

Hidden Markov Models

Temporal Models $x_0 \rightarrow x_1 \rightarrow x_2 \rightarrow \dots \rightarrow x_{t-1} \rightarrow x_t \rightarrow \dots$

Special case of Bayesian Networks x_t depends on x_{t-1}

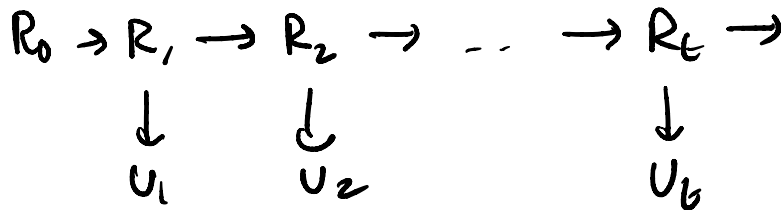
x_t : Rain on day t

$$P(x_t = 1 | x_{t-1} = 1)$$

$$P(x_t = 1 | x_{t-1} = 0)$$

x_{t-1}	$P(x_t)$
0	0.1
1	0.8

Observability - indoors, only see umbrella



Hidden Markov
Model

Notation: $x_{i:j} \stackrel{\text{def}}{=} x_i, x_{i+1}, \dots, x_j$ Hidden states
 $e_{i:j} \stackrel{\text{def}}{=} e_i, e_{i+1}, \dots, e_j$ Evidence = observations

Typical Inference Problems

$P(x_t | e_{1:t})$ Current state from evidence Filtering

$P(x_{t+k} | e_{1:t})$ Future state from evidence Prediction

$P(x_k | e_{1:t}), k < t$ Past state estimation Smoothing

$\operatorname{argmax} P(x_{1:t} | e_{1:t})$ Most likely sequence

FILTERING

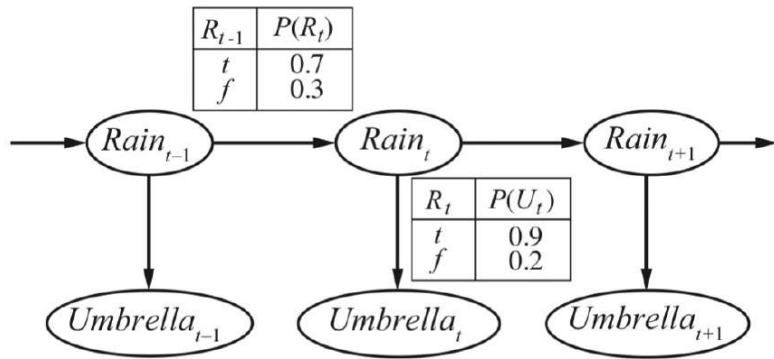
$$\underline{P(x_{t+1} | e_{1:t+1})} = P(x_{t+1} | e_{1:t}, e_{t+1})$$

$$= P(e_{t+1} | x_{t+1}, e_{1:t}) \cdot P(x_{t+1} | e_{1:t})$$

$$= P(e_{t+1} | x_{t+1}) \cdot P(x_{t+1} | e_{1:t})$$

$$= P(e_{t+1} | x_{t+1}) \cdot \sum_{x_t} P(x_{t+1} | x_t, e_{1:t}) P(x_t | e_{1:t})$$

$$= \underbrace{P(e_{t+1} | x_{t+1})}_{\text{blue}} \sum_{x_t} \underbrace{P(x_{t+1} | x_t)}_{\text{blue}} \underbrace{P(x_t | e_{1:t})}_{\text{blue}}$$



