

Last time: Specific noise models: ^{toy scenario} { • malicious noise (hard!) 😞
• random classification noise (easy!) 😊 }

- • Learning with malicious noise is hard 😞 → RCN
- Sad, sad positive result for
 - Start discussion of learning with RCN 😊

Today: Handling RCN.

- Revisit old PAC alg (for mon. conj.) that's not RCN-tolerant, adapt it to handle RCN
- Motivates new learning model: **learning from STATISTICAL QUERIES** (SQ)
- Any SQ learning alg. automatically yields a PAC alg. that can handle RCN.

Questions?

Let's consider... ^{PAC} learning mon. conj. over $\{0,1\}^n$,
with RCN.

Old alg:

- Start w/ hyp. $h(x) = x_1 \wedge x_2 \wedge \dots \wedge x_n$
Draw data set from $EX(c, \mathcal{D})$
- For each pos ex^x in data set,
if $x_i = 0$, remove x_i from h . (elim alg.)
- Works! (CHF thm, given large enough data set.)

This will fail badly under $EX^n(c, \mathcal{D})$: terrible idea to "trust" any given "pos. ex." in data set (might be noisy...)

Hi level idea of fix : use statistical properties of whole data set.

① • New version of alg, uses only ϕ , works w/ no noise
($EX(c, \mathcal{D})$)

② • adapt ① to setting of $EX^n(c, \mathcal{D})$.

①: Consider

(for $i \in [n]$)

$$p_i := \Pr_{(x, c(x)) \sim EX(c, \mathcal{D})} [c(x) = 1 \wedge x_i = 0]$$

(Note: if x_i in target $c(x)$, $p_i = 0$.)

Intuitively: if x_i not in $c(x)$, contrib. $\leq p_i$ to $err_{\mathcal{D}}(h, c)$.
having x_i in hyp. h

Goal: identify all x_i s.t. $p_i \geq \frac{\epsilon}{n}$.

If we cross all these off from h : tot. error of resulting h would be $\leq n \cdot \frac{\epsilon}{n} = \epsilon$.

So... to learn m conjs, suff. to get a $\pm \frac{\epsilon}{2m}$ acc. est. of each p_i , $i \in [n]$.

In noiseless setting (we have $EX(c, \mathcal{D})$):

easy! Draw m ex., take our estimate $\hat{p}_i$ of p_i to be:

$$\log(1/\delta) / \epsilon^2$$

(Frac. of the m ex. s.t. $c(x)=1, x_i=0$)

$$\tau = \frac{\epsilon}{10n}, \quad \sigma = \frac{\sigma}{n}$$

CB: take $m = O\left(\left(\frac{n}{\epsilon}\right)^2 \log\left(\frac{1}{\delta}\right)\right)$ tosses:
for each i ,

$$\Pr\{\text{est wrong by } > \frac{\epsilon}{10n}\} \leq \frac{\sigma}{n}$$

+ good :)

②: Goal: as before, est. $p_i := \Pr[c(x)=1 \mid x_i=0]$
 $(x, c(x)) \sim \text{EX}(c, \theta)$

(noiseless prob!), but we only have

$\text{EX}^n(c, \theta)$.

$$\Pr[A+B] = \Pr[A|B] \cdot \Pr[B]$$

Hummm... well, $p_i = \Pr[c(x)=1 \mid x_i=0] \cdot \Pr[x_i=0]$
 $(x, c(x)) \sim \text{EX}(c, \theta)$ $(x, c(x)) \sim \text{EX}(c, \theta)$

$= g_i$

• $\Pr[x_i=0]_{(x, c(x)) \sim \text{EX}(c, \theta)}$: event $x_i=0$ unaffected by noise :)

$$\Pr[x_i=0]_{(x, c(x)) \sim \text{EX}(c, \theta)}$$

$$\Pr[x_i=0]_{(x, c(x)) \sim \text{EX}^n(c, \theta)}$$

So can est.
easily, given
noisy data :)

• If $\Pr[x_i=0]_{(x, c(x)) \sim \text{EX}(c, \theta)}$ close to 0: good :): p_i also ≈ 0 .
 $(p_i \leq \text{this})$.

So can assume that not too small; say, $\geq \frac{\epsilon}{4n}$.

Remaining goal: estimate g_i .

