

Relaxing the realizability assumption

Data generating distribution $\mathcal{D}$ is over $\mathcal{X}, \mathcal{Y}$: $\mathbf{z} = (x, y) \sim \mathcal{D}$

True error: $L_{\mathcal{D}}(h) \triangleq \mathbb{P}_{\mathcal{D}}(h(x) \neq y)$

Empirical error: $L_S(h) \triangleq \frac{|\{i: i \in [n]: h(x_i) \neq y_i\}|}{n}$

Goal: to find h that minimizes $L_{\mathcal{D}}(h)$

Bayes Optimal Predictor has lower error than any other predictor [HW1] Prob 2

$$f_{\mathcal{D}}(x) = \begin{cases} 1 & \text{if } \mathbb{P}_{\mathcal{D}}(y=1|x) \geq 1/2 \\ 0 & \text{otherwise} \end{cases}$$

but requires knowledge of $\mathcal{D}$.

Agnostic PAC learnability.

A class of hypotheses, $\mathcal{H}$, is Agnostic PAC learnable

if $\exists m_{\mathcal{H}}$ and a learning algorithm s.t. for $\epsilon, \delta, \mathcal{D}$

running Algo $m \geq m_{\mathcal{H}}(\epsilon, \delta)$ samples returns h s.t. w.p $\geq 1-\delta$

$$L_{\mathcal{D}}(h) \leq \min_{h' \in \mathcal{H}} L_{\mathcal{D}}(h') + \epsilon$$

Rebinary Binary Classification

Multi-class classification $\mathcal{Y} = [K]$

Regression, $\mathcal{Y} \subseteq \mathbb{R}$

loss $\ell_f(x, y) = (h(x) - y)^2$

Risk function $L_{\mathcal{D}}(h) \triangleq \mathbb{E}_{\mathbf{z} \sim \mathcal{D}}[\ell(h, \mathbf{z})]$

Empirical Risk function $L_S(h) \triangleq \frac{1}{n} \sum_{i=1}^n \ell(h, (x_i, y_i))$

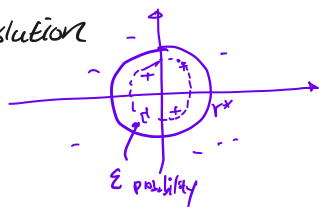
Q3.3 $\mathcal{X} = \mathbb{R}^2, \mathcal{Y} = \{0, 1\}, \mathcal{H} = \{h_r: r \in \mathbb{R}_+\}$

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

Prove that $\mathcal{H}$ is PAC learnable (assuming realizability), with

$$m_{\mathcal{H}}(\epsilon, \delta) \leq \left\lceil \frac{\log 1/\delta}{\epsilon} \right\rceil$$

Solution



fix $\mathcal{D}, f_{\mathcal{D}}$, define $\mathcal{H}_{\text{Bad}} = \{r: r \leq r^*, \mathcal{D}(\{x: \|x\| \leq r\}) \geq \epsilon\}$

Algo: return smallest circle that includes all +.

$$\mathbb{P}_{\mathcal{D}}(\text{Algo} \leq r) \leq (1-\epsilon)^m \leq e^{-\epsilon m}$$

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

chapter 4

Uniform Convergence is sufficient for learnability

Definition. (ϵ -representative sample)

Training set S is ϵ -representative if

$$\forall h \in \mathcal{H}, |L_S(h) - L_{\mathcal{D}}(h)| \leq \epsilon$$

Lemma. Assume we are given a training set S that is ϵ -representative. Then, any output of ERM $h_S \in \arg \min_{h \in \mathcal{H}} L_S(h)$, satisfies

Proof strategy
Show $\forall h \in \mathcal{H}$
- if $m \geq \frac{1}{\epsilon^2} \log \frac{1}{\delta}$
Use Lemma to show ERM is

$$L_{\mathcal{D}}(h_S) \leq \min_{h \in \mathcal{H}} L_{\mathcal{D}}(h) + \epsilon$$

Proof. by triangular inequality i.e. $a \leq b + |a-b|$

$$\begin{aligned}
 L_D(h_S) &\leq |L_D(h_S) - L_S(h_S)| + L_S(h_S) \\
 &\leq \frac{\epsilon}{2} + \underbrace{L_S(h_S) - L_S(h)}_{\leq 0 \sim \text{ERM}} + L_S(h) \\
 &\leq \frac{\epsilon}{2} + 0 + |L_S(h) - L_D(h)| + L_D(h) \\
 &\leq \frac{\epsilon}{2} + 0 + \frac{\epsilon}{2} + L_D(h) \quad \text{for } h \in \mathcal{H}
 \end{aligned}$$

To show ERM is agnostic PAC Learner, we are left to show that with probability $1-\delta$, S is ϵ -representative.

