

CSE 150A-250A AI: Probabilistic Models

Lecture 13

Fall 2025

Trevor Bonjour
Department of Computer Science and Engineering
University of California, San Diego

Slides adapted from previous versions of the course (Prof. Lawrence, Prof. Alvarado, Prof Berg-Kirkpatrick)

Agenda

Review

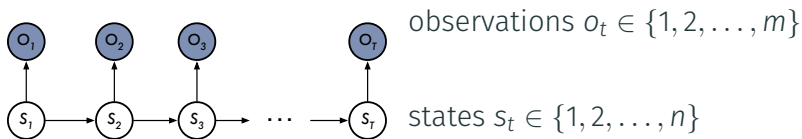
Learning in HMMs

Backward Algorithm

Review

Hidden Markov models

- Belief network



- Parameters

$$a_{ij} = P(S_{t+1}=j|S_t=i) \quad \text{transition matrix}$$

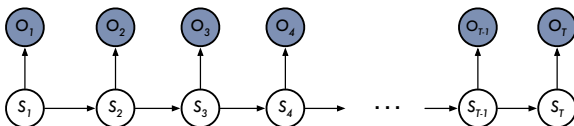
$$b_{ik} = P(O_t=k|S_t=i) \quad \text{emission matrix}$$

$$\pi_i = P(S_1=i) \quad \text{initial state distribution}$$

- Notation

Sometimes we'll write $b_i(k) = b_{ik}$ to avoid double subscripts.

Key computations in HMMs



Inference

1. How to compute the likelihood $P(o_1, o_2, \dots, o_T)$?
2. How to compute the most likely hidden states $\operatorname{argmax}_{\vec{s}} P(\vec{s} | \vec{o})$?
3. How to update beliefs by computing $P(s_t | o_1, o_2, \dots, o_t)$?

Learning

How to estimate parameters $\{\pi_i, a_{ij}, b_{ik}\}$ that maximize the log-likelihood of observed sequences?

Forward Algorithm

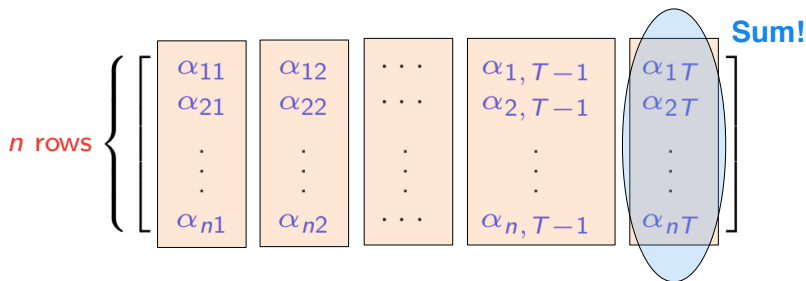
For a particular sequence of observations $\{o_1, o_2, \dots, o_T\}$, define the matrix with elements:

$$\alpha_{it} = P(o_1, o_2, \dots, o_t, S_t=i) \quad \begin{matrix} n \text{ rows} \end{matrix} \left\{ \begin{bmatrix} \alpha_{11} & \alpha_{12} & \cdots & \alpha_{1,T-1} & \alpha_{1T} \\ \alpha_{21} & \alpha_{22} & \cdots & \alpha_{2,T-1} & \alpha_{2T} \\ \vdots & \vdots & \vdots & \vdots & \vdots \\ \alpha_{n1} & \alpha_{n2} & \cdots & \alpha_{n,T-1} & \alpha_{nT} \end{bmatrix} \right.$$

The forward algorithm fills in the matrix of α_{it} elements one column at a time:

$$\begin{aligned} \alpha_{i1} &= \pi_i b_i(o_1) \\ \alpha_{j,t+1} &= \sum_{i=1}^n \alpha_{it} a_{ij} b_j(o_{t+1}) \end{aligned}$$

