

today: • Fisher LDA
• math tricks & MLE for Gaussian

generative model for classification: (Fisher) linear discriminant analysis

FLD (instead of LDA)

for classification $Y \in \{0,1\}$

$X \in \mathbb{R}^d$

generative approach: $p(x, y; \theta) = \overbrace{p(x|y; \theta)}^{\text{class conditional}} p(y; \theta)$

vs.

conditional approach $p(y|x; \theta)$

⊛ for Fisher model: we assume $p(x|y; \theta) = N(x | \mu_y, \Sigma)$

shared across classes

$\theta = (\mu_0, \mu_1, \Sigma, \pi)$
 μ_0 : mean for class 0
 μ_1 : mean for class 1
 Σ : shared covariance
 π : $p(y=1)$

as before (see exponential family argument)

can show that $p(y|x; \theta) = \sigma(w^T x)$ where w is a fct. of $(\mu_0, \mu_1, \Sigma, \pi)$

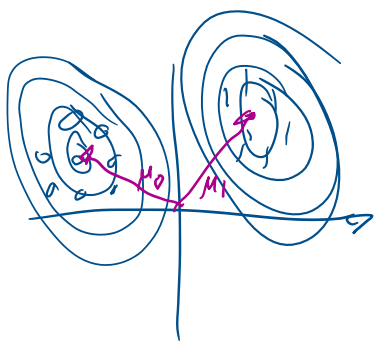
[note: if we use $\Sigma_0 \neq \Sigma_1$, get "quadratic discriminant analysis" (QDA)]

i.e. $\sigma(w^T \phi(x))$ where $\phi(x)$ is a quadratic fct. of x [see hwk. 2]

⊛ generative approach: do joint MLE to estimate

$$\hat{\theta} = \underset{\theta \in \Theta}{\text{argmax}} \sum_i \log p(x_i, y_i; \theta)$$

[vs. $\underset{w}{\text{argmax}} \sum_i \log p(y_i | x_i; w)$ for logistic regression]



sidenote: MLE for multivariate Gaussian

$X_i \stackrel{iid}{\sim} N(\mu, \Sigma)$

$\mu \in \mathbb{R}^d$

$$\Sigma \stackrel{\Delta}{=} \mathbb{E}[(x - \mu)(x - \mu)^T]$$

$\mu \in \mathbb{R}^d$
 $\Sigma \in \mathbb{R}^{d \times d}$
 Σ is symmetric
 $\Sigma \succeq 0$

$$\Sigma = \mathbb{E}[(x-\mu)(x-\mu)^T]$$

$$\Sigma^T = \Sigma$$

$$v^T \Sigma v = \mathbb{E}[v^T (x-\mu)(x-\mu)^T v] = \mathbb{E}[(x-\mu)^T v]^2 \geq 0 \Rightarrow \Sigma \succeq 0$$

$$p(x) = \frac{1}{\sqrt{(2\pi)^d |\Sigma|}} \exp\left(-\frac{1}{2} \underbrace{(x-\mu)^T \Sigma^{-1} (x-\mu)}_{\text{tr}((x-\mu)^T \Sigma^{-1} (x-\mu))}\right)$$

$$\text{tr}(AB) = \text{tr}(BA)$$

$$= \text{tr}(\Sigma^{-1} (x-\mu)(x-\mu)^T) = \langle \Sigma^{-1}, (x-\mu)(x-\mu)^T \rangle$$

$$\langle A, B \rangle \triangleq \sum_{i,j} A_{ij} B_{ij} = \text{tr}(A^T B)$$

log-likelihood: $\sum_{i=1}^n \log p(x_i; \Sigma) = \text{const.} - \frac{n}{2} \log |\Sigma| - \frac{1}{2} n \langle \Sigma^{-1}, \underbrace{\frac{1}{n} \sum_{i=1}^n (x_i - \mu)(x_i - \mu)^T}_{\triangleq \tilde{\Sigma}(\mu)} \rangle$

$$|\Sigma^{-1}| = |\Sigma|^{-1} = \frac{1}{|\Sigma|}$$

vector derivative review:

suppose $f: \mathbb{R}^m \rightarrow \mathbb{R}^n$

$h \in o(\|h\|)$ "little oh"
 $\Leftrightarrow \lim_{\|h\| \rightarrow 0} \frac{h(\|h\|)}{\|h\|} = 0$ e.g. $\|h\|^2 = o(\|h\|)$

f is differentiable at x_0 iff $\exists$ a linear operator $df_{x_0}: \mathbb{R}^m \rightarrow \mathbb{R}^n$
 s.t. $\forall \Delta \in \mathbb{R}^m \quad f(x_0 + \Delta) - f(x_0) = df_{x_0}(\Delta) + o(\|\Delta\|)$

"differential"

