

CSE 150A-250A AI: Probabilistic Models

Lecture 12

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

Viterbi algorithm

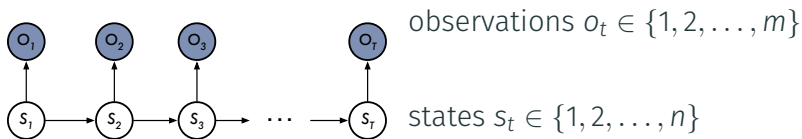
Learning in HMMs

Backward Algorithm

Review

Hidden Markov models

- Belief network



- Parameters

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

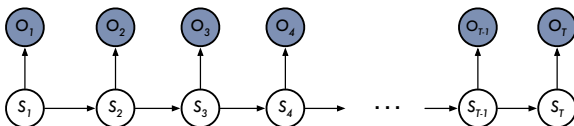
$$b_{ik} = P(O_t=k|S_t=i) \quad \text{emission matrix}$$

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

- 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

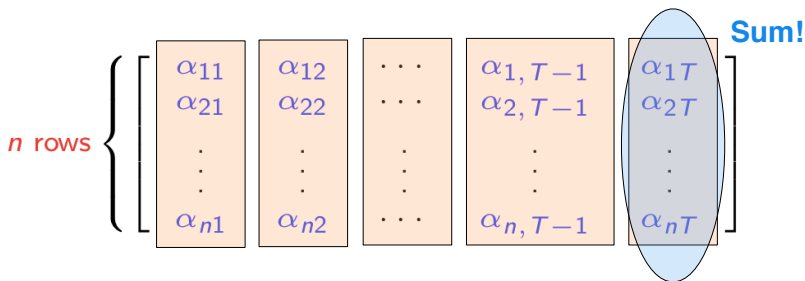
For a particular sequence of observations $\{o_1, o_2, \dots, o_T\}$, define the matrix with elements:

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

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



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

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

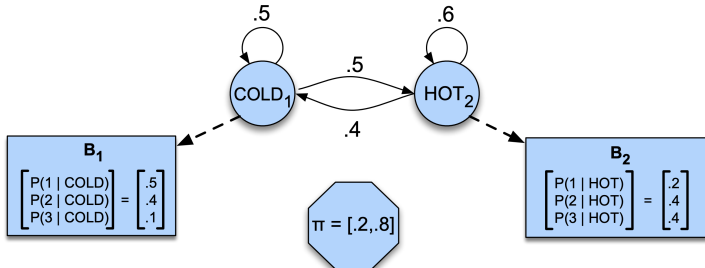
marginalization

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

sum of last column

Example¹

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

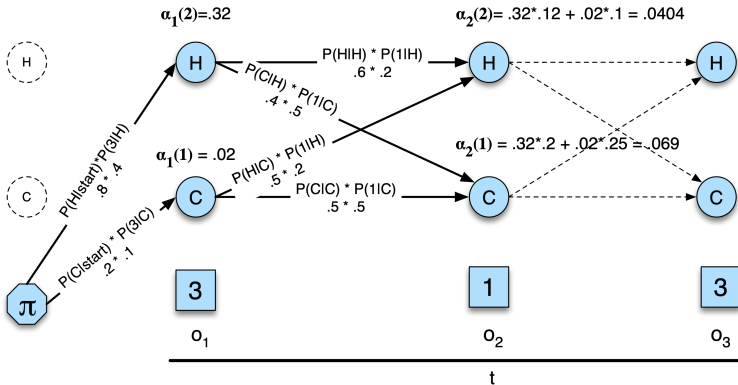


¹Eisner, J. 2002. An interactive spreadsheet for teaching the forward-backward algorithm.

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

The most likely sequence of hidden states

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

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

$$= \operatorname{argmax}_{s_1, s_2, \dots, s_T} \left[\frac{P(s_1, s_2, \dots, s_T, o_1, o_2, \dots, o_T)}{P(o_1, o_2, \dots, o_T)} \right] \quad \begin{array}{|c|} \hline \text{product} \\ \hline \text{rule} \\ \hline \end{array}$$

$$= \operatorname{argmax}_{s_1, s_2, \dots, s_T} P(s_1, s_2, \dots, s_T, o_1, o_2, \dots, o_T) \quad \begin{array}{|c|} \hline \text{no states in} \\ \hline \text{denominator} \\ \hline \end{array}$$

$$= \operatorname{argmax}_{s_1, s_2, \dots, s_T} \log P(s_1, s_2, \dots, s_T, o_1, o_2, \dots, o_T) \quad \begin{array}{|c|} \hline \text{log is} \\ \hline \text{monotonic} \\ \hline \end{array}$$

The matrix ℓ^*

- Definition

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

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

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

- Intuition

ℓ_{it}^*	log-probability of the t -step path of hidden states $s_1, s_2, \dots, s_t$ that best explains the observations $o_1, o_2, \dots, o_t$ and ends at state $s_t = i$ at time t
---------------	--

Computing the matrix ℓ^*

$$\ell_{it}^* = \max_{S_1, S_2, \dots, S_{t-1}} \log P(S_1, S_2, \dots, S_{t-1}, S_t = i, o_1, o_2, \dots, o_t)$$

