

CSCI567 Machine Learning (Fall 2025)

Haipeng Luo

University of Southern California

Oct 31, 2025

Administration

Reminder: HW3 is due on Nov 5th.

Outline

- 1 Principal Component Analysis (PCA)
- 2 Markov models
- 3 Hidden Markov Model

Outline

- 1 Principal Component Analysis (PCA)
- 2 Markov models
- 3 Hidden Markov Model

Dimensionality reduction

Dimensionality reduction is yet another important unsupervised learning problem.

Goal: reduce the dimensionality of a dataset so

- it is **easier to visualize and discover patterns**
- it **takes less time and space** to process for any applications (classification, regression, clustering, etc)
- **noise is reduced**
- ...

There are many approaches, we focus on a linear method:

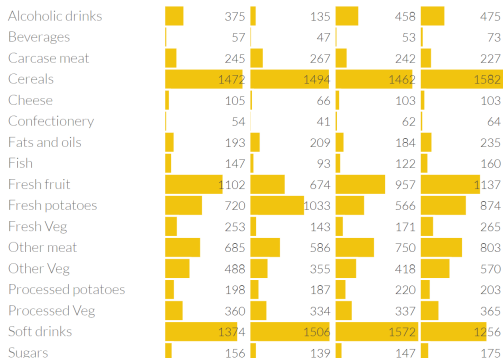
Principal Component Analysis (PCA)

Example

picture from here

Consider the following dataset:

- 17 features, each represents the average consumption of some food
- 4 data points, each represents some country



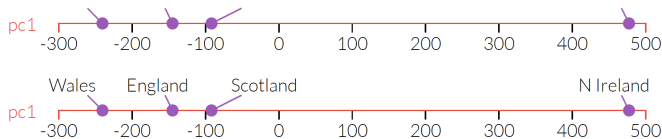
What can you tell?

Hard to say anything looking at all these 17 features.

Example

picture from here

PCA can help us! The **first principal component** of this dataset:



i.e. we **reduce the dimensionality** from 17 to just 1.

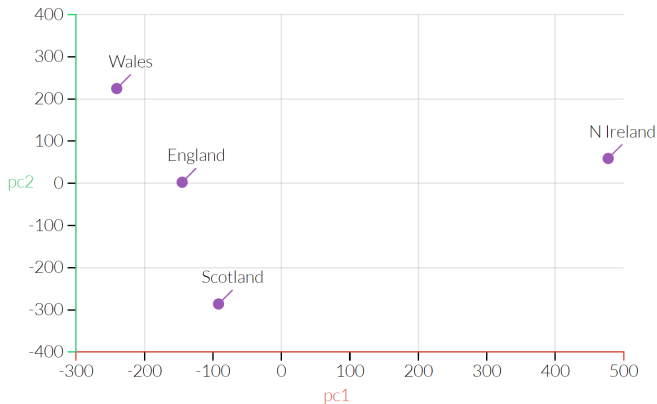
Now one data point is clearly different from the rest!

That turns out to be data from **Northern Ireland**, *the only country not on the island of Great Britain out of the 4 samples.*

Example

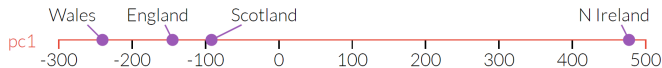
[picture from here](#)

PCA can find the **second (and more) principal component** of the data too:



High level idea

How does PCA find these principal components (PC)?



The first PC is in fact **the direction with the most variance**, i.e. the direction where the data is most spread out.

Finding the first PC

More formally, we want to find a direction $\mathbf{v} \in \mathbb{R}^D$ with $\|\mathbf{v}\|_2 = 1$, so that the **projection of the dataset on this direction has the most variance**, i.e.

$$\max_{\mathbf{v}: \|\mathbf{v}\|_2=1} \sum_{n=1}^N \left(\mathbf{x}_n^T \mathbf{v} - \frac{1}{N} \sum_m \mathbf{x}_m^T \mathbf{v} \right)^2$$

- $\mathbf{x}_n^T \mathbf{v}$ is exactly the **projection of $\mathbf{x}_n$ onto the direction $\mathbf{v}$**
- if we **pre-center the data**, i.e. let $\mathbf{x}'_n = \mathbf{x}_n - \frac{1}{N} \sum_m \mathbf{x}_m$, then the objective simply becomes

$$\max_{\mathbf{v}: \|\mathbf{v}\|_2=1} \sum_{n=1}^N \left(\mathbf{x}'_n{}^T \mathbf{v} \right)^2 = \max_{\mathbf{v}: \|\mathbf{v}\|_2=1} \mathbf{v}^T \left(\sum_{n=1}^N \mathbf{x}'_n \mathbf{x}'_n{}^T \right) \mathbf{v}$$

- we will simply assume $\{\mathbf{x}_n\}$ is centered (to avoid notation $\mathbf{x}'_n$)

Finding the first PC

With $\mathbf{X} \in \mathbb{R}^{N \times D}$ being the data matrix (as in Lec 2), we want

$$\begin{aligned} \max_{\mathbf{v}: \|\mathbf{v}\|_2=1} \mathbf{v}^\top (\mathbf{X}^\top \mathbf{X}) \mathbf{v} &= \max_{\mathbf{v}: \|\mathbf{v}\|_2=1} \mathbf{v}^\top \mathbf{U}^\top \begin{bmatrix} \lambda_1 & 0 & \cdots & 0 \\ 0 & \lambda_2 & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & \cdots & 0 & \lambda_D \end{bmatrix} \mathbf{U} \mathbf{v} \\ &= \max_{\mathbf{v}: \|\mathbf{v}\|_2=1} \sum_{i=1}^D (\mathbf{u}_i^\top \mathbf{v})^2 \lambda_i, \end{aligned}$$

where the columns of $\mathbf{U}^\top$ are unit **eigenvectors** $\mathbf{u}_1, \dots, \mathbf{u}_D$ of $\mathbf{X}^\top \mathbf{X}$ with corresponding **eigenvalues** $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_D \geq 0$.