$$U_1 = U_2 = t$$

$$P(R_2) = ?$$

$$\text{Assume } P(R_0) = 0.5$$

$$P(R_1) = \sum_{r_0} P(R_1 | r_0) P(r_0) = \langle 0.5, 0.5 \rangle$$

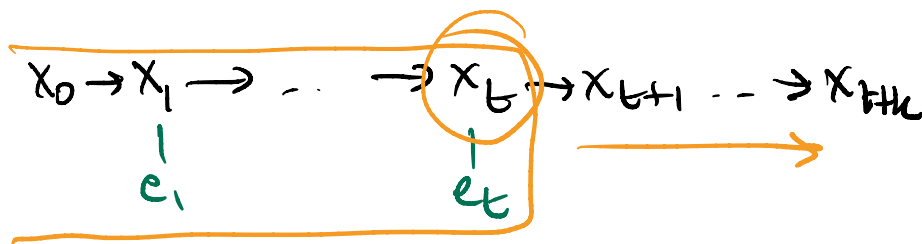
$$P(R_1 | U_1) = P(U_1 | R_1) P(R_1) = \langle 0.818, 0.182 \rangle$$

$$P(R_2 | U_1) = P(R_2 | R_1) = \langle 0.627, 0.373 \rangle$$

$$P(R_2 | U_1, U_2) = \langle 0.883, 0.117 \rangle$$

PREDICTION

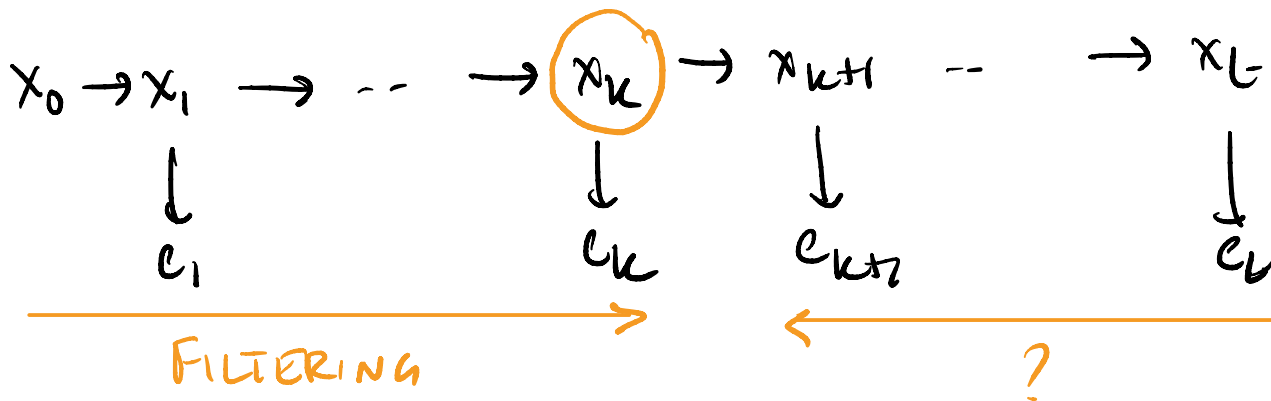
$$P(x_{t+k} | e_{1:t})$$



$$= \underbrace{P(x_t | e_{1:t})}_{\text{FILTER}} \cdot \underbrace{P(x_{t+1} | x_t)} \dots \underbrace{P(x_{t+k} | x_{t+k-1})}$$

SMOOTHING

$$P(x_k | e_{1:t}), k < t$$



$$P(x_k | e_{1:t}) = P(x_k | e_{1:k}, e_{k+1:t})$$

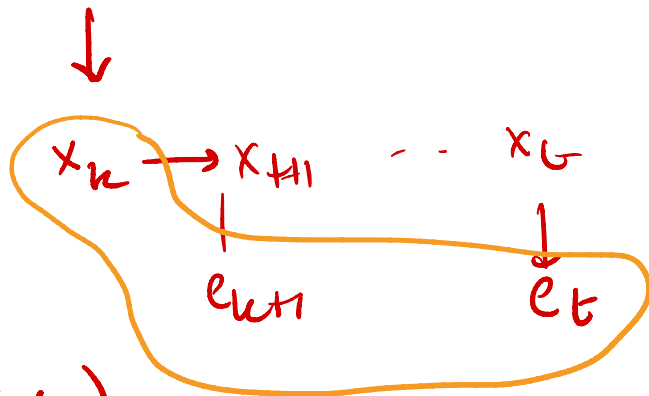
$$= P(x_k | e_{1:k}) \cdot P(e_{k+1:t} | x_k, e_{1:k})$$