Consider "noisy" version of g_i :

since —, can estimate this without too much slowdown!

$$\begin{aligned}
 \star &= \Pr[b=1 | x_i=0] \\
 &\quad (x, b) \sim \text{EX}(c, \theta) \\
 &= \underbrace{g_i}_{c=1} \cdot \underbrace{(1-n)}_{\text{no noise}} + \underbrace{(1-g_i)}_{c=0} \cdot \underbrace{n}_{\text{noise}} \\
 &= n + g_i(1-2n)
 \end{aligned}$$

So $g_i = \frac{\star - n}{1-2n}$: know est. of $\star$,
 know n ,
 so can est. g_i . 😊

What happened?

- Convinced ourselves: to learn, enough to est. some prob's. ←
- To est. desired prob. ^{noiseless} in noisy world:
 re-expressed $\downarrow$
 in terms of
 - part unaffected by noise
 - | predictably affected |

we'll do this more generally.

First, generalize

New Learning Model:

STATISTICAL ^{SQ} QUERY MODEL

(in + of itself... no noise in picture)

SQ model:

- no indiv. examples given to alg.
- alg. can estimate probabilities (of noiseless data).

Def: A predicate of a labeled ex (x, y) : a statement about (x, y) that's T or F.
 $x \in X, y \in \{0, 1\}$
 $X = \{0, 1\}^n$

Ex: "(x, y) has $x_1 = 1, x_2 = 0, \text{ and } y = 0$ "

Formally, a pred. is a fn $\chi: X \times \{0, 1\} \rightarrow \{0, 1\}$

Given a dist. $\mathcal{D}$ over X , a concept c , & a pred. χ , we write

$$P_\chi := \Pr_{(x, c(x)) \sim \text{EX}(c, \mathcal{D})} [\chi(x, c(x)) = 1]$$

Def: The oracle $\text{STAT}(c, \mathcal{D})$ takes 2 inputs from learner:

- a pred. χ , &
- a tolerance parameter $\tau > 0$.

$\text{STAT}(c, \mathcal{D})$ returns an est. $\hat{P}_\chi$ of P_χ s.t.

$$|\hat{P}_x - P_x| \leq \tau.$$

Ex: Learner calls $\text{STAT}(c, \mathcal{D})$ with
 $\pi = "x_i = 1 \vee y = 0"$, $\tau = 0.01$.

STAT returns 0.36: means

$$0.35 \leq \Pr_{(x, c(x)) \sim \text{EX}(c, \mathcal{D})} [c(x) = 0 \vee x_i = 1] \leq 0.37$$

Note: if had $\text{EX}(c, \mathcal{D})$:

can simulate having $\text{STAT}(c, \mathcal{D})$:

- given π, τ (input to STAT):
 - draw m ex $(x, c(x)) \sim \text{EX}(c, \mathcal{D})$.
 - evaluate π on each.
 - output $\hat{P}_x = \text{frac. of the } m \text{ ex. that satisfy } \pi$.

Taking $m = \Omega(\log(1/\delta) \cdot \frac{1}{\tau^2})$:

$$\Pr[|\hat{P}_x - P_x| > \tau] \leq \delta$$

(CB)

Cost of simulation: $\approx \frac{1}{\tau^2}$

involve cxity of
 evaluating π on $(x, c(x))$

So eff. SQ alg: doesn't call $\text{STAT}(c, \mathcal{D})$ with τ 's that are too small, or with κ 's that are too expensive to compute.

Def: Concept class $\mathcal{C}$ is SQ-learnable if there's a learning alg L s.t.:

- $\forall c \in \mathcal{C}$,
- $\forall \epsilon > 0$,
- $\forall \text{dist } \mathcal{D}$,

if L is given access to $\text{STAT}(c, \mathcal{D})$, then L outputs an h s.t. $\text{err}_{\mathcal{D}}(h, c) \leq \epsilon$.

We say $\mathcal{C}$ is efficiently SQ learnable by L if

- for every query (κ, τ) that L makes, $\tau \geq \frac{1}{\text{poly}(n, \frac{1}{\epsilon}, \text{size}(\mathcal{C}))}$ & $\kappa(x, y)$ is evaluable in time $\text{poly}(n, \frac{1}{\epsilon}, \text{size}(\mathcal{C}))$;
 - L runs in time $\text{poly}(n, \frac{1}{\epsilon}, \text{size}(\mathcal{C}))$, viewing each call to STAT as taking one time step.
-

We have:

Thm: Let $\mathcal{C}$ be eff. SQ learnable.

Then $\mathcal{C}$ is eff. PAC learnable.

Interesting thing (next time):

Thm: Let $\mathcal{C}$ be eff. SQ learnable.

Then $\mathcal{C}$ is eff. PAC learnable *even*
with RCN.
