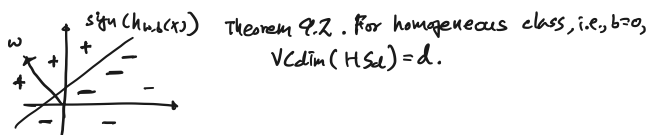


Chapter 9. Linear Predictors.

- class of affine functions $H_d = \{h_{w,b}(x) = \langle w, x \rangle + b : w \in \mathbb{R}^d, b \in \mathbb{R}\}$
 $\uparrow$ inner product
 $\langle w, x \rangle = w^T x = \sum_{i=1}^d w_i \cdot x_i$
- Prediction at x : $\hat{y} = \text{sign}(\langle w, x \rangle + b)$
- 0-1 loss: $h_{w,b}(x, y) = \mathbb{I}(\text{sign}(\langle w, x \rangle + b) \neq y)$
- Learning half spaces = Linear classifier for binary classification $\mathcal{Y} = \{-1, +1\}$

$$H_S = \text{sign} \circ H_d = \{ \text{sign}(h_{w,b}(x)) = \text{sign}(\langle w, x \rangle + b) : w \in \mathbb{R}^d, b \in \mathbb{R} \}$$

$\uparrow$
composition of a function and a class.



Proof $\triangleright$ ① $\exists C$ with $|C|=d$ s.t. H_S shatters C .

$C = \{e_1, e_2, \dots, e_d\}$ is shattered by H_S , because $\forall$ any $\{y_1, y_2, \dots, y_d\}$, $w = (y_1, y_2, \dots, y_d)$ gives $\text{sign}(\langle w, e_i \rangle) = y_i$

② $\forall C$ with $|C|=d+1$ cannot be shattered by H_S .

$C = \{x_1, x_2, \dots, x_d, x_{d+1}\}$, suppose C can be shattered by H_S
 Label: $y_1, y_2, \dots, y_d, y_{d+1}$
 s.t. $\text{sign}(a_1), \text{sign}(a_2), \dots, \text{sign}(a_{d+1}) = y_i$

$$\text{S.t. } a_1 \cdot x_1 + a_2 \cdot x_2 + \dots + a_{d+1} \cdot x_{d+1} = 0 \quad \leftarrow \text{this exists because } d+1 \text{ vectors are l.p.d.}$$

$I \triangleq \{i : a_i > 0\}$, $J \triangleq \{i : a_i < 0\}$, both cannot be empty.
 Suppose both non-empty.

$$\sum_{i \in I} a_i x_i = \sum_{j \in J} |a_j| x_j$$

$$0 < \sum_{i \in I} a_i \cdot \langle x_i, w \rangle = \left\langle \sum_{i \in I} a_i x_i, w \right\rangle = \left\langle \sum_{j \in J} |a_j| x_j, w \right\rangle = \sum_{j \in J} |a_j| \langle x_j, w \rangle < 0$$

$\uparrow$
 $0 < a_j \cdot \langle x_j, w \rangle$

• Hence, contradiction.

• If $\emptyset$ empty some contradiction with $\sum_{j \in J} |a_j| \langle x_j, w \rangle > 0$.

Theorem 9.3. $\text{VCdim}(H_S)$ with parameters w, b in $\mathbb{R}^d$ is $d+1$.

① $\exists C$ with $|C|=d+1$ shattered by H_S .

$$C = \{e_0, e_1, \dots, e_d\}$$

Label $y_0, y_1, \dots, y_d \rightarrow$ we let $b=y_0, w_1=y_1, \dots, w_d=y_d$. then $\forall x \in C, \langle w, x \rangle + b = y_i$

② Same proof in $\mathbb{R}^{d+1}$ and $d+2$ vectors.

$$\Rightarrow \text{ERM solves } L_D(H_S) \leq \epsilon$$

$\uparrow$
ERM

emp. loss

if $m \geq C \cdot \frac{d+1 \log \frac{1}{\delta}}{\epsilon}$, when $\mathcal{D}$ is ergodic *Hard if Not Realizable

*From Combinatorial version of Hoeffding's Theorem of statistical learning [Thm 6-8]

Q. How do you find ERM solution half spaces? (in the realizable, separable case) $\rightarrow \min_w L_S(w)$

① Linear Program: = optimization with linear objective and linear constraints.

$$\begin{aligned} \uparrow \\ \text{Convex optimization} \quad \max_{w \in \mathbb{R}^d} \quad & \langle u, w \rangle \\ \text{subject to} \quad & A \cdot w \leq v \quad \leftarrow \text{entrywise inequality} \end{aligned}$$

$$\min \begin{bmatrix} A \\ v \end{bmatrix} \geq y$$

- generally hard
- easy when realizable.

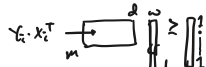
ERM ← L.P. solvers.

find w s.t. $\text{sign}(\langle w, x_i \rangle) = y_i$

$y_i \cdot \langle w, x_i \rangle > 0$. (such w exists under realizability)

claim $y_i \cdot \langle \tilde{w}, x_i \rangle \geq 1$, because we can use $\tilde{C} \cdot w$ instead with arbitrary C .

minimize 0
s.t. $y_i \cdot \langle w, x_i \rangle \geq 1, \forall i \in [m]$



② Another algorithm that finds an ERM for realizable linear classification.

