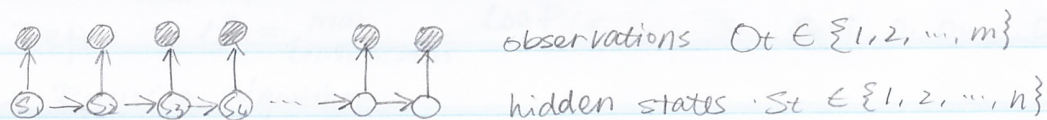


7-19

Review

* Hidden Markov models (HMMs)



* Parameters of CPTs

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

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

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

* Key questions

- 1) How to compute Likelihood $P(O_1, O_2, \dots, O_T)$?
 - 2) How to decode hidden states $\vec{S}^* = \text{argmax} P(\vec{S} | \vec{O})$?
 - 3) How to update beliefs $P(S_t = i | O_1, O_2, \dots, O_t)$?
 - 4) How to estimate/learn $\{\pi_i, a_{ij}, b_{jk}\}$ from data?
- $\left. \begin{matrix} 1-3 \text{ are} \\ \text{inference} \end{matrix} \right\}$
 $\left. \begin{matrix} 4 \text{ is} \\ \text{learning} \end{matrix} \right\}$

1. Computing Likelihood

$$\alpha_{it} = P(O_1, O_2, \dots, O_t, S_t = i)$$

$$i = 1 \dots n$$

$$t = 1 \dots T$$

Forward algorithm:

$$(t=1) \quad \alpha_{i1} = \pi_i b_i(O_1)$$

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

$$\text{Likelihood } P(O_1, O_2, \dots, O_T) = \sum_{i=1}^n \alpha_{iT}$$



2. Compute most likely state sequence.

$$S^* = \underset{S_1, S_2, \dots, S_T}{\operatorname{argmax}} \log P(S_1, S_2, \dots, S_T, O_1, O_2, \dots, O_T)$$

Define $l_{it}^* = \max_{\{S_1, S_2, \dots, S_{t-1}\}} \log P(S_1, S_2, \dots, S_{t-1}, S_t=i, O_1, O_2, \dots, O_t)$

* Recursive Algorithm

base case ($t=1$) left most column

$$\begin{aligned} l_{i1}^* &= \log P(S_1=i, O_1) \\ &= \log [P(S_1=i) P(O_1|S_1=i)] \quad \text{product rule} \\ &= \log [\pi_i b_i(O_1)] \\ &= \log \pi_i + \log b_i(O_1) \end{aligned}$$

recursive step from t to $t+1$

$$\begin{aligned} l_{j,t+1}^* &= \max_{S_1, S_2, \dots, S_t} [\log P(S_1, S_2, \dots, S_t, S_{t+1}=j, O_1, O_2, \dots, O_{t+1})] \\ &= \max_{S_1, S_2, \dots, S_t} \max_i [\log P(S_1, S_2, \dots, S_t=i, S_{t+1}=j, O_1, O_2, \dots, O_{t+1})] \\ &= \max_{S_1, S_2, \dots, S_t} \max_i [\log \{ P(S_1, \dots, S_t=i, O_1, \dots, O_t) \\ &\quad \times P(S_{t+1}=j | S_1, \dots, S_t=i, O_1, \dots, O_t) \\ &\quad \times P(O_{t+1} | S_1, \dots, S_t=i, S_{t+1}=j, O_1, \dots, O_t) \}] \\ &\quad \text{product rule} \\ &\quad \text{cond ind.} \\ &= \max_{S_1, \dots, S_t} \max_i [\log P(S_1, \dots, S_t=i, O_1, \dots, O_t) + \log P(S_{t+1}=j | S_t=i) \\ &\quad + \log P(O_{t+1} | S_{t+1}=j)] \end{aligned}$$

Note: The last two terms do not depend on $S_1, \dots, S_{t-1}$

$$\begin{aligned} l_{j,t+1}^* &= \max_i [\max_{S_1, \dots, S_{t-1}} \log P(S_1, \dots, S_{t-1}, S_t=i, O_1, \dots, O_t) \\ &\quad + \log P(S_{t+1}=j | S_t=i) + \log P(O_{t+1} | S_{t+1}=j)] \end{aligned}$$

Forward Pass

$$l_{j,t+1}^* = \max_i [l_{it}^* + \log a_{ij} + \log b_j(O_{t+1})]$$

fill in l_{it}^* column by column.

* Derive S^* from l^* .

Record most likely transitions:

$$\Phi_{t+1}(j) = \operatorname{argmax}_i [l_{it}^* + \log a_{ij}]$$

If I must end up in state j at time $t+1$, then what state i preceded at time t ?

