

10-301/601: Introduction to Machine Learning

Lecture 15 – Learning Theory

Geoff Gordon

with thanks to Matt Gormley & Henry Chai

Statistical learning theory setup

1. Data points are generated i.i.d. from some *unknown* distribution
 $\mathbf{x}^{(n)} \sim p^*(\mathbf{x})$
2. Labels are generated from some *unknown* function
 $y^{(n)} = c^*(\mathbf{x}^{(n)}) \in \{-1, +1\}$ note: **binary** classification
3. The learning algorithm chooses the hypothesis (classifier) with lowest *training* error rate from a specified hypothesis set, $\mathcal{H}$
4. Goal: return a hypothesis (or classifier) with low *true* error rate (measure on *test* set)

Types of Error

- True error rate
 - Actual quantity of interest for learning
 - How well your hypothesis will perform on new samples
- Training error rate
 - Used to choose $h \in \mathcal{H}$ (e.g., fit model parameters)
 - May be a very optimistic estimate of true error

Types of Error

- True error rate
 - Actual quantity of interest for learning
 - How well your hypothesis will perform on new samples
- Training error rate
 - Used to choose $h \in \mathcal{H}$ (e.g., fit model parameters)
 - May be a very optimistic estimate of true error
- Test error rate
 - Used to evaluate hypothesis performance
 - Good estimate of true error (w/ enough test data)
- Validation error rate
 - Used to help choose $\mathcal{H}$ (e.g., set hyperparameters)
 - Somewhat optimistic estimate of true error

Error rate is also
called risk

- True error rate = (true) *risk* — unknown

$$R(h) = \mathbb{E}$$

- Training error rate = *empirical risk* — we can measure this

$$\hat{R}(h) = \mathbb{E}$$

Three classifiers

1. The *true classifier*, c^* : best answer but may be unachievable
2. The (*true*) *risk minimizer* (best achievable answer):

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

3. The *empirical risk minimizer* (the only one of the three that we can actually know)

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

Three classifiers

1. The *true classifier*, c^* : best answer but may be unachievable
2. The *(true) risk minimizer* (best achievable answer):

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

3. The *empirical risk minimizer* (the only one of the three that we can actually know)

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

ERM = the learning method that picks a hypothesis by minimizing empirical risk

Overfitting

- Recall: **overfitting** = difference between true error rate and training error rate =
- Goal for today: **predict** and **control** overfitting for ERM
 - by finding (and proving) conditions that keep $\hat{R}(h)$ close to $R(h)$

Bound on overfitting

- PAC = Probably Approximately Correct criterion

$$P\left(\left|R(h) - \hat{R}(h)\right| \leq \epsilon\right) \geq 1 - \delta \quad \forall h \in \mathcal{H}$$

for some ϵ (difference between true and empirical risk)
and δ (probability of “failure”)

why the name?

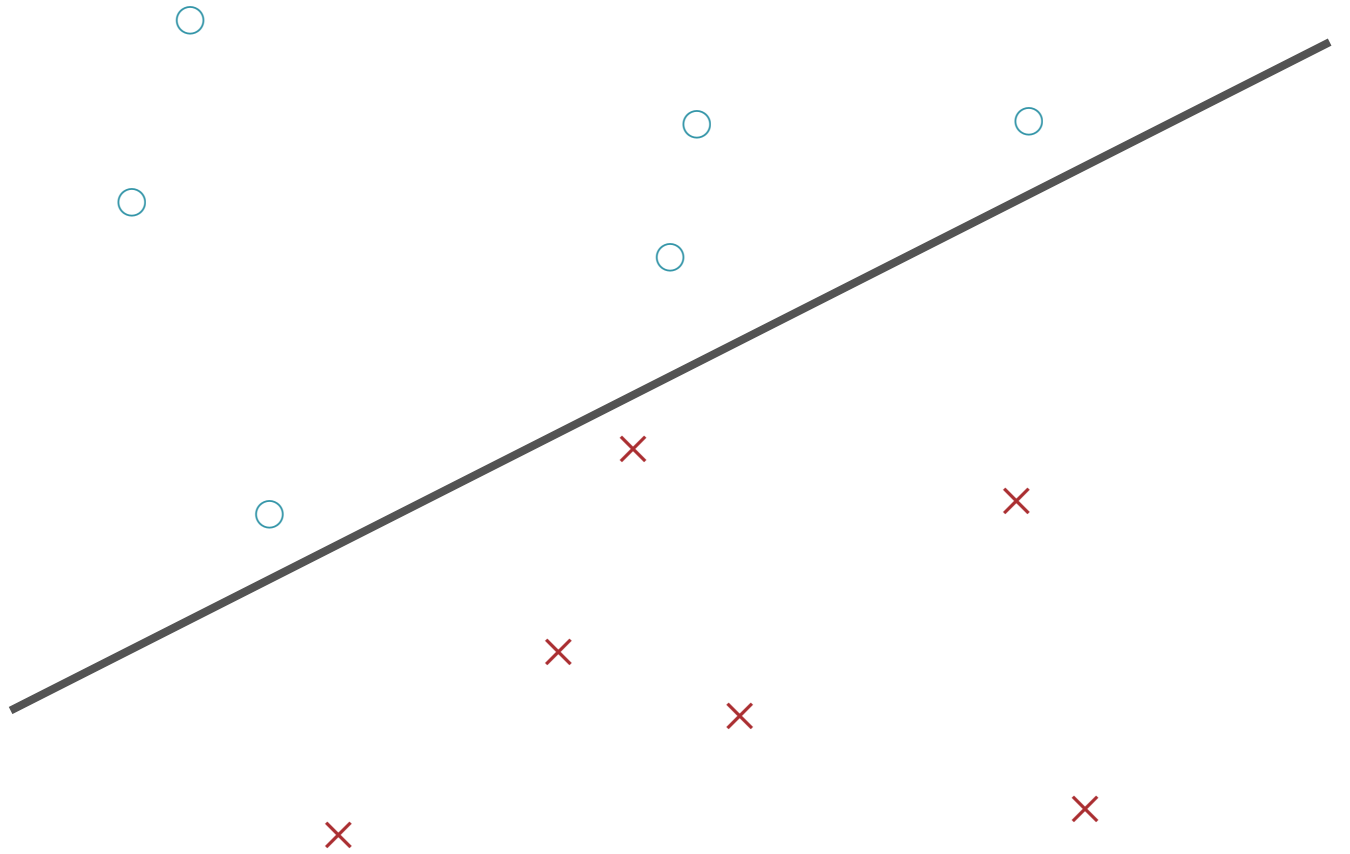
Sample Complexity

- We will do ERM on some $\mathcal{H}$ with M training examples
- We want to satisfy the PAC criterion with *small* ϵ and δ
- Chief levers: $\mathcal{H}$ and M
- Sample complexity (of ERM on $\mathcal{H}$) = the M we need in order to satisfy the PAC criterion for a given ϵ and δ

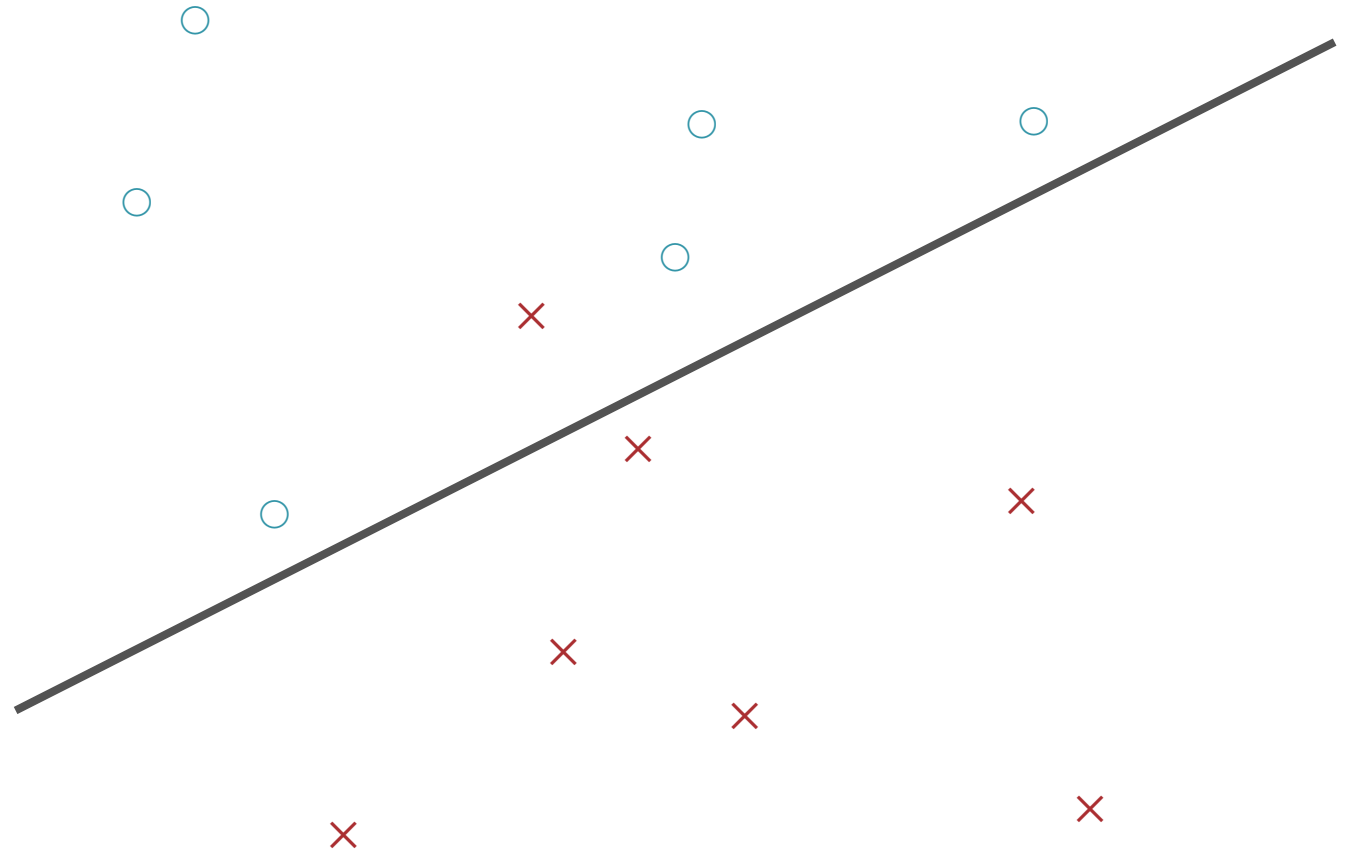
Four cases

- Realizable vs. Agnostic
 - Realizable $\rightarrow c^* \in \mathcal{H}$
 - Agnostic $\rightarrow c^*$ might or might not be in $\mathcal{H}$
- Finite vs. Infinite
 - Finite $\rightarrow |\mathcal{H}| < \infty$
 - Infinite $\rightarrow |\mathcal{H}| = \infty$