Computing the likelihood $P(o_1, o_2, \dots, o_T)$



$$P(o_1, o_2, \dots, o_T)$$

$$= \sum_{i=1}^n P(o_1, o_2, \dots, o_T, S_T = i)$$

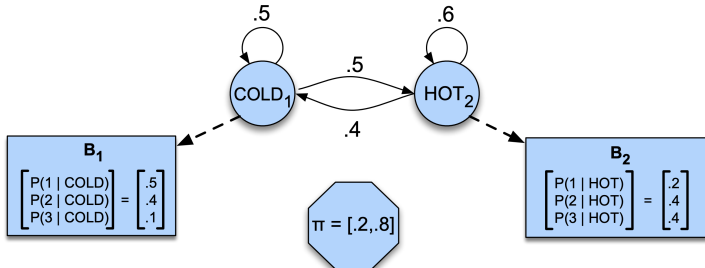
marginalization

$$= \sum_{i=1}^n \alpha_{iT}$$

sum of last column

Example¹

- Two hidden states (Weather): $\{H, C\}$
- Observations (Ice creams): $\{1, 2, 3\}$

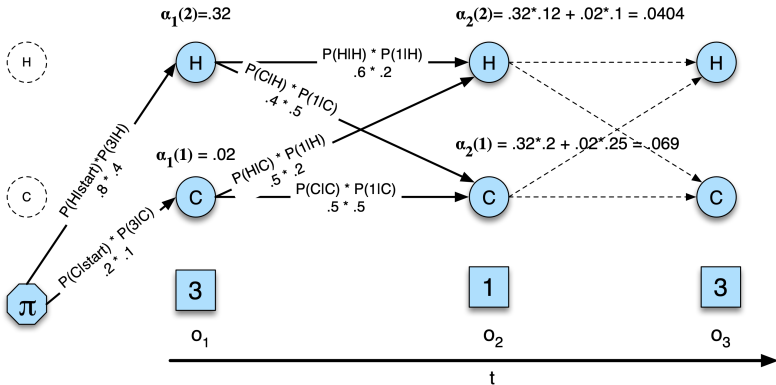


¹Eisner, J. 2002. An interactive spreadsheet for teaching the forward-backward algorithm.

Example - Forward Algorithm

$$\alpha_{i1} = \pi_i b_i(o_1)$$

$$\alpha_{j,t+1} = \sum_{i=1}^n \alpha_{it} a_{ij} b_j(o_{t+1})$$



Viterbi Algorithm

$$\{s_1^*, s_2^*, \dots, s_T^*\}$$

$$= \operatorname{argmax}_{s_1, s_2, \dots, s_T} P(s_1, s_2, \dots, s_T | o_1, o_2, \dots, o_T)$$

$$= \operatorname{argmax}_{s_1, s_2, \dots, s_T} \log P(s_1, s_2, \dots, s_T, o_1, o_2, \dots, o_T)$$

For a particular sequence of observations $\{o_1, o_2, \dots, o_T\}$, we define the following matrix:

$$\ell_{it}^* = \max_{s_1, s_2, \dots, s_{t-1}} \log P(s_1, s_2, \dots, s_{t-1}, s_t = i, o_1, o_2, \dots, o_t)$$

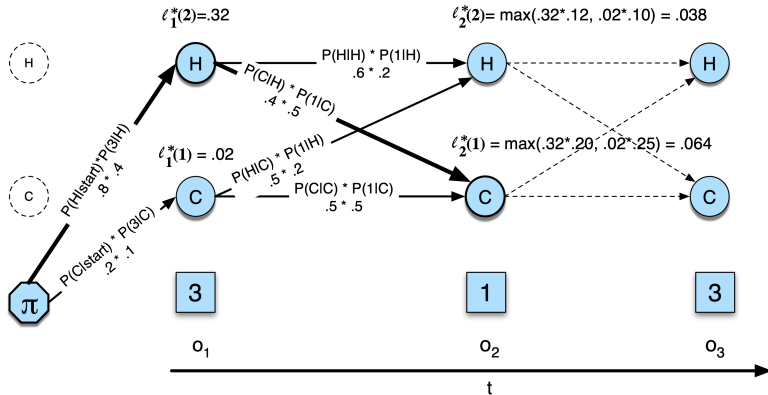
Viterbi Algorithm

We compute the matrix ℓ^* one column at a time, from left to right:

$$\begin{aligned}\ell_{i1}^* &= \log \pi_i + \log b_i(o_1) \\ \ell_{j,t+1}^* &= \max_i \left[\ell_{it}^* + \log a_{ij} \right] + \log b_j(o_{t+1})\end{aligned}$$