Note: $\sum_{i=1}^D (\mathbf{u}_i^\top \mathbf{v})^2 = \mathbf{v}^\top \left(\sum_{i=1}^D \mathbf{u}_i \mathbf{u}_i^\top \right) \mathbf{v} = \mathbf{v}^\top \mathbf{v} = 1$, so the maximum above is λ_1 , realized by $\mathbf{v} = \mathbf{u}_1$.

Conclusion: the first PC is the top eigenvector of the covariance matrix!

Finding the other PCs

If v_1 is the first PC, then the **second PC** is found via

$$\max_{v_2: \|v_2\|_2=1, v_1^T v_2=0} v_2^T (X^T X) v_2$$

i.e. the direction that maximizes the variance among all other dimensions

This is just the **second top eigenvector of the covariance matrix!**

Conclusion: the d -th principal component is the d -th eigenvector (sorted by the eigenvalue from largest to smallest).

PCA

Input: a dataset represented as $\mathbf{X}$, #components p we want

Step 1 Center the data by subtracting the mean

Step 2 Find the top p eigenvectors (with unit norm) of the covariance matrix $\mathbf{X}^T \mathbf{X}$, denoted by $\mathbf{V} \in \mathbb{R}^{D \times p}$

Step 3 Construct the new compressed dataset $\mathbf{XV} \in \mathbb{R}^{N \times p}$

(Optional) Step 4 “Reconstruct” original dataset as $\mathbf{XVV}^T \in \mathbb{R}^{N \times D}$

How many PCs do we want?

One common rule: pick p large enough so it **covers about 90% of the spectrum** (so 90% of variance is explained), i.e.

$$\frac{\sum_{d=1}^p \lambda_d}{\sum_{d=1}^D \lambda_d} \geq 90\%$$

where $\lambda_1 \geq \dots \geq \lambda_N$ are sorted eigenvalues.

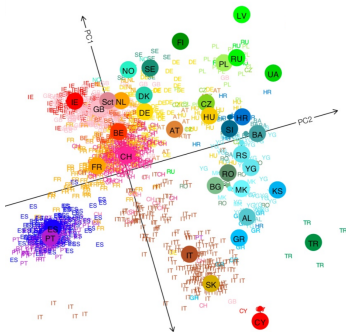
Note: $\sum_{d=1}^D \lambda_d = \text{Tr}(\mathbf{X}^T \mathbf{X})$, so no need to actually find all eigenvalues.

For **visualization**, also often pick $p = 1$ or $p = 2$.

Another visualization example

A famous study of **genetic map**

- dataset: **genomes of 1,387 Europeans**
- first 2 PCs shown below; *looks remarkably like the geographic map*



Genetic variation is highly structured by geography; top PCs correspond to geographic factors.

Compression Example

Dataset: 256×256 ($\approx 65K$ pixels)
dimensional images of about 2500
faces, all framed similarly
($N = 2500$, $D \approx 65,000$)

Even with $p \approx 100$, reconstructed
images $\mathbf{X}\mathbf{V}\mathbf{V}^T \in \mathbb{R}^{N \times D}$ look very
much similar to original images

The principal components (called
eigenfaces $\in \mathbb{R}^D$) are themselves
interpretable too!

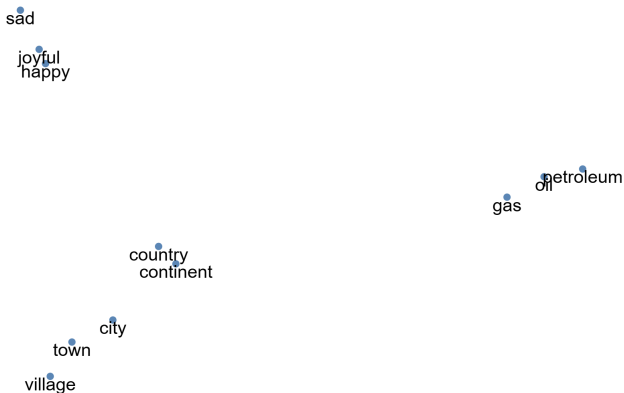
PCA finds the subspace (out of the
space of all 256×256 images) where
faces lie



Figure 2. Seven of the eigenfaces calculated from the input images of Figure 1.

Word embedding via PCA

Create **meaningful vector representation** of words; see project.



Outline

- 1 Principal Component Analysis (PCA)
- 2 Markov models
 - Markov chain
 - Learning Markov models
- 3 Hidden Markov Model

Sequential data and Markov models

Sequential data (i.e., each x_n is a *sequence*) are ubiquitous:

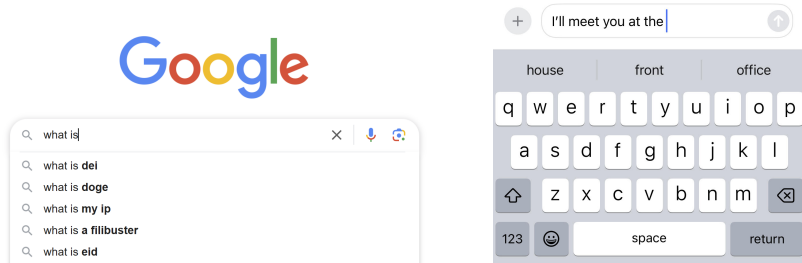
- text or speech data
- stock market data
- gene data

Markov models are powerful probabilistic tools to analyze them:

- not the state-of-the-art for e.g. Natural Language Processing (NLP)
- but still extremely useful in many areas (in particular, it is the foundation of *reinforcement learning*)

A motivating example

Next word prediction



Can be formulated as predicting the probability of the next word/state:

$$P(Z_{t+1} \mid Z_1, \dots, Z_t)?$$

abbreviated as $P(Z_{t+1} \mid \mathbf{Z}_{1:t})$ (i.e., $\mathbf{Z}_{1:t}$ denotes the sequence $Z_1, \dots, Z_t$).