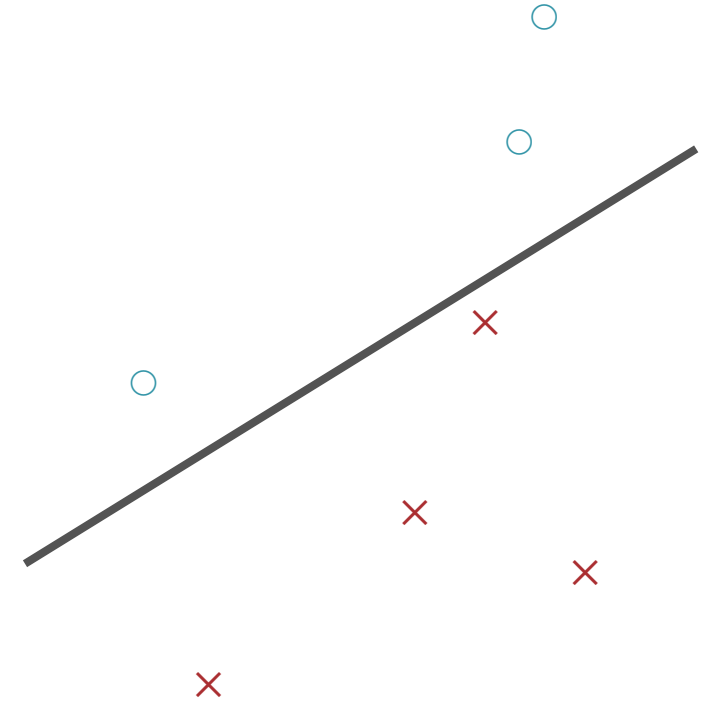
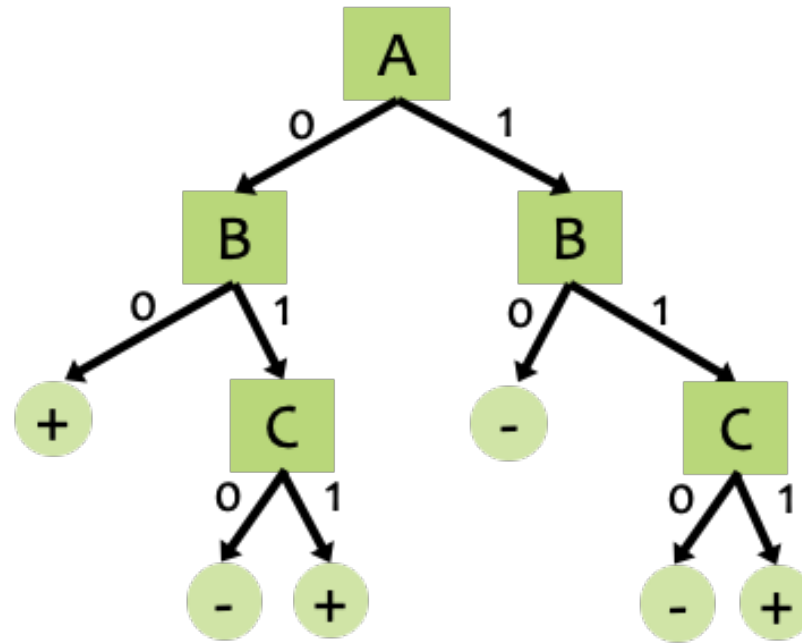
Realizable vs. agnostic



Realizable vs. agnostic



Finite vs.
infinite $|\mathcal{H}|$



- Decision trees of bounded depth on discrete attributes vs. linear separators in 2D

Theorem 1: Finite, Realizable Case

For a finite hypothesis set $\mathcal{H}$ s.t. $c^* \in \mathcal{H}$ and arbitrary distribution p^* , if the number of labelled training data points satisfies

$$M \geq \frac{1}{\epsilon} \left(\ln(|\mathcal{H}|) + \ln\left(\frac{1}{\delta}\right) \right)$$

then with probability at least $1 - \delta$, all $h \in \mathcal{H}$ with $\hat{R}(h) = 0$ have $R(h) \leq \epsilon$

We will prove this over the next few slides

Theorem 1: Finite, Realizable Case

Bound is
linear in $\frac{1}{\epsilon}$

hypothesis set $\mathcal{H}$ s.t. $c^* \in \mathcal{H}$ and arbitrary
of the number of labelled training data
points satisfies

$$M \geq \frac{1}{\epsilon} \left(\ln(|\mathcal{H}|) + \ln\left(\frac{1}{\delta}\right) \right)$$

then with probability at least $1 - \delta$, all $h \in \mathcal{H}$ with
 $\hat{R}(h) = 0$ have $R(h) \leq \epsilon$

We will prove this over the next few slides

Theorem 1: Finite, Realizable Case

Bound is
linear in $\frac{1}{\epsilon}$

but only logarithmic
in $|\mathcal{H}|$ and $\frac{1}{\delta}$

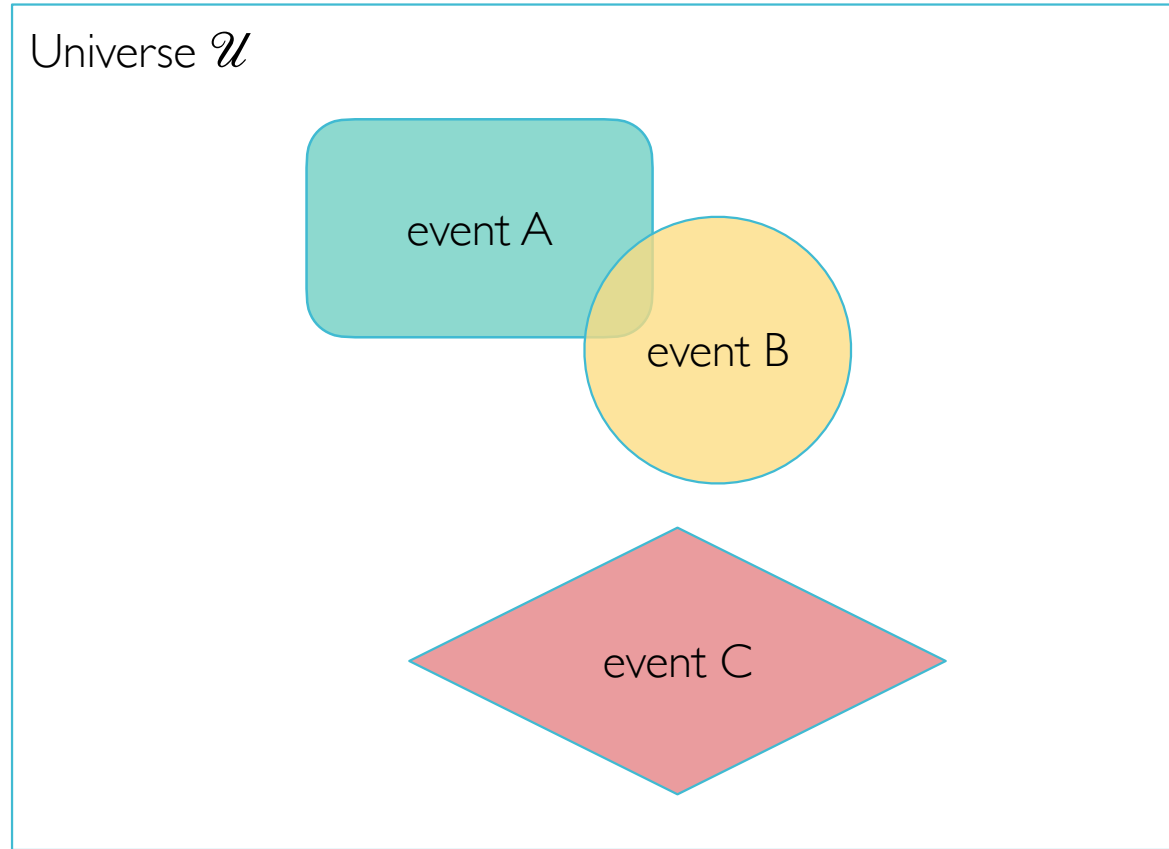
points satisfies

$$M \geq \frac{1}{\epsilon} \left(\ln(|\mathcal{H}|) + \ln\left(\frac{1}{\delta}\right) \right)$$

then with probability at least $1 - \delta$, all $h \in \mathcal{H}$ with $\hat{R}(h) = 0$ have $R(h) \leq \epsilon$

