

HIDDEN MARKOV MODELS

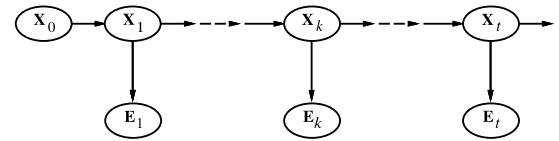
AIMA CHAPTER 15, SECTIONS 1–5

Hidden Markov Model (HMM)

Sensor Markov assumption: $P(E_t | X_{0:t}, E_{1:t-1}) = P(E_t | X_t)$

Stationary process: transition model $P(X_t | X_{t-1})$ and sensor model $P(E_t | X_t)$ fixed for all t

HMM is a special type of Bayes net, X_t is single discrete random variable:



with joint probability distribution

$$P(X_{0:t}, E_{1:t}) = P(X_0) \prod_{i=1}^t P(X_i | X_{i-1}) P(E_i | X_i)$$

AIMA Chapter 15, Sections 1–5 1

AIMA Chapter 15, Sections 1–5 4

Time and uncertainty

Consider a target tracking problem

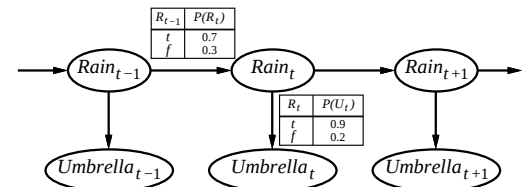
X_t = set of unobservable state variables at time t
e.g., *Position_t*, *Appearance_t*, etc.

E_t = set of observable evidence variables at time t
e.g., *Imagepixels_t*

This assumes **discrete time**; step size depends on problem

Notation: $X_{a:b} = X_a, X_{a+1}, \dots, X_{b-1}, X_b$

Example



First-order Markov assumption not exactly true in real world!

Possible fixes:

1. **Increase order** of Markov process
2. **Augment state**, e.g., add *Temp_t*, *Pressure_t*

Example: robot motion.

Augment position and velocity with *Battery_t*

AIMA Chapter 15, Sections 1–5 2

AIMA Chapter 15, Sections 1–5 5

Markov processes (Markov chains)

Construct a Bayes net from these variables:

Markov assumption: X_t depends on **bounded** subset of $X_{0:t-1}$

First-order Markov process: $P(X_t | X_{0:t-1}) = P(X_t | X_{t-1})$

Second-order Markov process: $P(X_t | X_{0:t-1}) = P(X_t | X_{t-2}, X_{t-1})$

First-order

Second-order

Stationary process: transition model $P(X_t | X_{t-1})$ fixed for all t

AIMA Chapter 15, Sections 1–5 3

AIMA Chapter 15, Sections 1–5 6

Inference tasks

Filtering: $P(X_t | e_{1:t})$

belief state—input to the decision process of a rational agent

Prediction: $P(X_{t+k} | e_{1:t})$ for $k > 0$

evaluation of possible action sequences;
like filtering without the evidence

Smoothing: $P(X_k | e_{1:t})$ for $0 \leq k < t$

better estimate of past states, essential for learning

Most likely explanation: $\arg \max_{x_{1:t}} P(x_{1:t} | e_{1:t})$

speech recognition, decoding with a noisy channel

Filtering

Aim: devise a **recursive** state estimation algorithm:

$$P(\mathbf{X}_{t+1}|\mathbf{e}_{1:t+1}) = f(\mathbf{e}_{t+1}, P(\mathbf{X}_t|\mathbf{e}_{1:t}))$$

Filtering

Aim: devise a **recursive** state estimation algorithm:

$$P(\mathbf{X}_{t+1}|\mathbf{e}_{1:t+1}) = f(\mathbf{e}_{t+1}, P(\mathbf{X}_t|\mathbf{e}_{1:t}))$$

$$\begin{aligned} P(\mathbf{X}_{t+1}|\mathbf{e}_{1:t+1}) &= P(\mathbf{X}_{t+1}|\mathbf{e}_{1:t}, \mathbf{e}_{t+1}) \\ &= \alpha P(\mathbf{e}_{t+1}|\mathbf{X}_{t+1}, \mathbf{e}_{1:t}) P(\mathbf{X}_{t+1}|\mathbf{e}_{1:t}) \\ &= \alpha P(\mathbf{e}_{t+1}|\mathbf{X}_{t+1}) P(\mathbf{X}_{t+1}|\mathbf{e}_{1:t}) \end{aligned}$$

