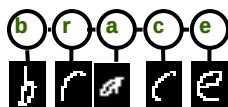


Advanced Machine Learning

Learning Graphical Models Max-margin learning of GM

Eric Xing

Lecture 16, August 13, 2009



Eric Xing

© Eric Xing @ CMU, 2006-2009

Reading:

1

Structured Prediction Problem

- Unstructured prediction



$$\mathbf{x} = \begin{pmatrix} x_{11} & x_{12} & \dots \\ x_{21} & x_{22} & \dots \\ \vdots & \vdots & \dots \end{pmatrix}$$

$$\mathbf{y} = \begin{pmatrix} y_1 \\ y_2 \\ \vdots \end{pmatrix}$$

- Structured prediction

- Part of speech tagging

$\mathbf{X}$ = "Do you want sugar in it?" $\Rightarrow$ $\mathbf{Y}$ = verb pron verb noun prep pron

- Image segmentation



$$\mathbf{x} = \begin{pmatrix} x_{11} & x_{12} & \dots \\ x_{21} & x_{22} & \dots \\ \vdots & \vdots & \dots \end{pmatrix}$$

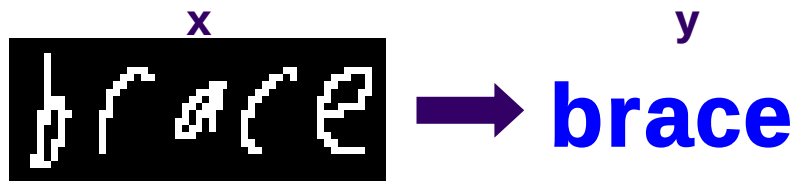
$$\mathbf{y} = \begin{pmatrix} y_{11} & y_{12} & \dots \\ y_{21} & y_{22} & \dots \\ \vdots & \vdots & \dots \end{pmatrix}$$

Eric Xing

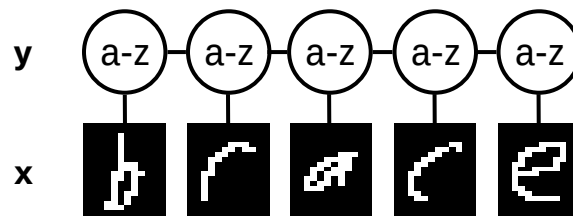
© Eric Xing @ CMU, 2006-2009

2

OCR example



Sequential structure

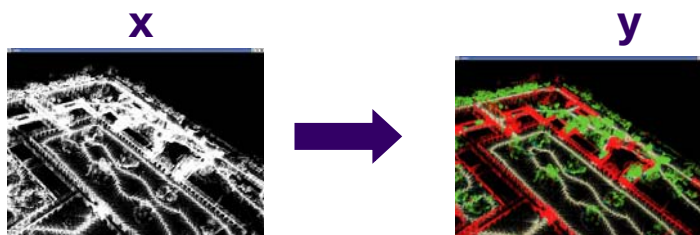


Eric Xing

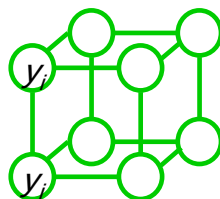
© Eric Xing @ CMU, 2006-2009

3

Image segmentation example



Spatial structure



Eric Xing

© Eric Xing @ CMU, 2006-2009

4

Classical Predictive Models



- Inputs:
 - a set of training samples $\mathcal{D} = \{(x^i, y^i)\}_{i=1}^N$ and $y^i \in C \triangleq \{c_1, c_2, \dots, c_L\}$
 - $x^i = [x_1^i, x_2^i, \dots, x_d^i]^\top$
- Outputs:
 - a predictive function $h(x) \quad y^* = h(x) \triangleq \arg \max_y F(x, y; \mathbf{w})$
- Examples: $F(x, y; \mathbf{w}) = g(\mathbf{w}^\top \mathbf{f}(x, y))$

Logistic Regression, Bayes classifiers

- Max-likelihood estimation

$$\text{E.g.: } \max_{\mathbf{w}} \mathcal{L}(\mathcal{D}; \mathbf{w}) \triangleq \sum_{i=1}^N \log p(y^i | x^i)$$

$$p(y|x) = \frac{\exp\{\mathbf{w}^\top \mathbf{f}(x, y)\}}{\sum_{y'} \exp\{\mathbf{w}^\top \mathbf{f}(x, y')\}}$$