We will prove this over the next few slides

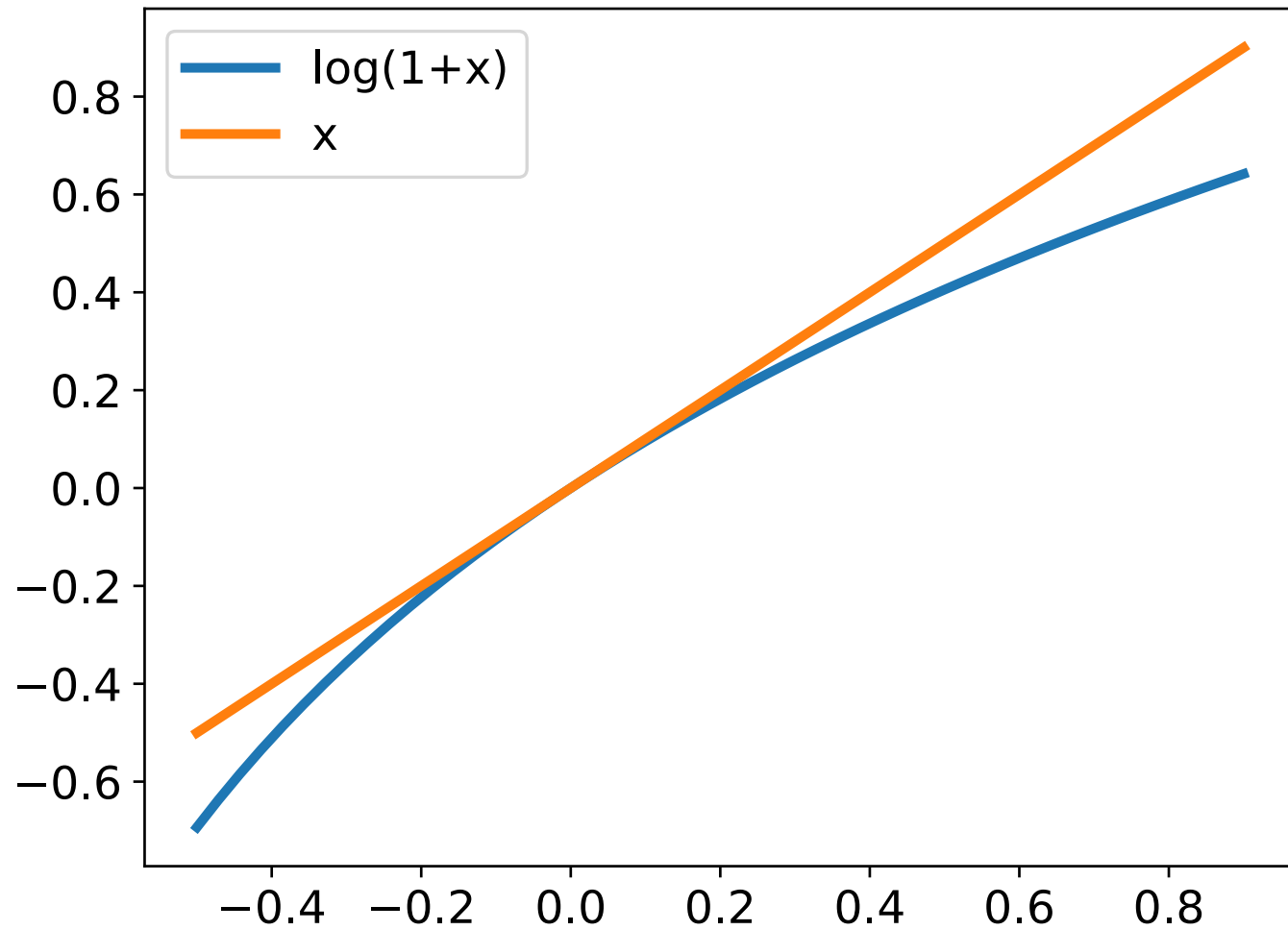
Union bound



$$P(A \text{ or } B \text{ or } C) \leq P(A) + P(B) + P(C)$$

$P(\text{some event happens}) \leq \text{sum of probabilities}$

Bound on log



$$\ln(1+x) \leq x$$

Notation for conclusion

- Theorem said: if

$$M \geq \frac{1}{\epsilon} \left(\ln(|\mathcal{H}|) + \ln\left(\frac{1}{\delta}\right) \right)$$

then with probability at least $1 - \delta$, all $h \in \mathcal{H}$ with $\hat{R}(h) = 0$ have $R(h) \leq \epsilon$

- Write E for event: $\exists h \in \mathcal{H}$ with $\hat{R}(h) = 0, R(h) > \epsilon$
- Theorem's conclusion is $P(E) < \delta$

Bound for one hypothesis

- Consider some h with $R(h) > \epsilon$. What's $P(\hat{R}(h) = 0)$?

Bound for k hypotheses

- Suppose there are k hypotheses with $R(h) > \epsilon$. What's $P(E)$, i.e., $P(\hat{R}(h_1) = 0 \text{ or } \hat{R}(h_2) = 0 \text{ or } \dots)$?

Theorem will be true if $P(E) \leq \delta$

Solve for M

we have

$$P(E) < |\mathcal{H}| (1 - \epsilon)^M$$

Theorem 1: corollary

- For a finite hypothesis set $\mathcal{H}$ s.t. $c^* \in \mathcal{H}$ and arbitrary distribution p^* , given a training data set S s.t. $|S| = M$, all $h \in \mathcal{H}$ with $\hat{R}(h) = 0$ have

$$R(h) \leq \frac{1}{M} \left(\ln(|\mathcal{H}|) + \ln\left(\frac{1}{\delta}\right) \right)$$

with probability at least $1 - \delta$.

Recall Theorem 1 said $M \geq \frac{1}{\epsilon} \left(\ln(|\mathcal{H}|) + \ln\left(\frac{1}{\delta}\right) \right)$

Theorem 2: finite, agnostic case

- For a finite hypothesis set $\mathcal{H}$ and arbitrary distribution p^* , if the number of labelled training data points satisfies

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

then with probability at least $1 - \delta$, all $h \in \mathcal{H}$ satisfy

$$\left| R(h) - \hat{R}(h) \right| \leq \epsilon$$

- Bound is inversely **quadratic** in ϵ , e.g., halving ϵ means we need four times as many labelled training data points

Theorem 2: finite, agnostic case

- For a finite hypothesis set $\mathcal{H}$ and arbitrary distribution p^* , if the number of labelled training data points satisfies

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

then with probability at least $1 - \delta$, all $h \in \mathcal{H}$ satisfy

$$\left| R(h) - \hat{R}(h) \right| \leq \epsilon$$

- Bound is inversely **quadratic** in ϵ , e.g., halving ϵ means we need four times as many labelled training data points
- Again, making the bound tight and solving for ϵ gives...

Theorem 2: corollary

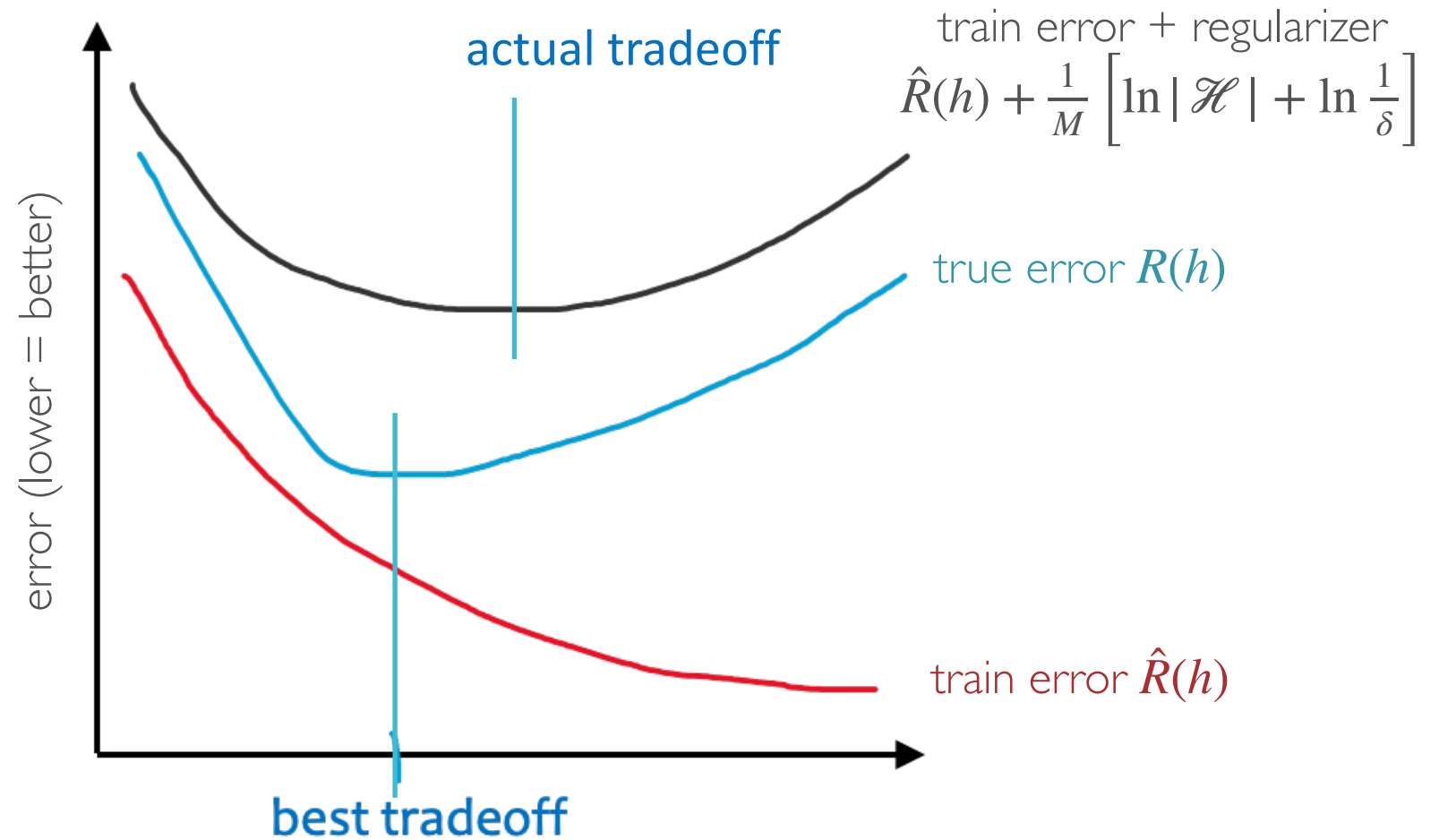
- For a finite hypothesis set $\mathcal{H}$ and arbitrary distribution p^* , given a training data set S s.t. $|S| = M$, all $h \in \mathcal{H}$ have

$$R(h) \leq \hat{R}(h) + \sqrt{\frac{1}{2M} \left(\ln(|\mathcal{H}|) + \ln\left(\frac{2}{\delta}\right) \right)}$$

with probability at least $1 - \delta$.

Learning theory & model selection

Key point: we want to trade
off low training error vs.
keeping $\mathcal{H}$ simple



Ex: $\mathcal{H}$ = conjunctions on d binary attributes: $\ln |\mathcal{H}| = d \ln 3$
Expert sorts attributes, most likely to be relevant first
We allow conjunctions on first d of them:
training error $\downarrow$ as d increases
regularizer $\uparrow$ as d increases
stop when PAC bound is smallest (best tradeoff)

What happens
when
 $|\mathcal{H}| = \infty$?

