

## Lecture 5 - scribbles - PAC-Bayes

Tuesday, January 22, 2019 14:33

today: PAC-Bayes  
probit loss  
review surrogate losses

### PAC-Bayes

Occam's bound  $\rightarrow$  we linked  $\hat{L}_n(w)$  with  $L_P(w)$   
uniformly over all  $w \in W$  (countable)  
using complexity  $\ln \pi$   
 $q$  "prior"

PAC-Bayes: generalizes this to:  
• arbitrary  $W$   
• general  $\ell(y, y') \in [0, 1]$

caveat: switch to a randomized predictor

ie. instead of  $\hat{w}$   $y = h_{\hat{w}}(x)$   
consider  $q$  distributions over  $W$

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

use  $\mathbb{E} \hat{q}[L(w)]$  as the generalization error for  $\hat{q}$  ie.  $\mathbb{E}_{(x,y) \sim P} \mathbb{E}_{w \sim \hat{q}} \ell(y, h_w(x))$   
 $\downarrow$   
empirical version

$\mathbb{E} \hat{L}_n(w)$  } structured prediction,  
this will yield probit surrogate loss (see soon)

PAC-Bayes Thm. [McAllester 1999, 2003]

(let  $(y, y') \in [0, 1]$ ) for any fixed prior  $\pi$  over  $\mathcal{W}$

and any dist.  $p$  on  $\mathcal{X} \times \mathcal{Y}$

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

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

new complexity term

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

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

probit loss for structured prediction: [WIPS 2011 McAllester & Keshet]

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

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

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

$$= \mathbb{E}_{(x,y) \sim P} \left[ \mathbb{E}_{\epsilon \sim N(0,1)} \ell(y, h_{w+\epsilon}(x)) \right]$$

$\triangleq$

$$\mathcal{J}_{\text{probit}}(x, y; w)$$

why name probit?

binary classification  $\mathcal{Y} = \{-1, +1\}$   
with  $C=1$  loss

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

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

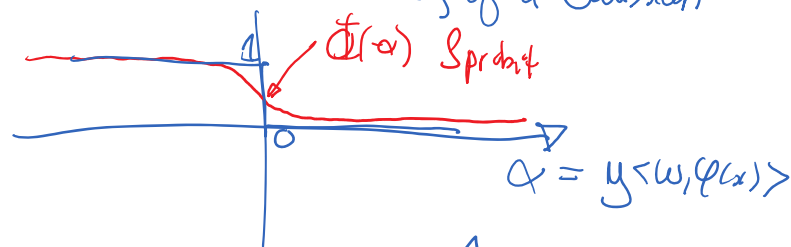
$$\text{then } \mathcal{J}_{\text{probit}}(x, y; w) = \mathbb{E}_{\epsilon \sim N(0,1)} \mathbb{1}_{\{y \neq h_{w+\epsilon}(x)\}}$$

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

$$(\text{supposing } \|\phi(x)\|=1) \Rightarrow \alpha < -y \langle \epsilon, \phi(x) \rangle$$

$$= P\{\epsilon_1 \geq \alpha\} = \Phi(-\alpha)$$

$\Phi$  cdf of a Gaussian



McClellan showed the consistency of the  $\hat{w}_{\text{probit}}$  which minimizes a  $\ell_2$ -reg,  $\hat{\mathcal{J}}_{\text{probit}}$

McAllester 2011 uses Catoni's PAC-Bayes version:

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

if we use  $\pi = N(0, I)$   
 $q_w = N(w, I)$  }  $\Rightarrow \hat{S}_{\text{probit}}(w)$   $\frac{1}{2} \|w\|^2$

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

Thm 1

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

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

McAllester calls this 'consistency'

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

(Lacoste-Julien unpublished fix: if  $L(w)$  is cts.

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

proof idea: use Catoni's PAC-Bayes bound

with prob  $\geq 1 - \delta_n$

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

$$\leq \hat{S}_{\text{probit}}(\alpha w) + \lambda_n \alpha^2 \|w\|^2$$

$$\leq \underbrace{\hat{J}_{\text{probit}}(\alpha w^*)}_{\leq J_{\text{probit}}(\alpha w^*)} + \frac{\ln \alpha^2 / \|w^*\|^2}{\partial n} \\ \hookrightarrow \leq J_{\text{probit}}(\alpha w^*) + \sqrt{\frac{\ln n}{n}} \quad \text{using Chernoff bound for } \alpha w^*$$

use  $\lim_{\alpha \rightarrow \infty} J_{\text{probit}}(\alpha w^*) \leq L(w^*)$

o.o

$$\lim_{n \rightarrow \infty} J_{\text{probit}}(\hat{w}_n) = L(w^*) \quad [\text{see paper for details}] //$$

