

today: consistency for convex surrogate losses

non-parametric viewpoint on scores

$$s(x, y; w) = \langle w, \phi(x, y) \rangle$$

$$k(x_i, \cdot; \tilde{y}, \cdot) \downarrow$$

$$w = \sum_{i, \tilde{y}} \alpha_i(\tilde{y}) \phi(x_i, \tilde{y})$$

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

often for simplicity: $k(x, x'; y, y') = k_x(x, x') k_y(y, y')$

"product kernel"

[is equivalent to having $\phi(x, y) \triangleq \phi_x(x) \otimes \phi_y(y)$]
 $\uparrow$
 Kronecker product

$$v \otimes w \quad v w^T$$

$$\begin{aligned} \langle v \otimes w, v' \otimes w' \rangle &= \text{tr}((v w^T)^T (v' w'^T)) \\ &= \text{tr}(w \underbrace{v^T v'}_{\langle v, v' \rangle} w'^T) \\ &= \langle w, w' \rangle \langle v, v' \rangle \end{aligned}$$

e.g. $k_x(x, x') = \exp(-\frac{\|x - x'\|^2}{2\sigma^2})$

RBF kernel (unwired)

$$\varphi_y: \mathcal{X} \rightarrow \mathbb{R}^d \quad d \ll k = |\mathcal{Y}| \quad k_y(y, y') = \langle \varphi(y), \varphi(y') \rangle$$

Back to consistency & surrogate losses

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

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

⊛ binary classification [Bartlett & al. 2004] characterized a whole family of consistent surrogate losses

↳ binary SVM
logistic regression

for multiclass classification [Lee & al. 2004] showed that multiclass
McAllester 2002 huge loss

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

is not consistent for 0-1 loss
when have no majority class (i.e. $p(y|x) < \frac{1}{|\mathcal{Y}|} \forall y$)

→ they propose a different surrogate loss that $\sum_{\tilde{y}} \dots$ instead of $\max_{\tilde{y}}$
which is consistent for 0-1 loss

exponential sum
→ could be intractable

for 0-1 loss

exponential sum
→ could be intractable

2 aspects of structured prediction which give a much richer theory than binary classification for consistency

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

2) $l(y, y')$ much richer

⊛ [Osaka in 4 ed. 2017] → we looked at effect of $p(y, y')$
for a easy to analyze convex surrogate loss
in the simplest possible setting
and we were careful about exponential constants

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

$$H_{\mathcal{L}, l, \mathcal{W}}(\varepsilon) \triangleq \inf_{\substack{w \in \mathcal{W} \\ q \in \Delta_{\mathcal{Y}}}} [\mathcal{L}_q(w) - \min_{w' \in \mathcal{W}} \mathcal{L}_q(w')] \quad \text{s.t.} \quad \mathcal{L}_q(w) - \min_{w' \in \mathcal{W}} \mathcal{L}_q(w') \geq \varepsilon$$

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

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

$$\mathcal{L}_q(w) \triangleq \mathbb{E}_{q(\tilde{y})} [l(\tilde{y}, h_w(x))] \quad \text{"conditional risk"}$$

→ smallest optimization surrogate regret (over all dist. q)
s.t. true regret $\geq \varepsilon$

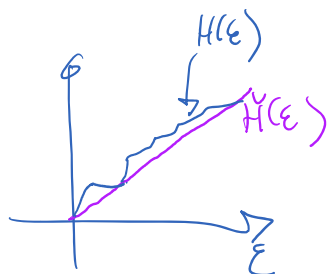
(conditional on x
version)

$$\forall q: g_q(w) \leq g_q^* + H(\epsilon) \Rightarrow L_q(w) \leq L_q^* + \epsilon$$

$$\boxed{\text{(thm. 2)} \quad \forall p: g(w) \leq g^* + \check{H}(\epsilon) \Rightarrow L(w) \leq L^* + \epsilon}$$

$\downarrow \mathbb{E}_{x,y} g(x,y,w)$ $\uparrow$

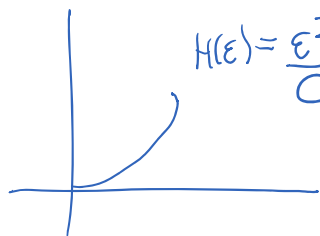
convex lower envelope of $H(\epsilon)$



$$\check{H}(\epsilon) \triangleq H^{**}(\epsilon) \quad f^*(z) \triangleq \sup_x x^T z - f(x) \quad \text{"Fenchel-Legendre" conjugate}$$

f is consistent iff $H(\epsilon) > 0 \quad \forall \epsilon > 0$
(and $H(\epsilon)$ is finite for some $\epsilon > 0$)

$$\boxed{\text{if } \check{H} \text{ is invertible} \\ L(w) - L^* \leq \check{H}^{-1}(g(w) - g^*)}$$



$$H(\epsilon) = \frac{\epsilon^2}{C} \Rightarrow L(w) - L^* \leq \sqrt{C(g(w) - g^*)}$$

you want small C ; for structural prediction $C = 137$ often?

14h24

note: scale of H is arbitrary ;
normalize it using optimization perspective (o.g. SGD)

[see soon]

simplest surrogate loss: square loss?