* Compute S^* by backtracking:

$$\text{At time } T: S_T^* = \operatorname{argmax}_i [l_{iT}^*]$$

$$\text{For times } t = T-1 \text{ to } 1: S_t^* = \Phi_{t+1}(S_{t+1}^*)$$

Backward Pass

* Jargon:

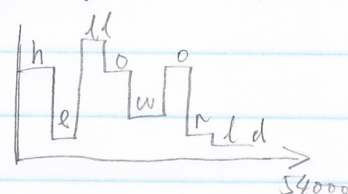
- S^* is known as "Viterbi" sequence or path

- Viterbi algorithm is example of dynamic programming

$$HW = n \geq 6 \quad (S_t \in \{a, b, c, \dots, z\})$$

$$m = 2 \quad (O_t \in \{0, 1\})$$

$$T = 54000$$



3. Belief updating

How to monitor hidden state in real time based on incoming evidence?

Key quantity: $q_{jt} = P(S_t = j | O_1, O_2, \dots, O_t)$

$$q_{jt} = P(S_t = j | O_1, \dots, O_{t-1}, O_t)$$

$$= \frac{P(O_t | S_t = j, O_1, \dots, O_{t-1}) P(S_t = j | O_1, \dots, O_{t-1})}{P(O_t | O_1, O_2, \dots, O_{t-1})}$$

$$b_{jt}(O_t) = \frac{P(O_t | S_t = j) P(S_t | O_1, \dots, O_{t-1})}{P(O_t | O_1, \dots, O_{t-1})} \quad \begin{matrix} \text{cond.} \\ \text{ind.} \end{matrix}$$

Focus on 2nd term in numerator:

$$P(S_t = j | O_1, \dots, O_{t-1}) = \sum_{i=1}^n P(S_t = j, S_{t-1} = i | O_1, O_2, \dots, O_{t-1})$$

$$= \sum_{i=1}^n P(S_{t-1} = i | O_1, \dots, O_{t-1}) \quad \begin{matrix} \text{product} \\ \text{rule} \end{matrix}$$

$$= \sum_{i=1}^n P(S_t = j | S_{t-1} = i, O_1, \dots, O_{t-1}) \quad \begin{matrix} \text{marginalization} \\ \text{cond. ind.} \end{matrix}$$

$$= \sum_{i=1}^n q_{i,t-1} a_{ij}$$

Putting into numerator =

$$g_{jt} = \frac{b_j(o_t) \sum_{i=1}^n g_{i,t-1} a_{ij}}{P(o_t | o_1, \dots, o_{t-1})}$$

normalization

$$g_{jt} = \frac{b_j(o_t) \sum_i g_{i,t-1} a_{ij}}{\sum_{j'} b_{j'}(o_t) \sum_i g_{i,t-1} a_{ij'}}$$

How much computation per time step? $O(n)$

Back from
break.

Key questions in HMMs

✓ 1) How to compute likelihood $P(o_1, o_2, \dots, o_T)$?

✓ 2) How to decode $\{s_1^*, s_2^*, \dots, s_T^*\} = \underset{s_1, \dots, s_T}{\operatorname{argmax}} P(s_1, \dots, s_T | o_1, \dots, o_T)$?

✓ 3) How to update beliefs $P(s_t = i | o_1, \dots, o_t)$?

Onto learning ...

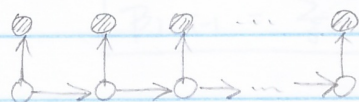
4. Learning in HMMs

Given = sequence of observations $\{o_1, o_2, \dots, o_T\}$
(just one for simplicity).

Goal: estimate $\{\pi_i, a_{ij}, b_{ik}\}$ to maximize $P(o_1, \dots, o_T)$

Assume: # hidden state values, n , is known.

* EM algorithm in HMMs:



CPTs to 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: compute posteriors

$$P(S_t = i | O_1, \dots, O_T)$$

$$P(S_t = i, S_{t+1} = j | O_1, \dots, O_T)$$

$$P(S_t = i, O_t = k | O_1, \dots, O_T) = I(O_t, k) P(S_t = i | O_1, \dots, O_T)$$

Assuming we can compute these posteriors efficiently...

* M-step: update CPTs

$$\pi_i \leftarrow P(S_t = i | O_1, \dots, O_T)$$

$$a_{ij} \leftarrow \frac{\sum_t P(S_t = i, S_{t+1} = j | O_1, \dots, O_T)}{\sum_t P(S_t = i | O_1, \dots, O_T)}$$

$$b_{ik} \leftarrow \frac{\sum_t P(S_t = i | O_1, \dots, O_T) I(O_t, k)}{\sum_t P(S_t = i | O_1, \dots, O_T)}$$

* How to compute these posterior probabilities?

