

SECCION VII.

La Interpolación Espacial y los Modelos Digitales de Terreno (MDT).

INTRODUCCIÓN

Representar matemáticamente fenómenos de la realidad no es fácil. Más aún si estos se componen de múltiples variables que interactúan entre sí como ocurre en el medio ambiente. La topografía por ejemplo, constituye una parte del paisaje. En él se insertan los conceptos de altitud, relieve, exposición e iluminación. A su vez, el paisaje está inserto en el clima, el cual está regulado por regímenes de temperaturas, humedad, presión y vientos. Pues bien, todos estos elementos constituyen parte de lo que llamamos variables continuas, es decir, que la acción continua sobre un espacio determinado forman una superficie, que en la mayoría de los casos es posible cuantificar a través de modelos predictivos. En el ámbito forestal por ejemplo, la distribución de diámetros cuadráticos medios, el área basal por hectárea en un conjunto de puntos de muestreo, los valores de índice de sitio y variables genéticas como la heredabilidad y patrones de crecimiento también son posibles de cuantificar y ser representados espacialmente como una superficie.

Los modelos de terreno son un ejemplo de cómo puede representarse la realidad mediante la representación espacial de una variable cuantitativa y continua. Y es por esto que los MDT presentan claras ventajas respecto al análisis tradicional de mapas o modelos análogos de terreno, tales como: la posibilidad de tratamiento numérico de los datos, y la posibilidad de realizar simulación de procesos, emulando el funcionamiento de un sistema dinámico real. Por tal razón un algoritmo sólidamente construido aplicado sobre datos fiables genera resultados aplicables al objeto real con un error moderado.

Por lo anterior, podríamos definir entonces a un modelo digital de terreno como una estructura numérica de datos que representa la distribución espacial de la altitud de la superficie del terreno. A su vez un terreno real puede describirse de forma genérica como una función bivariable continua del tipo $z = f(x, y)$, donde z representa la altitud del terreno en el punto de coordenadas (x, y) y z es una función que relaciona el valor de la variable con su localización geográfica.

La representación de una superficie, como veremos más adelante, puede ser construida mediante técnicas de interpolación. Sea cual sea el método, la génesis del MDT es un punto acotado: es el dato geográfico esencial que contiene el valor (su atributo) y sus coordenadas espaciales. El conjunto de puntos que cumplen con estas características permiten formar una estructura elemental de datos según dos modelos básicos: el modelo vectorial, conformado por contornos o polilíneas de altitud constante, puntos acotados, y red de triángulos irregulares adosados (que luego denominaremos TIN). El segundo modelo corresponde al raster, identificado por matrices regulares de datos, que almacenan en su interior las cotas o valores de altitud – o el valor de otra variable de interés si se trata de analizar un raster -, y el modelo jerárquico o quadtree conformado por matrices jerárquicas imbricadas.

En el caso de un modelo vectorial que representa las curvas de nivel, éstas están formadas por polilíneas, es decir, un vector de n pares de coordenadas (x, y) describe la trayectoria de cada curva, siendo el número de elementos de cada vector, variable en función de la modalidad y detalle de ingreso de los datos a la cobertura vectorial. Si estamos hablando de un archivo de curvas de nivel, entonces nuestro MDT está conformado por n elementos o cadenas, cada una de ellas conteniendo m pares de coordenadas o vértices (x, y) que describen la trayectoria de ese vector. Repasar la topología arco-nodo de Arc/Info, constituye un buen ejercicio para entender claramente la organización de los datos al interior de un fichero vectorial.

Los siguientes pasos a señalar, estarán enfocados a la construcción de un modelo digital de elevaciones, siendo perfectamente factible construir una variada gama de aplicaciones (ambientales por ejemplo), mediante las técnicas de interpolación. Revisaremos los antecedentes necesarios para la construcción de un MDT.

UN PASO PREVIO: LA CAPTURA DE DATOS

Para la construcción de un MDT, lógicamente se necesita una fuente de información. Esta puede provenir de una captura de datos en terreno, de información ya existente, o una situación mixta en donde muchas veces se requiere efectuar un re-muestreo de valores de altitud para la zona de interés.

Para el caso de la captura de datos que representan la altitud del terreno, existen dos métodos: el directo, en donde se efectúa una medida directa de los datos de altitud sobre el terreno (fuentes primarias), empleando los principios de la altimetría (altímetros transportados por plataformas aéreas), mediante GPS (Global Positioning System), o estaciones topográficas con grabación de datos.

Existe también formas indirectas tales como medir esta variable a partir de documentos previamente elaborados (fuentes secundarias), una fuente digital (como las imágenes satelitales), cámaras métricas, el escaneo, y la digitalización propiamente tal.

En general lo más evidente es transferir todas las curvas de nivel del mapa topográfico. Pero hay más cuestiones a tener en cuenta. La calidad de un MDT por ejemplo, puede mejorarse significativamente complementando las curvas de nivel con datos auxiliares de diversos tipos. En general, el listado de elementos que pueden aportar información para construir un MDT son:

- ❑ las curvas de nivel, descompuestas o no en puntos acotados -mass points- previa generalización o reducción de la densidad de vértices de la línea,
- ❑ los puntos acotados singulares o vips -de very important points- tales como cumbres de cerros y fondos de quebradas, las líneas estructurales, que definen elementos lineales con valores de altitud asociados a cada vértice,
- ❑ las líneas de inflexión o quiebre -breaklines-, que definen la posición de elementos lineales sin valores de altitud explícitos que rompen la continuidad de la superficie,
- ❑ las zonas de altitud constante, como polígonos que encierran una superficie de altitud única,
- ❑ las zonas de recorte, que definen los límites externos del MDT, y,
- ❑ las zonas vacías, en donde no es deseable asignar altitudes, sobre todo si no se conoce la situación real de estas zonas.

Todas estas observaciones son necesarias tomar en cuenta si se desea generar un MDT que permita aprovechar adecuada y eficazmente la información topográfica disponible.

EL MUESTREO Y DATOS NECESARIOS PARA LA INTERPOLACIÓN

Para la creación de superficies continuas, es necesario partir con una fuente de información (los datos); éstos pueden provenir de fotos aéreas estereoscópicas, utilizadas en fotogrametría, imágenes satélite, documentos escaneados; puntos de control medidos directa o indirectamente, transectos regulares o curvas de nivel, o polígonos digitalizados.

La localización de los datos en el espacio, siempre es un aspecto muy importante a la hora de tomarlos para generar una interpolación. Idealmente, los datos deben siempre cubrir el área de interés. Una malla regular de datos es siempre recomendable cuando es posible obtenerla, más aún si en ella es posible obtener información sobre puntos críticos del relieve tales como fondos de valle, líneas de altas cumbres y otros puntos de cambio importantes en la topografía.

En el caso de los datos puntuales, primeramente los puntos debiesen estar regularmente separados de modo de conformar una grilla regular de datos, cada uno de ellos identificados con su localización espacial y relativa respecto a los otros puntos. Este proceso de captura y ordenación de información se facilita enormemente si se dispone de un dispositivo GPS de captura de datos en terreno. Es necesario además, una aleatorización, de modo de no sesgar la distribución de ellos a zonas en las cuales la disponibilidad y acceso es mayor que en otros puntos de la zona de interés.

Este proceso de aleatorización normalmente presenta diferentes clases de agrupamiento, todas ellas válidas de acuerdo a la realidad en terreno. Las principales formas de agrupamiento de los datos para información puntual es la siguiente: Una malla regular de datos, un muestreo aleatorio simple, un muestreo por transectos, un muestreo aleatorio estratificado, un muestreo agrupado (con efecto cluster), y, un muestreo por contornos.

El muestreo cluster por ejemplo, es utilizado para examinar la variación espacial de los datos a diferentes escalas, mientras que los transectos regulares son frecuentemente utilizados para medir perfiles o contornos de ríos, líneas de costa, formas de laderas y contornos de líneas que comúnmente se emplean para la elaboración de cartografía topográfica y para la elaboración de modelos digitales de elevación.

En estudios demográficos, la información base es frecuentemente trabajada en forma de áreas, determinadas a través de censos o codificación de superficies. Se presenta por tanto, una variedad en tamaño y forma de estas superficies, lo que trae consigo dificultades al proceso de interpolación. Por otra parte, se presenta también el problema que los datos colectados no son correctamente medidos o atributados en su fuente original de captura. Es el caso por ejemplo, de los modelos digitales de elevación en donde la variable *z* de altitud es mal ingresada, o simplemente no es incorporada a los datos; lo mismo puede ocurrir en otros casos como la presión barométrica, indicadores socioeconómicos u otro tipo de variables expresadas en escala nominal (valor bruto) u ordinal (datos jerarquizados).

Las estadísticas de las diferencias (absolutas y cuadráticas) entre ambas técnicas de predicción representadas como un estimador del tipo $\bar{z}(x_i) - z(x_i)$ es frecuentemente usada como un indicador de calidad y exactitud de los métodos de interpolación.

En el caso de la interpolación a partir de puntos, la predicción de valores para la interpolación de variables climáticas representa el peor de los casos ya que, normalmente, la localización de los puntos de muestreo (estaciones meteorológicas) se decidió hace tiempo y resulta claramente imposible aumentar su número para mejorar la interpolación (aunque podría aumentarse el número para obtener mejores mapas en el futuro). En el caso de un muestreo de campo de variables del suelo, podemos decidir donde medir y, en función de los resultados, volver a muestrear en otros puntos, aumentando así, el tamaño muestral. Suponiendo que tenemos la posibilidad (y la responsabilidad) de hacer nuestro propio diseño de muestreo, para ello nos debemos basar en el conocimiento previo que tengamos acerca de la estructura de variación de la variable a interpolar. Los modelos básicos que podremos utilizar (o incluso combinar) son los siguientes:

- ☐ Muestreo regular
- ☐ Muestreo aleatorio
- ☐ Muestreo estratificado
- ☐ Muestreo por agregados

Los dos primeros son los más adecuados cuando no conocemos nada acerca de la estructura de variación. El muestreo regular puede dar problemas si la variable presenta un comportamiento rítmico, el aleatorio por su parte puede dejar áreas extensas sin muestras. Una solución de compromiso sería un muestreo aleatorio estratificado en el que el espacio a muestrear se divide en bloques que serán muestreados con un punto cuya ubicación dentro del bloque es aleatoria.

El muestreo estratificado es útil cuando tenemos una variable de apoyo, fácil de medir u observar, que sabemos que influye sobre la variable a interpolar, por ejemplo el tipo de suelo o la topografía van a condicionar el contenido en sales. Un muestreo estratificado dividiría el área de estudio en función de estas variables de apoyo para muestrear todos los posibles valores que aparezcan. Si la variable de apoyo es cualitativa la división se hace mediante polígonos y si es cuantitativa mediante un muestreo por transectos o isolíneas. Finalmente el muestreo por agregados se utiliza cuando el objetivo del muestreo no es tanto realizar un mapa como conocer la estructura de variabilidad de la variable ya que permite analizar la misma a diferentes escalas.