Definition of a Markov model/chain

A **Markov chain** is a stochastic process with **Markov property**: a sequence of random variables $Z_1, Z_2, \dots$ s.t.

$$P(Z_{t+1} \mid Z_{1:t}) = P(Z_{t+1} \mid Z_t) \quad (\text{Markov property})$$

i.e. *the current state only depends on the most recent state*

We only consider the following case:

- All Z_t 's take value from the same **discrete** set $\{1, \dots, S\}$
- $P(Z_{t+1} = s' \mid Z_t = s) = a_{s,s'}$, known as **transition probability**
- $P(Z_1 = s) = \pi_s$
- $(\{\pi_s\}, \{a_{s,s'}\}) = (\boldsymbol{\pi}, \mathbf{A})$ are **parameters of the model**

Examples

- Example 1 (**Language model**)

States $[S]$ represent a dictionary of words,

$$a_{\text{ice,cream}} = P(Z_{t+1} = \text{cream} \mid Z_t = \text{ice})$$

is an example of the transition probability.

- Example 2 (**Weather**)

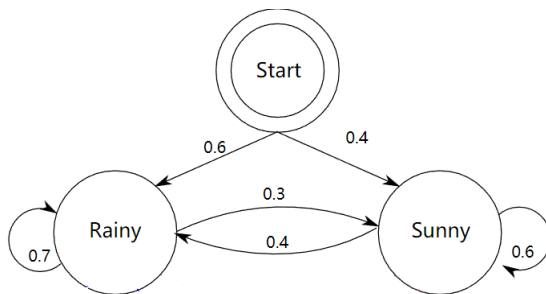
States $[S]$ represent weather at each day

$$a_{\text{sunny,rainy}} = P(Z_{t+1} = \text{rainy} \mid Z_t = \text{sunny})$$

Graph representation

picture from Wikipedia

It is intuitive to represent a Markov model as a **graph**



High-order Markov chain

Is the Markov assumption reasonable? Not completely for the language model for example:

$$\begin{aligned} P(Z_{t+1} = \text{office} \mid Z_{1:t} = \text{I'll meet you at the}) \\ = P(Z_{t+1} = \text{office} \mid Z_t = \text{the}) \end{aligned} \quad (\text{bigram, unreasonable})$$

Higher-order Markov chains make it more reasonable, e.g.

$$P(Z_{t+1} \mid Z_{1:t}) = P(Z_{t+1} \mid Z_t, Z_{t-1}) \quad (\text{second-order Markov})$$

i.e. the current word only depends on the last two words (the **trigram** model).

n -gram model

$$\begin{aligned} &P(Z_{t+1} = \text{office} \mid Z_{1:t} = \text{I'll meet you at the}) \\ &= P(Z_{t+1} = \text{office} \mid Z_{t-1:t} = \text{at the}) \end{aligned} \quad (\text{trigram})$$

$$\begin{aligned} &P(Z_{t+1} = \text{office} \mid Z_{1:t} = \text{I'll meet you at the}) \\ &= P(Z_{t+1} = \text{office} \mid Z_{t-2:t} = \text{you at the}) \end{aligned} \quad (\text{4-gram})$$

$$\begin{aligned} &P(Z_{t+1} = \text{office} \mid Z_{1:t} = \text{I'll meet you at the}) \\ &= P(Z_{t+1} = \text{office} \mid Z_{t-3:t} = \text{meet you at the}) \end{aligned} \quad (\text{5-gram})$$

Learning higher-order Markov chains is similar, but more *expensive*.

For simplicity, we will only consider standard Markov chains.

Learning from examples

Now suppose we have observed N sequences of examples:

- $z_{1,1}, \dots, z_{1,T}$
- $\dots$
- $z_{n,1}, \dots, z_{n,T}$
- $\dots$
- $z_{N,1}, \dots, z_{N,T}$

where

- for simplicity we assume each sequence has the same length T
- lower case $z_{n,t}$ represents the value of the random variable $Z_{n,t}$

From these observations how do we *learn the model parameters* $(\pi, \mathbf{A})$?

