

Tema 2. Análise discriminante linear e regresión linear

- LDA: linear discriminant analysis, 2 clases
- Proxecta a un espazo $v = \mathbf{w}^T \mathbf{x}$ unidimensional no cal:

a) Maximízase a separación entre as medias das 2 clases (entre clases): $m_i = \mathbf{w}^T \mathbf{m}_i$, $i=1,2$

b) Minimízase a separación entre os patróns da mesma clase (intraclase)

- Medida de separación: covarianza no espazo proxectado

- Intraclase s_i^2 : $s_i^2 = \sum_{y_j=i} (v_j - m_i)^2 = \sum_{y_j=i} (\mathbf{w}^T \mathbf{x}_j - \mathbf{w}^T \mathbf{m}_i)^2$, $i=1,2$

$$\mathbf{m}_i = \frac{1}{N_i} \sum_{y_j=i} \mathbf{x}_j$$

$N_i = n^\circ$ patróns de clase i

Proxección do patrón

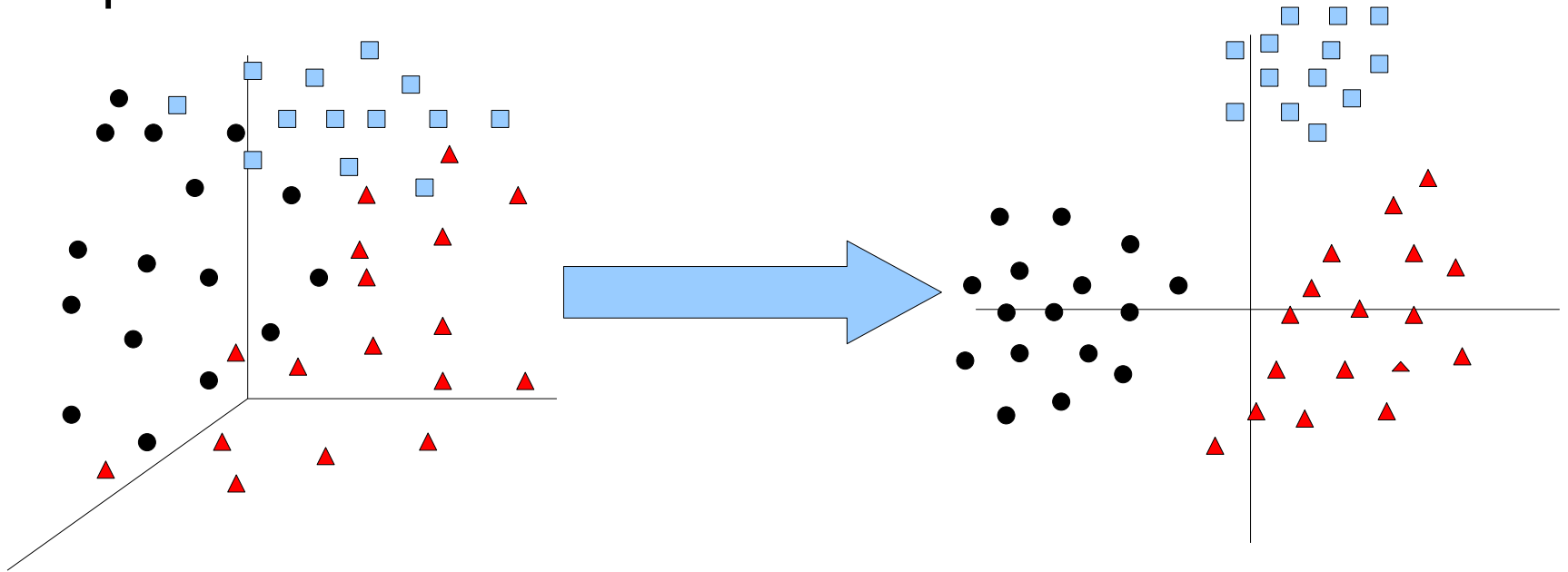
Proxección da media de clase

LDA para clasificación binaria (I)