Definition (Uniform Convergence)

A hypothesis class $\mathcal{H}$ has the uniform convergence property,

if $\exists m_{\mathcal{H}}^{UC}$ s.t. for every ϵ, δ and every D , if $m \geq m_{\mathcal{H}}^{UC}(\epsilon, \delta)$

the S is ϵ -representative with probability at least $1-\delta$.

Corollary 1 $\mathcal{H}$ is agnostically PAC learnable with sample complexity

$$m_{\mathcal{H}}(\epsilon, \delta) \leq m_{\mathcal{H}}^{UC}(\frac{\epsilon}{2}, \delta)$$

And ERM is the agnostic PAC Learner.

Next, we want to show finite $\mathcal{H}$ is agnostic PAC learnable.

↑ show uniform convergence.

Claim. Any finite $\mathcal{H}$, i.e., $|\mathcal{H}| < \infty$, is uniformly convergent with

$$m_{\mathcal{H}}^{UC}(\epsilon, \delta) \leq \left\lceil \frac{\log \frac{2 \cdot |\mathcal{H}|}{\delta}}{\epsilon} \right\rceil$$

Corollary 2. ERM is agnostic PAC learner with sample complexity

$$m_{\mathcal{H}}(\epsilon, \delta) \leq m_{\mathcal{H}}^{UC}(\frac{\epsilon}{2}, \delta) \leq \left\lceil \frac{2 \log \frac{2 \cdot |\mathcal{H}|}{\delta}}{\epsilon} \right\rceil$$

↑ Corollary 1.
 ↑ claim

We are left to prove the Claim.

= Q. What is the sample size that guarantees, assuming $\sum_{k=1}^D l(k, z_i) \leq 1$

$|L_S(h) - L_D(h)| \leq \epsilon$ w.p $1-\delta$?

n.t.s (need to show) $P_S(\{S: \forall h \in \mathcal{H}, |L_S(h) - L_D(h)| \leq \epsilon\}) \geq 1-\delta$

n.t.s $P_S(\{S: \exists h \in \mathcal{H}, |L_S(h) - L_D(h)| > \epsilon\}) < \delta$

$$\leq \sum_{h \in \mathcal{H}} P_S(\{S: |L_S(h) - L_D(h)| > \epsilon\}) \quad (*)$$

↑
Union bound

$$L_S(h) = \frac{1}{m} \sum_{i=1}^m l(h, z_i) \leftarrow \text{average of } z_i = (x_i, y_i) \text{ i.i.d R.V.}$$

$$L_D(h) = \mathbb{E}_{z \sim D}[l(h, z)] = \mathbb{E}[L_S(h)] \leftarrow \text{expectation of each term}$$

We need to show that the measure of R.V. $L_S(h)$

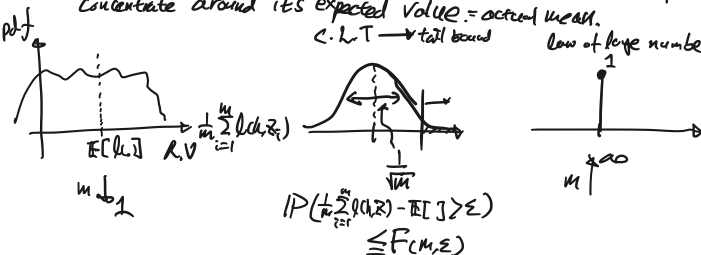
= distribution

concentrates around its mean.

* most important tool in statistical analysis for learning = concentration of measure

Q. How fast (in terms of # of samples) does the measure of empirical mean concentrate around its expected value = actual mean.

C.L.T → tail bound law of large numbers



* Hoeffding's Inequality

Let $\theta_1, \dots, \theta_m$ be a set of i.i.d random variables with $\mu \triangleq \mathbb{E}[\theta_i]$, satisfying

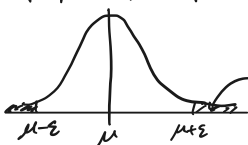
$$\mathbb{P}(a \leq \theta_i \leq b) = 1, \text{ for all } i \in [m] \quad [\text{Boundedness}]$$

then for any $\varepsilon > 0$,

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

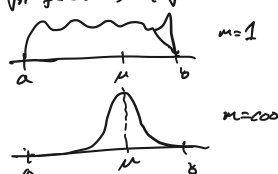
first example of $\leftarrow$ concentration of measure

