

Announcements

Assignments

- HW10 (programming + “written”)
 - Due Thu 4/30, 11:59 pm

Introduction to Machine Learning

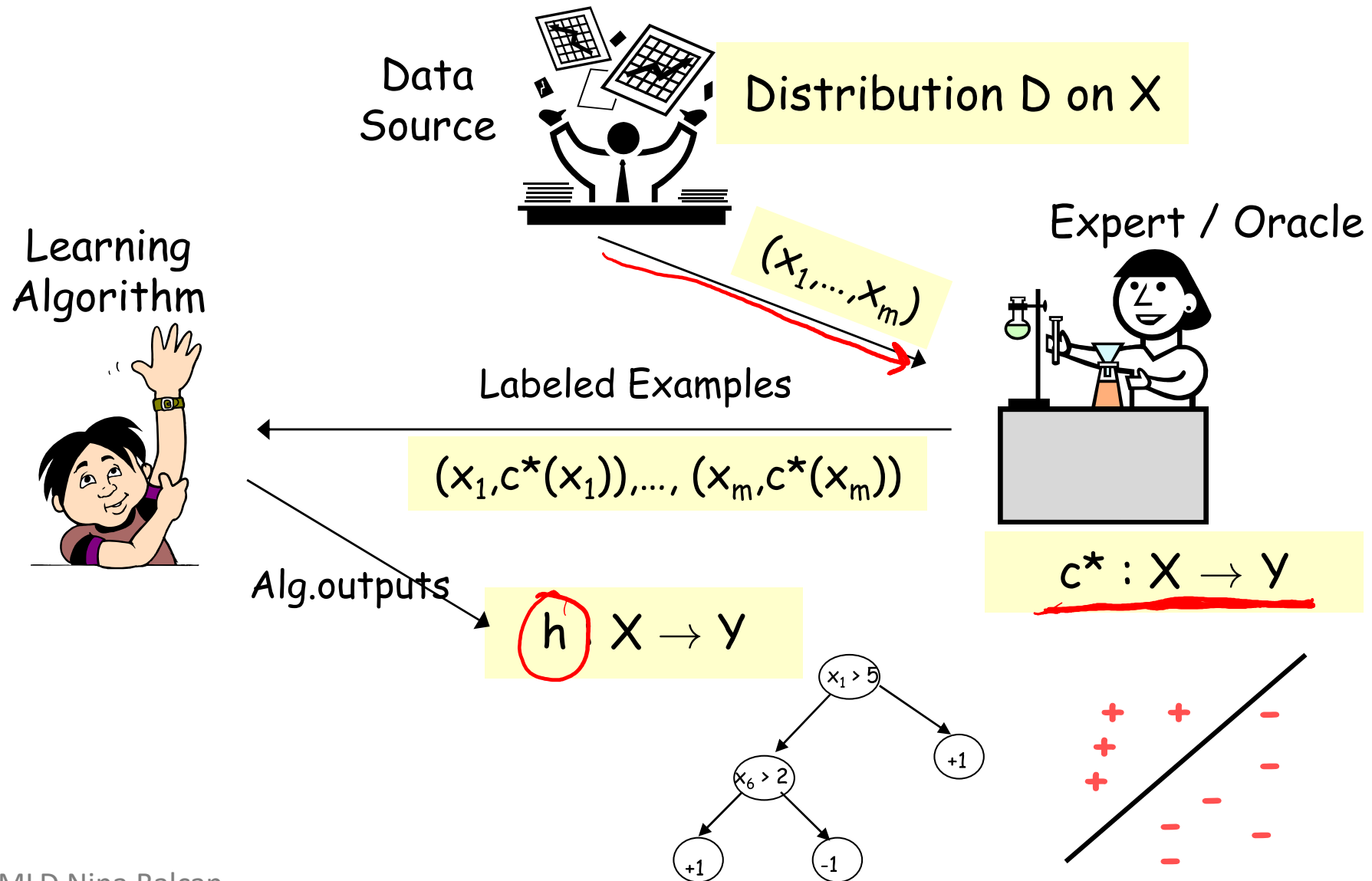
Learning Theory

Instructor: Pat Virtue

Questions For Today

1. Given a classifier with **zero training error**, what can we say about **true error** (aka. generalization error)?
(Sample Complexity, Realizable Case)
2. Given a classifier with **low training error**, what can we say about **true error** (aka. generalization error)?
(Sample Complexity, Agnostic Case)
3. Is there a **theoretical justification for regularization** to avoid overfitting?
(Structural Risk Minimization)

Model for Supervised Learning



Optimal Classification Function

Find the best $h(x) \rightarrow \hat{y}$ by searching in the space of hypothesis functions $h \in \mathcal{H}$.

Optimal classifier:

$$h^*(x) = \operatorname{argmax}_y \underbrace{P(Y = y \mid X = x)}_{\text{train}} \leftarrow$$

But why?

Goal: find a prediction function $h^*: \mathcal{X} \rightarrow \mathcal{Y}$ that minimizes the expected loss for randomly drawn test data (X, Y)

$$h^* = \operatorname{argmin}_h \underbrace{\mathbb{E}_{XY}}_{\text{train}} [L(Y, h(X))] \leftarrow R(h) = \mathbb{E} \text{ Loss}$$

$L(y, \hat{y})$ is the loss or cost of predicting $\hat{y}$ when the true value is y .

Loss Functions

$$h^* = \operatorname{argmin}_h \mathbb{E}_{XY} [\underline{L(Y, h(X))}]$$

Loss function:

$$L: \mathcal{Y} \times \mathcal{Y} \rightarrow \mathbb{R}$$

Classification:

