnock, *The Philosophy of Perception*, Oxford University Press, 1967, por donde cito. (*La filosofía de la percepción*, Fondo de Cultura Económica, México.)
- : «Utterer's Meaning, Sentence-Meaning and Word-Meaning», *Foundations of Language*, 1968; incluido en la recopilación de J. R. Searle, *The Philosophy of Language*, Oxford University Press, 1971, por donde cito.
- : «Utterer's Meaning and Intentions», *The Philosophical Review*, 1969.
- : «Logic and Conversation», en la recopilación de P. Cole y J. L. Morgan, *Syntax and Semantics*, vol. III, *Speech Acts*, Academic Press, Nueva York, 1975.
- Grice, H. P., y Strawson, P. F.: «In Defense of a Dogma», *The Philosophical Review*, 1956; incluido en la recopilación de L. Sumner y J. Woods, *Necessary Truth*, Random House, Nueva York, 1969.
- Griffin, J.: *Wittgenstein's Logical Atomism*, Oxford University Press, 1964.
- Gross, M.: *Mathematical Models in Linguistics*, Prentice-Hall, Englewood Cliffs, 1972. (*Modelos matemáticos en lingüística*, Gredos, Madrid, 1976.)
- Gross, M., y Lentin, A.: *Notions sur les grammaires formelles*, Gauthier-Villars, París, 1970, 2.^a ed. (*Nociones sobre las gramáticas formales*, Tecnos, Madrid, 1976.)
- Haack, S.: *Philosophy of Logics*, Cambridge University Press, 1978.
- Habermans, J.: *Erkenntnis und Interesse*, Suhrkamp, Frankfurt, 1968 («Epílogo» añadido en 1973).
- Hacker, P. M. S.: *Insight and Illusion. Wittgenstein on Philosophy and the Metaphysics of Experience*, Oxford University Press, 1972.
- Hacking, I.: *Why Does Language matter to Philosophy?*, Cambridge University Press, 1975.
- Hadlich, R. L.: *A Transformational Grammar of Spanish*, Prentice-Hall, Englewood Cliffs, 1971. (*Gramática transformativa del español*, Gredos, Madrid, 1975.)
- Hare, R. M.: *The Language of Morals*, Oxford University Press, 1952. (*El lenguaje de la moral*, Universidad Nacional Autónoma de México, 1975.)
- : *Freedom and Reason*, Oxford University Press, 1963.
- : «Austin's Distinction between Locutionary and Illocutionary Acts», en la obra del autor, *Practical Inferences*, MacMillan, Londres, 1971.
- Harnad, S., Steklis, H., y Lancaster, J. (recops.): *Origins and Evolution of Language and Speech*, New York Academy of Sciences, 1976.
- Hartnack, J.: *Wittgenstein og den moderne filosofi*, Gyldendalske Boghandel, Copenhagen, 1962. (*Wittgenstein y la filosofía contemporánea*, Ariel, Barcelona, 1972.)
- Heidegger, M.: *Was ist Metaphysik?*, Frankfurt am Main, 1929. («¿Qué es metafísica?», en *Cruz y Raya*, 1933; reeditado por Cruz del Sur, Santiago de Chile, 1963.)

- Hempel, C. G.: «Problems and Changes in the Empiricist Criterion of Meaning», *Revue Internationale de Philosophie*, 1950; incluido en la recopilación de A. J. Ayer. *Logical Positivism*, The Free Press, Nueva York, 1959. («Problemas y cambios en el criterio empirista de significado», en la traducción de la recopilación citada de Ayer, y también en la recopilación de M. Bunge, *Antología semántica*, véanse ambas en sus lugares respectivos.)
- : «The Concept of Cognitive Significance: a Reconsideration», *Proceedings of the American Academy of Arts and Sciences*, 1951.
- : «Empiricist Criteria of Cognitive Significance: Problems and Changes», en la obra del autor *Aspects of Scientific Explanation*, The Free Press, Nueva York, 1965.
- Herriot, P.: *An Introduction to the Psychology of Language*, Methuen, Londres, 1970. (*Introducción a la psicología del lenguaje*, Labor, Barcelona, 1977.)
- Hierro S. Pescador, J.: *Problemas del análisis del lenguaje moral*, Tecnos, Madrid, 1970.
- : *La teoría de las ideas innatas en Chomsky*, Labor, Barcelona, 1976.
- : «Lenguaje, ideología y clases sociales», *Sistema*, 1978.
- Hintikka, J.: «On the Limitations of Generative Grammar», *Proceedings of the Scandinavian Seminar on Philosophy of Language*, Upsala, 1975.
- : «Quantifiers in Natural Languages: Some Logical Problems», en la recopilación de E. Saarinen, *Game-Theoretical Semantics*, Reidel, Dordrecht, 1979.
- : «Language-Games», *Acta Philosophica Fennica*, 1976; incluido en la recopilación de E. Saarinen, *Game-Theoretical Semantics*, Reidel, Dordrecht, 1979.
- : «On Any-Thesis and the Methodology of Linguistics», en *Proceedings of the 6th International Congress of Logic, Methodology and Philosophy of Science*, North-Holland, Amsterdam.
- Hockett, Ch. F.: «The Problem of Universals in Language», en J. Greenberg (recop.), *Universals of Language*.
- : *The State of the Art*, Mouton. La Haya, 1968. (*El estado actual de la lingüística*, Akal, Madrid, 1974.)
- Holdcroft, D.: *Words and Deeds. Problems in the Theory of Speech Acts*, Oxford University Press, 1978.
- Houdebine, J. L.: *Langage et marxisme*, Klincksieck, París, 1977.
- Husserl, E.: «Philosophie als strenge Wissenschaft», *Logos*, 1911 (y separadamente, Klostermann, Frankfurt, 1965). (*La filosofía como ciencia estricta*, Nova, Buenos Aires, 1962.)
- Iglesias Rozas, María T.: *Linguistic Representation: A Study of B. Russell and L. Wittgenstein, 1912-1922*, tesis doctoral inédita, Universidad de Oxford, 1979.
- Ishiguro, H.: «Use and Reference of Names», en la recopilación de P. Winch, *Studies in the Philosophy of Wittgenstein*, Routledge and Kegan Paul, Londres, 1969. (*Estudios sobre las filosofías de Wittgenstein*, Editorial Universitaria de Buenos Aires, 1971.)
- Jackendoff, R. S.: *Semantic Interpretation in Generative Grammar*, M. I. T. Press, Cambridge, Mass., 1972.
- Jakobson, R.: «Lingüística y poética», en *Ensayos de lingüística general*, Seix Barral, Barcelona, 1975.
- Jakobson, R., y Halle, M.: *Fundamentals of Language*, Mouton, La Haya, 1956. (*Fundamentos del lenguaje*, Ciencia Nueva, Madrid, 1967.)
- Janik, A., y Toulmin, S.: *Wittgenstein's Vienna*, Simon and Schuster, Nueva York, 1973. (*La Viena de Wittgenstein*, Taurus, Madrid, 1974.)
- Kant, I.: *Kritik der reinen Vernunft*, 1781. (*Crítica de la razón pura*, Alfaguara, Madrid, 1978.)
- Kasher, A.: «Conversational Maxims and Rationality», en la recopilación del autor, *Language in Focus: Foundations, Methods and Systems*, Reidel, Dordrecht, 1976.
- Katz, J.: *The Philosophy of Language*, Harper & Row, Londres, 1966. (*Filosofía del lenguaje*, Martínez Roca, Barcelona, 1971.)

- : *Semantic Theory*, Harper & Row, Nueva York, 1972. (*Teoría semántica*, Aguilar, Madrid, 1979.)
- Katz, J., y Fodor, J.: «What's Wrong with the Philosophy of Language?», *Inquiry*, 1962.
- : «The Structure of a Semantic Theory», *Language*, 1963, en J. Fodor y J. Katz (recops.), *The Structure of Language*, Prentice-Hall, Englewood Cliffs, 1964. (*La estructura de una teoría semántica*, Siglo XXI, México, 1976.)
- Katz, J., y Postal, P.: *An Integrated Theory of Linguistic Descriptions*, M. I. T. Press, Cambridge, Mass., 1964.
- Katzner, K.: *The Languages of the World*, Routledge & Kegan Paul, Londres, 1977.
- Keenan, E. L.: «The Logical Diversity of Natural Languages», en la recopilación de Harnad, Steklis y Lancaster.
- Keenan, E. O.: «On the Universality of Conversational Implicatures», *Language in Society*, 1976.
- Kenny, A.: *Wittgenstein*, Penguin, Harmondsworth, 1973. (*Wittgenstein*, Alianza Universidad, Madrid, 1974.)
- Körner, S.: *Fundamental Questions of Philosophy*, Penguin, 1969. (*¿Qué es filosofía?*, Ariel, Barcelona, 1976.)
- Korsch, K.: *Marxismus und Philosophie*, Leipzig, 1923. (*Marxismo y filosofía*, Ariel, Barcelona.)
- Kraft, V.: *Der Wiener Kreis*, Springer, Viena, 1950. (*El Círculo de Viena*, Taurus, Madrid, 1966.)
- Kripke, S.: «Naming and Necessity», en la recopilación de D. Davidson y G. Harman, *Semantics of Natural Language*, Reidel, Dordrecht, 1972, por donde cito; como libro independiente, Blackwell, Oxford, 1980.
- Kutschera, F. von: *Sprachphilosophie*, Fink, Munich, 1971. (*Filosofía del lenguaje*, Gredos, Madrid, 1979.)
- Kuznetsov, I. V.: «Pero la filosofía es una ciencia», en Ayer, Gellner y Kuznetsov.
- Lakoff, G.: *Irregularity in Syntax*, Holt, Rinehart & Winston, Nueva York, 1970.
- : «Linguistics and Natural Logic», *Synthese*, 1970; recogido en Davidson y Harman, *Semantics of Natural Language*.
- : «On Generative Semantics», en la recopilación de Steinberg y Jakobovits, *Semantics*, Cambridge University Press, 1971. («Sobre la semántica generativa», en la recopilación de Sánchez de Zavala, *Semántica y sintaxis en la lingüística transformatoria*, vol. I.)
- Lambert, K., y Brittan, G.: *An Introduction to the Philosophy of Science*, Prentice-Hall, Englewood Cliffs, 1970. (*Introducción a la filosofía de la ciencia*, Guadarrama, Madrid, 1975.)
- Lambert, K., y van Fraassen, B. C.: *Derivation and Counterexample*, Dickenson, Encino (Cal.), 1972.
- Lázaro Carreter, F.: *Diccionario de términos filológicos*, Gredos, Madrid, 1974, 3.ª edición corregida.
- Lazerowitz, M.: «Moore's Paradox», en la recopilación de P. A. Schilpp, *The Philosophy of G. E. Moore*, Open Court, La Salle (Ill.), 1942.
- Leech, G.: *Semantics*, Penguin, Harmondsworth, 1974. (*Semántica*, Alianza, Madrid, 1977.)
- Lemmon, E. J.: «Sentences, Statements and Propositions», en la recopilación de B. Williams y A. Montefiore, *British Analytical Philosophy*, Routledge & Kegan Paul, Londres, 1966.
- Lenneberg, E. H.: *Biological Foundations of Language*, Wiley and Sons, Nueva York, 1967. (*Fundamentos biológicos del lenguaje*, Alianza Universidad, Madrid, 1975.)
- : «A Biological Perspective of Language», en Lenneberg (recop.).
- (recop.): *New Directions in the Study of Language*, M. I. T. Press, Cambridge, Mass., 1964. (*Nuevas direcciones en el estudio del lenguaje*, Revista de Occidente, Madrid.)
- Lepschy, G. C.: *La linguistica strutturale*, Einaudi, Turín, 1966. (*La lingüística estructural*, Anagrama, Barcelona, 1971.)