Es necesario señalar que para seleccionar el método de interpolación y las variables de apoyo, hay que tener en cuenta además la escala del trabajo ya que los factores que explican la distribución espacial de una misma variable pueden cambiar con la escala. Por ejemplo, a media escala las propiedades del suelo variarán fundamentalmente en función del material parental; pero a escala de detalle lo harán en relación con la topografía.

ELECCIÓN DE LA ESTRUCTURA DE DATOS

La adopción de una estructura de datos correcta, supone decidir el método de construcción del modelo y tipo de información que va a ser representada. Implica por lo tanto, un esquema concreto de almacenamiento y gestión informática de los datos, como también la necesidad de traducir los algoritmos a formas compatibles con la

estructura de datos elegida. Por último, supone aceptar las limitaciones que las aplicaciones informáticas puedan tener para gestionar la información en el formato elegido. En suma, la elección de la estructura de datos es importante porque condiciona el futuro manejo de la información.

Los SIG y algunas aplicaciones dedicadas expresamente al tratamiento de los MDT usan, en la práctica, dos alternativas: matrices regulares y TIN, existiendo por supuesto otras. Para esta sección, se mencionarán como principales estas dos.

La estructura TIN

Una forma de representación espacial muy frecuentemente utilizada en los programas SIG, es la llamada estructura TIN. Creada en 1978 por Peucker, y denominada "Triangulated Irregular Network". Es una estructura formada por triángulos irregulares, contruidos mediante un ajuste y calce sucesivo de planos limitados por tres puntos cercanos no colineales entre sí. El efecto que produce el conjunto es el de un mosaico, simulado la forma del terreno con distintos niveles de detalle, en función de la complejidad del relieve.

Esta construcción se basa esencialmente en puntos, conociendo sus tres coordenadas (x, y, z), aunque también pueden generarse a partir de curvas de nivel (Arc/View por ejemplo, permite crear estructuras TIN a partir de un vector de curvas de nivel). Los puntos se seleccionan de modo que aporten la mayor cantidad de información sobre la disposición de la topografía. Aquí, los llamados Puntos Críticos del Relieve, necesariamente deben estar representados para lograr un buen conjunto inicial de datos. También, se consideran de alta importancia las Líneas Críticas del Relieve: líneas de cumbres, y cauces de los ríos, principalmente.

La triangulación se realiza mediante la unión de los puntos no colineales mediante líneas rectas, formándose una red de triángulos denominados Triángulos de Delauney, sobre los cuales se pasa a la etapa de interpolación, en donde cada tres puntos que definen a un triángulo en la red, es posible calcular la ecuación del plano que los contiene. A partir de esta ecuación, ahora es posible calcular la altura de cada uno de los puntos interiores del triángulo. Este procedimiento supone que la altura varía linealmente al interior del triángulo, por lo tanto, las alturas interiores están entre la máxima y la mínima de los vértices.

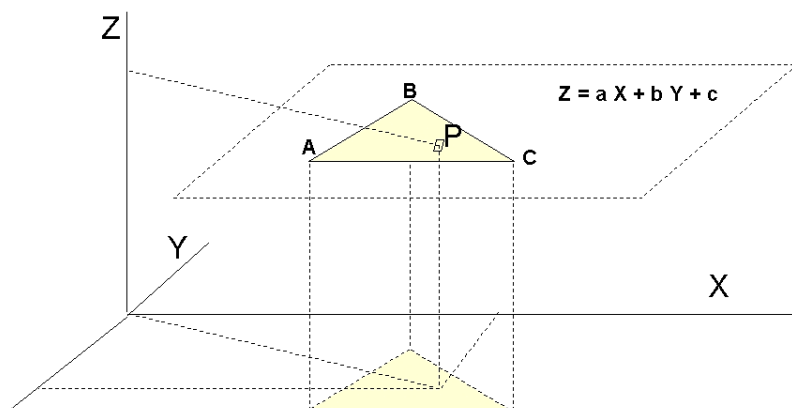


FIGURA VII-1. Representación espacial de la triangulación TIN

La estructura matricial

Estructura regular

Las matrices regulares son frecuentemente utilizadas para representar el relieve del terreno. Se construyen superponiendo una retícula sobre el terreno y extrayendo la altitud media de cada celda. Normalmente, la retícula es una red regular de malla cuadrada en donde la localización espacial de cada dato está determinada de forma implícita por su situación en la matriz, definidos el origen y el valor del intervalo entre filas y columnas. La matriz regular es la estructura más utilizada para construir los MDT debido a su cómodo manejo informático y simplicidad estructural, estando presente en casi la totalidad de los programas SIG.

Estructura anidada

Esta estructura, también es conocida como "quadtree". Aquí, las celdas de la matriz pueden contener datos elementales como en las matrices regulares. La diferencia radica en que el quadtree opera con una estructura de árbol jerárquico de matrices elementales y de profundidad arbitraria con n niveles, en cuyo caso, la resolución espacial se duplica en cada nivel.

Los quadtrees se han utilizado en el tratamiento de variables nominales. Los trabajos pioneros en MDT parecen corresponder a Ebner y Reinhardt (1984, 1988), que utilizan un modelo mixto de matrices jerárquicas y estructuras TIN.

LA INTERPOLACIÓN ESPACIAL Y LOS MÉTODOS DE INTERPOLACIÓN

El proceso de interpolación espacial consiste en la estimación de los valores que alcanza una variable Z en un conjunto de puntos definidos por un par de coordenadas (X,Y) , partiendo de los que adopta Z en una muestra de puntos situados en el mismo área de estudio. La estimación de valores fuera del área de estudio se denomina extrapolación.

Dicho de otro modo, y en términos simples, corresponde a la estimación del valor de una variable o propiedad en un sitio no muestreado y en un área cubierta por observaciones existentes. El proceso lleva a la creación de una superficie matemática, regionalizada y atributada.

En algunos casos pueden utilizarse otras variables de apoyo a la interpolación/extrapolación. El área de estudio vendría definida, aunque no de forma muy clara, por el entorno de los puntos en los que si se dispone de dato. Lo más habitual es partir de medidas puntuales (variables climáticas, variables del suelo) o de isolineas (curvas de nivel). Los métodos que se utilizan en uno u otro caso son bastante diferentes.

Todos los métodos de interpolación se basan en la presunción lógica de que cuanto más cercanos se encuentren dos puntos sobre la superficie terrestre más se parecerán, y por tanto los valores de cualquier variable cuantitativa que sea medida en ellos, serán más próximos a la realidad.

De acuerdo al principio básico de que los puntos más cercanos son más similares que puntos más lejanos (Ley Tobler de la geografía), la interpolación puede efectuarse de múltiples formas. La clasificación básica de métodos es la siguiente:

Clasificación de los métodos de interpolación

- ❑ Interpolación Puntual / Areal
- ❑ Interpoladores Globales / Locales
- ❑ Interpoladores Exactos / Aproximados
- ❑ Interpoladores Estocásticos / Determinísticos
- ❑ Interpoladores Graduales / Abruptos

Revisemos brevemente cada uno de estos métodos. Posteriormente se describirán con más detalle.

Interpolación Puntual / Areal

En este método, se dispone de un conjunto inicial de puntos cuya ubicación y valor son conocidos. Con ellos, es posible determinar los valores de otros puntos de la vecindad. Usualmente este procedimiento se denomina “punto a punto”, y se utiliza para datos que son posibles capturarlos en ubicaciones puntuales (estaciones hidrometeorológicas, pozos, puntos de altitud con GPS, etc.).

Un procedimiento muy utilizado en los procesos de interpolación es el traspaso de líneas a puntos, generalmente para convertir curvas de nivel a modelos de terreno en formato raster. También puede realizarse una interpolación a partir de áreas puntuales, en donde la estimación se extrapola a zonas vecinas.

Interpoladores Globales / Locales

Los interpoladores globales determinan una función que se usa en toda la región de interés. Los algoritmos entregan como resultado superficies suaves, asumiendo que no se presentan cambios abruptos en la topografía.

Si bien es cierto los resultados pueden resultar vistosamente llamativos, es necesario previamente plantearse una idea acerca de la forma de la superficie o una idea de la tendencia general. De no ser así, los resultados pueden modelar una superficie muy distinta a la real.

A su vez, los interpoladores locales aplican una función o algoritmo a una subzona dentro del mapa, y en donde el cambio en un dato de entrada solo afecta a la subzona dentro del mapa; no así en el caso de un interpolador global, donde el cambio de un valor de la variable afecta a los resultados del mapa completo.

Interpoladores Exactos / Aproximados

Los interpoladores exactos otorgan una mayor importancia a los puntos a partir de los cuales se está interpolando.

En esta categoría, es posible distinguir interpoladores denominados proximales. El más conocido es el B-Spline, el cual otorga importancia a los puntos de entrada. Este tipo de herramientas se utiliza cuando existe alguna incertidumbre sobre los valores de la superficie dados. Se parte del supuesto que en el conjunto de datos inicial existen

tendencias globales que varían muy lentamente, escondidas por fluctuaciones locales que varían muy rápidamente y producen incertidumbre (error) en los valores registrados. El efecto de suavización reducirá entonces los efectos del error en la superficie resultante.

Interpoladores Estocásticos / Determinísticos

Los métodos estocásticos incorporan el concepto de aleatoriedad. En ellos, la superficie interpolada es conceptualizada como una de las muchas que pudo ser observada y que pudieron haber producido los puntos con valores conocidos. Como se verá en una tabla resumen más adelante, los procesos matemáticos que dan lugar a una superficie de probabilidades son complejos, y de alto trabajo computacional; sin embargo, los resultados pueden ser altamente satisfactorios si el ajuste estadístico es el correcto y la superficie real a interpolar adecuado para este tratamiento. Los interpoladores estocásticos incluyen el análisis de tendencias de superficie (trend surface analysis), análisis de Fourier, y Kriging.

Interpoladores Graduales / Abruptos

Tal como señala su nombre, los interpoladores graduales interpolan superficies con cambios graduales en la topografía. Sin embargo, si el número de puntos usados se reduce a un pequeño número o a un punto, habrá cambios abruptos en la superficie. También está el caso donde se hace necesario incluir barreras en el proceso de interpolación, tales como precipicios u otro tipo de barreras naturales.

INTERPOLACIÓN ESPACIAL A PARTIR DE PUNTOS

Métodos Puntuales Exactos

Bajo esta clasificación, todos los valores se asumen iguales al valor del punto más cercano conocido. En general, la estructura de salida corresponde a polígonos de Thiessen con cambios abruptos en las fronteras. Corresponde a un interpolador de tipo local, en donde la carga computacional para el procesamiento es baja. Más detalles sobre las ventajas y desventajas de este método se encuentran en el cuadro VI-3.