$$n \text{ rows} \left\{ \begin{bmatrix} \ell_{11}^* \\ \ell_{21}^* \\ \vdots \\ \ell_{n1}^* \end{bmatrix} \begin{bmatrix} \ell_{12}^* \\ \ell_{22}^* \\ \vdots \\ \ell_{n2}^* \end{bmatrix} \begin{bmatrix} \dots \\ \dots \\ \vdots \\ \dots \end{bmatrix} \begin{bmatrix} \ell_{1,T-1}^* \\ \ell_{2,T-1}^* \\ \vdots \\ \ell_{n,T-1}^* \end{bmatrix} \begin{bmatrix} \ell_{1T}^* \\ \ell_{2T}^* \\ \vdots \\ \ell_{nT}^* \end{bmatrix} \right\}$$

Example - Viterbi (Fill ℓ^*)



Computing $\{s_1^*, s_2^*, \dots, s_T^*\}$

- Form one more matrix:

$$\Phi_{t+1}(j) = \arg \max_i \left[\ell_{it}^* + \log a_{ij} \right]$$

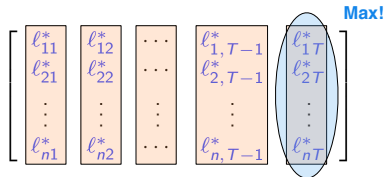
- Compute the most likely states by **backtracking**:

$$s_T^* = \arg \max_i \left[\ell_{iT}^* \right]$$

for $t = T-1$ to 1

$$s_t^* = \Phi_{t+1}(s_{t+1}^*)$$

end



Summary of Viterbi algorithm

- Fill ℓ^* matrix from left to right:

$$\boxed{t = 1} \quad \ell_{i1}^* = \log \pi_i + \log b_i(o_1)$$

$$\boxed{t > 1} \quad \ell_{j,t+1}^* = \max_i \left[\ell_{it}^* + \log a_{ij} \right] + \log b_j(o_{t+1})$$

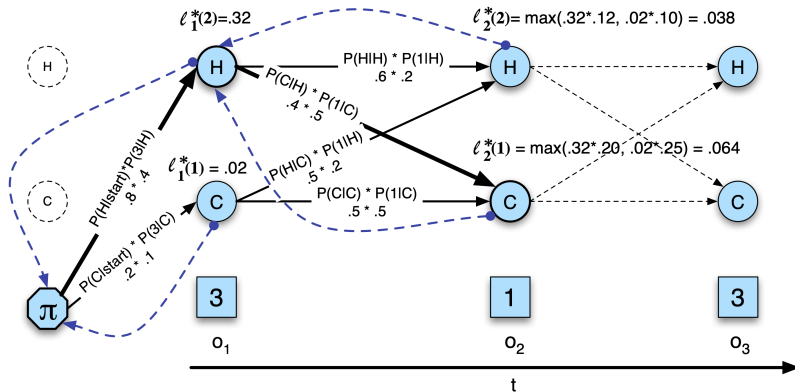
- Backtrack through ℓ^* from right to left:

$$\boxed{t = T} \quad s_T^* = \operatorname{argmax}_i \left[\ell_{iT}^* \right]$$

$$\boxed{t < T} \quad s_t^* = \operatorname{argmax}_i \left[\ell_{it}^* + \log a_{is_{t+1}^*} \right]$$

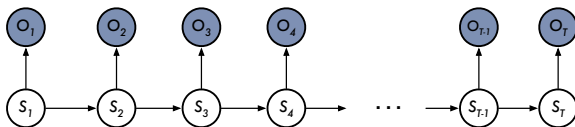
Sometimes $\{s_1^*, s_2^*, \dots, s_T^*\}$ is called the **Viterbi path**.

Example - Viterbi (Backtrack through ℓ^*)



Learning in HMMs

Learning in HMMs



Given: one or more sequences of observations $\{o_1, o_2, \dots, o_T\}$.

For simplicity, we'll assume just one.

Goal: estimate $\{\pi_i, a_{ij}, b_{ik}\}$ to maximize $P(o_1, o_2, \dots, o_T)$,
the likelihood of the observed data.

