

Maximum Margin Markov Networks and Constraint Generation

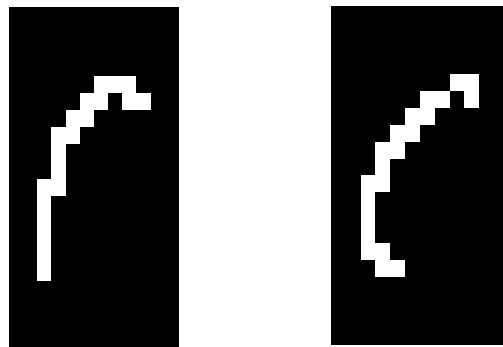
Optimization - 10725
Carlos Guestrin
Carnegie Mellon University

February 11th, 2008

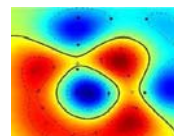
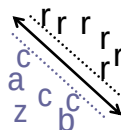
©2008 Carlos Guestrin

1

Handwriting Recognition



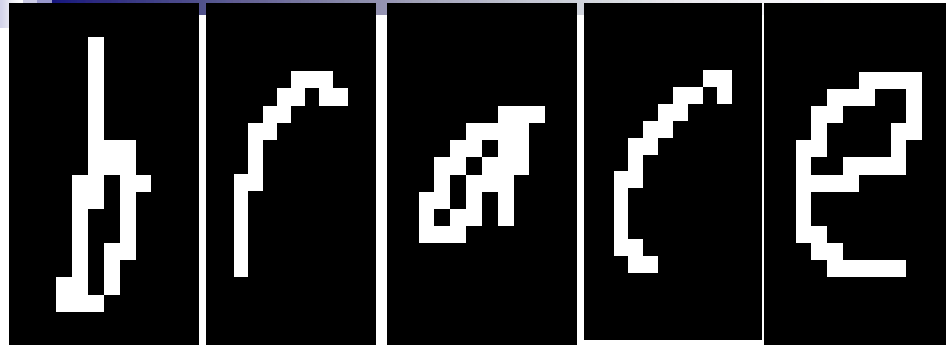
Character recognition: kernel SVMs



©2008 Carlos Guestrin

2

Handwriting Recognition 2



SVMs for sequences?

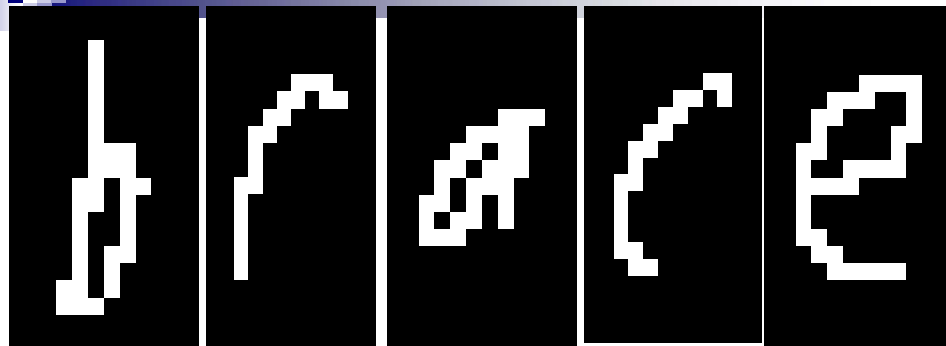
Problem: # of classes exponential in length

brare
 zzzzz brick
 aaaaa
 brake

©2008 Carlos Guestrin

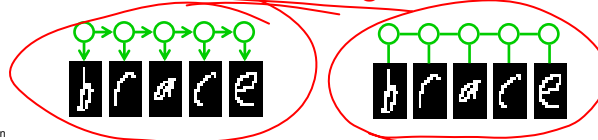
3

Handwriting Recognition 2



Graphical models: HMMs, MNs

Linear in length



©2008 Carlos Guestrin

4

Chain Markov Net (aka CRF*)

$$P(\mathbf{y}|\mathbf{x}) = \frac{1}{Z(\mathbf{x})} \prod_i \phi(\mathbf{x}_i, y_i) \prod_i \phi(y_i, y_{i+1})$$

$$\phi(\mathbf{x}_i, y_i) = \exp\{\sum_{\alpha} w_{\alpha} f_{\alpha}(\mathbf{x}_i, y_i)\}$$

