

Last time: • discuss Winnow 1 $\rightarrow$ F(5) " δ -sep. max LTFs" LTF's

- Winnow 2 alg: OCMB alg. for certain LTF's (with positive weights, $X = \{0,1\}^n$)
- Perceptron alg. for LTFs "with a margin" over $\mathbb{R}^n$

Today:

- finish PCT proof; example of using PCT
- dual Perceptron, kernelization End of our "concrete examples"
- start generic bounds: "halving algorithm"

Questions?

Recall PCT setup:

$v =$ target vector $\|v\|=1$
 $w =$ hyp. vector
 $\hookrightarrow$ initially 0

Thm: (Perc. Convergence Theorem):

Suppose run Perc. on seq of ex labeled by $\frac{\text{sign}(v \cdot x)}{a_2}$, where $\frac{\|v\|}{a_2}$ each $\|x\|$ are ≤ 1 .

Let $\delta := \min_{\substack{\text{all} \\ \text{ex. } x \\ \text{negat}}} |v \cdot x|$. (= min Euclidean dist. from x to v 's hyperplane)

Then #mist Perc. makes is $\leq \frac{1}{\delta^2}$.

L1: After M mistakes, $v \cdot w \geq \delta \cdot M$

L2: After M mist., $w \cdot w \leq M$.

$$(\|w\| = \sqrt{w \cdot w} \leq \sqrt{M})$$

Proof of PCT: $\left(\begin{smallmatrix} \text{length of} \\ \text{proj of } w \text{ in dir. } v \end{smallmatrix} \right) \leq \left(\begin{smallmatrix} \text{length} \\ \text{of } w \end{smallmatrix} \right)$

$$\delta M \stackrel{L1}{\leq} v \cdot w \leq \|w\| = \sqrt{w \cdot w} \stackrel{L2}{\leq} \sqrt{M}, \text{ so}$$

$$\delta \sqrt{M} \leq 1, \text{ so } M \leq \frac{1}{\delta^2}.$$

$$v \cdot w = \|w\| \cdot \cos(\alpha)$$

b/c $v = \text{unit vector}$



Discussion:

- + simple!
- + noise-tolerant
- + fast
- + "kernelizable" (more soon...)

(better MB) $-\frac{1}{\delta^2}$: not great.

Faster LTF learning known, with m.b. $\approx \log \frac{1}{\delta}$ but these lack nice properties.

Ex: Class of "reoriented n -var MAJ's"

$$(n \text{ odd}) \quad \mathcal{C} = \{ \text{MAJ}(l_1, l_2, \dots, l_n) : l_i \text{ either } x_i \text{ or } \bar{x}_i \}$$

$$|\mathcal{C}| = 2^n$$

$$X = \{0, 1\}^n$$

ex: $n=3$ $\text{MAJ}(x, \bar{x}_2, \bar{x}_3)$

To use PCT: want each ex $\|x\|=1$,
target $\|v\|=1$, $\text{sign}(v \cdot x)$

Remap: $x \in \{0, 1\}^n$ $0 \rightarrow -\frac{1}{\sqrt{n}}$
 $\downarrow$
 $z \in \{-\frac{1}{\sqrt{n}}, \frac{1}{\sqrt{n}}\}^n$ $1 \rightarrow \frac{1}{\sqrt{n}}$

every $z \in \{-\frac{1}{\sqrt{n}}, \frac{1}{\sqrt{n}}\}^n$ has $\|z\|=1$. $\cup$

Have $\text{sign}(v \cdot z) \equiv \text{MAJ}(l_1, \dots, l_n)$ ^{x_i or $\bar{x}_i$}

by $x \rightarrow z$ as above,

$$v_i = \begin{cases} 1 & \text{if } l_i = x_i \\ -1 & \text{if } l_i = \bar{x}_i \end{cases}$$

$v \cdot z$ adds up $\pm z_1, \pm z_2, \dots, \pm z_n$

+ iff $\text{MAJ}(l_1, \dots, l_n) = 1$

$v \leftarrow \frac{1}{\sqrt{n}}$ this v :

$v_i = \pm \frac{1}{\sqrt{n}}$ so $\|v\|=1$

What's γ ? $\gamma = \min |v \cdot z|$

$v \cdot z = \left(\pm \frac{1}{\sqrt{n}} \cdot \pm \frac{1}{\sqrt{n}} \right) + \dots$

n of these

$|v \cdot z| \geq \frac{1}{n}$ so $\delta = \frac{1}{n}$ + Perc. will make $\leq n^2$ mistakes.