15h3)

problem:  $J_{\text{probit}}(x, y; w)$  is non-convex  $\Rightarrow$  no optimization guarantee

now: convex surrogates  $s(\tilde{y}) \triangleq s(x, \tilde{y}; w)$  i.e.  $x$  &  $w$  are implicit

Review of convex surrogates mentioned so far:

$J_{\text{perceptron}}(x, y; w) = \max_{\tilde{y} \in \mathcal{Y}} s(\tilde{y}) - s(y)$   
 $= \max_{\tilde{y} \in \mathcal{Y}} [-m(\tilde{y})] = [\max_{\tilde{y} \in \mathcal{Y}} m(\tilde{y})]_+$

$J_{\text{hinge}}(\text{---}) = \max_{\tilde{y} \in \mathcal{Y}} [s(\tilde{y}) + \ell(y, \tilde{y})] - s(y)$   
 (structured sum)

"margin rescaling"  $\hookrightarrow = \max_{\tilde{y}} [\ell(y, \tilde{y}) - m(\tilde{y})]$   
 "slack rescaling"  $= \max_{\tilde{y}} \ell(y, \tilde{y}) [1 - m(\tilde{y})]$

$J_{\text{SCE}}(\text{---}) = \frac{1}{\beta} \log(\sum \exp(\beta s(\tilde{u}))) - s(u) \quad \left[ -\log P_w(y|x) \right]$

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

are both upper bounds on  $\ell(y, h_w(x))$

$$\begin{aligned}
 \mathcal{L}_{\text{CRF}}(\text{---}) &= \frac{1}{\beta} \log \left( \sum_{\tilde{y}} \exp(\beta s(\tilde{y})) \right) - s(y) \quad [-\log p(y|x)] \quad \min_{\tilde{y}} m(\tilde{y}) \\
 (\text{log-loss for CRF}) & \\
 &\quad \beta \rightarrow \infty \Rightarrow \text{perceptron loss} \\
 &\quad \frac{1}{\beta} \log \left( \sum_{\tilde{y}} \exp(\beta m(\tilde{y})) \right) \xrightarrow{\text{"smoothed hinge loss"}} \frac{1}{\beta} \log \left( \sum_{\tilde{y}} \exp(\beta [0(y, \tilde{y}) - m(\tilde{y})]) \right) \\
 &\quad \text{[e.g. Pletschen \& al. 2010]}
 \end{aligned}$$

note: slow rescaling more robust when have small  $0(y, \tilde{y})$  [e.g. 0]  
 but more computationally costly

what theoretical properties could we look at?

- a) generalization error bounds [today]
- b) consistency properties & calibration [next class]

why structured score functions?  $s(x, y) = \sum_{c \in \mathcal{C}} s_c(x, y_c)$   
 motivations similar to graphical models

- 1) statistical efficiency: less # of parameters (simpler score functions  $s_c$ )  
 $\Rightarrow$  easier to learn (generalization guarantees) [see Cortes \& al. NIPS 2016] (beg. of next class)
- 2) computational " : compute  $\arg \max_{\tilde{y} \in \mathcal{Y}} s(\tilde{y})$

BUT compare to what happens for Hamming loss:

BUT compare to what happens for Hamming loss:

given true conditional  $q_x(y) \triangleq p(y|x)$  generating data  $y = (y_1, \dots, y_p, \dots, y_L)$   
 expected error when using  $\tilde{y}$  as prediction is  $\mathbb{E}_{y \sim q_x(y)} [l(y, \tilde{y})] \triangleq l_{q_x}(\tilde{y})$  <sup>part</sup>

for Hamming loss:  $l_{q_x}(\tilde{y}) = \mathbb{E}_{q_x(y)} \left[ \sum_p \frac{\mathbb{1}\{y_p \neq \tilde{y}_p\}}{1 - \mathbb{1}\{y_p = \tilde{y}_p\}} \right] = \sum_p (1 - q_x(\tilde{y}_p))$   
 $\Rightarrow$  best decision  $y^* = \underset{\tilde{y} \in \mathcal{Y}}{\operatorname{argmin}} l_x(\tilde{y})$  marginal on  $y_p$   
i.e.  $\sum_{y: y_p = \tilde{y}_p} q_x(y)$

is just independent predictions  $\tilde{y}_p = \underset{\tilde{y}_p}{\operatorname{argmax}} p(\tilde{y}_p|x)$

max marginals decoding

⊗ thus; if there is no constraints, then can just train independently models for each part marginal  $p(y_p|x)$   
 i.e.  $S_p(y_p|x; w_p)$

but a) this function might be too complicated

b) statistically, could be beneficial to share "transfer learning" between parts  
 e.g.  $p(y_p|x) \propto \exp(S_p(y_p; w))$