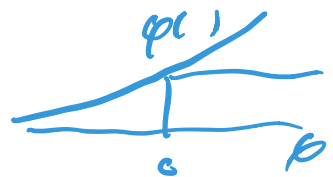


8.2 Binary classification w/surrogate loss

Thm 8.2 $(X, \mathcal{P}, \mathcal{F}, \phi)$ $\ell_\phi(y, f(x))$
 $\text{sgn}(f(x))$ $= \phi(-y f(x))$

(single version) w/p $\geq 1 - e^{-2t^2}$



$$L(f) = A_\phi(f) \leq A_{\phi,n}(f) + 4M_\phi \underbrace{ER_n(\mathcal{F}(x^n))}_{\forall f \in \mathcal{F}} + \frac{Bt}{\sqrt{n}}$$

holds for all functions $f_n \in \mathcal{F}$.

(double version) For $\hat{f}_n$ - ERM for surrogate loss

$$L(\hat{f}_n) = A_\phi(\hat{f}_n) \leq A_\phi^*(\mathcal{F}) + 8M_\phi \underbrace{ER_n(\mathcal{F}(x^n))}_{\forall f \in \mathcal{F}} + \frac{2Bt}{\sqrt{n}}$$

8.3 Weighted linear combinations of classifiers

$$(X, P, 0-1 \text{ loss})$$

$$g = g: X \rightarrow \{-1, 1\} \quad \text{"weak set of classifiers"}$$

For $\lambda > 0$, let

$$\mathcal{F}_\lambda = \left\{ f = \sum_{i=1}^N c_i g_i : N \geq 1, \sum_{i=1}^N |c_i| \leq \lambda, \begin{matrix} \parallel \\ g_1, \dots, g_N \in \mathcal{G} \end{matrix} \right\}$$

~~sgn(f)~~

$$= \lambda \text{absconv}(g)$$

$$\mathcal{F}_\lambda(x^n) = \lambda \text{absconv}(\underbrace{g(x^n)}_{\parallel})$$

$$\{(g(x_1), \dots, g(x_n)) : g \in \mathcal{G}\}$$

$$R_n(\mathcal{F}_\lambda(x^n)) = \lambda R_n(g(x_n))$$

$$\leq \lambda C \sqrt{\frac{V(g)}{n}}$$

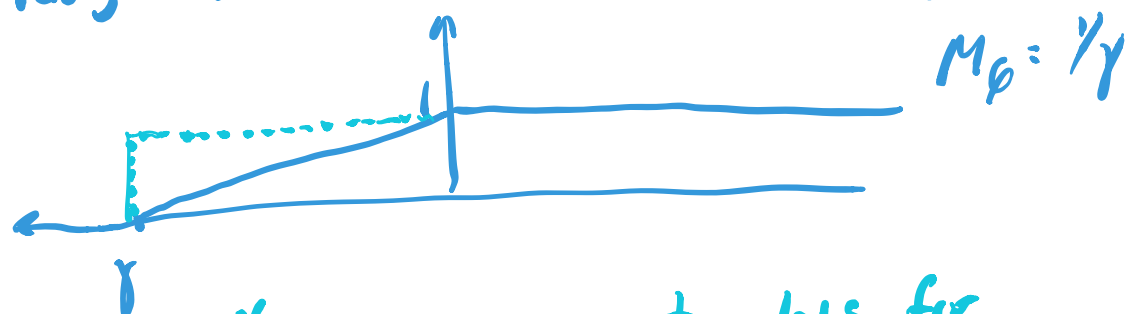
VC
dim
 Dudley
version.

Then For any A , w/p $\geq 1-\delta$:

$$L(\hat{f}_n) \leq A \phi_n(\hat{f}_n) + C \lambda M_\phi \sqrt{\frac{\log(1/\delta)}{n}} + \sqrt{\frac{1.5 C \lambda \log(1/\delta)}{2n}}$$

for all $\hat{f}_n \leftarrow$

Margin based risk bound. ramp function



Let $\underline{L}_n^\gamma(f)$ = surrogate loss for the upper bound
 $\tilde{\phi}(x) = 1, x \geq -\gamma.$

If $\underline{\gamma} f(x) > \gamma$ - no error, and
 even $\ell_p(\gamma, f(x)) = 0.$

If $0 < \gamma f(x) \leq \gamma$ - no error, but
 $\ell_p(\gamma, f(x)) > 0$
 $\tilde{\phi}(\gamma, f(x)) = 1$