① Dual Perceptron + ② Kernelization

↳ Different, $\equiv$ way to run Perceptron
Clunkier... but enables "kernel" - way to run
Perc. over very high-dim feature space super-efficiently!

① Key insight: all it takes to run Perc. is being able to compute $x \cdot x'$ where x, x' are example vectors.

Recall how Perc works: $w_{init} = (0, \dots, 0)$
After 1st update, on example x' ,
 w is either x' or $-x'$.

To ^{keep} running Perceptron: on next ex x , $w \cdot x$ is $x' \cdot x$ or $-x' \cdot x$

2nd mist is on, say, x^2 , + new w is (say) $x' - x^2$
So $w \cdot x = x' \cdot x - x^2 \cdot x$ (Can keep going!)

More formal descr. of dual Perc:

Hyp. maintained as list of pairs

$$\textcircled{\star} \begin{cases} (x^1, y^1) \\ \vdots \\ (x^k, y^k) \end{cases}$$

example
 $(x^i, y^i) = i^{\text{th}}$ example Perc. got wrong.
"label $\in \{-1, 1\}$ "

Given new example x , compute

$$w \cdot x = \sum_{i=1}^k y^i (x^i \cdot x).$$

Perfectly simulates Perceptron.

Why? • More memory-intensive, slower "

• " " " We can replace "•"
(inner prod. over original ex)

with any "kernel function".

② Kernel functions + "feature expansions"

Let $X =$ original feature space ($\{0, 1\}^n$ or $\mathbb{R}^n$).

$\left(\begin{array}{l} X' = \text{new "high dim" feature space} \\ (\{0, 1\}^N \text{ or } \mathbb{R}^N) \end{array} \right.$

$\rightarrow$ both spaces have inner products.

Def: A feature expansion Φ is a mapping $X \rightarrow X'$.

Ex: $X = \{0, 1\}^n$, $X' = \{0, 1\}^{3^n}$

Φ = "all-conjunctions" feature expansion

$\Phi(x)$ has a coord for each of the 3^n
conj. over $x_1, \dots, x_n$.

$n=2$: $x = (x_1, x_2) \in \{0, 1\}^2$

$\Phi(x)$ = 9-bit string

$$= \left(\underset{\substack{\uparrow \\ \text{empty conj}}}{1}, x_1, \bar{x}_1, x_2, x_1 x_2, \bar{x}_1 x_2, \bar{x}_2, x_1 \bar{x}_2, \bar{x}_1 \bar{x}_2 \right)$$

$\Phi(x)$ = "expanded" version of x .

A kernel function K corresp. to Φ is a fn

$$K: X \times X \rightarrow \mathbb{R}, \quad K(a, b) = \underbrace{\Phi(a)}_{\substack{\uparrow \\ \text{in } X}} \cdot \underbrace{\Phi(b)}_{\substack{\uparrow \\ \text{in } X}} \in \mathbb{R}$$

Back to learning:

Suppose ^{we} get ex. in $X \xrightarrow{\Phi} \mathbb{R}^N$ but we want to
run Perc. over $X' \xrightarrow{\Phi} \mathbb{R}^N$ using $\Phi(x)$ for each ex. x .

||
• Obvious approach: do it directly: 
for each x we get, write down $\Phi(x)$,
maintain $w \in \mathbb{R}^N$ as hyp. vector, run Perc. on.

Drawback: takes time $\geq N$ per example - slow!

||
• Better way: use dual Perc., replacing each $a \cdot b$
with $K(a, b)$.

This exactly simulates running Perc. on $\Phi(\cdot)$ -versions of examples!

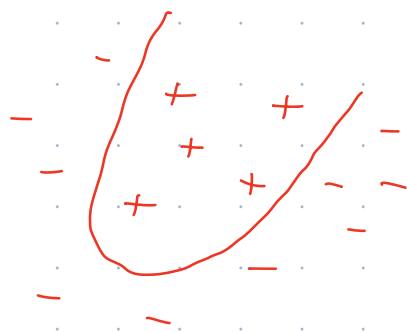
If we can evaluate $K(a,b)$ "fast", we can run Perc. over X' fast!