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

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

[can be seen as generalization of squared loss for binary classification to multi class
 $(1 - y_i + w_i \phi(x))^2$]

$$\mathcal{J}_q(s) \triangleq \mathbb{E}_{q(y)} \mathcal{J}(x, y; s)$$

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

does not depend on s

$$\mathbb{E}_{q(y)} \ell(y, \tilde{y}) \triangleq \ell_{q_x}(\tilde{y})$$

suppose s is unconstrained

$$\min_s \mathcal{J}_{q_x}(s)$$

$$s(\tilde{y}) + \ell_{q_x}(\tilde{y}) = 0 \quad \forall \tilde{y}$$

$$\Rightarrow s^*(\tilde{y}) = -\ell_{q_x}(\tilde{y})$$

$$\arg\max_{\tilde{y}} s^*(\tilde{y}) = \arg\min_{\tilde{y}} \ell_{q_x}(\tilde{y}) \quad \text{ie. you predict optimally (pt. wise)}$$

so here $\mathcal{J}$ is consistent

$$\text{ie. } s^* \in \arg\min_{\text{all } s} \mathcal{J}(s) \Rightarrow L(h_{s^*}) = \min_{\text{all } h} L(h)$$

$$h_{\tilde{y}}^*(x) = \arg \max_{\tilde{y}} \tilde{s}(x, \tilde{y})$$

$$\mathcal{J}_q(s) = \|s - (-l_q)\|^2 + \underline{cst}$$

$$\mathcal{J}_q(s) - \min_{s \in \mathbb{R}^k} \mathcal{J}_q(s) = \frac{1}{2k} \|s - (-l_q)\|_2^2$$

$$l_{q,x} = \sum_y q(y|x) l(y, \cdot)$$

$$\text{Def } \overset{\leftrightarrow}{L} \text{ is a } k \times k \text{ matrix where } \overset{\leftrightarrow}{L}_{\tilde{y}, y} = l(y, \tilde{y})$$

$$l_{q,x} = \overset{\leftrightarrow}{L} q_x$$

$$s^* = -l_{q,x} = -\overset{\leftrightarrow}{L} q_x \in \text{span}(\overset{\leftrightarrow}{L}) \text{ i.e. } \sum_{\tilde{y}} q_y \overset{\leftrightarrow}{L}(\cdot, \tilde{y})$$

to get consistency for l , it is sufficient to consider $s \in \text{span}(\overset{\leftrightarrow}{L})$
 or that $s \in \text{span}(F) \supseteq \text{span}(\overset{\leftrightarrow}{L})$
restriction on scores

$$F \in \mathbb{R}^{k \times r} \text{ matrix}$$

can be chosen cleverly depending on $\overset{\leftrightarrow}{L}$

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

$$\mathcal{J}_q(\theta) - \min_{\theta \in \mathbb{R}^r} \mathcal{J}_q(\theta) = \frac{1}{2k} \|F\theta - (\overset{\leftrightarrow}{L} q)\|_2^2$$

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

, Cover bound $\Rightarrow$ easiness result

thm. 7

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

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

Cover bound $\Rightarrow$ easiness result

this bad

$$\Delta_{ij} \triangleq \ell_i - \ell_j \in \mathbb{R}^k$$

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

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

thm. 8: if $\text{span}(F) = \mathbb{R}^k$ (i.e. no constraints) hardness result

then $H(\epsilon) \geq \frac{\epsilon^2}{2k}$ for any loss?

ie. for any loss, we need an exponential # of samples (in the worst case) to learn "well" } caveat $\rightarrow$ all these are bounds and worst case

⊗ but for hamming loss, if add constraints that $\sum_{P \in \mathcal{P}} \mathbb{P}(\tilde{y}) \leq \sum_{P \in \mathcal{P}} \mathbb{P}(y)$

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 for Hamming loss & separable score ...

optimization normalization

setup Let $s(x, \tilde{y})$ be of the form $F\Theta(x)$ $\Theta(x) \in \mathbb{R}^r$

$\Theta_j(\cdot) \in \mathcal{H} \subset \mathbb{R}^{\text{HKS}}$
optimization variables

$$J(\Theta) = \mathbb{E}_{(x,y) \sim p} s(x, y, \Theta)$$

run projected SGD on $J(\Theta)$ i.e. $\Theta^{(t+1)} = \mathcal{P}_{\Theta} \{ \Theta^{(t)} - \gamma \nabla_{\Theta} s(x^{(t)}, y^{(t)}, \Theta^{(t)}) \}$

a ball of radius D

$$\frac{(x^{(t)}, y^{(t)}) \tilde{\text{Id}}_p}{F^T \nabla_s s(x^{(t)}, y^{(t)}, s) \Phi(x^{(t)})^T} \quad \begin{array}{l} \text{feature map} \\ \text{on } \mathcal{H} \end{array}$$

$K(x^{(t)}, \cdot)$

Convergence result: if $\|\Theta^*\|_{\text{HKS}} \leq D$
 (thm. 5)

$$\text{if } \mathbb{E}_{(x,y) \sim p} \|\nabla_{\Theta} s(x, y, \Theta)\|_{\text{HKS}} \leq M^2$$

then averaged SGD with stepsize $\gamma = \frac{2D}{M\sqrt{n}}$

$$\mathbb{E}[J(\bar{\Theta}^{[n]})] - J(\Theta^*) \leq \frac{2DM}{\sqrt{n}} \quad (\text{convergence result})$$

$$\mathbb{E} \left[\frac{1}{n} \sum_{t=1}^n \ell(\sigma^{(t)}) - J(\sigma^*) \right] \leq \frac{2DM}{\sqrt{n}}$$

(convergence
result)

thm. 6 Learning complexity (to be ctd) ?