- Lewis, D.: *Convention*, Harvard University Press, 1969.
- : «General Semantics», *Synthèse*, 1970, recogido en Davidson y Harman, *Semantics of Natural Language*, Reidel, Dordrecht, 1972.
- : «Index, Context and Content», en la recopilación de S. Kanger y S. Oehman, *Philosophy and Grammar*, Reidel, Dordrecht, 1981.
- Lewy, C.: «A Note on the Text of the *Tractatus*», *Mind*, 1967.
- Lieberman, P.: «On the Evolution of Language: A Unified View», *Cognition*, 1973. («Un enfoque unitario de la evolución del lenguaje», en la recopilación de Sánchez de Zavala, *Sobre el lenguaje de los antropoides*.)
- : *On the Origins of Language*, MacMillan, Nueva York, 1975.
- List, G.: *Psycholinguistik. Eine Einführung*, Kohlhammer, Stuttgart, 1972. (*Introducción a la psicolingüística*, Gredos, Madrid, 1977.)
- Lyons, J.: *Introduction to Theoretical Linguistics*, Cambridge University Press, 1968. (*Introducción en la lingüística teórica*, Teide, Barcelona, 1971.)
- : *Chomsky*, Fontana, Londres, 1970. (*Chomsky*, Grijalbo, Barcelona, 1974.)
- Lledó, E.: *La filosofía, hoy*, Salvat, Barcelona, 1975.
- MacCorquodale, K.: «On Chomsky's Review of Skinner's *Verbal Behavior*», *Journal of the Experimental Analysis of Behavior*, 1970 (trad. al castellano en la recopilación de R. Bayés).
- Malcolm, N.: «Moore and Ordinary Language», en la recopilación de P. A. Schilpp, *The Philosophy of G. E. Moore*, Open Court, La Salle (Ill.), 1942. («Moore y el lenguaje común», en la recopilación de V. C. Chappell, *El lenguaje común*, Tecnos, Madrid, 1971.)
- Malmberg, B.: *Teckenlära*, 1973. (*Teoría de los signos*, Siglo XXI, México, 1977.)
- Marcellesi, J. B., y Gardin, B.: *Introduction à la sociolinguistique. La linguistique sociale*, Larousse, París, 1974. (*Introducción a la sociolingüística*, Gredos, Madrid.)
- Marías, J.: *Ortega. I. Circunstancia y vocación*, Alianza Universidad, Madrid, 1960.
- : *La realidad histórica y social del uso lingüístico*, Real Academia Española, Madrid, 1965.
- Marler, P.: «Animal Communication», en L. Krames, P. Pliner y T. Alloway (recops.), *Nonverbal Communication*, Plenum Press, Nueva York, 1974.
- Marr, N.: «Ueber die Entstehung der Sprache», *Unter dem Banner des Marxismus*, 1925-26.
- Marx, C., y Engels, F.: *Die deutsche Ideologie*, 1846. (*La ideología alemana*, Grijalbo, Barcelona, 1970.)
- Mattelart, A.: «El marco del análisis ideológico», *Cuadernos de la realidad nacional* (Chile), 1970.
- McCawley, J. D.: «Where do Noun Phrases Come from?», en la recopilación de Jacobs y Rosenbaum, *Readings in English Transformational Grammar*, Ginn, Waltham, Mass., 1970. («¿De dónde proceden los sintagmas nominales?», en la recopilación de Sánchez de Zavala, *Semántica y sintaxis en la lingüística transformatoria*, vol. I.)
- McNeill, D.: «The Creation of Language by Children», en Lyons y Wales (recops.), *Psycholinguistics Papers*, Edinburgh University Press, 1966.
- : *The Acquisition of Language*, Harper and Row, Nueva York, 1970.
- Mill, J. S.: *A System of Logic Ratiocative and Inductive*, 1843 (vols. VII y VIII de sus *Collected Works*, edición de J. M. Robson, University of Toronto Press-Routledge & Kegan Paul, Toronto y Londres, 1973-74).
- Miller, D.: «The Uniqueness of Atomic Facts in Wittgenstein's *Tractatus*», *Theoria*, 1977.
- Moliner, M.: *Diccionario de uso del español*, Gredos, Madrid, 1967, 2 vols.
- Montague, R.: «Universal Grammar», *Theoria*, 1970; incluido en *Formal Philosophy*. («Gramática universal», en *Ensayos de filosofía formal*, véase más abajo).
- : «English as a Formal Language», en el volumen colectivo *Linguaggi nella società e nella tecnica*, Edizioni di Comunità, Milán, 1970; incluido en *Formal Philosophy* (pero no en la traducción española de esta obra).

- : *Formal Philosophy*, Yale University Press, 1974. (*Ensayos de filosofía formal*, traducción parcial de la obra anterior, Alianza, Madrid, 1977.)
- Moore, G. E.: «A Defence of Common Sense», 1925; incluido en *Philosophical Papers*, Allen & Unwin, Londres, 1959; lo he citado por la edición de Collier Books, 1962. («Defensa del sentido común», en *Defensa del sentido común y otros ensayos*, Taurus, Madrid, 1972.)
- : «An Autobiography», en la recopilación de P. A. Schilpp, *The Philosophy of G. E. Moore*, Open Court, La Salle (Ill.), 1942.
- : «A Reply to my Critics», en la recopilación que se acaba de citar.
- : «Wittgenstein's Lectures in 1930-33», *Mind*, 1954-55; incluido en *Philosophical Papers*, citado más arriba («Las conferencias de Wittgenstein de 1930-33», en *Defensa del sentido común y otros ensayos*, citado más arriba.)
- Morris, Ch.: «Foundations of the Theory of Signs», *International Encyclopedia of Unified Science*, vol. 1, núm. 2, University of Chicago Press, 1938. (*Fundamentos de la teoría de los signos*, Universidad Nacional de México, 1958; los dos primeros capítulos en Gracia (recop.), *Presentación del lenguaje*.)
- : *Signs, Language and Behavior*, Prentice-Hall, Nueva York, 1946. (*Signos, lenguaje y conducta*, Losada, Buenos Aires, 1962.)
- : *Signification and Significance*, M. I. T. Press, Cambridge, Mass., 1964. (*La significación y lo significativo*, Alberto Corazón, Madrid, 1974.)
- Morton, J.: «What Could Possibly Be Innate?», en J. Morton (recop.), *Biological and Social Factors in Psycholinguistics*, Logos Press, Londres, 1971.
- Moulines, U.: «Lo analítico y lo sintético: dualismo admisible», *Teorema*, 1973.
- : *La estructura del mundo sensible (Sistemas fenomenalistas)*, Ariel, Barcelona, 1973.
- Mounce, H. O.: *Wittgenstein's Tractatus. An Introduction*, Blackwell, Oxford, 1981. (*Introducción al «Tractatus» de Wittgenstein*, Tecnos, Madrid, 1983.)
- Mounin, G.: *Dictionnaire de la linguistique*, Presses Universitaires de France, París, 1974. (*Diccionario de Lingüística*, Labor, Barcelona, 1979.)
- Muguerza, J. (recop.): *La concepción analítica de la filosofía*, Alianza, Madrid, 1974, 2 vols.
- Neurath, O.: «Protokollsätze», *Erkenntnis*, 1932-33; trad. inglesa en la recopilación de A. J. Ayer, *Logical Positivism*, The Free Press, Nueva York, 1959. («Proposiciones protocolares», en *El positivismo lógico*, Fondo de Cultura Económica, México, 1965.)
- Nowell-Smith, P. H.: *Ethics*, Penguin, Harmondsworth, 1954. (*Ética*, Verbo Divino, Estella, 1977.)
- : «Contextual Implication and Ethical Theory», *Aristotelian Society Supplementary Volume*, 1962.
- Occam, G. de: *Summa Totius Logicae*, 1488 (edición parcial de Philoteus Boehner, Nueva York, 1951-54).
- : *Super quattuor libros sententiarum*, 1495.
- Ogden, C. K., y Richards, I. A.: *The Meaning of Meaning*, Kegan Paul, Londres, 1923. (*El significado del significado*, Paidós, Buenos Aires, 1954.)
- Ortega y Gasset, J.: «Meditación del marco», *El espectador*, vol. III, 1921, *Obras Completas*, II, Revista de Occidente, Madrid, 1946.
- : «Ensayo de estética a manera de prólogo», 1914, en *Obras completas*, VI, Revista de Occidente, Madrid, 1947.
- : *El hombre y la gente*, Revista de Occidente, Madrid, 1957.
- : *¿Qué es filosofía?*, Revista de Occidente, Madrid, 1957.
- Parkin, F.: *Class Inequality and Political Order*, Londres, 1971. (*Orden político y desigualdades de clase*, Debate, Madrid, 1978.)
- Parkinson, G. (recop.): *The Theory of Meaning*, Oxford University Press, 1968. (*La teoría del significado*, Fondo de Cultura Económica, México.)
- Partee, B.: «Possible Worlds Semantics and Linguistic Theory», *The Monist*, 1977.
- Patterson, F.: «Conversations with a Gorilla», *National Geographic Magazine*, 1978.
- Pears, D.: *Wittgenstein*, Fontana-Collins, Londres, 1971. (*Wittgenstein*, Grijalbo, Barcelona, 1973.)

