

## Lecture 2 - examples

Thursday, January 12, 2023 1:28 PM

Reminder: do not forget to fill survey <http://bit.ly/IFT6132-W23> ASAP if not done!

today: • examples of structured prediction  
• structured perceptron { friends

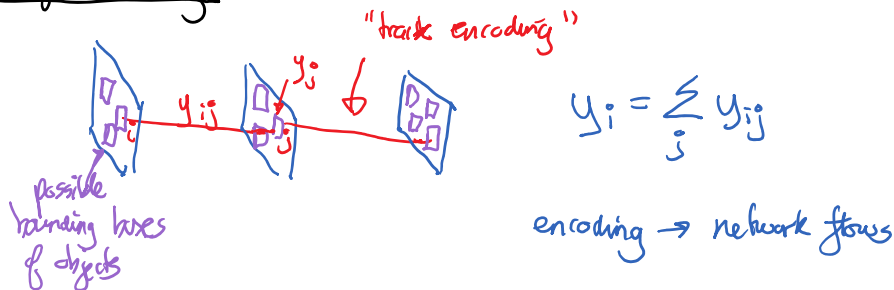
### Examples

#### I) word alignment (continuation)

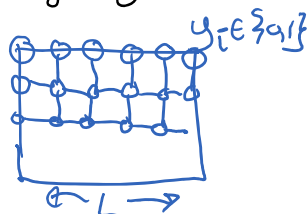
here  $x = (\underbrace{x_1^E, x_2^E, \dots, x_{L_E}^E}_{\text{English words}}; \underbrace{x_1^F, \dots, x_{L_F}^F}_{\text{French words}})$

$$\mathcal{Y}(x) = \{ y \in \{0,1\}^{L_E \times L_F} : \sum_j y_{ij} \leq 1, \sum_i y_{ij} \leq 1 \}$$

#### II) multi-object tracking:



#### III) image segmentation



$x =$  image of RGB values  $L \times L$  pixels

$$\mathcal{Y}(x) = \{0,1\}^{L \times L}$$

background foreground

#### prediction model $hw(x)$

standard:  $hw(x) \triangleq \underset{y \in \mathcal{Y}(x)}{\operatorname{argmax}} \begin{bmatrix} s(x,y;\omega) \\ -E(x,y;\omega) \end{bmatrix}$  compatibility score of  $y$  for  $x$   
energy fct.  $E$

linear model:  $s(x,y;\omega) = \langle \omega, \underbrace{\phi(x,y)}_{\text{"joint feature vector"} \in \mathbb{R}^d} \rangle$   $\phi: X \times \mathcal{Y} \rightarrow \mathbb{R}^d$

word alignment:  $\phi(x, y) = \sum_{i,j} y_{ij} \psi(x_i^E, x_j^F)$

features defined on a pair of English word  $x_i^E$  French "  $x_j^F$

string edit distance  $(x_i^E, x_j^F)$   
 • distance between  $i$  &  $j$   
 •  $\{x_i^E, x_j^F\}$  in dictionary etc.

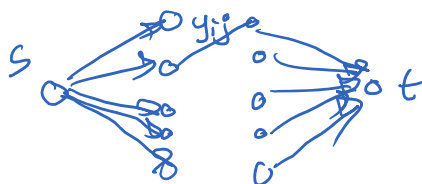
$$s(x, y; w) = \langle w, \phi(x, y) \rangle = \sum_{i,j} y_{ij} \langle w, \psi(x_i^E, x_j^F) \rangle$$

"score to match  $i$  to  $j$ "

$$h_w(x) = \arg \max_{y \in \mathcal{Y}(x)} s(x, y; w) \leadsto \max_y \sum_{i,j} y_{ij} s_{ij}(x)$$

s.t.  $y_{ij} \in \{0, 1\}$   
 $\sum_j y_{ij} \leq 1 \forall i$   
 $\sum_i y_{ij} \leq 1 \forall j$

can be solved exactly  
 as min linear cost matching problem



e.g. Hungarian alg  
 or more generally  
 min cost network flow alg

14h29

Learning w?

[x denote: integer program with LP relaxation]

good question about  
 tradeoff: tractable search  
 vs. more powerful models

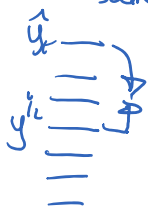
I) structured perceptron:

- initialize  $w_0$
- repeat for  $t=0, \dots$

$$\left[ \begin{array}{l} \text{• sample } i_t \\ \text{• let } \hat{y}_t = h_{w_t}(x^{(i_t)}) = \arg \max_{y \in \mathcal{Y}(x^{(i_t)})} \langle w_t, \phi(x^{(i_t)}, y) \rangle \\ \text{• } w_{t+1} = w_t + \eta \left( \underbrace{\langle \phi(x^{(i_t)}, y^{(i_t)}) \rangle}_{\substack{\uparrow \\ \text{step-size} \rightarrow \text{boost score of ground truth}}} - \underbrace{\langle \phi(x^{(i_t)}, \hat{y}_t) \rangle}_{\text{penalize incorrect prediction}} \right) \end{array} \right]$$

