

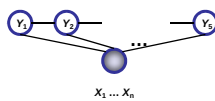
Machine Learning

10-701/15-781, Fall 2011

Conditional Random Fields

Eric Xing

Lecture 12, October 19, 2011



Reading: Chap. 13 CB

© Eric Xing @ CMU, 2006-2011

1

Definition (of HMM)

- **Observation space**

Alphabetic set: $C = \{c_1, c_2, \dots, c_K\}$
Euclidean space: $\mathbb{R}^d$

- **Index set of hidden states**

$I = \{1, 2, \dots, M\}$

- **Transition probabilities between any two states**

$$p(y_t^j = 1 | y_{t-1}^i = 1) = a_{i,j},$$

or $p(y_t | y_{t-1}^i = 1) \sim \text{Multinomial}(a_{i,1}, a_{i,2}, \dots, a_{i,M}), \forall i \in I.$

- **Start probabilities**

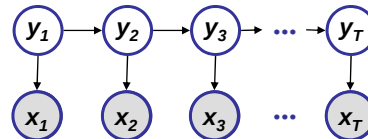
$$p(y_1) \sim \text{Multinomial}(\pi_1, \pi_2, \dots, \pi_M).$$

- **Emission probabilities associated with each state**

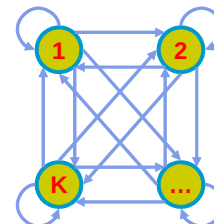
$$p(x_t | y_t^i = 1) \sim \text{Multinomial}(b_{i,1}, b_{i,2}, \dots, b_{i,K}), \forall i \in I.$$

or in general:

$$p(x_t | y_t^i = 1) \sim f(\cdot | \theta_i), \forall i \in I.$$



Graphical model



State automata

© Eric Xing @ CMU, 2006-2011

2

Viterbi decoding

- GIVEN $\mathbf{x} = x_1, \dots, x_T$, we want to find $\mathbf{y} = y_1, \dots, y_T$, such that $P(\mathbf{y}|\mathbf{x})$ is maximized:

$$\mathbf{y}^* = \operatorname{argmax}_{\mathbf{y}} P(\mathbf{y}|\mathbf{x}) = \operatorname{argmax}_{\pi} P(\mathbf{y}, \mathbf{x})$$

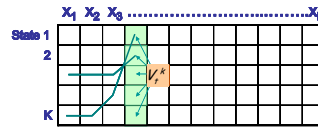
- Let

$$V_t^k = \max_{\{y_1, \dots, y_{t-1}\}} P(x_1, \dots, x_{t-1}, y_1, \dots, y_{t-1}, x_t, y_t^k = 1)$$

= Probability of most likely sequence of states ending at state $y_t = k$

- The recursion:

$$V_t^k = p(x_t | y_t^k = 1) \max_i a_{i,k} V_{t-1}^i$$



- Underflows are a significant problem

$$p(x_1, \dots, x_t, y_1, \dots, y_t) = \pi_{y_1} a_{y_1, y_2} \cdots a_{y_{t-1}, y_t} b_{y_1, x_1} \cdots b_{y_t, x_t}$$

- These numbers become extremely small – underflow
- Solution: Take the logs of all values: $V_t^k = \log p(x_t | y_t^k = 1) + \max_i (\log(a_{i,k}) + V_{t-1}^i)$

The Viterbi Algorithm – derivation

- Define the viterbi probability:

$$\begin{aligned} V_{t+1}^k &= \max_{\{y_1, \dots, y_t\}} P(x_1, \dots, x_t, y_1, \dots, y_t, x_{t+1}, y_{t+1}^k = 1) \\ &= \max_{\{y_1, \dots, y_t\}} P(x_{t+1}, y_{t+1}^k = 1 | x_1, \dots, x_t, y_1, \dots, y_t) P(x_1, \dots, x_t, y_1, \dots, y_t) \\ &= \max_{\{y_1, \dots, y_t\}} P(x_{t+1}, y_{t+1}^k = 1 | y_t) P(x_1, \dots, x_{t-1}, y_1, \dots, y_{t-1}, x_t, y_t) \\ &= \max_i P(x_{t+1}, y_{t+1}^k = 1 | y_t^i = 1) \max_{\{y_1, \dots, y_{t-1}\}} P(x_1, \dots, x_{t-1}, y_1, \dots, y_{t-1}, x_t, y_t^i = 1) \\ &= \max_i P(x_{t+1}, | y_{t+1}^k = 1) a_{i,k} V_t^i \\ &= P(x_{t+1}, | y_{t+1}^k = 1) \max_i a_{i,k} V_t^i \end{aligned}$$

