

Recap Lecture 3 (Chapter 4)

Hoeffding's ineq. $\underbrace{a \leq \theta_i \leq b}_{\text{bounded}}, \mathbb{E}[\theta_i] = \mu,$

$$\mathbb{P}\left(\left|\frac{1}{m} \sum_{i=1}^m \theta_i - \mu\right| > \varepsilon\right) \leq 2e^{-\frac{2m\varepsilon^2}{(b-a)^2}}$$

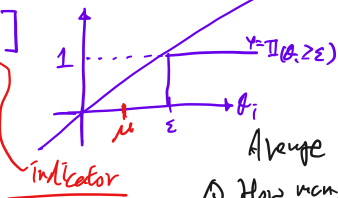
Concentration of measure

The more you know about the R.V. θ_i , the tighter concentration you can prove.

① [least information, Markov's inequality]

If $\theta_i \geq 0$ w.p.1 then $\mathbb{P}(\theta_i \geq \varepsilon) \leq \frac{\mu}{\varepsilon}$ no concentration, why? Tight?
 and $\mathbb{P}(\frac{1}{m} \sum_{i=1}^m \theta_i \geq \varepsilon) \leq \frac{\mu}{\varepsilon} = O(1)$ if $\theta_i = \frac{1}{\varepsilon}$

Proof: $\mathbb{P}(\theta_i \geq \varepsilon) = \mathbb{E}[\mathbb{I}(\theta_i \geq \varepsilon)]$
 $\int_{\varepsilon}^{\infty} f_{\theta_i}(t) dt$



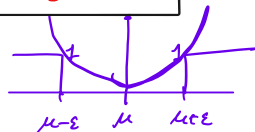
$\theta_i \geq 0 \rightarrow \leq \mathbb{E}[\frac{1}{\varepsilon} \theta_i]$
 $= \frac{\mu}{\varepsilon}$

Average age = 30, $m = 1,000,000$
 Q. How many over 70? $\frac{3}{4}$.

② [2nd order statistics, Chebyshev's inequality]

If $\text{Var}(\theta_i) \leq \sigma^2$ then $\mathbb{P}(|\theta_i - \mu| \geq \varepsilon) \leq \frac{\sigma^2}{\varepsilon^2}$ variance $\times \frac{1}{m}$
 and $\mathbb{P}(|\frac{1}{m} \sum_{i=1}^m \theta_i - \mu| \geq \varepsilon) \leq \frac{\sigma^2}{m\varepsilon^2} = O(\frac{1}{m})$

Proof: $\mathbb{P}(|\theta_i - \mu| \geq \varepsilon) = \mathbb{E}[\mathbb{I}(\theta_i \leq \mu - \varepsilon \text{ or } \theta_i \geq \mu + \varepsilon)]$



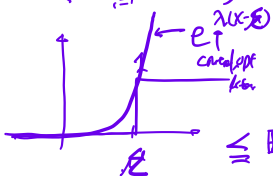
Guided by the fact that $\leq \mathbb{E}[\frac{1}{\varepsilon^2} (\theta_i - \mu)^2]$
 $\text{Var} \leq \sigma^2 \leq \frac{\sigma^2}{\varepsilon^2}$

③ [bounded domain, Hoeffding's inequality]

If $a \leq \theta_i \leq b$ then $\mathbb{E}[e^{\frac{\lambda(\theta_i - \mu)}{m}}] \leq e^{\frac{\lambda^2(b-a)^2}{8m^2}}$ (*)
 and $\mathbb{P}(|\frac{1}{m} \sum_{i=1}^m \theta_i - \mu| \geq \varepsilon) \leq 2 \cdot \exp(-\frac{2m\varepsilon^2}{(b-a)^2}) = o(e^{-m\varepsilon^2})$

Generic Recipe

$\mathbb{P}(\frac{1}{m} \sum_{i=1}^m \theta_i - \mu \geq \varepsilon) = \mathbb{E}[\mathbb{I}(\frac{1}{m} \sum_{i=1}^m \theta_i - \mu \geq \varepsilon)]$



