

Lecture 3 - Surrogate losses

Tuesday, January 24, 2023 2:49 PM

today: energy based methods & surrogate losses
multiclass

OCR - optical character recognition example

x : sequence of images of characters



$$x = (x_1, \dots, x_{L_x})$$

$$x_p \in \{0, 1\}^{6 \times 8}$$

y

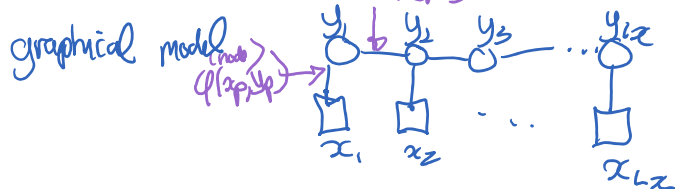
B R A C E

$$Y(x) = \sum_{L_x} x$$

$$\Sigma = \{A, \dots, Z\}$$

in max-margin Markov network (M³-net) papers:

$$\langle w, \ell(x, y) \rangle = \sum_{p=1}^{L_x} \langle w^{(node)}, \phi^{(node)}(x_p, y_p) \rangle + \sum_{p=1}^{L_x-1} \langle w^{(edge)}, \phi^{(edge)}(y_p, y_{p+1}) \rangle$$



$$p_w(y|x) = \frac{1}{Z_w(x)} \exp(\langle w, \ell(x, y) \rangle) = \frac{1}{Z_w(x)} \prod_{C \in \mathcal{C}} \psi_C(x_C, y_C)$$

where $\mathcal{C} = \{p, p+1\}$

notation $y_C \triangleq (y_i)_{i \in C} \Rightarrow$ can compute $\arg \max_y \langle w, \ell(x, y) \rangle$

dimension: $6 \cdot 8 \cdot 26$
of characters

node: $\phi^{(node)}(x_p, y_p) = \begin{pmatrix} 0 \\ 0 \\ \text{vector}(x_p) \end{pmatrix}$

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

w_a

w_b

using max-product dg, i.e. aka. Viterbi dg or max sum

$$\begin{aligned} \langle w, \phi^{(node)}(x_p, y_p) \rangle &= 0 + 0 + \langle w_{y_p}, x_p \rangle + 0 \\ &\quad \text{template for } y_p \end{aligned}$$

edge feature: $\phi(y_p, y_{p+1}) = \begin{pmatrix} 1 \text{ if } y_p = y_{p+1} \\ \vdots \\ \text{#}(y_p, y_{p+1}) \end{pmatrix} \quad 26^2$

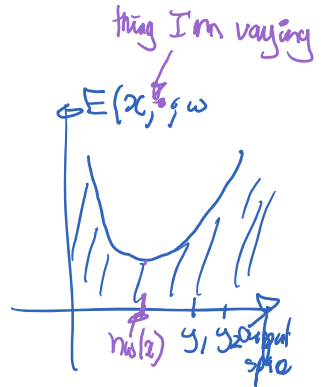
$$\langle w^{(edge)}, \phi(y_p, y_{p+1}) \rangle = w_{y_p, y_{p+1}}^{(edge)}$$

$$\langle w^{(\text{edge})}, \phi(y_p, y_{p+1}) \rangle = w_{y_p, y_{p+1}}^{(\text{edge})}$$

energy based methods : [Sutton & 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"



- ingredients :
- 1) what is $E(x, y; w)$? e.g. $s(x, y; w) = \langle w, \phi(x, y) \rangle$
or $E(x, y; w)$ scalar output of a NN with $x \oplus y$ as input
 - 2) how do you compute $\underset{y \in \mathcal{Y}(x)}{\operatorname{argmin}} E(x, y; w)$? $\rightarrow$ "decoding" / "inference"
 - 3) how to evaluate $E(x, y; w)$ on a training set? $\rightarrow$ surrogate loss
in general: $\tilde{J}(x^{(i)}, y^{(i)}; E(\cdot, \cdot, w))$
"loss functional"
 - 4) how to minimize $\tilde{J}(w)$ to learn $\hat{w}$? $\rightarrow$ optimization tricks

