

→ Learning CPTs from incomplete data

EM update:

$$P(x_i = \alpha \mid pa_i = \pi) \leftarrow \frac{\sum_t P(x_i = \alpha, pa_i = \pi / v^t)}{\sum_t P(pa_i = \pi / v^t)} \quad (\text{for roots with parents})$$

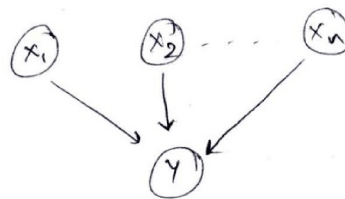
$$P(x_i = \alpha) \leftarrow \frac{1}{T} \sum_t P(x_i = \alpha / v^t) \quad (\text{root nodes})$$

where v^t denotes visible nodes.

→ HW4 | NOISY-OR.

$$P(Y=1 \mid \vec{x}) = 1 - \prod_{i=1}^n (1 - P_i)^{x_i}$$

we learn from data $\{(\vec{x}_t, y_t)\}_{t=1}^T$



EM update of this model is given by:

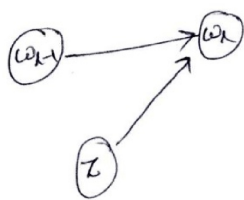
$$P_i \leftarrow \frac{1}{T_i} \sum_{t=1}^T \frac{y_t x_{it} P_i}{1 - \prod_{j=1}^n (1 - P_j)^{x_{jt}}}$$

where T_i is the number of examples in which $x_i = 1$

→ Linear interpolation (mixing) of n-gram models:

$$P_n(w_k / w_{1:k-1}) = \lambda P_1(w_k) + (1 - \lambda) P_2(w_k / w_{1:k-1})$$

To estimate λ , we rewrite this as a hidden variable model



$Z \in \{1, 2\}$

$$P(w_k / w_{1:k-1}, Z) = P_1(w_k) \text{ if } Z=1$$

$$P_2(w_k / w_{1:k-1}) \text{ if } Z=2$$

$$P(Z=1) = \lambda$$

$$P(Z=2) = 1 - \lambda$$

In this model,

$$\begin{aligned}
 P(\omega_k / \omega_{k-1}) &= \sum_{z'} P(\omega_k, z' / \omega_{k-1}) \\
 &= \sum_{z'} P(z' / \omega_{k-1}) P(\omega_k / z', \omega_{k-1}) \quad [\text{Product Rule}] \\
 &= \sum_{z'} P(z') P(\omega_k / z', \omega_{k-1}) \quad [\text{Conditional independence}] \\
 &= \lambda P_1(\omega_k) + (1-\lambda) P_2(\omega_k / \omega_{k-1})
 \end{aligned}$$

E-STEP

Compute posterior probability

$$* P(z / \omega_k, \omega_{k-1}) = \frac{P(\omega_k / z, \omega_{k-1}) P(z / \omega_{k-1})}{P(\omega_k / \omega_{k-1})} \quad [\text{Conditional Bayes' Rule}]$$

$$* P(z=1 / \omega_k, \omega_{k-1}) = \frac{\lambda P_1(\omega_k / \omega_{k-1})}{\lambda P_1(\omega_k) + (1-\lambda) P_2(\omega_k / \omega_{k-1})}$$

$$* P(z=2 / \omega_k, \omega_{k-1}) = 1 - P(z=1 / \omega_k, \omega_{k-1})$$

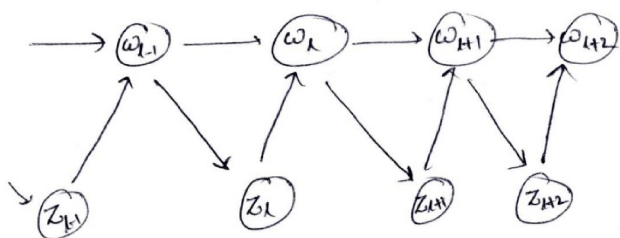
M-STEP

In this model

$$P(z=1) \leftarrow \frac{1}{L} \sum_{k=1}^L P(z=1 / \omega_{k-1}, \omega_k)$$

$\nearrow$
 λ

The above model over simplifies the belief network. In real world applications, smoothing parameter, λ would depend on the previous word.



$$P(z_t=1/\omega_{t-1}) = \lambda(\omega_{t-1}) \quad P(z_t=2/\omega_{t-1}) = 1 - \lambda(\omega_{t-1})$$

$$P(\omega_t/\omega_{t-1}) = \lambda(\omega_{t-1}) P_1(\omega_t) + 1 - \lambda(\omega_{t-1}) P_2(\omega_t/\omega_{t-1})$$

Using EM, we can estimate λ , which has the same size as the words in the vocabulary.

HIDDEN MARKOV MODELS

→ Random variables :

$S_t = \{1, 2, \dots, n\}$ state at time t

$O_t = \{1, 2, \dots, m\}$ observation at time t .

* Note that n, m have no relation ($n > m, n = m, n < m \leftarrow$ all possible)

* observations, O_t can be thought of as a noisy reflection of the hidden state S_t

→ Ex: puppy training.