$\leq \mathbb{E}[e^{\lambda(\frac{1}{m} \sum_{i=1}^m \theta_i - \mu) - \varepsilon}]$
 $\leq e^{-\lambda\varepsilon} \prod_{i=1}^m \mathbb{E}[e^{\lambda \frac{\theta_i - \mu}{m}}]$

$\leq e^{-\lambda\varepsilon} \prod_{i=1}^m e^{\frac{\lambda^2(b-a)^2}{8m^2}}$

opt over $\lambda \rightarrow \leq \exp(-\lambda\varepsilon + \frac{\lambda^2(b-a)^2}{8m})$

$-\varepsilon + \frac{2\lambda(b-a)^2}{8m} = 0 \rightarrow \leq \exp(-\frac{2m\varepsilon^2}{(b-a)^2})$

$\lambda^* = \frac{4\varepsilon m}{(b-a)^2}$

do the same for

$\frac{1}{m} \sum_{i=1}^m \theta_i < \mu - \varepsilon$

Special to Hoeffding's

(*) $\mathbb{E}[e^{\frac{\lambda(\theta_i - \mu)}{m}}] \leq e^{\frac{\lambda^2(b-a)^2}{8m^2}}$

$\mathbb{E}[e^{\frac{\lambda(\theta_i - \mu)}{m}}] \leq \max_{a \leq \theta_i \leq b} \mathbb{E}[e^{\frac{\lambda(\theta_i - \mu)}{m}}]$

convexity of $\exp(\cdot)$ $\rightarrow \leq \max_{a \leq \theta_i \leq b} P_a e^{\frac{\lambda(a-\mu)}{m}} + P_b e^{\frac{\lambda(b-\mu)}{m}}$

$P_a + P_b = 1$
 $a \leq \mu \leq b$
 $\leq e^{-\frac{\lambda\mu}{m}} \max_{a \leq \theta_i \leq b} P_a e^{\frac{\lambda a}{m}} + P_b e^{\frac{\lambda b}{m}}$

$\leq e^{-\frac{\lambda\mu}{m}} \cdot \exp\left[\frac{\lambda a P_a + P_b \lambda}{m} + \frac{(b-a)^2 \lambda^2}{8m^2}\right]$

$\leq \exp\left(\frac{(b-a)^2 \lambda^2}{8m^2}\right)$

* Chernoff's Bound

$\theta_1 \dots \theta_m$ independent Bernoulli dist. $\theta_i \in \{0, 1\}$ w.p. P_i $1-P_i$

but not identically distributed

$$IP\left(\frac{1}{m} \sum_{i=1}^m \theta_i > (1+\delta) \frac{1}{m} \sum_{i=1}^m P_i\right) \leq e^{-\underbrace{\sum_{i=1}^m P_i}_{O(m)} \underbrace{\left(\frac{\delta^2}{2+\frac{\delta^2}{3}}\right)}_{2 \text{ phase}}} \leq e^{-\frac{\delta^2}{2+\frac{\delta^2}{3}} \sum_{i=1}^m P_i}$$

multiplicative $\approx (1+\delta) \lg(1/\delta) - \delta \geq \frac{\delta^2}{2+\frac{\delta^2}{3}}$

Recap 2.

lecture 2 (ch 2,3): PAC learning. Realizable $\mathcal{Y} = \{0,1\}$ $|\mathcal{H}| < \infty$

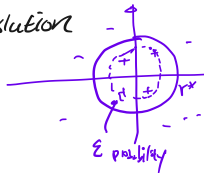
lecture 3 (ch 4): Agnostic PAC learning w.p. $1-\delta$.
 Technique: Union Bound.
 $L_D(h) \leq \min_{h' \in \mathcal{H}} L_D(h') + \epsilon$
 Technique: concentration of measure.