- Two-class, 0,1 loss ←

	$\hat{y}=0$	$\hat{y}=1$
$y=0$	0	1
$y=1$	1	0

- Two-class, arbitrary loss

Regression:

Optimal Classification Function

Expected loss is also called risk:

$$\begin{aligned} R(h) &= \mathbb{E}_{XY}[L(Y, h(X))] \\ h^* &= \operatorname{argmin}_h R(h) \\ &= \operatorname{argmin}_h \mathbb{E}_{XY}[L(Y, h(X))] \end{aligned}$$

For 0,1 loss classification, risk is also error:

$$\begin{aligned} R(h) &= \mathbb{E}_{XY}[L(Y, h(X))] \\ &= \sum_{x,y} p(x,y) L(y, h(x)) \end{aligned}$$

$$L(y, \hat{y}) = \begin{cases} 1 & \text{if } h(x) \neq y \\ 0 & \text{ow} \end{cases} = \mathbb{1}(h(x) \neq y)$$

$$\hat{R}(h) = \sum_{i=1}^N \frac{1}{N} \mathbb{1}(y^{(i)} \neq h(x^{(i)})) = \frac{1}{N} \sum \# \text{error}$$

Two Types of Error

1. True Error (aka. **expected risk**)

$$R(h) = P_{\mathbf{x} \sim p^*(\mathbf{x})}(\underline{c^*(\mathbf{x})} \neq h(\mathbf{x}))$$

This quantity
is always
unknown

2. Train Error (aka. **empirical risk**)

$$\hat{R}(h) = P_{\mathbf{x} \sim \mathcal{S}}(\underline{c^*(\mathbf{x})} \neq h(\mathbf{x}))$$

$$= \frac{1}{N} \sum_{i=1}^N \mathbb{1}(\underline{c^*(\mathbf{x}^{(i)})} \neq h(\mathbf{x}^{(i)}))$$

$$= \frac{1}{N} \sum_{i=1}^N \mathbb{1}(\underline{y^{(i)}} \neq h(\mathbf{x}^{(i)}))$$

We can
measure this
on the training
data

→ where $\mathcal{S} = \{\mathbf{x}^{(1)}, \dots, \mathbf{x}^{(N)}\}_{i=1}^N$ is the training data set, and $\mathbf{x} \sim \mathcal{S}$ denotes that $\mathbf{x}$ is sampled from the empirical distribution.

