

	x (input)	y (output)
Spam detection	document	spam or not
Sentiment analysis	product review	sentiment
Medical diagnosis	patient data	diagnosis
Credit scoring	financial history	loan decision

The cases (input-output) pairs are assumed to be *independent and identically distributed (i.i.d.)*.

Structured prediction

In many applications, the i.i.d. assumption is wrong

	x (input)	y (output)
POS tagging	word sequence	POS sequence
Parsing	word sequence	parse tree
OCR	image (array of pixels)	sequences of letters
Gene prediction	genome	genes

Structured/sequence learning is prevalent in NLP.

Sequence learning - a demonstration of the problem



- The most likely (local) prediction is not the correct prediction
- Individual predictions depend on each other
- Can we treat the whole sequence as a single label?

Recap: chain rule

We rewrite the relation between the joint and the conditional probability as

$$P(X, Y) = P(X | Y)P(Y)$$

We can also write the same quantity as,

$$P(X, Y) = P(Y | X)P(X)$$

In general, for any number of random variables, we can write

$$P(X_1, X_2, \dots, X_n) = P(X_1 | X_2, \dots, X_n)P(X_2, \dots, X_n)$$

Recap: (conditional) independence

If two variables X and Y are independent,

$$P(X | Y) = P(X) \quad \text{and} \quad P(X, Y) = P(X)P(Y)$$

If two variables X and Y are independent given another variable Z,

$$P(X, Y | Z) = P(X | Z)P(Y | Z)$$

An example: probability of a sentence

$$P(\text{The old man the boats}) = ?$$

- We cannot estimate this probability by counting all occurrences of the sentence, and dividing it to the total number of sentences in English
- We can (potentially) calculate its probability based on the probabilities of the words. Using chain rule

$$\begin{aligned} P(S) &= P(\text{boats} | \text{The old man the})P(\text{The old man the}) \\ &= P(\text{boats} | \text{The})P(\text{the} | \text{The old man})P(\text{The old}) \\ &= P(\text{boats} | \text{The old man the})P(\text{the} | \text{The old man})P(\text{man} | \text{The old})P(\text{old} | \text{The})P(\text{The}) \end{aligned}$$

- Did we solve the problem of probability estimation?

Markov chains

calculating probabilities

Given a sequence of events (or states), $q_1, q_2, \dots, q_t$,

- In a *first-order Markov chain*, the probability of an event q_t is

$$P(q_t | q_1, \dots, q_{t-1}) = P(q_t | q_{t-1})$$

- In higher order chains, the dependence of history is extended, e.g., *second-order Markov chain*:

$$P(q_t | q_1, \dots, q_{t-1}) = P(q_t | q_{t-2}, q_{t-1})$$

- The conditional independence properties simplify the probability distributions

Markov chains

definition

A Markov model is defined by,

- A set of states $Q = \{q_1, \dots, q_n\}$
- A special start state q_0
- A transition probability matrix

$$A = \begin{bmatrix} a_{01} & a_{02} & \dots & a_{0n} \\ a_{11} & a_{12} & \dots & a_{1n} \\ \vdots & \vdots & \ddots & \vdots \\ a_{n1} & a_{n2} & \dots & a_{nn} \end{bmatrix}$$

where a_{ij} is the probability of transition from state i to state j

Back to sentence probability example

- With a first-order Markov assumption,

$$\begin{aligned} P(S) &= P(\text{boats} | \text{The old man the})P(\text{The old man the}) \\ &= P(\text{boats} | \text{the})P(\text{the} | \text{man})P(\text{man} | \text{old})P(\text{old} | \text{The})P(\text{The}) \end{aligned}$$

- Now the probabilities are easier to estimate
- The above approach is an example of *n-gram language models* that we will get back to very soon

Hidden/latent variables

- In many machine learning problems we want to account for unobserved/unobservable *latent* or *hidden* variables
- Some examples
 - 'personality' in many psychological data
 - 'topic' of a text
 - 'socio-economic class' of a speaker
- Accounting for latent variables improve the accuracy of the models
- Since we cannot observe them, latent variables make learning algorithms difficult

Learning with hidden variables

An informal/quick introduction to the EM algorithm

- The EM algorithm (or its variants) is used in many machine learning models with latent/hidden variables

1. Randomly initialize the parameters

2. Iterate until convergence:

E-step: compute likelihood of the data, given the parameters

M-step: re-estimate the parameters using the predictions based on the E-step

Hidden Markov models (HMM)