Iterative algorithm:

input: $S = \{(x_1, y_1), \dots, (x_m, y_m)\}$

initialize: $w^{(1)} = (0, \dots, 0)$

for $t = 1, 2, \dots$

if $\exists i$ s.t. $y_i \cdot \langle w^{(t)}, x_i \rangle \leq 0$ then

$w^{(t+1)} \leftarrow w^{(t)} + y_i x_i$

else
output $w^{(t)}$.

We have $y_i \cdot \langle w, x_i \rangle > 0$.

$$\leq y_i \cdot \langle w^{(t)}, x_i \rangle$$

$$= y_i \cdot \langle w^{(t)} + y_i x_i, x_i \rangle$$

$$= y_i \cdot \langle w^{(t)}, x_i \rangle + y_i^2 \|x_i\|^2$$

$$> y_i \cdot \langle w^{(t)}, x_i \rangle$$

Theorem 9.1

Assume S is separable, $B = \min\{\|w\| : \forall i \in [m] y_i \cdot \langle w, x_i \rangle \geq 1\}$

$$R = \max_i \|x_i\|.$$

Dependent in B is suboptimal

then Perceptron Algorithm stops after at most $(RB)^2$ iterations.

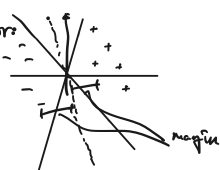
• when it stops it finds an ERM with $y_i \cdot \langle w^{(t)}, x_i \rangle > 0, \forall i \in [m]$

• perceptron is PAC learner with $m_{\mathcal{H}}(\epsilon, S) \geq \frac{\log \frac{d}{\epsilon}}{\epsilon}$

Proof

Let $w^* \in \arg \min_w \|w\|$ s.t. $y_i \cdot \langle w, x_i \rangle \geq 1, \forall i \in [m]$

$\Rightarrow$ max margin separator



$$\text{Claim: } \langle w^*, w^{(t+1)} \rangle$$

$$\geq \cos \angle w^*, w^{(t+1)} = \frac{\langle w^*, w^{(t+1)} \rangle}{\|w^*\| \cdot \|w^{(t+1)}\|} \geq \frac{T}{RB} \Rightarrow T \leq (RB)^2$$

$$\uparrow \quad \uparrow \quad \uparrow$$

$$\|w^*\| = B, \langle w^*, w^{(t+1)} \rangle \geq T, \|w^{(t+1)}\| \leq RT$$

$$(*) \quad \langle w^*, w^{(t+1)} \rangle - \langle w^*, w^{(t)} \rangle = \langle w^*, w^{(t+1)} - w^{(t)} \rangle$$

$$= \langle w^*, y_i x_i \rangle$$

$$= y_i \cdot \langle w^*, x_i \rangle \geq 1$$

$$\Rightarrow \langle w^*, w^{(t+1)} \rangle = \sum_{i=1}^t \langle w^*, w^{(i)} \rangle - \langle w^*, w^{(0)} \rangle \geq T$$

(***)

$$\begin{aligned} \|w^{(t+1)}\|^2 &= \|w^{(t)} + y_i x_i\|^2 \\ &= \|w^{(t)}\|^2 + y_i^2 \|x_i\|^2 + 2 y_i \langle w^{(t)}, x_i \rangle \\ &\leq \|w^{(t)}\|^2 + R^2 \end{aligned}$$

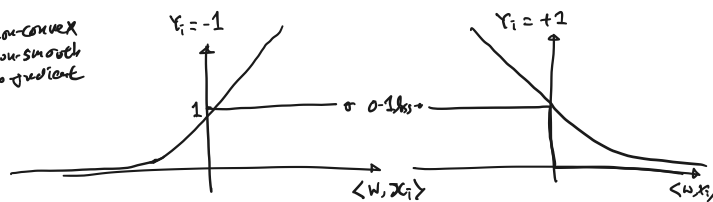
③ Logistic Regression: surrogate loss

True loss: 0-1 loss; $\ell(w, x) = \mathbb{I}(\text{sign}(\langle w, x \rangle) \neq y)$

$y_i = -1$

$y_i = +1$

• non-convex
• non-smooth
• no gradient



simple loss = logistic loss: $\ell(w; x, y) = \log(1 + \exp(-y \langle w, x \rangle))$
 convex ERM: minimize $\frac{1}{n} \sum_{i=1}^n \log(1 + \exp(-y_i \langle w, x_i \rangle))$
 smooth

* Online Perceptron Algorithm [much more on online learning at CSE 541, Interactive Machine Learning]

↓
 Data come in an online fashion.
 $1, 2, \dots, t, t+1, \dots$
 receive x_t
 choose h_t
 predict $h_t(x_t)$
 receive y_t
 pay loss: $\ell_t(x_t, y_t)$

Initialize $w^0 = (0, 0, \dots, 0)$
 for $t = 1 \dots T$
 receive x_t
 predict $p_t = \text{sign}(\langle w^{t-1}, x_t \rangle)$
 If $y_t \langle w^{t-1}, x_t \rangle \leq 0$
 $w^{(t)} \leftarrow w^{(t-1)} + y_t x_t$
 else $w^{(t)} \leftarrow w^{(t-1)}$

we care about, in an online learning problem, how many mistakes you make.

Theorem 21.16. $R = \max_t \|x_t\|$, if realizable, i.e. $\exists w^*$ st $y_t \langle w^*, x_t \rangle \geq 1 \forall t$

$$|\# \text{ of mistakes}| \leq R^2 \|w^*\|^2$$

proof [HW1 prob 6.]