Finding the MLE

Same story, find the **MLE**. The log-likelihood of a sequence $z_1, \dots, z_T$ is

$$\begin{aligned} \ln P(Z_{1:T} = z_{1:T}) \\ = \sum_{t=1}^T \ln P(Z_t = z_t \mid Z_{1:t-1} = z_{1:t-1}) \end{aligned} \quad (\text{always true})$$

$$= \sum_{t=1}^T \ln P(Z_t = z_t \mid Z_{t-1} = z_{t-1}) \quad (\text{Markov property})$$

$$= \ln \pi_{z_1} + \sum_{t=2}^T \ln a_{z_{t-1}, z_t}$$

$$= \sum_s \mathbb{I}[z_1 = s] \ln \pi_s + \sum_{s, s'} \left(\sum_{t=2}^T \mathbb{I}[z_{t-1} = s, z_t = s'] \right) \ln a_{s, s'}$$

Finding the MLE

Summing over all N sequences, we obtain the joint log-likelihood:

$$L(\boldsymbol{\pi}, \mathbf{A}) = \sum_s (\text{\#initial states with value } s) \ln \pi_s \\ + \sum_{s,s'} (\text{\#transitions from } s \text{ to } s') \ln a_{s,s'}$$

The MLE $\operatorname{argmax}_{\boldsymbol{\pi}, \mathbf{A}} L(\boldsymbol{\pi}, \mathbf{A})$ is:

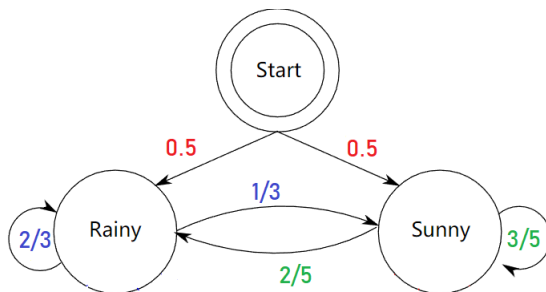
$$\pi_s \propto \text{\#initial states with value } s \\ a_{s,s'} \propto \text{\#transitions from } s \text{ to } s'$$

Example

Suppose we observed the following 2 sequences of length 5

- sunny, sunny, rainy, rainy, rainy
- rainy, sunny, sunny, sunny, rainy

MLE is the following model



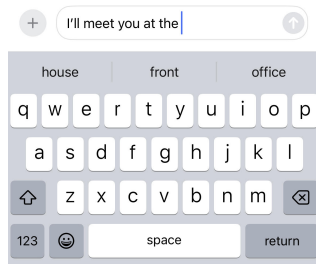
n -gram example

Recall: I'll meet you at the _____.

Suppose

- we use a 5-gram model
- “meet you at the” appears 1000 times in the training set:
 - “meet you at the house” appears 500 times
 - “meet you at the front” appears 300 times
 - “meet you at the office” appears 200 times

Then the model predicts the next word as “house” with prob. 0.5, “front” with prob. 0.3, and “office” with prob. 0.2.



Text generation

adapted from Stanford CS224n

After learning the model, can also **generate texts** by repeatedly sampling conditioning on what have been generated:

- today the _____
- today the price _____
- today the price of _____
- today the price of gold _____

final result: *today the price of gold per ton, while production of shoe lasts and shoe industry, the bank intervened just after it considered and rejected an IMF demand to rebuild depleted European stocks.*

Surprisingly grammatical! but *incoherent*...

company	0.153
bank	0.153
price	0.077
italian	0.039
emirate	0.039
...	

of	0.308
for	0.050
it	0.046
to	0.046
is	0.031
...	

the	0.072
18	0.043
oil	0.043
its	0.036
gold	0.018
...	

Outline

- 1 Principal Component Analysis (PCA)
- 2 Markov models
- 3 Hidden Markov Model
 - Inferring HMMs
 - Learning HMMs

Markov model with outcomes/observations

Now suppose each state Z_t also “emits” some **outcome/observation** $X_t \in [O]$ based on the following model

$$P(X_t = o \mid Z_t = s) = b_{s,o} \quad (\text{emission probability})$$

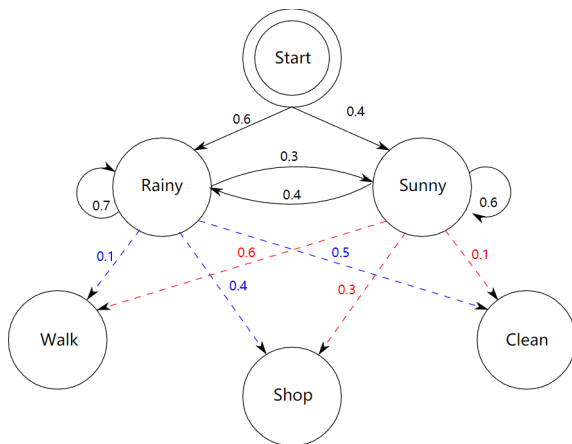
independent of anything else.

The model parameters are $(\{\pi_s\}, \{a_{s,s'}\}, \{b_{s,o}\}) = (\boldsymbol{\pi}, \boldsymbol{A}, \boldsymbol{B})$.

Another (toy) running example

picture from Wikipedia

On each day, we also observe **Bob's activity: walk, shop, or clean**, which only depends on the weather of that day.



Joint likelihood

The joint log-likelihood of a **state-outcome sequence** $z_1, x_1, \dots, z_T, x_T$ is

$$\begin{aligned}
 & \ln P(Z_{1:T} = z_{1:T}, X_{1:T} = x_{1:T}) \\
 &= \ln P(Z_{1:T} = z_{1:T}) + \ln P(X_{1:T} = x_{1:T} \mid Z_{1:T} = z_{1:T}) \quad (\text{always true}) \\
 &= \sum_{t=1}^T \ln P(Z_t = z_t \mid Z_{t-1} = z_{t-1}) + \sum_{t=1}^T \ln P(X_t = x_t \mid Z_t = z_t) \\
 & \hspace{15em} (\text{due to all the independence}) \\
 &= \ln \pi_{z_1} + \sum_{t=2}^T \ln a_{z_{t-1}, z_t} + \sum_{t=1}^T \ln b_{z_t, x_t}
 \end{aligned}$$

