

Notas clase 7/8 MDS7203 modelos discriminativos vs generativos

SEAN $C_1, C_2, \dots, C_K$ CLASES. CONSIDEREMOS QUE NUESTRO CONJUNTO DE D DATOS ESTÁ DADO POR

$$D = \{(x_i, t_i)\}_{i=1}^N \subset \mathbb{R}^M \times \{C_i\}_{i=1}^K$$

CONSIDEREMOS EL CASO EN EL QUE QUEREMOS MODELAR LA PROBABILIDAD CONDICIONAL SIGUIENTE:

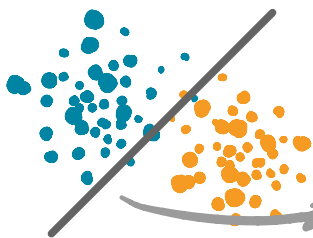
$$P(C_K | x)$$

AL PLANTEARLO DE ESTE MODO PODEMOS COMPUTAR EL ESTIMADOR MÁXIMO A POSTERIORI.

ACÁ QUEREMOS MAXIMIZAR LA PROBABILIDAD DE QUE LA ETIQUETA ASIGNADA A LOS DATOS SEA AQUELLA QUE CORRESPONDE. ESTO SE PUEDE ESCRIBIR, PARA EL CASO DE CLASIFICACIÓN BINARIA Y UN PUNTO DE DATO (x_i, t_i) :

$$\begin{aligned} L_i(\theta) = p(x_i, t_i | \theta) &= \begin{cases} p(x_i, C_1 | \theta) & \text{si } t_i = C_1 \\ p(x_i, C_2 | \theta) & \text{si } t_i = C_2 \end{cases} \rightarrow \text{ASUMAMOS QUE } C_1 = 1 \text{ Y } C_2 = 0 \\ &= p(x_i, C_1 | \theta)^{t_i} p(x_i, C_2 | \theta)^{1-t_i} \end{aligned}$$

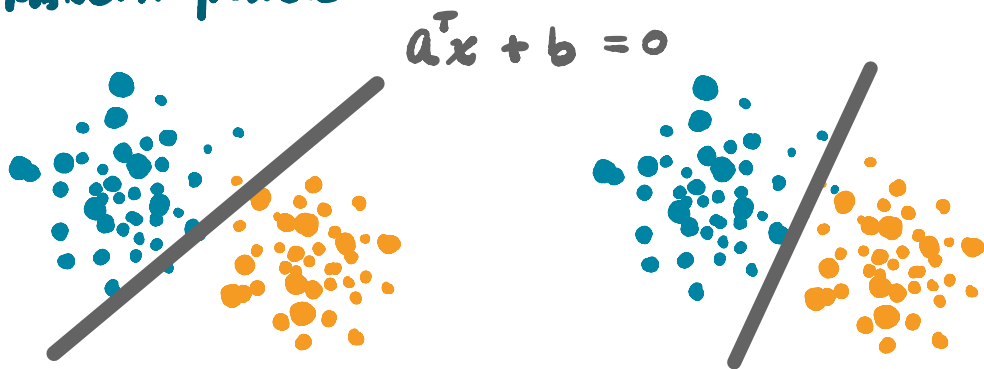
CONSIDEREMOS POR SIMPLICIDAD EL CASO DE CLASIFICACIÓN BINARIA $K=2$.



IMAGINEMOS QUE NUESTROS CONJUNTOS DE DATOS DISTRIBUYEN DE ACUERDO A UNA DISTRIBUCIÓN GAUSSIANA, (CON LA MISMA COVARIANZA)

→ BORDE DE DECISIÓN

TENEMOS LA INTUICIÓN DE ENCONTRAR UN
HIPERPLANO QUE SEPARA LOS DATOS DE LA MEJOR
MANERA POSIBLE:

