

Lecture 6 - statistical decision theory

Monday, September 23, 2024 12:32 PM

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 world is P } describes the goal/task

↳ often write $L(\theta, a)$ if we have a parametric model of world
ie. P has a pdf/pmf p_θ for some $\theta \in \Theta$

• $S: \underset{P \rightarrow \Omega_D}{D} \rightarrow A$ "decision rule"

examples: a) parameter estimation

$A = \Theta$ for some parametric family $\mathcal{P}_\Theta$

S is a parameter estimator from data

^{typically} $D = (x_1, x_2, \dots, x_n)$
[usually $x_i \stackrel{iid}{\sim} p_\theta$]

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

"squared loss"

but another loss is $KL(p_\theta \| p_a)$

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

S describes a statistical test

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

loss → usually 0/1 loss $L(\theta, a) = \mathbb{1}_{\{\theta \notin C\}}(a) (\triangleq \mathbb{1}_{\{\theta \neq a\}})$

$\{\theta\}^C = \Theta \setminus \{\theta\}$

→ related to false & true errors

$\{\theta\}^C$ "whole"

$$105 = 100 + 5$$

→ related to type I & type II errors

$$\{ \emptyset \}^C \text{ "whole world"}$$

$$A^C = B \setminus A$$

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

here $D = (x_i, y_i)_{i=1}^n$ $x_i \in X$ (input space) $y_i \in Y$ (output space)

$Y = \{0,1\} \rightarrow$ classification
 $Y = \mathbb{R} \rightarrow$ regression

let p_0 be joint on (X, Y)
 $\mathcal{F} = Y^X$ (set of functions from X to Y)

$D \sim P$ where $P = \underbrace{p_0 \otimes p_0 \otimes p_0 \dots \otimes p_0}_{n \text{ times}}$

$$p(\tilde{x}_1, \dots, \tilde{x}_n) = \prod_{i=1}^n p_0(\tilde{x}_i)$$

in ML $L(p_0, f) \triangleq \mathbb{E}_{(x,y) \sim p_0} [l(y, f(x))]$

↑
"prediction loss"

"generalization error"
 "classification error"

in ML, is often called the "risk"

eg classification $l(y, f(x)) = \mathbb{1}\{y \neq f(x)\}$
 0-1 error

Simon calls it the "Vapnik risk" to distinguish it from frequentist risk

$$\mathbb{E}_D [L(p_0, \delta(D))]$$

→ think cross-validation eg

* decision rule

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

↑ "learning dg"

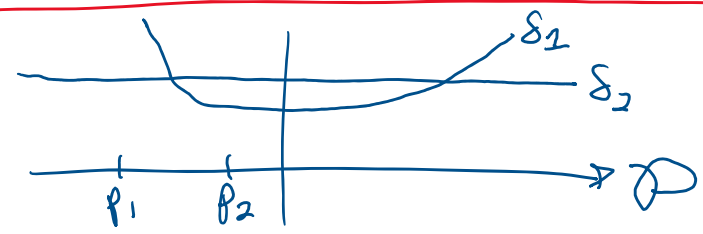
prediction fct. classifier/regressor/etc. ...

comparing procedures?

δ_1 vs δ_2

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

"risk profiles"



* transform to scalar

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

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

13h34

PAC theory vs. frequentist risk:

in ML, usually they look at test bound for dist. on $L(P, S(D))$
 where D is random

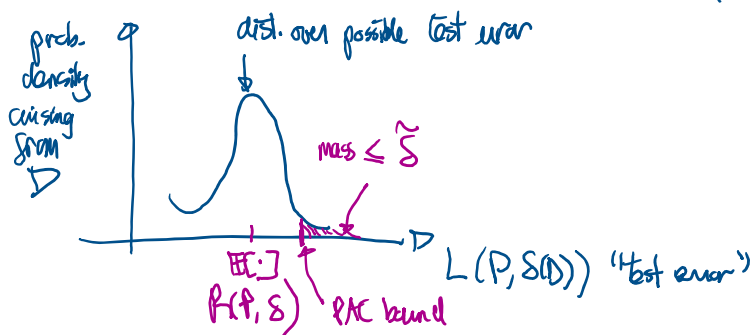
PAC theory
 ↳ "probably approx. correct"