PAC / SLT Model

1. Generate instances from *unknown* distribution p^*

$$\mathbf{x}^{(i)} \sim p^*(\mathbf{x}), \forall i \quad (1)$$

2. Oracle labels each instance with *unknown* function c^*

$$y^{(i)} = c^*(\mathbf{x}^{(i)}), \forall i \quad (2)$$

3. Learning algorithm chooses hypothesis $h \in \mathcal{H}$ with low(est) training error, $\hat{R}(h)$ *$h^* = \operatorname{argmin} R(h)$*

$$\hat{h} = \operatorname{argmin}_h \hat{R}(h) \quad (3)$$

4. Goal: Choose an h with low generalization error $R(h)$

Three Hypotheses of Interest

The **true function** c^* is the one we are trying to learn and that labeled the training data:

$$y^{(i)} = \underline{c}^*(\mathbf{x}^{(i)}), \forall i \quad (1)$$

The **expected risk minimizer** has lowest true error:



$$\underline{h}^* = \operatorname{argmin}_{h \in \mathcal{H}} R(h)$$

Question:
True or False:
 h^* and c^* are
always equal.

The **empirical risk minimizer** has lowest training error:

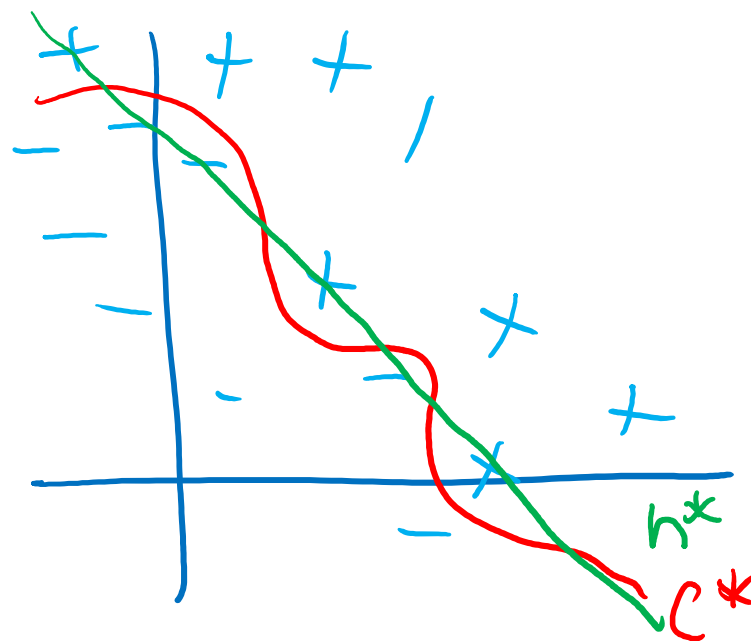
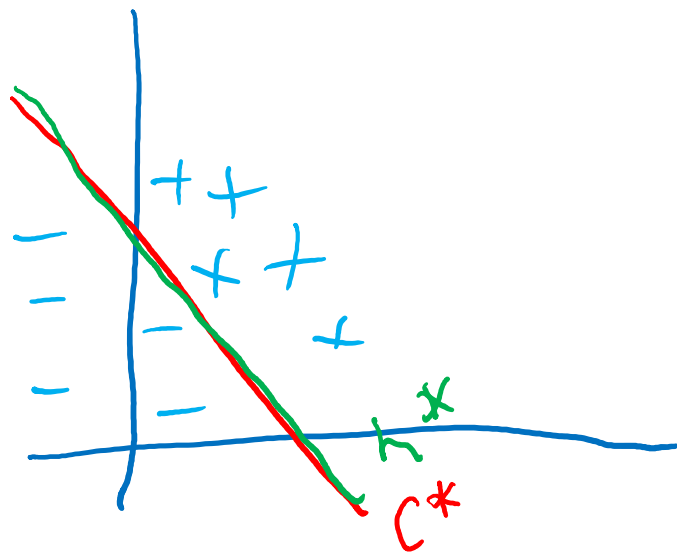
$$\hat{h} = \operatorname{argmin}_{h \in \mathcal{H}} \hat{R}(h) \quad (3)$$

Piazza Poll 1

H linear sep

True or False: h^* and c^* are always equal.