Assume: the cardinality n of the hidden state space is fixed.

$$s_t \in \{1, 2, \dots, n\}$$

How do we estimate $\{\pi_i, a_{ij}, b_{ik}\}$?

EM Algorithm!

How EM works in general:

To re-estimate $P(X_i=x|\text{pa}_i=\pi)$ in the M-step,
we must compute $P(X_i=x, \text{pa}_i=\pi|V)$ in the E-step.

EM algorithm for HMMs

- CPTs to re-estimate:

$$\pi_i = P(S_1=i)$$

$$a_{ij} = P(S_{t+1}=j|S_t=i)$$

$$b_{ik} = P(O_t=k|S_t=i)$$

- E-step in HMMs must compute:

$$\begin{aligned} &P(S_1=i|o_1, o_2, \dots, o_T) \\ &P(S_{t+1}=j, S_t=i|o_1, o_2, \dots, o_T) \quad \underbrace{\text{special case of below } (t=1)} \\ &P(O_t=k, S_t=i|o_1, o_2, \dots, o_T) = l(o_t, k) P(S_t=i|o_1, o_2, \dots, o_T) \end{aligned}$$

How to efficiently compute these posteriors?

Computing $P(S_t=i|o_1,\dots,o_T)$

$$P(S_t=i|o_1,\dots,o_T) = \frac{P(S_t=i, o_1, o_2, \dots, o_T)}{P(o_1, o_2, \dots, o_T)}$$

product rule

- Numerator

$$\begin{aligned} P(S_t=i, o_1, o_2, \dots, o_T) \\ &= P(o_1, \dots, o_t, S_t=i) P(o_{t+1}, \dots, o_T | S_t=i, o_1, \dots, o_t) \quad \text{product rule} \\ &= P(o_1, \dots, o_t, S_t=i) P(o_{t+1}, \dots, o_T | S_t=i) \quad \text{conditional independence} \\ &= \alpha_{it} P(o_{t+1}, \dots, o_T | S_t=i) \end{aligned}$$

Backward Algorithm

We need one more matrix ...

Analogous to $\alpha_{it} = P(o_1, o_2, \dots, o_t, S_t = i),$

define $\beta_{it} = P(o_{t+1}, o_{t+2}, \dots, o_T | S_t = i).$

$$n \text{ rows} \left\{ \begin{bmatrix} \beta_{11} & \beta_{12} & \cdots & \beta_{1,T-1} & \beta_{1T} \\ \beta_{21} & \beta_{22} & \cdots & \beta_{2,T-1} & \beta_{2T} \\ \vdots & \vdots & \vdots & \vdots & \vdots \\ \beta_{n1} & \beta_{n2} & \cdots & \beta_{n,T-1} & \beta_{nT} \end{bmatrix} \right.$$

Understand the **differences** between these matrices:

- α_{it} predicts observations up to and including time t .
- β_{it} predicts observations from **time $t + 1$** to time T .

Computing $\beta_{it} = P(o_{t+1}, o_{t+2}, \dots, o_T | S_t = i)$

- Last column ($t = T$)

$$\beta_{iT} = P(\text{---} | S_T = i) \quad \text{What does this mean?}$$

Note: β_{it} computes the probability of the future given $S_t = i$.

But we don't see *any* observations beyond time T .
Put another way, the future after time T is unspecified.

What is the probability of some unspecified future occurring?

By definition, we set:

$$\beta_{iT} = 1 \quad \text{for all } i \in \{1, 2, \dots, n\}$$

Computing $\beta_{it} = P(o_{t+1}, o_{t+2}, \dots, o_T | S_t = i)$

- Previous columns ($t < T$)

$$\begin{aligned}\beta_{it} &= P(o_{t+1}, o_{t+2}, \dots, o_T | S_t = i) \\&= \sum_{j=1}^n P(S_{t+1} = j, o_{t+1}, o_{t+2}, \dots, o_T | S_t = i) \quad \boxed{\text{marginalization}} \\&= \sum_{j=1}^n \left[P(S_{t+1} = j | S_t = i) \cdot \right. \\&\quad \left. P(o_{t+1} | S_t = i, S_{t+1} = j) \cdot \right. \\&\quad \left. P(o_{t+2}, \dots, o_T | S_t = i, S_{t+1} = j, o_{t+1}) \right] \quad \boxed{\text{product rule}} \\&= \sum_{j=1}^n \left[P(S_{t+1} = j | S_t = i) P(o_{t+1} | S_{t+1} = j) P(o_{t+2}, \dots, o_T | S_{t+1} = j) \right] \quad \boxed{\text{CI}} \\&= \sum_{j=1}^n a_{ij} b_j(o_{t+1}) \beta_{j,t+1} \quad \boxed{\text{CPTs}}\end{aligned}$$

