



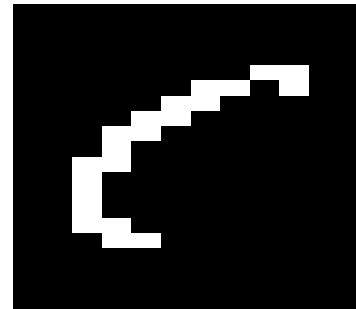
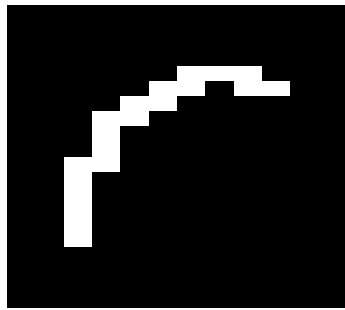
Hidden Markov Models

Machine Learning 10-601
April 15, 2009

Tom M. Mitchell
Machine Learning Department
Carnegie Mellon University

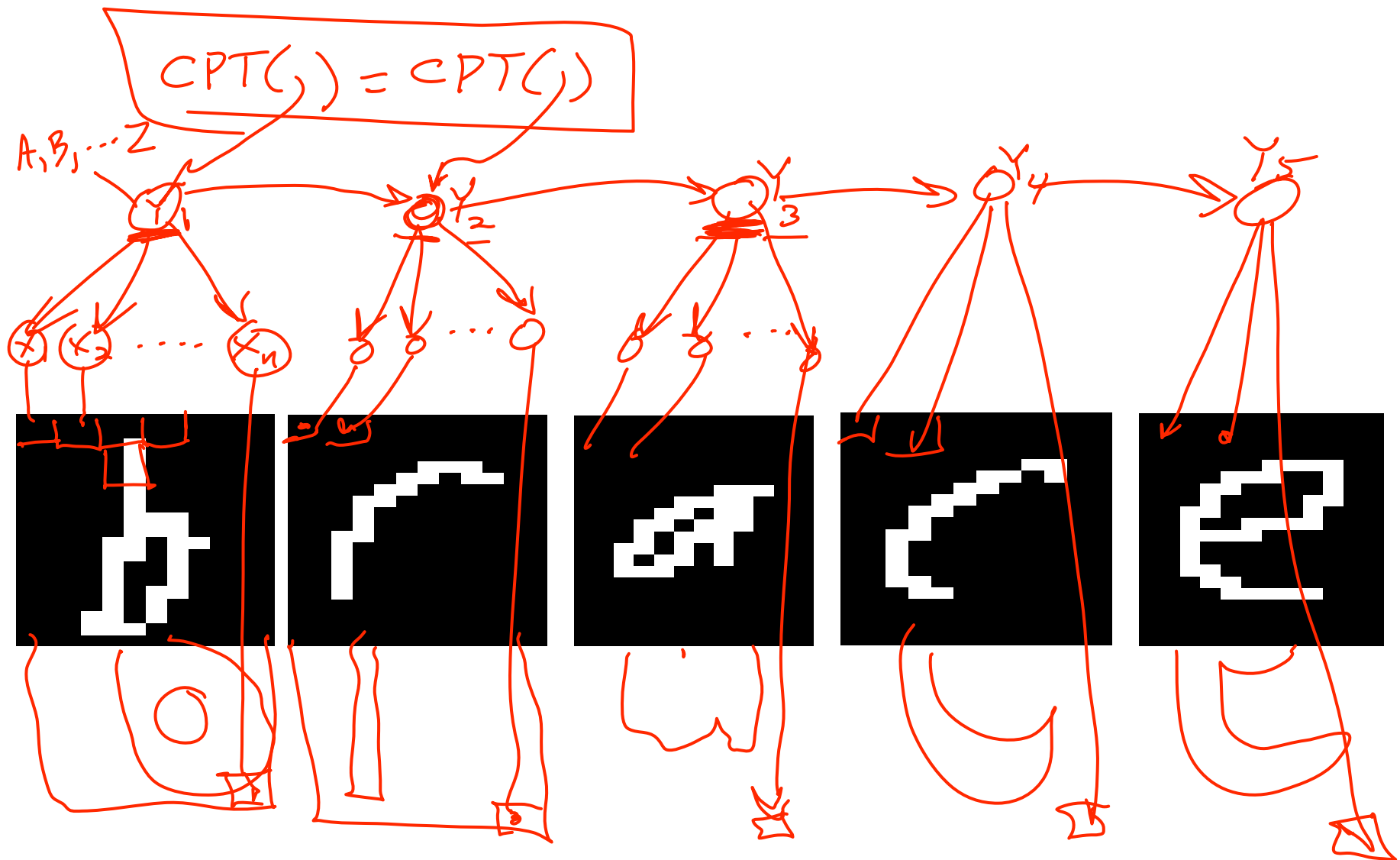
With thanks to Prof. Carlos Guestrin for some of these slides