The Viterbi Algorithm

- Input: $\mathbf{x} = x_1, \dots, x_T$

Initialization:

$$V_1^k = P(x_1 | y_1^k = 1) \pi_k$$

Iteration:

$$V_t^k = P(x_t | y_t^k = 1) \max_i a_{i,k} V_{t-1}^i$$

$$\text{Ptr}(k, t) = \arg \max_i a_{i,k} V_{t-1}^i$$

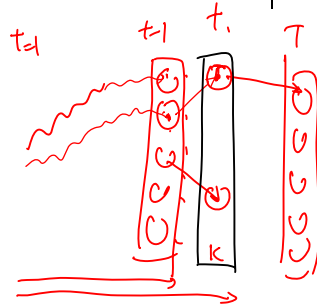
Termination:

$$P(\mathbf{x}, \mathbf{y}^*) = \max_k V_T^k$$

TraceBack:

$$y_T^* = \arg \max_k V_T^k$$

$$y_{t-1}^* = \text{Ptr}(y_t^*, t)$$



© Eric Xing @ CMU, 2006-2011

5

$\mathbf{x} = 1, 2, 1, 5, 6, 2, 1, 6, 2, 4$

V_t^k

0.6250	0.3750
0.7353	0.2647
0.8224	0.1776
0.8853	0.1147
0.7200	0.2800
0.8108	0.1892
0.8772	0.1228
0.7043	0.2957
0.7988	0.2012
0.8687	0.1313



$P(1|F) = 1/6$
 $P(2|F) = 1/6$
 $P(3|F) = 1/6$
 $P(4|F) = 1/6$
 $P(5|F) = 1/6$
 $P(6|F) = 1/6$

$P(1|L) = 1/10$
 $P(2|L) = 1/10$
 $P(3|L) = 1/10$
 $P(4|L) = 1/10$
 $P(5|L) = 1/10$
 $P(6|L) = 1/2$

$$\alpha_t^k = P(x_t | y_t^k = 1) \sum_i \alpha_{t-1}^i a_{i,k}$$

$$\beta_t^k = \sum_i a_{k,i} P(x_{t+1} | y_{t+1}^i = 1) \beta_{t+1}^i$$

$$V_t^k = p(x_t | y_t^k = 1) \max_i a_{i,k} V_{t-1}^i$$

© Eric Xing @ CMU, 2006-2011

6

Viterbi Vs. Posterior Decoding (individual)



$x = 1, 2, 1, 5, 6, 2, 1, 6, 2, 4$

V_t^k	$ptr(k, t)$	Veterbi	PD	$p(y_t^k = 1 x)$
0.6250 0.3750	N/A	1	1	0.8128 0.1872
0.7353 0.2647	1 2	1	1	0.8238 0.1762
0.8224 0.1776	1 2	1	1	0.8176 0.1824
0.8853 0.1147	1 2	1	1	0.7925 0.2075
0.7200 0.2800	1 2	1	1	0.7415 0.2585
0.8108 0.1892	1 2	1	1	0.7505 0.2495
0.8772 0.1228	1 2	1	1	0.7386 0.2614
0.7043 0.2957	1 2	1	1	0.7027 0.2973
0.7988 0.2012	1 2	1	1	0.7251 0.2749
0.8687 0.1313	1 2	1	1	0.7251 0.2749



© Eric Xing @ CMU, 2006-2011

7

Another Example



$x = 6, 2, 3, 5, 6, 2, 6, 3, 6, 6$

V_t^k	$ptr(k, t)$	Veterbi	PD	$p(y_t^k = 1 x)$
0.2500 0.7500	N/A	2	2	0.2733 0.7267
0.5263 0.4737	2 2	1	1	0.6040 0.3960
0.6494 0.3506	1 2	1	1	0.6538 0.3462
0.7143 0.2857	1 1	1	1	0.6062 0.3938
0.3333 0.6667	1 1	2	2	0.2861 0.7139
0.5263 0.4737	2 2	2	1	0.5342 0.4658
0.2703 0.7297	1 2	2	2	0.2734 0.7266
0.5263 0.4737	2 2	2	1	0.5226 0.4774
0.2703 0.7297	1 2	2	2	0.2252 0.7748
0.1818 0.8182	2 2	2	2	0.2159 0.7841

max



Same transition probabilities

© Eric Xing @ CMU, 2006-2011

8

Computational Complexity and implementation details



- What is the running time, and space required, for Forward, and Backward?

$$\alpha_t^k = p(x_t | y_t^k = 1) \sum_i \alpha_{t-1}^i a_{i,k}$$

$$\beta_t^k = \sum_i a_{k,i} p(x_{t+1} | y_{t+1}^i = 1) \beta_{t+1}^i$$

$$V_t^k = p(x_t | y_t^k = 1) \max_i a_{i,k} V_{t-1}^i$$

Time: $O(K^2 T)$

Space: $O(KM)$

- Useful implementation technique to avoid underflows

- Viterbi: sum of logs
- Forward/Backward: rescaling at each position by multiplying by a constant

