

today: consistency for convex surrogate losses

non-parametric viewpoint on scores

$$s(x, y, w) = \langle w, \phi(x, y) \rangle \quad \text{if } w = \sum_{i, y} \alpha_i(y) \phi(x_i, y)$$

$$\Rightarrow \langle w, \phi(x, y) \rangle = \sum_{i, y} \alpha_i(y) \underbrace{\langle \phi(x_i, y), \phi(x, y) \rangle}_{K(x_i, x; y, y)}$$

$$\text{often for simplicity: } K(x, x'; y, y') = K_x(x, x') K_y(y, y')$$

$$[\text{is equivalent to have } \phi(x, y) \triangleq \phi_x(x) \otimes \phi_y(y)]$$

$\otimes$ "product kernel"
Kronecker product

$$V \otimes w \quad V w^T$$

$$\langle V \otimes w, V' \otimes w' \rangle = \text{tr}((w w^T)^T (V' w^T))$$

$$= \text{tr}(w \underbrace{V'^T V}_{\langle V', V \rangle} w^T) = \langle V, V' \rangle \underbrace{\text{tr}(w w^T)}_{\langle w, w \rangle} = \langle V, V' \rangle \langle w, w \rangle //$$

$$\text{e.g. } K_x(x, x') = \exp\left(-\frac{\|x - x'\|^2}{2\sigma^2}\right) \quad \begin{matrix} \text{RBF kernel} \\ \text{(universal)} \end{matrix}$$

$$\phi_y: \mathcal{Y} \rightarrow \mathbb{R}^d \quad d \leq |\mathcal{Y}| \triangleq k \quad K_y(y, y') = \langle \phi_y(y), \phi_y(y') \rangle$$

$\uparrow$
 10^{23}

back to consistency & surrogate losses

$$\hat{w}_n \triangleq \arg\min_w \hat{L}_n(w) + \lambda_n \frac{\|w\|^2}{2}$$

$$\text{consistency: } L(\hat{w}_n) \xrightarrow{n \rightarrow \infty} \min_w L(w)$$

⊛ binary classification [Bartlett & al. 2004] characterized a whole family
of consistent (convex) surrogate losses
↳ binary SVM
logistic regression

for multiclass classification [Lee & al. 2004] showed that multiclass hinge loss

McAllister 2007

$$J_{\text{hinge}}(x, y; w) = \max_{\tilde{y}} (s(\tilde{y}) + \ell(y, \tilde{y})) - s(y)$$

is not consistent for 0-1 loss
when have no "majority" class

$$(i.e. p(\tilde{y}|x) < \frac{1}{2} \forall \tilde{y})$$

↑ true p

→ they propose a different surrogate loss that we $\sum_{\tilde{y}}$ instead of $\max_{\tilde{y}}$

which is consistent
for 0-1 loss

↓
exponential sum
→ could be intractable for
structured prediction

2 aspects of structured prediction
which give a richer theory than binary class. for consistency

1) "noise model" $p(y|x)$ is much richer

2) $\ell(y, y')$ much richer

⊛ [Ozaki & al. 2017] → we looked at effect of $\ell(y, y')$

for a easy to analyze convex surrogate loss is consistent
in the simplest possible setting

and we were careful about exponential constants (e.g. $|Y|=k$)

15h18

calibration function for a structured loss ℓ , surrogate loss $\mathcal{L}$ and set W

$$H_{\mathcal{L}, \ell, W}(\epsilon) \triangleq \inf_{\substack{w \in W \\ q \in \Delta_{|Y|}}} [\mathcal{L}_q(w) - \min_{\substack{w' \in W \\ \mathcal{L}_q^*}} \mathcal{L}_q(w')]$$

$$s.t. \mathcal{L}_q(w) - \min_{w' \in W} \mathcal{L}_q(w') \geq \epsilon$$

x is fixed
outside
 q is a potential
 $p(y|x)$

$$\mathcal{L}_q(w) \triangleq \mathbb{E}_{q(\tilde{y})} [\mathcal{L}(x, \tilde{y}; w)]$$

$$\mathcal{L}_q(w) \triangleq \mathbb{E}_{q(\tilde{y})} [\ell(\tilde{y}, h_w(x))]$$