- Bounds:

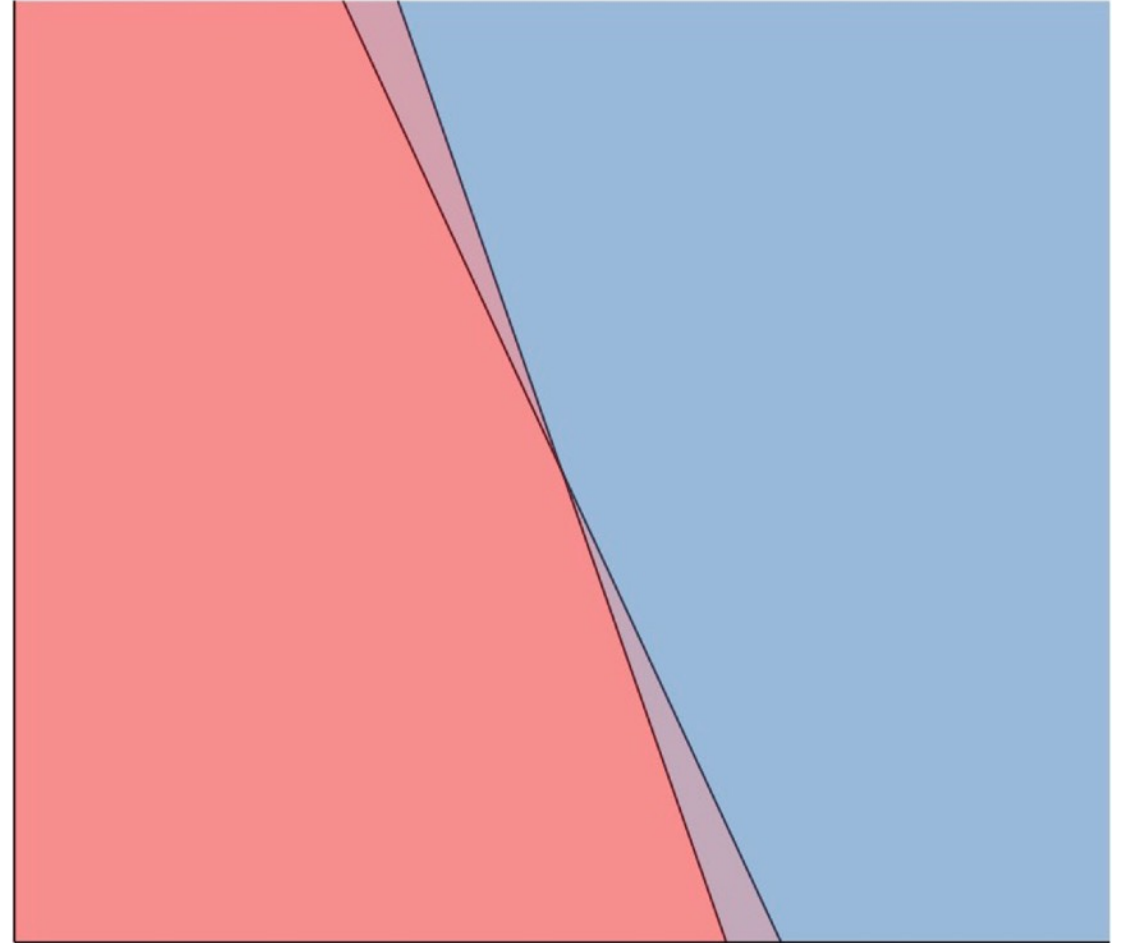
$$R(h) \leq \frac{1}{M} \left(\ln(|\mathcal{H}|) + \ln\left(\frac{1}{\delta}\right) \right)$$

$$R(h) \leq \hat{R}(h) + \sqrt{\frac{1}{2M} \left(\ln(|\mathcal{H}|) + \ln\left(\frac{2}{\delta}\right) \right)}$$

with probability at least $1 - \delta$.

Intuition

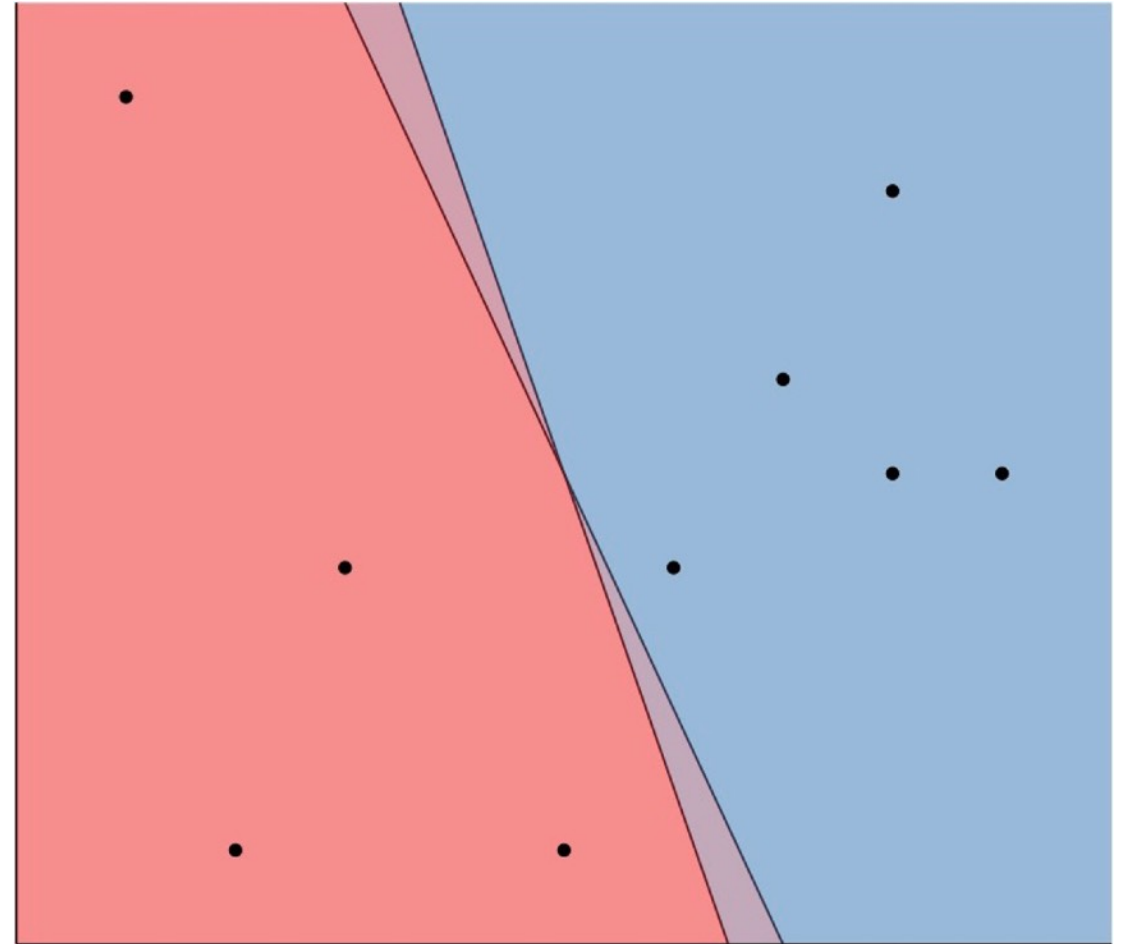
For “nice” infinite hypothesis sets $\mathcal{H}$, many hypotheses in $\mathcal{H}$ will behave similarly



Intuition

For “nice” infinite hypothesis sets $\mathcal{H}$, many hypotheses in $\mathcal{H}$ will behave similarly

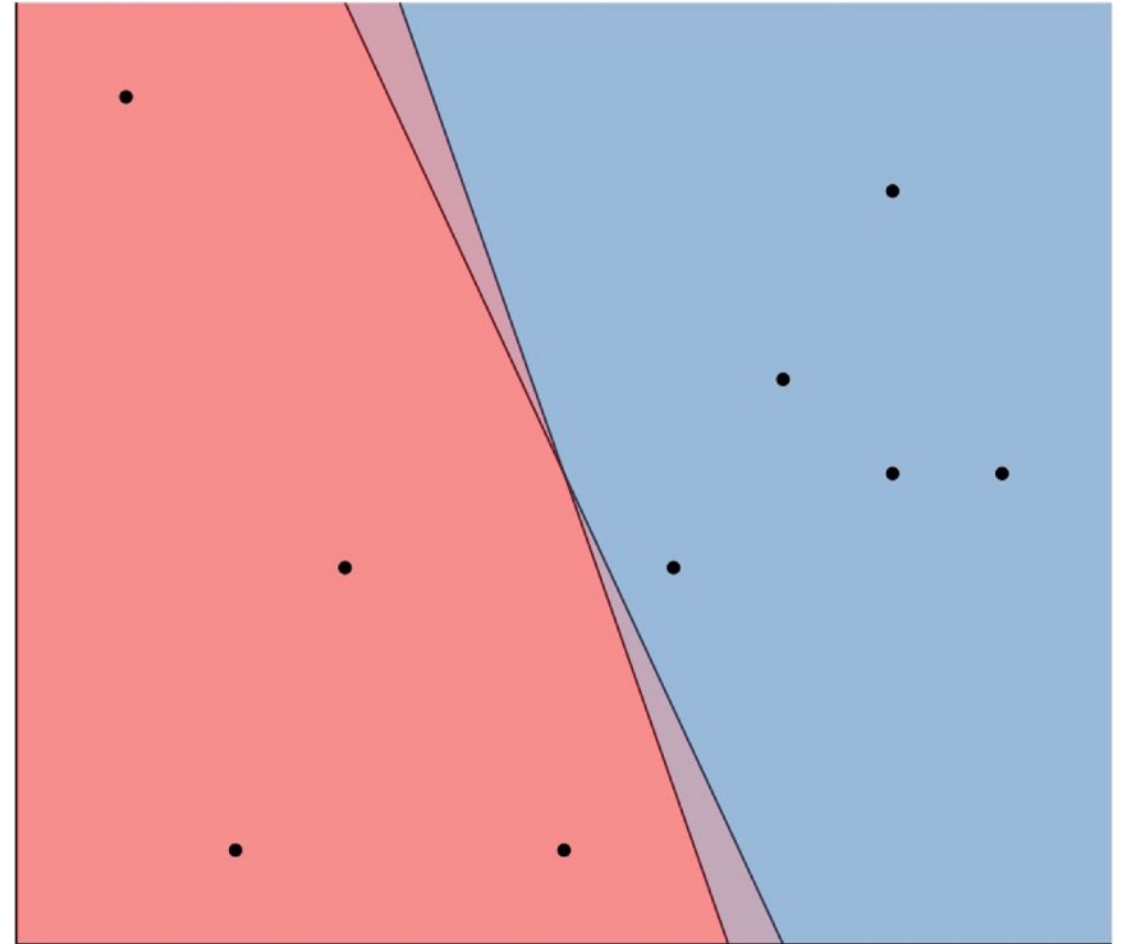
Relative to this dataset, these two hypotheses are *identical*!