I.e., **prediction** + **estimation**. Prediction by summing out $\mathbf{X}_t$:

$$\begin{aligned} P(\mathbf{X}_{t+1}|\mathbf{e}_{1:t+1}) &= \alpha P(\mathbf{e}_{t+1}|\mathbf{X}_{t+1}) \sum_{\mathbf{x}_t} P(\mathbf{X}_{t+1}, \mathbf{x}_t|\mathbf{e}_{1:t}) \\ &= \alpha P(\mathbf{e}_{t+1}|\mathbf{X}_{t+1}) \sum_{\mathbf{x}_t} P(\mathbf{X}_{t+1}|\mathbf{x}_t, \mathbf{e}_{1:t}) P(\mathbf{x}_t|\mathbf{e}_{1:t}) \\ &= \alpha P(\mathbf{e}_{t+1}|\mathbf{X}_{t+1}) \sum_{\mathbf{x}_t} P(\mathbf{X}_{t+1}|\mathbf{x}_t) P(\mathbf{x}_t|\mathbf{e}_{1:t}) \end{aligned}$$

$\mathbf{f}_{1:t+1} = \text{FORWARD}(\mathbf{f}_{1:t}, \mathbf{e}_{t+1})$ where $\mathbf{f}_{1:t} = P(\mathbf{X}_t|\mathbf{e}_{1:t})$
Time and space **constant** (independent of t)

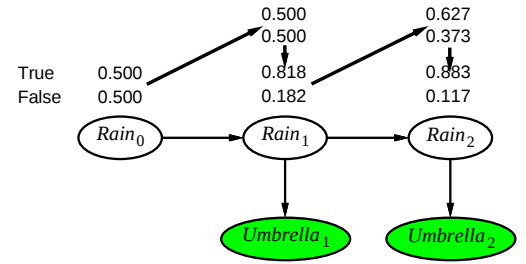
Filtering

Aim: devise a **recursive** state estimation algorithm:

$$P(\mathbf{X}_{t+1}|\mathbf{e}_{1:t+1}) = f(\mathbf{e}_{t+1}, P(\mathbf{X}_t|\mathbf{e}_{1:t}))$$

$$\begin{aligned} P(\mathbf{X}_{t+1}|\mathbf{e}_{1:t+1}) &= P(\mathbf{X}_{t+1}|\mathbf{e}_{1:t}, \mathbf{e}_{t+1}) \\ &= \alpha P(\mathbf{e}_{t+1}|\mathbf{X}_{t+1}, \mathbf{e}_{1:t}) P(\mathbf{X}_{t+1}|\mathbf{e}_{1:t}) \\ &= \alpha P(\mathbf{e}_{t+1}|\mathbf{X}_{t+1}) P(\mathbf{X}_{t+1}|\mathbf{e}_{1:t}) \end{aligned}$$

Filtering example



$$P(\mathbf{X}_{t+1}|\mathbf{e}_{1:t+1}) = \alpha P(\mathbf{e}_{t+1}|\mathbf{X}_{t+1}) \sum_{\mathbf{x}_t} P(\mathbf{X}_{t+1}|\mathbf{x}_t) P(\mathbf{x}_t|\mathbf{e}_{1:t})$$

R_{t-1}	$P(R_t)$	R_t	$P(U_t)$
t	0.7	t	0.9
f	0.3	f	0.2

Filtering

Aim: devise a **recursive** state estimation algorithm:

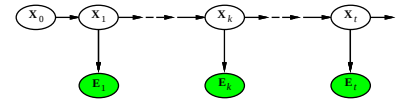
$$P(\mathbf{X}_{t+1}|\mathbf{e}_{1:t+1}) = f(\mathbf{e}_{t+1}, P(\mathbf{X}_t|\mathbf{e}_{1:t}))$$

$$\begin{aligned} P(\mathbf{X}_{t+1}|\mathbf{e}_{1:t+1}) &= P(\mathbf{X}_{t+1}|\mathbf{e}_{1:t}, \mathbf{e}_{t+1}) \\ &= \alpha P(\mathbf{e}_{t+1}|\mathbf{X}_{t+1}, \mathbf{e}_{1:t}) P(\mathbf{X}_{t+1}|\mathbf{e}_{1:t}) \\ &= \alpha P(\mathbf{e}_{t+1}|\mathbf{X}_{t+1}) P(\mathbf{X}_{t+1}|\mathbf{e}_{1:t}) \end{aligned}$$

I.e., **prediction** + **estimation**. Prediction by summing out $\mathbf{X}_t$:

$$\begin{aligned} P(\mathbf{X}_{t+1}|\mathbf{e}_{1:t+1}) &= \alpha P(\mathbf{e}_{t+1}|\mathbf{X}_{t+1}) \sum_{\mathbf{x}_t} P(\mathbf{X}_{t+1}, \mathbf{x}_t|\mathbf{e}_{1:t}) \\ &= \alpha P(\mathbf{e}_{t+1}|\mathbf{X}_{t+1}) \sum_{\mathbf{x}_t} P(\mathbf{X}_{t+1}|\mathbf{x}_t, \mathbf{e}_{1:t}) P(\mathbf{x}_t|\mathbf{e}_{1:t}) \\ &= \alpha P(\mathbf{e}_{t+1}|\mathbf{X}_{t+1}) \sum_{\mathbf{x}_t} P(\mathbf{X}_{t+1}|\mathbf{x}_t) P(\mathbf{x}_t|\mathbf{e}_{1:t}) \end{aligned}$$

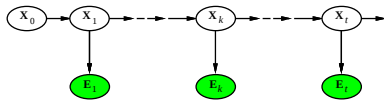
Smoothing



Divide evidence $\mathbf{e}_{1:t}$ into $\mathbf{e}_{1:k}$, $\mathbf{e}_{k+1:t}$:

$$\begin{aligned} P(\mathbf{X}_k|\mathbf{e}_{1:t}) &= P(\mathbf{X}_k|\mathbf{e}_{1:k}, \mathbf{e}_{k+1:t}) \\ &= \alpha P(\mathbf{X}_k|\mathbf{e}_{1:k}) P(\mathbf{e}_{k+1:t}|\mathbf{X}_k, \mathbf{e}_{1:k}) \\ &= \alpha P(\mathbf{X}_k|\mathbf{e}_{1:k}) P(\mathbf{e}_{k+1:t}|\mathbf{X}_k) \\ &= \alpha \mathbf{f}_{1:k} \mathbf{b}_{k+1:t} \end{aligned}$$

Smoothing



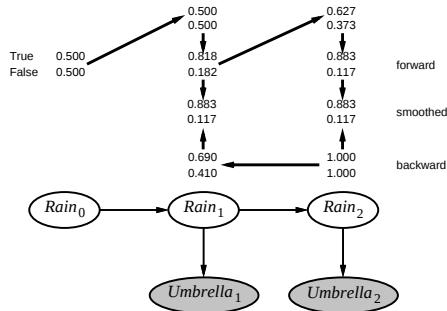
Divide evidence $e_{1:t}$ into $e_{1:k}$, $e_{k+1:t}$:

$$\begin{aligned} P(\mathbf{X}_k | e_{1:t}) &= P(\mathbf{X}_k | e_{1:k}, e_{k+1:t}) \\ &= \alpha P(\mathbf{X}_k | e_{1:k}) P(e_{k+1:t} | \mathbf{X}_k, e_{1:k}) \\ &= \alpha P(\mathbf{X}_k | e_{1:k}) P(e_{k+1:t} | \mathbf{X}_k) \\ &= \alpha f_{1:k} b_{k+1:t} \end{aligned}$$