for fixed m , varying ε



$$e^{-m\varepsilon^2/2} \cdot C = \frac{2}{(b-a)^2}$$

for fixed ε , varying m



Letting $\theta_i = f(h, z_i)$, and assuming $a \leq f(h, z_i) \leq b$, and $\mu = \mathbb{E}[f(h, z)]$

$$\mathbb{P}_S(|L_S(h) - L_0(h)| > \varepsilon) = \mathbb{P}\left(\left|\frac{1}{m} \sum_{i=1}^m \theta_i - \mu\right| > \varepsilon\right)$$

$$\text{Hoeffding's} \rightarrow \leq 2 \cdot e^{-2m\varepsilon^2}$$

$$(*) \Rightarrow \mathbb{P}_S(\{h \in \mathcal{H} : |L_S(h) - L_0(h)| > \varepsilon\}) \leq \sum_{h \in \mathcal{H}} 2e^{-2m\varepsilon^2} = 2|\mathcal{H}| \cdot e^{-2m\varepsilon^2}$$

$$\text{letting } m(\varepsilon, \delta) \geq \frac{\log(2|\mathcal{H}|/\delta)}{2\varepsilon^2}, \leq \delta$$

Summary

$$\text{if } m \geq \frac{\log(2|\mathcal{H}|/\delta)}{2\varepsilon^2}$$

then $\mathcal{H}$ has uniform convergence

(ε, δ)

Lemma: Uniform Convergence $\Leftrightarrow m_{\mathcal{H}}^{uc}(\varepsilon, \delta) \geq m_{\mathcal{H}}(\varepsilon, \delta)$

(ε, δ) -PAC learnable with ERM

$$\text{if } m \geq \left\lceil \frac{2 \cdot \log(2|\mathcal{H}|/\delta)}{\varepsilon^2} \right\rceil \geq m_{\mathcal{H}}^{uc}(\varepsilon/2, \delta) \geq m_{\mathcal{H}}(\varepsilon, \delta)$$

then $\mathcal{H}$ is (ε, δ) -PAC learnable with ERM.

End of Lecture 3.

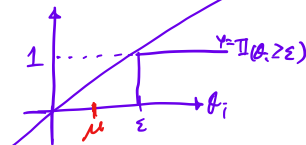
Concentration of measure

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

① [least information, Markov's inequality]

$$\text{If } \theta_i \geq 0 \text{ up to } \infty \text{ then } \mathbb{P}(\theta_i \geq \varepsilon) \leq \frac{\mu}{\varepsilon} \quad \text{and} \quad \mathbb{P}\left(\frac{1}{m} \sum_{i=1}^m \theta_i \geq \varepsilon\right) \leq \frac{\mu}{\varepsilon} = O(1)$$

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

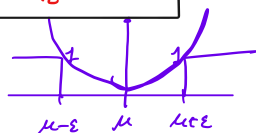


$$\theta_i \geq 0 \rightarrow \leq \mathbb{E}\left[\frac{1}{\varepsilon} \theta_i\right] = \frac{\mu}{\varepsilon}$$

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

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

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



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

③ [bounded domain, Hoeffding's inequality]

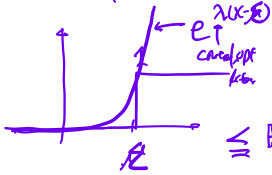
$$\text{If } a \leq \theta_i \leq b \text{ then } \mathbb{E}\left[e^{\lambda(\theta_i - \mu)}\right] \leq e^{\frac{\lambda^2(b-a)^2}{2}} \quad (*)$$

$$\text{and } \mathbb{P}\left(\left|\frac{1}{m} \sum_{i=1}^m \theta_i - \mu\right| \geq \varepsilon\right) \leq 2 \cdot \exp\left(-\frac{2m\varepsilon^2}{(b-a)^2}\right) = o(e^{-m\varepsilon^2})$$

Generic Recipe

Special to Hoeffding's

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



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

$$\leq e^{-\lambda \varepsilon} \prod_{i=1}^m e^{\lambda \frac{\theta_i - \mu}{m}}$$

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

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

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

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

do the same for

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

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

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

$$\text{Convexity of exp} \rightarrow \leq \max_{P_a, P_b} P_a e^{\frac{\lambda(a-\mu)}{m}} + P_b e^{\frac{\lambda(b-\mu)}{m}}$$

$$\text{Part 1: } a=P_a, b=P_b \rightarrow \leq e^{-\frac{\lambda \mu}{m}} \cdot \max_{P_a, P_b} P_a e^{\frac{\lambda a}{m}} + P_b e^{\frac{\lambda b}{m}}$$

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

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

* Chernoff's Bound

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

but not identically distributed

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

$$\geq \frac{\delta^2}{2+\frac{2}{3}\delta}$$