- Covarianza entre clases en 1D: $(m_1 - m_2)^2$
 - Criterio de Fisher: $J(\mathbf{w}) = \frac{(m_1 - m_2)^2}{s_1^2 + s_2^2}$
 - Substituyendo m_i e s_i^2 vs. $\mathbf{w}$: $J(\mathbf{w}) = \frac{\mathbf{w}^T \mathbf{S}_B \mathbf{w}}{\mathbf{w}^T \mathbf{S}_W \mathbf{w}}$ Cociente de Rayleigh
 - $\mathbf{S}_B$: matriz de covarianza entre clases no espacio orixinal
- $$\mathbf{S}_B = (\mathbf{m}_2 - \mathbf{m}_1)(\mathbf{m}_2 - \mathbf{m}_1)^T \quad \mathbf{S}_W = \sum_{i=1}^2 \sum_{y_j=i} (\mathbf{x}_j - \mathbf{m}_i)(\mathbf{x}_j - \mathbf{m}_i)^T$$
- $\mathbf{S}_W$: suma de matrices de covarianza intracalse
- $\mathbf{S}_B$ e $\mathbf{S}_W$ cadradas de orde n

Exemplo de LDA

- Proxéctase a un espazo $C-1$ -dimensional onde se maximiza a covarianza entre clases e minimízase a suma de covarianzas intraclase
- Exemplo con $C=3$ clases de $\mathbb{R}^3$ a $\mathbb{R}^2$



- Maximiza a covarianza entre clases e minimiza a covarianza intraclase

LDA para clasificación binaria (II)

- Para maximizar $J(\mathbf{w})$, fago $\nabla J(\mathbf{w}) = dJ(\mathbf{w})/d\mathbf{w} = \mathbf{0}$ e obteño:

$$(\mathbf{w}^T \mathbf{S}_B \mathbf{w}) \mathbf{S}_W \mathbf{w} = (\mathbf{w}^T \mathbf{S}_W \mathbf{w}) \mathbf{S}_B \mathbf{w}$$

- Consideremos que $\mathbf{S}_b \mathbf{w} = (\mathbf{m}_2 - \mathbf{m}_1)(\mathbf{m}_2 - \mathbf{m}_1)^T \mathbf{w}$.
- Como $(\mathbf{m}_2 - \mathbf{m}_1)^T \mathbf{w}$ é un escalar, isto pódese escribir como $\lambda(\mathbf{m}_2 - \mathbf{m}_1)$, obtendo

$$(\mathbf{w}^T \mathbf{S}_B \mathbf{w}) \mathbf{S}_W \mathbf{w} = (\mathbf{w}^T \mathbf{S}_W \mathbf{w}) \lambda(\mathbf{m}_2 - \mathbf{m}_1)$$

- Suprimo os escalares $\mathbf{w}^T \mathbf{S}_B \mathbf{w}$ e $\mathbf{w}^T \mathbf{S}_W \mathbf{w}$ e o vector $\mathbf{w}$ é paralelo a $\mathbf{S}_W^{-1}(\mathbf{m}_2 - \mathbf{m}_1)$
- Como $|\mathbf{w}|$ non importa, úsase $\mathbf{w} = \mathbf{S}_W^{-1}(\mathbf{m}_2 - \mathbf{m}_1)$
- Se non existe $\mathbf{S}_W^{-1}$ regularízase: $(\mathbf{S}_W + \alpha \mathbf{1})^{-1}$, con $\alpha = 0^+$
- Clasificación: $y(\mathbf{x}) = \text{sign}(\mathbf{w}^T \mathbf{x} - b)$: b é o umbral $b = \frac{\mathbf{w}^T (\mathbf{m}_1 + \mathbf{m}_2)}{2}$

Discriminante linear de Fisher

- A miúdo úsase como sinónimo de LDA, pero hai lixeiras diferencias: o DL de Fisher non supón clases con distribucións normais nin iguais covarianzas de clase
- Se $\mathbf{m}_1$, $\mathbf{m}_2$ e $\mathbf{S}_1$, $\mathbf{S}_2$ son as medias e covarianzas das dúas clases, $\mathbf{w}^T \mathbf{x}$ ten media $\mathbf{w}^T \mathbf{m}_i$ e covarianza $\mathbf{w}^T \mathbf{S}_i \mathbf{w}$ para $i=1,2$
- Fisher definiu a separación S entre as dúas distribucións das clases como o cociente da varianza entre clases e a varianza intraclase:
$$S = \frac{\sigma_B^2}{\sigma_W^2} = \frac{(\mathbf{w}^T \mathbf{m}_2 - \mathbf{w}^T \mathbf{m}_1)^2}{\mathbf{w}^T (\mathbf{S}_1 + \mathbf{S}_2) \mathbf{w}}$$
- É fácil demostrar que a máxima separación S acádase cando:
$$\mathbf{w} \propto (\mathbf{S}_1 + \mathbf{S}_2)^{-1} (\mathbf{m}_2 - \mathbf{m}_1)$$
- Se as distribucións son parecidas, un offset b axeitado é:
$$b = \frac{\mathbf{w}^T (\mathbf{m}_1 + \mathbf{m}_2)}{2} = \frac{\mathbf{m}_1^T \mathbf{S}_1^{-1} \mathbf{m}_1 + \mathbf{m}_2^T \mathbf{S}_2^{-1} \mathbf{m}_2}{2}$$
- Se $\mathbf{S}_1 = \mathbf{S}_2$ e as distribucións son normais, acádase o LDA