Advantages:

1. Full probabilistic semantics
2. Straightforward Bayesian or direct regularization
3. Hidden structures or generative hierarchy

Support Vector Machines (SVM)

- Max-margin learning

$$\min_{\mathbf{w}, \xi} \frac{1}{2} \mathbf{w}^\top \mathbf{w} + C \sum_{i=1}^N \xi_i;$$

$$\text{s.t. } \mathbf{w}^\top \Delta \mathbf{f}_i(y) \geq 1 - \xi_i, \quad \forall i, \forall y \neq y^i.$$

Advantages:

1. Dual sparsity: few support vectors
2. Kernel tricks
3. Strong empirical results

Eric Xing

© Eric Xing @ CMU, 2006-2009

5

Structured Prediction Models



Conditional Random Fields (CRFs) (Lafferty et al 2001)

- Based on Logistic Regression
- Max-likelihood estimation (point-estimate)

$$\max_{\mathbf{w}} \mathcal{L}(\mathcal{D}; \mathbf{w}) \triangleq \sum_{i=1}^N \log p(\mathbf{y}^i | \mathbf{x}^i)$$

$$p(\mathbf{y} | \mathbf{x}) = \frac{\exp\{\mathbf{w}^\top \mathbf{f}(\mathbf{x}, \mathbf{y})\}}{\sum_{\mathbf{y}'} \exp\{\mathbf{w}^\top \mathbf{f}(\mathbf{x}, \mathbf{y}')\}}$$

Max-margin Markov Networks (M³NS) (Taskar et al 2003)

- Based on SVM
- Max-margin learning (point-estimate)

$$P0(M^3N) : \min_{\mathbf{w}, \xi} \frac{1}{2} \|\mathbf{w}\|^2 + C \sum_{i=1}^N \xi_i$$

$$\text{s.t. } \forall i, \forall \mathbf{y} \neq \mathbf{y}^i : \mathbf{w}^\top \Delta \mathbf{f}_i(\mathbf{y}) \geq \Delta \ell_i(\mathbf{y}) - \xi_i, \quad \xi_i \geq 0,$$

where $\mathbf{w}^\top \Delta \mathbf{f}_i(\mathbf{y} | \mathbf{x}_i)$ denotes the margin and $\Delta \ell_i(\mathbf{y})$ is a loss function.

Eric Xing

© Eric Xing @ CMU, 2006-2009

6

Structured models



$$h(\mathbf{x}) = \arg \max_{\mathbf{y} \in \mathcal{Y}(\mathbf{x})} s(\mathbf{x}, \mathbf{y})$$

↑
space of feasible outputs

← scoring function

Assumptions:

$$\text{score}(\mathbf{x}, \mathbf{y}) = \mathbf{w}^\top \mathbf{f}(\mathbf{x}, \mathbf{y}) = \sum_p \mathbf{w}^\top \mathbf{f}(\mathbf{x}_p, \mathbf{y}_p)$$

linear combination of features

sum of part scores:

- index p represents a part in the structure

Eric Xing

© Eric Xing @ CMU, 2006-2009

7

Learning $\mathbf{w}$



- Training examples $(\mathbf{x}_i, \mathbf{y}_i)$

- Probabilistic approach:

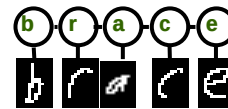
$$P_{\mathbf{w}}(\mathbf{y} \mid \mathbf{x}) = \frac{1}{Z_{\mathbf{w}}(\mathbf{x})} \exp\{\mathbf{w}^\top \mathbf{f}(\mathbf{x}, \mathbf{y})\}$$

- Computing $Z_{\mathbf{w}}(\mathbf{x})$ can be NP-complete

- Tractable models but intractable estimation

- Large margin approach:

- Exact and efficient when prediction is tractable



Eric Xing

© Eric Xing @ CMU, 2006-2009

8

Learning Strategy



- Recall that in CRF

- We predict based on:

$$y^* | x = \arg \max_y p_\theta(y | x) = \frac{1}{Z(\theta, x)} \exp \left\{ \sum_c \theta_c f_c(x, y_c) \right\}$$