Learning the model

If we observe N state-outcome sequences: $z_{n,1}, x_{n,1}, \dots, z_{n,T}, x_{n,T}$ for $n = 1, \dots, N$, the MLE is again very simple (verify yourself):

$$\pi_s \propto \text{\textcolor{blue}{\#initial states with value } } s$$

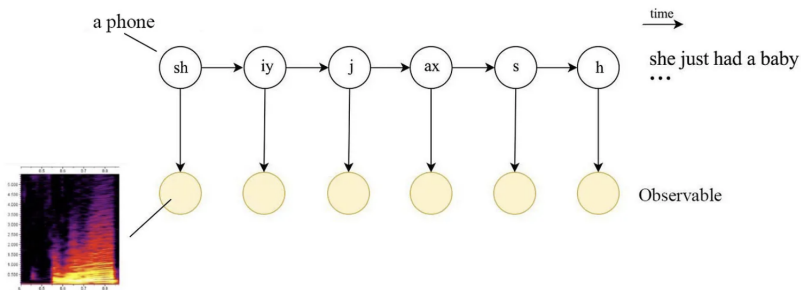
$$a_{s,s'} \propto \text{\textcolor{blue}{\#transitions from } } s \text{ to } s'$$

$$b_{s,o} \propto \text{\textcolor{blue}{\#state-outcome pairs } } (s, o)$$

Hidden Markov models

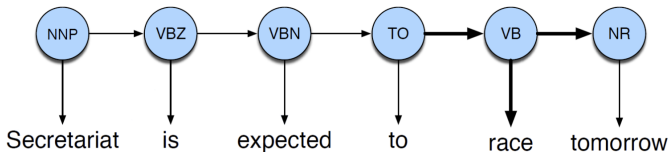
However, *most often we do not observe the states!* Hence the name **Hidden Markov Model (HMM)**

- **speech recognition**: observe the speech $X_{1:T}$ but not the underlying words/phones $Z_{1:T}$



Hidden Markov models (cont.)

- **Part of Speech (POS) tagging**: observe the sentence $X_{1:T}$ but not the underlying categories $Z_{1:T}$ (noun, verb, adjective, adverb, etc.)



See programming project!

Learning HMMs

How to learn HMMs? **Roadmap:**

- first discuss how to **infer** when the model is known (key: **dynamic programming**)
- then discuss how to **learn** the model (key: **EM**)

What can we infer about an HMM?

Knowing the parameter of an HMM, we can infer

- **most likely hidden states path, given an observation sequence**

$$\operatorname{argmax}_{z_{1:T}} P(Z_{1:T} = z_{1:T} \mid X_{1:T} = x_{1:T})$$

e.g. given Bob's activities for one week, what's the most likely weather for this week?

- **the probability of observing some sequence**

$$P(X_{1:T} = x_{1:T})$$

e.g. prob. of observing Bob's activities "walk, walk, shop, clean, walk, shop, shop" for one week

What can we infer for a known HMM?

Knowing the parameter of an HMM, we can infer

- **the state at some point, given an observation sequence**

$$P(Z_t = s \mid X_{1:T} = x_{1:T})$$

e.g. given Bob's activities for one week, how was the weather like on Wed?

- **the transition at some point, given an observation sequence**

$$P(Z_t = s, Z_{t+1} = s' \mid X_{1:T} = x_{1:T})$$

e.g. given Bob's activities for one week, how was the weather like on Wed and Thu?

Forward and backward messages

The key to infer all these is to compute two things:

- **forward messages**: for each s and t

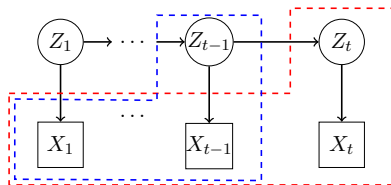
$$\alpha_s(t) = P(Z_t = s, X_{1:t} = x_{1:t})$$

- **backward messages**: for each s and t

$$\beta_s(t) = P(X_{t+1:T} = x_{t+1:T} \mid Z_t = s)$$

Computing forward messages

Key: *establish recursion*



$$\begin{aligned}
 & \alpha_s(t) \\
 &= P(Z_t = s, X_{1:t} = x_{1:t}) \\
 &= \sum_{s'} P(Z_t = s, Z_{t-1} = s', X_{1:t} = x_{1:t}) && \text{(marginalizing)} \\
 &= \sum_{s'} P(Z_{t-1} = s', X_{1:t-1} = x_{1:t-1}) \\
 &\quad \times P(Z_t = s, X_t = x_t \mid Z_{t-1} = s', X_{1:t-1} = x_{1:t-1}) \\
 &= b_{s,x_t} \sum_{s'} a_{s',s} \alpha_{s'}(t-1) && \text{(by the model definition)}
 \end{aligned}$$

Base case: $\alpha_s(1) = P(Z_1 = s, X_1 = x_1) = \pi_s b_{s,x_1}$

Forward procedure

Forward procedure

For all $s \in [S]$, compute $\alpha_s(1) = \pi_s b_{s,x_1}$.