Backward algorithm

The backward algorithm fills in the matrix of β_{it} elements one column at a time:

Learning in HMMs - EM Algorithm

Computing $P(S_t=i|o_1, \dots, o_T)$

$$P(S_t=i|o_1, \dots, o_T) = \frac{P(S_t=i, o_1, o_2, \dots, o_T)}{P(o_1, o_2, \dots, o_T)}$$

product rule

- Numerator

$$P(S_t=i, o_1, o_2, \dots, o_T)$$

$$= P(o_1, \dots, o_t, S_t=i) P(o_{t+1}, \dots, o_T | S_t=i, o_1, \dots, o_t)$$

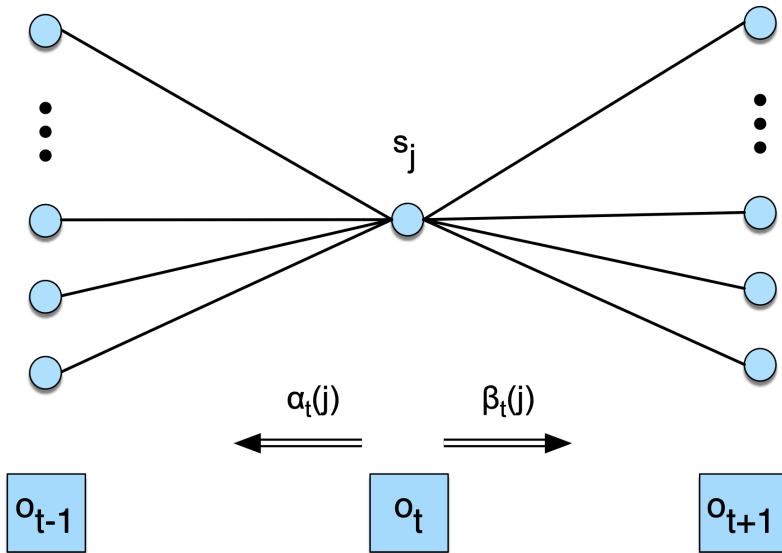
product rule

$$= P(o_1, \dots, o_t, S_t=i) P(o_{t+1}, \dots, o_T | S_t=i)$$

conditional independence

$$= \alpha_{it} P(o_{t+1}, \dots, o_T | S_t=i)$$

$$= \alpha_{it} \beta_{it}$$



Computing $P(S_t=i|o_1, \dots, o_T)$

$$P(S_t=i|o_1, \dots, o_T) = \frac{P(S_t=i, o_1, o_2, \dots, o_T)}{P(o_1, o_2, \dots, o_T)}$$

product rule

- Numerator

$$\begin{aligned} P(S_t=i, o_1, o_2, \dots, o_T) &= P(o_1, \dots, o_t, S_t=i) P(o_{t+1}, \dots, o_T | S_t=i, o_1, \dots, o_t) \quad \text{product rule} \\ &= P(o_1, \dots, o_t, S_t=i) P(o_{t+1}, \dots, o_T | S_t=i) \quad \text{conditional independence} \\ &= \alpha_{it} \beta_{it} \end{aligned}$$

- Denominator

$$\begin{aligned} P(o_1, o_2, \dots, o_T) &= \sum_k P(S_t=k, o_1, o_2, \dots, o_T) \quad \text{marginalization} \\ &= \sum_k \alpha_{kt} \beta_{kt} \quad \text{by above} \quad \text{Note: this holds for all values of } t. \end{aligned}$$