- HMMs are like Markov chains: probability of a state depends only a limited history of previous states

$$P(q_i | q_1, \dots, q_{i-1}) = P(q_i | q_{i-1})$$

- Unlike Markov chains, state sequence is hidden, they are not the observations
- At every state q_i , an HMM *emits* an output, o_i , whose probability depends only on the associated hidden state
- Given a state sequence $\mathbf{q} = q_1, \dots, q_T$, and the corresponding observation sequence $\mathbf{o} = o_1, \dots, o_T$,

$$P(\mathbf{o}, \mathbf{q}) = p(q_1) \prod_{i=2}^T P(q_i | q_{i-1}) \prod_{i=1}^T P(o_i | q_i)$$

Example: HMMs for POS tagging



- The tags are hidden
- Probability of a tag depends on the previous tag
- Probability of a word at a given state depends only on the current tag

HMMs: formal definition

An HMM is defined by

- A set of states $Q = \{q_1, \dots, q_n\}$
- The set of possible observations $O = \{o_1, \dots, o_m\}$
- A transition probability matrix

$$\mathbf{A} = \begin{bmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ \vdots & \vdots & \ddots & \vdots \\ a_{n1} & a_{n2} & \dots & a_{nn} \end{bmatrix}$$

a_{ij} is the probability of transition from state q_i to state q_j

- Initial probability distribution $\pi = [P(q_1), \dots, P(q_n)]$

- Probability distributions of

$$\mathbf{B} = \begin{bmatrix} b_{11} & b_{12} & \dots & b_{1m} \\ \vdots & \vdots & \ddots & \vdots \\ b_{n1} & b_{n2} & \dots & b_{nm} \end{bmatrix}$$

b_{ij} is the probability of emitting output o_i at state q_j

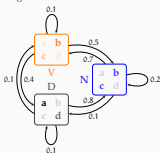
A simple example

- Three states: $\mathbb{N}, \mathbb{V}, \mathbb{D}$
- Four possible observations: a, b, c, d

$$\mathbf{A} = \begin{bmatrix} \mathbb{N} & \mathbb{V} & \mathbb{D} \\ 0.2 & 0.7 & 0.1 \\ 0.5 & 0.1 & 0.4 \\ 0.8 & 0.1 & 0.1 \end{bmatrix} \quad \mathbf{B} = \begin{bmatrix} \mathbb{N} & \mathbb{V} & \mathbb{D} \\ 0.1 & 0.1 & 0.5 \\ 0.4 & 0.5 & 0.1 \\ 0.4 & 0.3 & 0.1 \\ 0.1 & 0.1 & 0.3 \end{bmatrix} \quad \mathbf{a} = \begin{bmatrix} a \\ b \\ c \\ d \end{bmatrix}$$

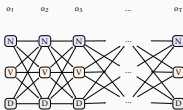
$$\pi = (0.3, 0.1, 0.6)$$

HMM transition diagram



Unfolding the states

HMM lattice (or trellis)



HMMs: three problems

Recognition/decoding

Calculating probability of state sequence, given an observation sequence

$$P(\mathbf{q} | \mathbf{o}; \mathbf{M})$$

Evaluation

Calculating likelihood of a given sequence

$$P(\mathbf{o} | \mathbf{M})$$

Learning

Given observation sequences, a set of states, and (sometimes) corresponding state sequences, estimate the parameters $(\pi, \mathbf{A}, \mathbf{B})$ of the HMM

Assigning probabilities to observation sequences

$$P(\mathbf{o} | \mathbf{M}) = \sum_{\mathbf{q}} P(\mathbf{o}, \mathbf{q} | \mathbf{M})$$

- We need to sum over an exponential number of hidden state sequences
- The solution is using a dynamic programming algorithm
 - for each node of the lattice, store *forward probabilities*

$$\alpha_{t,i} = \sum_{j=1}^N \alpha_{t-1,j} P(q_i | q_j) P(o_t | q_i)$$

Assigning probabilities to observation sequences

the forward algorithm

- Start with calculating all forward probabilities for $t = 1$

$$\alpha_{1,i} = \pi_i P(o_1 | q_i) \quad \text{for } 1 \leq i \leq |Q|$$

store the α values

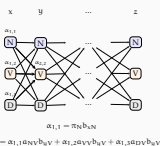
- For $t > 1$,

$$\alpha_{t,i} = \sum_{j=1}^{|Q|} \alpha_{t-1,j} P(q_i | q_j) P(o_t | q_i) \quad \text{for } 1 \leq i \leq |Q|, 2 \leq t \leq n$$

- Likelihood of the observation is the sum of the forward probabilities of the last step

$$P(\mathbf{o} | \mathbf{M}) = \sum_{i=1}^{|Q|} \alpha_{n,i}$$

Forward algorithm



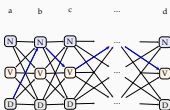
$$\alpha_{2,2} = \alpha_{1,1} a_{1V} b_{2V} + \alpha_{1,2} a_{2V} b_{2V} + \alpha_{1,3} a_{3V} b_{2V}$$