"conditional (logit) risk"
[conditional on x version]

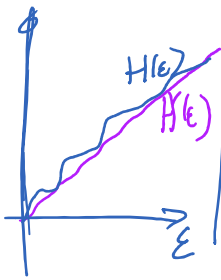
smallest "surrogate optimization regret"
(over all $\text{dist. } q$) s.t. true regret $\geq \epsilon$

$$i.e. \forall q: \mathcal{L}_q(w) < \mathcal{L}_q^* + H(\epsilon) \Rightarrow \mathcal{L}_q(w) \leq \mathcal{L}_q^* + \epsilon$$

L continuous on x reason

$$i.e. \{ \forall q \in \mathcal{Q}(w) : \forall q \in \mathcal{Q}(w) : L(q) \leq L^* + \epsilon \}$$

$$\Rightarrow L(q) \leq L^* + \epsilon$$



$$(thm. 2) \forall p : g(w) \leq g^* + H(\epsilon)$$

$$\mathbb{E}_{(x,y) \sim p} [L(g(w; x, y))] \Rightarrow L(w) \leq L^* + \epsilon$$

$$\Rightarrow L(w) \leq L^* + \epsilon$$

convex lower envelope of $H(\epsilon)$

(↑ basically shown using Jensen's ineq.)

$$\check{H}(\epsilon) \triangleq H^{**}(\epsilon)$$

$$f^*(z) \triangleq \sup_x x^T z - f(x) \rightarrow \text{Fenchel-legendre conjugate}$$

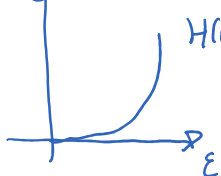
If $\check{H}$ is invertible

$$L(w) - L^* \leq \check{H}^{-1}(g(w) - g^*)$$

$\mathcal{L}$ is consistent iff $H(\epsilon) > 0 \forall \epsilon > 0$

(and $H(\epsilon)$ is finite for some $\epsilon > 0$)

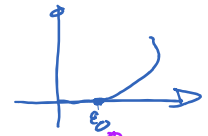
standard H is



$$H(\epsilon) = \frac{\epsilon^2}{C} \quad H^{-1}(z) = \sqrt{Cz}$$

$$L(w) - L^* \leq \sqrt{C(g(w) - g^*)}$$

you want small C ; for structured prediction $C = |\mathcal{S}|$ often? (bad)



min regret I can guarantee

note: scale of H is arbitrary

normalize it using stochastic optimization perspective (e.g. SGD) [next class]

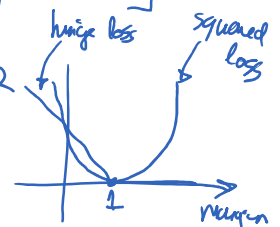
concrete example: simplest surrogate loss: square loss?

$$s(\cdot) \in \mathbb{R}^K \quad (s(x))$$

$$\mathcal{L}(x, y, s) \triangleq \frac{1}{2K} \|s - (-\ell(y; \cdot))\|_2^2 = \frac{1}{2K} \sum_{\tilde{y}} (s(x, \tilde{y}) + \ell(y, \tilde{y}))^2$$

