

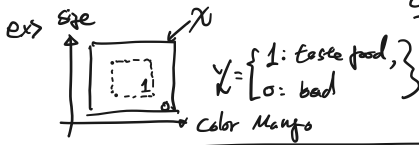
# Statistical learning framework.

learner:

Domain  $\mathcal{X}$ : set of objects you wish to label.

Label set  $\mathcal{Y}$ : the set of possible labels.

A prediction rule,  $h: \mathcal{X} \rightarrow \mathcal{Y}$ : used to label future examples.  
also called predictor, hypothesis, classifier.



## Batch learning

Learner's input: Training data,  $S = \{(\mathbf{x}_i, y_i)\}_{i=1}^m \in (\mathcal{X} \times \mathcal{Y})^m$   
Set indexed by  $i \in \{1, \dots, m\}$       Set of  $m$ -tuple of consistent labels

Learner's output:

- A prediction rule,  $h: \mathcal{X} \rightarrow \mathcal{Y}$

Goal: to get correct labels on future predictions

- suppose  $f$  is a correct classifier, then we want  $f \approx h$

- One approach: define error

$$L_{D,f}(h) \triangleq \mathbb{P}_{\mathbf{x} \sim D} [h(\mathbf{x}) \neq f(\mathbf{x})]$$

$D$  is some unknown distribution.

$\rightarrow \mathbb{P}_{\mathbf{x} \sim D}(A)$  is probability that  $\mathbf{x}$  falls in set  $A \subseteq \mathcal{X}$

- Goal: find  $h$  s.t.  $L_{D,f}(h)$  is small

We need some assumption on relation between  $D$  and  $f$ .  
simple data generation

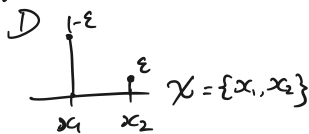
- Independent and identically distributed (i.i.d.):  
each  $\mathbf{x}_i$  is sampled independently from  $D$

- Realizable:  
for every  $i \in [m]$ ,  $y_i = f(\mathbf{x}_i)$   
 $\uparrow$   
 $\{1, \dots, m\}$

Claim: this can only be approximately correct

i.e., we cannot find  $h$  s.t.  $L_{D,f}(h) = 0$  for all  $D, f$ .

Proof:



$$\cdot \mathbb{P}(\text{no } x_2 \text{ in } S_m) = (1-\epsilon)^m \approx \underset{0}{e^{-\epsilon m}} \approx 1.$$

letting  $\epsilon \ll \frac{1}{m}$ ,  $m \uparrow$

$\Rightarrow$  you don't see any  $x_2$ , and cannot learn  $f(x_2)$ .

$\Rightarrow$  Relax goal to be approximate:  $L_{D,f}(h) \leq \epsilon$ .

Claim: this can only be Probably correct.

i.e., no algorithm can guarantee  $L_{D,f}(h) \leq \epsilon$  w.p.1  
almost surely

Recall that input is random

$\rightarrow$  there is a (small) chance that we see the same example again

$\Rightarrow$  Relax: we let algorithm fail with probability  $\delta$ ,

for some  $\delta \in (0, 1)$  chosen by the user

where the randomness is over the sampling of  $S$ .

## Probably Approximately Correct (PAC) learning

- learner does not know  $D$  or  $f$
- learner chooses accuracy parameter  $\epsilon$  and confidence parameter  $\delta$
- learner gets training data  $S$  of size  $m(\epsilon, \delta)$
- learner outputs a hypothesis  $h$  s.t.

$$\mathbb{P}(L_{D,f}(h) \leq \epsilon) \geq 1 - \delta$$

- then the learner is Probably Approximately Correct.  
w.p.  $1 - \delta$  up to  $\epsilon$

## Empirical Risk Minimization (ERM)

- Training error, empirical error, empirical risk

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

$$[m] = \{1, \dots, m\}$$

$$\min_h L_S(h)$$

- What can go wrong? overfitting

$$h_S^{(x)} = \begin{cases} y_i & \text{if } x = x_i \\ 0 & \text{otherwise} \end{cases} \rightarrow L_S(h_S) = 0.$$

Continuous  $D$ ,  $L_{D,f}(h_S) = \frac{1}{2}$  because  $y_i$ .

