

Summary of PAC bounds for finite model classes

With probability $\geq 1-\delta$,

1) For all $h \in \underline{H}$ s.t. $\text{error}_{\text{train}}(h) = 0$,

$$\text{error}_{\text{true}}(h) \leq \varepsilon = \frac{\ln |\underline{H}| + \ln \frac{1}{\delta}}{\underline{m}}$$

Haussler's bound

2) For all $h \in H$

$$|\text{error}_{\text{true}}(h) - \text{error}_{\text{train}}(h)| \leq \varepsilon = \sqrt{\frac{\ln |H| + \ln \frac{2}{\delta}}{2m}}$$

Hoeffding's bound

PAC bound and Bias-Variance tradeoff

$$P(|\text{error}_{\text{true}}(h) - \text{error}_{\text{train}}(h)| \geq \epsilon) \leq 2|H|e^{-2m\epsilon^2} \leq \delta$$

- Equivalently, with probability $\geq 1 - \delta$

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

- Fixed m

Model class		
complex	small	large
simple	large	small

What about the size of the model class?

$$2|H|e^{-2m\epsilon^2} \leq \delta$$

- Sample complexity

$$m \geq \frac{1}{2\epsilon^2} \left(\ln |H| + \ln \frac{2}{\delta} \right)$$


- How large is the model class?

Number of decision trees of depth k

$$m \geq \frac{1}{2\epsilon^2} \left(\ln |H| + \ln \frac{2}{\delta} \right)$$
 $\sim 2^k$
$$\Rightarrow |H| \sim 2^{2^k}$$

$$H_0 = 2$$

$$H_k = (\text{\#choices of root attribute})$$

*(# possible left subtrees)

$$*(\# \text{ possible right subtrees}) = n * H_{k-1} * H_{k-1}$$

Write $L_k = \log_2 H_k$

$$\rightarrow L_k = \log_2 n + \underbrace{2 \log_2 H_{k-1}}_{L_{k-1}}$$

$$L_0 = 1$$

$$\begin{aligned} L_k &= \log_2 n + 2L_{k-1} = \log_2 n + 2(\log_2 n + 2L_{k-2}) \\ &= \log_2 n + 2\log_2 n + 2^2\log_2 n + \dots + 2^{k-1}(\log_2 n + 2L_0) \end{aligned}$$

So $L_k = (2^k - 1)(1 + \log_2 n) + 1$

PAC bound for decision trees of depth k

$$\underline{m} \geq \frac{\ln 2}{2\epsilon^2} \left((\underline{2^k} - 1)(1 + \log_2 n) + 1 + \log_2 \frac{2}{\delta} \right)$$

- Bad!!!
 - Number of points is exponential in depth k !
- But, for m data points, decision tree can't get too big...

Number of leaves never more than number data points

Number of decision trees with k leaves

$$m \geq \frac{1}{2\epsilon^2} \left(\ln |H| + \ln \frac{2}{\delta} \right)$$

H_k = Number of binary decision trees with k leaves

$$H_1 = 2$$

→ $H_k = (\text{\#choices of root attribute}) * \begin{matrix} k-1 & \dots & 2 & 1 & \boxed{} & \boxed{} & k-1 & k-2 & \dots & 1 \end{matrix}$

$[(\text{\# left subtrees with 1 leaf}) * (\text{\# right subtrees with k-1 leaves})$
 $+ (\text{\# left subtrees with 2 leaves}) * (\text{\# right subtrees with k-2 leaves})$
 $+ \dots$
 $+ (\text{\# left subtrees with k-1 leaves}) * (\text{\# right subtrees with 1 leaf})]$

$$H_k = n \sum_{i=1}^{k-1} H_i H_{k-i} = n^{k-1} C_{k-1} \quad (C_{k-1} : \text{Catalan Number})$$