- And we learn based on:

$$\theta_c^* | \{y_n, x_n\} = \arg \max_{\theta_c} \prod_n p_\theta(y_n | x_n) = \prod_n \frac{1}{Z(\theta, x_n)} \exp \left\{ \sum_c \theta_c f_c(x_n, y_{n,c}) \right\}$$

- Max-Margin:

- We predict based on:

$$y^* | x = \arg \max_y \sum_c \theta_c f_c(x, y_c) = \arg \max_y w^T F(x, y)$$

- And we learn based on:

$$w^* | \{y_n, x_n\} = \arg \max_w \left(\max_{y'_n \neq y_n, \forall n} w^T (F(y_n, x_n) - F(y'_n, x_n)) \right)$$

Eric Xing

© Eric Xing @ CMU, 2006-2009

9

MLE of Feature Based UGMs



- Scaled likelihood function

$$\begin{aligned} \tilde{\ell}(\theta; \mathcal{D}) &= \ell(\theta; \mathcal{D}) / N = \frac{1}{N} \sum \log p(x_n | \theta) \\ &= \sum_x \tilde{p}(x) \log p(x | \theta) \\ &= \sum_x \tilde{p}(x) \sum_i \theta_i f_i(x) - \log Z(\theta) \end{aligned}$$

- Instead of optimizing this objective directly, we attack its lower bound

- The logarithm has a linear upper bound ...

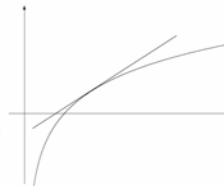
$$\log Z(\theta) \leq \mu^T Z(\theta) - \log \mu - 1$$

- This bound holds for all μ , in particular, for

$$\mu = Z^{-1}(\theta^{(r)})$$

- Thus we have

$$\tilde{\ell}(\theta; \mathcal{D}) \geq \sum_x \tilde{p}(x) \sum_i \theta_i f_i(x) - \frac{Z(\theta)}{Z(\theta^{(r)})} - \log Z(\theta^{(r)}) + 1$$



Eric Xing

© Eric Xing @ CMU, 2006-2009

10

Generalized Iterative Scaling (GIS)



- Lower bound of scaled loglikelihood

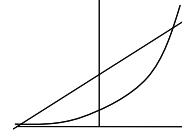
$$\tilde{\ell}(\theta; D) \geq \sum_{\mathbf{x}} \tilde{p}(\mathbf{x}) \sum_i \theta_i f_i(\mathbf{x}) - \frac{Z(\theta)}{Z(\theta^{(t)})} - \log Z(\theta^{(t)}) + 1$$

- Define $\Delta\theta_i^{(t)} \stackrel{\text{def}}{=} \theta_i - \theta_i^{(t)}$

$$\begin{aligned} \tilde{\ell}(\theta; D) &\geq \sum_{\mathbf{x}} \tilde{p}(\mathbf{x}) \sum_i \theta_i f_i(\mathbf{x}) - \frac{1}{Z(\theta^{(t)})} \sum_{\mathbf{x}} \exp\left\{\sum_i \theta_i f_i(\mathbf{x})\right\} - \log Z(\theta^{(t)}) + 1 \\ &= \sum_i \theta_i \sum_{\mathbf{x}} \tilde{p}(\mathbf{x}) f_i(\mathbf{x}) - \frac{1}{Z(\theta^{(t)})} \sum_{\mathbf{x}} \exp\left\{\sum_i \theta_i^{(t)} f_i(\mathbf{x})\right\} \exp\left\{\sum_i \Delta\theta_i^{(t)} f_i(\mathbf{x})\right\} - \log Z(\theta^{(t)}) + 1 \\ &= \sum_i \theta_i \sum_{\mathbf{x}} \tilde{p}(\mathbf{x}) f_i(\mathbf{x}) - \sum_{\mathbf{x}} p(\mathbf{x} | \theta^{(t)}) \exp\left\{\sum_i \Delta\theta_i^{(t)} f_i(\mathbf{x})\right\} - \log Z(\theta^{(t)}) + 1 \end{aligned}$$

