

CS 378 Lecture 10

HMMs, Viterbi

Announcements

- A2 due today
- A3 out today
- Midterm: in 2.5 weeks
- Final proj: if independent project, start thinking
- Bias response due today

Recap POS tagging

NNP	VBZ	NN	NNS
Fed	raises	interest	rates

How to do with a classifier?

$$P(y_i = \hat{y} \mid \bar{x}) = \frac{e^{\bar{w}_y^T f(\bar{x}, i)}}{\sum_{y \in \mathcal{Y}} e^{\bar{w}_y^T f(\bar{x}, i)}}$$

↑
prediction

$$f(\bar{x}, i) = \begin{cases} \text{Prev word} = \dots \\ \text{Curr word} = \dots \\ \text{Next word} = \dots \end{cases}$$

↙ index these +
treat as your
features

Problem: all y_i are predicted independently
How do we model the sequence
all at once?

Today Hidden Markov Models

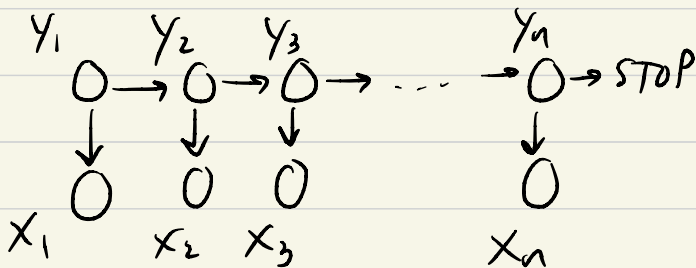
- Definition
- Training
- Inference (Viterbi)

HMM: generative model $P(\bar{y}, \bar{x})$

There are discriminative seq. models,
Conditional Random Fields $P(\bar{y} | \bar{x})$

HMMs Tags $y_i \in \overset{\text{tag set}}{\mathcal{T}}$, words $x_i \in \mathcal{V}$

$$P(\bar{y}, \bar{x}) = P(y_1) P(x_1 | y_1) P(y_2 | y_1) P(x_2 | y_2) \dots P(\text{STOP} | y_n)$$



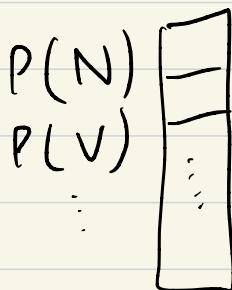
Generative story: draw y_1 to start tag
NNP
Sequence
Fed

draw $x_1 | y_1$ as the first word

draw $y_2 | y_1$ as the next tag

Parameters

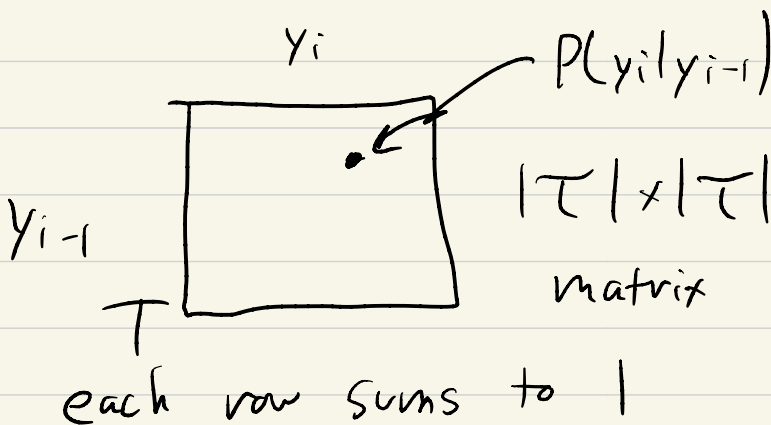
$P(y_i)$
initial distribution



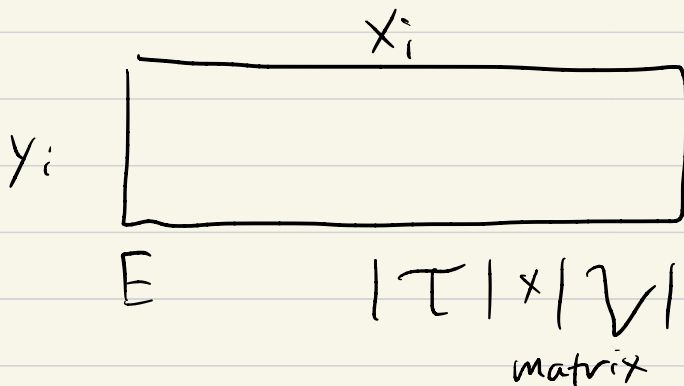
call this S

$|T|$ -len
vector
that sums
to 1

$P(y_i | y_{i-1})$
transitions



$P(x_i | y_i)$
emissions



why Markov model: $P(y_i | y_{i-1})$ y_i doesn't depend on y_{i-2}

Goals with HMMs

HMM $P(\bar{y}, \bar{x})$

What we want: $P(\bar{y} | \bar{x}) = \frac{P(\bar{y}, \bar{x})}{P(\bar{x})}$

Ex $\mathcal{T} = \{N, V, STOP\}$ $\mathcal{V} = \{they, can, fish\}$

S	$P(y_i)$	T	y_i		
			N	V	STOP
N	1.0	y_{i-1}	1/5	3/5	1/5
V	0.0		1/5	1/5	3/5
STOP	0.0				