Q3.3 $\mathcal{X} = \mathbb{R}^2, \mathcal{Y} = \{0,1\}, \mathcal{H} = \{h_r : r \in \mathbb{R}_+\}$

$$h_r(x) = \mathbb{I}_{\|x\| \leq r}$$

Prove that $\mathcal{H}$ is PAC learnable (assuming realizability), with $m_{\mathcal{H}}(\epsilon, \delta) \leq \left\lceil \frac{\lg 1/\delta}{\epsilon^2} \right\rceil$

solution



fix $D, f_{\mathcal{H}}$, define $\mathcal{H}_{\text{Bad}} = \{r : r \leq \tilde{r}, D_{\text{Bad}}(h_r) > \epsilon\}$

Algo: return smallest circle that includes all +.

$$IP_S(\text{Algo} < \tilde{r}) \leq (1-\epsilon)^m \leq e^{-\epsilon m}$$

$$\text{for } m \geq \frac{1}{\epsilon} \lg 1/\delta \leq 6.$$

Chapter 5.

the reason above problem has low sample complexity for PAC learning is that we had a prior knowledge [ground truth D has a circular decision boundary]
 $\mathcal{H}$ is collection of circular decision boundaries

Q. is such prior knowledge necessary? yes.

Learning Problem, $\mathcal{X}, \mathcal{Y}$ fixed

NATURE

D

S

LEARNER

$\mathcal{H}, \text{Algo}$

$\Downarrow$

h

Goal $\downarrow$
 $L_D(h)$

Q. Can there be a universal algorithm Algo

that is PAC learnable for all problem instances D ?

Q. is prior knowledge necessary for PAC learning?

Realizable w.p. $1-\delta$ (implicit rec 3) $\rightarrow$ should we use Agnostic PAC, otherwise.

How much prior knowledge is good to assume?

a lot

- learner knows $D, h^* \rightarrow \mathcal{H} = \{h^*\}$ achieves minimum risk
 e.g. knows optimal

- learner knows a class

- lecture 2: realizable + finite $\mathcal{H} \rightarrow$ PAC learnable $m \approx \frac{\lg |\mathcal{H}|}{\epsilon^2}$

- lecture 3: e.g. $\mathcal{H} = \{h_r : \mathbb{I}_{\|x\| \leq r}\} \rightarrow$ PAC learnable $m \approx \frac{\lg 1/\delta}{\epsilon^2}$

- no prior knowledge + finite $\mathcal{H} \rightarrow$ Agnostic PAC learnable $m \approx \frac{\lg |\mathcal{H}|}{\epsilon^2}$

- learner knows nothing: no-free-lunch Theorem.

no

No-free-lunch theorem

Theorem. let Algo be any learning algorithm. $\mathcal{Y} = \{0,1\}$, any $\mathcal{X}$, $m \leq \frac{1}{2\epsilon}$

Then exists a distribution D over $\mathcal{X} \times \mathcal{Y}$ s.t.

1. $\exists f: \mathcal{X} \rightarrow \{0,1\}$ with $L_D(f) = 0$

Constructive Proof 2. $P_S(L_D(A(S)) \geq \frac{1}{8}) \geq \frac{1}{4}$

Negative Result.

proof sketch: $|X| \geq 2m$, $C \subseteq X$, and $|C|=2m$, and we observe random S of m samples.

$m=2$

$|C|=4, |H|=16$

$h_0(1000)$
 $h_1(0001)$
 $h_2(0010)$
 $h_3(0100)$
 $h_4(1000)$
 $h_5(0001)$
 $h_6(0010)$
 $h_7(0100)$
 $h_8(1000)$
 $h_9(0001)$
 $h_{10}(0010)$
 $h_{11}(0100)$
 $h_{12}(1000)$
 $h_{13}(0001)$
 $h_{14}(0010)$
 $h_{15}(0100)$

ox_1
 ox_2
 ox_3
 ox_4

$|H| = 2^{2m}$ all possible ^{ways to} labels.
 $= \{h_0, h_1, \dots, h_{2^m-1}\}$

for any Algo, $\exists$ s.t. one cannot do better than random guess on $\frac{1}{2}$ of the samples not observed.

$\rightarrow \mathbb{E}_S [L_D(A(S))] \geq \frac{1}{4}$
 $\rightarrow$ claim by Markov's Ineq.

