

CS 378 Lecture 15:

Language Modeling, RNNs

Today

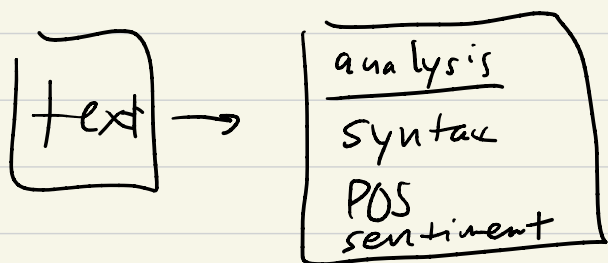
- Language modeling
- N-gram LMs
- Neural LMs
- (Start) RNNs

Announcements

- + midterm
- A4 + A5
- FP custom proposals

Where we are

First half of the course: analysis
of text



input: text

output: structured analysis

same models!
(or are they?)

Next ~3 weeks:

Task: text generation ★

dialogue
summarization

input: text

output: text

Machine translation

Language modeling: "autocomplete"
next word prediction

★ Technique: recurrent neural networks
(RNNs)
(and Transformers)

Final ~3 weeks:] BERT,
Question answering Transformers

End of the course: miscellaneous applications + topics

QA: question (text)



Language Modeling

Imagine we want a distribution over grammatical sent. in a language

$P(\bar{w})$

PCFGs give us this

Why?

the cat ran $\nwarrow$ higher
the the cat ran $\nearrow$ prob

Grammatical error correction:

I give you a sent $\bar{w}$
has a grammatical error

Then: $P(\bar{w})$ is low

Correct the error by finding
 $\bar{w}'$ close to $\bar{w}$ with
 $P(\bar{w}') \gg P(\bar{w})$

Another application: machine translation

Translate $\bar{w}_1$ in lang 1 $\Rightarrow \bar{w}_2$ in
lang 2. Maybe we want to
make $\bar{w}_2$ more natural

Find $\bar{w}_2'$ s.t. $P(\bar{w}_2') \gg P(\bar{w}_2)$

N-gram language models

$$P(\bar{w}) = P(w_1) P(w_2 | w_1) P(w_3 | w_1, w_2) \\ \dots P(w_n | w_1, \dots, w_{n-1})$$

True by chain rule of prob.

N-gram model: only look at the past $n-1$ words

$$P(\bar{w}) = P(w_1) P(w_2 | w_1) \dots \\ P(w_i | w_{i-n+1}, \dots, w_{i-1})$$

2-gram model: $\bar{w} =$ the cat ran

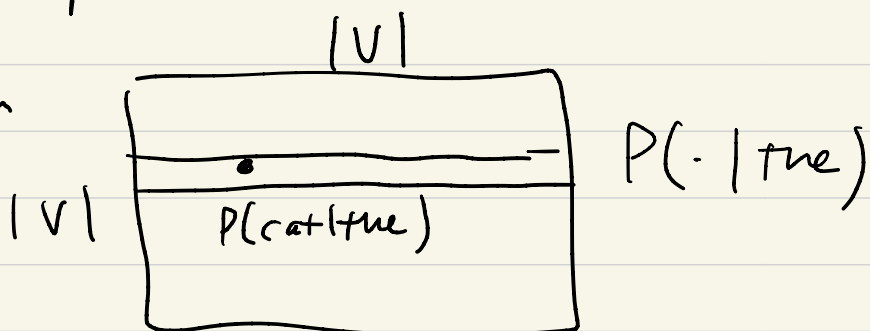
$$P_2(\bar{w}) = P(\text{the} | \langle S \rangle) P(\text{cat} | \text{the}) \cdot P(\text{ran} | \text{cat}) \\ \text{start of sent} \quad \cdot P(\text{STOP} | \text{ran})$$

$$\text{3-gram: } P(\text{the} | \langle S \rangle \langle S \rangle) \\ \cdot P(\text{cat} | \langle S \rangle \text{the}) \\ P(\text{ran} | \text{the cat}) \quad P(\text{STOP} | \dots)$$

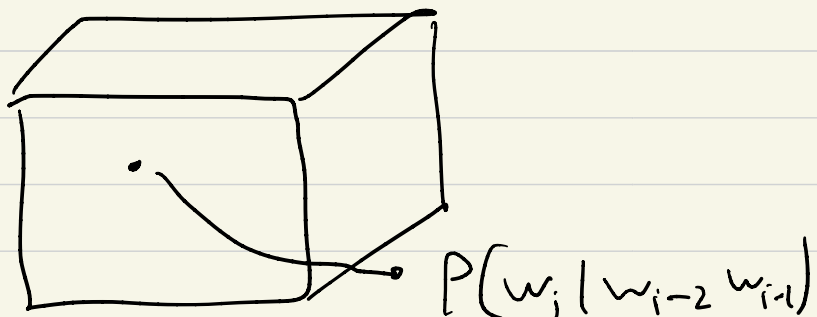
Parameters: transitions in HMM

Big lookup table

2-gram



3-gram



Get these params: count + normalize

5-gram: $P(\text{Maui} | \text{hate to go to}) = 0$

Nobody has said this before, but
it's still reasonable!