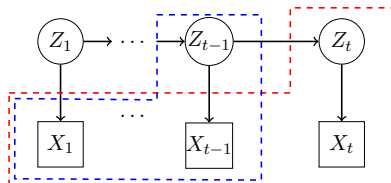
For $t = 2, \dots, T$

- for each $s \in [S]$, compute

$$\alpha_s(t) = b_{s,x_t} \sum_{s'} a_{s',s} \alpha_{s'}(t-1)$$

It takes

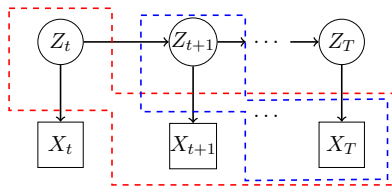
- $O(S^2T)$ time
- $O(ST)$ space



Computing backward messages

Again, establish a recursion

$$\begin{aligned}
 & \beta_s(t) \\
 &= P(X_{t+1:T} = x_{t+1:T} \mid Z_t = s) \\
 &= \sum_{s'} P(X_{t+1:T} = x_{t+1:T}, Z_{t+1} = s' \mid Z_t = s) && \text{(marginalizing)} \\
 &= \sum_{s'} P(Z_{t+1} = s' \mid Z_t = s) P(X_{t+1:T} = x_{t+1:T} \mid Z_{t+1} = s', Z_t = s) \\
 &= \sum_{s'} a_{s,s'} P(X_{t+1} = x_{t+1} \mid Z_{t+1} = s') P(X_{t+2:T} = x_{t+2:T} \mid Z_{t+1} = s') \\
 &= \sum_{s'} a_{s,s'} b_{s',x_{t+1}} \beta_{s'}(t+1)
 \end{aligned}$$



Base case: $\beta_s(T) = 1$

Backward procedure

Backward procedure

For all $s \in [S]$, set $\beta_s(T) = 1$.

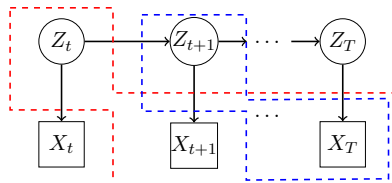
For $t = T - 1, \dots, 1$

- for each $s \in [S]$, compute

$$\beta_s(t) = \sum_{s'} a_{s,s'} b_{s',x_{t+1}} \beta_{s'}(t+1)$$

Again it takes

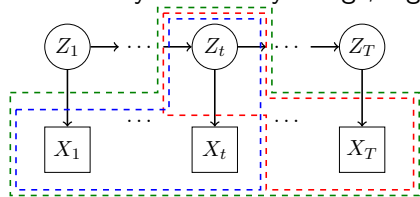
- $O(S^2T)$ time
- $O(ST)$ space



Using forward and backward messages

With forward and backward messages, we can easily infer many things, e.g.

$$\begin{aligned}
 \gamma_s(t) &= P(Z_t = s \mid X_{1:T} = x_{1:T}) \\
 &\propto P(Z_t = s, X_{1:T} = x_{1:T}) \\
 &= P(Z_t = s, X_{1:t} = x_{1:t}) P(X_{t+1:T} = x_{t+1:T} \mid Z_t = s, X_{1:t} = x_{1:t}) \\
 &= \alpha_s(t) \beta_s(t)
 \end{aligned}$$



What constant are we omitting in “ $\propto$ ”? It is exactly

$$P(X_{1:T} = x_{1:T}) = \sum_s \alpha_s(t) \beta_s(t),$$

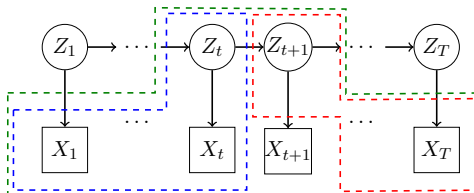
the probability of observing the sequence $x_{1:T}$.

This is true for any t ; a good way to check correctness of your code.

Using forward and backward messages

Another example: the conditional probability of transition s to s' at time t

$$\begin{aligned}
 & \xi_{s,s'}(t) \\
 &= P(Z_t = s, Z_{t+1} = s' \mid X_{1:T} = x_{1:T}) \\
 &\propto P(Z_t = s, Z_{t+1} = s', X_{1:T} = x_{1:T}) \\
 &= P(Z_t = s, X_{1:t} = x_{1:t}) P(Z_{t+1} = s', X_{t+1:T} = x_{t+1:T} \mid Z_t = s, X_{1:t} = x_{1:t}) \\
 &= \alpha_s(t) P(Z_{t+1} = s' \mid Z_t = s) P(X_{t+1:T} = x_{t+1:T} \mid Z_{t+1} = s') \\
 &= \alpha_s(t) a_{s,s'} P(X_{t+1} = x_{t+1} \mid Z_{t+1} = s') P(X_{t+2:T} = x_{t+2:T} \mid Z_{t+1} = s') \\
 &= \alpha_s(t) a_{s,s'} b_{s',x_{t+1}} \beta_{s'}(t+1)
 \end{aligned}$$



The normalization constant is in fact again $P(X_{1:T} = x_{1:T})$

Finding the most likely path

Can't use forward and backward messages directly to find the most likely path, but it is **very similar to the forward procedure**. Key: compute

$$\delta_s(t) = \max_{z_{1:t-1}} P(Z_t = s, Z_{1:t-1} = z_{1:t-1}, X_{1:t} = x_{1:t})$$

the probability of the most likely path for time $1 : t$ ending at state s

Computing $\delta_s(t)$

Observe

$$\begin{aligned}
 \delta_s(t) &= \max_{z_{1:t-1}} P(Z_t = s, Z_{1:t-1} = z_{1:t-1}, X_{1:t} = x_{1:t}) \\
 &= \max_{s'} \max_{z_{1:t-2}} P(Z_t = s, Z_{t-1} = s', Z_{1:t-2} = z_{1:t-2}, X_{1:t} = x_{1:t}) \\
 &= \max_{s'} P(Z_t = s \mid Z_{t-1} = s') P(X_t = x_t \mid Z_t = s) \cdot \\
 &\quad \max_{z_{1:t-2}} P(Z_{t-1} = s', Z_{1:t-2} = z_{1:t-2}, X_{1:t-1} = x_{1:t-1}) \\
 &= b_{s,x_t} \max_{s'} a_{s',s} \delta_{s'}(t-1) \quad (\text{recursion!})
 \end{aligned}$$

Base case: $\delta_s(1) = P(Z_1 = s, X_1 = x_1) = \pi_s b_{s,x_1}$

Exactly the same as forward messages except replacing “sum” by “max”!

