

Readings:

K&F: 11.3, 11.5
Yedidia et al. paper from the class website
Chapter 9 – Jordan
K&F: 16.2

Unifying Variational and GBP Learning Parameters of MNs EM for BNs

Graphical Models – 10708

Carlos Guestrin

Carnegie Mellon University

November 15th, 2006

1

Revisiting Mean-Fields

$$\ln Z = \underbrace{F[P_{\mathcal{F}}, Q]}_{\text{constant}} + \underbrace{D(Q||P_{\mathcal{F}})}_{\text{KL Divergence}} \quad F[P_{\mathcal{F}}, Q] = \sum_{\phi \in \mathcal{F}} E_Q[\ln \phi] + H_Q(\mathcal{X})$$

- Choice of Q:

- Optimization problem:

$$\max_{Q_i} \sum_{\phi \in \mathcal{F}} E_Q[\ln \phi] + \sum_i H_{Q_i}(X_i)$$

$$\text{s.t. } \sum_i Q_i(x_i) = 1$$

$$Q_i(x_i) \geq 0$$

decomposes
mean-fields
solution is
a stationary
point...

$$\max_Q F[P_{\mathcal{F}}, Q] = \sum_{\phi \in \mathcal{F}} E_Q[\ln \phi] + \sum_j H_{Q_j}(X_j), \quad \forall i, \sum_{x_i} Q_i(x_i) = 1$$

Interpretation of energy functional

■ Energy functional: $F[P_F, Q] = \sum_{\phi \in \mathcal{F}} E_Q[\ln \phi] + H_Q(\mathcal{X})$

■ Exact if $P=Q$: $\ln Z = F[P_F, Q] + D(Q||P_F)$
Q ∈ P *maximized*

■ View problem as an approximation of entropy term:

estimate expectation w.r.t. Q

approximating $H_Q(\mathcal{X}) \approx H_P(\mathcal{X})$

interpretation: $H_P(\mathcal{X}) \approx \sum_i H_{Q_i}(x_i)$

10-708 - ©Carlos Guestrin 2006

3

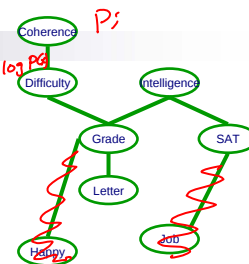
Entropy of a tree distribution

■ Entropy term: $H_P(\mathcal{X}) = -E_P[\log P(\mathcal{X})] = \sum_{\mathcal{X}} P(\mathcal{X}) \log \frac{1}{P(\mathcal{X})}$

■ Joint distribution: $P(\mathcal{X}) = \prod_{i,j} \psi_{ij}(x_i, x_j)$

■ Decomposing entropy term:

$$\begin{aligned} P(\mathcal{X}) &= \frac{P(CD) \cdot P(DG) \cdot P(GI) \cdot P(IG) \cdot P(LI)}{P(CD) \cdot P(G) \cdot P(G) \cdot P(LI)} \\ &= \sum_{\text{cliques}} H(C_i) - \sum_{\text{seps}} H(S_{ij}) \\ H_P(\mathcal{X}) &= H_P(CD) + H_P(DG) + H_P(GI) + H_P(IG) + H_P(LI) \\ &\quad - H_P(D) - H_P(G) - H_P(G) - H_P(I) \end{aligned}$$



CD
 | D
 DG — G — GI — IS
 | G
 GL

■ More generally: $H_P(\mathcal{X}) = \sum_{(i,j) \in E} H(X_i, X_j) - \sum_i (d_i - 1) H(X_i)$
 □ d_i number neighbors of X_i

10-708 - ©Carlos Guestrin 2006

4

Loopy BP & Bethe approximation

■ Energy functional: $F[P_{\mathcal{F}}, Q] = \sum_{\phi \in \mathcal{F}} E_Q[\ln \phi] + H_Q(\mathcal{X})$

■ Bethe approximation of Free Energy:

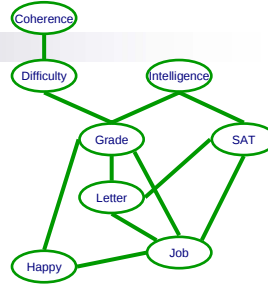
□ use entropy for trees, but loopy graphs:

$$\tilde{F}[P_{\mathcal{F}}, Q] = \sum_{(i,j) \in E} E_{\pi_{ij}}[\ln \phi_{ij}] + \sum_{(i,j) \in E} H_{\pi_{ij}}(X_i, X_j) - \sum_i (d_i - 1) H_{\pi_i}(X_i)$$