"derivative"

$$\lim_{\|\Delta\| \rightarrow 0} \frac{f(x_0 + \Delta) - f(x_0)}{\|\Delta\|} = \lim_{\|\Delta\| \rightarrow 0} \left(\underbrace{df_{x_0}(\Delta)}_{\frac{df_{x_0}(\Delta)}{\|\Delta\|}} + \underbrace{o(\|\Delta\|)}_{\frac{o(\|\Delta\|)}{\|\Delta\|}} \right)$$

$$\underbrace{\frac{df_{x_0}(\Delta)}{\|\Delta\|}}_{\frac{df_{x_0}(\frac{\Delta}{\|\Delta\|})}{1}} = df_{x_0}(\vec{d})$$

$\rightarrow$ directional derivative at x_0 in direction $\vec{d}$

if $n=1$;

$$df_{x_0}(\vec{d}) = \langle \nabla f(x_0), \vec{d} \rangle$$

$$\nabla f(x_0) = (df_{x_0})^T$$

df_{x_0} is linear

means that $df_{x_0}(\Delta + b\Delta_0)$

$$= df_{x_0}(\Delta) + b df_{x_0}(\Delta_0)$$

can represent as a $n \times m$ matrix called the Jacobian matrix

standard representation $(df_{x_0})_{ij} = \frac{\partial f_i}{\partial x_j}$

then $df_{x_0}(\Delta) = df_{x_0} \cdot \Delta$

1) this gives a way to get df_{x_0} for "anything" (matrix, tensor, k -dim fct., etc.)

2) ...

1) this gives a way to get df_{x_0} for "anything" (matrix, tensor, 10-dim fct., etc.)

2) be careful with dimensions

$$f: \mathbb{R}^m \rightarrow \mathbb{R} \quad df_{x_0} \text{ is a row vector } (1 \times m)$$

$$df_{x_0} = (\nabla f(x_0))^T$$

chain rule:

$$\begin{aligned} f: \mathbb{R}^m &\rightarrow \mathbb{R}^n \\ g: \mathbb{R}^n &\rightarrow \mathbb{R}^q \end{aligned}$$

$$d(g \circ f)_{x_0} = dg_{f(x_0)} \circ df_{x_0}$$

$$g(f(x_0))$$

$$= \begin{pmatrix} \downarrow \end{pmatrix} \cdot \begin{pmatrix} \downarrow \end{pmatrix}$$

matrix product of Jacobians

e.g. $f(\mu) = x - \mu$

$$g(v) = v^T A v$$

$$g \circ f(\mu) = (x - \mu)^T A (x - \mu)$$

$$df_{\mu_0} = -I$$

$$dg_{v_0} = v^T (A + A^T)$$

$$d(g \circ f)_{\mu_0} = dg_{f(\mu_0)} \circ df_{\mu_0}$$

$$= (x - \mu_0)^T (A + A^T) (-I)$$

for Gaussian: $-\frac{1}{2} \sum_i (x_i - \mu)^T \Sigma^{-1} (x_i - \mu)$

$$\begin{aligned} \nabla_{\mu} \\ \downarrow \\ (d(g \circ f)_{\mu})^T \end{aligned}$$

$$+\frac{1}{2} \sum_i 2 \Sigma^{-1} (x_i - \mu) \stackrel{\text{want}}{=} 0$$

$$\Rightarrow \hat{\mu}_{MLE} = \frac{1}{n} \sum_i x_i$$

15h29

example 2: derivative of $f(A) \triangleq \log \det(A)$ where assume A is symmetric $A > 0$

can represent the derivative of a fct. from matrix to scalar, as a matrix

$$f(A + \Delta) - f(A) = \text{tr}(f'(A)^T \Delta) + o(\|\Delta\|)$$

$$= \langle \underbrace{f'(A)}_{df_A}, \Delta \rangle + o(\|\Delta\|)$$

df_A

$$\log \det(A + \Delta) \sim \log \det(A) \quad \left\{ \begin{array}{l} A > 0 \Rightarrow \text{invertible; has unique square root } A^{1/2} \end{array} \right.$$

$$= \log \det(A^{1/2} (I + A^{-1/2} \Delta A^{-1/2}) A^{1/2}) - \log \det(A)$$

$$= \log \underbrace{|A|^{1/2} |I + A^{-1/2} \Delta A^{-1/2}| |A|^{1/2}}_{\text{e-values of } B} - \log |A|$$