- Relax again

- Assume $f_i(\mathbf{x}) \geq 0$, $\sum_i f_i(\mathbf{x}) = 1$
- Convexity of exponential: $\exp\left(\sum_i \pi_i x_i\right) \leq \sum_i \pi_i \exp(x_i)$



- We have:

$$\tilde{\ell}(\theta; D) \geq \sum_i \theta_i \sum_{\mathbf{x}} \tilde{p}(\mathbf{x}) f_i(\mathbf{x}) - \sum_{\mathbf{x}} p(\mathbf{x} | \theta^{(t)}) \sum_i f_i(\mathbf{x}) \exp(\Delta\theta_i^{(t)}) - \log Z(\theta^{(t)}) + 1 \stackrel{\text{def}}{=} \Lambda(\theta)$$

Eric Xing

© Eric Xing @ CMU, 2006-2009

11

GIS



- Lower bound of scaled loglikelihood

$$\tilde{\ell}(\theta; D) \geq \sum_i \theta_i \sum_{\mathbf{x}} \tilde{p}(\mathbf{x}) f_i(\mathbf{x}) - \sum_{\mathbf{x}} p(\mathbf{x} | \theta^{(t)}) \sum_i f_i(\mathbf{x}) \exp(\Delta\theta_i^{(t)}) - \log Z(\theta^{(t)}) + 1 \stackrel{\text{def}}{=} \Lambda(\theta)$$

- Take derivative: $\frac{\partial \Lambda}{\partial \theta_i} = \sum_{\mathbf{x}} \tilde{p}(\mathbf{x}) f_i(\mathbf{x}) - \exp(\Delta\theta_i^{(t)}) \sum_{\mathbf{x}} p(\mathbf{x} | \theta^{(t)}) f_i(\mathbf{x})$

- Set to zero

$$e^{\Delta\theta_i^{(t)}} = \frac{\sum_{\mathbf{x}} \tilde{p}(\mathbf{x}) f_i(\mathbf{x})}{\sum_{\mathbf{x}} p(\mathbf{x} | \theta^{(t)}) f_i(\mathbf{x})} = \frac{\sum_{\mathbf{x}} \tilde{p}(\mathbf{x}) f_i(\mathbf{x})}{\sum_{\mathbf{x}} p^{(t)}(\mathbf{x}) f_i(\mathbf{x})} Z(\theta^{(t)})$$

- where $p^{(0)}(\mathbf{x})$ is the unnormalized version of $p(\mathbf{x} | \theta^{(0)})$

- Update

$$\theta_i^{(t+1)} = \theta_i^{(t)} + \Delta\theta_i^{(t)} \Rightarrow p^{(t+1)}(\mathbf{x}) = p^{(t)}(\mathbf{x}) e^{\Delta\theta_i^{(t)} f_i(\mathbf{x})}$$

$$\begin{aligned} p^{(t+1)}(\mathbf{x}) &= \frac{p^{(t)}(\mathbf{x})}{Z(\theta^{(t)})} \prod_i \left(\frac{\sum_{\mathbf{x}} \tilde{p}(\mathbf{x}) f_i(\mathbf{x})}{\sum_{\mathbf{x}} p^{(t)}(\mathbf{x}) f_i(\mathbf{x})} Z(\theta^{(t)}) \right)^{f_i(\mathbf{x})} \\ &\Rightarrow \frac{p^{(t)}(\mathbf{x})}{Z(\theta^{(t)})} \prod_i \left(\frac{\sum_{\mathbf{x}} \tilde{p}(\mathbf{x}) f_i(\mathbf{x})}{\sum_{\mathbf{x}} p^{(t)}(\mathbf{x}) f_i(\mathbf{x})} \right)^{f_i(\mathbf{x})} (Z(\theta^{(t)}))^{\sum_i f_i(\mathbf{x})} \\ &= p^{(t)}(\mathbf{x}) \prod_i \left(\frac{\sum_{\mathbf{x}} \tilde{p}(\mathbf{x}) f_i(\mathbf{x})}{\sum_{\mathbf{x}} p^{(t)}(\mathbf{x}) f_i(\mathbf{x})} \right)^{f_i(\mathbf{x})} \end{aligned}$$

Eric Xing

© Eric Xing @ CMU, 2006-2009

12

Learning w

- Training examples $(\mathbf{x}_i, y_i)$

- Probabilistic approach:

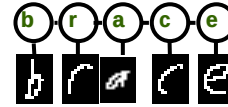
$$P_{\mathbf{w}}(\mathbf{y} \mid \mathbf{x}) = \frac{1}{Z_{\mathbf{w}}(\mathbf{x})} \exp\{\mathbf{w}^T \mathbf{f}(\mathbf{x}, \mathbf{y})\}$$

- Computing $Z_{\mathbf{w}}(\mathbf{x})$ can be NP-complete

- Tractable models but intractable estimation

- Large margin approach:

- Exact and efficient when prediction is tractable



OCR Example

- We want:

$$\operatorname{argmax}_{\text{word}} \mathbf{w}^T \mathbf{f}(\text{image}, \text{word}) = \text{"brace"}$$

- Equivalently:

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

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

...

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

} a lot!

Large Margin Estimation



- Given training example $(\mathbf{x}, \mathbf{y}^*)$, we want:

$$\arg \max_{\mathbf{y}} \mathbf{w}^\top \mathbf{f}(\mathbf{x}, \mathbf{y}) = \mathbf{y}^*$$

$$\mathbf{w}^\top \mathbf{f}(\mathbf{x}, \mathbf{y}^*) > \mathbf{w}^\top \mathbf{f}(\mathbf{x}, \mathbf{y}) \quad \forall \mathbf{y} \neq \mathbf{y}^*$$

$$\mathbf{w}^\top \mathbf{f}(\mathbf{x}, \mathbf{y}^*) \geq \mathbf{w}^\top \mathbf{f}(\mathbf{x}, \mathbf{y}) + \gamma \ell(\mathbf{y}^*, \mathbf{y}) \quad \forall \mathbf{y}$$

- Maximize margin γ
- Mistake weighted margin: $\gamma \ell(\mathbf{y}^*, \mathbf{y})$

$$\ell(\mathbf{y}^*, \mathbf{y}) = \sum_i I(y_i^* \neq y_i) \quad \# \text{ of mistakes in } \mathbf{y}$$

*Taskar et al. 03

Eric Xing

© Eric Xing @ CMU, 2006-2009

15

Large Margin Estimation



- Recall from SVMs:
 - Maximizing margin γ is equivalent to minimizing the square of the L2-norm of the weight vector $\mathbf{w}$:
- New objective function:

$$\begin{aligned} \min_{\mathbf{w}} \quad & \frac{1}{2} \|\mathbf{w}\|^2 \\ \text{s.t.} \quad & \mathbf{w}^\top \mathbf{f}(\mathbf{x}_i, \mathbf{y}_i) \geq \mathbf{w}^\top \mathbf{f}(\mathbf{x}_i, \mathbf{y}'_i) + \ell(\mathbf{y}_i, \mathbf{y}'_i), \quad \forall i, \mathbf{y}'_i \in \mathcal{Y}_i \end{aligned}$$

Eric Xing

© Eric Xing @ CMU, 2006-2009

16

Min-max Formulation

- Brute force enumeration of constraints:

$$\min \frac{1}{2} \|w\|^2$$

$$w^\top f(x, y^*) \geq w^\top f(x, y) + \ell(y^*, y), \quad \forall y$$

- The constraints are exponential in the size of the structure

- Alternative: min-max formulation

- add only the most violated constraint

$$y' = \arg \max_{y \neq y^*} [w^\top f(x^i, y) + \ell(y^i, y)]$$

$$\text{add to QP: } w^\top f(x^i, y^i) \geq w^\top f(x^i, y') + \ell(y^i, y')$$

- Handles more general loss functions
- Only polynomial # of constraints needed

Min-max formulation

$$\min \frac{1}{2} \|w\|^2$$

$$w^\top f(x, y^*) \geq \max_{y \neq y^*} w^\top f(x, y) + \ell(y^*, y)$$

- Key step: convert the maximization in the constraint from discrete to continuous

- This enables us to plug it into a QP

$$\max_{y \neq y^*} w^\top f(x, y) + \ell(y^*, y) \iff \max_{z \in \mathcal{Z}} (F^\top w + \ell)^\top z$$

discrete optim.

continuous optim.

- How to do this conversion?

