

PAC LEARNING

11

PROBABLY APPROXIMATELY CORRECT

CLASSIFICATION:

INPUT

m data points

$(x_1, y_1), \dots, (x_m, y_m)$

x_i chosen i.i.d. from some distribution

$y_i = c(x_i)$

UNKNOWN CONCEPT
 $c: X \mapsto Y$

GOAL

learn c

APPROACH

- consider a finite set of hypotheses $\mathcal{H}$
(decision trees of depth d)
- find h with
$$\text{err}_{\text{train}}(h) = 0$$

(we say that h is CONSISTENT with data)

WHAT IS THE PROBABILITY THAT

$$\text{err}_{\text{true}}(h) > \epsilon \quad ?$$

$$\text{Pr}[h(x) \neq c(x)]$$

PROBABILITY OF A MISTAKE

BAD HYPOTHESES: $h_1, \dots, h_K$
(TRUE ERROR $> \epsilon$)

STEP I:

What is the probability that a fixed BAD hypothesis gets all data right?

say h_3 :

$$\Pr[h_3 \text{ makes an error}] > \epsilon$$

[BAD HYPOTHESIS]

$$\Pr[h_3 \text{ gets a single point right}] \leq 1 - \epsilon$$

m points

$$\leq (1 - \epsilon)^m$$

STEP II:

Prob. that any bad hypothesis gets all m points right

$$\leq K \cdot (1 - \epsilon)^m$$

[UNION BOUND]

$$\leq |\mathcal{H}| \cdot (1 - \epsilon)^m \leq |\mathcal{H}| \cdot e^{-m\epsilon}$$

[by $1 - x \leq e^{-x}$]

WITH ENOUGH DATA, WE MANAGE TO EXCLUDE ALL BAD HYPOTHESES WITH HIGH PROBABILITY

THEOREM [Haussler 1988]

For a finite hypothesis space $\mathcal{H}$, m i.i.d. training examples, $0 < \epsilon < 1$, and any h consistent w/ data:

$$\Pr[\text{err}_{\text{true}}(h) > \epsilon] \leq |\mathcal{H}| e^{-m\epsilon}$$

3

SUPPOSE WE WANT THIS PROBABILITY $< \delta$:

- How many examples suffice?

$$\text{WANT: } |\mathcal{H}| e^{-m\epsilon} \leq \delta$$

$$\text{REQUIRE: } m \geq \frac{\ln |\mathcal{H}| + \ln(1/\delta)}{\epsilon}$$

↑
SAMPLE COMPLEXITY

NOTE:

- grows gently with $|\mathcal{H}|$ and amount of confidence $1-\delta$
- can be large if we want ϵ to be small

- What is the largest error we get after perfectly fitting m samples?

$$\text{err}_{\text{true}}(h) \leq \frac{\ln |\mathcal{H}| + \ln(1/\delta)}{m}$$

REWRITE ERROR BOUND: $\frac{\text{\#bits to identify } h}{\log_2 e} \cdot \frac{\log_2 |\mathcal{H}| + \log_2(1/\delta)}{m}$

$$\text{err}_{\text{true}}(h) \leq \frac{1}{\log_2 e} \cdot \frac{\log_2 |\mathcal{H}| + \log_2(1/\delta)}{m}$$

If $\mathcal{H}$ finite, encode each h by its index: $1, \dots, |\mathcal{H}|$ - requires $\lceil \log_2 |\mathcal{H}| \rceil$ bits for each h .

Other coding schemes possible and the result still holds (if not necessarily finite):

$$\text{err}_{\text{true}}(h) \leq \frac{1}{\log_2 e} \cdot \frac{\text{size}(h) + \log_2(1/\delta)}{m}$$

OCCAM'S RAZOR: SHORTER HYPOTHESES THAT EXPLAIN THE DATA HAVE BETTER GENERALIZATION

4

EXAMPLE:

DECISION TREES FOR n BINARY ATTRIBUTES $x_0, \dots, x_{n-1}$

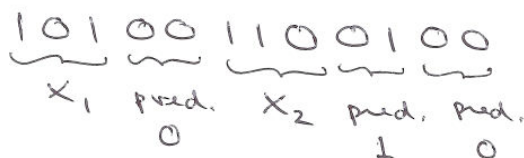
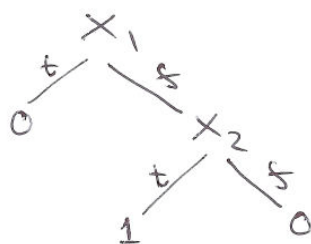
- define encoding recursively:

$$\text{ENCODE}(\text{leaf}) = \begin{cases} 00 & \text{if predicts 0} \\ 01 & \text{if predicts 1} \end{cases}$$

$$\text{ENCODE} \left(\begin{array}{c} x_i \\ \swarrow \quad \searrow \\ \text{LEFT} \quad \text{RIGHT} \end{array} \right) = 1 \cdot \langle i \text{ in binary} \rangle$$