LDA multiclase (I)

- Con $C > 2$ clases, LDA proyecta a espacio D -dimensional, con $D \leq C-1$: usualmente $D=C-1$
- Suponse que as C clases teñen medias $\mathbf{m}_i$, $i=1..C$, e a mesma matriz de covarianza media $\mathbf{S}_W$:
$$\mathbf{S}_W = \frac{1}{C} \sum_{i=1}^C \frac{1}{N_i} \sum_{y_j=i} (\mathbf{x}_j - \mathbf{m}_i)(\mathbf{x}_j - \mathbf{m}_i)^T$$
- A covarianza $\mathbf{S}_B$ entre clases é:
$$\mathbf{S}_B = \frac{1}{C} \sum_{i=1}^C (\mathbf{m}_i - \mathbf{m})(\mathbf{m}_i - \mathbf{m})^T$$
- A separación S entre clases na dirección de $\mathbf{w}$ está dada por:
$$\mathbf{m} = \frac{1}{N} \sum_{i=1}^N \mathbf{x}_i$$
$$S = \frac{\mathbf{w}^T \mathbf{S}_B \mathbf{w}}{\mathbf{w}^T \mathbf{S}_W \mathbf{w}}$$
- O vector $\mathbf{w}$ é un autovector de $\mathbf{S}_W^{-1} \mathbf{S}_B$, e a separación S é o seu autovalor asociado

LDA multiclase (II)

- O produto exterior de 2 vectores n-dimensionais ($\mathbf{x}\mathbf{x}^T$, sendo $\mathbf{x}$ un vector columna) é unha matriz $n \times n$ de rango 1
- O rango da matriz $\mathbf{S}_B$ é $\leq C-1$ porque $\mathbf{S}_B$ é a suma de C matrices que son produtos exteriores
- Como $\mathbf{S}_W^{-1}\mathbf{S}_B$ é diagonalizábel, a variabilidade das n entradas orixinais consérvase no subespazo de dimensión $C-1$ xerado polos $C-1$ meirandes autovalores de $\mathbf{S}_W^{-1}\mathbf{S}_B$
- No canto de por un vector $\mathbf{w}$ como na clasificación binaria, o LDA multiclase ven definido pola matriz $\mathbf{W}$, de orde $(C-1) \times n$, que proxecta $\mathbf{x}_i$ a $\mathbf{v}_i = \mathbf{W}\mathbf{x}_i \in \mathbb{R}^{C-1}$, cun vector $\mathbf{w}_i$, con $i=1..C-1$, en cada fila
- O LDA emprégase para clasificación e tamén para redución da dimensionalidade (de n a $D \leq C-1$)

LDA multiclase (III)

- No **programa** que veremos na clase **interactiva**:

- Matriz de covarianza de clase i : $\mathbf{S}_i = \frac{1}{N_i - 1} \sum_{y_j=i} (\mathbf{x}_j - \mathbf{m}_i)(\mathbf{x}_j - \mathbf{m}_i)^T$

- Matriz de covarianza total $\mathbf{S}_T$:

$$\mathbf{S}_T = \sum_{i=1}^C \frac{N_i - 1}{N - C} \mathbf{S}_i = \frac{1}{N - C} \sum_{i=1}^C \sum_{y_j=i} (\mathbf{x}_j - \mathbf{m}_i)(\mathbf{x}_j - \mathbf{m}_i)^T$$

- Offset b_i de clase i :

$$b_i = \log \frac{N_i}{N} - \frac{\mathbf{m}_i^T \mathbf{S}_T^{-1} \mathbf{m}_i}{2}, i = 1 \dots C$$

- Vector $\mathbf{w}_i$ de clase i , con $i = 1 \dots C$: $\mathbf{w}_i = \mathbf{S}_T^{-1} \mathbf{m}_i$

- Activación $L_i(\mathbf{x})$ de clase i para patrón de teste $\mathbf{x}$:

$$L_i(\mathbf{x}) = \mathbf{w}_i^T \mathbf{x} + b_i$$

- Clase predita $z(\mathbf{x})$: $z(\mathbf{x}) = \underset{i=1 \dots C}{\operatorname{argmax}} \left\{ \frac{e^{L_i(\mathbf{x})}}{\sum_{c=1}^C e^{L_c(\mathbf{x})}} \right\}$