- Peirce, C. S.: *Collected Papers*, Harvard University Press, 6 vols., 1931-1935.
- Peters, S.: «Why There Are Many "Universal" Bases», *Papers in Linguistics*, 1970.
- Peters, S., y Richtie, R.: «On Restricting the Base Component of Transformational Grammars», *Information and Control*, 1971.
- : «Nonfiltering and Local-filtering Transformational Grammars», en J. Hintikka, J. Moravcsik y P. Suppes (recops.), *Approaches to Natural Language*, Reidel, Dordrecht, 1973.
- Piaget, J.: *Le langage et la pensée chez l'enfant*, Delachaux et Niestlé, Neuchâtel, 1923. (*El lenguaje y el pensamiento en el niño*, Guadalupe, Buenos Aires, 1973.)
- : *La formation du symbole chez l'enfant*, Delachaux et Niestlé, Neuchâtel, 1945. (*La formación del símbolo en el niño*, Fondo de Cultura Económica, México, 1961.)
- : *Psychologie de l'intelligence*, Armand Colin, París, 1947. (*Psicología de la inteligencia*, Psique, Buenos Aires, 1960.)
- : *Six études de psychologie*, Gonthier, París, 1964. (*Seis estudios de psicología*, Barral, Barcelona, 1971.)
- : *Sagesse et illusions de la philosophie*, Presses Universitaires de France, 1965. (*Sabiduría e ilusiones de la filosofía*, Península, Barcelona, 1970.)
- : *Le structuralisme*, Presses Universitaires de France, París, 1968. (*El estructuralismo*, Proteo, Buenos Aires.)
- : *L'epistémologie génétique*, Presses Universitaires de France, París, 1970. (*La epistemología genética*, Redondo, Barcelona.)
- : *Problèmes de psychologie génétique*, Denoël-Gonthier, París, 1972. (*Problemas de psicología genética*, Ariel, Barcelona, 1975.)
- Piaget, J., e Inhelder, B.: *La psychologie de l'enfant*, Presses Universitaires de France, París, 1966. (*Psicología del niño*, Morata, Madrid, 1969.)
- Pinborg, J.: *Logik und Semantik im Mittelalter*, Fromman-Holzboog, Stuttgart, 1972.
- Pinillos, J. L.: «Lenguaje, individuo y sociedad», *Psicometría e investigación psicológica*, 1970.
- : *Principios de psicología*, Alianza, Madrid, 1975.
- : «Comunicación animal y lenguaje humano», en la obra colectiva *Comunicación y lenguaje*, Karpos, Madrid, 1977.
- Pitcher, G. (recop.): *Wittgenstein: The Philosophical Investigations*, Doubleday, Nueva York, 1966.
- Platón: *Parménides*.
- Platts, M. de B.: *Ways of Meaning. An Introduction to a Philosophy of Language*, Routledge & Kegan Paul, Londres, 1979.
- Ponzio, A.: *Produzione linguistica e ideologia sociale*, De Donato, Bari, 1973. (*Producción lingüística e ideología social*, Alberto Corazón, Madrid, 1974.)
- Popper, K. R.: *Logik der Forschung*, Springer, Viena, 1935; trad. inglesa ampliada, *The Logic of Scientific Discovery*, Hutchinson, Londres, 1959. (*La lógica de la investigación científica*, Tecnos, Madrid, 1962.)
- : «The Demarcation between Science and Metaphysics», en la recopilación de P. A. Schilpp, *The Philosophy of Rudolf Carnap*, Open Court, La Salle (Ill.), 1963; incluido en la obra del autor, *Conjectures and Refutations*, Routledge & Kegan Paul, Londres, 1963. («La demarcación entre la ciencia y la metafísica», en *El desarrollo del conocimiento científico*, Paidós, Buenos Aires, 1967.)
- : «Philosophical Comments on Tarski's Theory of Truth», en la obra del autor *Objective Knowledge. An Evolutionary Approach*, Oxford University Press, 1972. («Comentarios filosóficos sobre la teoría de la verdad de Tarski», en *Conocimiento objetivo*, Tecnos, Madrid, 1974.)
- Poulantzas, N.: *Les classes sociales dans le capitalisme aujourd'hui*, Du Seuil, París, 1974. (*Las clases sociales en el capitalismo actual*, Siglo XXI, México, 1976.)
- Premack, A. J., y Premack, D.: «Teaching Language to an Ape», *Scientific American*, 1972.
- Premack, D.: «Language in Chimpanzee?», *Science*, 1971.
- : «Some General Characteristics of a Method for Teaching Language to Organisms that do not ordinarily Acquire it», en L. Jarrard (recop.), *Cognitive Processes of*

- Nonhuman Primates*, Academic Press, Nueva York, 1971. («Algunas características generales de un método para enseñar el lenguaje a organismos que normalmente no lo adquieren», en Sánchez de Zavala (recop.), *Sobre el lenguaje de los antropoides*.)
- : *Intelligence in Ape and Man*, L. Erlbaum, Hillsdale, 1976.
- Putnam, H.: «The Analytic and the Synthetic», en la recopilación de H. Feigl y G. Maxwell, *Scientific Explanation, Space and Time*, University of Minnesota Press, 1962.
- : «The Innateness Hypothesis and Explanatory Models in Linguistics», *Synthèse*, 1967 (trad. castellana en *Teorema*, 1973).
- Quesada, D.: *La lingüística generativo-transformacional: supuestos e implicaciones*, Alianza, Madrid, 1974.
- : «Lógica y gramática en Richard Montague», *Convivium*, 1976.
- Quine, W. van O.: «On What There Is», *Review of Metaphysics*, 1948; incluido en *From a Logical Point of View*, véase más abajo. («Sobre lo que hay», en *Desde un punto de vista lógico*.)
- : *Mathematical Logic*, Harvard University Press, 1951 (ed. revisada). (*Lógica matemática*, Revista de Occidente, Madrid, 1972.)
- : «Two Dogmas of Empiricism», *The Philosophical Review*, 1951; incluido en *From a Logical Point of View*, véase más abajo. («Dos dogmas del empirismo», en *Desde un punto de vista lógico*.)
- : *From a Logical Point of View*, Harvard University Press, 1953. (*Desde un punto de vista lógico*, Ariel, Barcelona, 1963.)
- : *Word and Object*, M. I. T. Press, Cambridge (Mass.), 1960. (*Palabra y objeto*, Labor, Barcelona, 1968.)
- : «Ontological Relativity», *The Journal of Philosophy*, 1968; incluido en *Ontological Relativity and Other Essays*, véase más abajo. («La relatividad ontológica», en *La relatividad ontológica y otros ensayos*.)
- : «Epistemology Naturalized», en *Ontological Relativity and Other Essays*. («La epistemología naturalizada», en *La relatividad ontológica y otros ensayos*.)
- : «Existence and Quantification», en *Ontological Relativity and Other Essays*. («Existencia y cuantificación», en *La relatividad ontológica y otros ensayos*.)
- : *Ontological Relativity and Other Essays*, Columbia University Press, Nueva York, 1969. (*La relatividad ontológica y otros ensayos*, Tecnos, Madrid, 1974.)
- : «Reply to Chomsky», en la recopilación de D. Davidson y J. Hintikka, *Words and Objections*, Reidel, Dordrecht, 1969.
- : «Reply to Davidson», en la recopilación que se acaba de citar.
- : *Philosophy of Logic*, Prentice-Hall, Englewood-Cliffs, 1970. (*Filosofía de la lógica*, Alianza, Madrid, 1973.)
- : «Methodological Reflections on Current Linguistic Theory», *Synthèse*, 1970, recogido en Davidson y Harman (recops.).
- : «Reflexiones filosóficas sobre el aprendizaje del lenguaje», *Teorema*, 1972.
- : *The Roots of Reference*, Open Court, La Salle (Ill.), 1973. (*Las raíces de la referencia*, Revista de Occidente, Madrid, 1977.)
- Reguera, I.: *La miseria de la razón (El primer Wittgenstein)*, Taurus, Madrid, 1980.
- Richelle, M.: «Analyse Formelle et analyse fonctionnelle du comportement verbal», *Bulletin de Psychologie*, 1972-73 (trad. castellana en la recopilación de R. Bayés).
- : *L'acquisition du langage*, Dessart et Mardaga, Bruselas, 1976, cuarta edición. (*La adquisición del lenguaje*, Herder, Barcelona.)
- Rivero, M. L.: *Estudios de gramática generativa del español*, Cátedra, Madrid, 1977.
- Rivière, A.: «El análisis experimental de la conducta y el conductismo radical como filosofía», *Investigación y ciencia*, abril 1977.
- Rodríguez Adrados, F.: *Lingüística estructural*, Gredos, Madrid, 1974, segunda edición revisada y aumentada.
- Ronat, M.: *Dialogues avec Noam Chomsky*, Flammarion, París, 1977. (*Conversaciones con Chomsky*, Granica, Barcelona, 1978.)