$$\phi(y_i, y_{i+1}) = \exp\{\sum_{\beta} w_{\beta} f_{\beta}(y_i, y_{i+1})\}$$

$$f_{\beta}(y, y') = I(y='z', y'='a')$$

$$f_{\alpha}(\mathbf{x}, y) = I(x_p=1, y='z')$$

©2008 Carlos Guestrin *Lafferty et al. 01 5

CRF - short notation

$$P(\mathbf{y}|\mathbf{x}) = \frac{1}{Z(\mathbf{x})} \prod_i \phi(\mathbf{x}_i, y_i) \prod_i \phi(y_i, y_{i+1}) = \frac{1}{Z(\mathbf{x})} \exp\{\mathbf{w}^T \mathbf{f}(\mathbf{x}, \mathbf{y})\}$$

$$\phi(\mathbf{x}_i, y_i) = \exp\{\sum_{\alpha} w_{\alpha} f_{\alpha}(\mathbf{x}_i, y_i)\}$$

$$\phi(y_i, y_{i+1}) = \exp\{\sum_{\beta} w_{\beta} f_{\beta}(y_i, y_{i+1})\}$$

$$\mathbf{w} = \begin{bmatrix} w_{\alpha_1} \\ \vdots \\ w_{\alpha_n} \\ w_{\beta_1} \\ \vdots \\ w_{\beta_m} \end{bmatrix}$$

$$\mathbf{f}(\mathbf{x}, \mathbf{y}) = \begin{bmatrix} f_{\alpha_1}(\mathbf{x}_1, y_1) \\ \vdots \\ f_{\alpha_n}(\mathbf{x}_n, y_n) \\ f_{\beta_1}(y_1, y_2) \\ \vdots \\ f_{\beta_m}(y_n, y_{n+1}) \end{bmatrix}$$

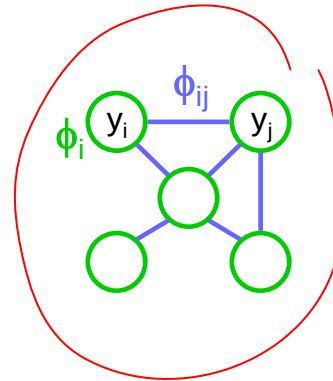
©2008 Carlos Guestrin *Lafferty et al. 01 6

Pairwise Markov Nets

$$P(\mathbf{y}|\mathbf{x}) = \frac{1}{Z(\mathbf{x})} \prod_i \phi_i(\mathbf{x}, y_i) \prod_{ij} \phi_{ij}(\mathbf{x}, y_i, y_j)$$

$$= \frac{1}{Z(\mathbf{x})} \exp\{\mathbf{w}^T \mathbf{f}(\mathbf{x}, \mathbf{y})\}$$

↑
normalized



$$\mathbf{w} = [\dots, \mathbf{w}_i, \dots, \mathbf{w}_{ij}, \dots]$$

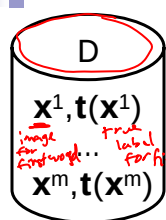
$$\mathbf{f}(\mathbf{x}, \mathbf{y}) = [\dots, \mathbf{f}_i(\mathbf{x}, y_i), \dots, \mathbf{f}_{ij}(\mathbf{x}, y_i, y_j), \dots]$$

©2008 Carlos Guestrin

7

Max (Conditional) Likelihood

$$\arg\max_y f(y) = \arg\max_y \log f(y)$$



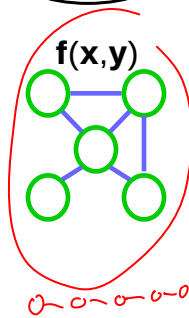
Estimation

$$\underset{\mathbf{w}}{\text{maximize}} \sum_{\mathbf{x} \in D} \log P_{\mathbf{w}}(\mathbf{t}(\mathbf{x}) | \mathbf{x})$$

↑
parameters $\mathbf{w}$

Classification

$$\text{test case } \mathbf{x}_q \rightarrow \arg\max_y P_{\mathbf{w}}(y | \mathbf{x}_q)$$



At classification time

$$\arg\max_y P_{\mathbf{w}}(y | \mathbf{x}) = \arg\max_y \frac{1}{Z(\mathbf{x})} e^{\mathbf{w}^T \mathbf{f}(\mathbf{y}, \mathbf{x})}$$

doesn't depend y
irrelevant for
decision

$$= \arg\max_y \log \frac{1}{Z(\mathbf{x})} e^{\mathbf{w}^T \mathbf{f}(\mathbf{y}, \mathbf{x})} = \arg\max_y \mathbf{w}^T \mathbf{f}(\mathbf{y}, \mathbf{x}) - \log Z(\mathbf{x})$$

$$= \arg\max_y \mathbf{w}^T \mathbf{f}(\mathbf{y}, \mathbf{x})$$

why do I care about $Z(\mathbf{x})$

©2008 Carlos Guestrin

8

OCR Example

- We want:

$$\arg\max_{\text{word}} \mathbf{w}^T \mathbf{f}(\text{brace} \text{word}) = \text{brace}$$