Función **softmax**:
aproximación
suavizada
(derivábel) á
función máximo

Regresión lineal: mínimos cuadrados (least mean squares, LMS)

- Función lineal para predecir la salida: $z = \mathbf{v}^T \mathbf{u} + b = \mathbf{w}^T \mathbf{x}$

onde $\mathbf{w} = (\mathbf{v}, 1)$, $\mathbf{x} = (\mathbf{u}, 1)$ dados $\{\mathbf{x}_i, y_i\}_{i=1}^N$

$$\mathbf{y} = \begin{bmatrix} y_1 \\ \vdots \\ y_N \end{bmatrix} \quad \mathbf{X} = \begin{bmatrix} x_{11} & \cdots & x_{1n} \\ \vdots & \vdots & \vdots \\ x_{N1} & \cdots & x_{Nn} \end{bmatrix}$$

- Minimizar el error cuadrático:

$$E(\mathbf{w}) = \sum_{i=1}^N (y_i - \mathbf{w}^T \mathbf{x}_i)^2 = \|\mathbf{y} - \mathbf{X} \mathbf{w}\|^2 = (\mathbf{y} - \mathbf{X} \mathbf{w})^T (\mathbf{y} - \mathbf{X} \mathbf{w})$$

- En forma matricial:

$$E(\mathbf{w}) = \mathbf{y}^T \mathbf{y} - \mathbf{y}^T \mathbf{X} \mathbf{w} - \mathbf{w}^T \mathbf{X}^T \mathbf{y} + \mathbf{w}^T \mathbf{X}^T \mathbf{X} \mathbf{w} = \mathbf{y}^T \mathbf{y} - 2 \mathbf{w}^T \mathbf{X}^T \mathbf{y} + \mathbf{w}^T \mathbf{X}^T \mathbf{X} \mathbf{w}$$

- Derivando e igualando a $\mathbf{0}$:

$$\nabla E(\mathbf{w}) = \frac{\partial E(\mathbf{w})}{\partial \mathbf{w}} = -2 \mathbf{X}^T \mathbf{y} + 2 \mathbf{X}^T \mathbf{X} \mathbf{w} = \mathbf{0} \Rightarrow \mathbf{X}^T \mathbf{X} \mathbf{w} = \mathbf{X}^T \mathbf{y}$$

- La solución es: $\mathbf{w} = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{y} = \mathbf{X}^\dagger \mathbf{y}$

- La matriz $(\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T$ es la pseudoinversa de Moore-Penrose ($\mathbf{X}^\dagger$) de $\mathbf{X}$: función `pinv(x)` en Octave

- Función **lm** de paquete **stats** en R

Regresión lineal regularizada (ridge regression)

- **Mínimos cadrados regularizados** ou **ridge** (Tikhonov): para evitar sobreaprendizaxe, pondérase tamén $|\mathbf{w}|$, que mide a complexidade do modelo:

$$E(\mathbf{w}) = \sum_{i=1}^N (y_i - \mathbf{w}^T \mathbf{x}_i)^2 + \lambda |\mathbf{w}|^2 = |\mathbf{y} - \mathbf{X}\mathbf{w}|^2 + \lambda |\mathbf{w}|^2$$

- Orientada a problemas colineares (ill-posed) nos que a matriz $\mathbf{X}$ é singular, o cal é frecuente con datos anchos (n alto).
- Suma un valor positivo (parámetro de regularización, λ) á diagonal da matriz.
- Solución: $\mathbf{w} = (\lambda \mathbf{I}_n + \mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{y}$ ← $\mathbf{I}_n$ é a matriz identidade cadrada de orde n
- Función **lm.ridge** en paquete **MASS** de R
- **Regresión loxística**: $y = \mathbf{w}^T \Phi(\mathbf{x}) + b$, onde $\Phi(\mathbf{x}) = (\Phi_1, \dots, \Phi_m)$ e $\Phi_i(\mathbf{x}) = f((x_i - m_i)/\sigma_i)$, sendo $f(t) = 1/(1 + e^{-t})$ a función sigmoide

Implementacións

- As implementacións en Octave non usan ningúna función ou obxecto

	Python	R	Weka-Java
LDA	discriminant_analysis/ LinearDiscriminantAnalysis	MASS/lda	weka.classifiers.functions. LDA
Regresión linear	linear_model/LinearRegression	MASS/lm	weka.classifiers.functions. LinearRegression
Regresion regularizada	linear_model/Ridge	MASS/ lm.ridge	