Handwriting recognition

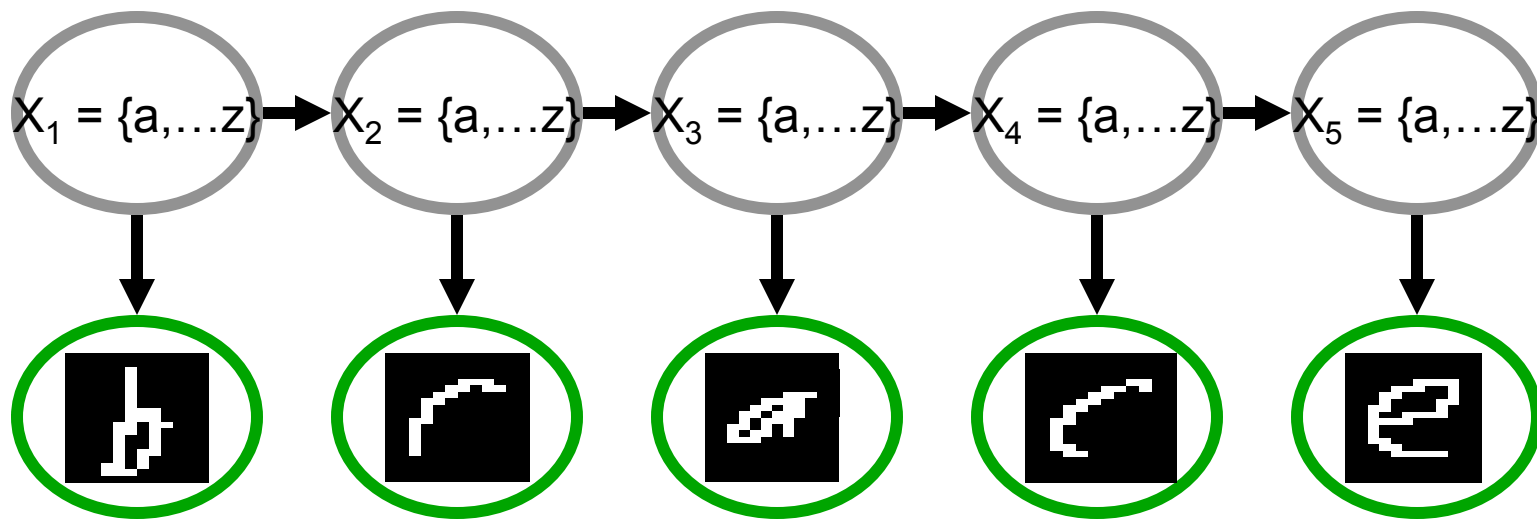


Character recognition, e.g., logistic regression

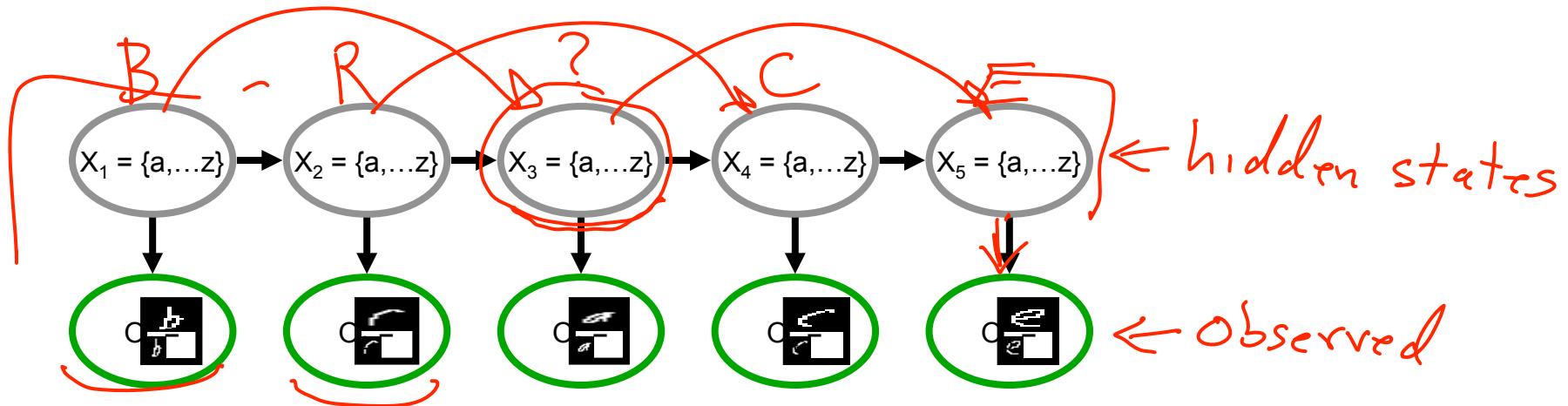
Example of a hidden Markov model (HMM)