- Equivalently:

$$\mathbf{w}^T \mathbf{f}(\text{brace}, \text{brace}) > \mathbf{w}^T \mathbf{f}(\text{brace}, \text{aaaaa})$$

$$\mathbf{w}^T \mathbf{f}(\text{brace}, \text{brace}) > \mathbf{w}^T \mathbf{f}(\text{brace}, \text{aaaab})$$

...

$$\mathbf{w}^T \mathbf{f}(\text{brace}, \text{brace}) > \mathbf{w}^T \mathbf{f}(\text{brace}, \text{zzzzz})$$

strictly greater than?
no w???

number of constraints is exponential in length of word

©2008 Carlos Guestrin

9

Max Margin Estimation

- Goal: find $\mathbf{w}$ such that

$$\mathbf{w}^T \mathbf{f}(\mathbf{x}, \mathbf{t}(\mathbf{x})) > \mathbf{w}^T \mathbf{f}(\mathbf{x}, \mathbf{y}) \quad \forall \mathbf{x} \in D \quad \forall \mathbf{y} \neq \mathbf{t}(\mathbf{x})$$

$$\mathbf{w}^T [\mathbf{f}(\mathbf{x}, \mathbf{t}(\mathbf{x})) - \mathbf{f}(\mathbf{x}, \mathbf{y})] > 0$$

$$\mathbf{w}^T \Delta \mathbf{f}_x(\mathbf{y}) > 0$$

©2008 Carlos Guestrin

10

Not all margins are equal

- Goal: find w such that

$$w^T \Delta f_x(y) \geq \gamma \quad \forall x \in D \quad \forall y \neq t(x)$$

- Gain over y grows with # of mistakes in y : $\Delta t_x(y)$

$\Delta t_{\text{brace}}(\text{"craze"})$

$\Delta t_{\text{brace}}(\text{"zzzzz"})$

$w^T \Delta f_{\text{brace}}(\text{"craze"})$

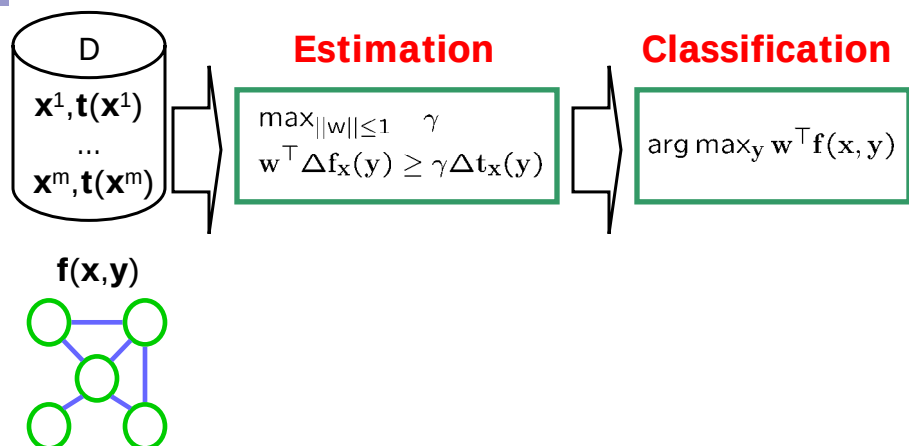
$w^T \Delta f_{\text{brace}}(\text{"zzzzz"})$

©2008 Carlos Guestrin

11

Maximum Margin Markov Nets

[Taskar, Guestrin, Koller '03]