'decoding oracle'

$$\text{score}_{t+1}(\hat{y}) = \text{score}_t(\hat{y}) + \eta [\langle \phi(y^{(i_t)}, \hat{y}) \rangle - \langle \phi(\hat{y}_t), \hat{y} \rangle]$$



for stability: output  $\hat{w}_T = \frac{1}{T+1} \sum_{t=0}^T w_t$  "Polyak averaging"

⊗ structured perceptron can be interpreted as

doing stochastic subgradient method (opt.) on the following non-smooth obj.:

$$\hat{J}(w) = \frac{1}{n} \sum_{i=1}^n g^{\text{percept.}}(x^{(i)}, y^{(i)}; w)$$

$$g^{\text{percept.}}(x, y; w) \triangleq \left[ \max_{\tilde{y} \in \mathcal{Y}} \langle w, \ell(x, \tilde{y}) \rangle - \langle w, \ell(x, y) \rangle \right]_+$$

where  $[a]_+ \triangleq \begin{cases} a & \text{if } a \geq 0 \\ 0 & \text{o.w.} \end{cases}$   
 if  $y^{(i)} \in \mathcal{Y}$ ; then this is always  $\geq 0$  and  $[\cdot]_+$  is not needed

## II) conditional random field

define  $p_w(y|x)$  or  $\exp(\langle w, \ell(x, y) \rangle)$

$$h_w(x) = \underset{y \in \mathcal{Y}(x)}{\text{argmax}} p_w(y|x) = \underset{y}{\text{argmax}} \langle w, \ell(x, y) \rangle$$

then maximum conditional likelihood on training set to learn  $\hat{w}$

$$\hat{J}^{\text{CAF}}(w) = \frac{1}{n} \sum_{i=1}^n J^{\text{CAF}}(x^{(i)}, y^{(i)}; w) + \underbrace{\lambda \|w\|^2}_2 \text{ regularizer}$$

$$J^{\text{CAF}}(x, y; w) \triangleq -\log p_w(y|x)$$

$$= \log \left( \sum_{\tilde{y} \in \mathcal{Y}(x)} \exp(\langle w, \ell(x, \tilde{y}) \rangle) \right) - \langle w, \ell(x, y) \rangle$$

issues:  $\ell(y, y')$  doesn't appear in it

$\sum_{\tilde{y} \in \mathcal{Y}} \exp(\langle w, \ell(x, \tilde{y}) \rangle)$  is often intractable

e.g. #complete problem for  $\mathcal{Y}$  = set of all matchings?

## III) structured SVM

intuition: want  $s(x^{(i)}, y^{(i)}; w) \geq s(x^{(i)}, \tilde{y}; w) + \ell(y^{(i)}, \tilde{y}) \quad \forall \tilde{y} \in \mathcal{Y}_i \triangleq \mathcal{Y}/x_i$   
 $\min \|w\|^2$  s.t.  $\rightarrow$  "hard margin structured SVM"

$$\left[ \text{binary SVM: } y \in \{-1, +1\}, h_w(x) = \text{sgn}(\langle w, \ell(x) \rangle) \right]$$

$$y_i \langle w, \ell(x^{(i)}) \rangle \geq 1$$

$$y_i \langle w, \phi(x^{(i)}) \rangle \geq 1$$

soft-margin structured SVM :  $R(w)$

$$\min_{w, \xi} \frac{\lambda \|w\|^2}{2} + \frac{1}{n} \sum_i \xi_i$$

with an exponential # of constraints

$$\xi_i + \langle w, \phi(x^{(i)}, y^{(i)}) \rangle \geq \langle w, \phi(x^{(i)}, \tilde{y}) \rangle + l(y^{(i)}, \tilde{y}) \quad \forall \tilde{y} \in \mathcal{Y}_i, \tilde{y} \neq y^{(i)}$$

equivalent (non-smooth) formulation:

$$\min_w \frac{\lambda \|w\|^2}{2} + \frac{1}{n} \sum_{i=1}^n \mathcal{J}^{\text{SVM}}(x^{(i)}, y^{(i)}; w)$$

loss augmented decoding

where  $\mathcal{J}^{\text{SVM}}(x, y; w) \triangleq \max_{\tilde{y} \in \mathcal{Y}(x)} [\langle w, \phi(x, \tilde{y}) \rangle + l(y, \tilde{y})] - \langle w, \phi(x, y) \rangle$

"structured hinge loss" (suppose that  $y \in \mathcal{Y}(x)$ )