Three Main Questions on HMMs



1. Evaluation

GIVEN an HMM M , and a sequence x ,
 FIND Prob ($x | M$)
 ALGO. Forward

$x_1 \dots x_T$
 $\frac{1}{T} \frac{A}{B}$
 μ, Σ

2. Decoding

GIVEN an HMM M , and a sequence x ,
 FIND the sequence y of states that maximizes, e.g., $P(y | x, M)$,
 or the most probable subsequence of states
 ALGO. Viterbi, Forward-backward

3. Learning

GIVEN an HMM M , with unspecified transition/emission probs.,
 and a sequence x ,
 FIND parameters $\theta = (\pi_i, a_{ij}, \eta_{ik})$ that maximize $P(x | \theta)$
 ALGO. Baum-Welch (EM)

Learning HMM: two scenarios

- **Supervised learning:** estimation when the “right answer” is known

- **Examples:**

GIVEN: a genomic region $x = x_1 \dots x_{1,000,000}$ where we have good (experimental) annotations of the CpG islands

GIVEN: the casino player allows us to observe him one evening, as he changes dice and produces 10,000 rolls

$$\left\{ \begin{array}{c} x_1 \dots x_T \\ \eta_1 \dots \eta_T \end{array} \right\}$$

- **Unsupervised learning:** estimation when the “right answer” is unknown

- **Examples:**

GIVEN: the porcupine genome; we don't know how frequent are the CpG islands there, neither do we know their composition

GIVEN: 10,000 rolls of the casino player, but we don't see when he changes dice

$$\{x_1 \dots x_T\}$$

- **QUESTION:** Update the parameters θ of the model to maximize $P(x|\theta)$ --- Maximal likelihood (ML) estimation

© Eric Xing @ CMU, 2006-2011

11

MLE

$$\begin{aligned} \{x_i, y_i\} & \quad \left(\begin{array}{c} y \rightarrow y \rightarrow y \\ \downarrow \quad \downarrow \quad \downarrow \\ x \quad x \quad x \end{array} \right) \\ \{\pi, A, B\} \arg \max_{\pi, A, B} \log P(x, y) = \log \prod_i P(x_i, y_i) \\ &= \log \prod_i P(y_{t+1} | y_t) \pi(P(x_t | y_t)) \\ &= \log \prod_{i,j} a_{ij}^{\delta(y_{t+1}=i) \delta(y_t=j)} \\ \text{conditional counting} & \quad \begin{array}{c} y_t \rightarrow y_{t+1} \rightarrow x_{t+1} \\ \downarrow \quad \downarrow \quad \downarrow \\ \pi_k \quad a_{ij} \quad b_{kj} \end{array} \\ &= \frac{\#(y_t = k)}{\sum \#(y_t = k')} = \frac{\#(y_{t-1}=i, y_t=j)}{\#(y_{t-1}=i)} \\ & \quad \pi_i = \begin{bmatrix} \pi_1 \\ \vdots \\ \pi_k \end{bmatrix} \quad \pi_{ik} = \frac{n_{ik}}{N} \end{aligned}$$

© Eric Xing @ CMU, 2006-2011

12

Supervised ML estimation



- Given $x = x_1 \dots x_N$ for which the true state path $y = y_1 \dots y_N$ is known,

- Define:

$$\begin{aligned} A_{ij} &= \# \text{ times state transition } i \rightarrow j \text{ occurs in } y \\ B_{ik} &= \# \text{ times state } i \text{ in } y \text{ emits } k \text{ in } x \end{aligned}$$

- We can show that the maximum likelihood parameters θ are:

$$\begin{aligned} a_{ij}^{ML} &= \frac{\#(i \rightarrow j)}{\#(i \rightarrow \bullet)} = \frac{\sum_n \sum_{t=2}^T y_{n,t-1}^i y_{n,t}^j}{\sum_n \sum_{t=2}^T y_{n,t-1}^i} = \frac{A_{ij}}{\sum_{j'} A_{ij'}} \\ b_{ik}^{ML} &= \frac{\#(i \rightarrow k)}{\#(i \rightarrow \bullet)} = \frac{\sum_n \sum_{t=1}^T y_{n,t}^i x_{n,t}^k}{\sum_n \sum_{t=1}^T y_{n,t}^i} = \frac{B_{ik}}{\sum_{k'} B_{ik'}} \end{aligned}$$

- What if y is continuous? We can treat $\{(x_{n,t}, y_{n,t}) : t = 1:T, n = 1:N\}$ as $N \times T$ observations of, e.g., a Gaussian, and apply learning rules for Gaussian ...

© Eric Xing @ CMU, 2006-2011

13

Supervised ML estimation, ctd.