- ENCODE(LEFT)
- ENCODE(RIGHT)

SAT: ATTRIBUTES x_0, x_1, x_2, x_3



k leaves, $k-1$ inner nodes: $\text{size}(h) = 2k + (k-1) \log n$ BITS

IF WE CAN FIT DATA WITH A TREE WHICH HAS k leaves:

$$\text{error}_{\text{tree}}(h) \leq \frac{1}{\log_2 2} \cdot \frac{2k + (k-1) \log_2 n + \log_2(V\delta)}{m}$$

< 1 (BOUND NON-TRIVIAL)
if $k < \frac{m}{\log(4n/\delta)}$

LIMITATIONS OF PAC LEARNING

(5)

- consistency requirement:

what if

concept $\notin \mathcal{H}$

or there is inherent noise in data?

- coding can be cumbersome

e.g.: how to encode half-spaces?

DROP CONSISTENCY REQUIREMENT:

- m i.i.d. data points,
but y_i not necessarily
a deterministic function of x_i

APPROACH

- find $h \in \mathcal{H}$ that makes fewest
mistakes on the data

NOW WE CAN SHOW:

$$\Pr[\text{err}_{\text{true}}(h) > \text{err}_{\text{train}}(h) + \epsilon] \leq |\mathcal{H}| \cdot e^{-2m\epsilon^2}$$

$\uparrow$ BEFORE: THIS = 0 $\uparrow$ BEFORE: THIS = $e^{-m\epsilon}$

PROVED USING THE UNION BOUND (OW "BAD" EVENTS)
AND Hoeffding's Inequality

Let $\mathcal{H} = \{h_1, \dots, h_{|\mathcal{H}|}\}$

[6]

BAD EVENT j : $j = 1, \dots, |\mathcal{H}|$

$$\text{err}_{\text{true}}(h_j) > \text{err}_{\text{train}}(h_j) + \epsilon$$

$$\Pr[\text{err}_{\text{true}}(h_j) - \text{err}_{\text{train}}(h_j) > \epsilon] \leq ???$$

=

DETOUR: CLASSIFICATION AS COIN FLIPS

INCORRECT $\hat{=}$ heads

CORRECT $\hat{=}$ tails

$\text{err}_{\text{true}}(h_j) \hat{=}$ PROBABILITY OF HEADS, θ

$\text{err}_{\text{train}}(h_j) \hat{=}$ MAXIMUM LIKELIHOOD ESTIMATE

$$\hat{\theta} = \frac{\# \text{misclassifications}}{m}$$

HOEFFDING INEQUALITY:

$$\Pr[\theta - \hat{\theta} > \epsilon] \leq e^{-2m\epsilon^2}$$

$\nwarrow$ RANDOM VARIABLE

GENERAL HOEFFDING BOUND:

$z_1, z_2, \dots, z_m$ indep. samples from
a distribution over $[0, 1]$
with mean μ

$$\Pr\left[\mu - \underbrace{\frac{1}{m} \sum_{i=1}^m z_i}_{\text{empirical average}} > \epsilon\right] \leq e^{-2m\epsilon^2}$$

7

SUMMARIZING INCONSISTENT HYPOTHESIS MODEL:

with prob $1-\delta$, for all $h \in \mathcal{H}$:

$$\text{err}_{\text{true}}(h) \leq \text{err}_{\text{train}}(h) + \sqrt{\frac{\ln|\mathcal{H}| + \ln(1/\delta)}{2m}}$$

or similarly as before

$$\underbrace{\text{err}_{\text{true}}(h)}_{\text{want small}} \leq \text{err}_{\text{train}}(h) + \sqrt{\frac{1}{\log 2} \cdot \frac{\text{size}(h) + \log_2(1/\delta)}{2m}}$$

fixed m, δ	for complex hypotheses	small	large
	for simpler hypotheses	large	small

"bias"

"variance"

DECREASES TO ZERO
w/ MORE DATA

NOTE: a worse dependence
on # samples (need more samples)
than the CONSISTENT CASE