↑
sum edges
↑
weighted
↑
node entropies

↑
Selfies

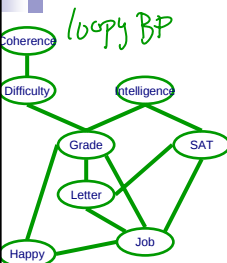
■ Theorem: If Loopy BP converges, resulting π_{ij} & π_i are stationary point (usually local maxima) of Bethe Free energy!



10-708 - ©Carlos Guestrin 2006

5

GBP & Kikuchi approximation



■ Exact Free energy: Junction Tree

$$F[P_{\mathcal{F}}, Q] = \sum_{(i,j) \in E} E_{\pi_{ij}}[\ln \phi_{ij}] + \sum_i H_{\pi_{C_i}}(C_i) - \sum_{(i,j) \in \mathcal{T}} H_{\pi_{S_{ij}}}(S_{ij})$$

↑
full cliques (exact)
↑
pairs

■ Bethe Free energy:

$$\tilde{F}[P_{\mathcal{F}}, Q] = \sum_{(i,j) \in E} E_{\pi_{ij}}[\ln \phi_{ij}] + \sum_{(i,j) \in E} H_{\pi_{ij}}(X_i, X_j) - \sum_i (d_i - 1) H_{\pi_i}(X_i)$$

↑
weights

■ Kikuchi approximation: Generalized cluster graph

$$H_{Q_{\mathcal{K}}}(\mathcal{X}) = \sum_{\text{cliques}} H(C_i) - \sum_{\text{seps}} c_{ij} H(S_{ij})$$

- spectrum from Bethe to exact
- entropy terms weighted by counting numbers
- see Yedidia et al.

■ Theorem: If GBP converges, resulting π_{C_i} are stationary point (usually local maxima) of Kikuchi Free energy!



10-708 - ©Carlos Guestrin 2006

6

What you need to know about GBP

- Spectrum between Loopy BP & Junction Trees:
 - More computation, but typically better answers
- If satisfies RIP, equations are very simple
- General setting, slightly trickier equations, but not hard
- Relates to variational methods: Corresponds to local optima of approximate version of energy functional

10-708 – ©Carlos Guestrin 2006

7

Announcements

- Tomorrow's recitation
 - Khalid on learning Markov Networks

10-708 – ©Carlos Guestrin 2006

8

Learning Parameters of a BN

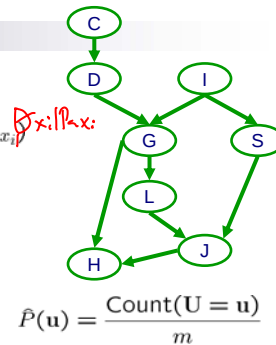
- Log likelihood decomposes:

$$\ell(\mathcal{D} : \theta) = \log P(\mathcal{D} | \theta) = m \sum_i \sum_{x_i, \text{Pa}_{x_i}} \underbrace{\hat{P}(x_i, \text{Pa}_{x_i})}_{\text{empirical dist}} \log P(x_i | \text{Pa}_{x_i})$$

$\hat{P}(x_i, \text{Pa}_{x_i})$

$$\theta_{x_i=t | \text{Pa}_{x_i}=u} = \frac{\text{Count}(x_i=t, \text{Pa}_{x_i}=u)}{\text{Count}(\text{Pa}_{x_i}=u)}$$

- Learn each CPT independently
- Use counts



10-708 - ©Carlos Guestrin 2006

9

Log Likelihood for MN

$$Z = \sum_x \prod_i \psi_i(x)$$

- Log likelihood of the data:

iid. m data points

$$\ell(\mathcal{D} : \theta) = \log P(\mathcal{D} | \theta) = \log \prod_{i \in \mathcal{D}} \frac{1}{Z} \prod_i \psi_i(c_i = c_i^{(i)})$$

