

today: Statistical decision theory
evaluation of estimators

Statistical decision theory - formal setup

1) • a random observation $D \sim P$ *unknown distribution which models the world / phenomenon*
(often P_θ)

2) • action space A

3) • ^(statistical) loss

$L(P, a) =$ statistical loss of doing action $a \in A$ when the world is P } describe the goal / task

↳ often write $L(\theta, a)$ if we have a parametric model of world i.e. P is a pdf / pmf P_θ for some $\theta \in \Theta$

• $\delta: \mathcal{D} \rightarrow A$ "decision rule"
 $\mathcal{D}$

examples: a) parameter estimation:

$\Omega = \Theta$ for parametric family P_θ

δ is a parameter estimator from data

typically $D = (X_1, \dots, X_n)$

typical loss

$$L(\theta, a) = \|\theta - a\|_2^2$$

[usually $X_i \text{ i.i.d. } P_\theta$]

"squared loss"

but another loss is $KL(P_\theta \| P_a)$

b) $A = \{0, 1\}$; this is hypothesis testing

δ describes a statistical test

$$\mathbb{1}_A(u) = \begin{cases} 1 & \text{if } u \in A \\ 0 & \text{o.w.} \end{cases}$$

loss $\rightarrow$ usually 0-1 loss $L(\theta, a) = \mathbb{1}_{\{\theta \neq a\}}(\delta) \left(\triangleq \mathbb{1}_{\{\theta \neq a\}} \right)$
 $\mathcal{D}^c = \Theta \setminus \{\theta\}$

$\rightarrow$ need to find a δ that minimizes the loss

$$\mathcal{Z}^c = \mathcal{H} \setminus \{0\}$$

c) prediction in ML: learn a prediction fct. in supervised learning (function estimation)

$$\text{here } D = (X_i, Y_i)_{i=1}^n$$

$$X_i \in \mathcal{X} \text{ (input space)} \\ Y_i \in \mathcal{Y} \text{ (output space)}$$

$$\mathcal{Y} = \{0, 1\} \rightarrow \text{classification}$$

$$\mathcal{Y} = \mathbb{R} \rightarrow \text{regression}$$

P_θ gives joint on (X, Y)

$$\Omega = \mathcal{Y}^{\mathcal{X}} \text{ (set of fct's from } \mathcal{X} \text{ to } \mathcal{Y})$$

$$|\Omega| = |\mathcal{Y}|^{|\mathcal{X}|}$$

$$D \sim P \text{ where } P = \underbrace{P_\theta \otimes P_\theta \otimes \dots \otimes P_\theta}_{n \text{ times}}$$

$$p(x_1, \dots, x_n) = \prod_{i=1}^n p_\theta(x_i)$$

in ML

$$L(P_\theta, f) \triangleq \mathbb{E}_{P_\theta} [l(Y, f(X))]$$

"generalization error"

"classification error"

$$(X, Y) \sim P_\theta$$

"prediction loss"

in ML, is often call the "risk"

→ e.g. classification

$$l(Y, f(X)) = \mathbb{1}_{\{Y \neq f(X)\}}$$

Shannon calls it "Vapnik risk"

to distinguish it from frequentist risk

$$\mathbb{E}_D [L(P_\theta, \delta(D))]$$

→ think cross-validation e.g.

* decision rule

$$\hat{f} = \delta(D)$$

prediction fct./
classifier/etc.

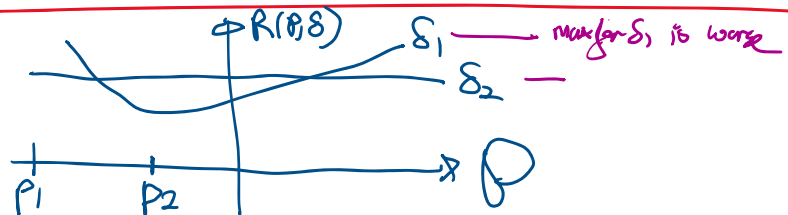
"learning alg."

comparing procedures?

δ_1 vs. δ_2

$$\text{(frequentist) risk } R(P, \delta) \triangleq \mathbb{E}_{D \sim P} [L(P, \delta(D))]$$

"risk profiles"



* transform to scalen

• "minimax" analysis: $\max_{D \in \mathcal{D}} R(P, \delta)$ "worst case"

• "minimax" analysis: $\max_{P \in \mathcal{P}} R(P, \delta)$ "worst case"

• weighted average $\int_{\mathcal{H}} R(f_\theta, \delta) \pi(\theta) d\theta$ (kind of a Bayesian feel)
 (weighting)

5h3)

PAC theory vs frequentist risk:

in ML, usually they look at tail bound for dist. of $L(P, \delta(D))$ where D is random