Loose bound (using Sterling's approximation):

$$\underline{H_k} \leq n^{k-1} \underline{2^{2k-1}}$$

Number of decision trees

- With k leaves $m \geq \frac{1}{2\epsilon^2} \left(\ln |H| + \ln \frac{2}{\delta} \right)$

$$\log_2 H_k \leq (k - 1) \log_2 n + 2k - 1 \quad \text{linear in } k$$

number of points m is linear in #leaves

- With depth k

$$\log_2 H_k = (2^k - 1)(1 + \log_2 n) + 1 \quad \text{exponential in } k$$

number of points m is exponential in depth

PAC bound for decision trees with k leaves – Bias-Variance revisited

With prob $\geq 1-\delta$ $\text{error}_{\text{true}}(h) \leq \text{error}_{\text{train}}(h) + \sqrt{\frac{\ln |H| + \ln \frac{2}{\delta}}{2m}}$

With $H_k \leq n^{k-1} 2^{2k-1}$, we get

$$\text{error}_{\text{true}}(h) \leq \text{error}_{\text{train}}(h) + \sqrt{\frac{(k-1) \ln n + (2k-1) \ln 2 + \ln \frac{2}{\delta}}{2m}}$$

		 
 $k = m$	0	large ($\sim > \frac{1}{2}$)
 $k < m$	> 0	small ($\sim < \frac{1}{2}$)

What did we learn from decision trees?

- Moral of the story:

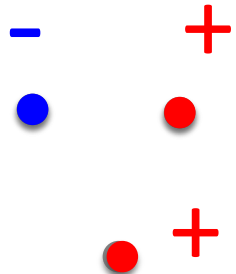
Complexity of learning not measured in terms of size of model space, but in maximum *number of points* that can be classified using a classifier from this model space

Rademacher Complexity

- Instead of all possible labelings, measure complexity by how accurately a model space can match a random labeling of the data.

For each data point i , draw random label

$$\sigma_i \quad \text{s.t.} \quad P(\sigma_i = +1) = \frac{1}{2} = P(\sigma_i = -1)$$



Then empirical Rademacher complexity of H is

$$\hat{R}_m(H) = \mathbb{E}_\sigma \left[\sup_{h \in H} \left(\frac{1}{m} \sum_{i=1}^m \sigma_i h(X_i) \right) \right]$$

Finite model class

- Rademacher complexity can be upper bounded in terms of model class size $|H|$:

$$\underbrace{\hat{R}_m(H)} \leq \underbrace{\sqrt{\frac{2 \ln |H|}{m}}}$$

- Often Rademacher bounds are significantly better, e.g. ...

Linear models with bounded norm

- Consider $h(X_i) = \langle w, X_i \rangle$ with fixed $\|w\|$, $\|X_i\| \leq R$

$$\hat{R}_m(H) = \mathbb{E}_\sigma \left[\sup_{h \in H} \left(\frac{1}{m} \sum_{i=1}^m \sigma_i h(X_i) \right) \right]$$

$\vdots$

$$\leq \frac{\|w\| R}{\sqrt{m}}$$

$w_1 \dots w_d$
Other norm bounds – HW5!

Complexity increases with number of parameters d and norm of weights

$$\hat{R}_m(H) = E_{\sigma} \left[\sup_{W: \|W\|_2 \leq B} \frac{1}{m} \sum_{i=1}^m \sigma_i W^T X_i \right] \quad h^{(X_i)} = W^T X$$

$$= E_{\sigma} \left[\sup_{W: \|W\|_2 \leq B} W^T \frac{1}{m} \sum_{i=1}^m \sigma_i X_i \right]$$

$$\leq \frac{B}{m} E_{\sigma} \left[\left\| \sum_{i=1}^m \sigma_i X_i \right\| \right]$$