Otro interpolador perteneciente a esta categoría es el B-Spline, el cual utiliza una ecuación polinómica con primera y segundas derivadas. Su planteamiento matemático asegura continuidad en la estimación de la elevación (continuidad de orden 0 si la superficie no tiene escarpes), pendiente (continuidad de orden 1 si la pendiente no cambia abruptamente), y curvatura del terreno (continuidad de orden 2 para una curvatura mínima). También corresponde a un interpolador local, con una carga computacional moderada, y es apto para las estimaciones en superficies suaves, sin grandes cambios en la topografía.

Métodos Puntuales Aproximados

Se encuentran aquí el análisis de tendencias para una superficie (trend surface analysis), y la ponderación por distancias.

En el análisis de tendencias, la superficie se aproxima a una función polinómica, dependiendo de grado de la ecuación predictiva. Es un interpolador global, en donde se asume que la tendencia general de la superficie es independiente de errores aleatorios propios de los datos de entrada. La carga computacional para este método es

relativamente baja, y su principal limitante es que presenta problemas de borde severos y genera superficies redondeadas. A su vez, en el IDW (ponderación por distancia), la interpolación del punto se realiza asignando pesos a los datos del entorno en función inversa de la distancia que los separa -inverse distance weighting-.

INTERPOLACIÓN ESPACIAL A PARTIR DE LÍNEAS

Interpolación a partir de curvas de nivel

La interpolación a partir de puntos resulta necesaria cuando, *a priori*, no se conoce nada acerca de la distribución espacial de la variable y es necesario medirla en una serie de puntos de muestreo a partir de los que estimar sus valores en toda el área de trabajo. En el caso de la topografía, si contamos con un mapa topográfico, el caso es algo diferente ya que lo que vamos a tener no son puntos sino isolíneas derivadas del análisis de pares de fotogramas estereoscópicos. El procedimiento va a ser en primer lugar digitalizar las curvas de nivel y en segundo lugar utilizar alguno de los programas que interpolan a partir de curvas. En general el fundamento de todos estos métodos consiste en hacer interpolaciones lineales o más complejas entre curva y curva. Los algoritmos que utilizan IDRISI o GRASS son bastante simples pero presentan una serie de problemas a tener en cuenta. Estos problemas se derivan directamente del tipo de algoritmo, que pueden ser resueltos con algo de esfuerzo adicional. El proceso de interpolación a partir de curvas de nivel consta de las siguientes fases:

- ❑ Digitalización (en tableta digitalizadora o en pantalla)
- ❑ Rasterización del vectorial (cuando sea necesario)
- ❑ Interpolación
- ❑ Análisis de errores

Existen tres problemas fundamentales que pueden dar lugar a errores y que a veces no son fáciles de corregir. En primer lugar, las curvas deben estar cerradas y deben cortar los límites de la capa raster creada; si esto no ocurre, se producen errores en la interpolación, frecuentemente representados por largos trazos o segmentos rectilíneos que cruzan desde una región de la imagen a otra. Para el caso de los ficheros vectoriales, es un punto relevante también a considerar (figuras VII-2 y VII-3).

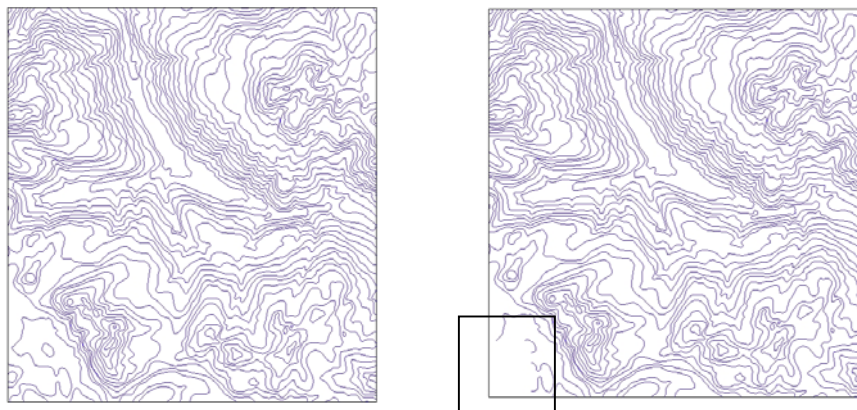
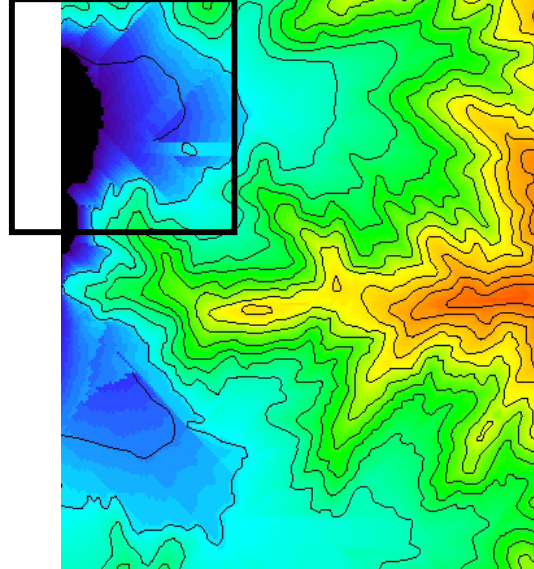


FIGURA VII-2. Fichero vectorial de curvas de nivel.
Izquierda: Todas las curvas llegan al borde; bien. Derecha: En el recuadro, se indican curvas "cortadas" que no llegan al borde de la cobertura; mal.

FIGURA VII-3. En el recuadro, se indican algunos errores típicos en la interpolación debido a cadenas cortadas.

Otro aspecto a considerar es que, si se rasterizan las curvas, éstas no deben superponerse ya que en la práctica equivale a que algunas curvas no se cierren. A su vez, las curvas de nivel rasterizadas mantienen su valor con lo que el modelo de terreno adquiere un aspecto escalonado.



Tanto en GRASS como en IDRISI (hasta la versión 2.0), el resultado es un modelo de terreno con valores enteros por lo que en las zonas llanas se puede producir un escalonamiento artificial (bandas continuas por intervalos de altitud) y trazos o rectas que cruzan la imagen si las unidades en que se mide la altitud no permiten una relación equidistancia de curvas de nivel/número de píxeles entre curvas de nivel adecuado.

Una solución a este problema sería retocar a mano las curvas de nivel rasterizadas pero además de muy trabajoso supone desplazar curvas arbitrariamente. Otra solución más adecuada sería generar modelo de terreno con áreas especialmente abruptas con tamaños de pixel más pequeños. Estos modelos pueden después degradarse al tamaño de pixel adecuado y superponerse al original. En el manual de IDRISI se recomienda utilizar un filtro de paso bajo (media aritmética) pero esto supone aplicar la solución a toda la imagen en lugar de sólo a los píxeles problemáticos.

ASPECTOS A CONSIDERAR EN LA VALIDACIÓN DE UN MODELO DE TERRENO

Para verificar la calidad de un mapa interpolado debe utilizarse un conjunto de validación formado por una serie de puntos de muestreo, idealmente que no se hayan utilizado para realizar la interpolación y en los que sea posible estimar los valores previamente medidos. La diferencia entre el valor medido y el estimado es el error de estimación en ese punto. De este modo a cada punto de validación se asigna un error. El conjunto de los errores debe tener las siguientes características:

- ❑ Media de errores y media de errores al cuadrado próximo a cero
- ❑ Los valores de error deben ser independientes de su localización en el espacio y no estar autocorrelacionados
- ❑ La función de distribución de los errores debe aproximarse a la distribución normal

El problema es que no siempre es posible ni conveniente disponer de un conjunto de validación independiente de los puntos de muestreo utilizados para interpolar. Por ello suele utilizarse la técnica de la validación cruzada en la que se estiman los valores en los puntos de muestreo, excluyéndolos del conjunto de interpolación.

Este método se utiliza para comprobar los resultados de un procedimiento de interpolación. Consiste en la estimación de los valores de Z en todos los puntos en los que se conoce *a priori* el resultado y el error cometido. La media de los errores cometidos en los diferentes puntos es un buen indicador del error medio global.

Un problema característico en la creación de MDT es la aparición de concavidades o pozos a lo largo de los fondos de valles. Este tipo de problemas puede generarse por el uso de funciones de interpolación de grado superior en zonas conflictivas, como ocurre en las superficies de tendencia. Es el caso de la simulación de procesos hidrológicos, en donde la presencia de concavidades condicionan las líneas de flujo de los torrentes.

La corrección de concavidades de los MDT ya existentes es posible mediante algoritmos que simulan la inundación y relleno de los pozos. Este proceso modifica el modelo de elevaciones original por lo que en algunos casos debe ser aplicado con enorme precaución. Una forma de evitar el problema es introducir la red hidrológica como líneas de rotura en la generación del modelo TIN. De esta forma, la triangulación utilizará estas líneas como lados de los triángulos trazando correctamente las líneas de flujo.

LA INTERPOLACIÓN. UNA REVISIÓN MÁS DETALLADA.

Como se señalaba anteriormente, se distinguen dos métodos generales de interpolación. El primero corresponde a la familia de los Métodos Globales, en donde se construyen superficies de tendencia a partir de coordenadas geométricas, métodos de regresión, y análisis espectrales. Un segundo método se denomina Local, en donde se distingue la construcción de polígonos de Thiessen, interpolación lineal inversa a la distancia y métodos splines. Todos estos sub-métodos son relativamente confiables, requiriendo sólo un conocimiento de métodos estadísticos simples, como la estadística descriptiva. Estos métodos han sido frecuentemente incluidos en los softwares SIG comerciales. El siguiente cuadro describe los principales métodos utilizados para efectuar interpolaciones, indicando en ellos los datos de entrada al proceso. Al final de este capítulo se describe en más detalle un resumen de las capacidades de éstos y otras características de interés.

CUADRO VII-1. Descripción de los principales métodos de interpolación

MÉTODO	DESCRIPCIÓN	DATOS DE ENTRADA
INTERPOL	Genera una superficie mediante la interpolación de datos puntuales. El procedimiento de interpolación puede ser por medias, medias móviles ponderadas por la distancia o por un modelo potencial.	Datos puntuales
INTERCON	Genera un Modelo Digital del Terreno raster mediante la interpolación de curvas de nivel digitalizadas.	Iso-líneas Curvas de nivel
TIN	Genera una red de triángulos irregulares a partir de contornos o una malla de puntos	Datos puntuales o lineales
TINSURF	Genera un raster a partir de un TIN	A partir de un TIN existente
KRIGING ESPACIAL	Primer paso de la geoestadística para la modelación de la variabilidad o continuidad espacial	Datos de precipitación, temperaturas, presiones, etc.
KRIGING Y SIMULACIÓN	Interpolación de los datos a partir de un modelo estocástico	Superficie ajustada a un modelo estocástico
THIESSEN	Construcción de polígonos (llamados Thiessen) a partir de un conjunto inicial de puntos. La construcción se genera mediante la asignación de puntos de control más próximos y representativos dentro de una zona de influencia. De esta forma el espacio se segmenta y representa por este conjunto de puntos	Datos puntuales
TREND	Calcula las ecuaciones polinómicas de superficies de tendencia lineal, cuadrática y cúbica para datos espaciales, e interpola superficies a partir de estas ecuaciones	Iso-líneas Datos puntuales

Métodos geoestadísticos utilizando autocorrelación espacial, conocido como kriging, requiere del conocimiento de los principios básicos de estadística espacial. En particular, este método presenta distintas variantes que luego serán descritas, siendo en todos los casos, utilizados cuando la variación de algunos atributos es irregular, y las muestras son tales, que métodos simples de interpolación pueden generar estimaciones muy alejadas de la realidad.

