

today:
 • review surrogate losses
 • generalization error bounds
 { structured SVM

Review of convex surrogate losses mentioned so far:

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

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

$$= \max_{\tilde{y} \in \mathcal{Y}} [-m(\tilde{y})] \quad (= [\max_{\tilde{y} \neq y} -m(\tilde{y})]_+ \text{ if } y \in \mathcal{Y})$$

$$\mathcal{J}_{\text{hinge}}(\quad) = \max_{\tilde{y}} [s(\tilde{y}) + \ell(y, \tilde{y})] - s(y)$$

$$\text{"margin rescaling"} \rightarrow \max_{\tilde{y} \in \mathcal{Y}} [\ell(y, \tilde{y}) - m(\tilde{y})]$$

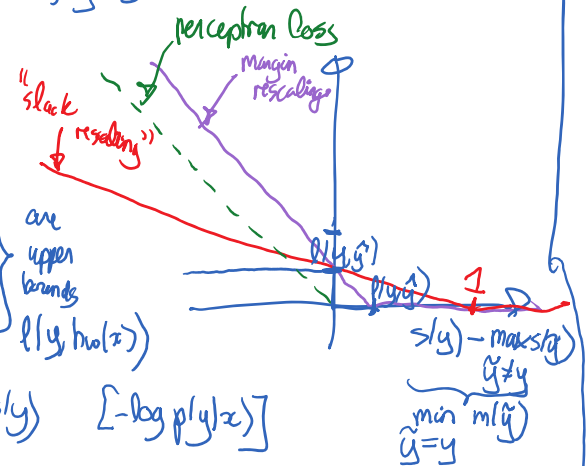
$$\text{"slack rescaling"} = \max_{\tilde{y}} \ell(y, \tilde{y}) [1 - m(\tilde{y})]$$

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

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

$$\frac{1}{\beta} \log \left(\sum_{\tilde{y}} \exp(-\beta m(\tilde{y})) \right) \xrightarrow{\text{suggests "smoothed hinge loss"}}$$

$$\frac{1}{\beta} \log \left(\sum_{\tilde{y}} \exp(\beta (\ell(y, \tilde{y}) - m(\tilde{y}))) \right) \quad \text{[e.g. Plotkin et al. 2010]}$$



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

but more computationally costly

what are theoretical properties could we look at?

a) generalization error bounds [today]

b) consistency properties { calibration function [next class]

$\hookrightarrow$ relationship between $L(w)$ & $\mathcal{J}(w)$

why structured 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. Σ)

$\Rightarrow$ easier to learn
(generalization guarantees)

[see Cortes & al. NeurIPS 2016]
(today)

2) computational " : compute $\arg\max_{y \in \mathcal{Y}} s(y)$

generalization error bounds:

for binary classification

a classical PAC bound is:

for any fixed dist. p on data
with prob. $\geq 1-\delta$ on D_n

$$\forall w \in W \quad L_0(w) \leq \hat{L}_n(w) + \frac{1}{\sqrt{n}} \sqrt{d \log \frac{1}{\delta} + \log \frac{1}{\delta}}$$

where d is VC-dimension of $H \triangleq \{h_w : w \in W\}$

VC-dimension of $H \triangleq \max \{m : \exists \text{ a set of } m \text{ points s.t.} \\ \forall \text{ labelings of these points} \\ \exists w \text{ s.t. } h_w \text{ gives the} \\ \text{correct label on these points}\}$
"challenging the set of points"

$\hookrightarrow$ # of binary functions on m points
 2^m

for $H = \{ \text{linear classifiers of } p \text{ parameters} \}$ VC-dim(H) = $p+1$

* one issue for this is that it's true for all distributions $\Rightarrow$ too loose bound

$\Rightarrow$ motivates going to data distribution dependent
measure of complexity

example: empirical Rademacher complexity

$$\hat{R}_n(H) \triangleq \mathbb{E}_{\sigma} \left[\sup_{h_w \in H} \left| \frac{2}{n} \sum_{i=1}^n \sigma_i \mathbb{1}\{y_i \neq h_w(x_i)\} \right| \right]$$

"correlation with random noise"

$\sigma_i = \begin{cases} +1 \\ -1 \end{cases}$ uniformly "Rademacher R.V."

bound : with prob $\geq 1-\delta$

$$\forall w \quad L_0(w) \leq \hat{L}_n(w) + \hat{R}_n(H) + \frac{1}{\sqrt{n}} 3 \sqrt{\log \frac{2}{\delta}}$$

$$\forall w \quad L(w) \leq \hat{L}_n(w) + \hat{R}_n(H) + \frac{1}{\sqrt{n}} 3 \sqrt{\log \frac{2}{\delta}}$$

complexity depends on D_n (implicitly on p)