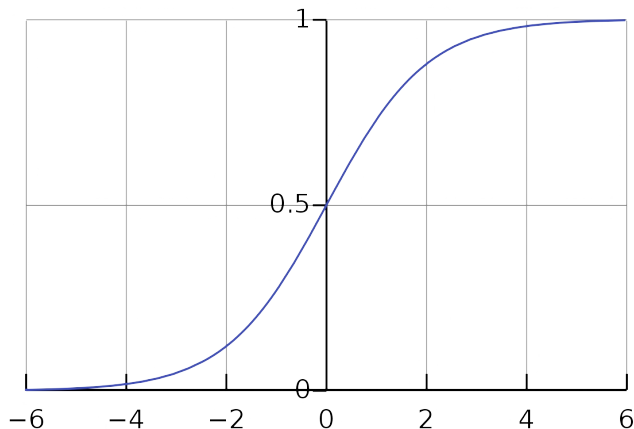


GEOMÉTRICAMENTE SABEMOS QUE DADOS a Y b , NUESTRA
RECTA ESTÁ DETERMINADA POR $a^T x + b = 0$, LUEGO
LOS CASOS $a^T x + b > 0$ Y $a^T x + b < 0$ DETERMINAN
A QUE LADO DEL BORDE DE DECISIÓN QUEDA EL PUNTO.

LA LEJANÍA A LA RECTA NOS DA LA INTUICIÓN DE MÁS
CERTIDUMBRE RESPECTO A LA CLASIFICACIÓN, PERO
DEBEMOS PASAR ESTO A PROBABILIDADES...

RECORDAMOS LA FUNCIÓN LOGÍSTICA:

$$\sigma: \mathbb{R} \rightarrow (0, 1)$$
$$r \mapsto \frac{1}{1 + e^{-r}}$$



ES UN Mapeo CONTINUO
DE LOS REALES AL
DOMINIO $(0, 1)$!!!

Esta función tiene además propiedades convenientes:

- Refleja: $\sigma(-r) = 1 - \sigma(r)$
- Diferenciable, más aún: $\frac{d}{dr} \sigma(r) = \sigma(r)(1 - \sigma(r))$
- Con inversa: $r(\sigma) = \log\left(\frac{\sigma}{1-\sigma}\right)$

Esto nos supone usar:

$$P(C_k | x) = \frac{1}{1 + e^{-a^T x - b}}$$

A esto lo llamamos **regresión logística**.

Llegado este punto, podemos optimizar directamente los parámetros a y b . Consideremos el siguiente abuso de notación:

$$w = \begin{pmatrix} a \\ b \end{pmatrix}, \quad \tilde{x} = \begin{pmatrix} x \\ 1 \end{pmatrix}$$

Recordando que la verosimilitud está dada por

$$\hat{\theta}_{\text{ML}} = \underset{\theta}{\text{argmáx}} \prod_{i=1}^N L_i(\theta) = \underset{\theta}{\text{argmáx}} \prod_{i=1}^N p(C_1 | x_i)^{t_i} p(C_2 | x_i)^{1-t_i}$$

Recordemos que $P(C_k | x_i) = \frac{1}{1 + e^{-w^T x_i}} = \sigma(w^T x_i)$

Denotando $\sigma_i = \sigma(w^T x_i)$ queda:

$$\hat{\theta}_{\text{ML}} = \underset{w}{\text{argmáx}} \prod_{i=1}^N \sigma_i^{t_i} (1 - \sigma_i)^{1-t_i}$$

Entonces derivando la log verosimilitud nos da:

$$\begin{aligned} \frac{\partial \log L}{\partial w} &= \frac{\partial}{\partial w} \left(\sum_{i=1}^N (t_i \log \sigma_i + (1-t_i) \log (1-\sigma_i)) \right) \\ &= \sum_{i=1}^N t_i \frac{\partial \log \sigma_i}{\partial w} + (1-t_i) \frac{\partial \log (1-\sigma_i)}{\partial w} \end{aligned}$$