E	y_i			
		they	can	fish
N	N	1	0	0
	V	0	1/2	1/2

Poll

① What is the probability of

$\begin{pmatrix} N & V & V \\ \text{they can fish} \end{pmatrix}$

$$P(\bar{y}, \bar{x}) = P(N) P(\text{they} | N) \dots$$

$$\begin{array}{ccccccc} 1.0 & \cdot & 1.0 & \dots & & & \\ 1.0 & \cdot & 3/5 & \cdot & 1/5 & \cdot & 3/5 \end{array} \quad \leftarrow P(\text{STOP} | V)$$
$$\begin{array}{ccc} 1.0 & 1/2 & 1/2 \end{array}$$

② $\begin{array}{ccc} N & V & V \\ \text{they fish can} \end{array}$

$\begin{array}{cc} N & V \\ \text{they can} \end{array}$

$P(y_i = V) = 0.0$
 ~~$\begin{array}{ccc} V & N & V \\ \text{can they fish} \end{array}$~~

③ other seqs for
they can fish?
No

Training Given labeled sequences

$$(\bar{x}^{(i)}, \bar{y}^{(i)})_{i=1}^D$$

Data: $N \ V$ $N \ V$
they can they fish

What parameters maximize the likelihood of the data?

Can find this exactly by
counting + normalizing

Biased coin w/probability p of heads
See H H H T

$p=3/4$ is the MLE

$$P(x_i | v) = \begin{matrix} \text{can: } 1 \\ \text{fish: } 1 \\ \text{they: } 0 \end{matrix} \Rightarrow \begin{matrix} \text{emissions for } v \\ \text{they } 0 \\ \text{can } 1/2 \\ \text{fish } 1/2 \end{matrix}$$

This is what "training" involves

$$P(y_i | y_{i-1} = v) = \begin{matrix} \text{counts} & \text{probs} \\ N: 0 & N: 0 \\ v: 0 & v: 0 \\ \text{stop: } 2 & \text{stop: } 1 \end{matrix} \Rightarrow$$

Smoothing: pretend we saw everything with some nonzero count

Add 0.1 to every count

$$\begin{matrix} N & 0 \\ v & 0 \\ \text{stop} & 2 \end{matrix} \Rightarrow \begin{matrix} N & 0.1 \\ v & 0.1 \\ \text{stop} & 2.1 \end{matrix} \Rightarrow \text{probs}$$

Inference in HMMs

We defined HMMs as $P(\bar{y}, \bar{x})$

What we really want is a tagger

$P(\bar{y} | \bar{x})$ "you give me a sent,
I tell you tags"

$\operatorname{argmax}_{\bar{y}} P(\bar{y} | \bar{x})$ Two problems:

① Space is exponentially
large $(|\mathcal{Y}|^n)$

② we don't have $P(\bar{y} | \bar{x})$ yet

Solving #2:

$$\begin{aligned}\operatorname{argmax}_{\bar{y}} P(\bar{y} | \bar{x}) &= \operatorname{argmax}_{\bar{y}} \frac{P(\bar{y} | \bar{x}) \cdot P(\bar{x})}{P(\bar{x})} \\ &= \operatorname{argmax}_{\bar{y}} \frac{P(\bar{y}, \bar{x})}{P(\bar{x})}\end{aligned}$$

$$P(\bar{x}) = \sum_{\bar{y}} P(\bar{y}, \bar{x})$$

We don't need this

$$\operatorname{argmax}_{\bar{y}} \frac{P(\bar{y}, \bar{x})}{P(\bar{x})} = \operatorname{argmax}_{\bar{y}} P(\bar{y}, \bar{x})$$

Conclusion: $\operatorname{argmax}_{\bar{y}} P(\bar{y}, \bar{x})$

to build our tagger #2 ✓

Solving #1: the Viterbi algorithm

First: move to log space

$$\operatorname{argmax}_{\bar{y}} P(\bar{y}, \bar{x}) = \operatorname{argmax}_{\bar{y}} \log P(\bar{y}, \bar{x})$$

$$\begin{aligned}
 & \arg \max_{\bar{y}} \log P(\bar{y}, \bar{x}) \\
 &= \arg \max_{\tilde{y}_1, \dots, \tilde{y}_n} \log P(\tilde{y}_1) + \log P(x_1 | \tilde{y}_1) \\
 & \quad + \log P(\tilde{y}_2 | \tilde{y}_1) + \dots
 \end{aligned}$$

Basic concept: dynamic programming

Best sequence of $\tilde{y}_1, \dots, \tilde{y}_n$
can be found from

best sequence from $\tilde{y}_1, \dots, \tilde{y}_{n-1}$
+ HMM probs

optimal seq for each
possible $\tilde{y}_{n-1}$