

Lecture 3 - scribes

Tuesday, January 15, 2019 14:12

today: energy based methods & surrogate losses
multiclass

energy based methods : [LeCun & al. 2006]

model: $h_w(x) = \underset{y \in \mathcal{Y}(x)}{\operatorname{argmin}} E(x, y; w)$ "energy fct."

$= \underset{y \in \mathcal{Y}(x)}{\operatorname{argmax}} S(x, y; w)$ "score / compatibility fct."



- ingredients:
modeling {
- 1) what is $E(x, y; w)$? e.g. $S(x, y; w) = \langle w, \phi(x, y) \rangle$
 - 2) how do you compute $\underset{y \in \mathcal{Y}(x)}{\operatorname{argmin}} E(x, y; w)$? $\rightarrow$ "inference" / "decoding"
- learning {
- 3) how to evaluate $E(x, y; w)$ on a training set? $\rightarrow$ surrogate loss $\hat{\mathcal{L}}(w)$
in general: $\hat{\mathcal{L}}(x^{(i)}, y^{(i)}, E(\cdot, \cdot))$ "loss functional"
 - 4) how to minimize $\hat{\mathcal{L}}(w)$ to learn w ? $\rightarrow$ optimization tricks

that multiclass case. :

$\in \mathbb{R}^d$

flat multiclass case :

"flat" setting $h_w(x) = \arg\max_y \langle w_y, \phi(x) \rangle$

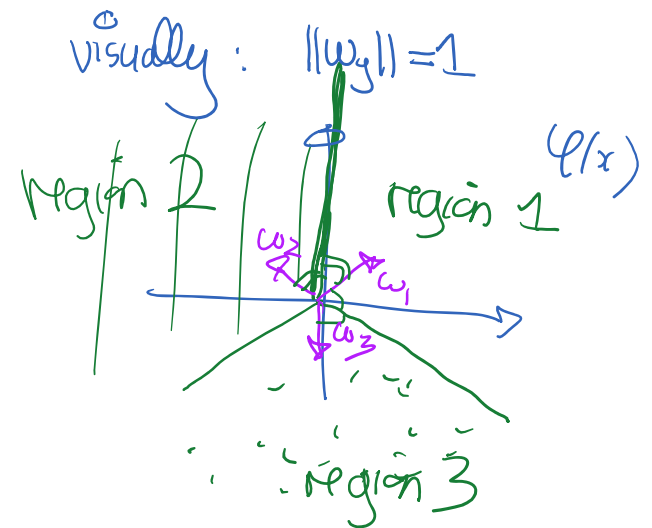
$\in \mathbb{R}^d$

equivalent to

$\phi(x, y) = \begin{pmatrix} 0 \\ 0 \\ \phi(x) \\ 0 \\ 0 \\ \vdots \end{pmatrix} \in \mathbb{R}^{d \times k}$

of classes

$\leftarrow y^{\text{th}} \text{ position}$



OR ^{node} feature map VS.

$\langle w, \phi(x, y) \rangle = \sum_p \underbrace{\langle w, \phi(x, y_p) \rangle}_{\sum_{y_p} \mathbb{1}\{y_p = y'\} \langle w_{y_p}, \phi(x) \rangle}$

$\rightarrow$ here "sharing" of parameters between different pieces of the labels $\rightarrow$ "structure"

aside: in structured prediction, usually absorb "bias" in parameters

standard binary classification: $\text{sign}(\langle w, x \rangle + b)$

$$\begin{pmatrix} \tilde{\phi}(x) \\ 1 \end{pmatrix}$$

$$\langle \tilde{w}, \tilde{\phi}(x) \rangle = \langle w, \phi(x) \rangle + b$$

$$\tilde{w} = \begin{pmatrix} w \\ b \end{pmatrix}$$

non question: regularization or not the bias

open question: regularizing or not the bias
in structured prediction, does it matter?

sumo gate losses:

$$\hat{J}(w) = \frac{1}{n} \sum_{i=1}^n \mathcal{L}(x^{(i)}, y^{(i)}; w) + R(w)$$

I) perceptron loss [Collins & d. 2002 EMNLP]

$$\mathcal{L}(x, y; w) = \max_{\tilde{y} \in \mathcal{Y}(x)} s(x, \tilde{y}; w) - \underbrace{s(x, y; w)}_{\text{score of ground truth}} \quad (\text{assume } y \in \mathcal{Y}(x))$$