Determining best sequence of latent variables

Decoding

- We often want to know the hidden state sequence given an observation sequence, $P(\mathbf{q} | \mathbf{o}; \mathbf{M})$
 - For example, given a sequence of tokens, find the most likely POS tag sequence
- The problem (also the solution, the *Viterbi algorithm*) is very similar to the forward algorithm
- Two major differences
 - we store maximum likelihood leading to each node on the lattice
 - we also store backlinks, the previous state that leads to the maximum likelihood

HMM decoding problem



Learning the parameters of an HMM supervised case

- We want to estimate π, A, B
- If we have both the observation sequence $\mathbf{o}$ and the corresponding state sequence, MLE estimate is

$$\pi_i = \frac{C(q_0 \rightarrow q_i)}{\sum_k C(q_0 \rightarrow q_k)}$$

$$a_{ij} = \frac{C(q_i \rightarrow q_j)}{\sum_k C(q_i \rightarrow q_k)}$$

$$b_j = \frac{C(q_i \rightarrow o_j)}{\sum_k C(q_i \rightarrow o_k)}$$

HMM variations

- The HMMs we discussed so far are called *ergodic* HMMs: all a_{ij} are non-zero
- For some applications, it is common to use HMMs with additional restrictions
- A well known variant (Bakis HMM) allows only forward transitions

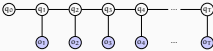


- The emission probabilities can also be continuous, e.g., $p(q|o)$ can be a normal distribution

Graphical models

- Graphical models define models involving multiple random variables
- It is generally more intuitive (compared to corresponding mathematical equations) to work with graphical models
- In a graphical model, by convention, the observed variables are shaded
- Graphs can also be undirected, which are also called *Markov random fields*

Conditional random fields



- A related model used in NLP is *conditional random field* (CRF)
- CRFs are *undirected models*
- CRFs model $P(q | o)$ directly

$$P(q_t | q_{t-1}, o) = \frac{1}{Z} e^{\sum_i w_i f_i(o, q_{t-1})}$$

Generative vs. discriminative models

- HMMs are *generative* models, they model the joint distribution
 - you can generate the output using HMMs
- CRFs are *discriminative* models they model the conditional probability directly
- It is easier to add arbitrary features on discriminative models
- In general:
 - Discriminative models (CRFs) do not model the joint distribution, and typically work better for classification tasks
 - Generative models can assign probabilities to (observed) input sequences (and generate them)

Additional reading, references, credits

 Jurafsky, Daniel and James H. Martin (2020). *Speech and Language Processing: An Introduction to Natural Language Processing, Computational Linguistics, and Speech Recognition with Language Models*. 3rd. Online manuscript released January 6, 2020. <https://web.stanford.edu/~jurafsky/slp3/>

Learning the parameters of an HMM

- Given a training set with observation sequence(s) $\mathbf{o}$ and state sequence $\mathbf{q}$, we want to find $\theta = (\pi, A, B)$

$$\arg \max_{\theta} P(\mathbf{o} | \mathbf{q}, \theta)$$

- Typically solved using EM
 - Initialize θ
 - Repeat until convergence
 - E-step: given θ , estimate the hidden state sequence
 - M-step: given the estimated hidden states, use 'expected counts' to update θ
- An efficient implementation of EM algorithm is called *Baum-Welch algorithm*, or *forward-backward algorithm*

Directed graphical models: a brief divergence

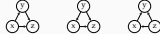
Bayesian networks

- We saw earlier that joint distributions of multiple random variables can be factorized different ways

$$P(x, y, z) = P(x)P(y|x)P(z|x, y) = P(y)P(x|y)P(z|x, y) = P(z)P(x|z)P(y|x, z)$$

- Graphical models display this relations in graphs,
 - variables are denoted by nodes,
 - the dependence between the variables are indicated by edges

- Bayesian networks are directed acyclic graphs

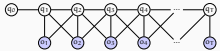


- A variable (node) depends only on its parents

HMM as a graphical model



CRFs can incorporate more dependencies



$$P(q_t | q_{t-1}, o) = \frac{1}{Z} e^{\sum_i w_i f_i(o, q_{t-1})}$$

Summary

- In many problems, e.g., POS tagging, I.L.d. assumption is wrong
- We need models that are aware of the effects of the sequence (or structure in general) in the data
- HMMs are generative sequence models:
 - Markov assumption between the hidden states (POS tags)
 - Observations (words) are conditioned on the state (tag)
- There are other sequence learning methods
 - Briefly mentioned: CRF
 - Coming soon: recurrent neural networks
- Reading: Jurafsky and Martin (2026, Chapter 17)
- Next
 - Recurrent and convolutional networks
 - Reading: Jurafsky and Martin (2026, Chapter 13)