- Intuition:

- When we know the underlying states, the best estimate of θ is the average frequency of transitions & emissions that occur in the training data

- Drawback:

- Given little data, there may be overfitting:
 - $P(x|\theta)$ is maximized, but θ is unreasonable
 - 0 probabilities – VERY BAD**

- Example:

- Given 10 casino rolls, we observe

$x = 2, 1, 5, 6, 1, 2, 3, 6, 2, 3$
 $y = F, F, F, F, F, F, F, F, F, F$

- Then:

$a_{FF} = 1; \quad a_{FL} = 0$
 $b_{F1} = b_{F3} = .2;$
 $b_{F2} = .3; b_{F4} = 0; b_{F5} = b_{F6} = .1$

© Eric Xing @ CMU, 2006-2011

14

- © Eric Xing @ CMU, 2006-2011

15

© Eric Xing @ CMU, 2006-2011

16

Unsupervised ML estimation

- Given $\mathbf{x} = x_1 \dots x_N$ for which the true state path $y = y_1 \dots y_N$ is unknown,

- EXPECTATION MAXIMIZATION**

0. Starting with our best guess of a model $\mathcal{M}$, parameters θ .
1. Estimate A_{ij}, B_{ik} in the training data
 - How? $A_{ij} = \sum_{n,t} \langle y_{n,t-1}^i y_{n,t}^j \rangle$ $B_{ik} = \sum_{n,t} \langle y_{n,t}^i x_{n,t}^k \rangle$,
 - Update θ according to A_{ij}, B_{ik}
 - Now a "supervised learning" problem
2. Repeat 1 & 2, until convergence

This is called the Baum-Welch Algorithm

We can get to a provably more (or equally) likely parameter set θ each iteration

The Baum Welch algorithm

- The complete log likelihood

$$\ell_c(\theta; \mathbf{x}, \mathbf{y}) = \log p(\mathbf{x}, \mathbf{y}) = \log \prod_n \left(p(y_{n,1}) \prod_{t=2}^T p(y_{n,t} | y_{n,t-1}) \prod_{t=1}^T p(x_{n,t} | y_{n,t}) \right)$$

- The expected complete log likelihood

$$\langle \ell_c(\theta; \mathbf{x}, \mathbf{y}) \rangle = \sum_n \left(\langle y_{n,1}^i \rangle_{p(y_{n,1} | \mathbf{x}_n)} \log \pi_i \right) + \sum_n \sum_{t=2}^T \left(\langle y_{n,t-1}^i y_{n,t}^j \rangle_{p(y_{n,t-1}, y_{n,t} | \mathbf{x}_n)} \log a_{i,j} \right) + \sum_n \sum_{t=1}^T \left(\langle x_{n,t}^k y_{n,t}^i \rangle_{p(y_{n,t} | \mathbf{x}_n)} \log b_{i,k} \right)$$

- EM

- The E step

$$\gamma_{n,t}^i = \langle y_{n,t}^i \rangle = p(y_{n,t}^i = 1 | \mathbf{x}_n)$$

$$\xi_{n,t}^{i,j} = \langle y_{n,t-1}^i y_{n,t}^j \rangle = p(y_{n,t-1}^i = 1, y_{n,t}^j = 1 | \mathbf{x}_n)$$

- The M step ("symbolically" identical to MLE)

$$\pi_i^{ML} = \frac{\sum_n \gamma_{n,1}^i}{N} \quad a_{ij}^{ML} = \frac{\sum_n \sum_{t=2}^T \xi_{n,t}^{i,j}}{\sum_n \sum_{t=1}^{T-1} \gamma_{n,t}^i} \quad b_{ik}^{ML} = \frac{\sum_n \sum_{t=1}^T \gamma_{n,t}^i x_{n,t}^k}{\sum_n \sum_{t=1}^T \gamma_{n,t}^i}$$

The Baum-Welch algorithm -- comments



Time Complexity:

iterations $\times O(K^2N)$

- Guaranteed to increase the log likelihood of the model
- Not guaranteed to find globally best parameters
- Converges to local optimum, depending on initial conditions
- Too many parameters / too large model: Overt-fitting

Summary: the HMM algorithms



Questions:

- **Evaluation:** What is the probability of the observed sequence? **Forward**
- **Decoding:** What is the probability that the state of the 3rd roll is loaded, given the observed sequence? **Forward-Backward**
- **Decoding:** What is the most likely die sequence? **Viterbi**
- **Learning:** Under what parameterization are the observed sequences most probable? **Baum-Welch (EM)**

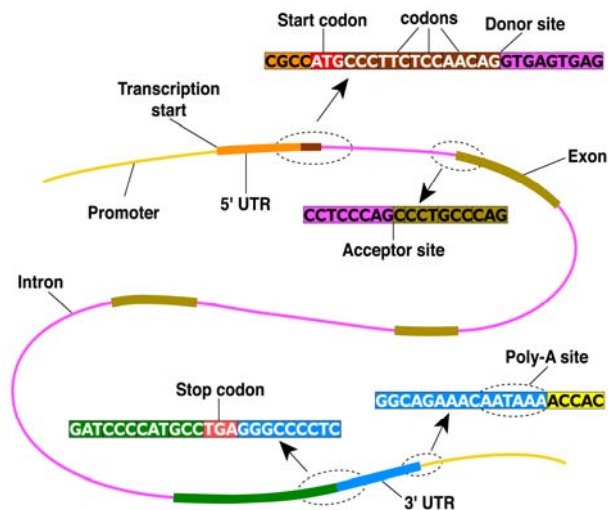
Applications of HMMs

- Some early applications of HMMs
 - finance, but we never saw them
 - speech recognition
 - modelling ion channels
- In the mid-late 1980s HMMs entered genetics and molecular biology, and they are now firmly entrenched.
- Some current applications of HMMs to biology
 - mapping chromosomes
 - aligning biological sequences
 - predicting sequence structure
 - inferring evolutionary relationships
 - finding genes in DNA sequence