Backward message computed by a backwards recursion:

$$\begin{aligned} P(e_{k+1:t} | \mathbf{X}_k) &= \sum_{\mathbf{x}_{k+1}} P(e_{k+1:t} | \mathbf{X}_k, \mathbf{x}_{k+1}) P(\mathbf{x}_{k+1} | \mathbf{X}_k) \\ &= \sum_{\mathbf{x}_{k+1}} P(e_{k+1:t} | \mathbf{x}_{k+1}) P(\mathbf{x}_{k+1} | \mathbf{X}_k) \\ &= \sum_{\mathbf{x}_{k+1}} P(e_{k+1} | \mathbf{x}_{k+1}) P(e_{k+2:t} | \mathbf{x}_{k+1}) P(\mathbf{x}_{k+1} | \mathbf{X}_k) \end{aligned}$$

Smoothing example



Forward-backward algorithm: cache forward messages along the way
Time linear in t (polytree inference), space $O(t|f|)$

Most likely explanation

Most likely explanation

Most likely sequence $\neq$ sequence of most likely states!!!!

Most likely path to each $\mathbf{x}_{t+1}$

= most likely path to **some** $\mathbf{x}_t$ plus one more step

$$\begin{aligned} \max_{\mathbf{x}_1 \dots \mathbf{x}_t} P(\mathbf{x}_1, \dots, \mathbf{x}_t, \mathbf{X}_{t+1} | e_{1:t+1}) \\ = P(e_{t+1} | \mathbf{X}_{t+1}) \max_{\mathbf{x}_t} \left(P(\mathbf{X}_{t+1} | \mathbf{x}_t) \max_{\mathbf{x}_1 \dots \mathbf{x}_{t-1}} P(\mathbf{x}_1, \dots, \mathbf{x}_{t-1}, \mathbf{x}_t | e_{1:t}) \right) \end{aligned}$$

Identical to filtering, except $f_{1:t}$ replaced by

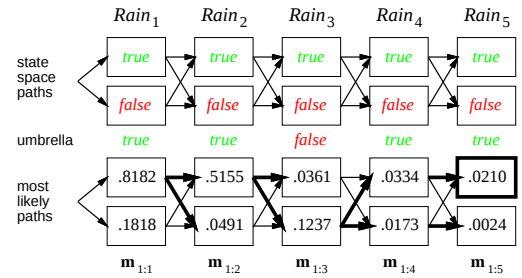
$$\mathbf{m}_{1:t} = \max_{\mathbf{x}_1 \dots \mathbf{x}_{t-1}} P(\mathbf{x}_1, \dots, \mathbf{x}_{t-1}, \mathbf{X}_t | e_{1:t}),$$

i.e., $\mathbf{m}_{1:t}(i)$ gives the probability of the most likely path to state i .

Update has sum replaced by max, giving the **Viterbi algorithm**:

$$\mathbf{m}_{1:t+1} = P(e_{t+1} | \mathbf{X}_{t+1}) \max_{\mathbf{x}_t} (P(\mathbf{X}_{t+1} | \mathbf{x}_t) \mathbf{m}_{1:t})$$

Viterbi example



Most likely explanation

Example Umbrella Problems

Filtering:

$$P(\mathbf{X}_{t+1} | e_{1:t+1}) = \alpha P(e_{t+1} | \mathbf{X}_{t+1}) \sum_{\mathbf{x}_t} P(\mathbf{X}_{t+1} | \mathbf{x}_t) P(\mathbf{x}_t | e_{1:t}) =: f_{1:t+1}$$

Smoothing:

$$\begin{aligned} P(\mathbf{X}_k | e_{1:t}) &= \alpha f_{1:k} b_{k+1:t} \\ P(e_{k+1:t} | \mathbf{X}_k) &= \sum_{\mathbf{x}_{k+1}} P(e_{k+1} | \mathbf{x}_{k+1}) P(e_{k+2:t} | \mathbf{x}_{k+1}) P(\mathbf{x}_{k+1} | \mathbf{X}_k) =: b_{k+1:t} \end{aligned}$$

R_{t-1}	$P(R_t)$	R_t	$P(U_t)$
t	0.7	t	0.9
f	0.3	f	0.2

$$\begin{aligned} P(R_3 | \neg u_1, u_2, \neg u_3) &= ? \\ \arg \max_{R_{1:3}} P(R_{1:3} | \neg u_1, u_2, \neg u_3) &= ? \\ P(R_2 | \neg u_1, u_2, \neg u_3) &= ? \end{aligned}$$