

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

generative model for classification: (Fisher) linear discriminant analysis

FLD (instead of LDA)

for classification $\mathcal{Y} = \{0, 1\}$
 $x \in \mathbb{R}^d$

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 $p(x|y; \theta) = N(x | \mu_y, \Sigma)$ sharded across class

$\theta = (\underbrace{\mu_0, \mu_1}_{\text{mean for class}}, \underbrace{\Sigma}_{\text{sharded covariance}}, \underbrace{\pi}_{p(y=1)})$

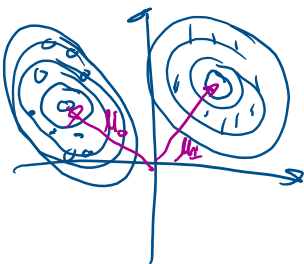
$\begin{pmatrix} \mu_0 \\ \mu_1 \\ \text{vec}(\Sigma) \\ \pi \end{pmatrix}$

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 use Σ_0, Σ_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 estimate θ :



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

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

exercise: MLE for multivariate Gaussian

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

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

$$\Sigma \in \mathbb{R}^{d \times d}$$

Σ is symmetric

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

$$\Sigma^T = \mathbb{E}[\text{ " " }] = \Sigma$$

$$V^T \Sigma V = \mathbb{E}[\underbrace{y^T (x - \mu)}_{\text{ " " }} \underbrace{(x - \mu)^T V}_{\text{ " " }}]$$

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

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

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

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

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

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

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

Vector derivative review:

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

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

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

"differentiable"

"derivative"

$$\lim_{\delta \rightarrow 0} \frac{f(x_0 + \delta \vec{d}) - f(x_0)}{\delta} = \lim_{\delta \rightarrow 0} \frac{df_{x_0}(\delta \vec{d}) + o(\delta \|\vec{d}\|)}{\delta} \\ = \frac{\delta df_{x_0}(\vec{d})}{\delta}$$

$\rightarrow$ directional derivative of f 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

$$df_{x_0}(a_1 + b a_2) \\ = df_{x_0}(a_1) + b df_{x_0}(a_2)$$

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 you a way to get df_{x_0} for "anything" (matrix, tensor, 2nd order fct. etc...)

2) be careful with dimensions

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

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

chain rule

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

$$g: \mathbb{R}^n \rightarrow \mathbb{R}^q$$

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

$$= \begin{pmatrix} \vdots \end{pmatrix} \begin{pmatrix} \vdots \end{pmatrix}$$

matrix product of Jacobians

eg $f(\mu) = x - \mu$

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

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

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

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

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

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

for Gaussian:

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

$$\frac{\partial}{\partial \mu} \left[\frac{1}{2} \sum_i (x_i - \mu)^T \Sigma^{-1} (x_i - \mu) \right] = 0 \Rightarrow \hat{\mu}_{MLE} = \frac{1}{n} \sum_{i=1}^n x_i$$

15h38

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

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

$$f(A + \Delta) - f(A) = \text{tr}(df_A^T \Delta) + o(\|\Delta\|)$$

$$= \langle df_A, \Delta \rangle + o(\|\Delta\|)$$

df_A

$A > 0 \Rightarrow$ invertible; has unique square root $A^{1/2}$

$$\log \det(A + \Delta) - \log \det(A)$$

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

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

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

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

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

$$= \sum_i \lambda_i(A^{-1/2} \Delta A^{-1/2}) + O(\|A^{-1/2} \Delta A^{-1/2}\|^2)$$

$$= \sum_i \lambda_i(A^{-1/2} \Delta A^{-1/2}) + O(\|A^{-1/2} \Delta A^{-1/2}\|^2)$$

$$\|A^{-1/2} \Delta A^{-1/2}\|^2$$

sup norm-1

est. with respect

because A is homogeneous fct.

$$Bv = \lambda v$$

$$(B \setminus v) = (\lambda \setminus v)$$

$$\frac{d}{dA} \log \det(A) = A^{-1}$$

$$\frac{h(\Delta)}{\|\Delta\|} \leq c \frac{\|\Delta\|^2}{\|\Delta\|} = c \|\Delta\|$$

$$= \text{tr}(A^{-1/2} \Delta A^{-1/2}) + o(\|\Delta\|)$$

$$= \text{tr}(A^{-1} \Delta) + o(\|\Delta\|)$$

$$\Rightarrow \frac{d}{dA} \log \det(A) = A^{-1}$$

$$\text{ie. } Bv = Av$$

$$\left(\frac{B}{b}\right)v = \left(\frac{A}{b}\right)v$$

$$\text{sup } \|\Delta\| = 1 \text{ const. with respect to } \|\Delta\|$$

$$O(\|\Delta\|^2) = o(\|\Delta\|)$$

$$\text{(recall } A \text{ is symmetric)}$$

see [Boyd's book](#) A.4.1 for the above proof

back to log-likelihood of Gaussian

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

take derivative w.r. to $\Sigma^{-1} = \Sigma$

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

$$\Rightarrow \hat{\Sigma}_{MLE} = \tilde{\Sigma}(\mu_{MLE})$$

$$= \frac{1}{n} \sum_{i=1}^n (x_i - \mu_{MLE})(x_i - \mu_{MLE})^T$$

(the empirical covariance matrix)

unsupervised learning

here X without any label



consider the Gaussian mixture model (GMM)
(can be obtained FLO)

[extension of FLO to multiple classes]

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

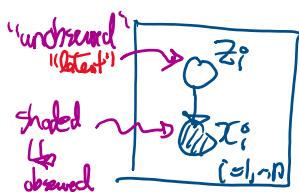
$$X | Y=j \sim N(\mu_j, \Sigma)$$

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

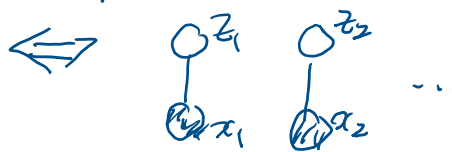
"GMM"

(more generally, can have Σ_j per class)

graphical model for this "latent variable model"



"plate" = repetition



$$\left. \begin{array}{l} z_1 \sim \text{Mult}(\pi) \\ x_i | z_i \sim N(x | \mu) \end{array} \right\} \text{GMM model}$$

observed latent

x_1 x_2

x_n $x_i | z_i \sim N(x | \mu_{z_i}, \Sigma_{z_i})$