Los métodos geoestadísticos proporcionan estimaciones probabilísticas sobre la calidad de la interpolación. Estos métodos también permiten efectuar cálculos adicionales sobre las características físicas del relieve, tales como mapas de sombreado y modelo de exposiciones.

Adicionalmente, estos métodos permiten interpolar distintas funciones continuas de una variable de interés, incorporando grupos de datos altamente asociados al conjunto inicial de información. Por lo general, los resultados son de alta fiabilidad si la base de la información a interpolar también lo es.

Algunos SIG incluyen métodos geoestadísticos simples, pero usualmente son limitados en cuanto al alcance y profundidad de los análisis, y las capacidades de trabajar con distintos formatos de intercambio con otros softwares SIG especializados.

Interpolación Global

Los métodos de interpolación pueden ser divididos en dos grandes grupos: los interpoladores globales, encargados de efectuar predicciones numéricas en un área de interés, generalmente extensa, e interpoladores locales destinados a trabajar sobre superficies más pequeñas, en cuyo caso los análisis de vecindad o estimación de valores intermedios en torno a los puntos de observación, proporcionan calidad en la micro-superficie a interpolar.

Los interpoladores globales son utilizados preferentemente para efectuar análisis macro en la distribución espacial de los datos, con el propósito de examinar tendencias, y efectos en la variación global de una o más variables, expresados a través de una superficie de tendencia. Particular importancia presenta el análisis de residuos, y varianza en la distribución de los valores presentes en una superficie generada por interpolación.

Los métodos de regresión incorporados como parte de los algoritmos de interpolación, permiten explorar las posibles relaciones entre atributos de una misma variable de medición, de modo que resulta sencillo predecir el valor de este atributo para una vecindad extendida, dentro de un intervalo de confianza. Estos métodos se basan solamente en coordenadas geográficas de los puntos de control (en cuyo caso el método se denomina análisis de superficies de tendencia), o bien en la relación entre atributos de una o más variables espaciales (conocido en el análisis de regresión como funciones de transferencia).

Métodos alternativos como las transformadas de Fourier también son utilizados por procedimientos para interpolar, particularmente en el análisis de imágenes mediante sensores remotos.

Predicciones Globales usando modelos de clasificación

Cuando la información es posible espacializarla, muchas veces conviene asumir que las observaciones que componen ese espacio pertenecen a una población estadísticamente estacionaria. Por supuesto, que también puede asumirse que la localización geográfica de los elementos geográficos a medir, presentan una variación espacial – en coordenadas – en el tiempo; casos al respecto se ven con frecuencia en el ámbito de los recursos naturales renovables y variables meteorológicas, representados por constantes variaciones en la expresión espacial de sus atributos. Si se opta por este último supuesto, entonces automáticamente estamos pensando en una metodología para clasificar, dentro del marco de la predicción espacial, lo cual implica que la estructura espacial de la variación es determinada por externalidades residentes en las unidades espaciales. Dicho de otro modo, si la opción es la clasificación, entonces será necesario computar las predicciones utilizando metodologías estadísticas como el análisis de varianza (ANOVA).

La clasificación mediante polígonos, asume que la variación intrapolígono es menor que entre polígonos vecinos, siendo los cambios más importantes – reflejados en el valor de atributo – en los bordes adyacentes a cada unidad homogénea. Este modelo conceptual es frecuentemente utilizado en el estudio de unidades de suelos, y en general en variables relacionadas con el paisaje, justamente para definir – entre otras variables – unidades homogéneas de suelos, unidades de paisaje y ecotipos, donde objetos como la tierra (expresado espacialmente como polígonos), terrazas marinas,

laderas, quiebres de pendientes y otros, son reconocidos como rasgos útiles para derivar a estudios más específicos acerca del paisaje.

Un modelo estadístico simple es el correspondiente a un análisis de varianza, dado por:

$$Z(x_0) = \mu + \Phi_k + \varepsilon$$

donde Z es el valor del atributo en la localización x_0 ; μ es el promedio general de Z sobre el dominio de interés; Φ_k es la desviación entre μ y el promedio de unidades k , y ε es el error residual, usualmente conocido como "ruido" en el modelo.

Este modelo asume que por cada clase k de atributos, se presenta una distribución gaussiana (normal), mientras que el promedio de valores por atributo en cada clase k es $\mu + \Phi_k$, cuando es estimado a partir de un conjunto de muestras independientes, y que por tanto se asume que existe independencia espacial. La variación al interior de cada clase está dada por ε , asumiéndose que es la misma para todas las clases.

La varianza relativa (σ_w^2 / σ_z^2), donde σ_w^2 es el agrupamiento al interior de cada clase y σ_z^2 es la variación de la muestra, es una medida de la bondad de ajuste de la clasificación. Ambos parámetros pueden ser estimados cuando todas las unidades espaciales que contienen las muestras, contienen más de 1 muestra. Una baja varianza relativa implicaría por lo tanto, un buen ajuste en la clasificación.

Este modelo consta como supuestos, que las variaciones entre el valor de z al interior de las unidades de mapa son aleatorias y no espacialmente contiguas. Por otra parte, todas las unidades de mapa presentan la misma varianza al interior de cada clase, la cual es uniforme al interior de cada polígono o unidad de mapa. Se asume además que todos los atributos están normalmente distribuidos.

Interpoladores globales utilizando superficies de tendencia

Cuando la variación de un atributo de interés se presenta en forma frecuente en un paisaje, podría asumirse en primera instancia que un modelo acorde a utilizar correspondería a una superficie matemática lisa. Hay varias maneras de hacer esto, siendo la más frecuente, el fabricar un polinomio de la superficie, tomada de los datos puntuales que componen esta unidad, y que permiten fabricar el contorno mediante el ingreso de coordenadas para estos datos.

Un modelo simple para abordar la variación espacial es la regresión múltiple de valores en relación a las localizaciones geográficas. La idea es crear un polinomio o una superficie si los datos están expresados en 2 dimensiones, minimizando la suma de cuadrados para $\hat{z}(x_i) - z(x_i)$. Se asume que las coordenadas espaciales (x, y) son variables independientes, que z es el atributo de interés, y que las variables dependientes están normalmente distribuidas.

Un simple ejemplo, sería el considerar el valor de un atributo ambiental z que ha sido medido mediante un transecto de puntos $x_1, x_2, x_3, \dots, x_n$. Si además de la menor variación, el valor de z se incrementa linealmente conforme varía la posición x , entonces el rango de variación es amplio, aproximándose al modelo de regresión

$$Z(x) = b_0 + b_1x + \varepsilon$$

Donde b_0 y b_1 son coeficientes del polinomio, conocidos respectivamente como la intersección y la pendiente en una regresión simple. El error residual (ε) se asume que presenta una distribución normal e independiente a los valores de x . En muchas circunstancias z no representa una función lineal de x , lo cual es muy común en variables dinámicas como las presentes en la naturaleza. Un modelo un tanto más complejo que responde a una función no lineal, es la correspondiente a un polinomio de grado superior como

$$Z(x) = b_0 + b_1x + b_2x^2 + \varepsilon$$

Por el incremento en el número de términos es frecuente encontrar numerosos grupos de puntos que dificultan establecer una clara tendencia de la curva. En dos dimensiones, los polinomios derivados por regresión múltiple para las coordenadas (x,y) , presentan la forma

$$F\{(x,y)\} = \sum_{r+s \leq p} (b_{rs} * x^r * y^s)$$

En esta expresión, podemos ver el comportamiento de los tres primeros términos:

b_0	plano
$b_0 + b_1 * x + b_2 * y$	lineal
$b_0 + b_1 * x + b_2 * y + b_3 * x^2 + b_4 * xy + b_5 * y^2$	cuadrático

El entero p que aparece en la sumatoria anterior corresponde al orden de la superficie de tendencia. Existen $p = (p+1)(p+2)/2$ coeficientes, los cuales se encargan de minimizar la expresión

$$\sum_{i=1}^n \{z(x_i) - f(x_i)\}^2$$

donde x es el vector notación para (x,y) . En un plano horizontal el orden de la ecuación es cero (b_0); en un plano inclinado es de primer orden, en una superficie cuadrática se segundo orden, mientras que en una superficie cúbica con diez parámetros es de tercer orden. Las superficies de tendencia pueden ser visualizadas por estimación del valor $z(x)$ para todos los puntos representados en una grilla regular.



FIGURA VII-4. Izquierda: ajuste de grado 1 para un set de puntos; centro: ajuste de grado 2; derecha: ajuste de grado 3.

La principal ventaja del análisis de superficies de tendencia es que es una técnica sencilla de implementar y entender, respecto a otros procedimientos de cálculo. Sin embargo, esta técnica presenta la desventaja que las superficies generadas son susceptibles a efectos de borde, deformando muchas veces los datos al exterior de cada polígono cuando el polinomio es de grado superior.

La fiabilidad estadística de una superficie de tendencia puede ser testeada mediante el uso de técnicas de análisis de varianza, particionando el análisis entre la varianza de la tendencia y los errores residuales. Por cada n observaciones, corresponden $(n-1)$ grados de variación respecto a la varianza total. Los grados de libertad (m) para la regresión son determinados por el número de términos en la ecuación polinómica de regresión, excluyendo el coeficiente (b_0).

Para una regresión lineal $z(x,y) = b_0 + b_1x + b_2y$, los grados de libertad corresponden a $m = 2$, mientras que para la desviación, los grados corresponden a $(n-1)-m$.

MÉTODOS LOCALES DETERMINÍSTICOS PARA LA INTERPOLACIÓN

Los métodos anteriormente descrito hasta ahora, han expuesto una estructura espacial global para realizar la interpolación. En todos los casos las variaciones locales han estado sujetas al "ruido" de la aleatoriedad y distribución de los datos. Por lo anterior se han desarrollado métodos alternativos que buscan examinar la vecindad de los puntos muestrales en forma más precisa, lo cual implica resolver en mejor forma la definición de un área de búsqueda para estimar los valores contiguos a esta vecindad, encontrar los puntos apropiados dentro de esta vecindad, elegir una función matemática que represente la variación entre valores originales y valores a predecir, y las evaluaciones posteriores de los datos estimados en una malla regular (o grilla de datos). Esto corresponde a un proceso iterativo, que se extiende a toda el área bajo análisis, y por tanto a todos los puntos que pertenecen a esta grilla interpolada.