© Eric Xing @ CMU, 2006-2011

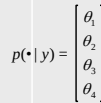
21

Typical structure of a gene

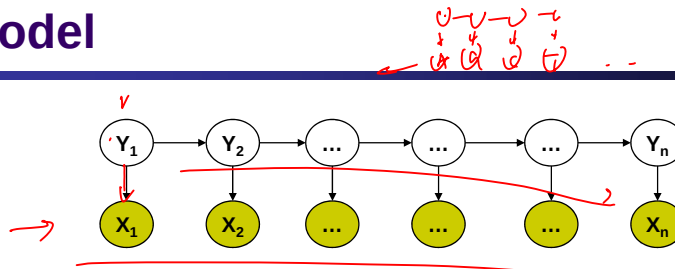


© Eric Xing @ CMU, 2006-2011

22

[illegible]

23



- 3

$$P(x|y) P(y) \rightarrow \frac{P(y|x)}{P(y)}$$

24

Recall Generative vs. Discriminative Classifiers

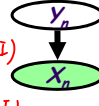
- Goal: Wish to learn $f: X \rightarrow Y$, e.g., $P(Y|X)$
- Generative classifiers (e.g., Naïve Bayes):
 - Assume some functional form for $P(X|Y)$, $P(Y)$
This is a '**generative**' model of the data!
 - Estimate parameters of $P(X|Y)$, $P(Y)$ directly from training data
 - Use Bayes rule to calculate $P(Y|X=x)$
- Discriminative classifiers (e.g., logistic regression)
 - Directly assume some functional form for $P(Y|X)$
This is a '**discriminative**' model of the data!
 - Estimate parameters of $P(Y|X)$ directly from training data

$$\frac{\exp(-\theta^T f(x, y))}{Z}$$

$$P(y) \sim \text{Multi}$$

$$P(x|y) \sim N(\mu, \Sigma)$$

$$P(y|x) = \frac{1}{1 + \exp(\theta^T x)}$$

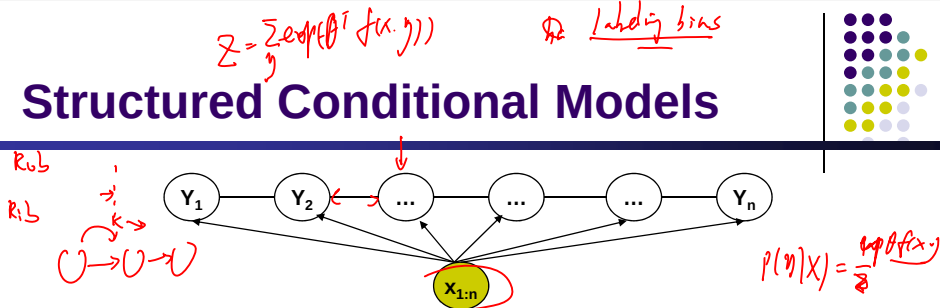


$$P(y|x) = \frac{1}{1 + \exp(-\theta^T x)} \Rightarrow \frac{\exp(\theta^T f(x, y))}{Z}$$

© Eric Xing @ CMU, 2006-2011

25

Structured Conditional Models



- Conditional probability $P(\text{label sequence } y \mid \text{observation sequence } x)$ rather than joint probability $P(y, x)$
 - Specify the probability of possible label sequences given an observation sequence
- Allow arbitrary, non-independent features on the observation sequence X
- The probability of a transition between labels may depend on **past** and **future** observations
- Relax strong independence assumptions in generative models

© Eric Xing @ CMU, 2006-2011

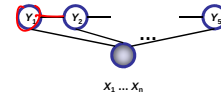
26

Conditional Distribution

- If the graph $G = (V, E)$ of Y is a tree, the conditional distribution over the label sequence $Y = y$, given $X = x$, by the Hammersley Clifford theorem of random fields is:

$$p_{\theta}(y|x) \propto \frac{1}{Z} \exp \left(\sum_{e \in E, k} \lambda_k \underbrace{f_k(e, y|_e, x)}_{f_k(x,y)} + \sum_{v \in V, k} \mu_k g_k(v, y|_v, x) \right)$$