$$P(e_{k+1:t} | x_k) =$$

$$\sum_{x_{k+1}} P(e_{k+1:t} | x_k, x_{k+1}) \cdot P(x_{k+1} | x_k)$$

$$= \sum_{x_{k+1}} P(e_{k+1}, e_{k+2:t} | x_{k+1}) \cdot P(x_{k+1} | x_k) \quad \text{MODEL}$$

$$\text{MODEL} \quad P(e_{k+1} | x_{k+1}) \cdot P(e_{k+2:t} | x_{k+1}) \quad \text{INDUCTION}$$



Earlier, by filtering

$$P(R_1 | U_1) = \langle 0.818, 0.182 \rangle$$

$$P(R_2 | U_1, U_2) = \langle 0.883, 0.117 \rangle$$

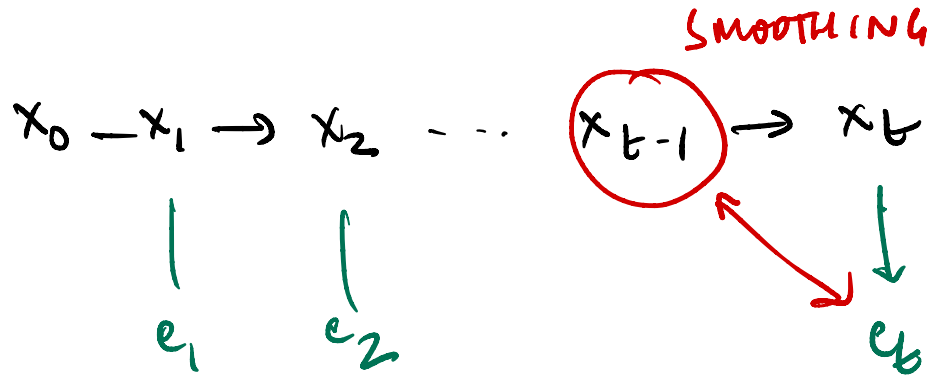
Smoothing

$$P(R_1 | U_1, U_2) = \langle 0.883, 0.117 \rangle$$

Better estimate !