Estos métodos implican resolver por tanto los siguientes problemas técnicos:

- ❑ El tipo de función a utilizar para la interpolación
- ❑ El tamaño, forma y orientación de la vecindad
- ❑ El número de puntos muestrales
- ❑ La distribución de los datos muestrales: en una grilla regular o irregular.

Aquí se describirán someramente diferentes funciones de interpolación tales como:

- ❑ Polígonos de Thiessen
- ❑ Pesos inversos a la distancia (IDW)
- ❑ Superficies Splines y otras funciones no lineales como los estimadores Laplacianos
- ❑ Funciones de Covarianza

Todas ellas son funciones locales que generan superficies suaves sobre un radio de búsqueda predefinido.

Polígonos de Thiessen (Dirichlet / Voronoi)

Los polígonos de Thiessen (conocidos también como Dirichlet o Voronoi), utilizan un modelo espacial de predicción sobre la base de puntos locales dentro de un área

determinada, generando una superficie continua, y subdividida en polígonos mediante cambios abruptos en los atributos presentes en los bordes de cada zona o polígono. La estructura de datos de salida es un raster conocido como Polígonos de Thiessen en torno a cada punto muestral.

Estos polígonos dividen una región de una manera tal que ella es totalmente definida por la configuración de los puntos, con una observación por celda. Si los puntos están definidos en una malla rectangular y cuadrada, los polígonos son todos iguales, con celdas regulares y de dimensiones iguales al lado del cuadrado de la malla original. Si los datos originales están organizados de una manera irregular, entonces los polígonos forman un enrejado irregular.

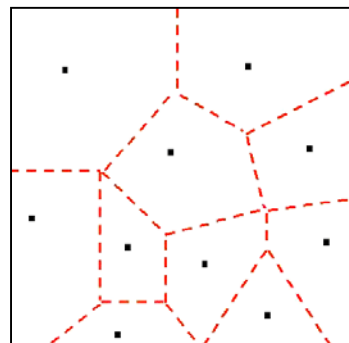


FIGURA VII-5. Representación esquemática

Los polígonos de Thiessen han sido frecuentemente utilizados en los SIG como un método rápido para relacionar puntos distribuidos en el espacio. Se han empleado por ejemplo para aplicaciones ecológicas, como determinación de territorios y áreas de influencia, en el análisis de observaciones meteorológicas donde se supone que los datos obtenidos en un lugar preciso (un aeropuerto por ejemplo) son válidos para toda la ciudad (polígono). Evidentemente esta suposición no se la puede aplicar a datos que varían de una manera continua como, por ejemplo, presión y temperatura del aire. Como existe una sola observación por polígono no hay manera de calcular variaciones al interior del polígono.

La principal ventaja de este método es que puede aplicarse fácilmente a datos cualitativos como tipos de vegetación o uso de la tierra.

Método de Tobler

El método anterior asume homogeneidad al interior de cada polígono, y que los atributos sólo varían en las intersecciones de ellos. En muchas aplicaciones esto podría constituir una aproximación grosera, particularmente porque los atributos de una variable continua presentan (valga la redundancia) continuidad y contigüidad espacial, como por ejemplo la distribución de densidad poblacional o distribución de las lluvias en una zona determinada (¿cuántas estaciones meteorológicas y a qué distancia debieran disponerse en una región para obtener real exactitud en la distribución de los datos?). Tobler (1979), abordó este problema inventando un interpolador que remueve los cambios abruptos entre bordes de polígonos. El método, conocido como Tobler asume una "masa" de datos que se agrupa en una región de origen. Indica que el "volumen" del atributo para esa masa (número de personas por ejemplo) en una entidad espacial (polígono o superficie administrativa), presenta un gradiente de similitudes con los valores de polígonos vecinos, independiente si la variación global entre polígonos es representada por superficies lisas u homogéneas claramente distinguibles entre sí.

El método pronostica una mayor similitud en la distribución de algunas variables como las demográficas, en donde se presentan variaciones en la densidad de los datos (personas, monto de precipitaciones), que en algunas subzonas es bajo, en otras de mediana y gran magnitud que el promedio.

La primera condición matemática para este procedimiento, para la preservación de masa es la condición de regionalización dada por

$$\iint_{R_i} f(x, y) dx dy = V_i$$

para todo i , donde V_i representa el valor (población, frecuencia, otro tipo atributo) para la región R_i . Esta expresión entrega una variación suave entre límites de polígonos

Al menos que se trate de obstáculos físicos como barreras, las densidades en áreas vecinas tienden a ser semejantes en las zonas de unión, por lo que las superficies lisas quedan interconectadas por cambios graduales entre el valor de un atributo de dos regiones contiguas. La condición de cambio la otorga la condición Laplaciana minimizada expresada como

$$\iint_R \left(\frac{\partial^2 f}{\partial x^2} + \frac{\partial^2 f}{\partial y^2} \right) dx dy$$

donde R es el dominio de todas regiones. La condición de borde general para esta expresión es

$$\frac{\partial f}{\partial \gamma} = 0$$

donde γ indica el borde entre regiones.

Método de medias ponderadas por el inverso de la distancia (IDW)

El IDW es un método que combina la idea de proximidad expuesta por los polígonos de Thiessen, con un cambio gradual en la superficie de tendencia.

Aquí la interpolación del punto problema se realiza asignando pesos a los datos del entorno en función inversa de la distancia que los separa (inverse distance weighting, IDW). Se asume que el valor de un atributo z varía en forma de gradiente respecto a la distancia entre puntos. Se supone que en cada localización, los puntos muestrales más próximos son los que tienen los valores más parecidos, y que además esta semejanza disminuye con la distancia entre el punto calculado y el punto muestral. Este procedimiento tiene el inconveniente de recoger el llamado efecto de cluster. Los puntos muestrales pueden estar localizados en una grilla (regular o irregular), perteneciente al área a interpolar. Otro inconveniente usual, es la generación de concavidades y picos al interior de la imagen, producto de la distribución espacial de los puntos muestrales y del radio de búsqueda elegido para la interpolación.

La expresión que denota el cambio en las medias ponderadas por la distancia corresponde a

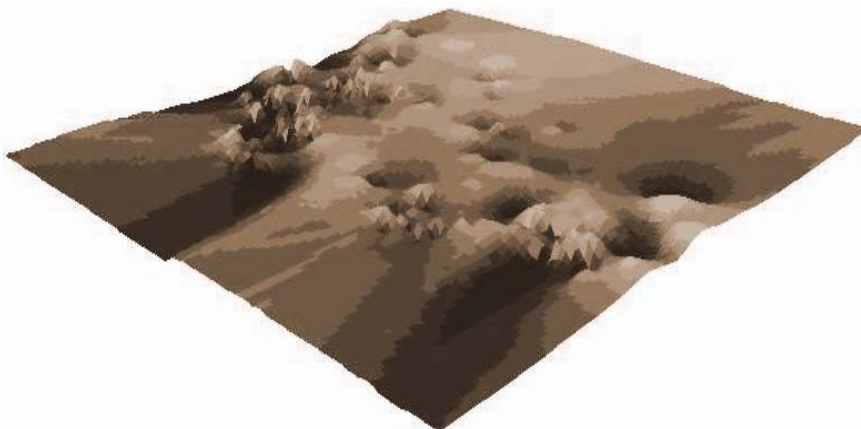
$$\hat{Z}(x_0) = \sum \lambda_i z(x_i); \quad \sum_{i=1}^n \lambda_i = 1$$

Donde los pesos λ_i están dados por $\zeta(d(x, x_i))$, donde d está sujeto a la tendencias exponenciales $e^{-(d)}$ y $e^{-(d^2)}$. La forma más común para $\zeta(d)$ es la función de proporcionalidad entre el peso y la distancia, expresada como

$$\hat{Z}(x_0) = \sum_{i=1}^n z(x_i) * d_{ij}^{-r} / \sum_{i=1}^n d_{ij}^{-r}$$

Donde x_j corresponde a los puntos donde la superficie está siendo interpolada y x_i son los puntos muestrales. Dado que en la expresión anterior $\zeta(d) \rightarrow \infty$ y $(d) \rightarrow 0$, el valor para un punto interpolado que coincide con la localización de un punto de origen, es copiado sobre sí mismo.

FIGURA VII-6. Superficie interpolada mediante IDW. Se distinguen imperfecciones en la superficie tales como cavidades y picos que no corresponden a la realidad. Este es un aspecto perfectamente corregible en función de los datos disponibles.



Este efecto causa una anomalía, donde el valor calculado para un punto, depende mucho de grupos de puntos muy cercanos entre sí, y por ende con valores muy parecidos. Su principal problema se origina en el hecho de que el rango de salida de los valores interpolados, está limitado por el rango de los datos de entrada. Por otra parte, es difícil determinar el número de puntos muestrales a considerar. Tampoco es posible operar sobre puntos irregularmente espaciados.

Splines

El método de los splines ajusta funciones polinómicas (como en una interpolación global mediante regresión) en una vecindad local. En general los resultados son muy buenos, con la ventaja de poder modificar una serie de parámetros en función del tipo de topografía.

El proceso se basa en una interpolación y unión continua de segmentos a partir puntos cuya función polinomial está representada por

$$p(x) = p_i(x) \quad x_i < x < x_{i+1}, \quad i = 0, 1, \dots, k-1 \quad i = 1, 2, \dots, k-1$$

donde los puntos $x_0, \dots, x_{(k-1)}$ que dividen el intervalo (x_0, x_k) en k sub-intervalos, denominados puntos de quiebre, mientras que los puntos de la curva en esas posiciones se los llama nodos.

Generalmente estas funciones polinomiales son de grado 1, 2 ó 3, llamadas sucesivamente lineales, cuadráticas o cúbicas. En estudios de las ciencias ambientales, se tiende a utilizar una variante al método dada las dificultades prácticas de cálculo de splines por un gran conjunto de subintervalos independientes como es el caso de una línea numerizada. Esta variante se encuentra en la gran mayoría de aplicaciones y softwares, denominada B-splines.

Los B-splines son sumas de splines que por definición poseen un valor cero al exterior del intervalo de interés. Emplean un polinomio que proporciona una serie de parches que determinan una superficie cuya primera y segunda derivada son continuas.

La aplicación de esta metodología asegura continuidad en:

- ❑ La Elevación: cuando la superficie no es escarpada
- ❑ La Pendiente: cuando no existen cambios abruptos
- ❑ Curvatura: se obtiene la mínima curvatura

Es un interpolador local, que genera superficies suaves. Si bien es cierto que en general los resultados son muy buenos, especialmente por esta ventaja de poder modificar una serie de parámetros en función del tipo de topografía, no es adecuado para superficies con fluctuaciones marcadas. Presenta además dificultades para el cálculo de áreas y perímetros. Es un método de interpolación muy popular en programas específicos de interpolación de superficies, pero no muy común en los SIG tradicionales.