Intuition

For “nice” infinite hypothesis sets $\mathcal{H}$, many hypotheses in $\mathcal{H}$ will behave similarly

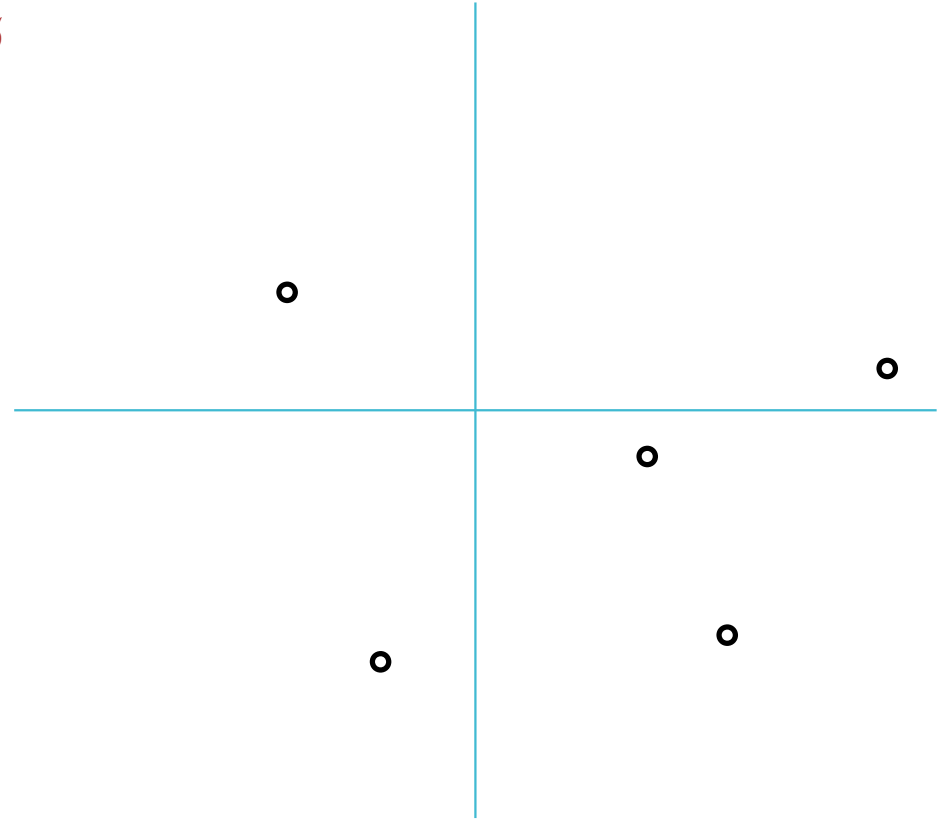
Relative to this dataset, these two hypotheses are *identical*!



Idea: instead of using full size of $\mathcal{H}$, count how many **actually distinct** hypotheses there are

How many
distinct $h \in \mathcal{H}$
can there be?

$$M = 5$$



- What's the largest possible number of distinct hypotheses on M points?

What does our
bound tell us if
 $|\mathcal{H}| = 2^m$?

$$R(h) \leq \frac{1}{M} \left(\ln(|\mathcal{H}|) + \ln\left(\frac{1}{\delta}\right) \right)$$

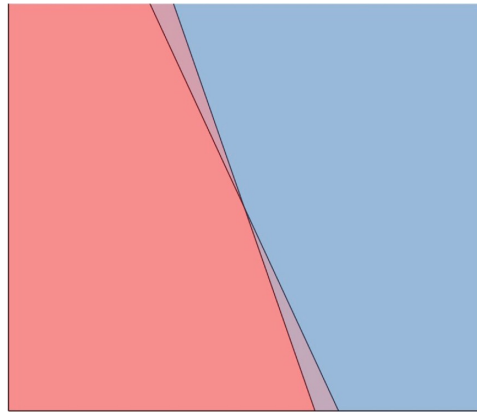
What does our
bound tell us if
 $|\mathcal{H}| = 2^m$?

$$R(h) \leq \frac{1}{M} \left(\ln(|\mathcal{H}|) + \ln\left(\frac{1}{\delta}\right) \right)$$

Vacuous! Need a tighter count of $|\mathcal{H}|$

Not surprising since $|\mathcal{H}| = 2^m$ is a kind of memorization learner

Counting h



$\mathcal{H}$ = halfspaces

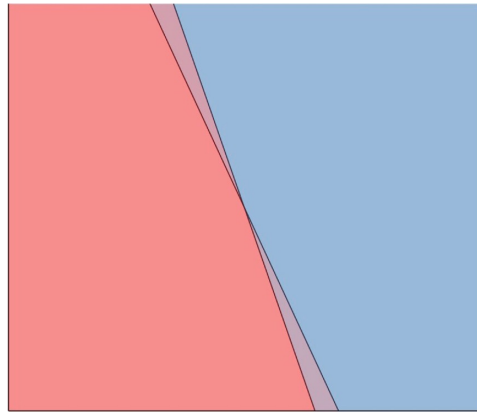
$$M = 1$$

$$\mathcal{D} =$$



- Fix $\mathcal{H}$
- Consider datasets $\mathcal{D}$ of size $M = 1, 2, \dots$
- For each dataset, count how many **actually distinct** hypotheses $h \in \mathcal{H}$ there are: $|\mathcal{H}(\mathcal{D})|$

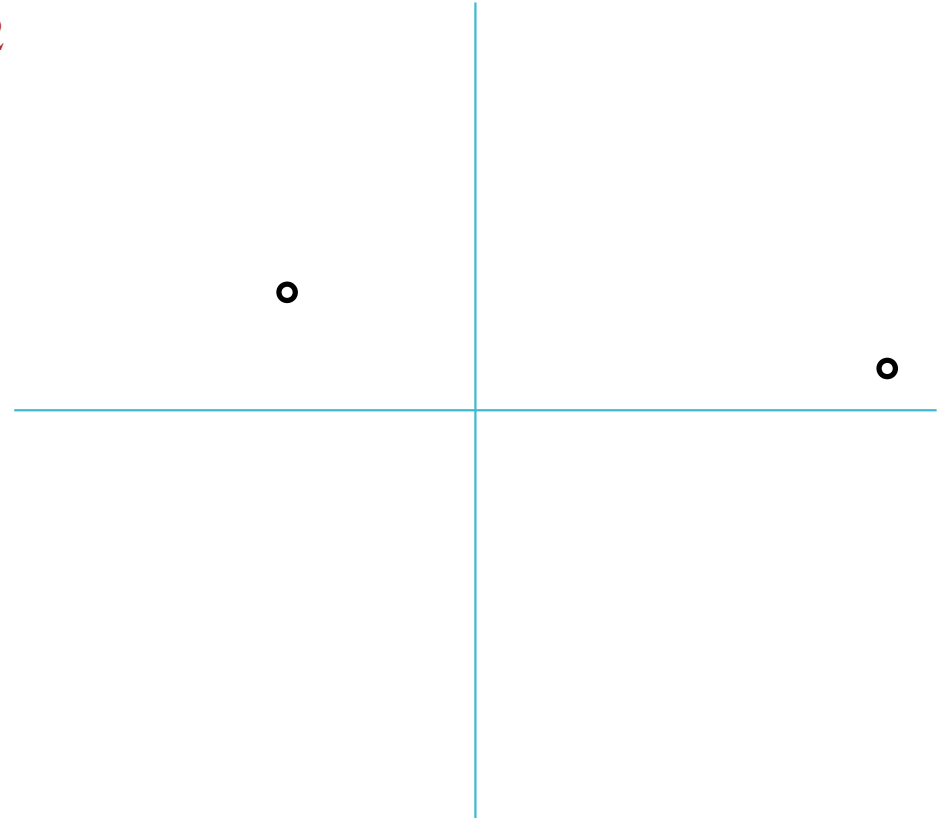
Counting h



$\mathcal{H}$ = halfspaces

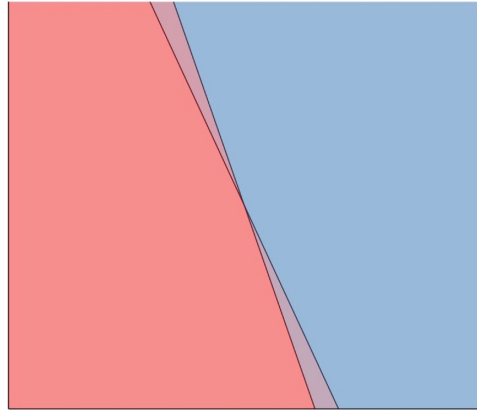
$$M = 2$$

$\mathcal{D} =$



- Fix $\mathcal{H}$
- Consider datasets $\mathcal{D}$ of size $M = 1, 2, \dots$
- For each dataset, count how many **actually distinct** hypotheses $h \in \mathcal{H}$ there are: $|\mathcal{H}(\mathcal{D})|$

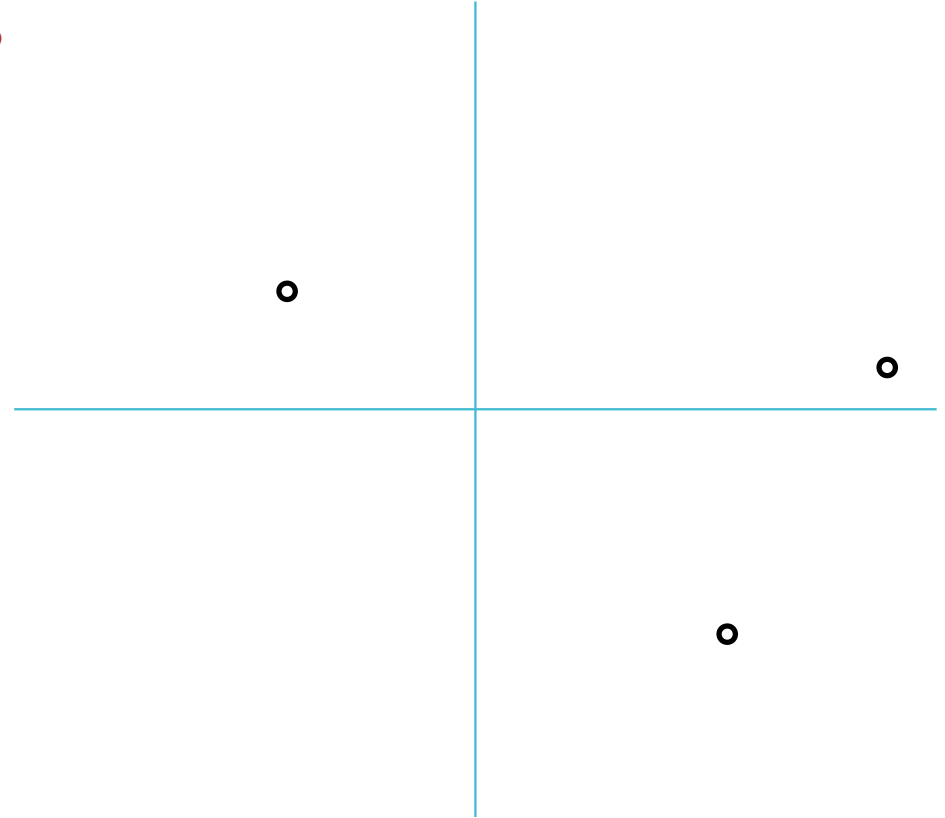
Counting h



$\mathcal{H}$ = halfspaces

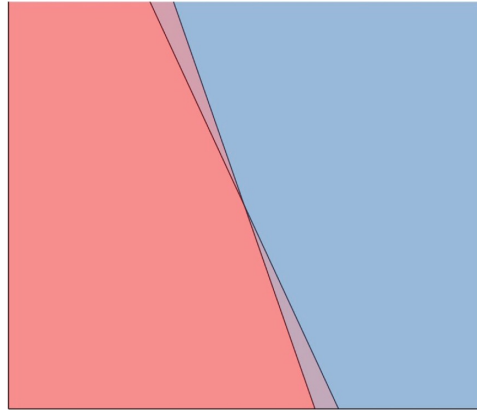
$$M = 3$$

$\mathcal{D} =$



- Fix $\mathcal{H}$
- Consider datasets $\mathcal{D}$ of size $M = 1, 2, \dots$
- For each dataset, count how many **actually distinct** hypotheses $h \in \mathcal{H}$ there are: $|\mathcal{H}(\mathcal{D})|$