14h30

structured prediction generalization bounds [Cortes et al. NeurIPS 2016]

general loss fct. $l(y, y')$ s.t. $l(y, y') \geq 0 \quad \forall y, y'$

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

$\hookrightarrow$ set of cliques of a graph model G / factor graph

thm. 7 with prob $\geq 1-\delta$

$$\forall w \in W \quad L(w) \leq \hat{J}_{\text{hinge}}(w) + 4\sqrt{2} \hat{R}_n^G(H_w) + 3 \frac{\max_{y, y'} l(y, y')}{\sqrt{n}} \sqrt{\log \frac{1}{\delta}}$$

where $\hat{R}_n^G \triangleq \frac{1}{n} \mathbb{E}_\sigma \left[\sup_{w \in W} \sum_{i=1}^n \sqrt{|\mathcal{C}_i|} \sum_{c \in \mathcal{C}_i} \sum_{y_c \in \mathcal{Y}_c} \sigma_{i,c,y_c} S_c(x_i, y_c; w) \right]$

actually only depends on $(x^{(i)})_{i=1}^n$ "empirical factor graph complexity" indep Rademacher R.V.

thm. 2 : If $S_c(x, y_c; w) = \langle w, \varphi_c(x, y_c) \rangle$

and consider $W_\Lambda \triangleq \{w : \|w\|_2 \leq \Lambda\}$; let $R = \max_{i,c,y_c} \|\varphi_c(x_i, y_c)\|_2$

$$\text{then } \left| \hat{R}_n^G(H_{W_\Lambda}) \leq \frac{R \Lambda}{\sqrt{n}} \max_i |\mathcal{C}_i| \sqrt{\max_{c \in \mathcal{C}_i} |\mathcal{Y}_c|} \right|$$

$\downarrow$ so want small cliques?

* plug thm. 2 in thm. 7 :

$$L(w) \leq \hat{J}_{\text{hinge}}(w) + \underbrace{\left(\frac{R \Lambda}{\sqrt{n}} \sqrt{\max_c |\mathcal{Y}_c|} \right)}_{\frac{\Lambda n}{2}} \underbrace{\|w\|_2}_{\frac{1}{2}} + \text{const.}$$

min of RHS suggests

$$\text{SVM struct dg. } \hat{w}_n = \arg \min_w \hat{J}_{\text{hinge}}(w) + \frac{\Lambda n}{2} \|w\|_2^2$$

missing link : (1) $\min_{\text{opt. } \|w\|_2 \leq 1} f(w)$ (if f is convex) , use Lagrangian

missing link: ① $\min_{\text{s.t. } \|w\|_2 \leq \sqrt{\lambda}} f(w)$

② $\min f(w) + \frac{\lambda}{2} \|w\|^2$

(if f is convex)

use Lagrangian

$\exists \lambda(\lambda) \leq 0$

then to ② same solution to ①

[note: constrained formulation can have solutions not achievable for ② when f is non-convex
but penalized/reg. formulation is less sensitive to choice of λ vs. constrained formulation

properties: • minimize upper bound, hope that $\min L(w)$

but no general guarantees



• Can evaluate bound to get guarantees

caveat:

also note here: no consistency guarantee

next: consistency + convex surrogate

consistency & calibration

need to relate $J(w)$ to $L(w)$: bad "calibration fct." [Steinwart]

relationship is usually very complicated

$\Rightarrow$ usual results look mainly at non-parametric setting (no # of parameters)

all functions $h: X \rightarrow Y$ are considered $\Rightarrow$ this evaluates the dependence on x of the analysis "pointwise analysis"

i.e. we suppose that $S(x, y; w)$ can be arbitrary for any x
(i.e. w is ∞ -dim.)

$\rightarrow$ can do this using a universal kernel

$$S(\cdot, \cdot, w) \in \mathcal{H}_{X \times Y}$$

RKHS:

motivation:



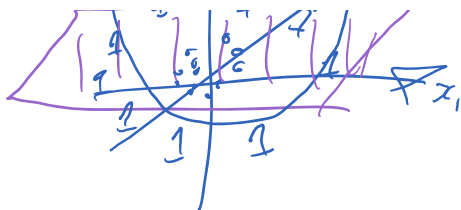
generalize linear structure

$\langle w, \phi(x) \rangle$ to higher dim.

+ kernel trick $\langle \phi(x), \phi(x') \rangle = k(x, x')$

$$\Phi: X \rightarrow \mathbb{R}^3$$

$\begin{matrix} \text{---} 1 & \text{---} 2 & \text{---} \end{matrix}$



$$\Phi: X \rightarrow \mathbb{R}^3$$

$$\Phi(x) = \begin{pmatrix} x_1^2 \\ \sqrt{2}x_1x_2 \\ x_2^2 \end{pmatrix}$$

$$\langle \Phi(x), \Phi(x') \rangle_{\mathbb{R}^3} = (\langle x, x' \rangle_{\mathbb{R}^2})^2 = K(x, x')$$

polynomial kernel e.g. $(\langle x, x' \rangle + 1)^p = k_{\text{poly}}(x, x')$

equivalent to mapping data to a space of dimension exponential in p

$$\langle \Phi(x), \Phi(x') \rangle$$

even have to dim e.g. $K(x, x') = \exp(-\frac{\|x - x'\|^2}{2})$

"RBF kernel"