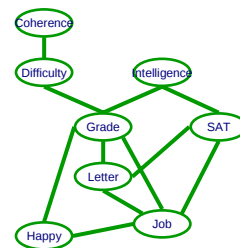
$$= \sum_{i \in \mathcal{D}} \sum_i \log \psi_i(c_i = c_i^{(i)}) - \sum_{i \in \mathcal{D}} \log Z$$

$m \log Z$

$$= \sum_i \sum_{i \in \mathcal{D}} \log \psi_i(c_i = c_i^{(i)}) - m \log Z$$

$$= \sum_{c_i} \sum_i \text{Count}(c_i = c_i) \log \psi_i(c_i = c_i) - m \log Z$$

$$= m \sum_i \hat{P}(c_i = c_i) \log \psi_i(c_i = c_i) - m \log Z(\psi_1, \dots, \psi_n)$$



10-708 - ©Carlos Guestrin 2006

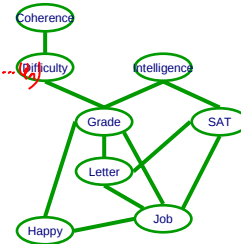
10

Log Likelihood doesn't decompose for MNs

$$\hat{P}(u) = \frac{\text{Count}(U = u)}{m}$$

Log likelihood:

$$\ell(\mathcal{D} : \theta) = \log P(\mathcal{D} | \theta, \mathcal{G}) = m \sum_i \sum_{c_i} \hat{P}(c_i) \log \psi_i(c_i) - m \log Z$$



A concave problem

- Can find global optimum!!

Term log Z doesn't decompose!!

10-708 - ©Carlos Guestrin 2006

11

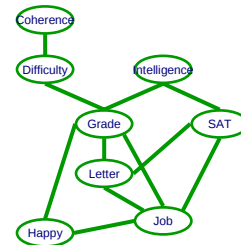
Derivative of Log Likelihood for MNs

$$\frac{\partial}{\partial x} \log(x) = \frac{1}{x} ; \quad \frac{\partial}{\partial x} \log z(x) = \frac{z'(x)}{z(x)}$$

$$\hat{P}(u) = \frac{\text{Count}(U = u)}{m}$$

$$\ell(\mathcal{D} : \theta) = \log P(\mathcal{D} | \theta, \mathcal{G}) = m \sum_i \sum_{c_i} \hat{P}(c_i) \log \psi_i(c_i) - m \log Z$$

$$\begin{aligned} \frac{\partial \ell}{\partial \psi_i(c_i)} &= \frac{\partial}{\partial \psi_i(c_i)} \left[m \sum_i \sum_{c_i} \hat{P}(c_i) \log \psi_i(c_i) \right] - \frac{\partial}{\partial \psi_i(c_i)} \log Z \\ &= \hat{P}(c_i) \frac{\partial}{\partial \psi_i(c_i)} \log \psi_i(c_i) \\ &= \frac{\hat{P}(c_i)}{\psi_i(c_i)} \end{aligned}$$



$$Z = \sum_x \prod_i \psi_i(c_i)$$

$$\frac{\partial}{\partial \psi_i(c_i)} \log Z = \frac{\frac{\partial Z}{\partial \psi_i(c_i)}}{Z} \dots$$

10-708 - ©Carlos Guestrin 2006

12

Derivative of Log Likelihood for MNs

$$\hat{P}(u) = \frac{\text{Count}(U = u)}{m}$$

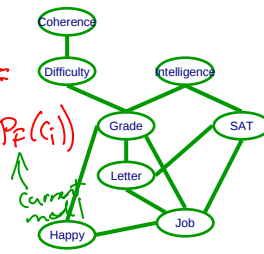
$$\ell(\mathcal{D} : \theta) = \log P(\mathcal{D} | \theta, \mathcal{G}) = m \sum_i \sum_{c_i} \hat{P}(c_i) \log \psi_i(c_i) - m \log Z$$

Derivative:

$$\frac{\partial \ell}{\partial \psi_i(c_i)} = \frac{m \hat{P}(c_i)}{\psi_i(c_i)} - \frac{m P_F(c_i)}{\psi_i(c_i)} = \frac{m}{\psi_i(c_i)} (\hat{P}(c_i) - P_F(c_i))$$

- Setting derivative to zero $\frac{\partial \ell}{\partial \psi_i(c_i)} = 0 \Rightarrow m \hat{P}(c_i) = m P_F(c_i)$
- Can optimize using gradient ascent $\frac{\partial \ell}{\partial \psi_i(c_i)}$

