

today: finish Occam's bound
 • PAC Bayes
 • probit loss

Occam's bound

for any fixed P ; with prob. $1-\delta$ over training set $D_n \sim P^n$

$$\forall w \in W \quad L_P(w) \leq \hat{L}_P(w) + \frac{1}{\sqrt{2n}} \Omega_\pi(w; \delta)$$

$$\text{where } \Omega_\pi(w; \delta) \triangleq \sqrt{(\ln 2) \underbrace{|w|_\pi}_{\text{complexity measure}} + \ln \frac{1}{\delta}}$$

* bound is useful only for dist. P s.t. $|w|_\pi$ is small
 $\hookrightarrow \argmin_{w \in W} L_P(w)$

example of learning alg: $\hat{w}_n = \argmin_{w \in W}$ RHS of bound $\Rightarrow$ universally consistent!

$$|w|_\pi = \log_2 \int \pi(w) \quad \text{if } \pi(w) \propto \exp(-\|w\|^2) \quad \Rightarrow \text{RHS is } l_2 \text{-regularized ERM}$$

then $|w|_\pi = \|w\|^2 + \text{const.}$

proof: use 3 things:

1) Chernoff bound (concentration inequality)

$$P\{\bar{D}_n : \hat{L}_n(w) \leq L(w) - \epsilon\} \leq \exp(-2n\epsilon^2) \quad \forall \epsilon \geq 0$$

2) union bound

$$P\{\exists x \text{ s.t. prop}(x) \text{ is true}\} \leq \sum_x P\{\text{prop}(x) \text{ is true}\}$$

3) "Kraft's ineq."

$$\sum_w 2^{-|w|_\pi} \leq 1$$

note: 0-1 loss assumption appears in constant of Chernoff

we say that w is haughty if bound fails $L - \epsilon > \hat{L}_n$

$$\text{bad}(w) = \mathbb{1}\left\{L(w) > \hat{L}_n(w) + \frac{1}{\sqrt{2n}} \Omega_\pi(w; \delta)\right\}$$

using Chernoff, $\hat{L}_n(w) \leq L(w) - \epsilon_n(w)$ with small prob. ...

$$P\{\text{bad}(w)\} \leq \exp(-2n\epsilon_n(w)^2) = \exp\left(-2n \frac{1}{2n} (\ln 2) |w|_\pi + \ln \frac{1}{\delta}\right)$$

$\propto -|w|_\pi$

$$\sum_{i=1}^n \exp(-\eta \frac{1}{\delta} (w_i x_i + w_i \frac{1}{\delta})) = \exp(-\eta \frac{1}{\delta} \sum_{i=1}^n (w_i x_i + w_i \frac{1}{\delta}))$$

$$= \delta 2^{-\frac{1}{\delta} \sum_{i=1}^n w_i x_i}$$

using union bound:

$$P\{\exists w : \text{bad}(w)\} \leq \sum_w P\{\text{bad}(w)\} \leq \sum_w \delta 2^{-\frac{1}{\delta} \sum_{i=1}^n w_i x_i} \stackrel{\text{Kraft's inequality}}{\leq} \delta$$

Surrogate loss: NP hard to minimize $\hat{L}_n(w)$; replace with $\hat{J}_n(w)$ which is "surrogate"
e.g. hinge loss
• log loss

next: countable $\rightarrow$ uncountable
"PAC-Bayes"

PAC-Bayes

Occam's bound $\rightarrow$ use linked $\hat{L}_n(w)$ with $L_p(w)$

uniformly over all $w \in W$ but countable

using complexity $|w|_{\mathcal{H}}(\text{prior})$

PAC-Bayes: generalize this to

- arbitrary W
- general $\ell(y, y') \in [0, 1]$

convent: switch to a randomized predictor

i.e. instead of learning $\hat{w}$, predicting $y = h_{\hat{w}}(x)$

consider $\hat{q}$ distribution over W

predict: first $w \sim \hat{q}(w)$; $y = h_w(x)$

$\Rightarrow$ use $\mathbb{E}_{\hat{q}}[L(w)]$ as the "generalization error" for $\hat{q}$

i.e. $\mathbb{E}_{(x, y) \sim p} \mathbb{E}_{w \sim \hat{q}}[\ell(y, h_w(x))]$

idea: use empirical version

$\mathbb{E}_{\hat{q}}[\hat{L}_n(w)]$
 $\uparrow$
optimize over q

$\rightarrow$ on structured prediction
will yield probit surrogate loss. (see soon)

PAC-Bayes thm. [McAllester 1999, 2003]

PAC-Bayes Thm. [McAllester 1999, 2003]

(let $l(y, y') \in [0, 1]$) ; for any fixed prior π over W

and any dist. p over $X \times Y$

then with $\geq 1 - \delta$ over $D_n \sim p^{\otimes n}$

it holds that for any dist. over W

$$\mathbb{E}_q [L_p(w)] \leq \mathbb{E}_q [\hat{L}_n(w)] + \frac{1}{\sqrt{2nm}} \left(\text{KL}(q||\pi) + \ln \frac{n}{\delta} \right)$$

new complexity term

note: if W is countable; let $q_{w_0} = \mathbb{1}_{\{w=w_0\}}$

$$\text{then } \text{KL}(q||\pi) = \sum_w q(w) \log \frac{q(w)}{\pi(w)} = \log \frac{1}{\pi(w)} = (\log 2) |w_0|_{\pi}$$