Recordando que $\frac{d}{dr} \sigma(r) = \sigma(r)(1 - \sigma(r))$

$$\begin{aligned}
 \frac{\partial \log \sigma_i}{\partial \omega} &= \frac{1}{\sigma_i} x_i (1 - \sigma_i) x_i \quad \text{y} \quad \frac{\partial \log(1 - \sigma_i)}{\partial \omega} = \frac{1}{1 - \sigma_i} \cdot -\sigma_i (1 - \sigma_i) \cdot x_i \\
 &= \sum_{i=1}^N \underbrace{t_i (1 - \sigma_i) x_i - (1 - t_i) \sigma_i x_i}_{t_i x_i - t_i \sigma_i x_i - \sigma_i x_i + t_i \sigma_i x_i = (t_i - \sigma_i) x_i} \\
 &= \sum_{i=1}^N (t_i - \sigma_i) x_i
 \end{aligned}$$

lamentablemente no es llevar ϵ igual a cero en este caso, pero podemos usar el método de Gradiente Estocástico:

$$\theta \mapsto \theta - \eta (t_i - \sigma_i) x_i$$

Se dice que regresión logística es un modelo discriminativo porque aprende las probabilidades condicionales

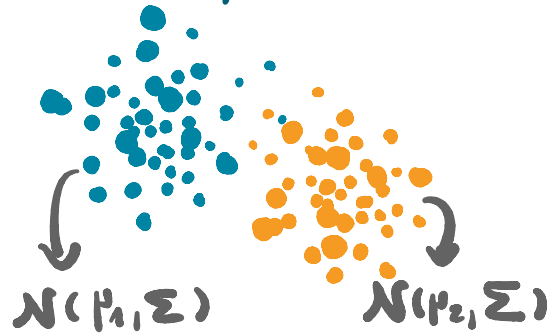
$p(C_k | x)$ directamente, a diferencia del approach GENERATIVO donde modelaremos $p(x, C_k | \theta)$ y aplicaremos Bayes para obtener un estimador.

Volvamos al caso inicial. Asumamos que nuestros datos son una mezcla GAUSSIANA, UNA GAUSSIANA para cada clase, con (priors): $p(C_1) = \alpha$, $p(C_2) = 1 - \alpha$

$$\text{y } p(x | C_k) = \mathcal{N}(x; \mu_k, \Sigma)$$

de este modo:

$$\begin{aligned}
 p(x, t) &= p(t = C_1) p(x | C_1) \\
 &\quad + p(t = C_2) p(x | C_2)
 \end{aligned}$$



Con esto nuestros parámetros a aprender son: $\mu_{1,2}, \Sigma$ y α

Recordando la densidad de una GAUSSIANA tenemos

$$p(x | C_k) = \frac{1}{(2\pi)^{\frac{N}{2}} |\Sigma|^{\frac{1}{2}}} \exp\left(-\frac{1}{2}(x - \mu_k) \Sigma^{-1} (x - \mu_k)^T\right)$$

y además,

$$\log(p(x|C_k)) = \frac{1}{2} \left[\log\left(\frac{1}{(2\pi)^{\frac{N}{2}} |\Sigma|^{\frac{1}{2}}}\right) - (x - \mu_k) \Sigma^{-1} (x - \mu_k)^T \right]$$

Junutando lo anterior tenemos:

$$\theta_{\text{MAP}} = \underset{\theta}{\text{argmax}} \times \prod_{i=1}^N L_i(\theta) = \underset{\theta}{\text{argmax}} \times \prod_{i=1}^N \underbrace{p(x_i, C_1 | \theta)}_{p(x_i | C_1) p(C_1)}^{t_i} \underbrace{p(x_i, C_2 | \theta)}_{p(x_i | C_2) p(C_2)}^{1-t_i}$$

$$\log() = \underset{\theta}{\text{argmax}} \times \sum_{i=1}^N t_i (\log(p(x_i | C_1)) + \log(p(C_1))) + (1-t_i) (\log(p(x_i | C_2)) + \log(p(C_2)))$$

$$\begin{aligned} &= \underset{\theta}{\text{argmax}} \times \sum_{i=1}^N t_i \log \alpha + (1-t_i) \log(1-\alpha) \\ &\quad + t_i \frac{1}{2} \left[\log\left(\frac{1}{(2\pi)^{\frac{N}{2}} |\Sigma|^{\frac{1}{2}}}\right) - (x - \mu_1) \Sigma^{-1} (x - \mu_1)^T \right] \\ &\quad + (1-t_i) \frac{1}{2} \left[\log\left(\frac{1}{(2\pi)^{\frac{N}{2}} |\Sigma|^{\frac{1}{2}}}\right) - (x - \mu_2) \Sigma^{-1} (x - \mu_2)^T \right] \end{aligned}$$

Tomamos el gradiente e igualamos a cero.