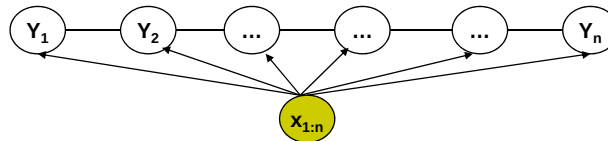
- x is a data sequence
- y is a label sequence
- v is a vertex from vertex set V = set of label random variables
- e is an edge from edge set E over V
- f_k and g_k are given and fixed. g_k is a Boolean vertex feature; f_k is a Boolean edge feature
- k is the number of features
- $\theta = (\lambda_1, \lambda_2, \dots, \lambda_n; \mu_1, \mu_2, \dots, \mu_n)$; λ_k and μ_k are parameters to be estimated
- $y|_e$ is the set of components of y defined by edge e
- $y|_v$ is the set of components of y defined by vertex v



© Eric Xing @ CMU, 2006-2011

27

Conditional Random Fields



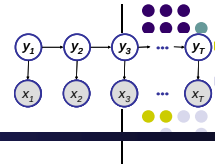
$$P(y_{1:n} | x_{1:n}) = \frac{1}{Z(x_{1:n})} \prod_{i=1}^n \phi(y_i, y_{i-1}, x_{1:n}) = \frac{1}{Z(x_{1:n}, w)} \prod_{i=1}^n \exp(w^T f(y_i, y_{i-1}, x_{1:n}))$$

- CRF is a ~~partially directed~~ model
 - Discriminative model
 - Usage of global normalizer $Z(x)$
 - Models the dependence between each state and the entire observation sequence

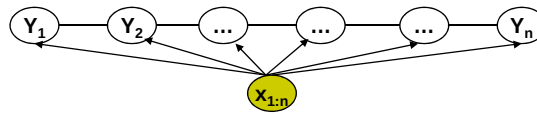
© Eric Xing @ CMU, 2006-2011

28

Conditional Random Fields



- General parametric form:



$$P(\mathbf{y}|\mathbf{x}) = \frac{1}{Z(\mathbf{x}, \lambda, \mu)} \exp\left(\sum_{i=1}^n \left(\sum_k \lambda_k f_k(y_i, y_{i-1}, \mathbf{x}) + \sum_l \mu_l g_l(y_i, \mathbf{x})\right)\right)$$

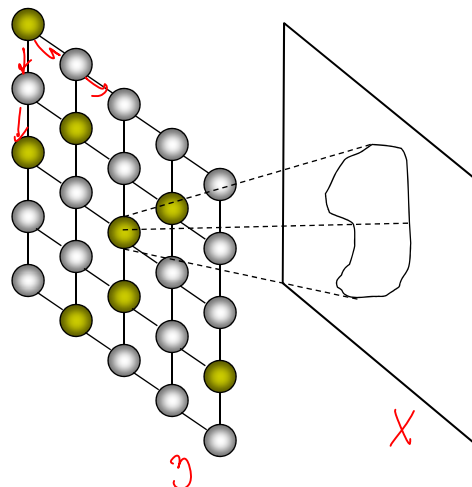
$$= \frac{1}{Z(\mathbf{x}, \lambda, \mu)} \exp\left(\sum_{i=1}^n (\lambda^T \mathbf{f}(y_i, y_{i-1}, \mathbf{x}) + \mu^T \mathbf{g}(y_i, \mathbf{x}))\right)$$

$$\text{where } Z(\mathbf{x}, \lambda, \mu) = \sum_{\mathbf{y}} \exp\left(\sum_{i=1}^n (\lambda^T \mathbf{f}(y_i, y_{i-1}, \mathbf{x}) + \mu^T \mathbf{g}(y_i, \mathbf{x}))\right)$$

© Eric Xing @ CMU, 2006-2011

29

Conditional Random Fields



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

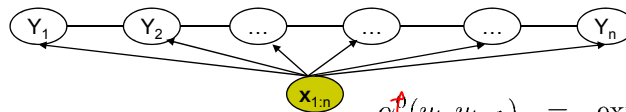
- Allow arbitrary dependencies on input
- Clique dependencies on labels
- Use approximate inference for general graphs

© Eric Xing @ CMU, 2006-2011

CRFs: Inference

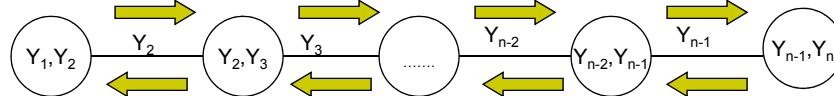
$$m_{i \rightarrow j}(S_{ij}) = \sum_{C_i \setminus S_{ij}} \psi_{C_i} \prod_{k \neq j} m_{k \rightarrow i}(S_{ki})$$

- Computing marginals using a message passing algorithm called “sum-product”:



$$\alpha^{\uparrow}(y_i, y_{i-1}) = \exp(\lambda^T \mathbf{f}(y_i, y_{i-1}, \mathbf{x}_d) + \mu^T \mathbf{g}(y_i, \mathbf{x}_d))$$