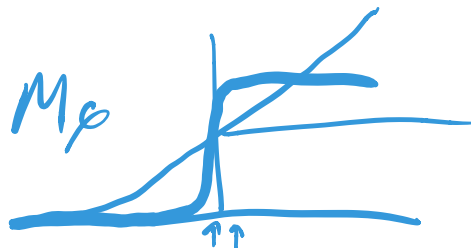
Thm 8.5 Margin based risk bound

For 170, with probability $\geq 1-\delta$:

$$L(\hat{f}_n) \leq L_n^\gamma(\hat{f}_n) + \underbrace{\frac{C\lambda}{\gamma} \sqrt{\frac{V(g)}{n}}}_{\substack{\uparrow \text{fraction of test} \\ \text{samples not strongly} \\ \text{correct (i.e. not} \\ \text{having } f(x)\gamma \geq \gamma}} + \underbrace{\sqrt{\frac{1.6511\lambda}{2n}}}_{\gamma\%}$$

Let $C_f = \{x : \text{sgn}(f) = 1\} - C$
 $\{\text{sgn}(f) : f \in \mathcal{F}\}$ VC dimension?

$$\mathcal{F}_\lambda = \left\{ f = \sum_{i=1}^N c_i g_i : N \geq 1, \sum_{i=1}^N |c_i| \leq \lambda, \begin{matrix} \parallel \\ g_1, \dots, g_N \in \mathcal{G} \end{matrix} \right\}$$



except on event w/p e^{-2t^2} $\varphi: \mathbb{R} \rightarrow [0,1]$

$$L(f) = A_p(f) \leq A_{p_n}(f) + 4M_p ER_n(\mathbb{E}(X^n)) + \frac{t + \sqrt{\log k}}{\sqrt{n}}$$

$$e^{-2(t + \sqrt{\log k})^2} = e^{-2t^2 - \log k^2 - 2t\sqrt{\log k}}$$

$$\leq \frac{1}{k^2} e^{-2t^2}$$

$$\sum_{k=1}^{\infty} \frac{1}{k^2} e^{-2t^2} = \frac{\pi^2}{6} e^{-2t^2} \leq \underline{\underline{2e^{-2t^2}}}$$

ADA Boost

$$\mathcal{F}_1 = \left\{ f = \sum_{k=1}^K \alpha_k g_k : K \geq 1, \alpha_1, \dots, \alpha_K \in \mathbb{R}, g_1, \dots, g_K \in \underline{\underline{\mathcal{G}}} \right\}$$

$$\phi(x) = \exp(x)$$

Algorithm - Find $g_1, g_2, \dots$ greedily attempting to minimize the surrogate loss.

Once $g_1, \dots, g_k$ are found, select β_{k+1} to minimize the empirical loss for a weighted version of training data $x_1, \dots, x_n$.

$$\phi(x) = e^x$$

$$y_i \in \{1, -1\}$$

$$A_{\theta, n}^{K+1} = \frac{1}{n} \sum_{i=1}^n \exp(-Y_i \underbrace{\sum_{k=1}^{K+1} \alpha_k g_k(x_i)}_n)$$

Suppose $\underbrace{g_1, \dots, g_K}_{\alpha_1, \dots, \alpha_K}$ | found. g_{K+1}? binary valued

$$A_{\theta, n}^{K+1} = \frac{1}{n} \sum_{i=1}^n w_i^{(K)} \exp(-Y_{K+1} \alpha_{K+1} g_{K+1}(x_i))$$

$$w_i^{(K)} = \exp(-Y_i \underbrace{\sum_{k=1}^K \alpha_k g_k(x_i)}_{\text{predictor of } Y_i \text{ based on } g_1, \dots, g_K(x_i)})$$

predictor of Y_i
based on $g_1, \dots, g_K(x_i)$

$$\begin{aligned} \exp(-Y_{K+1} \alpha_{K+1} g_{K+1}(x_i)) &= \exp(-\alpha_{K+1}) && \text{if } Y_i = g_{K+1}(x_i) \\ &= \exp(+\alpha_{K+1}) && \text{if } Y_i \neq g_{K+1}(x_i) \end{aligned}$$

es. select g_{K+1} to
minimize weighted empirical
0-1 loss using $w_i^{(K)}$ weights.

$\mathcal{G}$ is a set of simple binary classifiers