- Linear chain example in the next slides →

$y \Rightarrow z$ map for linear chain structures

OCR example: $y = 'ABABB'$;

z 's are the indicator variables for the corresponding classes (alphabet)

	$z_1(m)$	$z_2(m)$	$z_3(m)$	$z_4(m)$	$z_5(m)$
A	1	0	1	0	0
B	0	1	0	1	1
.	.	.	.	.	.
B	0	0	0	0	0

	$z_{12}(m, n)$	$z_{23}(m, n)$	$z_{34}(m, n)$	$z_{45}(m, n)$
A	0 1 . 0	0 0 . 0	0 1 . 0	0 0 . 0
B	0 0 . 0	1 0 . 0	0 0 . 0	0 1 . 0
.	. . . 0	. . . 0	. . . 0	. . . 0
B	0 0 0 0	0 0 0 0	0 0 0 0	0 0 0 0

A	B	.	B	A	B	.	B	A	B	.	B	A	B	.	B
---	---	---	---	---	---	---	---	---	---	---	---	---	---	---	---

Eric Xing

© Eric Xing @ CMU, 2006-2009

19

$y \Rightarrow z$ map for linear chain structures

Rewriting the maximization function in terms of indicator variables:

$$\begin{aligned}
 \max_z \quad & \sum_{j,m} z_j(m) [\mathbf{w}^\top \mathbf{f}_{\text{node}}(\mathbf{x}_j, m) + \ell_j(m)] \\
 & + \sum_{jk,m,n} z_{jk}(m, n) [\mathbf{w}^\top \mathbf{f}_{\text{edge}}(\mathbf{x}_{jk}, m, n) + \ell_{jk}(m, n)] \quad \left. \vphantom{\sum_{jk,m,n}} \right\} (\mathbf{F}^\top \mathbf{w} + \ell)^\top \mathbf{z} \\
 & z_k(n) \quad z_j(m) \geq 0; \quad z_{jk}(m, n) \geq 0; \\
 & \left. \begin{aligned} & \text{normalization} \quad \sum_m z_j(m) = 1 \\ & \text{agreement} \quad \sum_n z_{jk}(m, n) = z_j(m) \end{aligned} \right\} \mathbf{Az} = \mathbf{b} \\
 \max_{\mathbf{Az}=\mathbf{b}} \quad & (\mathbf{F}^\top \mathbf{w} + \ell)^\top \mathbf{z}
 \end{aligned}$$

Eric Xing

© Eric Xing @ CMU, 2006-2009

20

Min-max formulation



- Original problem:

$$\min \frac{1}{2} \|\mathbf{w}\|^2$$

$$\mathbf{w}^\top \mathbf{f}(\mathbf{x}, \mathbf{y}^*) \geq \max_{\mathbf{y}} \mathbf{w}^\top \mathbf{f}(\mathbf{x}, \mathbf{y}) + \ell(\mathbf{y}^*, \mathbf{y})$$

- Transformed problem:

$$\min \frac{1}{2} \|\mathbf{w}\|^2$$

$$\mathbf{w}^\top \mathbf{f}(\mathbf{x}, \mathbf{y}^*) \geq \max_{\substack{\mathbf{z} \geq 0; \\ \mathbf{A}\mathbf{z} = \mathbf{b};}} \mathbf{q}^\top \mathbf{z} \quad \text{where } \mathbf{q}^\top = \mathbf{w}^\top \mathbf{F} + \ell^\top$$

- Has integral solutions $\mathbf{z}$ for chains, trees
- Can be fractional for untriangulated networks

Min-max formulation



- Using strong Lagrangian duality:
(beyond the scope of this lecture)

$$\max_{\substack{\mathbf{z} \geq 0; \\ \mathbf{A}\mathbf{z} = \mathbf{b};}} \mathbf{q}^\top \mathbf{z} = \min_{\mathbf{A}^\top \boldsymbol{\mu} \geq \mathbf{q}} \mathbf{b}^\top \boldsymbol{\mu}$$