- Initialization:



- After calibration:

$$P(y_i, y_{i-1} | \mathbf{x}_d) \propto \alpha(y_i, y_{i-1})$$

$$\Rightarrow P(y_i, y_{i-1} | \mathbf{x}_d) = \frac{\alpha(y_i, y_{i-1})}{\sum_{y_i, y_{i-1}} \alpha(y_i, y_{i-1})} = \alpha'(y_i, y_{i-1})$$

Also called forward-backward algorithm

© Eric Xing @ CMU, 2006-2011

31

CRFs: Inference

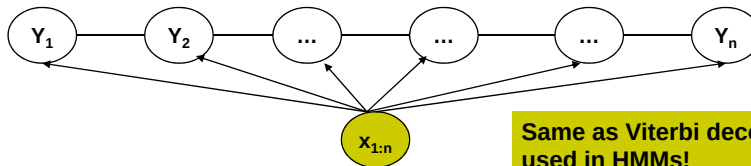
$$m_{i \rightarrow j}(S_{ij}) = \max_{C_i \setminus S_{ij}} \psi_{C_i} \prod_{k \neq j} m_{k \rightarrow i}(S_{ki})$$

- Given CRF parameters λ and μ , find the $\mathbf{y}^*$ that maximizes $P(\mathbf{y} | \mathbf{x})$

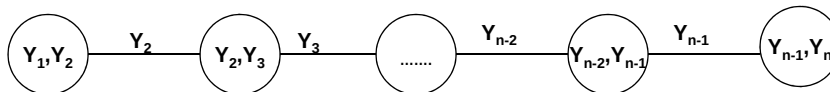
$$\mathbf{y}^* = \arg \max_{\mathbf{y}} \exp\left(\sum_{i=1}^n (\lambda^T \mathbf{f}(y_i, y_{i-1}, \mathbf{x}) + \mu^T \mathbf{g}(y_i, \mathbf{x}))\right)$$

- Can ignore $Z(\mathbf{x})$ because it is not a function of $\mathbf{y}$

- Again run a message-passing algorithm called “max-product”:



Same as Viterbi decoding used in HMMs!



© Eric Xing @ CMU, 2006-2011

32

CRF learning

- Given $\{(\mathbf{x}_d, \mathbf{y}_d)\}_{d=1}^N$, find λ^*, μ^* such that

$$\begin{aligned}\lambda^*, \mu^* &= \arg \max_{\lambda, \mu} L(\lambda, \mu) = \arg \max_{\lambda, \mu} \prod_{d=1}^N P(\mathbf{y}_d | \mathbf{x}_d, \lambda, \mu) \\ &= \arg \max_{\lambda, \mu} \prod_{d=1}^N \frac{1}{Z(\mathbf{x}_d, \lambda, \mu)} \exp\left(\sum_{i=1}^n (\lambda^T \mathbf{f}(y_{d,i}, y_{d,i-1}, \mathbf{x}_d) + \mu^T \mathbf{g}(y_{d,i}, \mathbf{x}_d))\right) \\ &= \arg \max_{\lambda, \mu} \sum_{d=1}^N \left(\sum_{i=1}^n (\lambda^T \mathbf{f}(y_{d,i}, y_{d,i-1}, \mathbf{x}_d) + \mu^T \mathbf{g}(y_{d,i}, \mathbf{x}_d)) - \log Z(\mathbf{x}_d, \lambda, \mu) \right)\end{aligned}$$

Gradient of the log-partition function in an exponential family is the expectation of the sufficient statistics.

- Computing the gradient w.r.t λ :

$$\nabla_{\lambda} L(\lambda, \mu) = \sum_{d=1}^N \left(\sum_{i=1}^n \mathbf{f}(y_{d,i}, y_{d,i-1}, \mathbf{x}_d) - \sum_{\mathbf{y}} (P(\mathbf{y} | \mathbf{x}_d) \sum_{i=1}^n \mathbf{f}(y_{d,i}, y_{d,i-1}, \mathbf{x}_d)) \right)$$

© Eric Xing @ CMU, 2006-2011

33

CRF learning

$$\nabla_{\lambda} L(\lambda, \mu) = \sum_{d=1}^N \left(\sum_{i=1}^n \mathbf{f}(y_{d,i}, y_{d,i-1}, \mathbf{x}_d) - \sum_{\mathbf{y}} (P(\mathbf{y} | \mathbf{x}_d) \sum_{i=1}^n \mathbf{f}(y_{d,i}, y_{d,i-1}, \mathbf{x}_d)) \right)$$

- Computing the model expectations:

- Requires exponentially large number of summations: Is it intractable?

$$\begin{aligned}\sum_{\mathbf{y}} (P(\mathbf{y} | \mathbf{x}_d) \sum_{i=1}^n \mathbf{f}(y_{d,i}, y_{d,i-1}, \mathbf{x}_d)) &= \sum_{i=1}^n \left(\sum_{\mathbf{y}} \mathbf{f}(y_{d,i}, y_{d,i-1}, \mathbf{x}_d) P(\mathbf{y} | \mathbf{x}_d) \right) \\ &= \sum_{i=1}^n \sum_{y_i, y_{i-1}} \mathbf{f}(y_i, y_{i-1}, \mathbf{x}_d) P(y_i, y_{i-1} | \mathbf{x}_d)\end{aligned}$$

Expectation of $\mathbf{f}$ over the corresponding marginal probability of neighboring nodes!!