Cauchy-Schwarz
 $W^T Z \leq \|W\|_2 \|Z\|_2$

$$\begin{array}{cc} \downarrow & \downarrow \\ \|W\|_P & \|W\|_q \\ P & q \\ \frac{1}{P} + \frac{1}{q} = 1 \end{array}$$

$$\leq \frac{B}{m} \sqrt{E_{\sigma} \left[\left\| \sum_{i=1}^m \sigma_i X_i \right\|^2 \right]}$$

Jensen's Inequality

$$E Z \leq \sqrt{E Z^2}$$

$$= \frac{B}{m} \sqrt{E_{\sigma} \left[\sum_{i=1}^m \sigma_i^2 \|X_i\|^2 + \sum_{i \neq j} \sigma_i \sigma_j X_i X_j \right]}$$

$$\|X_i\| \leq R$$

$$= \frac{B}{m} \sqrt{\sum_i \|X_i\|^2} \leq \frac{B}{m} \sqrt{m R^2} = \frac{B R}{\sqrt{m}}$$

$$K(x_i, x_i) = \phi(x_i)^T \phi(x_i) \longleftarrow x_i^T x_i \leq R^2$$

$$\uparrow$$

$$\|x_i\| \leq R$$

Kernel SVMs

- Consider $h(X_i) = \langle w, \underline{\phi(X_i)} \rangle$ with fixed $\|w\|, K_{x,x} \leq R^2$

$$\hat{R}_m(H) = \mathbb{E}_\sigma \left[\sup_{h \in H} \left(\frac{1}{m} \sum_{i=1}^m \sigma_i h(X_i) \right) \right]$$

$$\vdots$$

$$\leq \frac{\|w\| R}{\sqrt{m}}$$

Independent of dim d if
kernel is bounded!

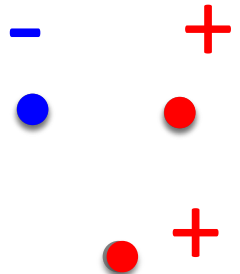
Complexity of kernel SVMs measured by $\|w\| \propto 1/\text{margin}$
whereas $\underline{|H|}, \underline{\text{VC}(\text{dim})} = \infty$ (for Gaussian kernels)

Rademacher Complexity

- Define for any function class $\mathcal{F}$ - empirical Rademacher complexity of $\mathcal{F}$ is

$$\sigma_i \sim \pm 1$$

$$\hat{R}_m(\mathcal{F}) = \mathbb{E}_{\sigma} \left[\sup_{f \in \mathcal{F}} \left(\frac{1}{m} \sum_{i=1}^m \sigma_i f(X_i) \right) \right]$$



Max correlation possible with random labels

→ Distribution and data dependent

- True Rademacher Complexity – distribution dependent

$$R_m(\mathcal{F}) = \mathbb{E}_X \left[\hat{R}_m(\mathcal{F}) \right]$$

Rademacher Concentration Bounds

- With probability $\geq 1-\delta$, bound deviation of true and empirical averages for any f in F using true and empirical Rademacher complexity of F

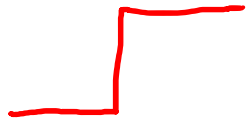
$$\sup_{f \in \mathcal{F}} \underbrace{\mathbb{E}_Z[f(Z)]} - \underbrace{\frac{1}{m} \sum_{i=1}^m f(Z_i)} \leq \underbrace{2R_m(\mathcal{F}) + \sqrt{\frac{\ln 1/\delta}{m}}}_{\substack{\text{BR} \\ \sqrt{m}}}$$

$$\sup_{f \in \mathcal{F}} \underbrace{\mathbb{E}_Z[f(Z)]} - \underbrace{\frac{1}{m} \sum_{i=1}^m f(Z_i)} \leq 2\hat{R}_m(\mathcal{F}) + 3\sqrt{\frac{2 \ln 1/\delta}{m}}$$

