

7-12

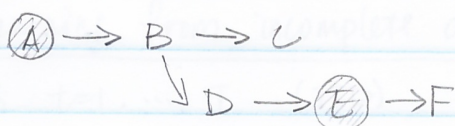
Review - Inference

- * Tractable in polytrees — singly connected networks ; no loops
- * Key tools -

Bayes Rule
Marginalization
Product Rule
Conditional Independence

(and conditionalized version of these)

Example



How to compute $P(B|E, A)$?

$$P(B|E, A) = \frac{P(E|B, A) P(B|A)}{P(E|A)}$$

Bayes rule = use it to rewrite LHS in terms of probabilities where nodes are conditioned on ancestors

$$P(E|B, A) = \sum_d P(E, D=d | B, A)$$

use marginalization to introduce parents b/c CPTs always express $P(\text{child} | \text{parents})$

$$P(E, D=d | B) = P(D=d | B) P(E | D=d, B)$$

product rule gets you even closer to CPTs, conditional independence b/c CPTs predict one node at time.

Review - Learning from complete data

* examples $t=1, 2, \dots, T$

variables $X_1, X_2, \dots, X_n$

data set $\{(X_1^t, X_2^t, \dots, X_n^t)\}_{t=1}^T$ (IID)

* choose CPTs to maximize log-likelihood $\mathcal{L} = \sum_t \log P(X_1^t, X_2^t, \dots, X_n^t)$

* ML estimates

$$P_{ML}(X_i = x | \text{par} = \pi) = \frac{\text{count}(X_i = x, \text{par} = \pi)}{\text{count}(\text{par} = \pi)} = \frac{\sum_{t=1}^T I(X_i^t = x) I(\text{par}_i = \pi)}{\sum_{t=1}^T I(\text{par}_i = \pi)}$$

Review - Learning from incomplete data

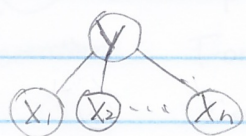
* examples $t=1, \dots, T$ (IID)

hidden nodes $H^{(t)}$

visible nodes $V^{(t)}$

Why useful?

Ex: document clustering



$Y \in \{1, 2, \dots, K\}$ document topic $H^{(t)} = \{Y\}$

$x \in \{0, 1\}$ Does i^{th} word in dictionary

appear in document? $V^{(t)} = \{X_1, X_2, \dots, X_n\}$

Question: can we learn such a model from a data set of unlabeled documents?

Data:

t	X_1	X_2	$\dots$	X_n	Y
1	0	1	$\dots$	1	?
2	0	0	$\dots$	1	?
3	0	1	$\dots$	0	?
$\vdots$					
T	1	0	$\dots$	0	?

* Choose CPTs to maximize log-likelihood

$$\begin{aligned}
 \mathcal{L} &= \sum_{t=1}^T \log P(V^{(t)}) \\
 &= \sum_{t=1}^T \log \sum_h P(V^{(t)}, H=h) \\
 &= \sum_{t=1}^T \log \sum_h \prod_{i=1}^n P(X_i = x_i | \text{pa}_i = \pi) \quad \left| \begin{array}{l} \text{evaluate these} \\ V^{(t)}, H=h \end{array} \right.
 \end{aligned}$$

/ Much more difficult to compute and optimize

No "closed-form" solution for CPTs.

Alternative = seek an iterative solution.

Iterative solution is known as EM algorithm.

* EM algorithm:

1) Initialization

Choose values for elements of CPTs in BN

(should be valid probabilities, could be poor guesses)

2) Iterative steps — repeat until convergence.

(E) E-step — Inference

expectation For each example $t=1, 2, \dots, T$ compute posterior distribution $P(H=h | V^{(t)})$

In particular, for each node, compute

$$P(X_i = x, \text{pa}_i = \pi | V = V^{(t)})$$

(M) M-step = update CPTs to increase log-likelihood.

maximization $P(X_i = x | \text{pa}_i = \pi) \leftarrow \frac{\sum_{t=1}^T P(X_i = x, \text{pa}_i = \pi | V^{(t)})}{\sum_{x^t} \left[\sum_{t=1}^T P(X_i = x, \text{pa}_i = \pi | V^{(t)}) \right]}$

Simplify =

① nodes w/ parents:

$$P(X_i = x | \text{pa}_i = \pi) \leftarrow \frac{\sum_{t=1}^T P(X_i = x, \text{pa}_i = \pi | V^{(t)})}{\sum_{t=1}^T P(\text{pa}_i = \pi | V^{(t)})}$$

② root node (no parents)

$$P(X_i = x) \leftarrow \frac{1}{T} \sum_{t=1}^T P(X_i = x | V^{(t)})$$

Note:

- 1) RHS of these update depends on current values of CPTs
- 2) Inference is key subroutine of learning.