- Tractable!

- Can compute marginals using the sum-product algorithm on the chain

© Eric Xing @ CMU, 2006-2011

34

CRF learning

- Computing feature expectations using calibrated potentials:

$$\sum_{y_i, y_{i-1}} \mathbf{f}(y_i, y_{i-1}, \mathbf{x}_d) P(y_i, y_{i-1} | \mathbf{x}_d) = \sum_{y_i, y_{i-1}} \mathbf{f}(y_i, y_{i-1}, \mathbf{x}_d) \alpha'(y_i, y_{i-1})$$

- Now we know how to compute $\nabla_{\lambda} L(\lambda, \mu)$:

$$\begin{aligned} \nabla_{\lambda} L(\lambda, \mu) &= \sum_{d=1}^N \left(\sum_{i=1}^n \mathbf{f}(y_{d,i}, y_{d,i-1}, \mathbf{x}_d) - \sum_{\mathbf{y}} (P(\mathbf{y} | \mathbf{x}_d) \sum_{i=1}^n \mathbf{f}(y_i, y_{i-1}, \mathbf{x}_d)) \right) \\ &= \sum_{d=1}^N \left(\sum_{i=1}^n (\mathbf{f}(y_{d,i}, y_{d,i-1}, \mathbf{x}_d) - \sum_{y_i, y_{i-1}} \alpha'(y_i, y_{i-1}) \mathbf{f}(y_i, y_{i-1}, \mathbf{x}_d)) \right) \end{aligned}$$

- Learning can now be done using gradient ascent:

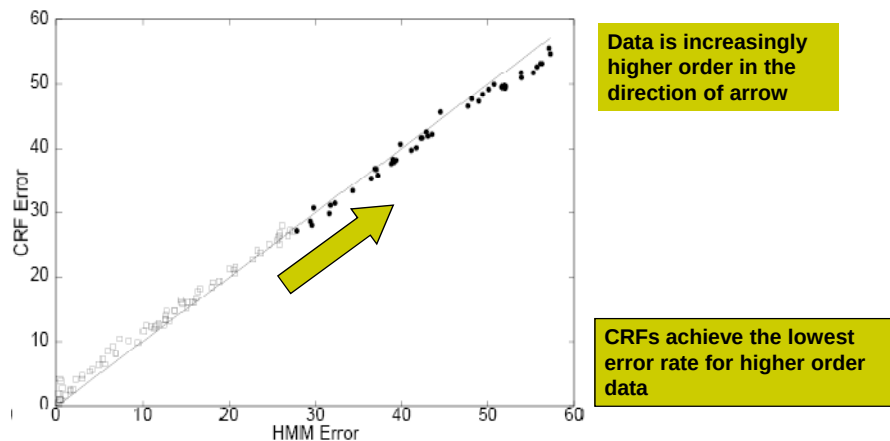
$$\begin{aligned} \lambda^{(t+1)} &= \lambda^{(t)} + \eta \nabla_{\lambda} L(\lambda^{(t)}, \mu^{(t)}) \\ \mu^{(t+1)} &= \mu^{(t)} + \eta \nabla_{\mu} L(\lambda^{(t)}, \mu^{(t)}) \end{aligned}$$

© Eric Xing @ CMU, 2006-2011

35

CRFs: some empirical results

- Comparison of error rates on synthetic data



© Eric Xing @ CMU, 2006-2011

36

CRFs: some empirical results

- Parts of Speech tagging

<i>model</i>	<i>error</i>	<i>oov error</i>
HMM	5.69%	45.99%
MEMM	6.37%	54.61%
CRF	5.55%	48.05%
MEMM ⁺	4.81%	26.99%
CRF ⁺	4.27%	23.76%

⁺Using spelling features

- Using same set of features: HMM $\geq$ CRF
- Using additional overlapping features: CRF⁺ $\gg$ HMM

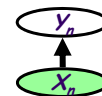
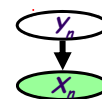
© Eric Xing @ CMU, 2006-2011

37

Summary

- Conditional Random Fields is a discriminative Structured Input Output model!
- HMM is a generative structured I/O model
- Complementary strength and weakness:

- 1.
- 2.
- 3.
- ...



© Eric Xing @ CMU, 2006-2011

38