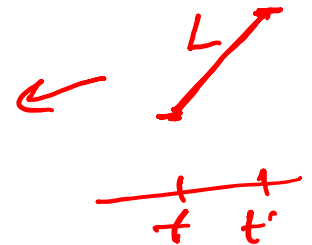
$\downarrow$
 $\text{loss}(h(Z_i)) = \text{loss} \circ h$

Contraction Property of Rademacher Complexity

- Talagrand Lemma/Lipschitz composition: If ψ is L-Lipschitz, i.e. for any t, t' in domain



$$|\psi(t) - \psi(t')| \leq L \|t - t'\|$$



then

$$\widehat{R}_m(\psi \circ H) \leq L \widehat{R}_m(H)$$

- Allows us to map between F (loss of model class) and H (model class) for Lipschitz losses

$$f(z) = \text{loss}(h(z))$$

$$\underline{f} = \text{loss} \circ \underline{h}$$

Other properties – HW5!

Rademacher Error Bounds

- For Lipschitz losses: With probability $\geq 1-\delta$,

$$\text{error}_{true}(h) \leq \text{error}_{train}(h) + \underbrace{2L\hat{R}_m(H)}_{2\hat{R}_m(\mathcal{F})} + 3\sqrt{\frac{\log(2/\delta)}{m}}$$

where empirical Rademacher complexity of H

$$\hat{R}_m(H) = \mathbb{E}_{\sigma} \left[\sup_{h \in H} \left(\frac{1}{m} \sum_{i=1}^m \sigma_i h(X_i) \right) \right]$$

is purely data-dependent.

Lipschitz losses

$$(1 - \langle w, x_i \rangle Y_i)_+$$

- Hinge loss

$$\begin{aligned} \text{loss}(\underline{h(Z_i)}) &= (1 - h(X_i)Y_i)_+ \\ &= \max(0, 1 - h(X_i)Y_i) \end{aligned}$$

→ $\psi(t) = \max(0, 1 - t)$ is 1-Lipschitz

$$\text{since } |\psi(t) - \psi(t')| \leq L \|t - t'\|$$

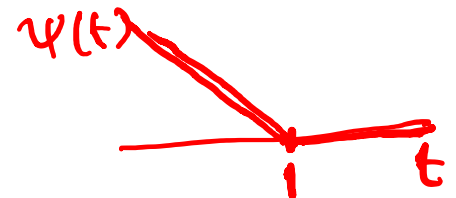
$$|\max(0, 1 - t) - \max(0, 1 - t')| \leq L \|t - t'\|$$

$L=1$

$$L=1.$$

$$\underline{\hat{R}_m(\mathcal{F})} \leq \underline{\hat{R}_m(H)}$$

$$\begin{aligned} \hat{R}_m(f) &\leq L \hat{R}_m(H) \\ f \in \mathcal{F} \quad f &= \text{loss}_{\text{hinge}} \circ h \end{aligned}$$



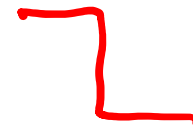
Binary 0/1 loss

- Binary (0/1) loss (in general, not Lipschitz over real domain but Lipschitz over its finite domain)

$$\text{loss}(h(Z_i)) = (1 - \underline{h(X_i)Y_i})/2$$

$$Z_i = (X_i, Y_i)$$

$$h(X_i) \in \{-1, 1\}$$



→ $\psi(t) = (1 - t)/2$ is $\frac{1}{2}$ -Lipschitz

$$\text{since } |\psi(t) - \psi(t')| \leq \frac{1}{2} \|t - t'\|$$

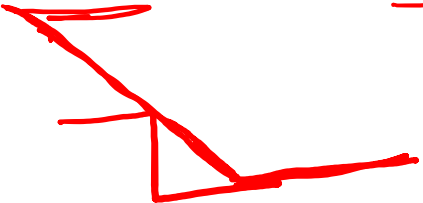
$$\underline{\hat{R}_m(\mathcal{F})} \leq \frac{1}{2} \underline{\hat{R}_m(H)}$$