En el ámbito de los interpoladores, variados son los programas que realizan superficies mediante los métodos descritos anteriormente. Según los datos de entrada, el nivel de error tolerable, la escala y los propósitos del estudio, el usuario podría elegir entre uno y otro método. Naturalmente, existen notables diferencias en el desempeño de algunos métodos y es por ello que frecuentemente se encuentran sorpresas en los resultados, pues se esperaría que la estimación se acercase lo mejor posible a la realidad, y esto no siempre ocurre.

Las razones de lo anterior podrían atribuirse a errores propios del modelo utilizado y a los datos de entrada, por lo cual el usuario debe comprender claramente en qué casos puede aplicar uno u otro modelo.

En este manual se mencionan sólo algunos de muchos softwares que realizan estas tareas. Mencionemos el caso de kriging, un interpolador estocástico incorporado a programas tales como IDRISI y Arc/View.

Algoritmos en IDRISI

Otro de los softwares más empleados en el ámbito de la generación de superficies y su tratamiento es el programa IDRISI, aunque si por especialización se trata, SURFER es el programa más indicado para realizar modelos de terreno. Creado por Clark University Graduate School of Geography, se dispone en el mercado de cuatro versiones de este programa, perfeccionando en su última versión muchos módulos, entre ellos el de la generación de superficies.

Una de las herramientas utilizadas en IDRISI para realizar modelos de terreno a partir de puntos e isolíneas rasterizadas es el módulo de interpolación, el cual contiene las alternativas INTERPOL, INTERCON, TIN, TINSURF, KRIGING, THIESSEN y TREND. Mencionaremos sólo algunos métodos. Pueden encontrarse mejores descripciones de estas opciones en los manuales correspondientes.

INTERCON por ejemplo, genera un Modelo Digital del Terreno raster mediante la interpolación de curvas de nivel digitalizadas. Utiliza una interpolación lineal entre curvas de nivel que produce un modelo anguloso, por lo que es recomendable utilizar después un filtro para suavizar la superficie.

Este algoritmo genera perfiles en torno a los cuatro márgenes del mapa para producir un espacio completamente cerrado por altitudes conocidas. Se crea una serie de perfiles horizontales a través de cada fila utilizando los perfiles de los márgenes y

cualquier curva de nivel que cruce la horizontal trazada. En cada celda de altitud desconocida se registra la altitud del perfil interpolado junto con la pendiente del perfil en ese punto. Producto de sucesivas iteraciones se generan perfiles verticales (a lo largo de las columnas) y diagonales de izquierda a derecha y de derecha a izquierda cruzando la imagen. La altitud final registrada en cada celda será la del perfil de máxima pendiente. Este algoritmo es una modificación del algoritmo CONSURF desarrollado por David Douglas en la Universidad de Ottawa, Canadá.

El algoritmo ha sido validado mediante la utilización de puntos aleatorios de control procedentes de un mapa topográfico, encontrándose mediante pruebas estadísticas que es posible obtener niveles de error que no superan el 1% respecto a la probabilidad que exceda a la mitad del intervalo entre curvas de nivel.

Lógicamente, y como se mostrará en los cuadros finales de este capítulo, la fiabilidad de los resultados dependerá en gran medida del tipo de superficie a interpolar, si poseen o no cambios bruscos en la topografía y si existen o no líneas críticas del relieve tales como fondos de valle, precipicios o picos pronunciados.

Otra variante para la interpolación en IDRISI es el comando INTERPOL, el cual permite generar una superficie mediante la interpolación de datos puntuales. El procedimiento de interpolación puede ser por medias, medias móviles ponderadas por la distancia o por un modelo potencial.

El modelo opera con un radio de búsqueda para efectuar la interpolación de la vecindad, y por lo tanto determina los valores de cada celda sólo en función de los valores de los puntos de control más cercanos. El radio de búsqueda y la cantidad de puntos de control serán esenciales para lograr calidad en las aproximaciones. Por ejemplo, si se encuentran menos de 4 puntos de control, el radio de búsqueda se incrementa temporalmente hasta que se encuentra un número suficiente. Si se encuentran más de 8 puntos de control, el radio de búsqueda es reducido hasta que se encuentran solamente entre 4 - 8 puntos de control.

El proceso de interpolación para INTERPOL es bastante lento, ya que el proceso de llenado de la matriz de datos es directamente proporcional al número de celdas y geométricamente proporcional al número de puntos de control. Se presenta una breve descripción de las capacidades de este interpolador en el cuadro VI-2.

Es importante señalar que muchos métodos de interpolación son solamente apropiados para superficies suaves, y otros para superficies con cambios abruptos en la topografía. La técnica de medias móviles ponderadas utilizada en INTERPOL, produce superficies suavizadas, con los valores mínimos y máximos en las localizaciones de los puntos de control. Cuanto más lejos se encuentra el píxel a interpolar respecto a los puntos de control, la superficie tenderá hacia la altitud media de la zona definida por el radio de búsqueda, lo cual podría llevar a errores en la validación de los resultados.

Los polígonos de Thiessen corresponden a otro método de interpolación descrito en párrafos anteriores, y que lo emplea IDRISI para construir una superficie de polígonos a partir de una serie de puntos. Los polígonos de Thiessen dividen el espacio de forma que cada punto se asigna al punto de control más próximo, definiendo la zona de influencia de cada punto. Este tipo de división espacial se denomina también Teselación Voronoi.

Por último, TREND corresponde a otro algoritmo que calcula ecuaciones polinómicas de superficies de tendencia lineal, cuadrática y cúbica para series de datos espaciales, e interpola superficies a partir de estas ecuaciones. El método ha sido explicado anteriormente.

ARCVIEW de ESRI, es otro software que utiliza procedimientos para interpolar. Básicamente opera con métodos basados en puntos y curvas de nivel mediante un proceso de triangulación TIN. La extensión 3D Analyst provee estas herramientas, incorporándose una extensión adicional para realizar interpolaciones estocásticas mediante kriging, del mismo modo que el IDRISI 32 bits e IDRISI Kilimanjaro. Es necesario señalar que existe una variada gama de programas que permiten efectuar operaciones de interpolación.

MÉTODOS ESTOCÁSTICOS: EL MÉTODO KRIGING

Cuando los datos son abundantes, muchas técnicas de interpolación llevan a similares resultados. Por otra parte, cuando se plantean supuestos para la elección de uno u otro método de interpolación, los resultados son obtenidos satisfactoriamente, y el usuario tendería a pensar que el análisis para los datos disponibles fue el correcto. Esto sin embargo podría llevar a resultados engañosos, sobre todo cuando no se adopta una análisis previo de los datos y las características reales del terreno antes de aplicar un método de interpolación.

Existen métodos geoestadísticos para la interpolación, en particular los adscritos al método "kriging", diseñado para optimizar la interpolación mediante la división del espacio de variación en tres componentes: una variación determinística (diferentes niveles de tendencia), una autocorrelación espacial, y un análisis del "ruido" en las estimaciones.

El proceso de variación espacial en las correlaciones está regulado por algoritmos denominados autocovariogramas y semivariogramas. Ambas herramientas proveen de información para una optimización en la interpolación y ajustar los radios de búsqueda.

Los métodos geoestadísticos proveen una gran flexibilidad en el proceso de interpolación, permitiendo el tratamiento de gran cantidad de información tanto en superficie como en volumen. Variantes del método han sido implementadas para la interpolación de datos binarios, la incorporación de tendencias, o estratificaciones. Todos estos métodos proveen superficies de variación acompañadas por una estimación en la varianza de los datos para la superficie. En contraste a otros métodos que operan sobre valores locales, los métodos de simulación condicional, permiten entregar un conjunto valioso de datos muy útiles para realizar análisis más acabados en un ambiente raster para aplicaciones en modelos ambientales.

La regionalización de la variable de interés

Los métodos geoestadísticos para la interpolación parten con un reconocimiento de la variación espacial del atributo para una variable continua, generalmente complicada para ser modelada con una función matemática simple. Para abordar este problema, es recomendable describir la variación mediante un proceso estocástico para la generación de una superficies. El atributo en este caso, es reconocido como una "variable regionalizada", aplicándose este término para distintas áreas, tales como la

variación en la presión atmosférica, el nivel de variación de las mareas, o la distribución de los indicadores demográficos.

La teoría de la regionalización asume que la variación espacial de cualquier variable puede ser expresada como la suma de tres grandes componentes: a) un componente estructural encargado de computar la media y la tendencia, b) un componente aleatorio, pero espacialmente correlacionado y conocido como la variación de la variable regionalizada, y c) una ruidó espacial no correlacionado, también conocido como error. Por lo cual la función estocástica puede expresarse como:

$$Z(x) = m(x) + \varepsilon'(x) + \varepsilon''$$

donde $m(x)$ es una función determinística que describe el componente estructural de Z como " x "; $\varepsilon'(x)$ es el término estocástico que varía localmente pero que depende de los errores residuales de $m(x)$. Por último, $\varepsilon''(x)$ representa el error espacial residual, espacialmente independiente y conocido también como ruido gaussiano, con media cero y varianza σ^2 .

El primer paso en el proceso de cálculo es decidir la función apropiada para $m(x)$. En el caso más simple, cuando no existe una tendencia, $m(x)$ es igual al valor promedio del área de búsqueda, mientras que las diferencias esperadas entre dos o más localizaciones - x y $x+h$ - separadas por una distancia vectorial h , será igual a cero:

$$E \{ Z(x) - Z(x+h) \} = 0.$$

donde $Z(x)$ y $Z(x+h)$ son los valores de una variable Z randomizada para localizaciones x , $x+h$. La varianza de las diferencias depende únicamente de las distancias entre sitios, de modo que:

$$E \{ [Z(x) - Z(x+h)]^2 \} = E \{ [\varepsilon'(x) - \varepsilon'(x+h)]^2 \} = 2\gamma(h)$$

donde $\gamma(h)$ es conocido como la semivarianza. Las dos condiciones, diferencia estacionaria y varianza de las diferencias, definen los requerimientos para la hipótesis intrínseca de regionalización de la variable. De este modo la ecuación kriging para $Z(x)$ puede representarse ahora como

$$Z(x) = m(x) + \gamma(h) + \varepsilon''$$

en donde $\gamma(h)$ representa ahora a ε' . Si las condiciones especificadas por hipótesis intrínsecas con las correctas, entonces la semivarianza puede ser estimada ahora por los datos muestrales de la forma

$$\hat{\gamma}(h) = \frac{1}{2n} \sum_{i=1}^n \{ z(x_i) - z(x_i + h) \}^2$$

donde n es el número de pares de puntos muestrales de observación para los valores del atributo z , separados por una distancia h . Una representación gráfica de $\hat{\gamma}(h)$, corresponde al variograma experimental, el cual provee de información útil para la optimización y determinación de patrones espaciales de distribución. Corresponde al primer modelo teórico y ajustado que representa a un variograma experimental.

