

today: theory basics

theory basics

decision theory setup

estimate: $h_w: X \rightarrow \mathcal{Y}$

generalization error = $L_P(w) \triangleq \mathbb{E}_{(x,y) \sim P} [\overset{\text{task loss}}{\ell(y, h_w(x))}]$

ultimate goal is to find $w^* = \underset{w \in \mathcal{W}}{\operatorname{argmin}} L_P(w)$

problem: do not know P (^{"true"} distribution on (x,y))

suppose $\underbrace{(x^{(i)}, y^{(i)})}_{\triangleq P_n \text{ training data}}_{i=1}^n \overset{\text{i.i.d.}}{\sim} P$

→ we could look at

$$\hat{L}_n(w) = \frac{1}{n} \sum_{i=1}^n \ell(y^{(i)}, h_w(x^{(i)}))$$

from statistics / prob. theory

$$\hat{L}_n(w) \xrightarrow[n \rightarrow \infty]{a.s.} L_P(w) \quad \forall w$$

→ this is weaker
than $\sup_w |\hat{L}_n(w) - L_P(w)| \xrightarrow[n \rightarrow \infty]{a.s.} 0$

$$\left[\begin{array}{c} X_n \xrightarrow{a.s.} X \\ \sim \quad \quad \quad \sim \end{array} \right]$$

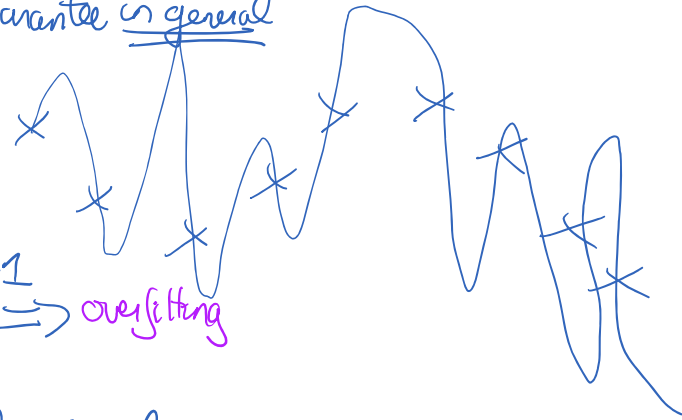
$$\mathbb{P}\{w \in \Omega : X_n(w) \xrightarrow{n \rightarrow \infty} X(w)\} = 0$$

note: minimizing training error gives no guarantee in general

e.g. polynomial regression

for n points, can
get zero training error
with poly. of degree $n-1$

$\Rightarrow$ overfitting



in learning theory, study properties of learning alg

in particular, what can we say about $L_p(A(D_n))$?

different approaches:

a) "frequentist risk"

$$R_{P,n}^F(A) \triangleq \mathbb{E}_{D_n \sim p^{\otimes n}} [L_p(A(D_n))] \quad \text{where } D_n \text{ is random}$$

b) PAC framework
"probably approximately correct"

$$\mathbb{P}\{L_p(A(D_n)) > \text{some bound}\} \leq \delta$$

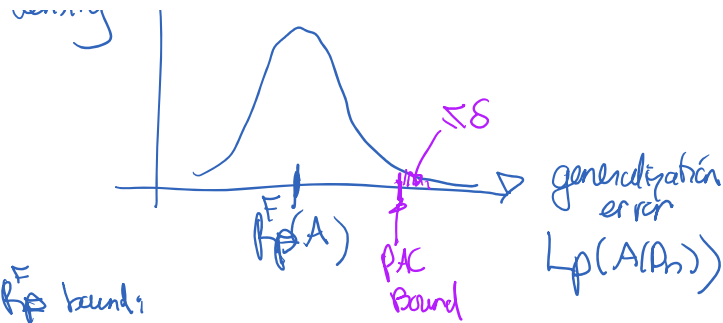
("tail bound")

$$\text{i.e. } L_p(A(D_n)) \leq \text{" " " with prob. } \geq 1 - \delta$$

generalization error bound

prob.
density



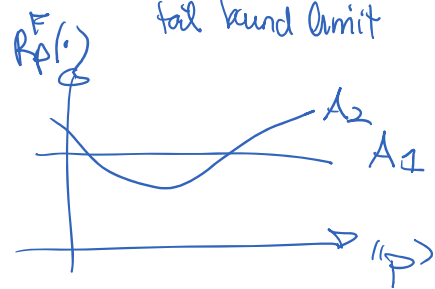


from PAC bound to R_p^F bound:

$$\mathbb{E}X = \underbrace{\mathbb{E}X \mathbb{1}_{\{X \leq B\}}}_{\leq B \cdot (1-\delta)} + \underbrace{\mathbb{E}X \mathbb{1}_{\{X > B\}}}_{\leq B \cdot \delta}$$

fail bound limit

issue with $R_p^F \rightarrow$ depends on D



weighted frequentist risk
 $\mathbb{E}_{\theta \sim p(\theta)} [R_{p_\theta}^F(A)]$

c) "Bayesian posterior risk"

$$R^{\text{post}}(w | D_n) \triangleq \mathbb{E}_{\theta \sim p(\theta | D_n)} [L_{p_\theta}(w)]$$

- prior $p(\theta)$ over distributions
- observation model $p(D_n | \theta)$
- posterior $\Rightarrow p(\theta | D_n)$

Bayesian estimate $\hat{w}_n^{\text{Bayes}} = \arg \min_w R^{\text{post}}(w | D_n)$

