

today: • PAC-Bayes
• probit loss
• review surrogate loss

PAC-Bayes:

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

uniformly over all $w \in W \rightarrow$ but (countable)
using complexity $|W|/\pi$ "prior"

PAC-Bayes $\circ$ generalize this to

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

current: 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. $\mathbb{E}_{(x,y) \sim p} \mathbb{E}_{w \sim \hat{q}}[l(y, h_w(x))]$
 $\rightarrow$ empirical version

$\mathbb{E}_{\hat{q}}[\hat{L}_n(w)] \rightarrow$ on structured prediction will yield probit surrogate loss (see soon)
 $\nearrow$ optimize over q

PAC-Bayes theorem [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 prob. $\geq 1 - \delta$ over $D_n \sim p^{\otimes n}$

it holds that $\forall q$ dist. over W $\mathbb{E}_q[L_p(w)] \leq \mathbb{E}_q[\hat{L}_n(w)]$

it holds that $\forall q$ dist. over $\mathcal{W}$ $\mathbb{E}_q[L_p(w)] \leq \mathbb{E}_q[L_n(w)] + \frac{1}{\sqrt{2(n-1)}} \sqrt{KL(q||\pi) + \ln n}$

note: if $\mathcal{W}$ is countable; let $q_{w_0} = \mathbb{1}\{w=w_0\}$ new complexity term

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

probit loss for structured prediction [NIPS 2011 McAllester & Keshet]

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

then $\mathbb{E}_{q_w}[L(w')]$ $\xrightarrow{w'=w+\epsilon}$ $\mathbb{E}_{(x,y) \sim p}[l(y, h_{w+\epsilon}(x))]$ where $\epsilon \sim N(0, I)$

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

why name probit?

binary classi. : $\mathcal{Y} = \{-1, +1\}$ with 0-1 loss

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

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

then $\text{probit}(x, y; w) = \mathbb{E}_{\epsilon \sim N(0, I)} \mathbb{1}\{y \neq h_{w+\epsilon}(x)\}$

$y \langle w+\epsilon, \phi(x) \rangle < 0$

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

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

same prob. as

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

$\hookrightarrow \mathbb{P}\{\epsilon_1 < -\alpha\}$

(suppose $1/\phi(x) = 1$)

$\text{probit} = \mathbb{P}\{\epsilon_1 < -\alpha\} = \Phi(-\alpha)$

$\uparrow$ cdf of standard Gaussian



⊗ define $\hat{w}_n^{(\text{probit})} = \arg \min_{w \in \mathcal{W}} \hat{\text{probit}}(w) + \frac{\lambda_n}{2n} \|w\|^2$

⊗ define $\hat{w}_n^{(probit)} = \arg \min_{w \in W} \left[\hat{S}_{probit}(w) + \frac{\lambda_n}{2n} \|w\|^2 \right] \quad (*)$

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

14h23

McAllister 2011 uses Caloni's PAC-Bayes version

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

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

$\hat{S}_{probit}(w)$ $\frac{\lambda_n}{n} \|w\|^2$

motivates $\hat{w}_n^{(probit)}$

thm 9

in paper: let $\lambda_n \rightarrow 0$ slowly enough so that $\frac{\lambda_n}{n \ln n} \rightarrow 0$

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

McAllister calls this 'consistency'

but true consistency would be $L(\hat{w}_n) \xrightarrow[n \rightarrow \infty]{a.s.} L^*$

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

then $\hat{S}_{probit}(\hat{w}_n) \xrightarrow[n \rightarrow \infty]{a.s.} L^*$

$\Rightarrow L(\hat{w}_n) \xrightarrow[n \rightarrow \infty]{a.s.} L^* = L(w^*)$

proof idea: use Caloni's PAC Bayes bound

with prob. $\geq 1 - \delta_n$

by def. of $\hat{w}_n$

$$\hat{S}_{probit}(\hat{w}_n) \leq \left(\frac{1}{1-\frac{1}{2\lambda_n}} \right) \left(\underbrace{\hat{S}_{probit}(\hat{w}_n) + \frac{\lambda_n}{2n} \|\hat{w}_n\|^2}_{\text{pkc } \alpha} + \ln \frac{1}{\delta_n} \right)$$

$$\leq \underbrace{\hat{S}_{probit}(\alpha w^*) + \frac{\lambda_n}{2n} \alpha^2 \|w^*\|^2}_{2n}$$