15h33

probit loss for structured prediction [NeurIPS 2011 McAllester & Kesht]

if $q_w(w') \triangleq N(w' | w, I)$

$w' = w + \epsilon$ where $\epsilon \sim N(0, I)$

$$\text{then } \mathbb{E}_{q_w} [L(w')] = \mathbb{E}_{w' \sim q_w} \mathbb{E}_{(x,y) \sim p} [l(y, h_{w'}(x))]$$

$$= \mathbb{E}_{(x,y) \sim p} \left[\underbrace{\mathbb{E}_{\epsilon \sim N(0, I)} [l(y, h_{w+\epsilon}(x))]}_{\text{probit}(x, y; w)} \right]$$

why called probit?

binary classification: $y = \{-1, +1\}$ with 0-1 loss

$$h_w(x) = \text{sgn}(\langle w, \phi(x) \rangle)$$

$$\text{let } \alpha = \text{margin} = y \langle w, \phi(x) \rangle$$

$$\text{then } \text{probit}(x, y; w) = \mathbb{E}_{\epsilon \sim N(0, I)} \mathbb{1}_{\{y \neq h_{w+\epsilon}(x)\}} \\ \text{i.e. } y \langle w + \epsilon, \phi(x) \rangle < 0$$

$$\underbrace{y \langle w, \phi(x) \rangle}_{\alpha} < -y \langle \epsilon, \phi(x) \rangle$$

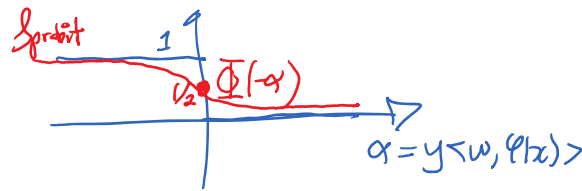
1 when

$$-\alpha > \langle \epsilon, \phi(x) \rangle$$

$$\underbrace{N(0, \|\phi(x)\|^2)}_{\text{assume } \|\phi(x)\| = 1}$$

$$\mathcal{J}_{\text{probit}} = \mathbb{P}\{\sum \varepsilon_i < -\alpha\} = \Phi(-\alpha)$$

↑ cdf of standard Gaussian



⊛ define
$$\hat{w}_n^{(\text{probit})} = \underset{w \in \mathcal{W}}{\text{argmin}} \left[\mathcal{J}_{\text{probit}}(w) + \frac{\lambda_n}{2n} \|w\|^2 \right] \quad (*)$$

McAllester showed the consistency of $\hat{w}_n^{(\text{probit})}$

McAllester 2011 uses Catoni's PAC-Bayes' version

$$\forall q, \mathbb{E}_q[L(w)] \leq \left(\frac{1}{1 - \frac{1}{2\lambda_n}} \right) \left[\mathbb{E}_q[\hat{L}_n(w)] + \frac{\lambda_n}{n} [KL(q||\pi) + \ln \frac{1}{\delta}] \right]$$

If we use $\pi = N(0, I)$
 $q_w = N(w, I)$

$$\left\{ \begin{array}{l} \hat{\mathcal{J}}_{\text{probit}}(w) + \frac{\lambda_n}{n} \frac{\|w\|^2}{2} \end{array} \right\}$$

Motivates $\hat{w}_n^{(\text{probit})}$ (*)

thm 1
in paper

Let $\lambda_n \nearrow \infty$ slowly enough so that $\frac{\lambda_n \ln n}{n} \rightarrow 0$

then $\mathcal{J}_{\text{probit}}(\hat{w}_n) \xrightarrow[n \rightarrow \infty]{\text{a.s.}} L^* = \min_{w \in \mathcal{W}} L_p(w)$

McAllester
call this
"consistency"

but true consistency would be $L(\hat{w}_n) \xrightarrow{\text{a.s.}} L^*$

[Lacoste-Julien unpublished result :
if $L(w)$ is cts.]

then $\mathcal{J}_{\text{probit}}(\hat{w}_n) \xrightarrow{\text{a.s.}} L^*$

$\Rightarrow L(\hat{w}_n) \xrightarrow{\text{a.s.}} L^* = L(w^*)$

proof idea: use Catoni's PAC Bayes bound

with prob $\geq 1 - \delta_n$

$$\mathcal{J}_{\text{probit}}(\hat{w}_n) \leq \left(\frac{1}{1 - \frac{1}{2\lambda_n}} \right) \left[\underbrace{\mathcal{J}_{\text{probit}}(\hat{w}_n) + \frac{\lambda_n}{2n} (\|\hat{w}_n\|)^2 + \ln \frac{1}{\delta_n}}_{\text{by def. of } \hat{w}_n} \right]$$

$$\leq \underbrace{\mathcal{J}_{\text{probit}}(w^*) + \frac{\lambda_n}{2n} \|w^*\|^2}_{\downarrow}$$

set $\delta_n = \frac{1}{n^2}$

$$\leq \mathbb{E} \text{probit}(\alpha w^*) + \sqrt{\frac{\ln 2}{n}} \quad \text{using Chernoff bound for } \alpha w^*$$

with prob $\geq 1 - \frac{1}{n^2}$

$$\text{get } \lim_{n \rightarrow \infty} \mathbb{E} \text{probit}(\hat{w}_n) \leq \mathbb{E} \text{probit}(\alpha w^*) \quad \text{with prob } 1$$

* also show that $\lim_{\alpha \rightarrow 0} \mathbb{E} \text{probit}(\alpha w^*) \leq L(w^*)$ [see paper for details]

problem: $\mathbb{E} \text{probit}(x, y; w)$ is non-convex in $w \Rightarrow$ no optimization guarantees

now: convex surrogates $\int$