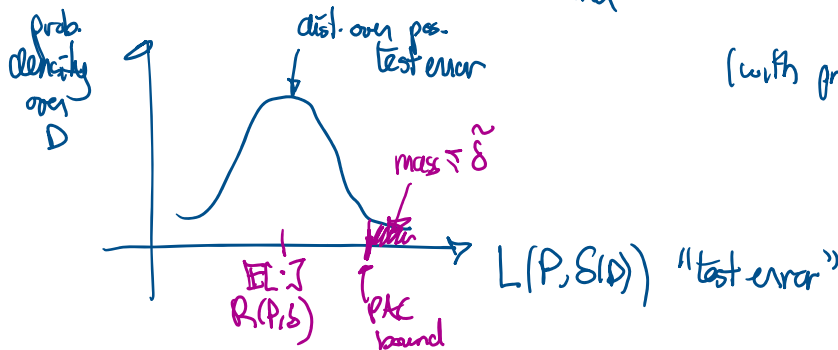
PAC theory
 $\hookrightarrow$ "probably approx. correct"

$$P\{L(P, \delta(D)) \geq \text{stuff}\} \leq \tilde{\delta}$$

example of
 generalization error
 bound

$$\text{test-error}(\hat{f}) \leq \text{train-error}(\hat{f}) + \frac{1}{\sqrt{n}} \sqrt{\text{complexity}(\hat{f})} + \log \frac{1}{\tilde{\delta}}$$

(with prob $\geq 1 - \tilde{\delta}$ on D)



Bayesian decision theory

$\rightarrow$ condition on data D

Bayesian posterior risk $R_B(a|D) = \int_{\mathcal{H}} L(\theta, a) p(\theta|D) d\theta$
 (posterior over "possible worlds")
 $\propto p(\theta) p(D|\theta)$

Bayesian optimal action: $S_{\text{Bayes}}(D)$
 $\triangleq \arg\min_{a \in \mathcal{A}} R_B(a|D)$

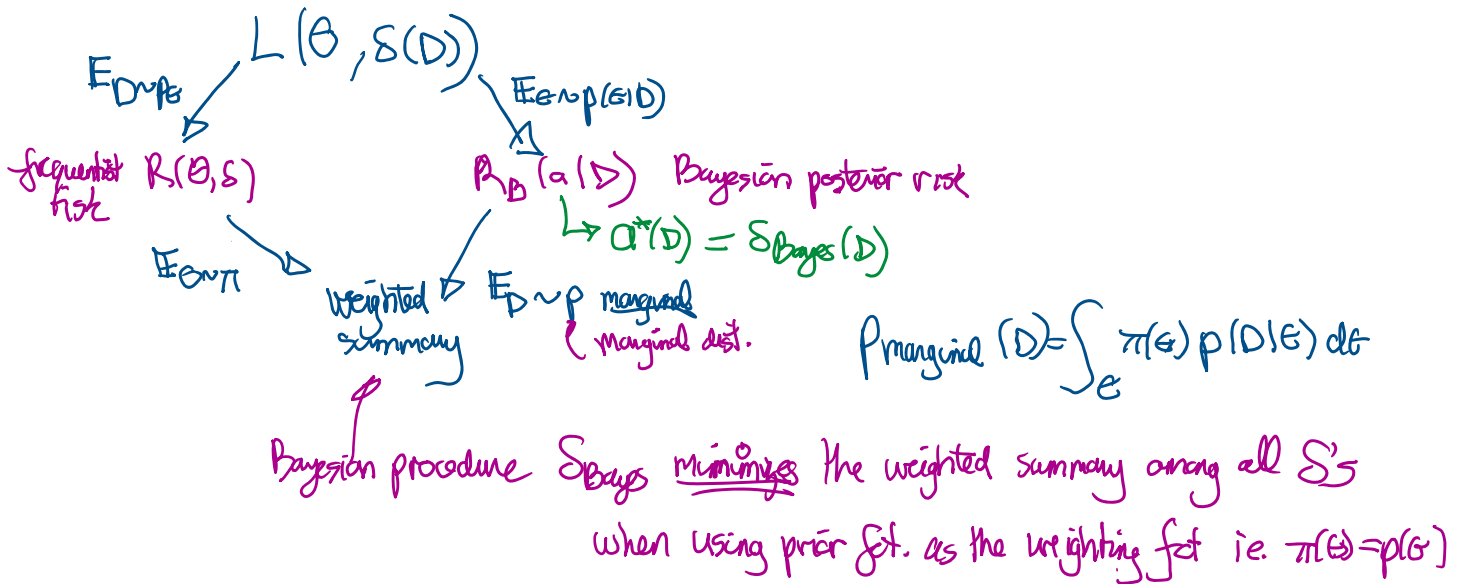
example: if $\mathcal{A} = \mathbb{R}$ ("estimation")

$$L(\theta, a) = \|\theta - a\|_2^2$$

then (exercise) $\delta_{\text{Bayes}}(D) = \mathbb{E}[\xi | D]$ (posterior mean)

but if use $L(\xi, a) = |\xi - a|$ (1D)

then $\delta_{\text{Bayes}}(D) = \underline{\text{posterior median}}$



Examples of estimators: $\delta: \mathcal{D} \rightarrow \mathcal{H}$

- 1) MLE
- 2) MAP
- 3) method of moments (MoM)

idea: find an injective mapping from $\mathcal{H}$ to "moments" of R.V.