- Rorty, R.: «Epistemological Behaviorism and the De-Transcendentalization of Analytical Philosophy», *Neue Hefte für Philosophie*, 1978.
- Ross, A.: *Directives and Norms*, Routledge & Kegan Paul, Londres, 1968. (*Lógica de las normas*, Tecnos, Madrid, 1971.)
- Rossi-Landi, F.: «Per un uso marxiano di Wittgenstein», *Nuovi Argumenti*, 1966; en *Il linguaggio come lavoro e come mercato*, Bompiani, Milán, 1968. (*El lenguaje como trabajo y como mercado*, Monte Avila, Caracas, 1970.)
- : *Linguistics and Economics*, Mouton, La Haya, 1977.
- : *Ideologia*, Isedi, Milán, 1978. (*Ideología*, Labor, Barcelona, 1980.)
- Rumbaugh, D., y Gill, T.: «The Mastery of Language-Type Skills by the Chimpanzee (Pan)», en Harnad, Steklis y Lancaster (recops.), 1976.
- Rumbaugh, D., Gill, T., y Glasersfeld, E. von: «Reading and Sentence Completion by a Chimpanzee (Pan)», *Science*, 1973. («La lectura y el completado de oraciones realizados por un chimpancé», en Sánchez de Zavala (recop.), *Sobre el lenguaje de los antropoides*.)
- Russell, B.: «On Denoting», *Mind*, 1905; incluido en *Logic and Knowledge*, véase más abajo. («Sobre la denotación», en *Lógica y conocimiento*.)
- : *The Problems of Philosophy*, Oxford University Press, 1912. (*Los problemas de la filosofía*, Labor, Barcelona.)
- : *Theory of Knowledge*, 1913, Collected Works, vol. 7.
- : «The Philosophy of Logical Atomism», *The Monist*, 1918; en *Logic and Knowledge* (trad. castellana no sólo en la de esta obra, sino también en Muguerza (recop.), *La concepción analítica de la filosofía*, vol. I.)
- : *Introduction to Mathematical Philosophy*, Allen & Unwin, Londres, 1919.
- : «Logical Atomism», en la recopilación de J. H. Muirhead, *Contemporary British Philosophy*, serie 1, Allen & Unwin, Londres, 1924; incluido en *Logic and Knowledge*, y también en la recopilación de Ayer, *Logical Positivism*, véase. («Atomismo lógico», en *Lógica y conocimiento*, así como en *El positivismo lógico*.)
- : *Logic and Knowledge*, Allen & Unwin, Londres, 1956. (*Lógica y conocimiento*, Taurus, Madrid, 1966.)
- : *My Philosophical Development*, Allen & Unwin, Londres, 1959. (*La evolución de mi pensamiento filosófico*, Aguilar, Madrid, 1960.)
- Russell, B., y Whitehead, A. N.: *Principia Mathematica*, Cambridge University Press, 1910-1913, 3 vols.
- Ruwet, N.: *Introduction à la grammaire générative*, Plon, París, 1968. (*Introducción a la gramática generativa*, Gredos, Madrid, 1974.)
- Ryle, G.: «Philosophical Arguments», conferencia, Oxford University Press, 1945; incluido en la recopilación de Ayer, *Logical Positivism*, véase. («Argumentos filosóficos», en *El positivismo lógico*.)
- : *The Concept of Mind*, Hutchinson, Londres, 1949; hay edición en Penguin. (*El concepto de lo mental*, Paidós, Buenos Aires, 1967.)
- : «Ordinary Language», *The Philosophical Review*, 1953; incluido en la recopilación de Ch. Caton, *Philosophy and Ordinary Language*, University of Illinois Press, Urbana, 1963. («El lenguaje común», en la recopilación de V. Chappell, *El lenguaje común*, Tecnos, Madrid, 1971.)
- : *Dilemmas*, Cambridge University Press, 1954. («Dilemas», traducción de algunos extractos en la recopilación de J. Muguerza, *La concepción analítica de la filosofía*, vol. 2, Alianza, Madrid, 1974.)
- Sacristán, M.: *Introducción a la lógica y al análisis formal*, Ariel, Barcelona, 1964.
- : *Sobre el lugar de la filosofía en los estudios superiores*, Nova Terra, Barcelona, 1968.
- Sampson, G.: *The Form of Language*, Weidenfeld and Nicolson, Londres, 1975.
- Sánchez de Zavala, V.: *Indagaciones praxiológicas*, Siglo XXI, Madrid, 1973.
- (recop.): *Semántica y sintaxis en la lingüística transformatoria*, Alianza, Madrid, volumen I, 1974; vol. II, 1976.
- (recop.): *Estudios de gramática generativa*, Labor, Barcelona, 1976.
- (recop.): *Sobre el lenguaje de los antropoides*, Siglo XXI, Madrid, 1976.

- Sapir, E.: «The status of Linguistics as a Science», *Language*, 1929.
- Saussure, F. de: *Cours de linguistique générale*, Payot, París, 1916. (*Curso de lingüística general*, Losada, Buenos Aires, 1945.)
- Schaf, A.: *Wstęp do Semantiki*, 1960. (*Introducción a la semántica*, Fondo de Cultura Económica, México, 1966.)
- : «Sobre la necesidad de una investigación lingüística marxista», versión original en *Kultura i Społeczeństwo*, 1960; en *Essays über die Philosophie der Sprache*, Europa Verlag, Viena, 1968. (*Ensayos sobre filosofía del lenguaje*, Ariel, Barcelona, 1973.)
- : *Lenguaje y conocimiento*, Grijalbo, México, 1967 (original polaco, hacia 1964).
- Searle, J. R.: «Austin on Locutionary and Illocutionary Acts», *The Philosophical Review*, 1968; incluido en la recopilación de varios autores, *Essays on J. L. Austin*, Oxford University Press, 1973.
- : *Speech Acts*, Cambridge University Press, 1969. (*Actos de habla*, Cátedra, Madrid, 1980.)
- : «A Taxonomy of Illocutionary Acts», en la recopilación de K. Gunderson, *Language, Mind and Knowledge*, University of Minnesota Press, Minneapolis, 1975. («Una taxonomía de los actos ilocucionarios», *Teorema*, 1976.)
- : «Indirect Speech Acts», en la recopilación de P. Cole y J. Morgan, *Syntax and Semantics*, vol. 3, Academic Press, Nueva York, 1975. («Actos de habla indirectos», *Teorema*, 1977.)
- : *Expression and Meaning*. Cambridge U. Press, 1979, incluyendo los dos últimos artículos citados.
- Skinner, B. F.: *Verbal Behavior*, Appleton-Century-Crofts, Nueva York, 1957.
- Simpson, T. M.: *Formas lógicas, realidad y significado*, Editorial Universitaria de Buenos Aires, 1964; 2.ª ed. corregida y aumentada, 1975.
- : recopil.: *Semántica filosófica: Problemas y discusiones*, Siglo XXI, Buenos Aires, 1973.
- Specht, E. K.: *Die sprachphilosophischen und ontologischen Grundlagen im Spätwerk Ludwig Wittgensteins*, Kölner Universitätsverlag, Colonia, 1963 (he citado por la traducción inglesa, *The Foundations of Wittgenstein's Late Philosophy*, Manchester University Press, 1969).
- Stalin, J.: «Sobre el marxismo en la lingüística», en *El marxismo, la cuestión nacional y la lingüística*, Akal, Madrid, 1977 (original ruso de 1950).
- Stegmüller, W.: «A Model Theoretic Explication of Wittgenstein's Picture Theory», en *Collected Papers on Epistemology, Philosophy of Science, and History of Philosophy*, vol. I, Reidel, Dordrecht, 1977 (original alemán del artículo, 1966).
- Steinberg, D., y Jakobovits, L. (recops.): *Semantics*, Cambridge University Press, 1971.
- Strawson, P. F.: «On Referring», *Mind*, 1950; incluido en la recopilación de A. Flew, *Essays in Conceptual Analysis*, MacMillan, Londres, 1956. («Sobre el referir», en la recopilación de T. M. Simpson, *Semántica filosófica: problemas y discusiones*, véase.)
- : *Introduction to Logical Theory*, Methuen, Londres, 1952. (*Introducción a la teoría de la lógica*, Nova, Buenos Aires, 1969.)
- : «Austin and "Locutionary Meaning"», en la recopilación de varios autores, *Essays on J. L. Austin*, Oxford University Press, 1973.
- Tarski, A.: «Der Wahrheitsbegriff in den Formalisierten Sprachen», *Studia Philosophica*, 1936 (trad. inglesa en la recopilación de trabajos de Tarski, *Logic, Semantics, Metamathematics*, Oxford University Press, 1956; trad. francesa en *Logique, sémantique, métamathématique*, Armand Colin, París, 1972, recopilación más amplia que la inglesa).
- : «The Semantic Conception of Truth and the Foundations of Semantics», *Philosophy and Phenomenological Research*, 1944; incluido en la recopilación de H. Feigl y W. Sellars, *Readings in Philosophical Analysis*, Appleton-Century-Crofts, Nueva York, 1949. («La concepción semántica de la verdad y los fundamentos de la semántica», en la recopilación de M. Bunge, *Antología semántica*, Nueva Visión, Buenos Aires, 1960.)

- Tato, J. L.: *Semántica de la metáfora*, Instituto de Estudios Alicantinos, Alicante, 1975.
- Terrace, H., y Bever, T.: «What Might Be Learned from Studying Language in the Chimpanzee?», en la recopilación de Harnad, Steklis y Lancaster, 1976.
- Thiel, Ch.: *Sinn und Bedeutung in der Logik Gottlob Freges*, A. Hain, Meisenheim am Glan, 1965 (*Sentido y referencia en la lógica de Gottlob Frege*, Tecnos, Madrid, 1972).
- Thorpe, W. H.: «The Comparison of Vocal Communication in Animals and Man», en R. A. Hinde (recop.), *Non-Verbal Communication*, Cambridge U. Press, 1972.
- Trentman, J.: «Ockham on Mental», *Mind*, 1970.
- Trujillo, R.: *El silbo gomero, análisis lingüístico*, Editorial Interinsular Canaria, 1979.
- Turing, A. M.: «On Computable Numbers, with an Application to the Entscheidungsproblem», *Proceedings of the London Mathematical Society*, 1936-37.
- Ullman, S.: *Semantics*, Blackwell, Oxford, 1962 (*Semántica*, Aguilar, Madrid, 1967, 2.ª ed.).
- Urmson, J. O.: *Philosophical Analysis*, Oxford University Press, 1956. (*El análisis filosófico*, Ariel, Barcelona, 1978.)
- Urmson, J., y otros: *Concise Encyclopedia of Western Philosophy and Philosophers*, Hutchinson, Londres, 1968, edición revisada. (*Enciclopedia concisa de filosofía y filósofos*, Cátedra, Madrid, 1979.)
- Valls, A.: *Introducción a la antropología*, Labor, Barcelona.
- Varios autores: *Sobre el Tractatus Logico-Philosophicus*, Teorema, Valencia, 1972.
- : *Essays on J. L. Austin*, Oxford U. Press, 1973.
- : *Aspectos de la filosofía de W. V. Quine*, Teorema, Valencia, 1976.
- Verón, E.: «Ideología y comunicación de masas: la semantización de la violencia política», en el volumen colectivo *Lenguaje y comunicación social*, Nueva Visión, Buenos Aires, 1971.
- Voloshinov, V.: Véase Bachtin, M.
- Vygotsky, L. S.: *Thought and Language*, M. I. T. Press, Cambridge, Mass., 1962 (edición original rusa: 1934; *Pensamiento y lenguaje*, Pléyade, Buenos Aires, 1964).
- Waismann, F.: «How I See Philosophy», en H. D. Lewis (recop.): *Contemporary British Philosophy*, 3.ª serie, Allen & Unwin, Londres, 1956, y también en Ayer (recop.), *Logical Positivism*, y en el libro del autor, *How I See Philosophy*, McMillan, Londres, 1968. («Mi visión de la filosofía», en Muguerza (recop.), *La concepción analítica de la filosofía*, vol. 2.)
- Wasow, T.: «On Constraining the Class of Transformational Languages», *Synthese*, 1978.
- Weinberg, J. R.: *An Examination of Logical Positivism*, Kegan Paul, Londres, 1936. (*Examen del positivismo lógico*, Aguilar, Madrid, 1959.)
- Weinreich, H.: «Explorations in Semantic Theory», en la recopilación de T. Sebeok, *Current Trends in Linguistics*, vol. 3, Mouton, La Haya, 1966; como libro: Mouton, La Haya, 1972.
- Whorf, B. L.: *Language, Thought and Reality*, M. I. T. Press, Cambridge, Mass., 1956. (*Lenguaje, pensamiento y realidad*, Barral, Barcelona.)
- Winch, P. (recop.): *Studies in the Philosophy of Wittgenstein*, Routledge & Kegan Paul, Londres, 1969. (*Estudios sobre las filosofías de Wittgenstein*, Editorial Universitaria de Buenos Aires, 1971.)
- Wisdom, J.: «Logical Constructions», *Mind*, 1931-1933.
- : «Gods», *Proceedings of the Aristotelian Society*, 1944; incluido en la obra del autor *Philosophy and Psychoanalysis*, véase más abajo.
- : «Philosophy, Metaphysics and Psychoanalysis», en *Philosophy and Psychoanalysis*, véase a continuación. («Filosofía, metafísica y psicoanálisis», en la recopilación de J. Muguerza, *La concepción analítica de la filosofía*, vol. 2, Alianza, Madrid, 1974.)
- : *Philosophy and Psychoanalysis*, Blackwell, Oxford, 1953.
- Wittgenstein, L.: «Notes on Logic», 1913; incluido en *Notebooks 1914-1916*, véase a