Let's look at a simpler solution



10-708 - ©Carlos Guestrin 2006

13

Iterative Proportional Fitting (IPF)

$$\hat{P}(u) = \frac{\text{Count}(U = u)}{m}$$

$$\frac{\partial \ell}{\partial \psi_i(c_i)} = \frac{m \hat{P}(c_i)}{\psi_i(c_i)} - \frac{m P_F(c_i)}{\psi_i(c_i)} = 0$$

Setting derivative to zero:

$$\frac{m \hat{P}(c_i)}{\psi_i^{(t)}(c_i)} = \frac{m P_F^{(t)}(c_i)}{\psi_i^{(t)}(c_i)}$$

Fixed point equation:

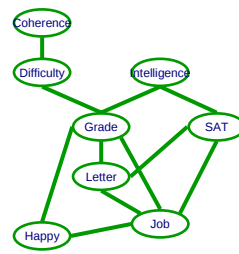
$$\psi_i^{(t+1)}(c_i) = \psi_i^{(t)}(c_i) \cdot \frac{\hat{P}(c_i)}{P_F^{(t)}(c_i)}$$

Iterate and converge to optimal parameters

Each iteration, must compute:

$$P_F^{(t)}(c_i) = \sum_{X=c_i} P(X)$$

inference



tabular representation

$$\psi_i(G, H) = \frac{G}{t} \frac{H}{f}$$

G	t	f
t		
f		

10-708 - ©Carlos Guestrin 2006

14

Log-linear Markov network (most common representation)

- **Feature** is some function $\phi[\mathbf{D}]$ for some subset of variables $\mathbf{D}$

□ e.g., indicator function

$$\psi(G, I, D) = \begin{cases} 1, & \text{if } G=c, I=t, D=f \\ 0, & \text{otherwise} \dots \end{cases}$$

- **Log-linear model** over a Markov network H :

□ a set of features $\phi_1[\mathbf{D}_1], \dots, \phi_k[\mathbf{D}_k]$

■ each $\mathbf{D}_i$ is a subset of a clique in H

■ two ϕ 's can be over the same variables

□ a set of weights $w_1, \dots, w_k$

■ usually learned from data

$$P(X_1, \dots, X_n) = \frac{1}{Z} \exp \left[\sum_{i=1}^k w_i \phi_i(\mathbf{D}_i) \right]$$

learn these
in Jordan Chapter 9: $w_i = \theta_i$
 $\phi_i = f_i$

10/708 - ©Carlos Guestrin 2006

15

Learning params for log linear models 1 – Generalized Iterative Scaling

$$P(X_1, \dots, X_n) = \frac{1}{Z} \exp \left[\sum_{i=1}^k w_i \phi_i(\mathbf{D}_i) \right]$$

- IPF generalizes easily if:

□ $\phi_i(\mathbf{x}) \geq 0$
□ $\sum_i \phi_i(\mathbf{x}) = 1$

- Update rule: $w_i^{(t+1)} = w_i^{(t)} + \log \left(\frac{\sum_{\mathbf{x}} P(\mathbf{x}) \phi_i(\mathbf{x})}{\sum_{\mathbf{x}} P^{(t)}(\mathbf{x}) \phi_i(\mathbf{x})} \right)$

- Must compute: $\text{Scope}[\phi_i] = \mathbf{D}_i$
if $\phi_i[\mathbf{D}_i]$: inference $P_F^{(t)}(\mathbf{D}_i)$

$$E_{P_F^{(t)}}[\phi_i] = \sum_{d_i} P_F^{(t)}(d_i) \phi_i(d_i)$$

$\sum_{\mathbf{x}} P(\mathbf{x}) \phi_i(\mathbf{x}) = \text{expected value of } \phi_i \text{ w.r.t. data}$

$$\text{e.g., } \frac{1}{m} \sum_{j \in D} \phi_i(x^{(j)})$$

$\sum_{\mathbf{x}} P^{(t)}(\mathbf{x}) \phi_i(\mathbf{x}) = \text{expected v. of } \phi_i \text{ w.r.t. current model}$
 $= E_{P^{(t)}}[\phi_i]$

$\sum_i \phi_i(\mathbf{x}) = 1 \quad \forall \mathbf{x}$
not easy...

- If conditions violated, equations are not so simple...