Ojo: esto debemos hacerlo para cada componente (tomamos en cada caso sólo los términos relevantes)

$$\bullet \frac{\partial \log L}{\partial \alpha} = \frac{\partial}{\partial \alpha} \sum_{i=1}^N t_i \log \alpha + (1-t_i) \log(1-\alpha)$$

$$= \sum_{i=1}^N t_i \cdot \frac{1}{\alpha} - (1-t_i) \frac{1}{1-\alpha}$$

$$\text{Imponiendo que } \frac{\partial L}{\partial \theta} = 0 \text{ queda } \frac{1}{\alpha} \sum_{i=1}^N t_i = \frac{1}{1-\alpha} \sum_{i=1}^N (1-t_i)$$

$$\Leftrightarrow (1-\alpha) \sum_{i=1}^N t_i = \alpha \sum_{i=1}^N (1-t_i)$$

$$\Leftrightarrow \underbrace{\sum_{i=1}^N t_i}_{N_1} - \cancel{\alpha \sum_{i=1}^N t_i} = \alpha N - \cancel{\alpha \sum_{i=1}^N t_i}$$

N_1 = cantidad de datos con etiqueta 1

$$\Leftrightarrow \alpha = \frac{N_1}{N} \text{ aka, proporción de datos } C_1$$

Esto justifica la intuición de que nuestro prior está determinado por la frecuencia relativa por clase.

$$\begin{aligned} \bullet \frac{\partial \text{Log}}{\partial \mu_1} &= \frac{\partial}{\partial \mu_1} \left(-\frac{1}{2} \sum_{i=1}^N t_i (x_i - \mu_1) \Sigma^{-1} (x_i - \mu_1)^T \right) \\ &= \sum_{i=1}^N t_i (\Sigma^{-1} (x_i - \mu_1)) = \Sigma^{-1} \sum_{i=1}^N (t_i x_i - t_i \mu_1) \end{aligned}$$

imponiendo que vuelva cero queda:

$$\underbrace{\sum_{i=1}^N t_i x_i}_{\sum_{x_i \in C_1} x_i} = \mu_1 \underbrace{\sum_{i=1}^N t_i}_{N_1} \Leftrightarrow \mu_1 = \frac{1}{N_1} \sum_{x_i \in C_1} x_i \quad \left. \vphantom{\sum_{i=1}^N} \right\} \text{La media empírica!}$$

Análogamente obtenemos $\mu_2 = \frac{1}{N_2} \sum_{x_i \in C_2} x_i$.

Cuando modelamos la distribución de los datos, i.e., damos una expresión para la probabilidad de un punto de dato dado, llamamos a esto un **Modelo generativo**.

La ventaja de lo que hicimos es que ahora conocemos la estructura probabilística de los datos, gracias a lo cual podemos generar nuevos puntos de datos con métodos como Box-Muller. Muchas bibliotecas python como numpy, scipy y pytorch tienen incorporados métodos para samplear Gaussians.

Para cerrar con las ventajas de modelar nuestros datos de manera generativa, vemos que podemos deducir una forma cerrada de un clasificador con la misma forma que lo que tenemos en regresión logística:

$$P(C_k | x) = \frac{1}{1 + e^{-a^T x - b}}$$

En vez de usar el método de gradiente estocástico, encontremos una expresión buena para a y b recordando el teorema de Bayes:

$$\begin{aligned} P(C_k | x) &= \frac{P(x | C_k) P(C_k)}{P(x)} \\ &= \frac{P(x | C_k) P(C_k)}{\sum_{j=1}^K P(x | C_j) P(C_j)} \end{aligned}$$

para el caso de clasificación binaria queda

$$\begin{aligned} P(C_1 | x) &= \frac{P(x | C_1) P(C_1)}{P(x | C_1) P(C_1) + P(x | C_2) P(C_2)} \\ &= \frac{1}{1 + \frac{P(x | C_2) P(C_2)}{P(x | C_1) P(C_1)}} \end{aligned}$$