Intuition:

- Use posterior distribution $P(H | V^{(t)})$ to "fill in" missing values of hidden nodes
- expected statistics of posterior distribution $P(H | V^{(t)})$ substitute for observed counts that we had in complete data case

- Compare to ML estimate in complete data case

$$\begin{aligned} P_{ML}(X_i = x | p_{ai} = \pi) &= \frac{\text{count}(X_i = x, p_{ai} = \pi)}{\text{count}(p_{ai} = \pi)} \\ &= \frac{\sum_{t=1}^T I(X_i^{(t)} = x) I(p_{ai}^{(t)} = \pi)}{\sum_{t=1}^T I(p_{ai}^{(t)} = \pi)} \end{aligned}$$

special case of EM when all data is observed.

Key properties of Expectation Maximization (EM) algorithm

- 1) Monotonic convergence

Each iteration of EM improves log-likelihood

$$L = \sum_{t=1}^T \log P(V^{(t)})$$

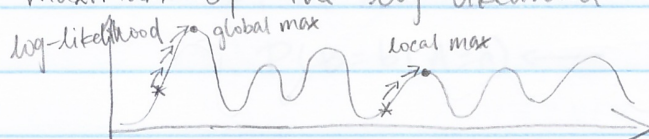
If L_k is log-likelihood at k^{th} iteration, then

$$L_k \geq L_{k-1}$$

Really helpful debugging diagnostic

- 2) Converges in general to local (but not global)

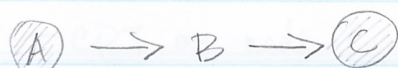
maximum of the log-likelihood $L = \sum_{t=1}^T \log P(V^{(t)})$



Final output of EM depends on model initialization

3) No tuning parameters, no learning rates, no back-tracking

Example



A, C are observed (visible) nodes

B is hidden node

— First, consider inference in this BN.

$$\begin{aligned}
 \text{Posterior: } P(B|A, C) &= \frac{P(C|B, A) P(B|A)}{P(C|A)} \quad \begin{array}{l} \text{Bayes} \\ \text{Rule /} \\ \text{condition} \\ \text{Independ} \end{array} \\
 &= \frac{P(C|B) P(B|A)}{\sum_{b'} P(C|B=b') P(B=b'|A)} \quad \text{normalization}
 \end{aligned}$$

Bottom line: We can compute $P(b|a, c)$ in terms of model's CPTs.

* How do we learn from incomplete data $\{(a_t, c_t)\}_{t=1}^T$

— EM algorithm

$$\begin{aligned}
 \text{Log-likelihood } \mathcal{L} &= \sum_t \log P(A=a_t, C=c_t) \\
 &= \sum_t \log \sum_b P(A=a_t, B=b, C=c_t) \\
 &= \sum_t \log \sum_b P(a_t) P(b|a_t) P(c_t|a_t)
 \end{aligned}$$

Reminder: general step of EM algorithm

$$P(X_i = x | \{p_{a_i} = \pi\}) \leftarrow \frac{\sum_{t=1}^T P(X_i = x, p_{a_i} = \pi | V^{(t)})}{\sum_{t=1}^T P(p_{a_i} = \pi | V^{(t)})}$$

Now specialize to our example $V^{(t)} = (a_t, c_t)$

— CPT at node B:

$$P(B=b | A=a) \leftarrow \frac{\sum_{t=1}^T P(B=b, A=a | A=a_t, C=c_t)}{\sum_{t=1}^T P(A=a | A=a_t, C=c_t)}$$