- continuación. («Notas sobre lógica», en *Sobre el Tractatus Logico-Philosophicus*, Teorema, Valencia, 1972.)
- : *Notebooks 1914-1916*, Blackwell, Oxford, 1961. (*Diario filosófico 1914-1916*, Ariel, Barcelona, 1982.)
- : *Tractatus Logico-Philosophicus*, *Annalen der Naturphilosophie*, 1921 (texto alemán con traducción inglesa, Routledge & Kegan Paul, Londres, 1922 y 1961; con traducción castellana, Revista de Occidente, Madrid, 1957, y Alianza, Madrid, 1973).
- : *The Blue and Brown Books*, Blackwell, Oxford, 1958. (*Los cuadernos azul y marrón*, Tecnos, Madrid, 1968.)
- : *Philosophische Untersuchungen*, Blackwell, Oxford, 1953.
- : *Letters from Ludwig Wittgenstein with a Memoir*, recopilación de P. Engelmann, Blackwell, Oxford, 1967.
- : «Notes for Lectures on Private Experience and Sense Data», *The Philosophical Review*, 1968.
- : *Letters to C. K. Ogden*, Blackwell, Oxford, 1973.
- Wright, G. H. von: «Wittgenstein. Biographical Sketch», en N. Malcolm, *Ludwig Wittgenstein. A Memoir*, Oxford University Press, 1958.

INDICE ANALITICO

- Acero, J. J., 21
- Actos
- de comunicación, 52
 - de habla (Austin), 65, 313
 - de habla: clasificación (Austin), 318-19
 - de habla: clasificación (Searle), 321
 - fonético, fático y rético (Austin), 315
 - ilocucionarios (Austin), 316-17
 - ilocucionarios (Searle), 320-21
 - ilocucionarios: criterios de clasificación (Searle), 321-27
 - ilocucionarios: clasificación (Searle), 322 s
 - ilocucionarios y decir (Searle), 320 s
 - ilocucionarios y funciones del lenguaje (Searle, Bühler, Jakobson), 327 s
 - ilocucionarios, funciones del lenguaje y tipos de discurso (Searle, Jakobson), 340
 - ilocucionarios y verbos ilocucionarios (Searle), 323 s
 - locucionarios (Austin), 315
 - locucionarios (Searle), 321
 - perlocucionarios (Austin), 316
 - proposicionales (Searle), 321
- Actuación, 54-61
- Adquisición del lenguaje, 136-45
- Alarcos, E., 112
- Alcina, J., 111
- Ajdukiewicz, K., 126, 441
- Algoritmo, 122-24
- Alienación lingüística (Bachtin, Voloshinov), 475
- Alston, W., 20, 264
- Ambito de una oración (Carnap), véase reglas de ámbito
- Análisis reductivo, 234, 282 s
- Analogía, 64
- Analógico (Razonamiento), véase Filosofía y su método (Wisdom)
- Andrés, T. de, 42-7
- Anscombe, E., 223, 266
- Antropología filosófica, 13
- Argumento ontológico, 202
- Aristóteles, 31, 38, 45, 49
- «Ascenso Semántico» (Quine), 425
- Atomismo lógico (Russell), 192, 269
- Austin, J. L., 15, 172, 311-24, 329, 339, 341, 358, 416
- Autoincrustación, 82
- Avenarius, R., 361
- Axioma de finitud, 238
- Ayer, A. J., 13, 20, 367 s, 402, 456
- Bach, E., 80 s, 87, 124-28, 131, 134
- Bachtin, M., 472 s, 475-77
- Baldinger, K., 463
- Bayés, R., 59, 172
- Beard, R. W., 267
- Bellugi-Klima, U., 162
- Berlin, B., 74-6
- Bever, T. G., 128, 168
- Biología del lenguaje, 65, 156-58
- Black, M., 20, 141, 169, 178, 385
- Blasco, J. L., 359, 427-29
- Blecua, J. M., 111
- Bloomfield, L., 17
- Boecio, 45

- Boehner, Ph., 42
 Brittan, G., 76
 Brown, R., 139, 159, 161-65
 Bruner, J. S., 155
 Bueno, G., 10 s, 20
 Bühler, K., 328
 Bunge, M., 432, 456
 Busnel, R. G., 49
 Bustos, E., 21
- Cadena preterminal, 96
 Cadena terminal, 84
 Cambridge, Escuela de, 302
 Campbell, R., 59 s, 140
 Capacidad generativa débil, 80
 Capacidad generativa fuerte, 80
 Características segmentales, 109
 Características suprasegmentales, 109
 Carnap, R., 9, 13, 15, 21, 39 s, 46, 102, 172, 177, 179, 186, 245, 258, 262, 269, 360, 378, 386-402, 404, 416 s, 422, 425, 429, 442, 445, 456, 462
 Casares, J., 100, 103
 Cassirer, E., 15, 27s
 Castilla del Pino, C., 461
 Categorema, 72
 Categoremático y sincategoremático (Términos), 174, 189
 Categorical, 94-6
 Categorías marcadas o débiles, 79
 Categorías no marcadas o fuertes, 79
 Cerezo, P., 20
 Cicerón, 29
 Ciencia, 10-3
 de la naturaleza y ciencias formales (Wittgenstein), 243
 de la naturaleza y filosofía (Wittgenstein), 244
 unidad de la (Carnap), 362
 y filosofía (Wittgenstein), 243
 Ciencias del lenguaje, 17-9
 Círculo de Viena, 245, 258, 269, 272, 360, 362, 366, 368
 Circunstancia (Ortega), 246
 Clack, R. J., 265
 Clase social
 dialectos y jergas de (Stalin), 472
 y comunidad semiótica (Bachtin, Voloshinov), 473
 Classe, A., 49
 Código, 36 s
 Cognization, 148
 Cognize, 148
 Competencia, 54-61, 124, 136, 138, 143, 148
 Competencia de comunicación, 60s
 Competencia gramatical, 60
 Competencia pragmática, 60
 Componente fonológico, 94s, 108-14, 117
 Componente semántico, 93-5, 98-108, 117
 Componente sintáctico, 93-8, 119
 Comentario, 73
 Comprensión (Wittgenstein), 288s
 Compromiso óptico (Quine), 423, 427-29
 Concepto, 41-3
 (Frege), 181
 Véase Mental
 Conductismo, 35, 61
 lógico (Wittgenstein), 288
 Connotación, 73
 (Mill), 175
 Subjetiva o fuerza expresiva, 460
 Conocimiento del lenguaje, 145-49
 Conocimiento (formas de), 145 s
 Conocimiento por familiaridad y por descripción (Russell), 190, 450
 Convencionalidad, 69 s
 Cooper, D. E., 62
 Cooperación (máximas de) (Grice), 354
 Copérnico, N., 262
 Copi, I. M., 234 s, 267
 Corominas, J., 29
 Coseriu, E., 50-3, 65 s, 475
 Creatividad del lenguaje, 61-5, 69, 124
 Criterios y síntomas (Wittgenstein), 288
- Chapel, U. C., 359
 Chomsky, N., 14, 16 s, 19, 44-6, 54 ss, 68 s, 76-8, 80, 84, 87 s, 91 s, 98, 104 s, 107, 115-21, 126-29, 133-51, 155-59, 171 s, 173, 348, 364, 416, 426
 Christensen, N. E., 264
 Church, A., 125, 178 s, 186
- Dahrendorf, R., 474
 Danto, A., 20
 Davidson, D., 16, 172, 385, 436-41, 466
 Davis, M., 125
 Deaño, A., 217
 Decir y mostrar (Wittgenstein), 241, 260-62
 Definición del lenguaje, 36-8, 67-72
 Definición léxica, 104
 Definición teórica, 104
 Demonte, V., 135
 Denotación, 70
 (Mill), 175 s
 Derwing, B. L., 59, 138, 142
 Derrida, J., 15
 Descartes, R., 173, 262, 295, 306
 Descripción de estado (Carnap), véase reglas de ámbito

- Descripción definida
 (Russell), 201, 205-10, 307
 (Strawson), 307 s, 310
 Designador
 (Carnap), 388
 rígido y no rígido o accidental (Kripke), 447-51
 Desplazamiento, 70
 D'Introno, F., 135
 Discurso
 (Modo formal del) (Carnap), 372 s, 377
 (Modo material del) (Carnap), 373
 Discurso (tipos de)
 clasificación (Morris), 333-36
 clasificación (Hierro), 337-39
 criterios de clasificación (Morris), 332-334
 descriptivo y prescriptivo (Hare), 330-331
 indicativo y directivo (Ross), 331 s
 Doble articulación, 34, 69
 Dogma del fantasma en la máquina (Ryle), 306
 Dorfman, A., 476
 Dubois, J., 22
 Ducrot, O., 22, 26
 Dummett, M., 174, 187, 265
 Duns Escoto, 45
 Durrell, L., 33

