

Lecture 11 FML March
30, 2021

multiclass classification

M classes $H_j, j=0, \dots, M-1$

$y \in \mathcal{Y}$

$\underbrace{\quad}_{\text{observations}}$

Bayesian Framework

- If data is coming from class j then y is a r.v. with distribution $f(y|H_j)$
- we know $f(y|H_j), j=0, \dots, M-1$

Prior probs of class H_j
 $= \pi_j$

C_{ij} = cost of deciding class i for observation y but it comes from class j .

≥ 0 and

$$\underline{\underline{C_{jj}}} < \underline{\underline{C_{ij}}}, i \neq j$$

Decision rule
 $\underline{\underline{\delta(y)}} \in 0, \dots, M-1$

If $y \in \mathcal{Y}_j$ then

$$\underline{\underline{\delta(y) = H_j}}$$

$\mathcal{Y}_j$ are disjoint



$$\mathcal{Y}_j = \{y : \delta(y) = j\}$$

Bayes' Risk

$$\underline{\underline{R(\delta) = \sum_{j=0}^{M-1} \sum_{i=0}^{M-1} c_{ij} \cdot \underline{\underline{P\{y \in \mathcal{Y}_i | H_j\}} \pi_j}}}$$

$$\underline{\underline{P\{y \in \mathcal{Y}_i | H_j\}}} \pi_j$$

$y \in \mathbb{R}^d$
then
 $\mathcal{Y}_j \subseteq \mathbb{R}^d$

$$\bigcup_j \mathcal{Y}_j \subseteq \mathbb{R}^d$$

— we want to choose
decision rule δ s.t.

$R(\delta)$ is minimized

$$R(\delta) = \sum_{i=0}^{M-1} \sum_{j=0}^{M-1} c_{ij} \cdot \underline{\underline{P(y_i | H_j) \pi_j}} \\ = \sum_{i=0}^{M-1} \sum_{j=0}^{M-1} c_{ij} \int_{\mathcal{Y}_i} f(y | H_j) \pi_j dy$$

we know
 $f(y | H_j)$,
 $P(H_j) = \pi_j$
 $\underline{\underline{f(y) = \sum_j f(y | H_j) \pi_j}}$

$$\underline{\underline{P(H_j | y) f(y)}} \uparrow \text{prob } y$$

$$\begin{aligned}
 \underline{R(\delta)} &= \sum_i \sum_j a_{ij} \int_{y_i} f(y|H_j) \pi_j dy \\
 &= \sum_i \sum_j \underline{c_{ij}} \int_{y_i} P(H_j|y) f(y) dy \quad \text{P(H}_j|y) f(y) \\
 &= \sum_i \int_{y_i} \left(\sum_j \underline{a_{ij}} P(H_j|y) \right) f(y) dy \\
 &= \sum_{i=0}^{M-1} \int_{y_i} \underline{c_i(y)} f(y) dy \\
 \underline{c_i(y)} &= \sum_{j=0}^{M-1} \underline{a_{ij}} P(H_j|Y=y) \\
 &\quad \text{Average cost of } \underline{\delta} \text{ given observation } y.
 \end{aligned}$$

- finding decision rule δ minimizing $R(\delta)$ is equivalent to finding regions y_i which will minimize $R(\delta)$.

$$\begin{aligned}
 \underline{y_i} &= \{y : \underline{c_i(y)} \\
 &= \min_{j=0, \dots, M-1} c_j(y)\}
 \end{aligned}$$

optimal decision rule (Bayes decision rule)

$$\underline{\delta_B(y)} = \underset{0 \leq i \leq M-1}{\operatorname{argmin}} c_i(y)$$

Special cases:

$$c_{ij} = 1 \{i \neq j\}$$

$$c_{jj} = 0.$$

This gives Bayes rule that minimizes prob of error

$$C_i(y) = \sum_{j=0}^{m-1} c_{ij} P(H_j | Y=y)$$

$$= \sum_{j \neq i} P(H_j | Y=y)$$

$$= 1 - P(H_i | Y=y)$$

$$\text{a posteriori} = \arg \max_i \boxed{P(H_i | Y=y)}$$

$\pi_j = P(H_j)$ a priori prob of H_j

$$P(y | H_j) \quad j = \dots$$

$$P(H_j | y)$$

$$= \frac{P(H_j, y)}{P(y)} = \frac{P(y | H_j) \pi_j}{P(y)}$$

Bayes rule

$$f_B(y) = \arg \min_i C_i(y)$$

MAP $\leftarrow \arg \min_i (1 - P(H_i | Y=y))$

$$\begin{aligned}
 f_{\text{MAP}}(y) &= \arg \max_i P[H_i | Y=y] \\
 &= \arg \max_i \frac{P[H_i, Y=y]}{P(Y=y)} \\
 &= \arg \max_i P[H_i, Y=y] \\
 &= \arg \max_i P[Y=y | H_i] \pi_i \\
 &\rightarrow \text{If } \pi_i \text{ are all same} \\
 &\quad \therefore \pi_i = \frac{1}{M} \\
 &= \arg \max_i P[Y=y | H_i]
 \end{aligned}$$

ML classifier

Ex

M classes

Y vector of dim m

$$p(y | H_i) = \mathcal{N}(m_i, \Lambda_i)$$

MAP classifier

$$= \arg \max_i \frac{\exp(-\frac{1}{2}(y-m_i)^T \Lambda_i^{-1} (y-m_i))}{(2\pi)^{m/2} \det(\Lambda_i)^{1/2}}$$

$$\Lambda_i = \sigma^2 I, \forall i$$

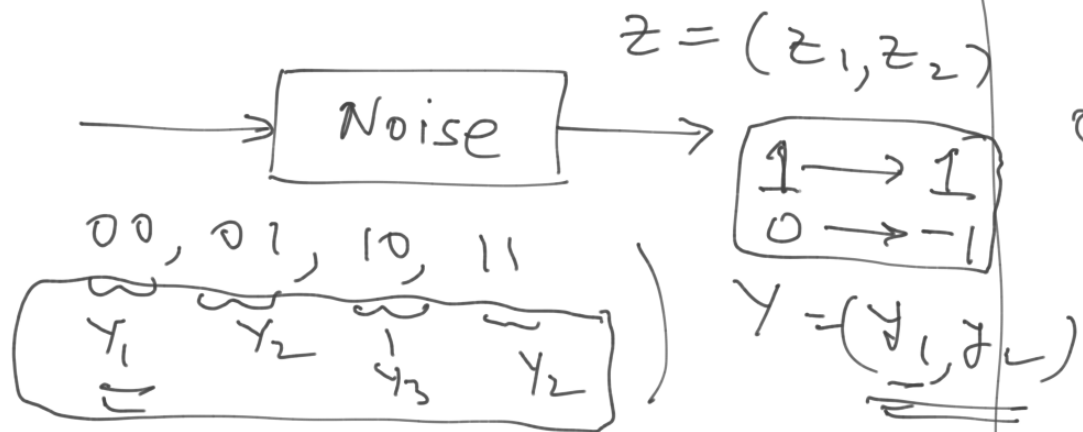
$$\pi_i = 1/M$$

$$\arg \max_i \left(\exp(-\frac{1}{2\sigma^2} (y-m_i)^T (y-m_i)) \right)$$

$$= \arg \min_i (y-m_i)^T (y-m_i)$$

$$= \arg \min_i \|y-m_i\|^2$$

min distance classifier
ex communication



$$\underline{z_i} = y_i + \underline{N_i}, i=1,2.$$

Receiving $\underline{z}$,
 we want to know
what was transmitted

$$P\{z|H_i\}$$

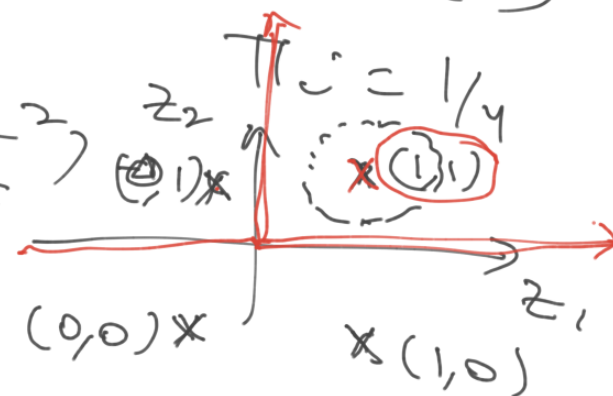
$$00 \rightarrow (N_1, N_2) \rightarrow (0,0)$$

$$\sigma^2 I = \sigma^2 \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix}$$

$$01 \rightarrow (N_1, 1+N_2)$$

$$\text{cov. } \sigma^2 I$$

$$\text{mean } (0, 1)$$



ML class

In ML we may

not know $P(y|H_i)$

We know prior dist

$\pi_i = P(H_i) \leftarrow$ prior.