Understanding HMM Semantics



HMMs semantics: Details



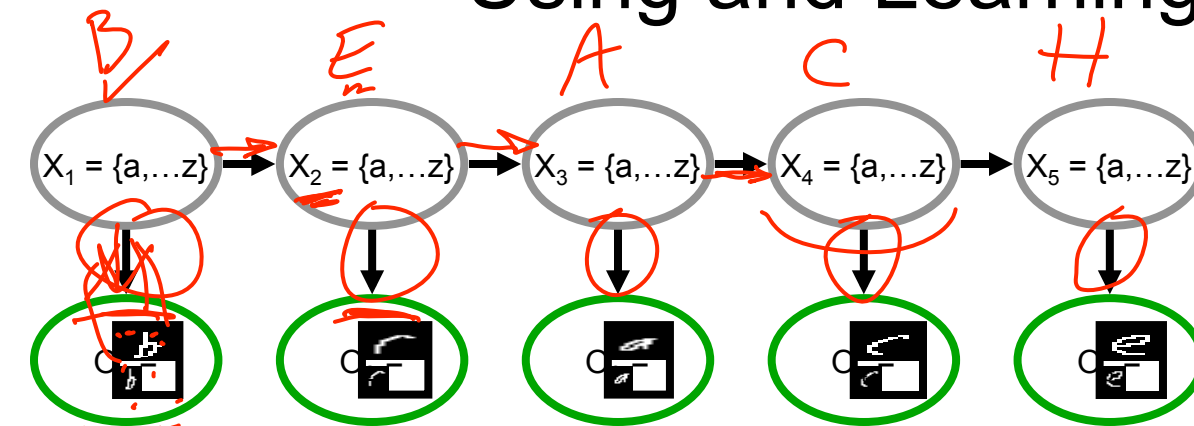
Just 3 distributions:

$$P(X_1)$$

$$P(X_i \mid X_{i-1})$$

$$P(O_i \mid X_i)$$

Using and Learning HMM's



$$P(X_i | o_{1..n})$$

$P(O_i | x_i)$ $\forall o, x, P(O_i = o | X_i = x)$
Core HMM questions:

$$P(o_1, \dots, o_5 | x_1 = B, x_2 = R, \dots, ACE)$$

1. How do we calculate $P(o_1, o_2, \dots, o_n)$?

2. How do we calculate argmax over $x_1, x_2, \dots, x_n$
of $P(x_1, x_2, \dots, x_n | o_1, o_2, \dots, o_n)$?

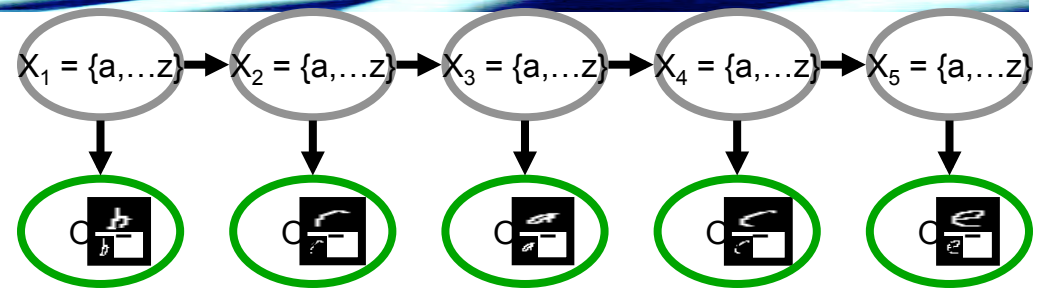
3. How do we train the HMM, given its structure and

3a. Fully observed training examples: $\langle x_1, \dots, x_n, o_1, \dots, o_n \rangle$