Utilizando el variograma para un análisis espacial

La posibilidad de obtener un variograma matemáticamente inestable es cierta si los resultados requeridos para computar $\hat{\gamma}(h)$ son insuficientes. Una regla para evitar este problema, es la posibilidad de tomar entre 50 a 100 valores muestrales para lograr un variograma estable. Dependiendo de la clase de variación encontrada, algunas superficies requieren más o menos puntos en función de la variación irregular del espacio y al área de búsqueda. Un variograma suave, por ejemplo, puede ser obtenido mediante el incremento del área de búsqueda para cada punto.

El rango de variación de cada variograma provee valiosa información acerca del tamaño de la ventana de búsqueda que podría ser utilizada, dado que este parámetro puede ser perfectamente modificable en función de las estimaciones obtenidas.

Si la distancia desde un punto conocido hacia otras localizaciones con valores desconocidos excede el rango del área de búsqueda, entonces estas distancias pueden ser modificadas mediante anisotropía, a través de la modificación de la forma de la vecindad de búsqueda desde un círculo a una elipse.

El modelo original de covarianza simple, asume que la variación espacial de un atributo de interés puede ser dividida en tres componentes básicos; ellos son, a) un promedio general de la tendencia, b) un variograma, y c) un error residual. En algunas situaciones, particularmente si los datos son insuficientes, es posible y útil distinguir más de un componente del variograma; por ejemplo, cuando se presentan dos o más rangos de valores muy separados entre sí, produciendo un quiebre de la tendencia. En este caso, el variograma $\gamma_T(h)$ permite unir y ajustar estos sub-rangos en un único modelo de correogionalización, expresado como:

$$\gamma_T(h) = \gamma_1(h) + \gamma_2(h) + \gamma_3(h) + \dots$$

de modo que la ecuación $Z(x)$ anterior, puede expresarse nuevamente como:

$$Z(x) = m(x) + \gamma_1(h) + \gamma_2(h) + \gamma_3(h) + \varepsilon$$

Cada uno de los sub-variogramas es definido por estos parámetros.

La confiabilidad de los variogramas se basan en el supuesto que todos los datos son localizados en la misma cobertura y en el mismo dominio espacial. En este caso se computa el variograma como un modelo digital que determina los pesos necesarios para la interpolación dentro del área de búsqueda. Sin embargo, esto no siempre es conveniente hacerlo; en algunos casos distintas partes del área presentan importantes diferencias en la cobertura de interés: drenaje, salinidad de los suelos, pH; por lo cual no sería válido utilizar una estructura correlacional global.

Por ello es importante analizar muy bien la superficie a representar, de modo de aplicar el mejor modelo, que en muchos casos puede ser un conjunto de variogramas específicos en vez de un modelo global. El principal problema que podría presentarse es la falta de datos, o bien que los disponibles se utilicen inadecuadamente.

Utilizando el variograma para la interpolación: el kriging ordinario

Una forma adecuada para la estimación del valor Z en un modelo de pesos λ necesarios para una interpolación local, está dada por un análisis geoestadístico del dato en un análisis general. Aquí, el valor de $Z(x_0)$ está dado por:

$$\hat{Z}(x_0) = \sum_{i=1}^n \lambda_i * z(x_i), \text{ con } \sum_{i=1}^n \lambda_i = 1$$

Los pesos λ_i , han sido elegidos de modo tal que la estimación de $\hat{Z}(x_0)$ es totalmente imparcial, con una varianza σ_e^2 muy inferior a otras combinaciones lineales para los valores observados.

La varianza mínima de $[\hat{Z}(x_0) - Z(x_0)]$, la predicción del error, o la “varianza kriging”, es representada como:

$$\sigma_e^2 = \sum_{i=1}^n \lambda_i * \gamma(x_i, x_0) + \psi, \text{ obtenido cuando}$$

$$\sum_{i=1}^n \lambda_i * \gamma(x_i, x_j) + \psi = \gamma(x_j, x_0), \text{ para todo } j.$$

El valor numérico $\gamma(x_i, x_j)$ es la semivarianza de Z entre los puntos de muestra x_i y x_j ; $\gamma(x_i, x_0)$ es la semivarianza entre los puntos de muestra x_i y el punto no visitado x_0 . Ambas cantidades son obtenidas de un variograma ajustado. La cantidad ψ es un Multiplicador de Lagrange requerido para la minimización.

Este método es conocido como “kriging ordinario”; es un interpolador exacto, en el sentido que cuando las ecuaciones anteriores son utilizadas, los valores interpolados, o las mejores medias locales, serán coincidentes con los valores de los puntos muestrales.

Gráficamente, los valores serán interpolados dentro de una grilla regular, en donde el espaciamiento entre celdas es utilizado como un indicador de la fineza en el cálculo de valores interpolados. Desde aquí, puede efectuarse una conversión a un mapa de contornos. De forma similar, la estimación del error σ_e^2 , conocido como varianza del kriging, puede ser utilizada para determinar la precisión de los valores interpolados para el área de interés. A menudo, la varianza es representada también como desviación kriging estándar (o error kriging).

Kriging por bloques

Considerando los extensos y pequeños rangos de variación en las variables de muchos fenómenos naturales, tales como las propiedades químicas de los suelos, calidad del agua, modelos de distribución de contaminantes y otros, un kriging ordinario de puntos podría generar resultados compuestos por múltiples “hoyos” o “picos” en la superficie estimada, si es que no existe una consideración previa acerca del número y distribución de los puntos muestrales. Esto puede superarse mediante una

modificación en las ecuaciones kriging, específicamente para la obtención de un estimador del valor $Z(B)$, de la variable Z sobre un bloque de superficie B . Esto es útil cuando se desea estimar el promedio de los valores de Z por medio del trazado de un plano para el área de interés, o la interpolación de valores para una grilla de un tamaño específico de celda.

Los valores promedio para Z sobre un bloque B , estaría dado entonces por

$$Z(B) = \int_B \frac{z(x)dx}{\text{area}B}, \text{ estimado mediante la expresión}$$

$$\hat{Z}(B) = \sum_{i=1}^n \lambda_i * z(x_i), \text{ con } \sum_{i=1}^n \lambda_i = 1$$

La varianza mínima es ahora

$$\hat{\sigma}^2(B) = \sum_{i=1}^n \lambda_i * \bar{\gamma}^*(x_i, B) + \psi - \bar{\gamma}^*(B, B), \text{ obtenida de}$$
$$\sum_{i=1}^n \lambda_i * \gamma^*(x_i, x_j) + \psi = \bar{\gamma}^*(x_j, B), \text{ para todo } j$$

La estimación de varianza obtenidas por bloques kriging es usualmente muy inferior respecto al kriging por puntos. Cuando estas ecuaciones son utilizadas, los resultados alisan la superficie interpolada, eliminando las cavidades y picos resultantes de una interpolación kriging ordinaria a partir de puntos.

OTRAS FORMAS DE KRIGING

Kriging simple

El kriging simple es la predicción por medio de regresión lineal, bajo el supuesto estacionario de segundo orden con una media conocida. Sin embargo, este supuesto planteado se transforma frecuentemente en una severa restricción para la mayoría de los datos procedentes de variables involucradas en el análisis del ambiente físico.

Kriging no lineal

Un kriging logaritmico es una interpolación basada en la distribución logarítmica. Los datos son primeramente transformados a logaritmo naturales, o logaritmos de base-10, de tal modo que el variograma y la interpolación se basan en la transformación $y(p) = \ln z(p)$, generándose el estimador $\hat{y}(p)$, para $\ln z(p)$. Los valores predecidos pueden ser transformados posteriormente mediante interpolación, pero con la precaución que la transformación inversa $\exp \Phi^*(p)$ es un estimador parcial de $z(p)$. N estimador imparcial propuesto sería el siguiente:

$$Z^*(p) = \exp [\Phi^*(p) + (\sigma_{sk}^2(p) / 2)]$$

Donde $\sigma_{sk}^2(p)$ es la varianza logaritmica del kriging. Una dificultad de utilizar esta expresión es la extrema sensibilidad en la transformación antilogaritmica, dificultando su aplicación e interpretación de los resultados. Se recomienda en este caso, utilizar un

kriging multigaussiano como una alternativa. Para el caso de este manual, no nos detendremos a analizar en profundidad este sub-método.

Kriging estratificado

Cuando se dispone de información suficiente para efectuar una clasificación de áreas en subáreas plenamente representativas de la realidad, y estos datos también son suficientes para computar variogramas para cada área de interés, entonces la interpolación podría ser tratada por puntos o bloques separados, ajustando las ecuaciones de kriging para evitar posibles discontinuidades en las vecindades de las distintas áreas tratadas.

En general, un estratificación de la desviación estándar reduce los errores en la interpolación, reduciendo en la mayoría de los casos la dependencia matemática en la distribución de los puntos muestrales. En este caso, la estratificación también preserva la estructura espacial de las áreas más pequeñas, por lo que la transición en la continuidad matemática de valores queda asegurada. Se soluciona también el problema de llenado de información en sectores donde el interpolador podría generar cavidades o picos en la superficie estimada.

Co-kriging

Frecuentemente los datos pueden estar disponibles para la mayoría de las variables de interés, aunque en distintas magnitudes y distribuciones en el espacio. Un set de datos (U) por ejemplo, podría ser muy dificultoso y caro de medir, optándose entonces por medir un conjunto alternativo de datos (V), el cual en la práctica resultaría más fácil y barato que el conjunto U. Si U y V se encuentran espacialmente correlacionados, entonces sería posible utilizar la información relacionada a la variación espacial de V, para complementar el análisis de datos obtenidos del conjunto U. Esto puede ser factible de realizar mediante un análisis conocido como Co-kriging. El método permite analizar los dos conjuntos puntuales de datos, a través de un variograma cruzado y asumiendo que ambos conjuntos de datos pertenecen al mismo dominio espacial.

Kriging probabilístico

Considerando las eventuales grandes variaciones existentes en los resultados obtenidos a partir de estos distintos métodos kriging de interpolación, algunos usuarios podrían decidir que para sus propósitos, no sería tan relevante conocer buenas estimaciones de $z(x_0)$, sino más bien, que la probabilidad del valor de un atributo en cuestión, no exceda un límite catalogado como umbral de tolerancia matemática. Desde este punto de vista, interesaría más la probabilidad de obtener un valor dentro de ciertos rangos, más que la estimación del dato para la variable de interés. Esto puede ser factible de realizar mediante indicadores kriging, en cuyo caso el conjunto de datos original es transformado desde una escala continua a una escala binaria (1 si $z \leq P_j$, y 0 en otro caso, o viceversa). Variogramas son computados para datos binarios, procediendo el kriging ordinario con la transformación de estos datos. Los mapas resultantes de este método muestran datos continuos en el rango 0-1, indicando la probabilidad que P_j haya excedido o no el rango para la variable de interés.

