

today: consistency for convex surrogate losses

non-parametric viewpoint on scores

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

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

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

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

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

"product kernel"
↑
Kronecker product

$$v \otimes w \rightarrow vw^T$$

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

$$\text{e.g. } k_2(x, x') = \exp\left(-\frac{\|x - x'\|^2}{2\sigma^2}\right) \quad \text{RBF kernel (universal)}$$

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

↑
10²³

back to consistency & surrogate losses

$$\hat{w}_n \triangleq \underset{w}{\text{argmin}} \left\{ J_n(w) + \lambda_n \frac{\|w\|^2}{2} \right\}$$

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

⊛ binary classification [Bartlett et 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
McAllester 2007

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

is not consistent for 0-1 loss when have no "majority" class (i.e. $p(\tilde{y}|x) < \frac{1}{2} \forall \tilde{y}$)

→ they propose a different surrogate loss that uses $\sum_y$ 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

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

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

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

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

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

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

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

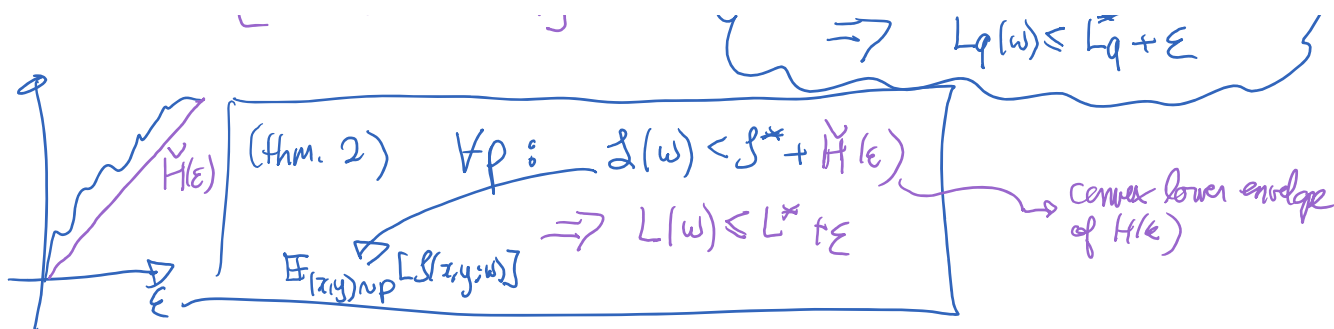
$$\mathcal{L}q(w) \triangleq \mathbb{E}_{q(y)} [\ell(y, h_w(x))]$$

"conditional (bipolar) risk"

[conditional on x version]

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

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



(↑ basically shown using Jensen's ineq.)

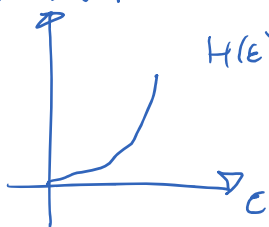
$$H(\epsilon) \triangleq H^{**}(\epsilon)$$

$$f^*(z) \triangleq \sup_z z^T z - f(z) \leftrightarrow \text{"Fenchel-legendre conjugate"}$$

if H is invertible

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

standard H :



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

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

you want small C ; for structural prediction $C = |S|$ often? (bad)

$\mathcal{L}$ is consistent

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



note: scale of H is arbitrary

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

14h34

concrete example: simplest surrogate loss: square loss?

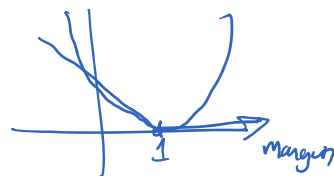
$$s(\cdot) \in \mathbb{R}^k \text{ (for } x)$$

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

$\mathcal{L}$ can be seen as generalization of squared loss for binary class, to multiclass

$$y_i = 1 \quad [1 - y_i \leq w, \ell(x_i)]^2$$

$$= [y_i - \langle w, \ell(x_i) \rangle]^2$$



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

$$= \frac{1}{2k} \sum_{(x,y)} [s(y)]^2 + 2s(y)\ell(y,\tilde{y}) + \text{const.}$$

does not depend on s

$$= \frac{1}{2k} \sum_y \mathbb{E}_{q(y)} [s(\tilde{y})^2 + 2s(\tilde{y})\ell(y, \tilde{y})] + \text{cst.}$$

$$= \frac{1}{2k} \|s + l_{q_x}\|_2^2 + \text{cst.} \quad l_{q_x}(\tilde{y}) \triangleq \mathbb{E}_{q(y)} [\ell(y, \tilde{y})]$$

suppose s is unconstrained, $\min_s \mathcal{J}_q(s) \Rightarrow s^*(\tilde{y}) = -l_{q_x}(\tilde{y})$

$$\arg\max_y s^*(\tilde{y}) = \arg\min_y l_{q_x}(\tilde{y})$$

ie. you predict optimally partwise on x

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

$$\Rightarrow L(h_{s^*}) = \min_{\text{all } h} L(h)$$

$$\mathcal{J}_q(s) - \overbrace{\min_{s' \in \mathbb{R}^k} \mathcal{J}_q(s')}^{\mathcal{J}_q^*} = \frac{1}{2k} \|s - (-l_{q_x})\|_2^2$$

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

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

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

⊗ 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 scores

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

can be chosen depending on $\tilde{L}$

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

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

thm. 7

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

$H_{\text{span}(F), \ell, F}(\epsilon)$

lower bound $\Rightarrow$ existence result

$$\geq \frac{\epsilon^2}{2k \max_{i,j} \|F \Delta_{ij}\|_2^2}$$

$$\geq \frac{\epsilon^2}{4k} \quad \text{this is bad}$$

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

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 $\text{span}(F) = \mathbb{R}^k$ (i.e. no constraints) hardness result
 then $H(\epsilon) \leq \frac{\epsilon^2}{4k}$ for any loss?

i.e. for any loss, we need an exp. # of samples (in worst case)
 to learn "well"
 [correct $\rightarrow$ all these bounds are worst case]

$\rightarrow$ Without structure on F , you're doomed in worst case over $p(x, y)$

⊗ but for Hamming loss, if add constraints that $s(\tilde{y}) = \sum_{\text{reports}} s_p(\tilde{y}_p)$
 over T binary variables, $H(\epsilon) = \frac{\epsilon^2}{8T}$ } not too big $\Rightarrow$ we can learn?

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

(polynomial sample complexity) (11)

$\rightarrow$ efficient to compute for

Hamming loss & separable score fct. ex. (15)