- Hypothesis class  $\mathcal{H} \rightarrow$  common measure of complexity  $|\mathcal{H}|$

$$h_S = \text{ERM}_{\mathcal{H}}(S) \triangleq \arg \min_{h \in \mathcal{H}} L_S(h)$$

Assumption 1. (Realizability)  $\exists h^* \in \mathcal{H}$  s.t.  $L_{D,f}(h^*) = 0$   
 $\Rightarrow L_S(h^*) = 0$  w.p. 1

Assumption 2. (iid samples)  $\begin{cases} x_i \sim \text{iid } D \\ y_i = f(x_i) \end{cases} \Rightarrow S \sim D^m$

Claim. For a finite hypothesis class  $\mathcal{H}$ , and some  $\delta \in (0, 1)$ , and  $\epsilon > 0$   
 if  $m \geq \frac{\log(|\mathcal{H}|/\delta)}{\epsilon}$ , then for any  $f$ , and  $D$ ,  
 that are realizable,

$$L_{D,f}(h_S) \leq \epsilon \quad \text{w.p. } \geq 1 - \delta.$$

$$\rightarrow \mathbb{P}(L_{D,f}(h_S) > \epsilon) \leq \delta$$

- $\text{ERM}_{\mathcal{H}}$  is Probably Approximately Correct learning.

Define. "bad" hypotheses in  $\mathcal{H}$ .

$$\mathcal{H}_B \triangleq \{h \in \mathcal{H} : L_{D,f}(h) > \epsilon\}$$

$$\mathbb{P}_S(L_{D,f}(h_S) > \epsilon) \leq \mathbb{P}_S\left(\bigcup_{h \in \mathcal{H}_B} \{S : L_S(h) = 0\}\right) \quad (*)$$

$$\leq \sum_{h \in \mathcal{H}_B} \mathbb{P}_S(L_S(h) = 0)$$

Union bound

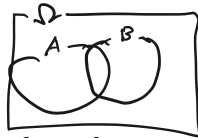
$$\leq |\mathcal{H}_B| \max_{h \in \mathcal{H}_B} \prod_{i=1}^m \mathbb{P}(h(x_i) = y_i)$$

$$\begin{aligned} & \leq |\mathcal{H}_B| \cdot (1-\epsilon)^m \\ & \uparrow \\ & \text{def } \mathcal{H}_B \\ & \leq |\mathcal{H}| \cdot e^{-\epsilon m} \leq \delta \\ & \uparrow \\ & \mathcal{H}_B \subseteq \mathcal{H}, \quad m \geq \frac{1}{\epsilon} \log \frac{|\mathcal{H}|}{\delta} \end{aligned}$$

\* Union Bound

$$P(A \cup B) \leq P(A) + P(B)$$

$$P\left(\bigcup_{i \in [m]} A_i\right) \leq \sum_{i \in [m]} P(A_i)$$



$$P(A \cup B) + \underbrace{P(A \cap B)}_{\geq 0} = P(A) + P(B)$$

$$(*) \quad \underset{\text{ERM}}{P_S(h_S)} \leq P_S\left(\bigcup_{h \in \mathcal{H}_B} \{s: h_S(s) = 0\}\right)$$

$$\begin{aligned} \{s: h_S(s) > \epsilon\} & \subseteq \{s: L_{D,f}(h) > \epsilon, h_S(s) = 0\} \\ & \cup_{h \in \mathcal{H}_B} \{s: h_S(s) = 0\} \end{aligned}$$

Define. Sample complexity  $m_{\mathcal{H}}: (0,1)^2 \rightarrow \mathbb{N}$   
 $(\epsilon, \delta) \rightarrow m_{\mathcal{H}}(\epsilon, \delta)$

minimum number of samples needed to guarantee PAC learning.

Corollary. any finite hypothesis class  $\mathcal{H}$  is PAC learnable with sample complexity

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

complexity of  $\mathcal{H}$  is measured by  $|\mathcal{H}|$ .  $\rightarrow$  VC-dimensions

End lecture 2.

Relaxing the realizability assumption

Data generating distribution  $\mathcal{D}$  is over  $\mathcal{X}, \mathcal{Y}$

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

Empirical error:  $L_S(h) \triangleq \frac{|\{i: h(x_i) \neq y_i\}|}{m}$

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

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

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

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

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

Relaxing Binary Classification  $\rightarrow$  Multi-class classification  $\mathcal{Y} = [K]$   
 $h(x, y_i) = \Pi(h(x), y_i)$   
 Regression,  $\mathcal{Y} \subseteq \mathbb{R}$

$$\text{loss } L_h(x, y_i) = (h(x) - y_i)^2$$

$$\text{Risk function } L_{\mathcal{D}}(h) \triangleq \mathbb{E}_{(x,y) \sim \mathcal{D}} [L_h(x, y)]$$

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