- Use the result above to minimize jointly over $\mathbf{w}$ and $\boldsymbol{\mu}$:

$$\min_{\mathbf{w}, \boldsymbol{\mu}} \frac{1}{2} \|\mathbf{w}\|^2$$

$$\text{s.t. } \mathbf{w}^\top \mathbf{f}(\mathbf{x}, \mathbf{y}^*) \geq \mathbf{b}^\top \boldsymbol{\mu};$$

$$\mathbf{A}^\top \boldsymbol{\mu} \geq \mathbf{q};$$

Min-max formulation

$$\begin{aligned} \min_{\mathbf{w}, \mu} \quad & \frac{1}{2} \|\mathbf{w}\|^2 \\ \text{s.t.} \quad & \mathbf{w}^\top \mathbf{f}(\mathbf{x}, \mathbf{y}^*) \geq \mathbf{b}^\top \mu; \\ & \mathbf{A}^\top \mu \geq (\mathbf{w}^\top \mathbf{F} + \ell)^\top \end{aligned}$$

- Formulation produces compact QP for
 - Low-treewidth Markov networks
 - Associative Markov networks
 - Context free grammars
 - Bipartite matchings
 - Any problem with compact LP inference

Eric Xing

© Eric Xing @ CMU, 2006-2009

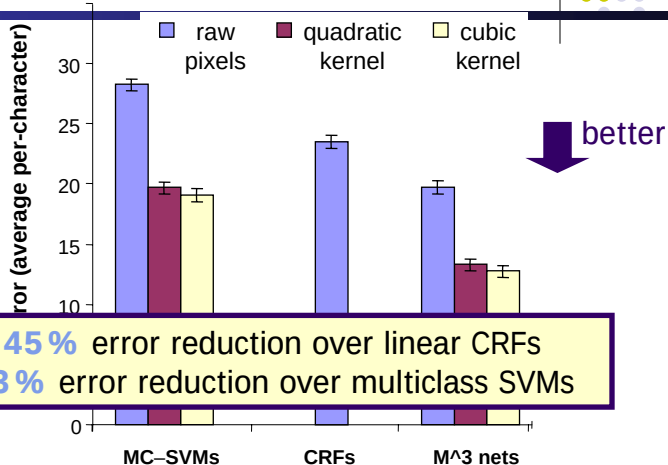
23

Results: 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

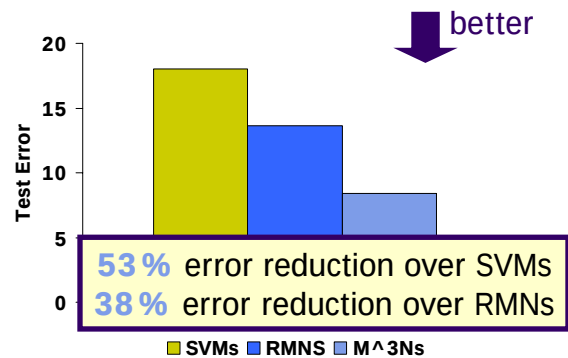
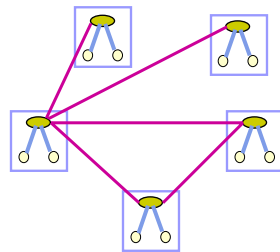
© Eric Xing @ CMU, 2006-2009

24

Results: 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



*Taskar et al 02

© Eric Xing @ CMU, 2006-2009

25

Summary

- Structured Maximum Margin Networks

- Concise representation
- Efficient QP algorithms
- Applications:
 - Sequences
 - Trees
 - Matchings
- Strong empirical results

- Acknowledgments:

- Adopted from two different presentations given by Ben Taskar.

Eric Xing

© Eric Xing @ CMU, 2006-2009

26

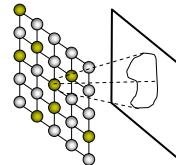
Open Problems



- Unsupervised CRF learning and MaxMargin Learning

- Only X , but not Y (sometimes part of Y), is available

- We want to recognize a pattern that is maximally different from the rest!



- What does margin or conditional likelihood mean in these cases?
Given only $\{X_n\}$, how can we define the cost function?

$$\text{margin} = w^T (F(y_n, x_n) - F(y'_n, x_n))$$

$$p_\theta(y | x) = \frac{1}{Z(\theta, x)} \exp \left\{ \sum_c \theta_c f_c(x, y_c) \right\}$$

- Algorithmic challenge
- Stay tuned for lecture 19!