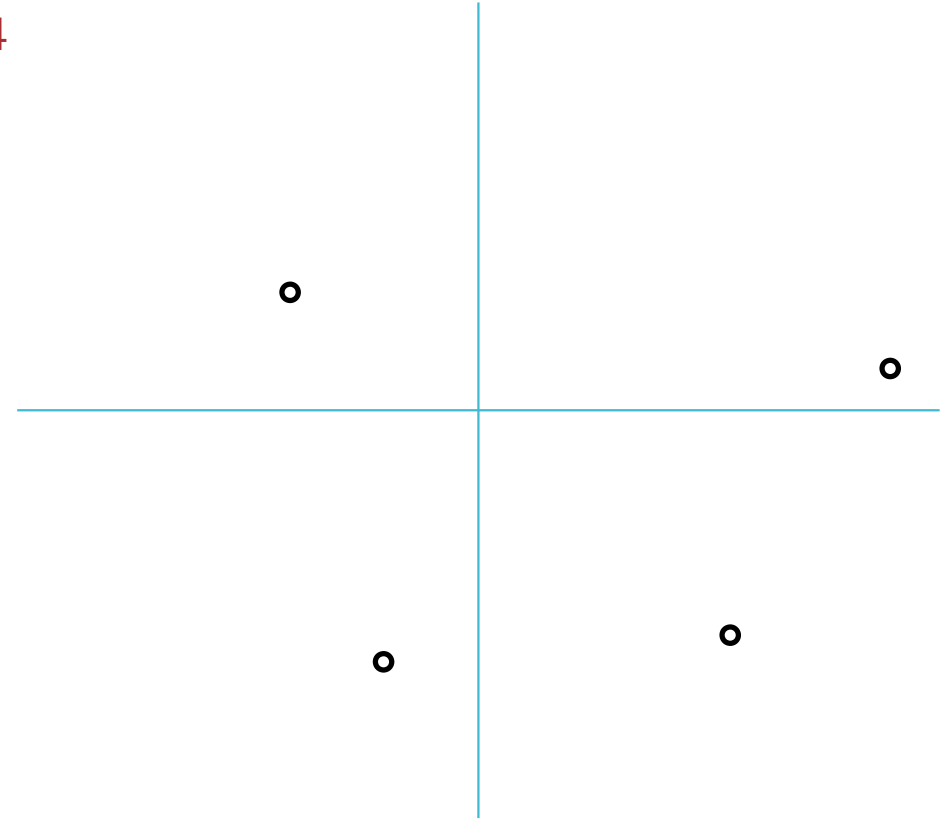
Counting h



$\mathcal{H}$ = halfspaces

$$M = 4$$

$\mathcal{D} =$



- Fix $\mathcal{H}$
- Consider datasets $\mathcal{D}$ of size $M = 1, 2, \dots$
- For each dataset, count how many **actually distinct** hypotheses $h \in \mathcal{H}$ there are: $|\mathcal{H}(\mathcal{D})|$

Growth function

- Def'n: the **growth function** $S_{\mathcal{H}}(M)$ is the maximum number of **distinct** $h \in \mathcal{H}$ for a dataset of size M

- for halfspaces in 2D,

M	1	2	3	4	5	...
$S_{\mathcal{H}}(M)$	2	4	8	14	≤ 26	...

- for larger M , it turns out $S_{\mathcal{H}}(M) = O(M^3)$

not obvious!

Growth function

Def'n: **shattering**

M	1	2	3	4	5	...
$S_{\mathcal{H}}(M)$	2	4	8	14	≤ 26	...

← shattered | not shattered →

$\mathcal{H}$ **shatters** a set of points if it can classify them all possible ways

Two kinds of behavior

- For many hypothesis classes $\mathcal{H}$, similar behavior:

$$S_{\mathcal{H}}(M) = \begin{cases} 2^M & M \leq d \quad \text{shattered} \\ \ll 2^M & M > d \quad \text{not shattered} \end{cases}$$

- e.g., intervals (or rectangles or hyperrectangles)
 - e.g., bounded-depth decision trees
 - e.g., fixed-architecture neural networks
- For many other classes $\mathcal{H}$, instead $S_{\mathcal{H}}(M) = 2^M$ for all M
 - e.g., unbounded-depth decision trees
 - e.g., unbounded-size neural networks

can shatter a
set of each size

Two kinds of behavior

Learnable (can't memorize more than d points)

- For many hypothesis classes $\mathcal{H}$, similar behavior:

$$S_{\mathcal{H}}(M) = \begin{cases} 2^M & M \leq d \quad \text{shattered} \\ \ll 2^M & M > d \quad \text{not shattered} \end{cases}$$

- e.g., intervals (or rectangles or hyperrectangles)
- e.g., bounded-depth decision trees
- e.g., fixed-architecture neural networks

- For many other classes $\mathcal{H}$, instead $S_{\mathcal{H}}(M) = 2^M$ for all M
can shatter a set of each size
 - e.g., unbounded-depth decision trees
 - e.g., unbounded-size neural networks

Not learnable (can memorize at any $|\mathcal{D}|$)

Sauer's lemma

- Suppose $S_{\mathcal{H}}(M) = 2^M$ for $M \leq d$, but $S_{\mathcal{H}}(d+1) < 2^{d+1}$
 → Then $S_{\mathcal{H}}(M) = O(M^d)$

“Suppose we grow exponentially (i.e., shatter) only up to $M = d$. Then for $M > d$ we grow polynomially, with degree d .”

related results derived multiple times: Sauer, Shelah, Perles, Vapnik/Chervonenkis

Sauer's lemma

d is called the **VC-dimension** of $\mathcal{H}$

- Suppose $S_{\mathcal{H}}(M) = 2^M$ for $M \leq d$, but $S_{\mathcal{H}}(d+1) < 2^{d+1}$

→ Then $S_{\mathcal{H}}(M) = O(M^d)$

“Suppose we grow exponentially (i.e., shatter) only up to $M = d$. Then for $M > d$ we grow polynomially, with degree d .”

related results derived multiple times: Sauer, Shelah, Perles, Vapnik/Chervonenkis

What do our
bounds tell us
w/ Sauer's
lemma?

- Finite realizable case:

$$R(h) \leq \frac{1}{M} \left(\ln(|\mathcal{H}|) + \ln\left(\frac{1}{\delta}\right) \right)$$

- Infinite realizable case:

$$R(h) \leq \frac{1}{M} \left(\ln(S_{\mathcal{H}}(M)) + \ln\left(\frac{1}{\delta}\right) \right)$$

What do our
bounds tell us
w/ Sauer's
lemma?

- Finite agnostic case:

$$R(h) \leq \hat{R}(h) + \sqrt{\frac{1}{2M} \left(\ln(|\mathcal{H}|) + \ln\left(\frac{2}{\delta}\right) \right)}$$

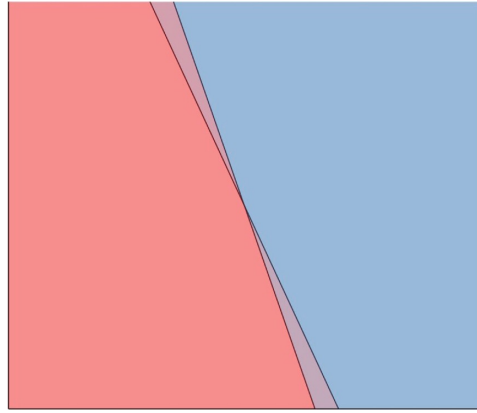
- Infinite agnostic case:

$$R(h) \leq \hat{R}(h) + \sqrt{\frac{1}{2M} \left(\ln(S_{\mathcal{H}}(M)) + \ln\left(\frac{2}{\delta}\right) \right)}$$

Finding the VC-dimension

- We defined VC-dimension of $\mathcal{H}$, $VC(\mathcal{H})$, as the size of the largest set S that $\mathcal{H}$ can shatter
 - If $\mathcal{H}$ can shatter arbitrarily large sets, $VC(\mathcal{H}) = \infty$
- To prove that $VC(\mathcal{H}) = d$, need to show
 1. $\exists$ some set of d data points that $\mathcal{H}$ can shatter and
 2. $\nexists$ a set of $d + 1$ data points that $\mathcal{H}$ can shatter