Viterbi Algorithm (!)

Viterbi Algorithm

For each $s \in [S]$, compute $\delta_s(1) = \pi_s b_{s,x_1}$.

For each $t = 2, \dots, T$,

- for each $s \in [S]$, compute

$$\delta_s(t) = b_{s,x_t} \max_{s'} a_{s',s} \delta_{s'}(t-1),$$

$$\Delta_s(t) = \operatorname{argmax}_{s'} a_{s',s} \delta_{s'}(t-1).$$

Backtracking: let $z_T^* = \operatorname{argmax}_s \delta_s(T)$.

For each $t = T, \dots, 2$: set $z_{t-1}^* = \Delta_{z_t^*}(t)$.

Output the most likely path $z_1^*, \dots, z_T^*$.

Example: $x_{1:4} = \text{clean, shop, walk, walk}$

Viterbi Algorithm

For each $s \in [S]$, compute $\delta_s(1) = \pi_s b_{s,x_1}$.

...

$$\delta_{\text{sunny}}(1) = \pi_{\text{sunny}} b_{\text{sunny}, \text{clean}}$$

$$\delta_{\text{rainy}}(1) = \pi_{\text{rainy}} b_{\text{rainy}, \text{clean}}$$

$$\delta_{\text{sunny}}(1) = 0.25$$

$$\delta_{\text{rainy}}(1) = 0.4$$

(numbers are made up here)

Example: $x_{1:4} = \text{clean, shop, walk, walk}$

Viterbi Algorithm

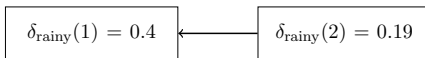
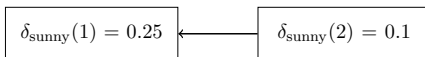
...

For each $t = 2, \dots, T$ and each $s \in [S]$, compute

$$\delta_s(t) = b_{s,x_t} \max_{s'} a_{s',s} \delta_{s'}(t-1), \quad \Delta_s(t) = \operatorname{argmax}_{s'} a_{s',s} \delta_{s'}(t-1).$$

$$\delta_{\text{sunny}}(2) = b_{\text{sunny},\text{shop}} \max_{s' \in \{\text{sunny}, \text{rainy}\}} a_{s',\text{sunny}} \delta_{s'}(1)$$

$$\delta_{\text{rainy}}(2) = b_{\text{rainy},\text{shop}} \max_{s' \in \{\text{sunny}, \text{rainy}\}} a_{s',\text{rainy}} \delta_{s'}(1)$$



Arrows represent the “argmax”, i.e. $\Delta_s(t)$.

Example: $x_{1:4} = \text{clean, shop, walk, walk}$

Viterbi Algorithm

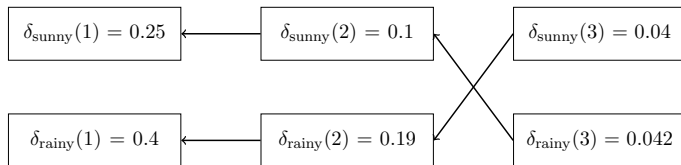
...

For each $t = 2, \dots, T$ and each $s \in [S]$, compute

$$\delta_s(t) = b_{s,x_t} \max_{s'} a_{s',s} \delta_{s'}(t-1), \quad \Delta_s(t) = \operatorname{argmax}_{s'} a_{s',s} \delta_{s'}(t-1).$$

$$\delta_{\text{sunny}}(3) = b_{\text{sunny},\text{walk}} \max_{s' \in \{\text{sunny}, \text{rainy}\}} a_{s',\text{sunny}} \delta_{s'}(2)$$

$$\delta_{\text{rainy}}(3) = b_{\text{rainy},\text{walk}} \max_{s' \in \{\text{sunny}, \text{rainy}\}} a_{s',\text{rainy}} \delta_{s'}(2)$$



Arrows represent the "argmax", i.e. $\Delta_s(t)$.

Example: $x_{1:4} = \text{clean, shop, walk, walk}$

Viterbi Algorithm

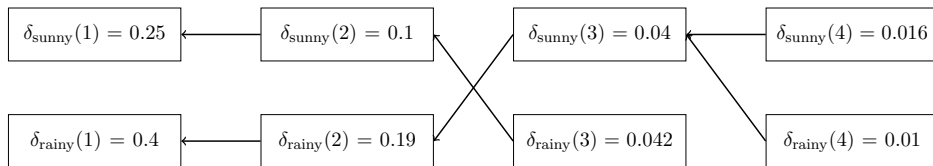
...

For each $t = 2, \dots, T$ and each $s \in [S]$, compute

$$\delta_s(t) = b_{s,x_t} \max_{s'} a_{s',s} \delta_{s'}(t-1), \quad \Delta_s(t) = \operatorname{argmax}_{s'} a_{s',s} \delta_{s'}(t-1).$$

$$\delta_{\text{sunny}}(4) = b_{\text{sunny},\text{walk}} \max_{s' \in \{\text{sunny}, \text{rainy}\}} a_{s',\text{sunny}} \delta_{s'}(3)$$

$$\delta_{\text{rainy}}(4) = b_{\text{rainy},\text{walk}} \max_{s' \in \{\text{sunny}, \text{rainy}\}} a_{s',\text{rainy}} \delta_{s'}(3)$$



Arrows represent the “argmax”, i.e. $\Delta_s(t)$.