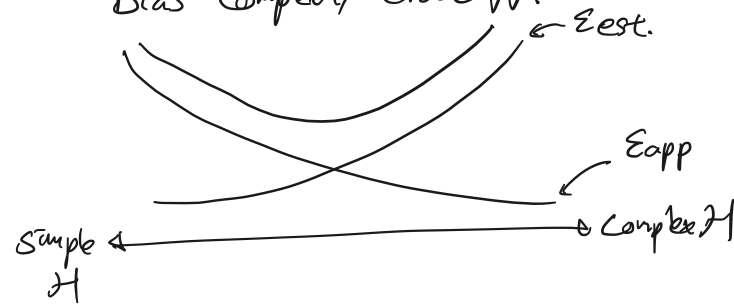
$\underbrace{\mathbb{E}_S [L_D(A(S))]}_{\text{positive R.V.}} \geq \frac{1}{4}$
 $\rightarrow$ $\frac{1}{2}$ of samples error $\frac{1}{2}$.

Corollary. X is infinite & H is all possible functions from X to $\{0,1\}$ then H is not PAC learnable.

Q. How do we choose the right hypothesis class?

$$L_D(h_S) = \underbrace{(L_D(h_S) - \min_{h' \in H} L_D(h'))}_{\substack{\uparrow \\ \text{ERM} \\ \epsilon_{\text{est}} = \text{Estimation error} \\ \cdot \text{ goes down with } m \\ \cdot \text{ depends on complexity } H}} + \underbrace{\min_{h' \in H} L_D(h')}_{\substack{\epsilon_{\text{app}} = \text{approximation error} \\ \cdot \text{ does not depend on } m \\ \cdot \text{ property of } D, H, \text{ complex}}}$$

Bias-Complexity Tradeoff.



chapter 6.

realizable + finite $H \rightarrow$ PAC learnable (lec 2)
 finite $H \rightarrow$ Agnostic PAC learnable (lec 3).

realizable + infinite $H = \{h_r: \mathbb{I}_{(x \in r)}\} \rightarrow$ PAC learnable (lect)

$C \subseteq X, |C| \geq 2m, H$ contains all possible 2^m labeling $\rightarrow$ no learning (lec 4)

Q. What class H is PAC learnable? Vladimir Vapnik & Alexey Chervonenkis 1970

Definition. (Restriction of H to C)

H is class of functions $X \rightarrow \{0,1\}$, and $C = (c_1, \dots, c_m) \subseteq X$,

the Restriction of H to C

$$H_C \triangleq \{ \tilde{h}: C \rightarrow \{0,1\} \mid h \in H \}$$

$c_i \mapsto h(c_i)$

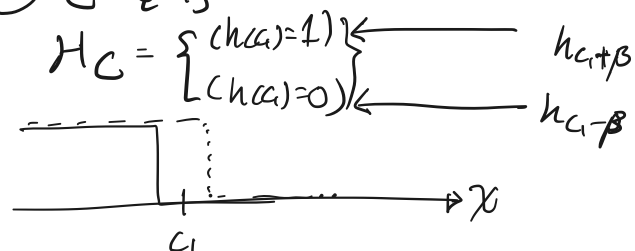
equivalently we can represent h by a vector $(h(c_1), \dots, h(c_m))$.

Definition (Shattering)

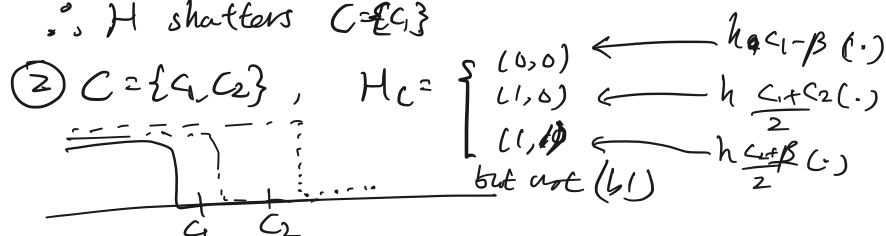
A hypothesis class H shatters a set $C \subseteq X$ if the restriction of H to C is the set of all functions from C to $\{0, 1\}$. That is $|H_C| = 2^{|C|}$.

Example. $H = \{h_a : \mathbb{R} \rightarrow \mathbb{R} = \mathbb{I}(x \leq a)\}$ class of all threshold functions.