Simplify:

$$P(B=b | A=a) \leftarrow$$

$$\frac{\sum_{t=1}^T I(a, a_t) P(b | a_t, c_t)}{\sum_{t=1}^T I(a, a_t)}$$

We've computed already
in terms of CPTs

↑
read off
from data

- CPT at node C

$$P(C=c | B=b) \leftarrow$$

$$\frac{\sum_{t=1}^T P(B=b, C=c | a_t, c_t)}{\sum_{t=1}^T P(B=b | a_t, c_t)}$$

Simplify:

$$P(C=c | B=b) \leftarrow$$

$$\frac{\sum_{t=1}^T I(c, c_t) P(b | a_t, c_t)}{\sum_{t=1}^T P(b | a_t, c_t)}$$

Example Noisy-OR Model

$X_1 \quad X_2 \quad \dots \quad X_n$

diseases $X_i \in \{0, 1\}$



symptom $Y \in \{0, 1\}$

$$P(Y=1 | X_1, X_2, \dots, X_n) = 1 - \prod_{i=1}^n (1 - p_i)^{X_i}$$

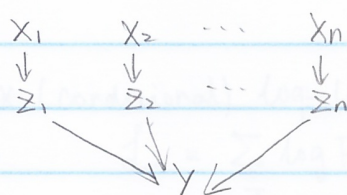
with $p_i \in [0, 1]$

* From complete data $\{(\vec{X}_t, Y_t)\}_{t=1}^T$ how do we estimate $p_i \in [0, 1]$?

Note: Noisy-OR is a "parametric" model of CPT; there is no simple, closed-form ML estimate for $p_i \in [0, 1]$ that maximizes.

$$\sum_t \log P(Y = y^{(t)} | X_1 = x_1^{(t)}, \dots, X_n = x_n^{(t)})$$

* Alternative formulation =



$X_i \in \{0, 1\}$

$Z_i \in \{0, 1\}$

$Y \in \{0, 1\}$

noisy copy of parent

$$P(Z_i=0 | X_i=0) = 1$$

$$P(Z_i=1 | X_i=0) = 0$$

$$P(Z_i=1 | X_i=1) = p_i$$

$$Y = \text{OR}(Z_1, Z_2, \dots, Z_n)$$

$$P(Y=1 | Z_1, \dots, Z_n) = \begin{cases} 1 & \text{if } \sum_{i=1}^n Z_i > 0 \\ 0 & \text{if } Z_i = 0 \text{ for all } i. \end{cases}$$

$$P(z_i=0 | x_i=0) = 1, \quad P(z_i=1 | x_i=0) = 0$$

$$P(z_i=1 | x_i=1) = p_i$$

Equivalently:

$$P(z_i=0 | x_i) = (1-p_i)^{x_i} = \begin{cases} 1-p_i & \text{if } x_i=1 \\ 1 & \text{if } x_i=0 \end{cases}$$

What is $P(Y=1 | \vec{X})$ in this new model?

$$\begin{aligned} P(Y=1 | \vec{X}) &= \sum_{\vec{z} \in \{0,1\}^n} P(Y=1, \vec{z} | \vec{X}) \quad \text{marginalization} \\ &= \sum_{\vec{z}} P(Y=1 | \vec{z}, \vec{X}) P(\vec{z} | \vec{X}) \quad \text{product rule} \\ &= \sum_{\vec{z} \in \{0,1\}^n} P(Y=1 | \vec{z}) P(\vec{z} | \vec{X}) \quad \text{cond. independence} \\ &= \sum_{\vec{z} \neq (0,0,0,\dots,0)} P(\vec{z} | \vec{X}) + \underbrace{P(Y=1 | \vec{z}=\vec{0})}_{0} P(\vec{z}=\vec{0} | \vec{X}) \\ &\quad \text{logical-OR, } Y = \text{OR}(\vec{z}) \\ &= 1 - P(\vec{z}=\vec{0} | \vec{X}) \quad \text{normalization} \\ &= 1 - \prod_{i=1}^n P(z_i=0 | x_i) \quad \text{product rule c.i.} \\ &= 1 - \prod_{i=1}^n (1-p_i)^{x_i} \end{aligned}$$

Key insight:

alternative BN has same expression for $P(Y=1 | \vec{X})$!!