$$\mathbb{E}X, \mathbb{E}X^2, \dots$$

and then invert it from empirical moments to get $\hat{\theta}$

$$\hat{\mathbb{E}}[X] \triangleq \frac{1}{n} \sum_{i=1}^n X_i$$

$$\hat{\mathbb{E}}[X^2] \triangleq \frac{1}{n} \sum_{i=1}^n X_i^2$$

example: for Gaussian $X \sim N(\mu, \sigma^2)$

$$\begin{aligned} \mathbb{E}X &= \mu \\ \mathbb{E}X^2 &= \sigma^2 + \mu^2 \end{aligned}$$

$$f(\mu, \sigma^2) \triangleq \begin{pmatrix} \mu \\ \sigma^2 + \mu^2 \end{pmatrix}$$

$$\begin{pmatrix} \hat{\mu} \\ \hat{\sigma}^2 \end{pmatrix} \triangleq f\left(\begin{pmatrix} \mathbb{E}[x] \\ \mathbb{E}[x^2] \end{pmatrix}\right) = \begin{pmatrix} \hat{\mathbb{E}}[x] \\ \hat{\mathbb{E}}[x^2] - (\hat{\mathbb{E}}[x])^2 \end{pmatrix}$$

(here, this estimator is same as MLE)
[general property in exponential family $\rightarrow$ see later]

⊛ MoM is quite used for latent variable models

$\rightarrow$ ("spectral methods" e.g.)



4) prediction example: $\mathcal{A} = \{f: \mathcal{X} \rightarrow \mathcal{Y}\}$ $\mathcal{X}$ = input space
 $\mathcal{Y}$ = output space

example of $\mathcal{S}: \mathcal{D} \rightarrow \mathcal{A}$

is using **empirical "risk" minimization (ERM)**
 $\rightarrow$ Vapnik risk i.e. generalization / test error

$$\text{i.e. } L(P, f) = \mathbb{E}_{(x,y) \sim P} [l(y, f(x))]$$

$$\text{replace this with } \hat{\mathbb{E}}[l(y, f(x))] = \frac{1}{n} \sum_{i=1}^n l(y_i, f(x_i))$$

$$\hat{\mathcal{S}}_{\text{ERM}} = \underset{f \in \mathcal{H}}{\text{argmin}} \hat{\mathbb{E}}[l(y, f(x))]$$

$\mathcal{H}$ = hypothesis class

James-Stein estimator:

estimator to estimate the mean of $N(\vec{\mu}, \sigma^2 I)$ $\leftarrow d$ independent Gaussian variable
 $X_i \sim N(\mu_i, \sigma^2)$
 $\hat{\mathcal{S}}_{\text{JS}}$ is biased, but much lower variance than MLE

$$\text{recall bias-variance decomposition: } R(\theta, \hat{\theta}) = \mathbb{E}[\|\theta - \hat{\theta}\|_2^2] \\ = \underbrace{\|\mathbb{E}\hat{\theta} - \theta\|_2^2}_{\text{bias}^2} + \underbrace{\mathbb{E}\|\hat{\theta} - \mathbb{E}\hat{\theta}\|_2^2}_{\text{variance}}$$

$\hat{\mathcal{S}}_{\text{JS}}$ actually strictly dominates $\hat{\mathcal{S}}_{\text{MLE}}$

for $d \geq 3$
 $\uparrow$ dimension of $\vec{\mu}$

$$\text{i.e. } R(\theta, \hat{\mathcal{S}}_{\text{JS}}) \leq R(\theta, \hat{\mathcal{S}}_{\text{MLE}}) + \epsilon$$

for $d \geq 3$
↑ dimension of $\vec{\mu}$

$$\text{ie. } R(\theta, \delta_{JS}) \leq R(\theta, \delta_{MLE}) + \frac{1}{n}$$

$$\text{and } \exists \theta \text{ st. } R(\theta, \delta_{JS}) < R(\theta, \delta_{MLE})$$

→ MLE is inadmissible in this case [note $n=1$]
here

← $d \geq 3$

(can interpret the δ_{JS} as an "empirical" Bayes method)