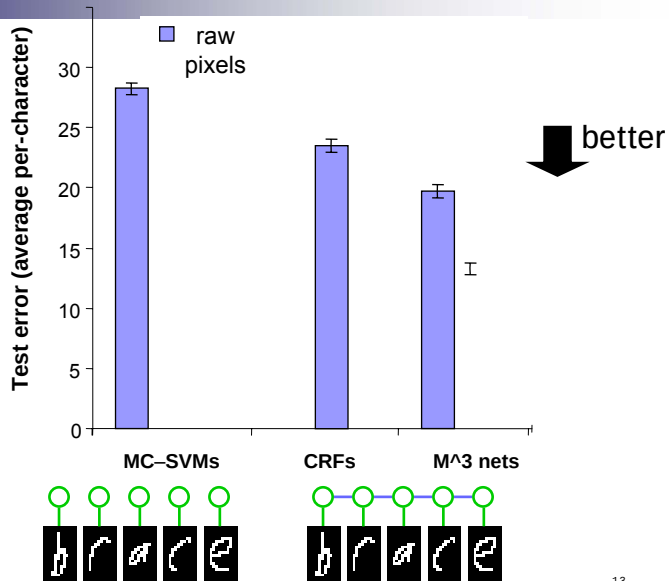
©2008 Carlos Guestrin

12

Handwriting Recognition

Length: ~8 chars
Letter: 16x8 pixels
10-fold Train/Test
5000/50000 letters
600/6000 words

Models:
Multiclass-SVMs*
CRFs
M³ nets

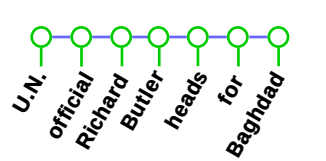


*Crammer & Singer 01

13

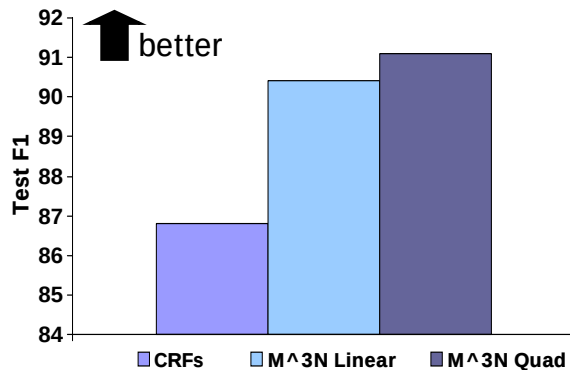
Named Entity Recognition

- Locate and classify named entities in sentences:
 - 4 categories: organization, person, location, misc.
 - e.g. "U.N. official Richard Butler heads for Baghdad".
- CoNLL 03 data set (200K words train, 50K words test)



$y_i = \text{org/per/loc/misc/none}$

$f(y_i, x) = [\dots,$
 $I(y_i=\text{org}, x_i=\text{"U.N."}),$
 $I(y_i=\text{per}, x_i=\text{capitalized}),$
 $I(y_i=\text{loc}, x_i=\text{known city}),$



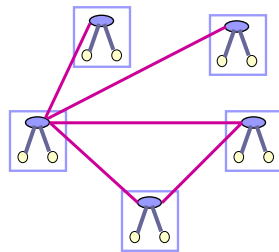
*2008 Carlos Guestrin

14

Hypertext Classification

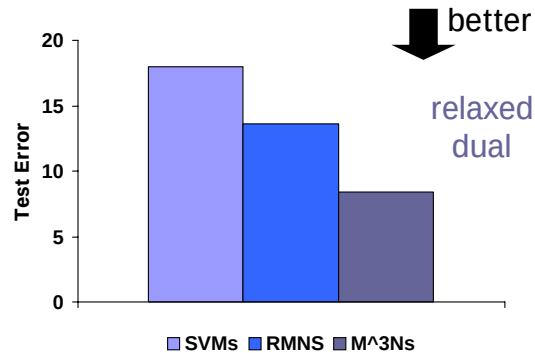
WebKB dataset

- Four CS department websites: 1300 pages/3500 links
- Classify each page: faculty, course, student, project, other
- Train on three universities/test on fourth



loopy belief propagation

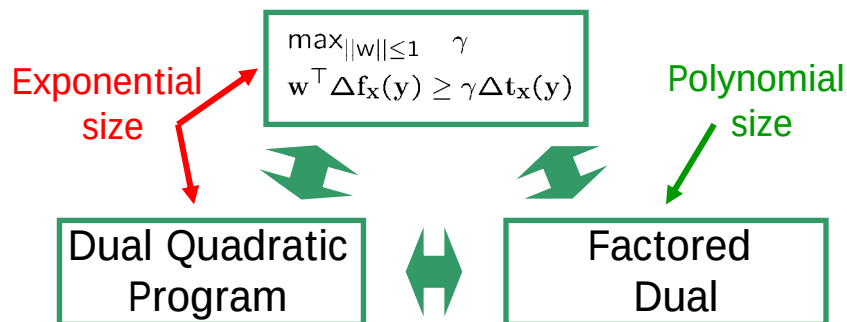
*Taskar et al 02



15

Solving M³Ns

Estimation



©2008 Carlos Guestrin

16