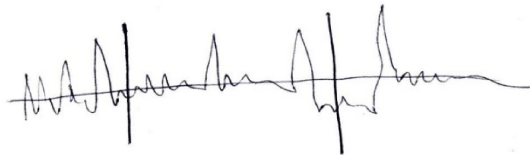
$S = \{\text{"have-to-go"}, \text{"doesn't-have-to-go"}, \text{"went"}\}$

$O = \{\text{"wagging tail"}, \text{"squeaking"}, \text{"running in circles"}, \text{"hiding in corners"}\}$

→ Ex: Speech recognition:

s_t = units of language: words, syllables, phonemes

o_t = acoustic measurement: energy in a window, no. of peaks in a window...



→ Ex: Robotics

s_t = location

o_t = sensor reading

MARKOV ASSUMPTIONS IN HMM!

finite Context:

$$P(s_t / s_1, s_2, \dots, s_{t-1}) = P(s_t / s_{t-1})$$

$$P(o_t / s_1, s_2, \dots, s_{t-1}, s_t, \dots, s_T) = P(o_t / s_t)$$

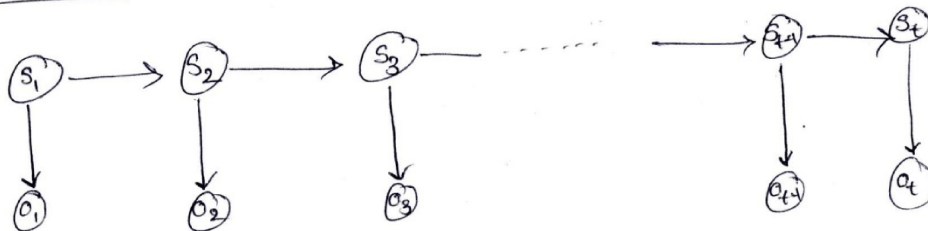
shared CPTs

$$P(s_{t+1} = s' / s_t = s) = P(s_t = s' / s_{t-1} = s)$$

$$P(o_t = o / s_t = s) = P(o_{t+1} = o / s_{t+1} = s)$$

BELIEF NETWORK OF HMM

"polytree"



PARAMETERS OF HMM:

$$\pi_i = \pi(s_i = i) \quad [\text{Initial State Distribution}]$$

$n \times n$ matrix $a_{ij} = P(S_{t+1}=j / S_t=i)$ [Transition matrix] $i=1,2,\dots,n$
 $j=1,2,\dots,n$

$$n \times m \text{ matrix } b_{ik} = p(O_t = k / S_t = i) \quad [\text{Emission matrix}] \quad \begin{matrix} i = 1, 2, \dots, n \\ k = 1, 2, \dots, m \end{matrix}$$

JOINT DISTRIBUTION OF HMM!

JOINT DISTRIBUTION OF THE

$$P(\vec{s}, \vec{o}) = P(s_1) \prod_{t=2}^T P(s_t | s_{t-1}) \prod_{t=1}^T P(o_t | s_t)$$

$s_1, s_2, \dots, s_T$ $o_1, o_2, \dots, o_T$

→ Ex: Isolated word speech recognition

* Consider the problem of recognising the word "CAT"

- * Consider the problem of recognising
- * We have to build HMM that assigns high probability to "CAT" utterances
low probability to other utterances

1. *Phragmites australis* (Cav.) Trin. ex Steud.

→ uses HMM with 5 states

state #	sound
1	initial silence
2	"c"
3	"A"
4	"T"
5	final silence

$$\pi_i = \begin{bmatrix} 1 \\ 0 \\ 0 \\ 0 \end{bmatrix}$$

← always start with state 1

$$a_{ij} = \begin{bmatrix} 0.9 & 0.1 & 0 & 0 & 0 \\ 0 & 0.9 & 0.1 & 0 & 0 \\ 0 & 0 & 0.99 & 0.01 & 0 \\ 0 & 0 & 0 & 0.4 & 0.6 \\ 0 & 0 & 0 & 0 & 1 \end{bmatrix}$$

→ This is an upper diagonal transition matrix
 → Special Case: left to right HMM

KEY QUESTIONS FOR HMMs

perform inference → given $\{\pi_i, a_{ij}, b_{ik}\}$ parameters

1) How to Compute likelihood $P(o_1, o_2, \dots, o_T)$

2) How to Compute most likely state sequence

$$(s_1^*, s_2^*, \dots, s_T^*) = \underset{s_1, s_2, \dots, s_T}{\operatorname{argmax}} \underbrace{P(s_1, s_2, \dots, s_T / o_1, o_2, \dots, o_T)}_{n^T \text{ possible settings.}}$$

3) How to compute $P(s_t = i / o_1, o_2, \dots, o_t)$?

↳ this is can be thought of as updating belief in real time.

learning → given $\{o_1, o_2, \dots, o_T\}$ observations.

4) How to estimate parameters $\{\pi_i, a_{ij}, b_{ik}\}$ that maximises likelihood $P(o_1, o_2, \dots, o_T)$ } EM algorithm.

$$\begin{aligned} \text{1) } P(o_1, o_2, \dots, o_T) &= \sum_{s_1, s_2, \dots, s_T} \underbrace{P(s_1, s_2, \dots, s_T)}_{\text{hidden}} \underbrace{P(o_1, o_2, \dots, o_T)}_{\text{observed}} \\ &= \sum_{\vec{s}} P(s_1) \prod_{t=2}^T P(s_t | s_{t-1}) \prod_{t=1}^T P(o_t | s_t) \end{aligned}$$

← sum over n^T hidden state sequences