□ c.f., Improved Iterative Scaling [Berger '97]

10/708 - ©Carlos Guestrin 2006

16

Learning params for log linear models 2 – Gradient Ascent

$$P(X_1, \dots, X_n) = \frac{1}{Z} \exp \left[\sum_{i=1}^k w_i \phi_i(\mathbf{D}_i) \right]$$

- Log-likelihood of data:

$$\begin{aligned} \ell(\mathbf{D}; \mathbf{w}) &= \log P(\mathbf{D} | \mathbf{w}) = \log \prod_{j \in \mathcal{D}} \frac{1}{Z} e^{\sum_{i=1}^k w_i \phi_i(x_j)} \\ &= \log e^{\sum_{j \in \mathcal{D}} \sum_{i=1}^k w_i \phi_i(x_j)} - \log Z \\ &= \sum_{j \in \mathcal{D}} \sum_{i=1}^k w_i \phi_i(x_j) - m \log Z \end{aligned}$$

- Compute derivative & optimize
 - usually with conjugate gradient ascent
 - You will do an example in your homework! ☺

w_i's are free!!

10-708 – ©Carlos Guestrin 2006

17

What you need to know about learning MN parameters?

- BN parameter learning easy
- MN parameter learning doesn't decompose!
- Learning requires inference!
- Apply gradient ascent or IPF iterations to obtain optimal parameters
 - applies to both tabular representations and log-linear models

10-708 – ©Carlos Guestrin 2006

18

Thus far, fully supervised learning

- We have assumed fully supervised learning:
- Many real problems have missing data:

10-708 – ©Carlos Guestrin 2006

19

The general learning problem with missing data

- Marginal likelihood – $\mathbf{x}$ is observed, $\mathbf{z}$ is missing:

$$\begin{aligned}\ell(\theta : \mathcal{D}) &= \log \prod_{j=1}^m P(\mathbf{x}_j | \theta) \\ &= \sum_{j=1}^m \log P(\mathbf{x}_j | \theta) \\ &= \sum_{j=1}^m \log \sum_{\mathbf{z}} P(\mathbf{x}_j, \mathbf{z} | \theta)\end{aligned}$$

10-708 – ©Carlos Guestrin 2006

20

E-step

- $\mathbf{x}$ is observed, $\mathbf{z}$ is missing
- Compute probability of missing data given current choice of θ
 - $Q(\mathbf{z}|\mathbf{x}_j)$ for each $\mathbf{x}_j$
 - e.g., probability computed during classification step
 - corresponds to "classification step" in K-means

$$Q^{(t+1)}(\mathbf{z} | \mathbf{x}_j) = P(\mathbf{z} | \mathbf{x}_j, \theta^{(t)})$$

10-708 – ©Carlos Guestrin 2006

21

Jensen's inequality

$$\ell(\theta : \mathcal{D}) = \sum_{j=1}^m \log \sum_{\mathbf{z}} P(\mathbf{z} | \mathbf{x}_j) P(\mathbf{x}_j | \theta)$$

- **Theorem:** $\log \sum_{\mathbf{z}} P(\mathbf{z}) f(\mathbf{z}) \geq \sum_{\mathbf{z}} P(\mathbf{z}) \log f(\mathbf{z})$

10-708 – ©Carlos Guestrin 2006

22

Applying Jensen's inequality

- Use: $\log \sum_{\mathbf{z}} P(\mathbf{z}) f(\mathbf{z}) \geq \sum_{\mathbf{z}} P(\mathbf{z}) \log f(\mathbf{z})$

$$\ell(\theta^{(t)} : \mathcal{D}) = \sum_{j=1}^m \log \sum_{\mathbf{z}} Q^{(t+1)}(\mathbf{z} | \mathbf{x}_j) \frac{P(\mathbf{z}, \mathbf{x}_j | \theta^{(t)})}{Q^{(t+1)}(\mathbf{z} | \mathbf{x}_j)}$$

10-708 – ©Carlos Guestrin 2006

23

The M-step maximizes lower bound on weighted data

- Lower bound from Jensen's:

$$\ell(\theta^{(t)} : \mathcal{D}) \geq \sum_{j=1}^m \sum_{\mathbf{z}} Q^{(t+1)}(\mathbf{z} | \mathbf{x}_j) \log P(\mathbf{z}, \mathbf{x}_j | \theta^{(t)}) + H(Q^{(t+1)})$$