[can be seen as a generalization of squared loss for binary class. to multiclass]

$$[1 - y_i \cdot w_i \cdot \ell(x, i)]^2$$



$$\mathcal{L}_Q(s) \triangleq \mathbb{E}_{Q(y)} \mathcal{L}(x, y, s)$$

$$= \frac{1}{2K} \sum_{\tilde{y}} \mathbb{E}_{Q(y)} [s(\tilde{y})^2 + 2s(\tilde{y})\ell(y, \tilde{y}) + \ell(y, \tilde{y})^2] \quad (\text{const.})$$

does not depend on s

$$\begin{aligned} \mathbb{E}_{q(y)} l(y, \tilde{y}) &\triangleq l_{q_x}(\tilde{y}) \\ &= \frac{1}{2k} \|s + l_{q_x}\|_2^2 + \text{cst.} \end{aligned}$$

Suppose s is unconstrained $\min_s l_{q_x}(s) \Rightarrow s^*(\tilde{y}) = -l_{q_x}(\tilde{y})$

$$\arg\max_{\tilde{y}} s^*(\tilde{y}) = \arg\min_{\tilde{y}} l_{q_x}(\tilde{y})$$

ie. you predict optimally pointwise on $\tilde{x}$

so here l is consistent ie. $s^* \in \arg\min_{s: X \rightarrow \mathbb{R}^k} l(s)$

$$l_q(s) = \min_{s' \in \mathbb{R}^k} l_q(s') = \min_{s'} \frac{1}{2k} \|s - (-l_{q_x})\|_2^2 \Rightarrow L(h_{s^*}) = \min_{\text{all } h}$$

let $\tilde{L}$ be a $k \times k$ matrix where $\tilde{L}_{\tilde{y}, y} = l(y, \tilde{y})$ $l_{q_x} = \sum_y q(y|x) l(y, \cdot)$

$$l_{q_x} = \tilde{L} q_x$$

recall: $s^* = -l_{q_x} = -\tilde{L} q_x \in \text{span}(\tilde{L})$ ie. $\sum_y \alpha_y \tilde{L}(y, \cdot)$

⊛ to get consistency for l , it is sufficient to consider $s \in \text{span}(\tilde{L})$

or that $s \in \text{span}(F) \supseteq \text{span}(\tilde{L})$

restriction on errors

$F \in \mathbb{R}^{k \times r}$ matrix
can be chosen clearly depending on $\tilde{L}$

$$S = F\theta \quad \theta \in \mathbb{R}^r$$

$$l_q(\theta) = \min_{\theta \in \mathbb{R}^r} l_q(\theta) = \frac{1}{2k} \|F\theta - (-\tilde{L}q)\|_2^2$$

Thm. 7

if $\text{span}(F) \supseteq \text{span}(\tilde{L})$

$$H_{\text{square}, l, F}(\epsilon) \geq \frac{\epsilon^2}{2k \max_{i \neq j} \|P_F \Delta_{ij}\|_2^2} \geq \frac{\epsilon^2}{4k}$$

lower bound $\Rightarrow$ robustness result

this is bad

$$\Delta_{ij} \triangleq e_i - e_j \in \mathbb{R}^k$$

P_F is orthogonal projection on $\text{span}(F)$ $P_F = F(F^T F)^+ F^T$

• in paper, we show that for 0-1 loss, $H(\epsilon) = \frac{\epsilon^2}{4K}$

thm. 8: if $\text{span}(F) = \mathbb{R}^K$ (ie no constraints) hardness result
 then $H(\epsilon) \leq \frac{\epsilon^2}{4K}$ for any loss?

ie. for any loss, we need an exp. # of samples (in worst case)
 to learn well [concat $\rightarrow$ all these bounds] one worst case

⊛ but for Hamming loss, if add constraints that $s(\tilde{y}) = \sum_{p \in \text{params}} s_p(\tilde{y}_p)$

over T binary variables, $H(\epsilon) = \frac{\epsilon^2}{8T}$ } not too big $\Rightarrow$ we can learn?

note: computation how to compute $\sum_{\tilde{y}} \ell(y, \tilde{y}) s(\tilde{y})$

$\rightarrow$ efficient to compute for
 Hamming loss & separable score fct, eg