$\leq \hat{S}_{probit}(\alpha w^*) + \sqrt{\frac{\ln n}{n}}$ using Chernoff bound for αw^*

* also use $\lim_{\alpha \rightarrow \infty} \hat{S}_{probit}(\alpha w^*) \leq L(w^*)$

$$\lim_{w \rightarrow \infty} \mathcal{L}(w) = L(w^*) \quad [\text{see paper for details}]$$

problem: $\mathcal{L}(x, y; w)$ is non-convex in $w \Rightarrow$ no optimization guarantees

how: convex surrogates $s(\tilde{y}) \triangleq s(x, \tilde{y}; w)$ i.e. $x \neq w$ is implicit

Review of convex surrogates mentioned so far:

$$\mathcal{L}_{\text{perceptron}}(x, y; w) = \max_{\tilde{y} \in \mathcal{Y}} s(\tilde{y}) - s(y)$$

$$\text{let } m(\tilde{y}) = s(y) - s(\tilde{y})$$

$$= \max_{\tilde{y} \in \mathcal{Y}} [-m(\tilde{y})] = \left[\max_{\tilde{y} \in \mathcal{Y}} -m(\tilde{y}) \right]_+$$

$$\mathcal{L}_{\text{hinge}}(\text{structured sum}) = \max_{\tilde{y}} [s(\tilde{y}) + l(y, \tilde{y})] - s(y)$$

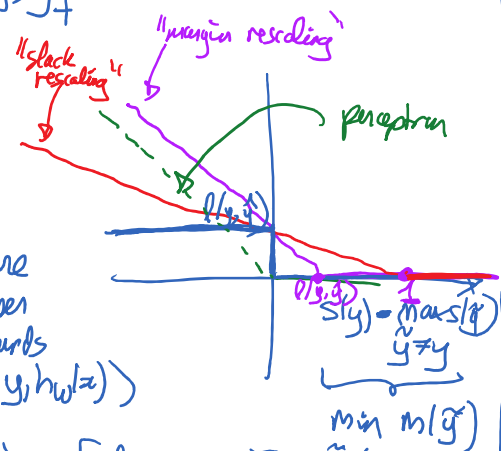
"margin rescaling"

$$= \max_{\tilde{y}} [l(y, \tilde{y}) - m(\tilde{y})]$$

vs.
"slack rescaling"

$$= \max_{\tilde{y}} l(y, \tilde{y}) [1 - m(\tilde{y})]$$

are upper bounds $l(y, h(w|x))$



$$\mathcal{L}_{\text{CRF}}(\) = \frac{1}{\beta} \log \left(\sum_{\tilde{y}} \exp(\beta s(\tilde{y})) \right) - s(y) \quad [-\log p_{\theta}(y|x)]$$

$\beta \rightarrow \infty \Rightarrow$ perceptron loss

$$\frac{1}{\beta} \log \left(\sum_{\tilde{y}} \exp(-\beta m(\tilde{y})) \right)$$

suggests "smashed hinge loss"

$$\frac{1}{\beta} \log \left(\sum_{\tilde{y}} \exp(\beta [l(y, \tilde{y}) - m(\tilde{y})]) \right)$$

[e.g. Platts et al. 2010]

note: slack rescaling more robust when have small $l(y, \tilde{y})$ [e.g. 0]

but more computationally costly

what theoretical properties could we look at?

a) generalization error bounds [next class]

b) consistency properties & calibration facts [next² class]
 $\hookrightarrow$ relationship between $L(w)$ & $\mathcal{L}(w)$

why structural score functions?

$$s(x, y) = \sum_{c \in \mathcal{C}} S_c(x, y_c)$$

Motivation similar to graphical models

- 1) statistical efficiency: less # of parameters (simpler score fct. S)
 $\Rightarrow$ easier to learn
(generalization guarantees) [see Cortis & al. NIPS 2006 next class]
- 2) computational " : compute $\arg\max_{\tilde{y} \in \mathcal{Y}} s(\tilde{y})$