- Corresponds to weighted dataset:

- $\langle \mathbf{x}_1, \mathbf{z}=1 \rangle$ with weight $Q^{(t+1)}(\mathbf{z}=1 | \mathbf{x}_1)$
- $\langle \mathbf{x}_1, \mathbf{z}=2 \rangle$ with weight $Q^{(t+1)}(\mathbf{z}=2 | \mathbf{x}_1)$
- $\langle \mathbf{x}_1, \mathbf{z}=3 \rangle$ with weight $Q^{(t+1)}(\mathbf{z}=3 | \mathbf{x}_1)$
- $\langle \mathbf{x}_2, \mathbf{z}=1 \rangle$ with weight $Q^{(t+1)}(\mathbf{z}=1 | \mathbf{x}_2)$
- $\langle \mathbf{x}_2, \mathbf{z}=2 \rangle$ with weight $Q^{(t+1)}(\mathbf{z}=2 | \mathbf{x}_2)$
- $\langle \mathbf{x}_2, \mathbf{z}=3 \rangle$ with weight $Q^{(t+1)}(\mathbf{z}=3 | \mathbf{x}_2)$
- ...

10-708 – ©Carlos Guestrin 2006

24

The M-step

$$\ell(\theta^{(t)} : \mathcal{D}) \geq \sum_{j=1}^m \sum_{\mathbf{z}} Q^{(t+1)}(\mathbf{z} | \mathbf{x}_j) \log P(\mathbf{z}, \mathbf{x}_j | \theta^{(t)}) + H(Q^{(t+1)})$$

- Maximization step:

$$\theta^{(t+1)} \leftarrow \arg \max_{\theta} \sum_{j=1}^m \sum_{\mathbf{z}} Q^{(t+1)}(\mathbf{z} | \mathbf{x}_j) \log P(\mathbf{z}, \mathbf{x}_j | \theta)$$

- Use expected counts instead of counts:
 - If learning requires $\text{Count}(\mathbf{x}, \mathbf{z})$
 - Use $E_{Q^{(t+1)}}[\text{Count}(\mathbf{x}, \mathbf{z})]$

10-708 – ©Carlos Guestrin 2006

25

Convergence of EM

- Define potential function $F(\theta, Q)$:

$$\ell(\theta : \mathcal{D}) \geq F(\theta, Q) = \sum_{j=1}^m \sum_{\mathbf{z}} Q(\mathbf{z} | \mathbf{x}_j) \log \frac{P(\mathbf{z}, \mathbf{x}_j | \theta)}{Q(\mathbf{z} | \mathbf{x}_j)}$$

- EM corresponds to coordinate ascent on F
 - Thus, maximizes lower bound on marginal log likelihood
 - As seen in Machine Learning class last semester

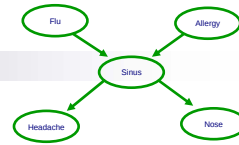
10-708 – ©Carlos Guestrin 2006

26

Data likelihood for BNs

- Given structure, log likelihood of fully observed data:

$$\log P(\mathcal{D} \mid \theta_{\mathcal{G}}, \mathcal{G})$$



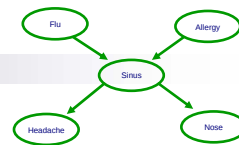
10-708 – ©Carlos Guestrin 2006

27

Marginal likelihood

- What if S is hidden?

$$\log P(\mathcal{D} \mid \theta_{\mathcal{G}}, \mathcal{G})$$



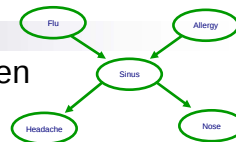
10-708 – ©Carlos Guestrin 2006

28

Log likelihood for BNs with hidden data

- Marginal likelihood – $\mathbf{O}$ is observed, $\mathbf{H}$ is hidden

$$\begin{aligned}\ell(\theta : \mathcal{D}) &= \sum_{j=1}^m \log P(\mathbf{o}^{(j)} | \theta) \\ &= \sum_{j=1}^m \log \sum_{\mathbf{h}} P(\mathbf{h}, \mathbf{o}^{(j)} | \theta)\end{aligned}$$