$$= \log |I + A^{-1/2} \Delta A^{-1/2}|$$

$$\text{use } \det(B) = \prod_i \lambda_i(B)$$

$$Bv = \lambda v$$

$$= \sum \log \lambda_i(I + A^{-1/2} \Delta A^{-1/2})$$

$$\text{use } \lambda(I + B) = 1 + \lambda(B)$$

$$(I + B)v = (1 + \lambda)v$$

$$\begin{aligned}
& - \log |I + A^{-1/2} \Delta A^{-1/2}| \\
& = \sum_i \log \lambda_i(I + A^{-1/2} \Delta A^{-1/2}) \quad \text{use } \lambda(I+B) = 1 + \lambda(B) \quad \Delta v = \lambda v \\
& = \sum_i \log(1 + \lambda_i(A^{-1/2} \Delta A^{-1/2})) \quad \text{use } \lambda(I+B) = 1 + \lambda(B) \quad (I+B)v = (1+\lambda)v \\
& = \sum_i \lambda_i(A^{-1/2} \Delta A^{-1/2}) + O\left(\sum_i \lambda_i^2(A^{-1/2} \Delta A^{-1/2})\right) \quad \log(1+x) = x + O(x^2) \text{ for } |x| < 1 \\
& = \text{tr}(A^{-1/2} \Delta A^{-1/2}) + O(\|A^{-1/2} \Delta A^{-1/2}\|_F^2) \quad \text{because } \lambda \text{ is homogeneous fct.} \\
& \quad \text{ie. } Bv = \lambda v \\
& \quad \left(\frac{B}{b}\right)v = \left(\frac{\lambda}{b}\right)v \\
& = \text{tr}(A^{-1} \Delta) + O(\|A^{-1} \Delta\|_F^2) \quad \text{see Boyd's book A.4.1 for the above proof} \\
& = \text{tr}(A^{-1} \Delta) + o(\|A^{-1} \Delta\|) \quad \Rightarrow \boxed{\frac{d}{dA} \log \det(A) = A^{-1}} \\
& \quad \text{(recall } A \text{ is symmetric)}
\end{aligned}$$

see Boyd's book A.4.1 for the above proof

back to log-likelihood of Gaussian

$$+\frac{n}{2} \log |\Sigma^{-1}| - \frac{n}{2} \langle \Sigma^{-1}, \tilde{\Sigma}(\mu) \rangle \quad (\text{concave fct. of } \Delta = \Sigma^{-1})$$

take derivative w.r. to

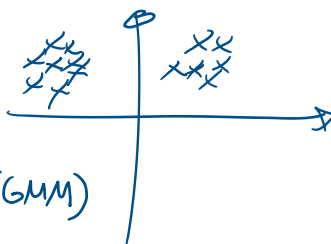
$$\Sigma^{-1} = \Delta \quad \frac{n}{2} \underbrace{(\Sigma^{-1})^{-1}}_{\Sigma} - \frac{n}{2} \tilde{\Sigma}(\mu) \stackrel{\text{want}}{=} 0$$

$$\Rightarrow \hat{\Sigma}_{MLE} = \tilde{\Sigma}(\mu_{MLE}) = \frac{1}{n} \sum_{i=1}^n (x_i - \mu_{MLE})(x_i - \mu_{MLE})^T$$

(the empirical covariance matrix)

unsupervised learning

have X without any label Y



consider the Gaussian mixture model (GMM)
(can be obtained from FLD)

$$Y \sim \text{Mult}(\pi) \quad \pi \in \Delta_K$$

[extension of FLD to multiple classes]

$Y \sim \text{Mult}(\pi) \quad \pi \in \Delta_K$ [extension of FLD to multiple classes]

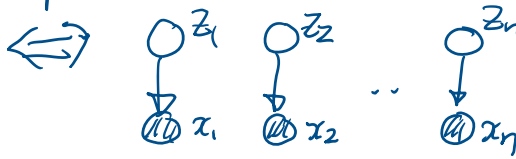
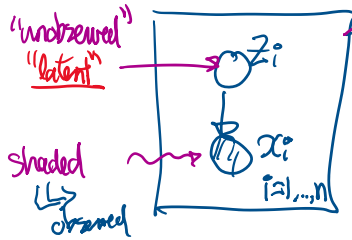
$X|Y=j \sim N(\mu_j, \Sigma)$

$$p(x) = \sum_y p(x|y) = \sum_y p(x|y)p(y) = \sum_{j=1}^K \pi_j N(x|\mu_j, \Sigma)$$

"GMM"

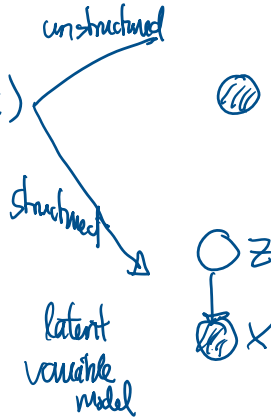
(more generally, can have Σ_j per class)

graphical model for this "latent variable model"



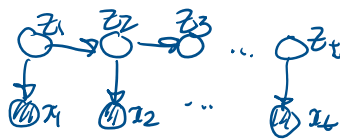
GMM model:
 $z_i \sim \text{Mult}(\pi)$
 $x_i|z_i \sim N(x_i|\mu_{z_i}, \Sigma_{z_i})$

two views on $p(x)$

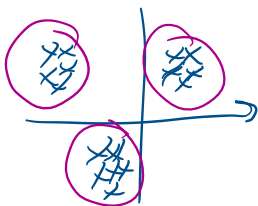


$$p(x) = \sum_z p(x|z) p(z)$$

(later in class, we will add time structure HMM)



k-means → to do clustering i.e. group data



we want to get a cluster assignment for every data point x_i

represent $z_{ij} = 1$ to mean that x_i belongs to cluster j

$j=1, \dots, K$
 K # of clusters (specified in advance for k-means)

applications: • vector quantization (compression)

• in computer vision: use k-means to get

"bag of visual words" representation
of image patches

- many many others?