- A. True
- ☒ B. False
- C. Calamity



PAC Learning

unknown

known

Can we bound $R(h)$ in terms of $\hat{R}(h)$?

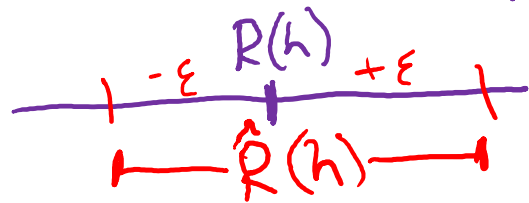
Yes! PAC: Probably Approximately Correct

PAC Learner yield $h \in H$ which is approximately correct: $R(h) \approx 0$
with high probability

$$P(R(h) \approx 0) \approx 1$$

Definition: PAC Criterion:

$$P(\forall h, |R(h) - \hat{R}(h)| \leq \epsilon) \geq 1 - \delta$$



Random? $\hat{R}(h)$ based $x_{\text{train}} \sim P^*(x)$

PAC Learning

Definition: sample complexity is min. num. of training examples N
s.t. the PAC criterion is satisfied for ϵ and δ

$1 - \delta = 99\%$
 $\epsilon = 0.1\%$

Definition: consistent hypothesis

a $h \in H$ is consistent with training data if $\hat{R}(h) = 0$

PAC Learning

The **PAC criterion** is that our learner produces a high accuracy learner with high probability:

$$P(|R(h) - \hat{R}(h)| \leq \epsilon) \geq 1 - \delta \quad (1)$$

Suppose we have a learner that produces a hypothesis $h \in \mathcal{H}$ given a sample of N training examples. The algorithm is called **consistent** if for every ϵ and δ , there exists a positive number of training examples N such that for any distribution p^* , we have that:

$$P(|R(h) - \hat{R}(h)| > \epsilon) < \delta \quad (2)$$

The **sample complexity** is the minimum value of N for which this statement holds. If N is finite for some learning algorithm, then $\mathcal{H}$ is said to be **learnable**. If N is a polynomial function of $\frac{1}{\epsilon}$ and $\frac{1}{\delta}$ for some learning algorithm, then $\mathcal{H}$ is said to be **PAC learnable**.