Simulación

Para muchas aplicaciones SIG, el usuario sólo requiere interpolar datos, generalmente puntuales, de modo de generar resultados, visualizarlos en pantalla y eventualmente

combinarlos con otros para perfeccionar los resultados. De este modo, el uso más común que se hace de los SIG es generar datos para modelos cuantitativos que representen procesos del medio ambiente, tales como cambio climático, contaminación del agua y aire, patrones de distribución de las plantas, animales y el hombre, entre los temas más globales. Para muchos modelos, sólo las mejores predicciones obtenidas a nivel de celda es necesaria. Para otros investigadores es necesario obtener mayores antecedentes acerca de cómo se comporta un cierto modelo de interés, frente a variaciones en los rangos de la variable de interés. Dado que en la práctica no es posible conocer en forma directa todos los valores de las celdas, es que resulta necesario realizar un procedimiento de simulación estocástica.

Con el propósito de examinar cómo un modelo numérico reacciona frente a importantes cambios en los valores de entrada, o tan pequeños que sean inferiores al promedio local, se debe aplicar una metodología estocástica que permita simular el valor para cada celda en forma individual. El método de Monte Carlo permite evaluar el valor probabilístico de cada celda mediante la igualdad $z(x) = \Pr(Z)$, donde Z es la variable espacial independiente y normalmente distribuida con media μ y varianza σ^2 . En este tratamiento, cada celda es espacialmente independiente de la vecina, por lo cual los errores en las estimaciones también son independientes entre sí.

Si el variograma es conocido, entonces diversos métodos tales como el vecino más cercano, o las bandas de cambio, pueden ser utilizados para simular una superficie que posea características espacialmente similares respecto a la superficie original. La simulación condicional combina los datos de observaciones puntuales con información procedente del variograma para obtener los probables mejores resultados por celda, y en función del variograma original.

El método requiere más esfuerzo computacional en los cálculos debido a las múltiples iteraciones que es necesario realizar para lograr una estabilización probabilística de la imagen resultado, respecto al kriging o co-kriging. Puede ser utilizado sin problemas combinándolo con información de bajo rango de variación para obtener mejores estimaciones.

Resumen de aspectos generales a considerar en el kriging

Como corolario, se menciona un listado de aspectos necesarios a considerar si se desea trabajar con la metodología kriging:

- ❑ Primero, es necesario examinar los datos en cuanto a la distribución estadística de ellos, sus variaciones y tendencias espaciales, de modo de elegir el modelo de transformación adecuado. Si se utiliza un indicador kriging, entonces deberá transformarse los datos a valores binarios (0-1).
- ❑ Calcular el variograma experimental y ajustarlo de acuerdo a la variación espacial. Si los datos no se encuentran autocorrelacionados, entonces la interpolación no será sensible a cambios en los valores de la variable de interés.
- ❑ Chequear el modelo mediante validación cruzada.
- ❑ Seleccionar un método kriging o simulación condicional.
- ❑ Si se selecciona kriging, entonces utilizar el variograma para interpolar sobre una grilla regular los valores mediante un kriging por puntos o kriging por bloques (si se desea operar con áreas más grandes). Si se opta por simulación condicional, entonces deberá realizarse no menos de 100 iteraciones sobre la grilla regular. En cada caso, se debe computar el promedio y desviación estándar de las superficies generadas.

- ❑ Representar los resultados en una grilla regular, o bien utilizando una estructura de contornos (no para el caso de la simulación condicional).
- ❑ Integrar estos resultados en el SIG y complementar el análisis con otros niveles de información.

CUADRO VII-2. Análisis comparativo para distintos métodos de interpolación

METODOS			COMPARACION	
IDRISI	SURFER	ARCVIEW	VENTAJAS	DESVENTAJAS
INTERPOL	Inverse Distance	MakeIDW	Aplicable a cualquier vector de puntos, y de manera sencilla	Dada la imposibilidad de tratar en detalle las discontinuidades topográficas, el método no es recomendable para tratar con altitudes
INTERCON			Opera bien para vectores lineales (como las curvas de nivel), por lo que se aplica para la realización de DEMs	No funciona bien para cambios bruscos del relieve, y para curvas de nivel muy distanciadas entre sí; es mejor emplear un TIN
TIN	Triangulation with Linear Interpolation	3D Analyst	El proceso de triangulación es bastante rápido y conservando las líneas críticas del relieve	
TINSURF		Convert to Grid	La estructura de datos es bastante simple, y de procesamiento liviano	Ocupa mucho espacio en memoria
KRIGING	Kriging	MakeKriging	Es un procedimiento muy flexible	No es apropiado cuando existen picos de valores (o cumbres pronunciadas para el caso de la topografía) o extensos espacios sin cambios de valores
THIESSEN			Las áreas de influencia son mucho más precisas cuando se cuenta con un mayor número de puntos	Cuando no hay suficientes puntos muestrales las áreas de influencia son mayores, lo que genera una interpolación adicional fuera de los puntos
TREND	Polynomial Regression	Make Trend	Se obtienen tres tipos de tendencias: lineal, cuadrática y cúbica.	Superficie con muchos errores cuando los puntos de muestra son escasos

CUADRO VII-3. UNA COMPARACIÓN DE DISTINTOS MÉTODOS DE INTERPOLACIÓN

METODO	DETERMINÍSTICO / ESTOCÁSTICO	VECINDAD LOCAL / GLOBAL	TRANSICIÓN ABRUPTA / GRADUAL	INTERPOLADOR EXACTO	LIMITANTES DEL MÉTODO	RECOMENDABLE PARA	PROCESAMIENTO COMPUTACIONAL	ESTRUCTURA DE SALIDA DE LOS DATOS	SUPUESTOS PLANTEADOS AL MODELO
Clasificación	Determinístico	Global	Abrupta si es que se usa sin otro método complementario	No	La delimitación de áreas y clases puede ser subjetiva, principalmente por efectos de la desviación estándar.	Para estimaciones rápidas cuando los datos son muy dispersos. Para remover diferencias sistemáticas antes de efectuar una interpolación a partir de datos puntuales	Liviano	Polígonos clasificados	Existe homogeneidad en los datos con los bordes
Superficies de Tendencia	Esencialmente determinístico (modelos empíricos)	Global	Gradual	No	El efecto físico de la tendencia no siempre es claro. Características de los bordes de polígonos pueden distorsionar el aspecto general de la superficie y puede ocasionar interpretaciones erróneas.	Para cálculos rápidos de tendencias generales o gruesas.	Liviano	Grilla continua representada en una superficie	Datos normalmente distribuidos
Modelos de Regresión	Esencialmente determinístico (modelos empíricos y estadísticos)	Global con ajustes locales	Gradual si los datos de entrada presentan variación gradual	No	Los resultados dependen fuertemente del modelo de regresión y de la calidad de los datos de entrada.	Como buena alternativa de modelación cuando otros modelos más apropiados no están disponibles o no hay presupuesto para utilizarlos	Liviano	Polígonos o grilla continua representada en una superficie	Modelo extrapolable a toda la superficie
Polígonos de Thiessen	Determinístico	Local	Abrupta	Sí	Sólo actúa un solo punto por polígono. El modelo de teselas depende de la distribución de los datos	Trabajar con datos de observaciones puntuales y nominales	Liviano	Polígonos o grilla	Excelente predicción mientras más cercanos sean los puntos

CUADRO VII-3. UNA COMPARACIÓN DE DISTINTOS MÉTODOS DE INTERPOLACIÓN (cont.)

METODO	DETERMINÍSTICO / ESTOCÁSTICO	VECINDAD LOCAL / GLOBAL	TRANSICIÓN ABRUPTA / GRADUAL	INTERPOLADOR EXACTO	LIMITANTES DEL MÉTODO	RECOMENDABLE PARA	PROCESAMIENTO COMPUTACIONAL	ESTRUCTURA DE SALIDA DE LOS DATOS	SUPUESTOS PLANTEADOS AL MODELO
Método de Tobler	Determinístico	Local	Gradual	No	Los datos de entrada corresponden a cuentas (frecuencia) o densidades	Pasar de modelos con resultados de distribución espacial discreta a superficies continuas	Liviano a moderado	Grilla continua representada en una superficie, o contornos rasters	Continuidad espacial y variación gradual de los datos dentro de la superficie
Interpolación lineal	Determinístico	Local	Gradual	Sí	No presenta limitantes	Para interpolar a partir de puntos cuando la densidad de datos es alta y para convertir una grilla de un sistema de proyección a otro	Liviano	Grilla continua representada en una superficie	Soporta una gran densidad de datos de entrada y salida; muchos más que una aproximación lineal
IDW	Determinístico	Local	Gradual	No con superficies lisas, pero puede ser forzado	Los resultados dependen del tamaño de la ventana de búsqueda y la elección de los pesos para los parámetros.	Una interpolación rápida de la dispersión de datos puntuales en una grilla regular o irregularmente espaciada	Liviano	Grilla continua representada en una superficie, o contornos rasters	Datos equitativamente distribuidos para evitar el efecto cluster
Spline	Determinístico con componentes locales estocásticos	Local	Gradual	Sí, con límites graduales	Es posible obtener muy buenos resultados	Una interpolación rápida (univariada o multivariada) de datos digitales de elevación y atributos relacionados con la creación de DEMs con un moderado nivel de detalle	Liviano	Grilla continua representada en una superficie, o contornos rasters	Toda la superficie sobre la cual opera el modelo presenta suaves variaciones

CUADRO VII-3. UNA COMPARACIÓN DE DISTINTOS MÉTODOS DE INTERPOLACIÓN (cont.)

METODO	DETERMINÍSTICO / ESTOCÁSTICO	VECINDAD LOCAL / GLOBAL	TRANSICIÓN ABRUPTA / GRADUAL	INTERPOLADOR EXACTO	LIMITANTES DEL MÉTODO	RECOMENDABLE PARA	PROCESAMIENTO COMPUTACIONAL	ESTRUCTURA DE SALIDA DE LOS DATOS	SUPUESTOS PLANTEADOS AL MODELO
Kriging	Estocástico	Local con variogramas globales Local con variogramas locales cuando se estratifica Local con tendencias globales	Gradual	Sí	El error en las estimaciones depende del variograma, de la distribución de los datos puntuales, y del tamaño de los bloques interpolados. Requiere cuidado el elegir qué modelo de correlación espacial se querrá obtener	Cuando los datos son suficientes para computar variogramas y kriging como un método para tratar con datos dispersos. Datos binarios y nominales pueden ser interpolados con indicadores de kriging. Información suave y gradual podría ser también incorporada como tendencias o estratificaciones. Datos multivariados también pueden ser interpolados mediante co-kriging.	Moderado	Grilla continua representada en una superficie	La superficie interpolada es lisa Existen hipótesis que operan bajo estadísticas estacionarias
Simulación Condicional	Estocástico	Local con variogramas globales Local con variogramas locales cuando se estratifica Local con tendencias globales	Irregular	No		Se desea obtener excelentes estimaciones sobre el rango posible de valores para un atributo o localizaciones no muestrales. En este caso se hace necesario un análisis Monte Carlo u otro tipo de modelo numérico	Moderado a fuerte	Grilla continua representada en una superficie	Existen hipótesis que operan bajo estadísticas estacionarias