Esta última expresión nos sugiere plantear lo siguiente:

$$\frac{P(x | C_2) P(C_2)}{P(x | C_1) P(C_1)} = \exp(-a^T x - b)$$

pues

$$P(C_k | x) = \frac{1}{1 + \underbrace{e^{-a^T x - b}}} = \frac{1}{1 + \underbrace{\frac{P(x | C_2) P(C_2)}{P(x | C_1) P(C_1)}}}$$

Resolvamos $\frac{P(x | C_2) P(C_2)}{P(x | C_1) P(C_1)} = \exp(-a^T x + b)$

Primero asumimos distribuciones Gaussianas para $P(x | C_k)$, con matriz de covarianza común:

$$P(x | C_k) = \frac{1}{(2\pi)^{\frac{N}{2}} |\Sigma|^{\frac{1}{2}}} \exp\left(-\frac{1}{2}(x - \mu_k)^T \Sigma^{-1} (x - \mu_k)\right)$$

La gracia de asumir covarianza común es que se simplifica la expresión:

$$\frac{P(x|C_2)P(C_2)}{P(x|C_1)P(C_1)} = \frac{\cancel{(2\pi)^{-\frac{n}{2}}|\Sigma|^{\frac{1}{2}}} \exp(-\frac{1}{2}(x-\mu_2)\Sigma^{-1}(x-\mu_2)^T) P(C_2)}{\cancel{(2\pi)^{-\frac{n}{2}}|\Sigma|^{\frac{1}{2}}} \exp(-\frac{1}{2}(x-\mu_1)\Sigma^{-1}(x-\mu_1)^T) P(C_1)}$$

Luego tomando logaritmo tenemos:

$$\log\left(\frac{\exp(-\frac{1}{2}(x-\mu_2)\Sigma^{-1}(x-\mu_2)^T) P(C_2)}{\exp(-\frac{1}{2}(x-\mu_1)\Sigma^{-1}(x-\mu_1)^T) P(C_1)}\right) = -a^T x - b$$

$$\Leftrightarrow \underbrace{-\frac{1}{2}(x-\mu_2)\Sigma^{-1}(x-\mu_2)^T + \log\left(\frac{P(C_2)}{P(C_1)}\right) + \frac{1}{2}(x-\mu_1)\Sigma^{-1}(x-\mu_1)^T}_{= -a^T x - b}$$

$$\frac{1}{2}(x^T \Sigma^{-1}(\mu_1 - \mu_2) + (\mu_1 - \mu_2)^T \Sigma^{-1} x - \mu_1^T \Sigma^{-1} \mu_1 + \mu_2^T \Sigma^{-1} \mu_2 + \log\left(\frac{P(C_1)}{P(C_2)}\right))$$

$$\Leftrightarrow (\mu_1 - \mu_2)^T \Sigma^{-1} x + \frac{1}{2}(\mu_2^T \Sigma^{-1} \mu_2 - \mu_1^T \Sigma^{-1} \mu_1) + \log\left(\frac{P(C_1)}{P(C_2)}\right) = -a^T x - b$$

con esto, $-a = (\mu_1 - \mu_2)^T \Sigma^{-1}$

$$-b = \frac{1}{2}(\mu_2^T \Sigma^{-1} \mu_2 - \mu_1^T \Sigma^{-1} \mu_1) + \log\left(\frac{P(C_1)}{P(C_2)}\right)$$

concluimos con esto la forma cerrada de nuestro clasificador, dada la estructura generativa de los datos.

Observación: la generalización multiclase de regresión logística es la función softmax (o exponencial generalizada).

$$P(C_k|x) = \frac{\exp(s_i)}{\sum_{j=1}^K \exp(s_j)}, \text{ con } s_i = \log(P(x|C_i)P(C_i))$$