PAC Learning

Four types of problems

Two cases c^*

- Realizable: $c^* \in H$
- Agnostic: $c^* \notin H$ or $c^* \in H$

Two cases $|H|$

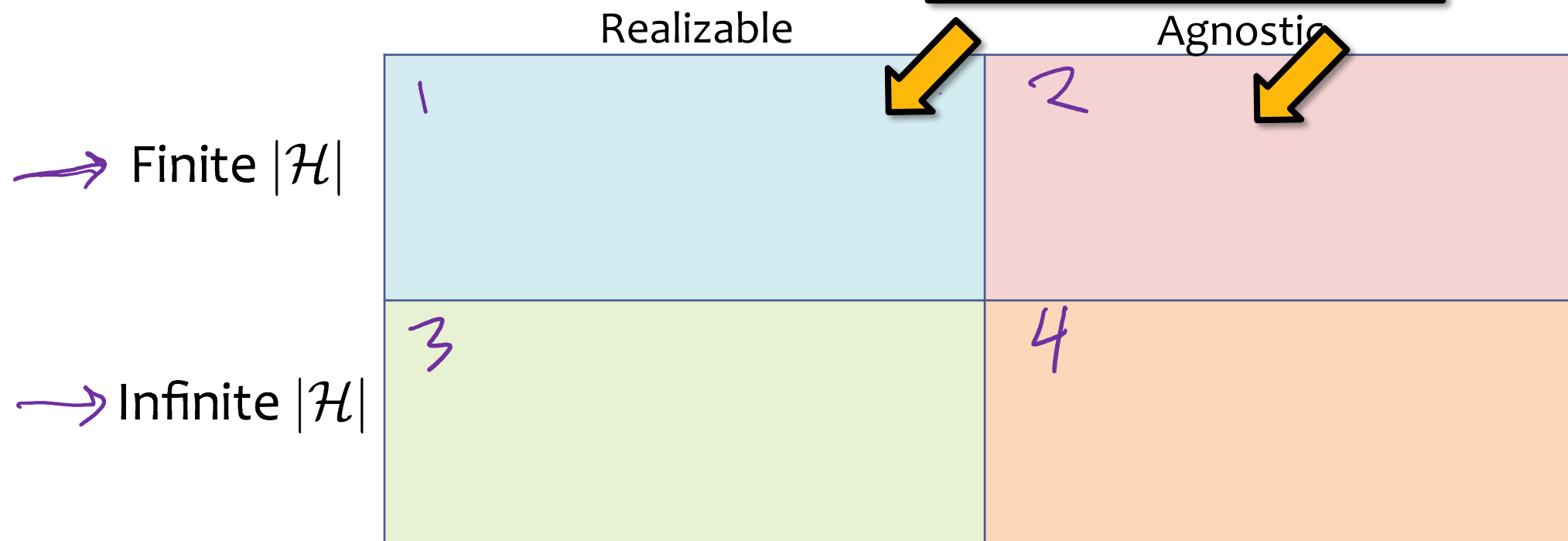
- Finite $|H| < \infty$
- Infinite $|H| = \infty$

Sample Complexity Results

Definition 0.1. The **sample complexity** of a learning algorithm is the number of examples required to achieve arbitrarily small error (with respect to the optimal hypothesis) with high probability (i.e. close to 1).

Four Cases we care about...

We'll start with the finite case...



PAC Learning

Theorem 1: Sample Complexity (Realizable, Finite $|\mathcal{H}|$)

 $N \geq \frac{1}{\epsilon} (\log |\mathcal{H}| + \log \frac{1}{\delta})$ labelled examples
are sufficient to ensure that with prob $1-\delta$
all $h \in \mathcal{H}$ with $\hat{R}(h)=0$ have $R(h) \leq \epsilon$

PAC Learning

Proof of Theorem 1

{To the whiteboard...}

Sample Complexity Results

Definition 0.1. The **sample complexity** of a learning algorithm is the number of examples required to achieve arbitrarily small error (with respect to the optimal hypothesis) with high probability (i.e. close to 1).

Four Cases we care about...

	Realizable	Agnostic
Finite $ \mathcal{H} $	Thm. 1 $N \geq \frac{1}{\epsilon} [\log(\mathcal{H}) + \log(\frac{1}{\delta})]$ labeled examples are sufficient so that with probability $(1 - \delta)$ all $h \in \mathcal{H}$ with $\hat{R}(h) = 0$ have $R(h) \leq \epsilon$.	
Infinite $ \mathcal{H} $		

Piazza Poll 2

Question:

Suppose H = class of conjunctions over $\mathbf{x}$ in $\{0,1\}^M$

Example hypotheses:

$$\rightarrow h(\mathbf{x}) = x_1 (1-x_3) x_5 \quad x_1 \wedge \neg x_3 \wedge x_5$$

$$\rightarrow h(\mathbf{x}) = x_1 (1-x_2) x_4 (1-x_5)$$

If $M = 10$, $\epsilon = 0.1$, $\delta = 0.01$, how many examples suffice according to Theorem 1?