Example: $x_{1:4} = \text{clean, shop, walk, walk}$

Viterbi Algorithm

...

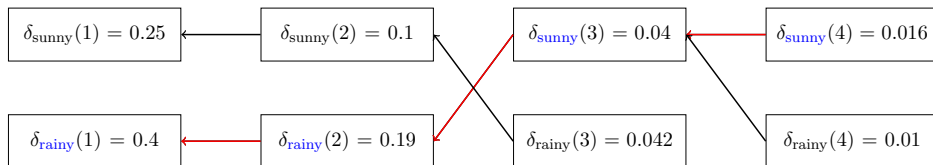
Backtracking: let $z_T^* = \operatorname{argmax}_s \delta_s(T)$.

For each $t = T, \dots, 2$: set $z_{t-1}^* = \Delta_{z_t^*}(t)$.

(just follow the arrow!)

Output the most likely path $z_1^*, \dots, z_T^*$.

$$\begin{aligned} z_4^* &= \operatorname{argmax}_{s \in \{\text{sunny}, \text{rainy}\}} \delta_s(4) = \text{sunny}, & z_3^* &= \Delta_{\text{sunny}}(4) = \text{sunny} \\ z_2^* &= \Delta_{\text{sunny}}(3) = \text{rainy}, & z_1^* &= \Delta_{\text{rainy}}(2) = \text{rainy} \end{aligned}$$



The most likely path is “rainy, rainy, sunny, sunny”

Learning the parameters of an HMM

All previous inferences depend on **knowing the parameters** $(\pi, \mathbf{A}, \mathbf{B})$.

How do we learn the parameters based on N observation sequences $x_{n,1}, \dots, x_{n,T}$ for $n = 1, \dots, N$?

MLE is **intractable due to the hidden variables** $Z_{n,t}$'s (similar to GMMs)

Need to apply **EM** again! Known as the **Baum–Welch algorithm**.

Applying EM: E-Step

Recall in the E-Step we fix the parameters and find the **posterior distributions q of the hidden states** (for each sample n), which leads to the complete log-likelihood:

$$\begin{aligned}
 & \mathbb{E}_{z_{1:T} \sim q} [\ln P(Z_{1:T} = z_{1:T}, X_{1:T} = x_{1:T})] \\
 &= \mathbb{E}_{z_{1:T} \sim q} \left[\ln \pi_{z_1} + \sum_{t=1}^{T-1} \ln a_{z_t, z_{t+1}} + \sum_{t=1}^T \ln b_{z_t, x_t} \right] \\
 &= \sum_s \gamma_s(1) \ln \pi_s + \sum_{t=1}^{T-1} \sum_{s, s'} \xi_{s, s'}(t) \ln a_{s, s'} + \sum_{t=1}^T \sum_s \gamma_s(t) \ln b_{s, x_t}
 \end{aligned}$$

We have discussed how to compute

$$\begin{aligned}
 \gamma_s(t) &= P(Z_t = s \mid X_{1:T} = x_{1:T}) \\
 \xi_{s, s'}(t) &= P(Z_t = s, Z_{t+1} = s' \mid X_{1:T} = x_{1:T})
 \end{aligned}$$

Applying EM: M-Step

The maximizer of complete log-likelihood is simply doing **weighted counting** (compared to the unweighted counting on Slide 36):

$$\pi_s \propto \sum_n \gamma_s^{(n)}(1) = \mathbb{E}_q [\text{\textcolor{blue}{\#initial states with value } } s]$$

$$a_{s,s'} \propto \sum_n \sum_{t=1}^{T-1} \xi_{s,s'}^{(n)}(t) = \mathbb{E}_q [\text{\textcolor{blue}{\#transitions from } } s \text{ to } s']$$

$$b_{s,o} \propto \sum_n \sum_{t: x_t=o} \gamma_s^{(n)}(t) = \mathbb{E}_q [\text{\textcolor{blue}{\#state-outcome pairs } } (s, o)]$$

where

$$\gamma_s^{(n)}(t) = P(Z_{n,t} = s \mid X_{n,1:T} = x_{n,1:T})$$

$$\xi_{s,s'}^{(n)}(t) = P(Z_{n,t} = s, Z_{n,t+1} = s' \mid X_{n,1:T} = x_{n,1:T})$$

(Recall how EM for GMM updates parameters using weighted averages.)

Baum–Welch algorithm

Step 0 Initialize the parameters (π, A, B)

Step 1 (E-Step) Fixing the parameters, **compute forward and backward messages for all sample sequences**, then use these to compute $\gamma_s^{(n)}(t)$ and $\xi_{s,s'}^{(n)}(t)$ for each n, t, s, s' (see Slides 47 and 48).

Step 2 (M-Step) Update parameters:

$$\pi_s \propto \sum_n \gamma_s^{(n)}(1), \quad a_{s,s'} \propto \sum_n \sum_{t=1}^{T-1} \xi_{s,s'}^{(n)}(t), \quad b_{s,o} \propto \sum_n \sum_{t: x_t=o} \gamma_s^{(n)}(t)$$

Step 3 Return to Step 1 if not converged

Summary

Very important models: **Markov chains**, **hidden Markov models**

Several algorithms:

- forward and backward procedures
- inferring HMMs based on forward and backward messages
- Viterbi algorithm and variants (see today's discussion)
- Baum–Welch algorithm