Nice to run Perc. over feature-expanded space:
• over $\mathbb{R}^n$, Perc. hyp. is like $\text{sign}(3x_1 + 4x_2 - 5x_3)$

• If we did a feature exp. to

$$\Phi(x) = (1, x_1, \dots, x_n, x_1^2, x_1x_2, \dots, x_1x_n, x_2^2, x_2x_3, \dots, x_n^2)$$

Perc hyp. can be $\text{sign}(5x_1^2 - 4x_1x_4 + 3x_2^2)$



Good news: For some interesting Φ 's, can eval. the associated $K(a,b)$ very fast!

Ex: "all-conj feature exp" Φ from above.

$$X = \{0,1\}^n \quad X' = \{0,1\}^{3^n}$$

Φ = "all-conjunctions" feature expansion

$\Phi(x)$ has a coord for each of the 3^n

conj. over $x_1, \dots, x_n$.

$$K(a, b) = \underline{\Phi(a) \cdot \Phi(b)}$$

Given $a, b \in \{0, 1\}^n$, define $\text{SAME}(a, b) \subseteq [n]$
 $\text{SAME}(a, b) = \{i \in [n] : a_i = b_i\}.$

Ex: $n = 16$

$$a = \overset{x_1}{1} \overset{x_2}{1} \overset{x_3}{1} \overset{x_4}{1} \overset{x_5}{1} \overset{x_6}{1} \overset{x_7}{1} \overset{x_8}{1} \overset{x_9}{0} \overset{x_{10}}{0} \overset{x_{11}}{0} \overset{x_{12}}{0} \overset{x_{13}}{0} \overset{x_{14}}{0} \overset{x_{15}}{0} \overset{x_{16}}{0}$$

$$b = \overset{x_1}{1} \overset{x_2}{1} \overset{x_3}{0} \overset{x_4}{0} \overset{x_5}{0} \overset{x_6}{0} \overset{x_7}{0} \overset{x_8}{1} \overset{x_9}{1} \overset{x_{10}}{1} \overset{x_{11}}{1} \overset{x_{12}}{0} \overset{x_{13}}{0} \overset{x_{14}}{0} \overset{x_{15}}{0} \overset{x_{16}}{0}$$

$\text{SAME}(a, b) = \{1, 2, 14, 15, 16\}.$

Claim: $K(a, b) = \underline{\sum |\text{SAME}(a, b)|}.$

Pf:

$$\Phi(a) = (1, \underset{\text{not in SAME}}{\underset{1}{a_1}}, \bar{a}_1, a_2, \bar{a}_2, \dots, \overset{1}{a_3}, \overset{1}{a_4}, \dots, \boxed{\bar{a}_{14} a_{15} a_{16}}, \dots, \bar{a}_{15} \bar{a}_{16}, \dots)$$

$$\Phi(b) = (1, \underset{\text{not in SAME}}{\underset{0}{b_1}}, \bar{b}_1, b_2, \bar{b}_2, \dots, \overset{1}{b_3}, \overset{1}{b_4}, \dots, \boxed{\bar{b}_{14} b_{15} b_{16}}, \dots, \bar{b}_{15} \bar{b}_{16}, \dots)$$

$\Phi(a) \cdot \Phi(b) = \# \text{ posit. where } \Phi(a) + \Phi(b) \text{ both are 1 in that pos.}$

Any conj containing a var outside $\text{SAME}(a, b)$:

gives 1.0 or 0.1 or 0.0 $\rightarrow$ 0

What abt conj. only over vars in SAME ?

For each $i \in \text{SAME}$

Must either not include x_i or include

$$\begin{cases} x_i & \text{if } a_i = b_i = 1 \\ \bar{x}_i & \text{if } a_i = b_i = 0. \end{cases}$$

So $\sum |SAME(a,b)|$ conj. that are sat.

both by a & by b ;

so $K(a,b) = \sum |SAME(a,b)| \rightarrow O(n)$ time

So... can run Perc. over space of all conj.,
& after k mistakes, it takes $O(k \cdot n)$
time per trial.

End of "specific" algo

Next time: generic algorithms