10-708 – ©Carlos Guestrin 2006

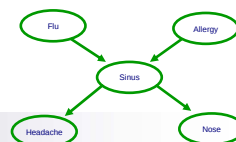
29

E-step for BNs

- E-step computes probability of hidden vars $\mathbf{h}$ given $\mathbf{o}$

$$Q^{(t+1)}(\mathbf{x} | \mathbf{o}) = P(\mathbf{x} | \mathbf{o}, \theta^{(t)})$$

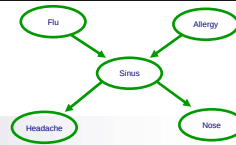
- Corresponds to inference in BN



10-708 – ©Carlos Guestrin 2006

30

The M-step for BNs



■ Maximization step:

$$\theta^{(t+1)} \leftarrow \arg \max_{\theta} \sum_{\mathbf{x}} Q^{(t+1)}(\mathbf{h} \mid \mathbf{o}) \log P(\mathbf{h}, \mathbf{o} \mid \theta)$$

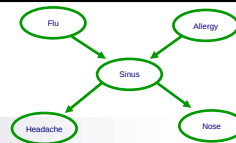
■ Use expected counts instead of counts:

- If learning requires Count(**h,o**)
- Use $E_{Q^{(t+1)}}[\text{Count}(\mathbf{h}, \mathbf{o})]$

10-708 - ©Carlos Guestrin 2006

31

M-step for each CPT



■ M-step decomposes per CPT

- Standard MLE:

$$P(X_i = x_i \mid \text{Pa}_{X_i} = \mathbf{z}) = \frac{\text{Count}(X_i = x_i, \text{Pa}_{X_i} = \mathbf{z})}{\text{Count}(\text{Pa}_{X_i} = \mathbf{z})}$$

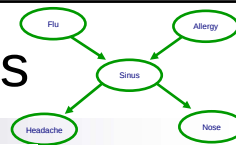
- M-step uses expected counts:

$$P(X_i = x_i \mid \text{Pa}_{X_i} = \mathbf{z}) = \frac{\text{ExCount}(X_i = x_i, \text{Pa}_{X_i} = \mathbf{z})}{\text{ExCount}(\text{Pa}_{X_i} = \mathbf{z})}$$

10-708 - ©Carlos Guestrin 2006

32

Computing expected counts



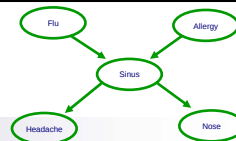
$$P(X_i = x_i | \text{Pa}_{X_i} = z) = \frac{\text{ExCount}(X_i = x_i, \text{Pa}_{X_i} = z)}{\text{ExCount}(\text{Pa}_{X_i} = z)}$$

- M-step requires expected counts:
 - For a set of vars **A**, must compute $\text{ExCount}(\mathbf{A}=\mathbf{a})$
 - Some of **A** in example j will be observed
 - denote by $\mathbf{A}_O = \mathbf{a}_O^{(j)}$
 - Some of **A** will be hidden
 - denote by $\mathbf{A}_H$
- Use inference (E-step computes expected counts):
 - $\text{ExCount}^{(t+1)}(\mathbf{A}_O = \mathbf{a}_O^{(j)}, \mathbf{A}_H = \mathbf{a}_H) \leftarrow P(\mathbf{A}_H = \mathbf{a}_H | \mathbf{A}_O = \mathbf{a}_O^{(j)}, \theta^{(t)})$

10-708 – ©Carlos Guestrin 2006

33

Data need not be hidden in the same way



- When data is fully observed
 - A data point is
- When data is partially observed
 - A data point is
- But unobserved variables can be different for different data points
 - e.g.,
- Same framework, just change definition of expected counts
 - $\text{ExCount}^{(t+1)}(\mathbf{A}_O = \mathbf{a}_O^{(j)}, \mathbf{A}_H = \mathbf{a}_H) \leftarrow P(\mathbf{A}_H = \mathbf{a}_H | \mathbf{A}_O = \mathbf{a}_O^{(j)}, \theta^{(t)})$

10-708 – ©Carlos Guestrin 2006

34

What you need to know

- EM for Bayes Nets
- E-step: inference computes expected counts
 - Only need expected counts over X_i and $\mathbf{Pa}_{x_i}$
- M-step: expected counts used to estimate parameters
- Hidden variables can change per datapoint
- Use labeled and unlabeled data → some data points are complete, some include hidden variables