① $C = \{c_1\}$



$\therefore H$ shatters $C = \{c_1\}$



$\therefore H$ does not shatter $C = \{c_1, c_2\}$

Corollary (No-free-lunch thm restated)

for some $H, X, Y = \{0, 1\}$, if $\exists C \subseteq X, |C| = 2m$ & shattered by H , then for any learning algorithm, $\exists D$ & $h \in H$ st.

$$L_D(h) = 0 \text{ but } \mathbb{P}(L_D(A(S)) \geq \frac{1}{8}) \leq \frac{1}{7}.$$

Definition (VC-dimension)

The VC-dimension of H , $\triangleq \text{VCdim}(H)$,

is the maximal size of $C \subseteq X$ that can be shattered by H .

$\rightarrow$ NO set of size $|C|+1$ can be shattered.

Example. need to show that $\exists C$ of $|C|=d$ shattered say,

$\forall C$ of $|C|=d+1$, not shattered.

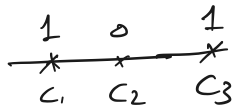
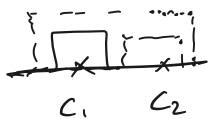
Exmple (Threshold: $H = \{\mathbb{I}(x \leq a)\}$)

$\exists C = \{c_1\}$ can be shattered.

all $C = \{c_1, c_2\}$ can not be shattered.

$$VCdim(H) = 1.$$

Example (Intervals: $H = \{ \Pi(a \leq x \leq b) : a \leq b \in \mathbb{R} \}$)

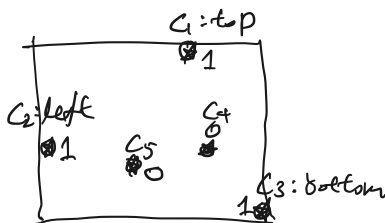
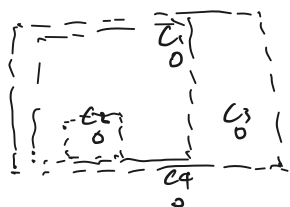


$$VCdim(H) = 2.$$

existence of shattered set all set not shattered.

Example (axis aligned rectangles: $H = \{ \Pi(a_1 \leq x_1 \leq b_1, a_2 \leq x_2 \leq b_2) : a_1 \leq b_1, a_2 \leq b_2 \}$)

$$VCdim(H) = 4$$



existence of shattered set. all set of 5 not shattered.

Example (finite classes: $|H| < \infty$)

shattering requires at least $2^{|C|}$ hypothesis. so

$$2^{|C|} \leq |H|$$

$$\rightarrow |C| \leq \log_2 |H|$$

$$\rightarrow VCdim \leq \log_2 |H|$$

tight in the worst case
VCdim can be much smaller.
c.f. finite & threshold.

Example (Number of parameters).

above examples have # parameters match VCdim,

Not always true: $H = \{h_\theta : \theta \in \mathbb{R}\}, VCdim(H) = 200$ [HW1 Problem 4]
" $\lceil \log \sin(\theta x) \rceil$

Theorem [Fundamental Theorem of Statistical Learning]

Let H be a hypothesis class, from $\mathcal{X}$ to $\{0, 1\}$,
and the loss function is a 0-1 loss, then the
following are equivalent.

1. H has the uniform convergence property
2. Any ERM rule is a successful agnostic PAC learner of H .
3. H is agnostic PAC learnable.
4. H is PAC learnable
5. Any ERM rule is a successful PAC learner for H

6. H has a finite VC-dimension.

Theorem [Fundamental theorem, Quantitative version]

$VCdim(H) = d$, then

1. H has uniform convergence property with

$$C_1 \cdot \frac{d + \log \frac{1}{\delta}}{\epsilon^2} \leq m_H^{VC}(\epsilon\delta) \leq C_2 \cdot \frac{d + \log \frac{1}{\delta}}{\epsilon^2}$$

2. H is agnostic PAC learnable with

$$C_1 \cdot \frac{d + \log \frac{1}{\delta}}{\epsilon^2} \leq m_H(\epsilon\delta) \leq C_2 \cdot \frac{d + \log \frac{1}{\delta}}{\epsilon^2}$$

3. H is PAC learnable with,

$$C_1 \cdot \frac{d + \log \frac{1}{\delta}}{\epsilon} \leq m_H(\epsilon\delta) \leq C_2 \cdot \frac{d \log \frac{1}{\epsilon} + \log \frac{1}{\delta}}{\epsilon}$$