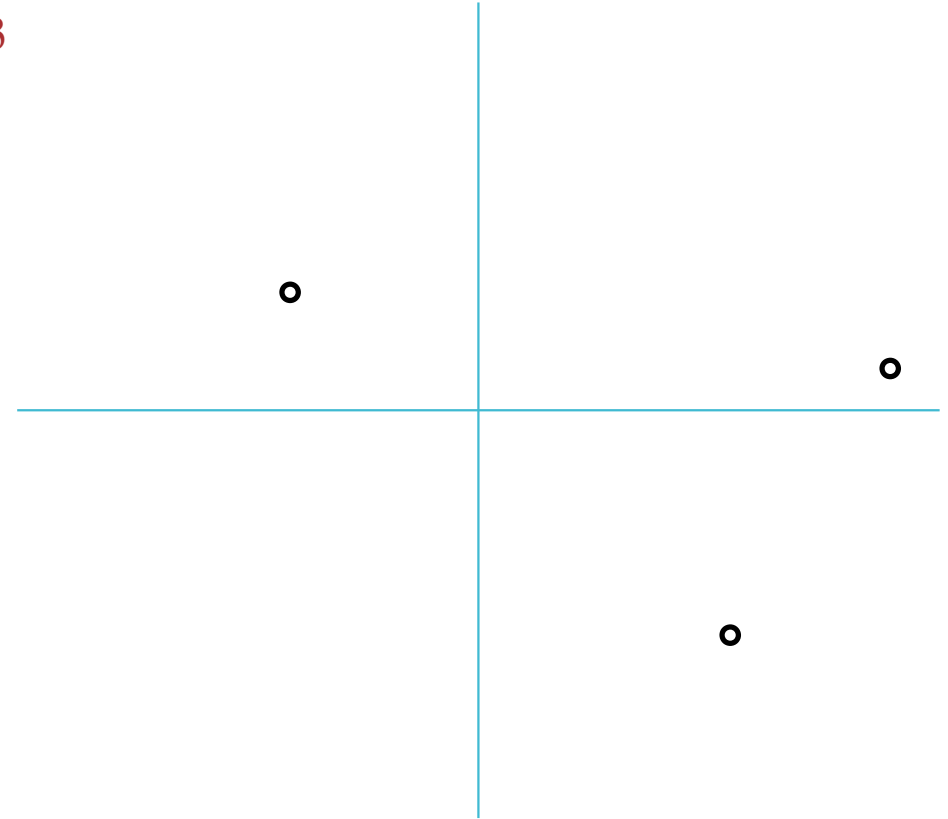
VC-dimension example



$\mathcal{H}$ = halfspaces

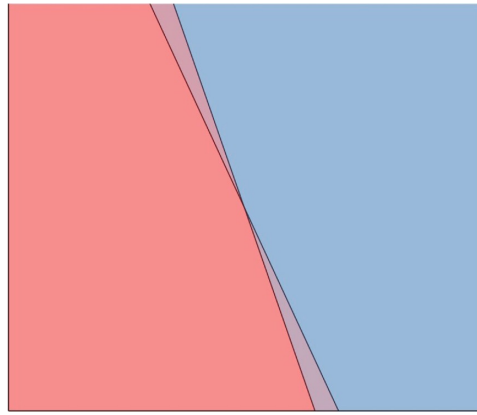
$$M = 3$$

$$\mathcal{D} =$$



- Before, we looked at this dataset of size 3

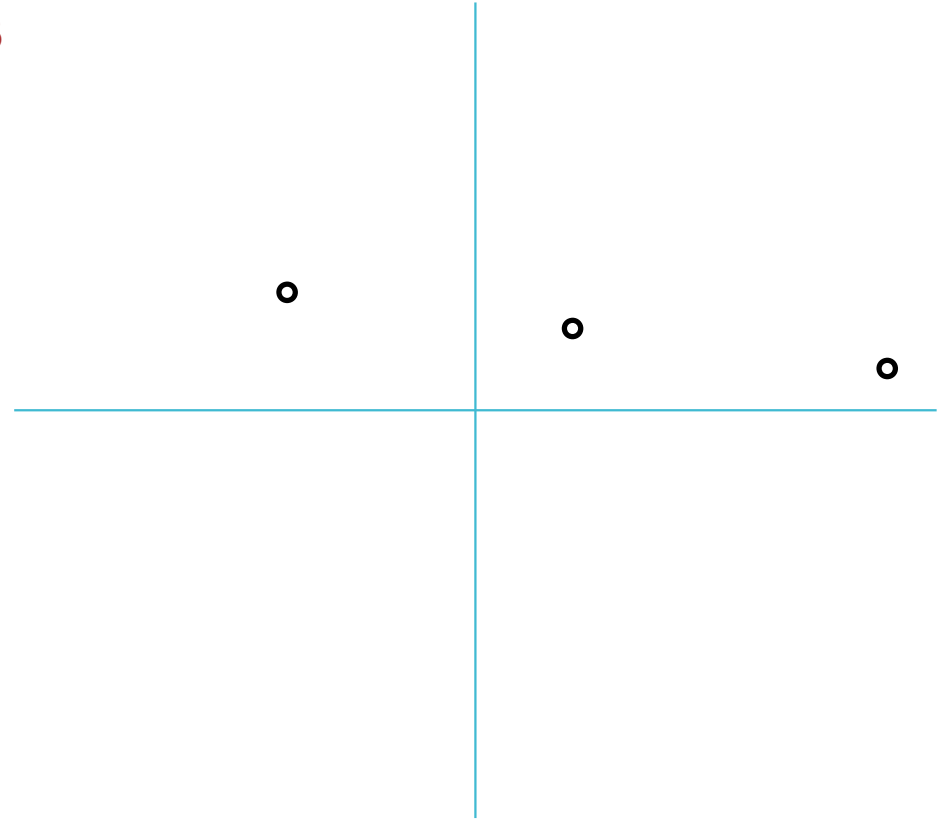
VC-dimension example



$\mathcal{H}$ = halfspaces

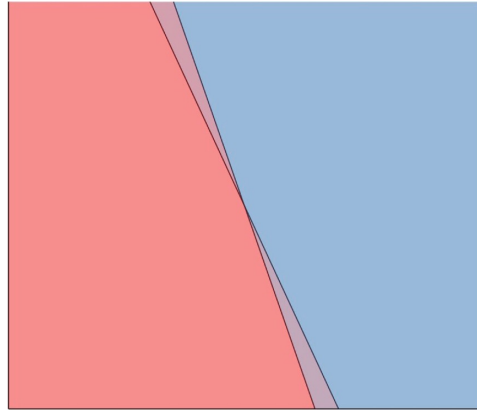
$$M = 3$$

$$\mathcal{D} =$$



- But what if we had looked at this one?

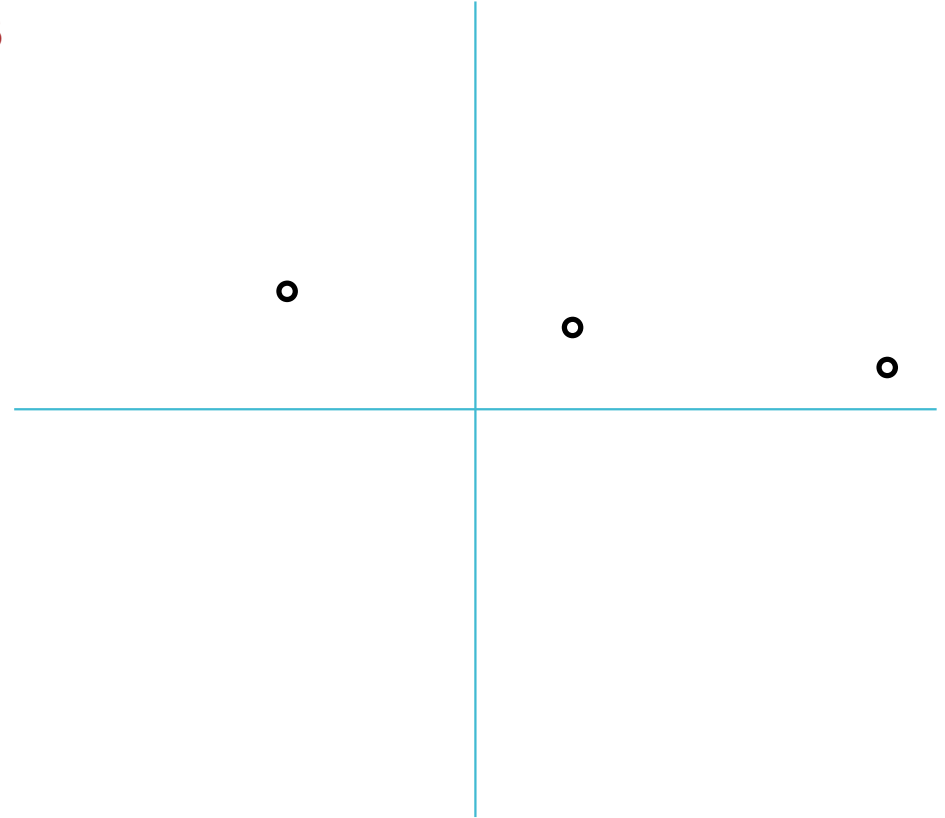
VC-dimension example



$\mathcal{H}$ = halfspaces

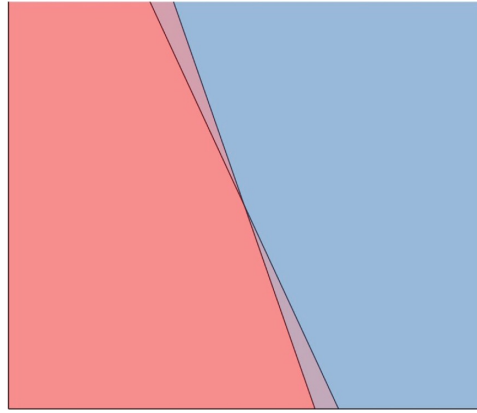
$$M = 3$$

$$\mathcal{D} =$$



- Only 6 distinct hypotheses $< 2^3$
- Tempting to say $\text{VC}(\mathcal{H}) < 3$ — but would be **wrong**