3b. Partially observed training examples: $\langle o_1, \dots, o_n \rangle$

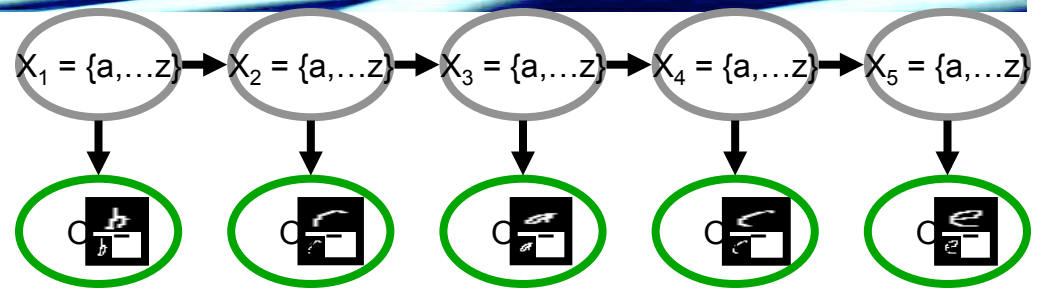
$E: \text{est } P(x_1, x_2, \dots, x_k | \hat{o})$
 $M: \text{update CPTs}$

How do we generate a
random output
sequence following
the HMM $P(o_1, o_2, \dots, o_T)$



How do we compute

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



1. brute force:

$$P(x_i = k)$$

2. Forward algorithm (dynamic progr., variable elimination):

define $\alpha_t(k) = P(o_1, o_2, \dots, o_t, X_t = k)$

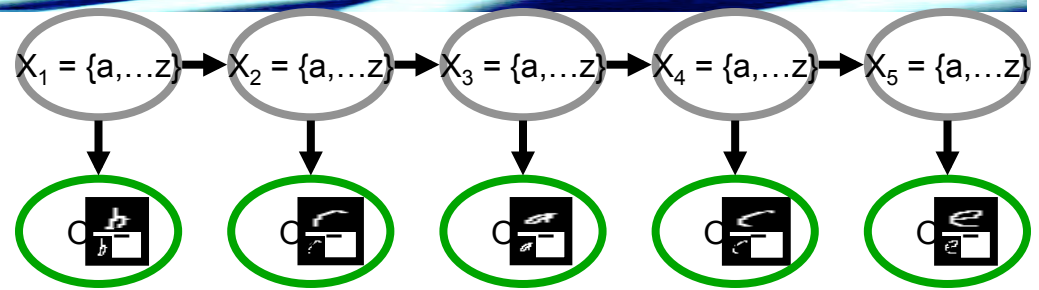
$$\alpha_1(k) = P(x_1 = k) P(o_1 | x_1 = k)$$

$$\alpha_{t+1}(k) = \left[\sum_{j=1}^N \alpha_t(j) P(x_{t+1} = k | x_t = j) \right] P(o_{t+1} = o_{t+1} | x_{t+1} = k)$$

$$P(o_1, o_2, \dots, o_T) = \sum_k \alpha_T(k)$$

k k th value of x

How do we compute
 $P(X_t = k | o_1, o_2, \dots, o_T)$

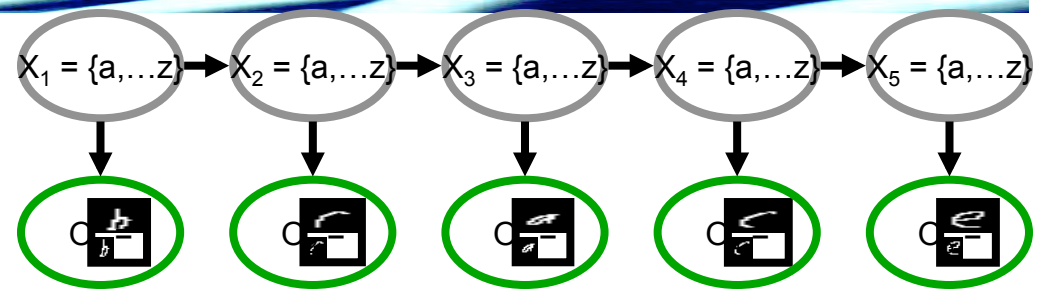


2. Backward algorithm (dynamic progr., variable elimination):

$$\alpha_t(k) = P(o_1, o_2, \dots, o_t, X_t = k)$$

define $\beta_t(k) = P(o_{t+1}, o_{t+2}, \dots, o_T | X_t = k)$

$$P(X_t = k | o_1, o_2, \dots, o_T) = \frac{P(X_t = k, o_1, o_2, \dots, o_T)}{P(o_1, o_2, \dots, o_T)}$$