Answer:

A. $10 * (2 * \ln(10) + \ln(100)) \approx 92$

B. $10 * (3 * \ln(10) + \ln(100)) \approx 116$

C. $10 * (10 * \ln(2) + \ln(100)) \approx 116$

D. $10 * (10 * \ln(3) + \ln(100)) \approx 156$

E. ~~$100 * (2 * \ln(10) + \ln(10)) \approx 691$~~

F. ~~$100 * (3 * \ln(10) + \ln(10)) \approx 922$~~

G. ~~$100 * (10 * \ln(2) + \ln(10)) \approx 924$~~

H. ~~$100 * (10 * \ln(3) + \ln(10)) \approx 1329$~~

$$|\mathcal{H}| = 3^{10}$$

$\rightarrow$ **Thm. 1** $N \geq \frac{1}{\epsilon} [\log(|\mathcal{H}|) + \log(\frac{1}{\delta})]$ labeled examples are sufficient so that with probability $(1 - \delta)$ all $h \in \mathcal{H}$ with $\hat{R}(h) = 0$ have $R(h) \leq \epsilon$.

Sample Complexity Results

Definition 0.1. The **sample complexity** of a learning algorithm is the number of examples required to achieve arbitrarily small error (with respect to the optimal hypothesis) with high probability (i.e. close to 1).

Four Cases we care about...

	Realizable	Agnostic
Finite $ \mathcal{H} $	Thm. 1 $N \geq \frac{1}{\epsilon} [\log(\mathcal{H}) + \log(\frac{1}{\delta})]$ labeled examples are sufficient so that with probability $(1 - \delta)$ all $h \in \mathcal{H}$ with <u>$\hat{R}(h) = 0$</u> have $R(h) \leq \epsilon$.	Thm. 2 $N \geq \frac{1}{2\epsilon^2} [\log(\mathcal{H}) + \log(\frac{2}{\delta})]$ labeled examples are sufficient so that with probability $(1 - \delta)$ for all $h \in \mathcal{H}$ we have that <u>$R(h) - \hat{R}(h) \leq \epsilon$</u> .
Infinite $ \mathcal{H} $		

1. Bound is **inversely linear in epsilon** (e.g. halving the error requires double the examples)
2. Bound is **only logarithmic in $|\mathcal{H}|$** (e.g. quadrupling the hypothesis space only requires double the examples)

1. Bound is **inversely quadratic in epsilon** (e.g. halving the error requires 4x the examples)
2. Bound is **only logarithmic in $|\mathcal{H}|$** (i.e. same as Realizable case)



Realizable



Agnostic

Thm. 1 $N \geq \frac{1}{\epsilon} [\log(|\mathcal{H}|) + \log(\frac{1}{\delta})]$ labeled examples are sufficient so that with probability $(1 - \delta)$ all $h \in \mathcal{H}$ with $\hat{R}(h) = 0$ have $R(h) \leq \epsilon$.

Thm. 2 $N \geq \frac{1}{2\epsilon^2} [\log(|\mathcal{H}|) + \log(\frac{2}{\delta})]$ labeled examples are sufficient so that with probability $(1 - \delta)$ for all $h \in \mathcal{H}$ we have that $|R(h) - \hat{R}(h)| \leq \epsilon$.

Finite $|\mathcal{H}|$

Infinite $|\mathcal{H}|$

Using a PAC bound

$$|H|e^{-\underline{m}\epsilon} \leq \delta$$

- Given ϵ and δ , yields sample complexity

#training data, $m \geq \frac{\ln |H| + \ln \frac{1}{\delta}}{\epsilon}$



- Given m and δ , yields error bound

error, $\epsilon \geq \frac{\ln |H| + \ln \frac{1}{\delta}}{m}$

Summary of PAC bounds for finite model classes

With probability $\geq 1-\delta$,

1) For all $h \in H$ s.t. $\text{error}_{\text{train}}(h) = 0$,

$$\text{error}_{\text{true}}(h) \leq \underline{\varepsilon} = \frac{\ln |H| + \ln \frac{1}{\delta}}{m}$$