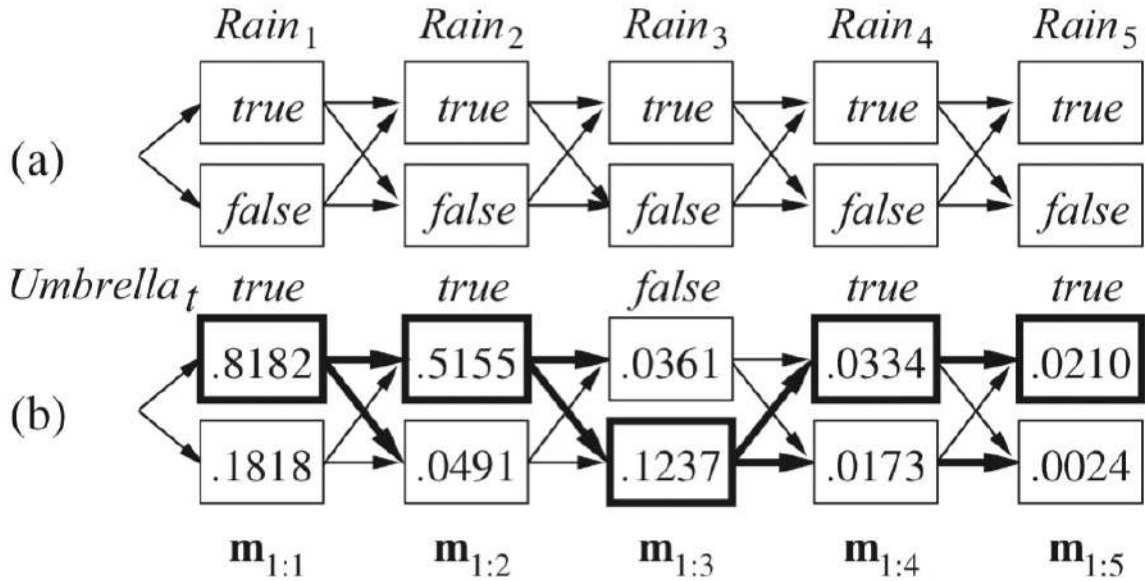
Most likely explanation

Given $e_{1:t}$, find best $x_{1:t}$



$$\text{Best}(x_{1:t+1} | e_{1:t+1}) = \max_{x_{t+1}} P(x_{t+1} | x_t, e_{t+1})$$

$$\cdot \text{Best}(x_{1:t} | e_{1:t})$$



Dynamic Programming

$B(x, t+1)$ - Best prob of $t+1$ length sequences
ending in state x

$$= \max_y P(x | e_{t+1}, y) B(y, t)$$

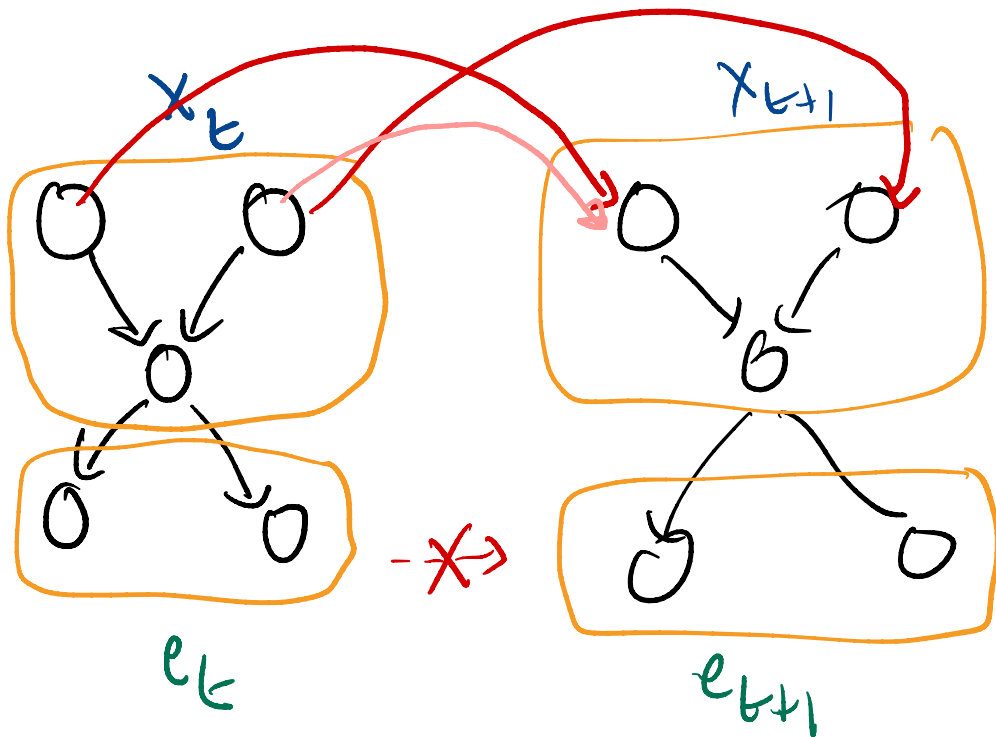
Viterbi

HMM

- Linear unrolling

$$P(x_{t+1} | x_t)$$

$$P(e_t | x_t)$$



"Global"

DYNAMIC
BAYESIAN
NETWORK
(DBN)