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

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 $\geq 1 - \tilde{\delta}$)
 ↳ randomness of D



Bayesian decision theory

→ condition on data D

Bayesian posterior risk $R_B(a|D) = \int_{\Theta} L(\theta, a) p(\theta|D) d\theta$

Bayesian optimal action $S_{\text{Bayes}}(D)$

$$\triangleq \arg\min_{a \in \mathcal{A}} R_B(a|D)$$

↳ posterior over "possible worlds"
 or $p(\theta) p(D|\theta)$

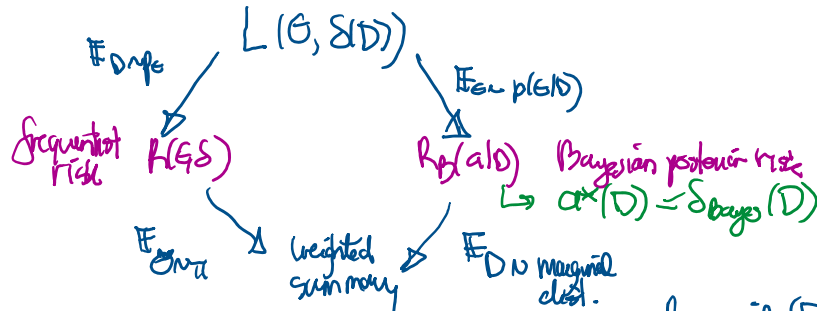
example: if $\mathcal{A} = \Theta$ "estimation"

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

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

but if use $L(G, a) = (G - a)^2$

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



$$p_{\text{marginal}}(D) = \int_{\Theta} \pi(\theta) p(D|\theta) d\theta$$

Bayesian procedure δ_{Bayes} minimizes the weighted frequentist risk among all δ 's

when using prior π as the weighting π , i.e. $\pi(\theta) = p(\theta)$

Examples of estimators: $\delta: D \rightarrow \Theta$

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

idea: find an injective mapping from Θ to "moments" of R.V.

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

$$\mathbb{E}[X] \cong \frac{1}{n} \sum_{i=1}^n x_i$$

$$\mathbb{E}[X^2] \cong \frac{1}{n} \sum_{i=1}^n x_i^2$$

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

$$\mathbb{E}X = \mu$$

$$\mathbb{E}X^2 = \sigma^2 + \mu^2$$

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

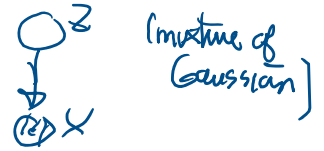
$$\begin{pmatrix} \hat{\mu} \\ \hat{\sigma}^2 \end{pmatrix} \triangleq f^{-1} \left(\begin{pmatrix} \mathbb{E}[X] \\ \mathbb{E}[X^2] \end{pmatrix} \right) = \begin{pmatrix} \mathbb{E}[X] \\ \mathbb{E}[X^2] - (\mathbb{E}[X])^2 \end{pmatrix}$$

(here, this estimator gives same answer as MLE)

[general property in exponential family $\rightarrow$ see later]

(*) MoM is quite useful for latent variable models

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



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

$\mathcal{Y} \leftarrow$ output space

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

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

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

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

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

$\nwarrow$ hypothesis class

James-Stein estimator:

estimator to estimate the mean of $N(\vec{\mu}, \sigma^2 I)$

S_{JS} is baised, but much lower variance than MLE

$\leftarrow d$ independent Gaussian variables
 $X_i \sim N(\mu_i, \sigma^2)$

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

$$= \underbrace{\mathbb{E}\|\hat{\theta} - \theta\|_2^2}_{\text{bias}^2} + \underbrace{\mathbb{E}\|\hat{\theta} - \mathbb{E}(\hat{\theta})\|_2^2}_{\text{variance}}$$

S_{JS} actually strictly dominates S_{MLE}

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

$$\text{i.e. } R(\theta, S_{JS}) \leq R(\theta, S_{MLE}) \forall \theta$$

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

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

$\leftarrow d \geq 3$

(... statement that the ... no ... method)