Haussler's bound



2) For all $h \in H$

$$\underbrace{|\text{error}_{\text{true}}(h) - \text{error}_{\text{train}}(h)|}_{\text{red bracket}} \leq \varepsilon = \sqrt{\frac{\ln |H| + \ln \frac{2}{\delta}}{2m}}$$

Hoeffding's bound

PAC bound and Bias-Variance tradeoff

$$P(|\text{error}_{\text{true}}(h) - \text{error}_{\text{train}}(h)| \geq \epsilon) \leq 2|H|e^{-2m\epsilon^2} \leq \delta$$

- Equivalently, with probability $\geq 1 - \delta$

$$\overset{R(h)}{\text{error}_{\text{true}}(h)} \leq \overset{\hat{R}(h)}{\text{error}_{\text{train}}(h)} + \sqrt{\frac{\ln |H| + \ln \frac{2}{\delta}}{2m}}$$

- Fixed $|H|$

~~Model class~~ Train Size

m small

m large

small

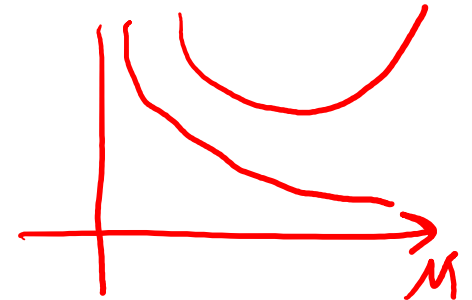
large

large

small

bias

variance



PAC bound and Bias-Variance tradeoff

$$P(|\text{error}_{\text{true}}(h) - \text{error}_{\text{train}}(h)| \geq \epsilon) \leq 2|H|e^{-2m\epsilon^2} \leq \delta$$

- Equivalently, with probability $\geq 1 - \delta$

$$\text{error}_{\text{true}}(h) \leq \text{error}_{\text{train}}(h) + \sqrt{\frac{\ln |H| + \ln \frac{2}{\delta}}{2m}}$$

- Fixed m

Model class

|H| large (complex)

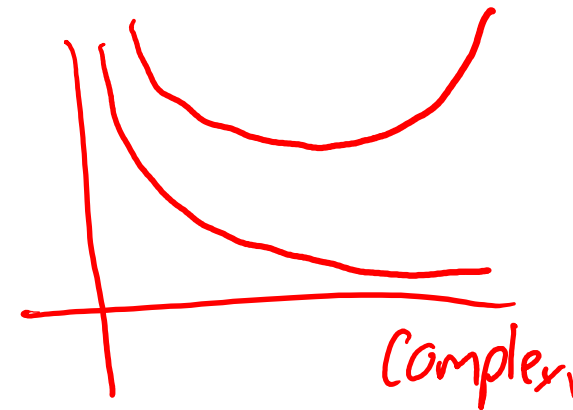
|H| small (simple)

small

large

large

small



Number of decision trees of depth k

Recursive solution:

$$m \geq \frac{1}{2\epsilon^2} \left(\ln |H| + \ln \frac{2}{\delta} \right)$$

Given n **binary** attributes

H_k = Number of **binary** decision trees of depth k

$$H_0 = 2$$

H_k = (#choices of root attribute)

* (# possible left subtrees)

* (# possible right subtrees) = $n * H_{k-1} * H_{k-1}$

Write $L_k = \log_2 H_k$

$$L_0 = 1$$

$$L_k = \log_2 n + 2L_{k-1} = \log_2 n + 2(\log_2 n + 2L_{k-2})$$

$$= \log_2 n + 2\log_2 n + 2^2\log_2 n + \dots + 2^{k-1}(\log_2 n + 2L_0)$$

$$\text{So } L_k = (2^k - 1)(1 + \log_2 n) + 1$$

PAC bound for decision trees of depth k

$$m \geq \frac{\ln 2}{2\epsilon^2} \left((2^k - 1)(1 + \log_2 n) + 1 + \log_2 \frac{2}{\delta} \right)$$

- Bad!!!
 - Number of points is exponential in depth k !
- But, for m data points, decision tree can't get too big...