iid observations m sample
from a class $0, \dots, M-1$

$\mathcal{H} = \{0, 1, \dots, M-1\}$

We have posterior
dist (based on m iid
sample)

$Q(h) = P(H_i | y_1, \dots, y_m)$

(loss) cost fn

$\ell(h)$

$$\underline{\underline{R(Q)}} = E_Q[\ell(h(x), y)]$$

$(x_1, y_1), \dots, (x_m, y_m)$
label $\in \mathcal{H}$.

Theorem with prob $\geq 1-\delta$

$$\underline{\underline{R(Q)}} \leq \underline{\underline{R(Q)}} + \left[D(Q||P) + \ln \frac{m}{\delta} \right] / 2(m-1)^{1/2} - \textcircled{X}$$

Find classifier that

min RHS in above theorem

$$D(Q|P) = \sum_x Q(x) \log Q(x) - \sum_x Q(x) \log P(x)$$

if support of Q is in P
o.w. $= \infty$.

ML Setup

X input space

Y output space

$\rightarrow k$ no of classes $x_i \in X$

Sample $S = (x_1, y_1), \dots, (x_m, y_m)$
iid $\sim D$, D not known.

Assumption $\exists f :$

$$y_i = f(x_i)$$

$$Y = \{1, 2, \dots, k\}$$

$\mathcal{H}$ Hypothesis set.

- We estimate f from some function $h \in \mathcal{H}$

$$\underline{h} : \underline{X} \rightarrow Y$$

$$R(h) = E_D [1\{h(x) \neq f(x)\}]$$

= expected error for $h \in \mathcal{H}$.

Generalization error

- Based on m samples we find $h \in \mathcal{H}$ which minimizes $R(h)$.
- We don't know $R(h)$ because we do not know the dist D of (x, y) .
- From the sample we can find out empirical error:

$$\hat{R}(h) = \frac{1}{m} \sum_{i=1}^m \mathbb{1}_{\{h(x_i) \neq y_i\}}$$

For the binary classification using SVM:

$$y = \{-1, 1\}$$

$h(x) = \text{sgn}(w \cdot x + b)$
 (w, b) define a hyperplane
 optimal sol. using dual problem

$$= \text{sgn}\left(\sum_{i=1}^m \alpha_i y_i (x_i \cdot x) + b\right)$$

$$h(x) = \arg \max_{y \in \{-1, 1\}} (y \cdot (w \cdot x + b))$$

-ve

$$h(x) = \text{sgn}(w \cdot x + b)$$

$$= \text{sgn} \left(\sum_{i=1}^m \alpha_i y_i (x_i \cdot x + b) \right)$$

$$= \underset{y \in \{-1, 1\}}{\text{argmax}} (y (w \cdot x + b))$$

$$= \underset{y \in \{-1, 1\}}{\text{argmax}} \left(y \cdot \left(\sum_{i=1}^m \alpha_i y_i (x_i \cdot x + b) \right) \right)$$

In kernel svm

$$K: \mathcal{X} \times \mathcal{X} \rightarrow \mathbb{R}$$

$$\Phi: \mathcal{X} \rightarrow \mathcal{H} \quad \text{Hilbert space}$$

$$K(x, x') = \langle \Phi(x), \Phi(x') \rangle$$

for k class classifier

$$h(x) = \underset{y \in \{1, \dots, k\}}{\text{argmax}}$$

$$h(x, y)$$

score function.

replace

$$x_i \cdot x \rightarrow K(x_i, x)$$

$$= \langle \Phi(x_i), \Phi(x) \rangle$$

We absorb b in kernel.

$$\underset{y \in \{-1, 1\}}{\text{argmax}} \left(y \sum_{i=1}^m \alpha_i y_i \langle \Phi(x_i), \Phi(x) \rangle \right)$$

$$= \underset{y \in \{-1, 1\}}{\text{argmax}} \left(\langle \sum_{i=1}^m y \alpha_i y_i \Phi(x_i), \Phi(x) \rangle \right)$$

α_i picked optimal from dual optimization problem

$$h(x) = \underset{y \in \{-1, 1\}}{\text{argmax}} \langle w_y, \Phi(x) \rangle$$