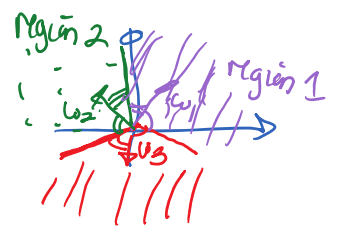
flat multiclass case

"flat" (i.e. non structured) setting $h_w(x) = \underset{y}{\operatorname{argmax}} \langle w_y, \phi(x) \rangle$

equivalent to $\phi(x, y) = \begin{pmatrix} 0 \\ 0 \\ \phi(x) \\ 0 \end{pmatrix}$ a y^{th} position $\in \mathbb{R}^{d \cdot K}$ # of classes

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

visually: $\|w_y\| = 1$



contrast this flat case
with

structured case

e.g. OCR: node feature map $\langle w, \phi(x, y) \rangle = \sum_p \langle w, \phi^{(node)}(x_p, y_p) \rangle$

→ have "sharing" of parameters between pieces of the joint labels
→ "structure"

aside: in structured prediction, usually absorb "bias" in parameters $\tilde{\phi}(x)$

$$\langle \tilde{w}, \tilde{\phi}(x) \rangle = \langle w, \phi(x) \rangle + b \quad \begin{pmatrix} \phi(x) \\ 1 \end{pmatrix}$$

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

open question: regularizing or not the bias

does it matter in structured pred.?

15h 45

Sumogette losses

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

I) percept loss [Collins et al 2002 EMNLP]

$$J_{\text{percep}}(x, y; w) = \left[\max_{\tilde{y} \in \mathcal{Y}(x)} s(x, \tilde{y}; w) - s(x, y; w) \right]$$

score of ground truth

not needed if assume $y \in \mathcal{Y}(x)$

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

by using $\tilde{y} = y$

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

observations: 1) degenerate solution to $J(w)$ with $w=0$ or constant score over y
2) overused reconstruction de.

observations: 1) degenerate solution to $J(w)$ with $w=0$ or constant over y

2) averaged perceptron alg.:

• amounts to running constant step-size stochastic subgradient method on $J(w)$

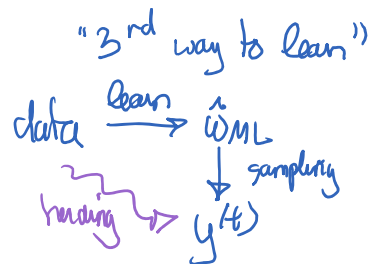
• output $\hat{w}_T = \frac{1}{T+1} \sum_{t=0}^T w_t$ (Polyak avg.)

↳ will converge to $w^*=0$ when data is not separable

comments: 1) Collins's paper → he gives error bound

and generalization error guarantees for perceptron

2) (aside) connection with the "harding" alg. by Weidinger & al. [ICML 2012]



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

suppose $p(y|x;w) \propto \exp(\beta s(x,y;w))$

↑
"inverse temperature" parameter

Boltzmann dist

$$\beta = \frac{1}{k_B \cdot T}$$

temperature

MCL → log-loss

$$\begin{aligned} \mathcal{L}(x,y;w) &= \underbrace{-\frac{1}{\beta}}_{\text{rescaling}} \log p(y|x;w) = -\frac{1}{\beta} \log \left(\frac{\exp(\beta s(x,y;w))}{\sum_y \exp(\beta s(x,y;w))} \right) \\ &= \underbrace{-\frac{1}{\beta} \log \left(\sum_y \exp(\beta s(x,y;w)) \right)}_{\text{"log-sum-exp" → "soft max" why?}} - s(x,y;w) \end{aligned}$$

partition fun.

let $\hat{y} = \text{argmax}_y s(x,y)$

$$\text{let } \hat{y} = \underset{y}{\text{argmax}} s(y)$$

$$\frac{1}{\beta} \log(\exp(\beta s(\hat{y})) \left[\sum_y \exp(\beta(s(y) - s(\hat{y}))) \right])$$

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

≤ 0
 ≈ 1.51

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

$$\frac{1}{\beta} \log\left(\sum_y \exp(\beta \cdot s(y))\right) \xrightarrow{\beta \rightarrow \infty} \max_y s(y)$$

note:
in deep learning
they call "soft max"

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

I call this
"soft-argmax"

thus $\lim_{\beta \rightarrow \infty} \log\text{-loss}(\beta) \rightarrow \text{perceptron loss}$