Number of leaves never more than number data points

Number of decision trees with k leaves

$$m \geq \frac{1}{2\epsilon^2} \left(\ln |H| + \ln \frac{2}{\delta} \right)$$

H_k = Number of binary decision trees with k leaves

$$H_1 = 2$$

$$H_k = (\text{\#choices of root attribute}) *$$

$$\begin{aligned} & [(\text{\# left subtrees with 1 leaf}) * (\text{\# right subtrees with k-1 leaves}) \\ & + (\text{\# left subtrees with 2 leaves}) * (\text{\# right subtrees with k-2 leaves}) \\ & + \dots \\ & + (\text{\# left subtrees with k-1 leaves}) * (\text{\# right subtrees with 1 leaf})] \end{aligned}$$

$$H_k = n \sum_{i=1}^{k-1} H_i H_{k-i} = n^{k-1} C_{k-1} \quad (C_{k-1} : \text{Catalan Number})$$

Loose bound (using Sterling's approximation):

$$H_k \leq n^{k-1} 2^{2k-1}$$

Number of decision trees

- With k leaves $m \geq \frac{1}{2\epsilon^2} \left(\ln |H| + \ln \frac{2}{\delta} \right)$

$$\log_2 H_k \leq \underline{(k-1)} \log_2 n + 2k - 1 \quad \text{linear in } k$$

number of points m is linear in #leaves

- With depth k

$$\log_2 H_k = (2^k - 1)(1 + \log_2 n) + 1 \quad \text{exponential in } k$$

number of points m is exponential in depth

PAC bound for decision trees with k leaves – Bias-Variance revisited

With prob $\geq 1-\delta$ $\text{error}_{\text{true}}(h) \leq \text{error}_{\text{train}}(h) + \sqrt{\frac{\ln |H| + \ln \frac{2}{\delta}}{2m}}$

With $H_k \leq n^{k-1} 2^{2k-1}$, we get

$$\text{error}_{\text{true}}(h) \leq \text{error}_{\text{train}}(h) + \sqrt{\frac{(k-1) \ln n + (2k-1) \ln 2}{2m} + \ln \frac{2}{\delta}}$$

	$\downarrow$	$\downarrow$
$k = m$	0	large ($\sim > \frac{1}{2}$)
$k < m$	> 0	small ($\sim < \frac{1}{2}$)

What did we learn from decision trees?

- Moral of the story:

Complexity of learning not measured in terms of size of model space, but in maximum *number of points* that allows consistent classification

Sample Complexity Results

Definition 0.1. The **sample complexity** of a learning algorithm is the number of examples required to achieve arbitrarily small error (with respect to the optimal hypothesis) with high probability (i.e. close to 1).

Four Cases we care about...

	Realizable	Agnostic
Finite $ \mathcal{H} $	Thm. 1 $N \geq \frac{1}{\epsilon} [\log(\mathcal{H}) + \log(\frac{1}{\delta})]$ labeled examples are sufficient so that with probability $(1 - \delta)$ all $h \in \mathcal{H}$ with $\hat{R}(h) = 0$ have $R(h) \leq \epsilon$.	Thm. 2 $N \geq \frac{1}{2\epsilon^2} [\log(\mathcal{H}) + \log(\frac{2}{\delta})]$ labeled examples are sufficient so that with probability $(1 - \delta)$ for all $h \in \mathcal{H}$ we have that $ R(h) - \hat{R}(h) \leq \epsilon$.
Infinite $ \mathcal{H} $		