Computing $P(S_t=i, S_{t+1}=j|o_1, \dots, o_T)$

$$P(S_t=i, S_{t+1}=j|o_1, \dots, o_T) = \frac{P(S_t=i, S_{t+1}=j, o_1, o_2, \dots, o_T)}{\underbrace{P(o_1, o_2, \dots, o_T)}_{\text{already computed}}}$$

product rule

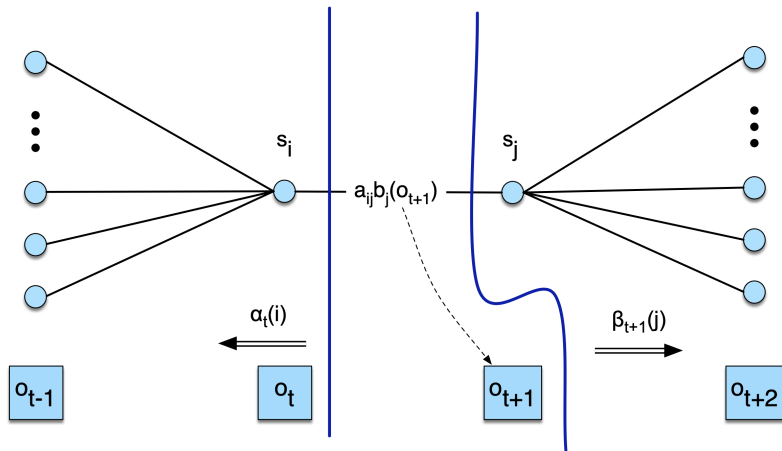
Express the numerator $P(S_t=i, S_{t+1}=j, o_1, o_2, \dots, o_T)$ in terms of α , β , and parameters of the model a , b .

Computing $P(S_t=i, S_{t+1}=j|o_1, \dots, o_T)$

$$P(S_t=i, S_{t+1}=j, o_1, \dots, o_T) = ?$$

How are you progressing?

- A. Not sure where to start.
- B. Making progress, but not there yet.
- C. I got an answer, but I am not sure if it is right.
- D. I finished and feel pretty confident about it.
- E. I got lost and wandered off into virtual space.



- Summary of E-step:

$$P(S_t = i | o_1, \dots, o_T) = \frac{\alpha_{it} \beta_{it}}{\sum_j \alpha_{jt} \beta_{jt}}$$
$$P(S_t = i, S_{t+1} = j | o_1, \dots, o_T) = \frac{\alpha_{it} a_{ij} b_j(o_{t+1}) \beta_{j,t+1}}{\sum_k \alpha_{kt} \beta_{kt}}$$

EM algorithm for HMMs

- CPTs to re-estimate:

$$\pi_i = P(S_1=i)$$

$$a_{ij} = P(S_{t+1}=j|S_t=i)$$

$$b_{ik} = P(O_t=k|S_t=i)$$

- M-step updates:

$$\pi_i \leftarrow P(S_1=i|o_1, o_2, \dots, o_T)$$

$$a_{ij} \leftarrow \frac{\sum_t P(S_{t+1}=j, S_t=i|o_1, o_2, \dots, o_T)}{\sum_t P(S_t=i|o_1, o_2, \dots, o_T)}$$

$$b_{ik} \leftarrow \frac{\sum_t I(o_t, k) P(S_t=i|o_1, o_2, \dots, o_T)}{\sum_t P(S_t=i|o_1, o_2, \dots, o_T)}$$

(for one sequence of observations)

What is the running time for **each** iteration of the update?

- A. $O(n)$
- B. $O(n^2)$
- C. $O(Tn^2)$
- D. $O(T^2n^4)$
- E. $O(n^T)$

Time complexity of HMM computations

T	length of observation sequence $(o_1, o_2, \dots, o_T)$
n	cardinality of state space $s_t \in \{1, 2, \dots, n\}$
m	cardinality of observation space $o_t \in \{1, 2, \dots, m\}$

- All of the following computations are $O(n^2T)$:

(a) computing the likelihood $P(o_1, o_2, \dots, o_T)$

(b) decoding $\operatorname{argmax}_{s_1, \dots, s_T} P(s_1, \dots, s_T | o_1, \dots, o_T)$

(c) re-estimating $\{\pi_i, a_{ij}, b_{ik}\}$ by **one update of EM**

(d) updating beliefs $P(S_t = i | o_1, \dots, o_t)$ **for T steps**

That's all folks!