 Eco, U., 23 s, 31, 37, 47, 107
 Edwards, P., 264
 Elementos categoremáticos, 72
 Elementos deícticos, 72
 Elementos del signo, 30 s
 Elementos indéxicos, 72
 Elementos sincategoremáticos, 72
 Emisor, 31, 38, 70
 Empirismo fenomenalista (Russell), 211
 Empirismo lógico, 245
 Engelmann, P., 263
 Engels, F., 16, 470-74
 Ernout, A., 29
 Error categorial (Ryle), 306
 Escepticismo (Wittgenstein), 221, 232
 Espacio lógico (Wittgenstein), 221, 232
 Estado de cosas o situación (Wittgenstein), 227, 229, 231, 233, 257
 existente e inexistente (Wittgenstein), 230
 Estética (Wittgenstein), *véase* *Ética*
 (doctrina escolástica de la), 250
 Estructura funcional, 117 s
 Estructura lingüística latente, 158
 Estructura modal, 117 s

 Estructura profunda, 88-95, 98, 115-21, 142
 Estructura superficial, 88-95, 115-21, 142
 Estructuras-D, 122
 Estructuras-S, 122
 Ética
 y Estética (Wittgenstein), 250
 Estética y eternidad (Wittgenstein), 250
 negativa (Wittgenstein), 249
 y religión (Wittgenstein), 251
 y vida (Wittgenstein), 250
 Excusas y conceptos filosóficos (Austin), 312 s
 Expresiones tipo
 (significado intemporal de las) (Grice), 343-45
 y procedimiento resultante (Grice), 344
 Extensión e intensidad
 (Carnap), 387, 393 s, 397
 (Lewis), 442
 Euclides, 431

 Facultad lingüística, 50, 159, 169
 Falacia de la referencia, 204
 Fann, K. T., 266
 Feigl, H., 269
 Fenomenalismo (Carnap), 361
 Véase *Empirismo fenomenalista*
 Fenomenología lingüística (Austin), 312
 Feromonas, 23 s, 33
 Ferrater Mora, J., 20, 22, 217
 Ficker, L., 263
 Filosofía, 9, 14
 (Carnap), 371, 386
 como actividad (Wittgenstein), 244
 como sinsentido (unsinnig) (Wittgenstein), 244
 su método (Wittgenstein), 252, 297-99
 su método (Wisdom), 303
 su método (Austin), 312
 Wittgenstein, 288, 296 s
 y análisis lógico (Carnap), 365 s
 y modo material y formal del discurso (Carnap), 373 s
 y psicoanálisis (Wisdom), 302
 Véase *Ciencia*
 Filosofía analítica, 15
 Filosofía del lenguaje, 14-7
 Filosofía dialéctica, 16
 Filosofía especulativa, 15
 Filmore, C. J., 115
 Fioravanti, E., 474
 Fishman, J., 169-71
 Fisicalismo
 (Carnap), 362

- (Wittgenstein), 258
y Ontología (Quine), 424
Foco, 117 s
Fodor, J. A., 18, 99
Fonema, 69, 93, 109
Fonética, 109
Fonología, 74, 109
Forma
 canónica (Quine), 417
 lógica, 120-22
 de la proposición existencial (Russell), 202
 (Wittgenstein), 221 s, 241, 259
Foster, J. A., 439 s
Fouts, R., 162
Fraassen, B. C. van, 210
Frankfurt (Escuela de), 16
Fraser, C., 139
Frege, G., 15, 103, 171, 173 s, 177-189, 192, 194, 205, 207, 209, 213, 223 s, 255, 263 s, 310, 330, 360, 386, 394, 416, 436, 440, 447, 460
Frisch, K. v., 24, 35
Función cognoscitiva, 157
Función designativa, 36
Función intensional (Russell), 198
Función rememorativa, 41
Función significativa, 36 s, 40
- Gadamer, H. G., 15, 27 s
Galileo, 262
Galmiche, M., 116, 135
Gallego, A., 23
García Suárez, A., 246, 357
Gardin, B., 471
Gardner, R. A., 160-63, 167
Gazdar, G., 356, 358
Geach, P., 46, 178
Gellner, E., 20
Gestalt (psicología de la), 362
Gill, T., 166, 168
Giro lingüístico, 173
Glaserfeld, E. v., 165
Goodman, N., 139, 141, 172
Gracia, F., 22, 47, 135, 385
Gramática categorial (Lewis), 441
Gramática de estados finitos, 80, 82, 85
Gramática de estructura sintagmática, 82 s, 85, 87
Gramática filosófica
 (Russell), 189
 (Wittgenstein), 296
Gramática generativa, 55, 79-87
Gramática libre de contexto, 85
Gramática ordinaria
 y gramática lógica (Wittgenstein), 254
 y gramática profunda (Wittgenstein), 296
Gramática sensible al contexto, 85 s
Gramática transformacional, 87-131, y *passim*
Gramática universal, 65, 107, 122-31, 137 s, 141, 145
Greenberg, J. H., 74-6, 79
Greene, J., 58 s, 142
Grice, H. P., 16, 172, 341, 358, 430, 460, 462
Griffin, J., 234, 267
Gross, M., 125
- Haack, S., 385
Haber y existir (Ortega), 203 s
Habermas, J., 173
Habla, 48-54, 65
Hacker, P. M. S., 266, 359
Hadlich, R. L., 92, 108-12, 135
Halle, M., 111
Hare, R. M., 320, 330-32
Hartnack, J., 266
Harris, Z., 17
Hayes, K. J., 159
Hecho
 o acontecimiento (Wittgenstein), 229 s, 233
 de actitud proposicional
 (Russell), 200
 (Wittgenstein), 239 s
 atómico (Russell), 192 s
 general y particular
 (Russell), 199 s
 (Wittgenstein), 237 s
 mental (Russell), 199
 negativo
 (Russell), 197 s
 (Wittgenstein), 237-39, 283
Heidegger, M., 9, 15, 30, 364, 367
Hempel, C. G., 368-70, 402
Herriot, P., 21, 56, 140, 142, 145, 151, 170, 172
Hintikka, J., 128, 130 s, 234
Hjelmlev, L., 17
Hochstetter, E., 42 s
Hockett, Ch. F., 64, 68, 70-6, 128-30, 171
Holdcroft, D., 329
Holismo
 (Quine), 402
 concepción holista del significado (Davidson), 436
Houdebine, J. L., 470-72
Humboldt, W. v., 21, 54
Husserl, E., 10-2, 15

- Icono, 24, 28, 41 s, 48, 70
 Idea, 10
 Ideas innatas, 136-56
 Ideología
 y lenguaje, 53 s
 y signo (Bachtin, Voloshinov), 472 s
 (Marx, Engels), 474
 Iglesias, M. T., 221, 254
 Illocucionario
 (Fuerza) implícita y explícita, 459
 (efecto) y efecto perlocucionario (Grice, Searle), 347
 Véase Acto
 Imagen, 34-6, 41 s
 Implicación
 analítica, 351, 353, 356
 contextual (Nowell-Smith), 351 s, 356-357, 460
 lógica, 350, 356
 lógica y conceptos-L (Carnap), 391
 material, 350, 356
 pragmática (Nowell-Smith), 352, 354 s
 Indicador gramatical, 100
 Indicador semántico, 99-106
 Indicador sintagmático inicial, 120
 Véase también Marcador sintagmático
 Indicadores semánticos (Katz), 429
 Índice, 24, 27 s, 41 s
 (Lewis), 442, 463
 Indicio, 24
 Instrumento de adquisición del lenguaje, 136-43
 Intención del hablante y expresiones declarativas e imperativas (Grice), 342
 Intensión, *véase* Extensión (Carnap).
 Intensionalidad isomórfica (Carnap), 395
 Interpretación radical (Teoría de la) (Davidson), 440
 Intuiciones del hablante nativo (Austin), 313
 Ishiguro, H., 236

 Jackendoff, R. S., 116-18
 Jakobson, R., 17, 77, 111 s, 328 s
 Janik, A., 266
 Juego
 de lenguaje (Wittgenstein), 277-79
 (Wittgenstein), 274
 y filosofía (Wittgenstein), 299 s
 Véase lenguaje privado, Uso

 Kant, I., 173, 202, 215, 263, 403, 429
 Kaplan, D., 445
 Kasher, A., 356
 Katz, J., 16, 18-20, 98-108, 117, 121, 135, 396, 416, 429 s, 436
 Katzner K., 37
 Kay, P., 746
 Keenan, E. L., 78, 355
 Kellogg, W. N., 159
 Kenny, A., 217, 219, 238, 265, 272, 358
 Körner, S., 20
 Kraft, V., 456
 Kripke, S., 172, 177, 188, 446-455, 463
 Kuhn, T. S., 115
 Kutschera, F. v., 21
 Kuznetsov, I. V., 10 s, 20

 Lakoff, G., 115 s, 126
 Lambert, K., 76, 210
 Lázaro Carreter, F., 22
 Lazerowitz, M., 270-71
 Lenguaje de clase, 53
 Lectura derivada, 99, 104
 Lectura léxica, 99, 104, 106 s
 Lectura metafórica, 104
 Leech, G., 76, 463
 Legisigno, 31
 Leibniz, G. W., 388, 390, 433
 Lemmon, E. J., 311
 Lengua, 48-54, 64 s, 122
 Lenguaje, 36-9, 48-52, 65, 67-72 y *passim*
 como instrumento (Wittgenstein), 279
 convencional, 41, 44, 46
 creatividad semántica del, 256
 empírico (Hempel), 369
 escrito, 34, 37, 44-8
 externo, 45
 formalizado, 40
 formalmente especificable (Tarski), 380
 funciones del, 327-29, 340
 hermenéutica del (Castilla del Pino), 461
 ideal, 46
 mental, 44-6, 49
 natural, 40, 44, 46
 oral, 34, 44-6, 50
 origen del (Marr), 471
 pictográfico, 37, 49
 semánticamente cerrado (Tarski), 380
 silbado, 49
 teoría formal del (Carnap), 377
 universal, 46
 usos del (Austin), 313
 verbal, 36-40, 48-50, 65, 67 s
 (Wittgenstein), 273
 y conocimiento, 149-156
 Lenguaje ordinario
 Ambigüedad del (Wittgenstein), 254 s
 (Austin), 312
 (Russell-Wittgenstein), 253
 (Wittgenstein), 269