Recall $\alpha_{it} = P(O_1, \dots, O_t, S_t = i)$

Define $\phi_{it} = P(O_{t+1}, O_{t+2}, \dots, O_T | S_t = i)$

Starts at time $t+1$

conditioned on $S_t = i$

* Recursion for filling in Φ matrix

(1) base case

(a) $\Phi_{iT} = P(- | S_T = i) = 1$ for all $i = 1 \dots n$

(b) $\Phi_{i,T-1} = P(O_T | S_{T-1} = i)$

$$= \sum_{j=1}^n P(S_T = j, O_T | S_{T-1} = i) \quad \text{marginalization}$$

$$= \sum_j P(S_T = j | S_{T-1} = i) P(O_T | S_T = j, S_{T-1} = i)$$

product rule
cond. ind.

$$\boxed{\Phi_{i,T-1} = \sum_j a_{ij} b_j(O_T)}$$

(2) backward - step (from time $t+1$ to time t)

$$\phi_{it} = P(O_{t+1}, \dots, O_T | S_t = i)$$

$$= \sum_{j=1}^n P(O_{t+1}, \dots, O_T, S_{t+1} = j | S_t = i) \quad \text{marginalization}$$

$$= \sum_{j=1}^n P(S_{t+1} = j | S_t = i) P(O_{t+1} | S_{t+1} = j, S_t = i) \quad \text{product rule}$$

$$P(O_{t+2}, \dots, O_T | S_{t+1} = j, S_t = i, O_{t+1}) \quad \text{cond. ind.}$$

$$= \sum_{j=1}^n P(S_{t+1}=j | S_t=i) P(O_{t+1} | S_{t+1}=j) P(O_{t+2}, \dots, O_T | S_{t+1}=j)$$

$$\boxed{\beta_{it} = \sum_{j=1}^n a_{ij} b_j(O_{t+1}) \beta_{j,t+1}}$$

"Forward-backward" algorithm in HMMs computes α_{it} and β_{it} matrices.

Last question:

How do α_{it} and β_{it} give us posteriors for E-step of EM algorithm?

* Return to posteriors =

1) $P(S_t=i | O_1, \dots, O_T)$ needed to π_i update

2) $P(S_t=i, S_{t+1}=j | O_1, \dots, O_T)$ needed for a_{ij} update

3) $P(S_t=i | O_1, \dots, O_T)$ needed for b_{it} update

* Focus on middle probability =

$$P(S_t=i, S_{t+1}=j | O_1, \dots, O_T) = \frac{P(S_t=i, S_{t+1}=j, O_1, \dots, O_T)}{P(O_1, O_2, \dots, O_T)}$$

Denominator = $P(O_1, O_2, \dots, O_T)$ is just the likelihood.

$$P(O_1, O_2, \dots, O_T) = \sum_{i=1}^n \alpha_{iT}$$

Numerator: $P(S_t=i, S_{t+1}=j, O_1, \dots, O_T)$

$$= P(O_1, \dots, O_t, S_t=i) P(S_{t+1}=j | S_t=i, O_1, \dots, O_t)$$

$$P(O_{t+1} | S_{t+1}=j, \cancel{S_t=i, O_1, \dots, O_t}) P(\cancel{O_{t+2}, \dots, O_T} | S_{t+1}=j, \cancel{S_t=i, O_1, \dots, O_{t+1}})$$

product rule, cond. ind.

$$= P(O_1, \dots, O_t, S_t=i) P(S_{t+1}=j | S_t=i) P(O_{t+1} | S_{t+1}=j) P(O_{t+2}, \dots, O_T | S_{t+1}=j)$$

$$= \alpha_{it} \cdot a_{ij} \cdot b_j(O_{t+1}) \cdot \beta_{j,t+1}$$

Denominator: (by normalization)

$$P(O_1, \dots, O_T) = \sum_{i,j} P(S_t=i, S_{t+1}=j, O_1, \dots, O_T)$$

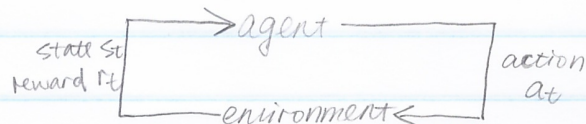
$$\text{Also correct.} \quad = \sum_{i,j} \alpha_{it} \cdot a_{ij} \cdot b_j(O_{t+1}) \cdot \beta_{j,t+1}$$

$$= \sum_i \alpha_{it} \beta_{it} = \sum_j \alpha_{j,t+1} \beta_{j,t+1}$$

HW. Compute $p(S_t=i | S_{t+1}=j, O_1, \dots, O_T)$

Reinforcement Learning

Q: How should embodied / embedded / decision-making agents act and learn from experience in the world?



Ex. robot navigation, game playing AI

Textbook: Sutton / Barto

* Challenges

- handling uncertainty
- exploration vs. exploitation dilemma
- delayed vs. immediate rewards — "temporal credit assignment"
- evaluative vs. instructive feedback
- complex worlds, computational guaranteed