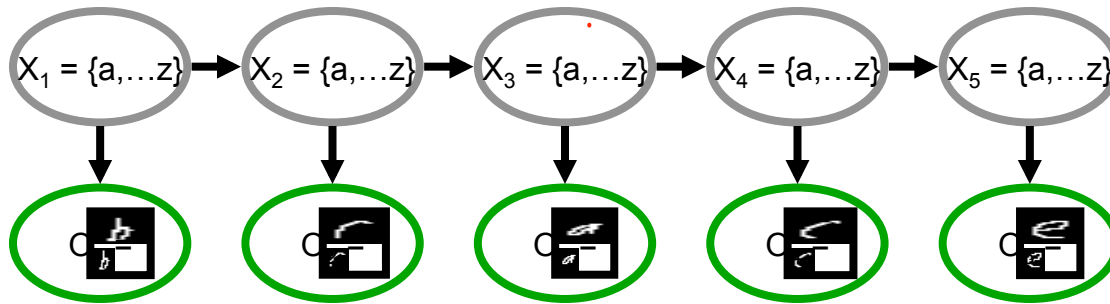
How do we compute

$$\arg \max_{x_1, \dots, x_T} P(\underbrace{x_1, \dots, x_T}_{\text{red bracket}} | o_1, o_2, \dots, o_T)$$

Viterbi algorithm, based on recursive computation of

$$\delta_t(k) = \max_{x_1, \dots, x_{t-1}} P(x_1, \dots, x_{t-1} X_t = k, o_1, o_2, \dots, o_t)$$

Learning HMMs from fully observable data: easy



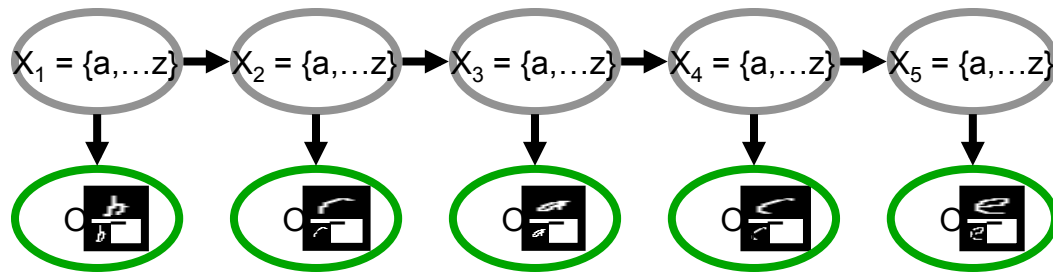
Learn 3 distributions:

$$P(X_1)$$

$$P(O_i \mid X_i)$$

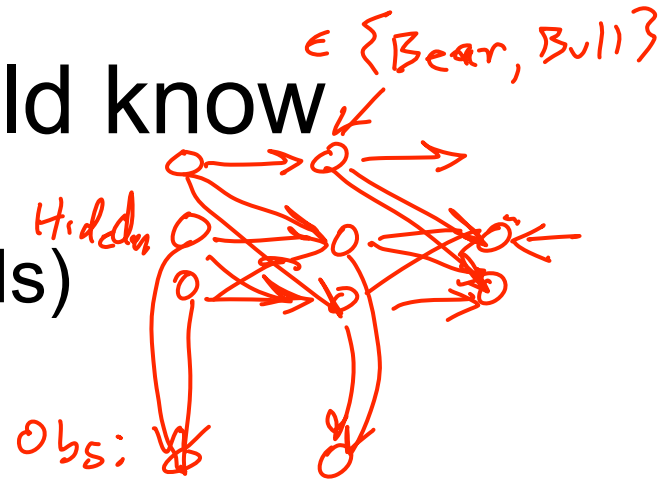
$$P(X_i \mid X_{i-1})$$

Learning HMMs when only observe $o_1 \dots o_T$



What you should know

- Hidden Markov models (HMMs)
 - Very useful, very powerful!
 - Speech, OCR, time series, ...
 - Parameter sharing, only learn 3 distributions
 - Dynamic programming/variable elimination reduces inference costs
 - Special case of Bayes net
 - Special case of Dynamic Bayesian Networks



Stationary vs Non-stationary

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

unchanging
across time

Thanks to Carlos Guestrin for many slides