Rademacher Error Bounds

- **Hinge Error:** With probability $\geq 1-\delta$,

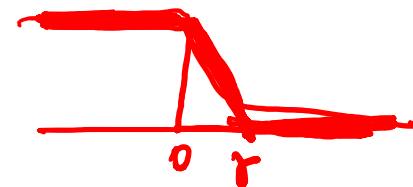
$$\text{error}_{\text{true}}^{\text{hinge}}(h) \leq \text{error}_{\text{train}}^{\text{hinge}}(h) + 2\hat{R}_m(H) + 3\sqrt{\frac{\log(2/\delta)}{m}}$$


- **0/1 Error:** With probability $\geq 1-\delta$,

$$\text{error}_{\text{true}}(h) \leq \text{error}_{\text{train}}(h) + \hat{R}_m(H) + 3\sqrt{\frac{\log(2/\delta)}{m}}$$


γ -margin loss

- γ -margin loss (upper bounds 0/1 loss)



$$\psi(t) = \begin{cases} 0 & t \geq \gamma \\ 1 & t \leq 0 \\ 1 - t/\gamma & \text{otherwise} \end{cases} \quad \text{is } 1/\gamma\text{-Lipschitz}$$

$$\text{since } |\psi(t) - \psi(t')| \leq \frac{1}{\gamma} \|t - t'\|$$

L

$$f = \text{loss}^r \circ h$$

$$\underline{\hat{R}_m(\mathcal{F})} \leq \frac{1}{\gamma} \underline{\hat{R}_m(H)}$$

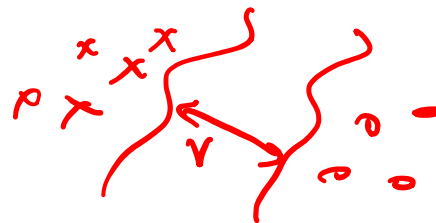
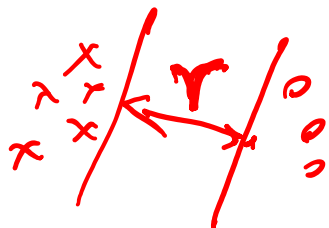
Kernel SVM Margin Bound

With probability $\geq 1-\delta$,

$$\begin{aligned} \text{error}_{\text{true}}^{0/1} &\leq \text{error}_{\text{true}}^{\gamma} \\ &\leq \text{error}_{\text{train}}^{\gamma} + \frac{1}{\gamma} \hat{R}_m(H) + \sqrt{\frac{\log(2/\delta)}{m}} \end{aligned}$$

$$\text{error}_{\text{true}}^{0/1}(h) \leq \underbrace{\text{error}_{\text{train}}^{\gamma}(h)}_{=0} + \underbrace{\frac{\|w\| R}{\gamma \sqrt{m}}}_{\text{small}} + 3 \sqrt{\frac{\log(2/\delta)}{m}}$$

- If $\text{error}_{\text{train}}^{\gamma}$ is small for large γ (f has large margin on data), then true 0/1 error is small.



Summary of PAC bounds

With probability $\geq 1-\delta$,

1) for all $h \in H$ s.t. $\text{error}_{\text{train}}(h) = 0$,

$$\text{error}_{\text{true}}(h) \leq \varepsilon = \frac{\ln |H| + \ln \frac{1}{\delta}}{m}$$

Finite
hypothesis
space

2) for all $h \in H$,

$$|\text{error}_{\text{true}}(h) - \text{error}_{\text{train}}(h)| \leq \varepsilon = \sqrt{\frac{\ln |H| + \ln \frac{2}{\delta}}{2m}}$$

3) For all $h \in H$,

Infinite hypothesis space

$$|\text{error}_{\text{true}}(h) - \text{error}_{\text{train}}(h)| \leq \varepsilon = \hat{R}_m(H) + 3\sqrt{\frac{\log(2/\delta)}{m}}$$