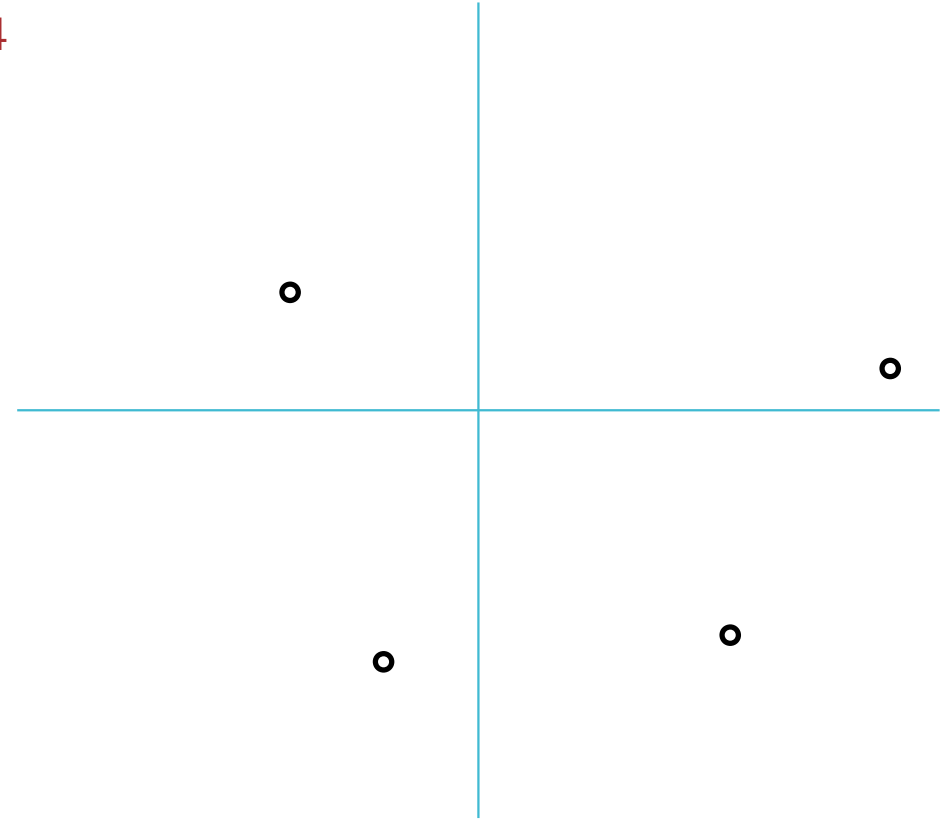
VC-dimension example



$\mathcal{H}$ = halfspaces

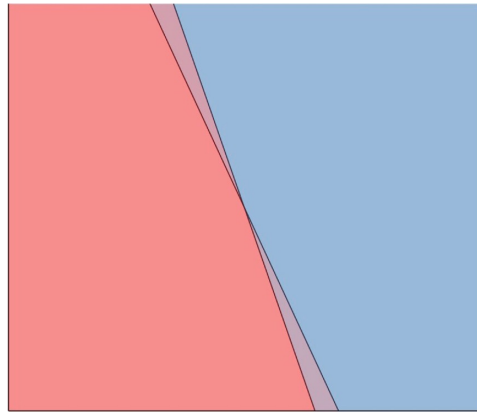
$$M = 4$$

$$\mathcal{D} =$$



- Similarly, looked at this dataset of size 4

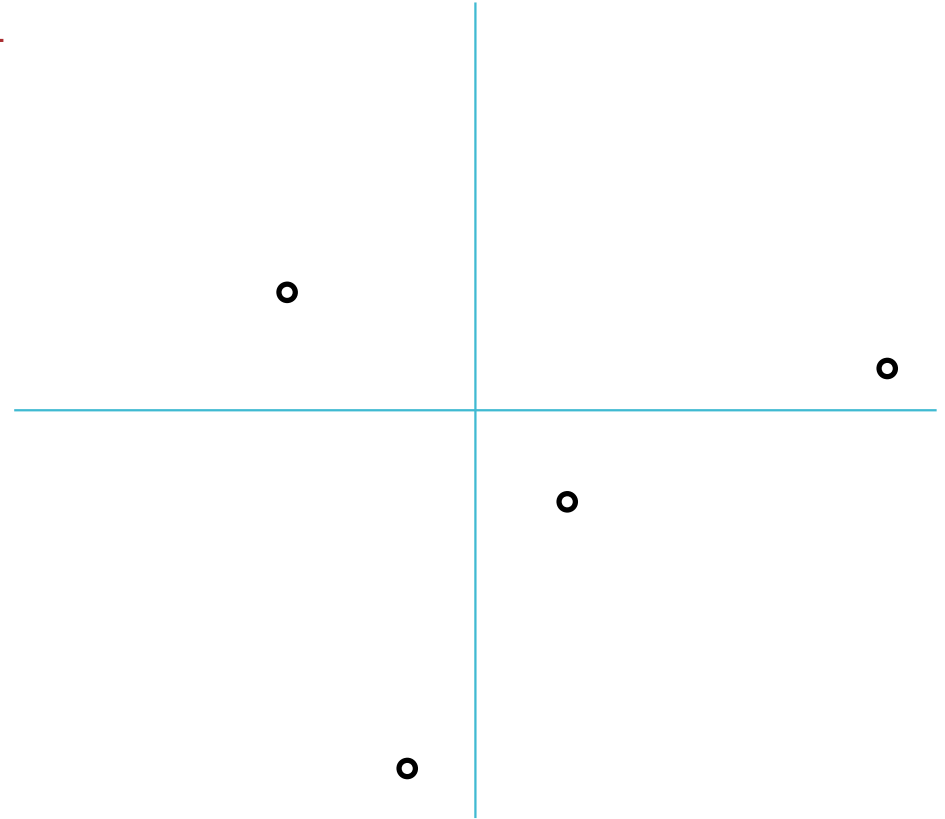
VC-dimension example



$\mathcal{H}$ = halfspaces

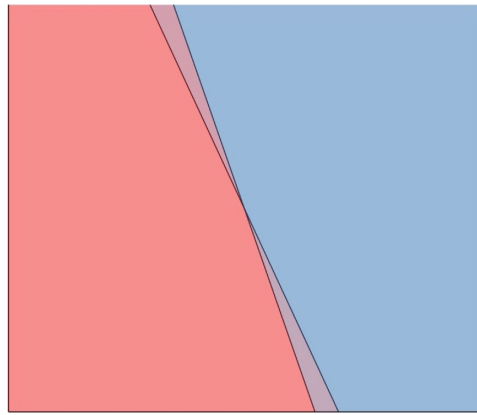
$$M = 4$$

$$\mathcal{D} =$$



- But really should have checked this one as well

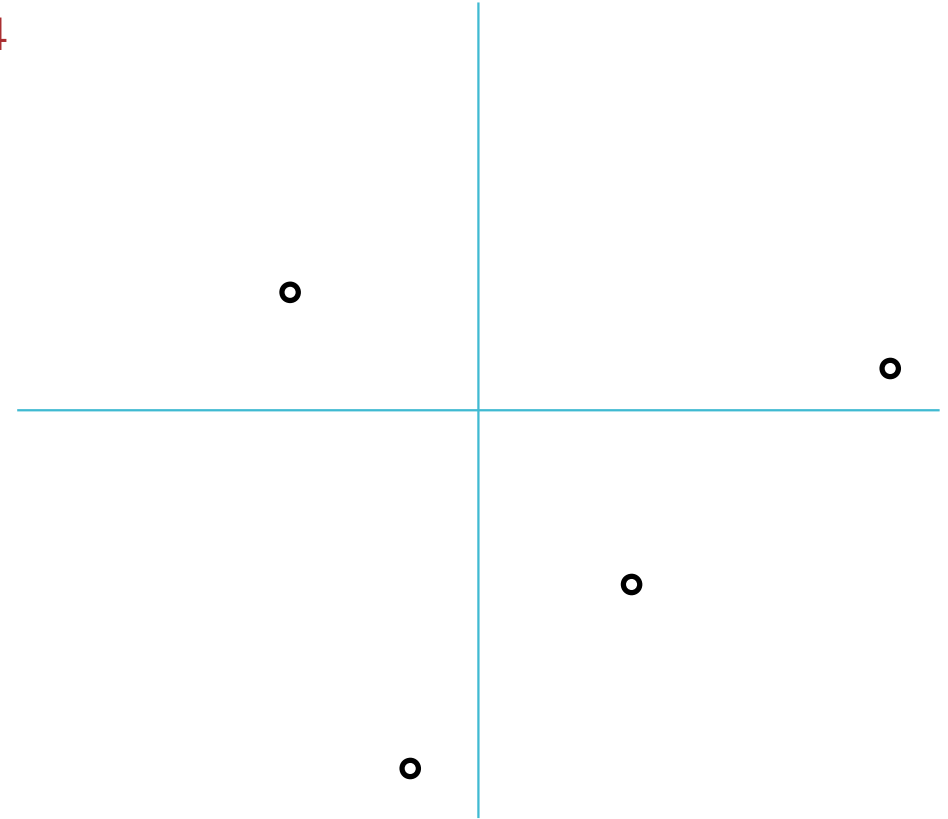
VC-dimension example



$\mathcal{H}$ = halfspaces

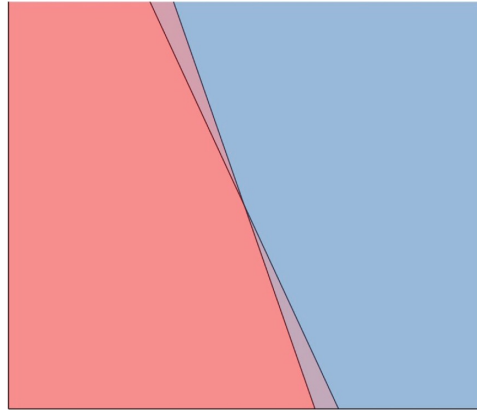
$$M = 4$$

$$\mathcal{D} =$$



- And this one

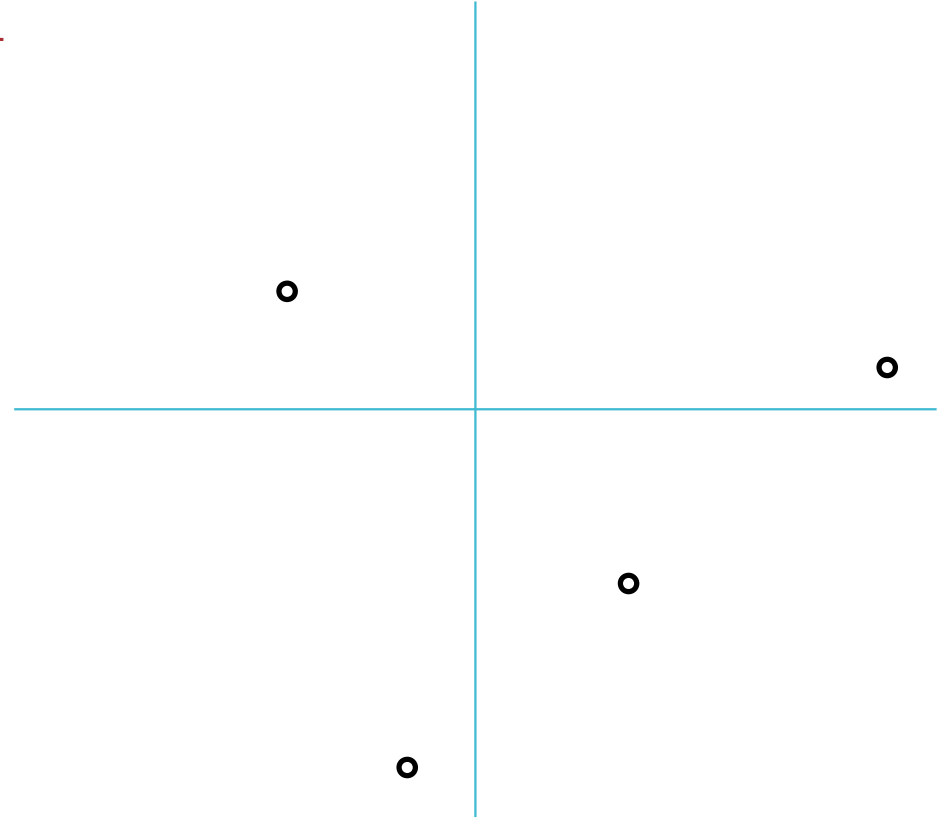
VC-dimension example



$\mathcal{H}$ = halfspaces

$$M = 4$$

$$\mathcal{D} =$$



- And this one

fortunately it turns out we're OK: no matter how we set up the points, can't shatter 4 of them

Halfspaces (linear separators)

- Just argued that halfspaces in 2D have $VC = 3$
- In general, halfspaces in d dimensions: $VC = d + 1$

More VC- dimension examples

- Try this at home: what is $VC(\mathcal{H})$ for
 - $\mathcal{H}$ = half-lines *where positive class is on right*
 - $\mathcal{H}$ = real intervals, positive when $x \in (a, b)$
 - $\mathcal{H}$ = axis-parallel rectangles in 2D (+ on interior)

Learning objectives

- You should be able to...
 - Identify properties of a learning setting, assumptions needed to ensure low generalization error
 - Distinguish true error, train error (and test, validation errors)
 - Define PAC: what is approximately correct and what occurs with high probability
 - Apply sample complexity bounds to real-world learning examples
 - Theoretically motivate regularization