→ estimate p_i parameters for noisy-OR CPT

using EM in "expanded" network

First, let's do inference in this model: $\begin{cases} 0 & \text{if } x_i=0 \\ p_i & \text{if } x_i=1 \end{cases}$

$$\begin{aligned} P(z_i=1 | \vec{X}, Y) &= \frac{P(Y | z_i=1, \vec{X}) P(z_i=1 | \vec{X})}{P(Y | \vec{X})} \\ &\quad \begin{cases} 0 & \text{if } Y=0 \\ 1 & \text{if } Y=1 \end{cases} \quad \text{only case if } Y=1, \text{ b/c otherwise numerator vanish.} \\ &= \frac{Y \cdot p_i \cdot x_i}{1 - \prod_{i=1}^n (1-p_i)^{x_i}} \\ &\quad \text{true for } x_i \in \{0,1\}, Y \in \{0,1\} \end{aligned}$$

* (conditional) log-likelihood of data $\{(\vec{x}_t, y_t)\}_{t=1}^T$

$$\begin{aligned} \mathcal{L} &= \sum_{t=1}^T \log P(Y_t | \vec{X}_t) \\ &= \sum_{t=1}^T [Y_t \log P(Y=1 | \vec{X}_t) + (1-Y_t) \log P(Y=0 | \vec{X}_t)] \\ &= \sum_{t=1}^T [Y_t \log (1 - \prod_{i=1}^n (1-p_i)^{x_i^{(t)}}) + (1-Y_t) \log \prod_{i=1}^n (1-p_i)^{x_i^{(t)}}] \end{aligned}$$

$$\mathcal{L}(p_1, p_2, \dots, p_n) = \sum_{t=1}^T y_t \log \left(1 - \prod_{i=1}^n (1-p_i)^{x_{it}} \right) + \sum_{t=1}^T \sum_{i=1}^n (1-y_t) x_{it} \log(1-p_i)$$

y_t, x_{it} are data observed

p_i is parameter — how to choose?

Note-

Complicated nonlinear expression of p_i .

How to optimize? EM to the rescue

Shorthand: let T = total # of examples

$$\text{let } T_i = \sum_{t=1}^T x_{it} = \text{count}(x_i=1)$$

EM update rule for $p_i = P(z_i=1 | x_i=1)$

$$P(z_i=1 | x_i=1) \leftarrow \frac{\sum_{t=1}^T P(z_i=1, x_{it}=1 | \vec{x}_t, y_t)}{\sum_{t=1}^T P(x_{it}=1 | \vec{x}_t, y_t)}$$

Simplify =

$$P(z_i=1 | x_i=1) \leftarrow \frac{\sum_{t=1}^T \overbrace{I(x_{it}, 1)}^{x_{it}} \overbrace{P(z_i=1 | \vec{x}_t, y_t)}^{\text{have already from Bayes rule}}}{\underbrace{\sum_{t=1}^T I(x_{it}, 1)}_{T_i}}$$

EM update
for Noisy-OR

$$p_i \leftarrow \frac{p_i}{T_i} \sum_{t=1}^T \frac{y_t x_{it}}{1 - \prod_{i=1}^n (1-p_i)^{x_{it}}}$$

$$\frac{y_t \cdot p_i \cdot x_{it}}{1 - \prod_{i=1}^n (1-p_i)^{x_{it}}}$$

This update rule, applied in parallel to all $\{p_i\}_{i=1}^n$,
will monotonically increase $\mathcal{L} = \sum_t \log P(y_t | \vec{x}_t)$

Recall from HW4.

$P_{\text{unigram}}(w)$

$P_{\text{bigram}}(w'|w)$

$$P_{\text{mixture}}(w'|w) = (1-\lambda) P_{\text{unigram}}(w') + \lambda P_{\text{bigram}}(w'|w) \text{ where } \lambda \in [0, 1]$$

More generally -

$$P_{\text{mixture}}(w'|w) = (1-\lambda(w)) P_{\text{unigram}}(w') + \lambda(w) P_{\text{bigram}}(w'|w) \text{ where}$$

Next Lecture: This is also a hidden variable model!

$$\lambda(w) \in [0, 1]$$

We can use EM to estimate λ .