- regimentación lógica del (Quine), 417-421
 y cálculo lógico (Ryle), 303-05
 y lógica (Strawson), 308, 310
 y sentido común (Moore), 270 s
 y su estructura lógica subyacente (Wittgenstein), 254
 y verdad (Tarski-Davidson), 384 s
- Lenguaje privado
 (Russell), 191
 (Wittgenstein), 246, 295
 y fenómenos mentales (Wittgenstein), 286 s, 290
 y juego lingüístico (Wittgenstein), 292
 y referencia (Wittgenstein), 288
 y uso (Wittgenstein), 292
- Lenneberg, E. H., 21, 56, 134, 156-59, 171 s
- Lentin, A., 125
- Lepschy, G. C., 111
- Lexema, 98 s
- Léxico, 94, 96, 98
- Lewis, D., 16, 126, 172, 433, 441, 460, 462 s, 463 s, 468
- Lewy, C., 246
- Lieberman, P., 157
- Lingüística, 17-9
- Lipps, H., 15
- List, G., 151
- Locucionario, *véase* Acto
- Lógica del discurso común y lógica formal (Strawson), 307
- Lógica informal (Ryle), 303-05, 307
- Lógica libre, 210
- Logicismo (Russell), 244
- Lyons, J., 21, 66, 82, 87, 126, 135
- Lledó, E., 20
- MacCorquodale, K., 150, 172
- Mach, E., 361
- Malcolm, N., 270 s
- Malinowski, B., 328
- Malmberg, B., 47
- Mallarmé, S., 203
- Máquina de Turing, 122-26, 128
- Marcador sintagmático, 86-93
- Marcellesi, J. M., 471
- Marías, J., 248, 305, 461
- Marler, P., 35, 140
- Marr, N., 16, 470, 472, 474
- Martinet, A., 34, 69, 98
- Marx, K., 16
- Marxismo, 16 s
- Matriz
 biológica, 156
 fonológica, 96-8, 111, 114
- Mattelart, A., 476
- McCawley, J. D., 115 s, 143 s
- McGuinness, B., 214
- McNeill, D., 139, 156
- Meillet, A., 29
- Meinong, A., 103, 203, 210
- Mental
 conceptos (Ryle), 306
 hecho (Russell), 199
Véase Lenguaje privado
- Mentalismo, 37, 51, 55 s
- Merleau-Ponty, M., 15
- Metafilosofía, 14
- Metafísica, 12
- Metalinguaje
 (Wittgenstein), 259 s
 (Russell), 260
- Mill, J. S., 171, 174-77, 179 s, 187, 194, 447, 451-53, 463
- Miller, D., 258
- Miller, G. A., 56, 134
- Modo convencional de significar, 41, 44
- Modo natural de significar, 42
- Moliner, M., 29, 32
- Monema, 69, 98
- Montague, R., 16, 441, 446
- Moore, G. E., 15, 213, 247, 264, 269-71
- Morfema, 69, 98
- Morfología, 74
 derivacional, 108
 inflexional, 108
- Morfofonémica, 108
- Moro Simpson, T., 179, 264, 456
- Morris, Ch., 25 s, 28, 35, 38-40, 47, 332-334, 466
- Morton, J., 140
- Moulines, U., 179, 362, 431
- Mounin, G., 22
- Muguerza, J., 265, 359, 456
- Mundo (Wittgenstein), 229 s, 232
- Mundo posible
 (Wittgenstein), 219, 441
 (Kripke), 447
 y significado (Foster), 439 s
 y verdad analítica (Leibniz, Lewis), 433
- Mundo y vida (Wittgenstein), 246, 249, 260
- Necesidad
 (Wittgenstein), 242
 (Kripke), 453
- Neopositivismo, 245, 257, 259-60, 262, 272, 286, 301, 330, 360, 371, 407
- Neurath O., 362 s
- Newton, I., 431

Nombre

- (Mill), 174
- común o de clase (Kripke), 461
- de experiencias internas (Wittgenstein), 291-93
- propio (Mill), 179, 194, 447
 - su sentido y referencia (Frege), 179 s, 447
 - en sentido lógico (Russell), 195 s, 208, 308, 447
 - ordinario y descripción definida (Russell), 201
 - Wittgenstein, 223 ss, 258, 276, 281
 - Kripke, véase Designador
- tipos de (Mill), 175-77
- Norma, 52-4, 65
- Nowell-Smith, P. H., 351 s, 356, 460

Objeto

- constituyentes del (Wittgenstein), 281
- Frege, 180, 182
- Inexistente (Russell), 203
- Russell, 191; véase Particulares
- Wittgenstein, 231-36, 258
- Occam, G., 31, 34, 40-7, 49
 - navaja de, 212
- Ogden, C. K., 214, 462 s
- Ontología, 12
 - y semántica (Carnap), 400 s
 - Véase Compromiso óntico
- Opacidad referencial (Quine), 426
- Oposiciones funcionales, 52
- Oración
 - eterna (Quine), 422
 - no declarativa
 - su sentido y referencia (Frege), 185
 - Lewis, 443 s
 - observacional (Quine), 407, 412
 - ocasional (Quine), 405 s
 - permanente (Quine), 405 s
 - subordinada en estilo indirecto: su sentido y referencia (Frege), 184
 - su sentido y referencia (Frege), 182
 - uso y preferencia o emisión (sentence, use and utterance) (Strawson), 308
- Ortega y Gasset, J., 10, 30, 203, 248, 305, 461
- Ostensión (Russell), 208
- Oxford (Escuela de), 302 s, 307, 311, 330

Paradoja

- de la identidad, 177, 185
- del mentiroso (Tarski), 379 s
- Parecidos de familia (Wittgenstein), 274, 303
- Parkin, F., 474

- Parkinson, G., 359
- Partee, B., 102
- Particulares (Russell), 193, 211
- Patterson, F., 167
- Peano, G., 431
- Pears, D., 214, 266
- Peirce, C. S., 24-8, 34, 41 s, 70
- Pensamiento
 - constitutivos del (Wittgenstein), 223
 - Frege, 183
 - límites del (Wittgenstein), 263
 - Wittgenstein, 221 s, 276
 - y lenguaje (Wittgenstein), 222, 254
- Perlocucionario, véase Acto
- Peters, S., 126-29, 141 s
- Piaget, J., 10, 20, 26, 28, 151-55, 158
- Pinborg, J., 46
- Pinillos, J. L., 30, 161, 171
- Pitcher, G., 359
- Platón, 10
- Platts, M. de B., 347-49, 440
- Ponzio, A., 17, 51 s, 54, 475
- Popper, R., 367-69, 385
- Postal, P., 117
- Postulados de significado, 102
- Poulantzas, N., 474
- Pragmática, 39 s
- Predicado: su sentido y referencia (Frege), 181
- Premack, D., 163-65
- Presuposición, 117 s
 - Strawson, 307 s, 351, 356
- Principio
 - de bivalencia y tercio excluido (Tarski), 384
 - de caridad (Davidson), 437
 - de cooperación (Grice), 354
 - de extensionalidad, 196, 212, 228
 - de falsabilidad (Popper), 368
 - de familiaridad (Russell), 193
 - de indeterminación de la traducción radical (Quine), 413
 - de isomorfía, 189, 356
 - (Wittgenstein), 220, 223, 235
 - de Leibniz (sustituibilidad), 183
 - de verificabilidad (Wittgenstein), 258
 - Ayer, Hempel, 367-69
- Procedimiento (Grice)
 - y competencia (Chomsky), 348
- Proposición
 - atómica y molecular (Russell), 192, 196
 - cuantificada: general y particular (Russell), 199 s
 - de actitudes proposicionales (Russell), 198
 - Mill, 174

- negativa (Wittgenstein), 228
 negativa y hecho negativo (Russell), 197 s
 poder lógico de la (Ryle), 303, 305
 protocolar (Carnap), 362
 sentido de la (Wittgenstein), 223, 225
 simple y compleja (Wittgenstein), 229
 sobre objetos inexistentes (Russell), 204 (Lewis), 443
 tipos de (Wittgenstein), véase uso lingüístico
 Wittgenstein, 225 s, 276, 280
 y discurso declarativo o asertórico (Russell), 192
 y hecho (Russell), 191
 y oración (Strawson), 309
 Pseudoproposiciones (Wittgenstein), 240
 lógicas, 241
 filosóficas, 243
 sobre solipsismo, 245 s (Carnap), 363
 metafísicas, 364 s
 de objeto, 372 s
 Psicolingüística, 57, 65
 Putnam, H., 126, 139, 141, 172, 431
- Quesada, D., 21, 80, 87, 125, 134, 149, 446
 Quine, W. v. O., 16, 37, 46, 63, 103 s, 141, 158, 172, 241, 350, 385, 396, 398, 401-32, 436-39, 456 s, 461, 465 s
- Ramsey, F. P., 200, 269
 Rasgos definitorios del lenguaje, 67-72
 Rasgos fonéticos distintivos, 109-114
 realidad (Wittgenstein), 232 s
 realizativo
 preferencias y expresiones (Austin), 313-15
 y constatativo (Austin), 314
 Véase Acto
 receptor, 31, 35, 38
 reconstrucción racional (Carnap), 399
 recursividad, 122-26
 reducción husserliana, 248
 referencia
 contextual, 459
 Frege, 178
 ontogénesis de la (Quine), 426
 Strawson, 310
 reflexividad, 705
 regla
 de ámbito (Carnap), 388 s
 de selección (Chomsky), 364
 de traducción o ley de proyección (Wittgenstein), 226
 de un lenguaje-L (Carnap), 387
 de verdad, véase Verdad
- reglas
 de inserción léxica, 96-8, 115, 121
 de interpretación, 120-22
 de proyección, 99, 104, 106 s
 de ramificación, 94-6
 de redundancia, 106
 de segmentación, 109, 114
 de subcategorización, 94-6
 de transformación, 78, 88-93, 109
 fonológicas, 109
 relación (Russell), 193
 relatividad lingüística, 168-171
 representación (Bild)
 forma de (Form der Abbildung) (Wittgenstein), 219
 forma propia de cada (Form der Darstellung) (Wittgenstein), 220
 relación de (Wittgenstein), 219
 relación interna (Wittgenstein), 326
 Wittgenstein, 216, 218
 sentido de la (Wittgenstein), 221
 representación enactiva, 155
 representación icónica, 155
 representación simbólica, 157
 restricciones de selección, 104-06
 Ricoeur, P., 15
 Richards, I. A., 462 s
 Richelle, M., 150, 153-56, 172
 Richtie, R., 126-29, 141 s
 Rivero, M. L., 135
 Rivière, A., 58
 Roberts, J. M., 171
 Rodríguez Adrados, F., 21
 Ronat, M., 107
 Rorty, R., 301
 Ross, A., 331, 334
 Ross, J. R., 115
 Rossi-Landi, F., 16, 475
 Rumbaugh, D., 165 s, 168
 Russell, B., 13, 15, 21, 46, 121, 171, 177, 179, 185 s, 188-215, 221, 223-40, 244, 246, 253-56, 257-60, 262-66, 268 s, 273, 281 s, 288, 307-09, 311, 330, 350, 360 s, 367, 407, 416-21, 443, 447, 450, 465, 467
 Ruwet, N., 21
 Ryle, G., 15, 146, 294, 303 s, 311 s
- Sacristán, M., 20, 218
 San Agustín, 41, 277, 282, 298
 Sampson, G., 59, 62, 64, 162, 165
 Sánchez de Zavala, V., 61, 135, 172
 Sapir, E., 17, 21, 169