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

- First column ($t = 1$)

$$\ell_{i1}^* = \log P(S_1 = i, o_1)$$

$$= \log \left[P(S_1 = i) P(o_1 | S_1 = i) \right]$$

product rule

$$= \log \pi_i + \log b_i(o_1)$$

CPTs

Computing the matrix ℓ^*

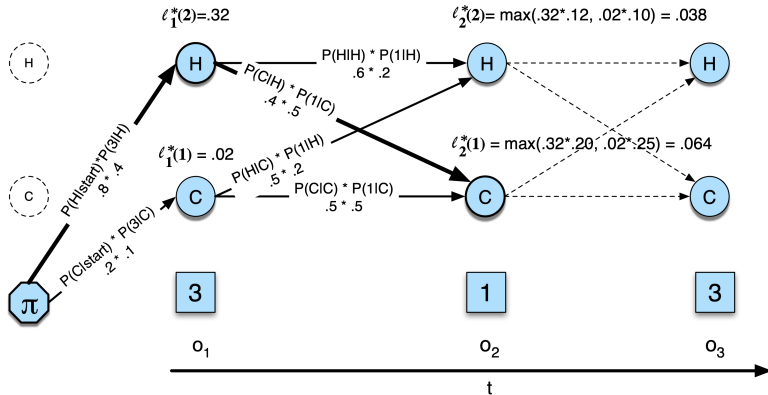
- Next columns ($t > 1$)

$$\begin{aligned}\ell_{j,t+1}^* &= \max_{s_1, \dots, s_t} \log P(s_1, \dots, s_t, S_{t+1}=j, o_1, \dots, o_{t+1}) \\&= \max_i \max_{s_1, \dots, s_{t-1}} \log P(s_1, \dots, s_{t-1}, S_t=i, S_{t+1}=j, o_1, \dots, o_{t+1}) \\&= \max_i \max_{s_1, \dots, s_{t-1}} \log \left[P(s_1, \dots, s_{t-1}, S_t=i, o_1, \dots, o_t) \cdot \boxed{\text{product rule}} \right. \\&\quad \left. P(S_{t+1}=j | s_1, \dots, s_{t-1}, S_t=i, o_1, \dots, o_t) \cdot \right. \\&\quad \left. P(o_{t+1} | s_1, \dots, s_{t-1}, S_t=i, S_{t+1}=j, o_1, \dots, o_t) \right] \\&= \max_i \max_{s_1, \dots, s_{t-1}} \log \left[P(s_1, \dots, s_{t-1}, S_t=i, o_1, \dots, o_t) \cdot \right. \\&\quad \left. P(S_{t+1}=j | S_t=i) \cdot P(o_{t+1} | S_{t+1}=j) \right] \boxed{\text{CI}} \\&= \max_i \max_{s_1, \dots, s_{t-1}} \left[\log P(s_1, \dots, s_{t-1}, S_t=i, o_1, \dots, o_t) + \log a_{ij} + \log b_j(o_{t+1}) \right] \\&= \max_i \left[\ell_{it}^* + \log a_{ij} \right] + \log b_j(o_{t+1})\end{aligned}$$

Summary

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

Example - Viterbi (Fill ℓ^*)



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

- Form one more matrix:

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

Intuitively, given the observations $o_1, o_2, \dots, o_{t+1}$, record the most likely state at time t given that $S_{t+1}=j$.

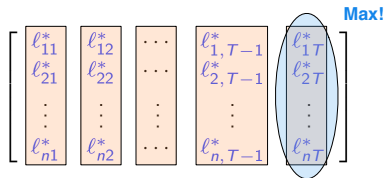
- Compute the most likely states by **backtracking**:

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

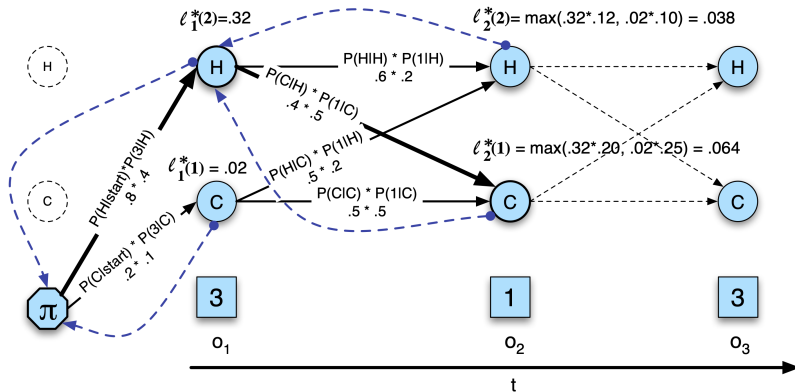
for $t = T-1$ **to** 1

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

end



Example - Viterbi (Backtrack through ℓ^*)



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:

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

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

- Where you've seen this before:

This algorithm is an instance of **dynamic programming**.
Sometimes $\{s_1^*, s_2^*, \dots, s_T^*\}$ is called the **Viterbi path**.

Which of the following statements is true(in terms of asymptotic time complexity)?

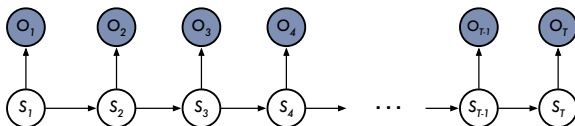
- A. The Viterbi algorithm generally runs faster than the Forward algorithm because it finds a single most likely sequence.
- B. The Forward algorithm generally runs faster than the Viterbi algorithm because it computes probabilities for all possible sequences.
- C. Both algorithms have similar running times.
- D. The Viterbi algorithm runs slower than the Forward algorithm due to its additional backtracking step.

What is the running time for Viterbi algorithm?

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

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

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

- 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:

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

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

- Numerator

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

Next lecture: EM with the forward-backward algorithm.

That's all folks!