A^{Bayesian} is optimal for weighted frequentist risk using $p(\theta)$

14/12/20

No free lunch

frequentist risk analysis of learning algorithm A

let $\mathcal{P}$ be a set of distributions $X \times Y$

sample complexity of A with respect to $\mathcal{P}$

is the smallest $n(\mathcal{P}, A, \epsilon)$ s.t. $\forall n \geq n(\mathcal{P}, A, \epsilon)$

we have $\sup_{P \in \mathcal{P}} [R_P^F(A; n) - L_P(h_P^*)] < \epsilon$

"uniform result"

$h_P^* = \argmin_{h: X \rightarrow Y} L_P(h)$

terminology: A is consistent for dist. P

if $\lim_{n \rightarrow \infty} R_P^F(A; n) - L_P(h_P^*) = 0$

A is uniformly consistent for a family $\mathcal{P}$

if $\lim_{n \rightarrow \infty} \left[\sup_{P \in \mathcal{P}} [R_P^F(A; n) - L_P(h_P^*)] \right] = 0$

Binary classification $Y = \{-1, +1\}$

I) if X is finite, then the "voting procedure" (assign the most frequent label to an input x)

is uniformly and universally consistent

$\rightarrow$ i.e. $\mathcal{P}$ is all distributions on $X \times Y$

with (universal) sample complexity $n(\mathcal{P}, \epsilon, A_{\text{voting}}) \leq \frac{|X|}{\epsilon^2}$ (free lunch? ")

II) if X is infinite

no free lunch thm

(for binary class with 0-1 loss)

for any n and any learning alg. A

$$\text{then } \sup_{P \text{ all dist.}} [R_P^F(A; n) - L_P(h_P^*)] \geq \frac{1}{2}$$

i.e. $\exists$ always a dist. P st. your alg. A is worse than random prediction (?)

NFLT II:

[thm. 7.2 in Devroye et al. 1996]

Let ϵ_n be any non-increasing sequence converging to zero (could be arbitrarily slowly) (e.g. $\frac{1}{\lg(\lg(\dots(n)))}$)

then $\exists P$ st. $R_P^F(A; n) - L_P(h_P^*) \geq \epsilon_n \quad \forall n$

⊗ ⊗ consequence: we need assumptions on P to say anything

Occam's generalization error bound

- binary class, and 0-1 loss
- consider W to be a countable set

Let's define a prior prob. over W : $\pi(w)$

$|w|_{\pi}$ = "description length" of w

$$\text{i.e. } \sum_{w \in W} \pi(w) = 1 \quad \pi(w) \geq 0 \quad \forall w$$

$$\triangleq \log_2 \frac{1}{\pi(w)}$$

Occam's bound

for any fixed P ; with prob $\geq 1-\delta$ over training set $D_n \sim P^{\otimes n}$

$$\forall w \in W \quad L_P(w) \leq \hat{L}_P(w) + \frac{1}{\sqrt{2n}} \Omega_{\pi}(w; \delta)$$

$$\text{where } \Omega_{\pi}(w; \delta) \triangleq \sqrt{(2n\delta) \underbrace{|w|_{\pi}}_{\text{complexity measure}} + \ln \frac{1}{\delta}}$$

(*) bound is only useful for distribution P st. $|w_P^*|_{\pi}$ is small
 $\hookrightarrow w_P^* = \arg \min_{w \in W} L_P(w)$

$$|w|_{\pi} = \log_2 \frac{1}{\pi(w)}$$

$$\text{if } \pi(w) \propto \exp(-\|w\|^2)$$

$$\text{then } |w|_{\pi} = \|w\|^2 + \text{const.}$$

note: 0-1 loss assumption appears in constants of Chernoff bound

proof: use 3 things

1) Chernoff bound.
(concentration inequality)

2) union bound

3) "Kraft's inequality"

"bad"

$$P\{D_n: \hat{L}_n(w) \leq L(w) - \epsilon\} \leq \exp(-2n\epsilon^2) \quad \forall \epsilon \geq 0$$

$$P\{\exists x \text{ st. } p(x) \text{ is true}\} \leq \sum_x P\{p(x) \text{ is true}\}$$

$$\sum_w 2^{-|w|_{\pi}} \leq 1$$

$$\Rightarrow 1 - \epsilon > \hat{L}_n$$

we say w is "bad" if bound fails

$$\text{bad}(w) = \mathbb{1}\left\{ L(w) > \hat{L}_n(w) + \underbrace{\frac{\epsilon_n(w)}{\sqrt{2n}}}_{\epsilon_n(w)} \right\}$$

$L - \epsilon > \hat{L}_n$

using Chernoff, $\hat{L}_n(w) \leq L(w) - \epsilon_n(w)$ with small prob.

$$\begin{aligned} \mathbb{P}\{\text{bad}(w)\} &\leq \exp(-2n\epsilon_n(w)^2) = \exp\left(-2n \frac{1}{2n} \left((\ln 2) |w|_n + \ln \frac{1}{8}\right)\right) \\ &= 8 2^{-|w|_n} \end{aligned}$$

using union bound

$$\mathbb{P}\{\exists w: \text{bad}(w)\} \leq \sum_w \mathbb{P}\{\text{bad}(w)\} \leq \sum_w 8 2^{-|w|_n} \leq 8 //$$

Surrogate Loss:

NP hard to minimize $\hat{L}_n(w)$; replace with $\hat{S}_n(w)$ which is "surrogate"
 e.g. hinge loss
 • log-loss
 ...