- Saussure, F. de, 14, 17, 26-8, 31, 37, 50-4, 66, 69, 152, 462, 470, 475
- Schaff, A., 16s, 25 s, 28, 31 s, 35, 37, 172, 324
- Schlick, M., 269, 360
- Searle, J. R., 172, 313, 320, 329, 339-41, 346 s, 349, 358 s
- Sema, 102
- Semántica generativa, 115 s
- Semántica interpretativa, 118
- Semanticidad, 70
- Semema, 102
- Semiótica, 39 s
- teoría conductista de la (Morris), 332
- Sentido, 99, 102
- Frege, 178
- matización o coloración del (Frege), 184, 460
- Strawson, 310
- Seña, 29
- Señal, 23-9, 31 s, 35
- significación, 26, 31, 34, 39, 43, 45 s
- Significado, 24, 26 s, 31, 33, 37 s, 40-2, 44, 94, 98 s
- concepción holista del (Davidson), 436
- cognitivo y significado
- expresivo o emotivo (Carnap), 366
- cognitivo de una expresión (Hempel) 369 s
- criterios de (Carnap), 363 s
- (criterio de) y criterio de demarcación (Carnap, Popper, Ayer), 367-371
- estimulativo (Quine), 405
- como función
- Lewis, 442
- gramatical, 458
- como función, 468
- y mundos posibles, *véase* Mundo posible
- natural y no natural (Grice), 341 s
- postulado de (Carnap), 396
- teoría pragmática del (Grice), 341, 343
- teoría referencialista del, 174
- Russell, 191
- Wittgenstein, 223, 281
- teoría social y pragmática del (Quine), 405
- teoría veritativa del (Davidson), 440
- y verificabilidad del (Carnap), 365
- Significante, 24, 28 s, 30 s, 36-8
- Significatividad
- condición general de y convención T (Davidson), 437
- y tipos de discurso (Davidson), 438
- Signo
- acontecimiento, 31, 51, 65
- artificial, 25 s, 28
- cultural, 33-6
- escrito, 41 s, 44
- hablado, 41 s, 44
- humano, 33, 35
- icónico, 24 s, 35
- indéxico, 24, 35
- lingüístico, 26, 28, 31, 37, 41-4, 51
- lógico, 32, 35
- matemático, 32, 35
- mental, 31, 41-4
- natural, 25, 28, 33-6, 40, 42
- plurisituacional, 38
- prelingüístico, 139
- proposicional (Wittgenstein), 222 s
- rememorativo, 41 s, 44
- representativo, 41
- simbólico, 25
- sustitutivo, 26, 28, 31 s
- tipo, 31, 51, 65
- y signo acontecimiento, 239, 309, 344, 461
- verbal, 26, 28, 34-8
- y símbolo (Wittgenstein), 254
- Símbolo, 25-30, 32-8
- Sincategorema, 72
- Sinonimia
- de oraciones y términos (Quine), 407
- estimulativa y social (Quine), 410
- estructural (Quine), 421-23
- y conducta lingüística (Quine), 405
- y expresiones coextensivas (Quine), 410 s
- y significado estimulativo (Quine), 406
- y verdad analítica (Quine), 403
- véase* Intensionalidad isomórfica (Carnap)
- Sinsentido (unsinnig) y carente de sentido (sinnlos), (Wittgenstein), 242, 244, 252
- Sinsigno, 31
- Sintaxis, 39 s, 74
- y semántica (Carnap), 378, 386
- Síntoma, 25, 28
- Sistema, 52-6, 58, 65
- Sistema de reescritura sin restricciones, 80 s, 85
- Skinner, B. F., 57-9, 61, 150 s, 172
- Sociología del lenguaje, 65
- Solipsismo
- Carnap, 361, 367
- Wittgenstein, 245, 262, 291, 295
- y lenguaje privado (Wittgenstein), 290
- y realismo (Wittgenstein), 248
- Specht, E. K., 278, 359
- Stalin, J., 16, 471 s, 475
- Stalnaker, R., 445
- Stegmüller, W., 218
- Strawson, P. F., 15, 185, 209, 307, 311, 315, 320, 351, 416, 430, 443

- Subcomponente
 - de base, 88, 94-8, 117 s, 120
 - transformacional, 88, 94 s, 117 s, 120
- Subsistema
 - fonológico, 69
 - gramatical, 69
 - semántico, 69
 - sintáctico, 69
- Sujeto (Wittgenstein), 240, 247 s
- Tarski, A., 172, 378-82, 384 s, 390, 417, 437, 439, 445, 456, 468
- Tato J. L., 105
- Tautología
 - y contradicción (Wittgenstein), 241 s, 249
 - Lewis, 442
- Tautonomía (Bunge), 432
- Teodicea negativa (Wittgenstein), 251
- Teoría, 9-14
 - de las huellas, 119
- Terrace, H., 168
- Tesis de Church, 125
- Tesis sobre *any*, 130
- Thiel, Ch., 186, 265
- Thorpe, W. H., 71
- Tierno Galván, E., 214
- Todorov, T., 22, 26, 28
- Toulmin, S. E., 266
- Traducción de
 - «Sinn» y «Bedeutung» (Frege), 178 s
 - «Bild» (Wittgenstein), 216-18
 - «Satz» (Wittgenstein), 222
 - «Bedeuten» y «Bedeutung» (Wittgenstein), 223 s
 - «Sachverhalt» (Wittgenstein), 227
 - «Sachlage» (Wittgenstein), 227
 - «der Sprache, die allein ich verstehe» (Wittgenstein), 246
 - «P entails q», «Entailment», 350
 - «regimentación» (Quine), 417
 - «Utterance», 461
- Traducción radical
 - Quine, 406, 415
 - indeterminación de la (Quine), 411
 - y mentalismo (Quine), 415 s
- Trascendental
 - lógica (Wittgenstein), 243, 248
 - sujeto (Wittgenstein), 248
 - ética (Wittgenstein), 250
 - estética (Wittgenstein), 251
- Trascendentalismo lingüístico (Wittgenstein), 295 s
 - destrascendentalización, proceso de (Wittgenstein), 301
- Trentman, J., 46
- Trubetzkoy, N. S., 17
- Trujillo, R., 49
- Turing, A. M., 125
- Ullman, S., 462 s
- Universal condicional, 74, 76, 79
 - de institución voluntaria, 42
 - fonológico, 72-4
 - formal, 77 s
 - lingüístico, 46, 71-9, 138, 140 s
 - natural, 42
 - semántico, 72, 76
 - sintáctico, 74, 126
 - sustantivo, 77, 126
- Urmson, J. O., 200, 211, 236, 264
- Uso
 - y significado (Wittgenstein), 276, 311
 - y nombre propio (Wittgenstein), 277, 282
 - lingüístico, tipos de (Wittgenstein), 283 s
 - lingüístico y actos de habla (Austin), 315
 - lingüístico y actos ilocucionarios (Wittgenstein, Austin, Searle), 341
 - y usanza (use and usage) (Ryle), 304 s, 311
 - y verdad (Strawson), 308
- Valor cognoscitivo de una expresión, 178, 186
- Valor veritativo, 181, 183, 224
 - y expresiones realizativas (Austin), 318
- Valls, A., 27
- Verbos proposicionales (Russell), 198
- Veracidad (convención de) (Lewis), 445
- Verdad
 - concepto de (Tarski), 378
 - convención T (Tarski), 379
 - condiciones de adecuación material y corrección formal (Tarski), 379
 - definición aristotélica de, 379, 385 (definición de) y satisfacción (Tarski), 380 s
 - y lenguaje ordinario (Tarski, Davidson), 384 s
 - reglas de (Carnap), 387
 - verdad necesaria o analítica y verdad empírica, 350 s
 - Carnap, 365, 390-92, 395, 398
 - Frege, 178, 186
 - Katz, 429
 - Kripke, 450-52, 455
 - Leibniz, 391
 - Moulines, 431
 - Quine, 402, 409-11
- Véase Tautología

- Verificación y confirmación (Carnap), 369
 Verón, E., 476
 Vestigio, 34-6, 41 s
 Vida
 forma de (Wittgenstein), 286, 292, 299
 sentido de la (Wittgenstein), 252, 263
 Vocabulario auxiliar, 83
 Vocabulario terminal, 83 s
 Voloshinov, V. N., 16, 472 s, 475
 Voluntad (Wittgenstein), 251
 Vygotsky, L. S., 154
 Waismann, F., 11, 269, 360
 Wales, R., 59 s, 140
 Wasow, T., 128, 131, 142
 Weinberg, J. R., 456
 Weinreich, H., 116
 Whitehead, A. N., 255, 350, 416
 Whitman, W., 335
 Whorf, B. L., 21, 169-71
 Winch, P., 266
 Wisdom, J., 211, 302 s
 Wittgenstein, L., 9, 12, 14-6, 21, 27, 46, 64, 171-3, 177, 186, 188, 199, 212-305, 309, 311-13, 316 s, 324, 329 s, 338, 341, 358, 360, 365 s, 372, 377, 390, 395, 416 s, 420, 428, 439, 441, 458
 Wright, G. H. von, 269
 Yerkes, R. M., 167