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

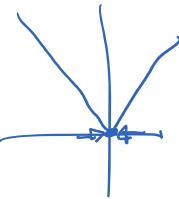
$$\max_{\tilde{y}} \langle w, \underbrace{\phi(x, \tilde{y}) - \phi(x, y)}_{-\phi(x, \tilde{y})} \rangle \geq 0 \quad \text{by using } \tilde{y} = y$$

observations: 1) degenerate solution $w=0$ or constant score over y

2) averaged perceptron alg.: → does not converge in general

• run constant step-size stochastic subgradient method on $\hat{J}(w)$

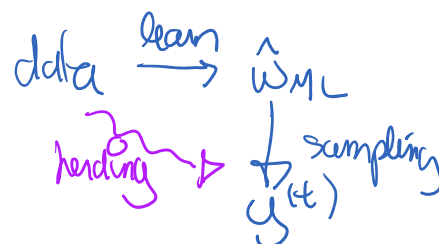
• output $\hat{w}_T = \frac{1}{T+1} \sum_{t=0}^T w_t$ (Polyak avg.) okay
→ will converge to $w^*=0$ when data is not separable



Comments: 1) Collins' paper $\rightarrow$ he gives error bound
and generalization error
guarantees for perceptron

2) (aside) connection with the "herding" alg. by Welling et al.
"3rd way to learn" [ICML 2012]

15h33



II) log-loss (CRF) (probabilistic interpretation)

Boltzmann dist. in physics

suppose $p(y|x;w) \propto \exp(\underbrace{\beta s(x,y;w)}_{\text{inverse temperature parameter}})$

$$(\beta = \frac{1}{k_B T})$$

MCL $\rightarrow$ log-loss

$$\mathcal{L}(x,y;w) = \underbrace{-\frac{1}{\beta}}_{\dots} \log p(y|x;w) = -\frac{1}{\beta} \log \left(\frac{\exp(\beta s(x,y;w))}{\sum_{\tilde{y}} \exp(\beta s(\tilde{y}))} \right) \quad \left\{ Z_{\beta}(x;w) \text{ partition} \right.$$

β
rescaling

$\beta \left\{ \sum_{\tilde{y}} \exp(\beta s(\tilde{y})) \right\} Z_{\beta}(x; w)$ partition set,

$$= \frac{1}{\beta} \log \left(\sum_{\tilde{y}} \exp(\beta s(\tilde{y})) \right) - \frac{\beta}{\beta} s(y)$$

"log-sum-exp" $\leadsto$ "soft-max" why?
let $\hat{y} = \underset{\tilde{y}}{\operatorname{argmax}} s(\tilde{y})$

NOTE:
in deep learning book

$$\left(\frac{\exp(s(y))}{\sum_{\tilde{y}} \exp(s(\tilde{y}))} \right)_{y \in \mathcal{Y}}$$

I call this
"soft-argmax"

$$\frac{1}{\beta} \log \left[\exp(\beta s(\hat{y})) \underbrace{\left[\sum_{\tilde{y}} \exp(\beta (s(\tilde{y}) - s(\hat{y}))) \right]}_{\leq |\mathcal{Y}|} \right]$$

$$= s(\hat{y}) + \frac{1}{\beta} \log(\text{stuff})$$

$\leq |\mathcal{Y}|$

$\beta \rightarrow \infty$ (ie. zero temp limit)

$$\frac{1}{\beta} \log \left(\sum_{\tilde{y}} \exp(\beta s(\tilde{y})) \right) \xrightarrow{\beta \rightarrow \infty} \max_{\tilde{y}} s(\tilde{y})$$

thus: $\lim_{\beta \rightarrow \infty} \text{log-loss}(\beta) \rightarrow \text{perception loss?}$

III) structured hinge loss

$$\mathcal{L}(x, y; w) = \max_{\tilde{y} \in \mathcal{H}(x)} [s(x, \tilde{y}; w) + \ell(y, \tilde{y})] - s(x, y; w)$$

$$\mathcal{L}(x, y; w) = \max_{\tilde{y} \in \mathcal{Y}(x)} [s(x, \tilde{y}; w) + \ell(y, \tilde{y})] - s(x, y; w)$$

"loss-augmented decoding"

a)
carlson:

$$\begin{aligned} & \geq \ell(y, y_{\max}) \\ & \quad \begin{array}{c} \{ s(y) \\ s(y_{\max}) \\ \vdots \end{array} \Rightarrow \mathcal{L}^{\text{sum}}(x, y; w) = 0 \end{aligned}$$

b) $\mathcal{L}(x, y; w) \geq \ell(y, h_w(x))$

why?

$$\mathcal{L}(x, y; w) = \max_{\tilde{y}} [s(\tilde{y}) + \ell(y, \tilde{y})] - s(y)$$

$$\geq s(\hat{y}) + \ell(\hat{y}) - s(y)$$

using $\hat{y} = \arg\max_{\tilde{y} \in \mathcal{Y}(x)} s(\tilde{y}) = h_w(x)$

if $y \in \mathcal{Y}(x) \Rightarrow s(\hat{y}) \geq s(y)$

$$\geq \ell(\hat{y}) = \ell(y, h_w(x)) //$$

binary case:

$$y \in \{-1, +1\}$$

$$w = \begin{pmatrix} w_+ \\ w_- \end{pmatrix}$$

$$\ell(x, +1) = \begin{pmatrix} \ell(x) \\ 0 \end{pmatrix}$$

$$h_w(x) = \arg\max \{ \langle w_+, x \rangle, \langle w_-, x \rangle \}$$

predict +1 if $\langle w_+, x \rangle \geq \langle w_-, x \rangle$

$$\Leftrightarrow \langle w_+ - w_-, x \rangle \geq 0$$

$$\tilde{w} \triangleq w_+ - w_-$$

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

$$\mathcal{L}^{\text{SVM}}(x, y; w) = \max \left\{ \underbrace{\langle w_+, x \rangle + \ell(y, +)}_{\mathbb{1}\{y \neq +1\}}, \underbrace{\langle w_-, x \rangle + \ell(y, -)}_{\mathbb{1}\{y \neq -1\}} \right\} - \langle w_y, x \rangle$$

$$\max \left\{ \langle \tilde{w}, x \rangle + \langle w_-, x \rangle + \mathbb{1}\{y \neq +1\}, \langle w_-, x \rangle + \mathbb{1}\{y \neq -1\} \right\} - \langle w_y, x \rangle$$

$$= \max \left\{ \langle \tilde{w}, x \rangle + \mathbb{1}, \mathbb{1} - \mathbb{1} \right\} + \langle w_-, x \rangle - \langle w_y, x \rangle$$

$$\begin{aligned} &\xrightarrow{y=+1} \max \left\{ \langle \tilde{w}, x \rangle, \mathbb{1} \right\} - \langle \tilde{w}, x \rangle = [1 - \langle \tilde{w}, x \rangle]_+ \\ &\xrightarrow{y=-1} \max \left\{ \langle \tilde{w}, x \rangle + \mathbb{1}, 0 \right\} + 0 = \max \left\{ 0, 1 - y \langle \tilde{w}, x \rangle \right\} \end{aligned}$$

structured SVM

$$\mathcal{L}(x, y; w) = [1 - y \langle \tilde{w}, x \rangle]_+$$

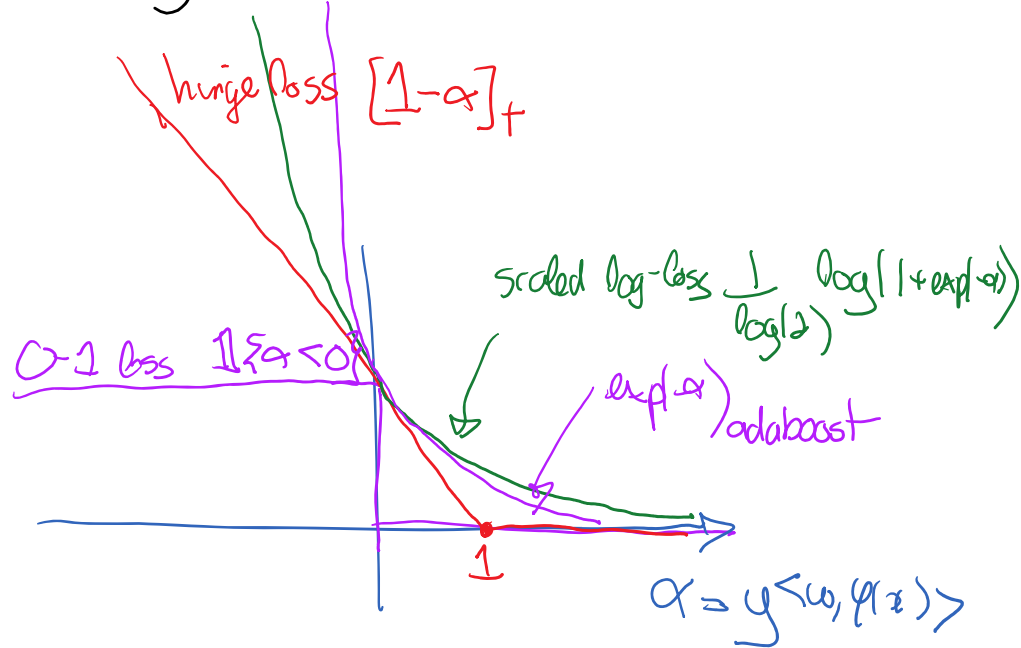
where $\tilde{w} = w_+ - w_-$

overall: $[1 - y \langle \tilde{w}, x \rangle]_+$

i.e. structured hinge loss
reduces to binary SVM hinge loss
when using $\ell(y, y) = \mathbb{1}\{y \neq y\}$
and $\mathcal{Y} = \{-1, +1\}$

Binary surrogate losses

Many surrogate losses



$y \langle w, \phi(x) \rangle$ "margin" $\geq 0 \Rightarrow$ make no mistake
 $\rightarrow \in \{-1, +1\}$

log loss: $\log(1 + \exp(-\alpha))$
 [logistic regression]

[BenRett et al. 2006] $\rightarrow$ showed all these methods are consistent