Smoothing

Very important in n-gram LMs

$$P(w_i | w_{i-4}, w_{i-3}, w_{i-2}, w_{i-1})$$

$$\begin{aligned} \Rightarrow & \lambda_1 P^{\text{raw}}(w_i | w_{i-4}, w_{i-3}, w_{i-2}, w_{i-1}) \\ & + \lambda_2 P^{\text{raw}}(w_i | w_{i-3}, w_{i-2}, w_{i-1}) \\ & + \lambda_3 P^{\text{raw}}(w_i | w_{i-2}, w_{i-1}) \\ & + \lambda_4 P^{\text{raw}}(w_i | w_{i-1}) \\ & + \lambda_5 P(w_i) \end{aligned}$$

4-gram!

$$\hookrightarrow > 0! \quad P(\text{Mavi})$$

$$\lambda_1 = 0 + \lambda_2 = 0 + \lambda_3 = 10^{-4} + \lambda_4 = 10^{-2} + \dots$$

$\Rightarrow$ reasonable prob.

Set λ carefully so this is a valid prob. dist.

Neural language models

Use a neural net to model

$$P(w_i | w_1, \dots, w_{i-1})$$

whole context, not just $n-1$ words

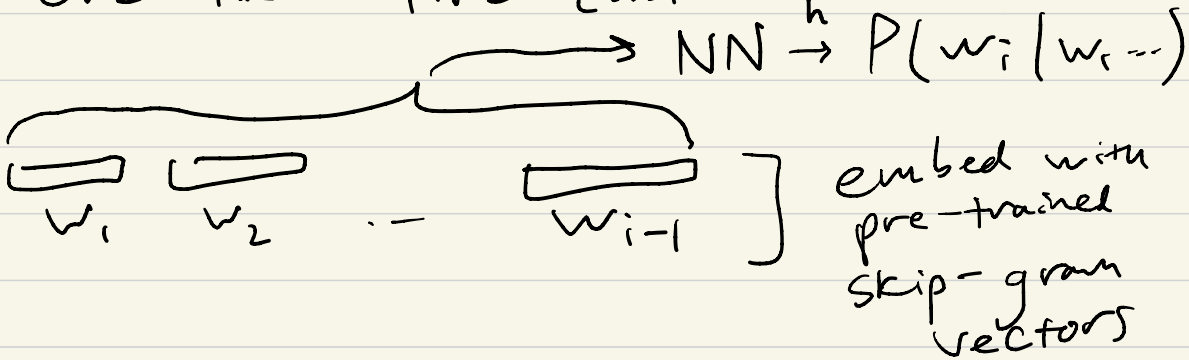
For now: simplify, consider

$$P(w_i | w_{i-1})$$

↑
vectors

Skip-gram:
 $P(\text{context} | \text{word})$

One idea! take context $\vec{h}$



$P(w_i | w_1, \dots, w_{i-1})$ dist over $|V|$ words

d -dim

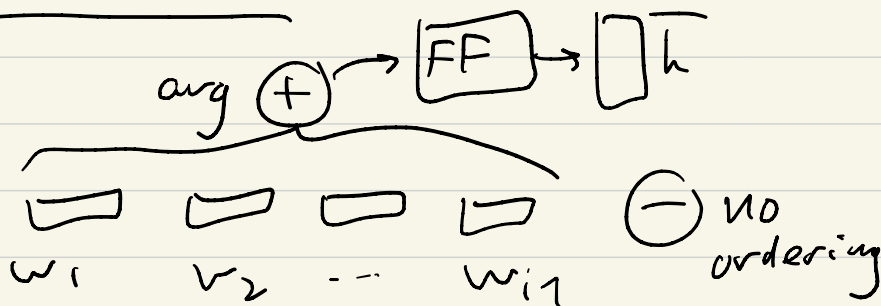
$\bar{h} = \text{neural net}(w_1, \dots, w_{i-1})$

$P(w_i | w_1, \dots, w_{i-1}) = \text{softmax}(W\bar{h})$

W : param matrix $|V| \times d$

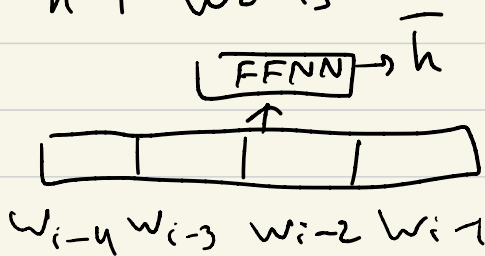
Ways to do this

DAN:



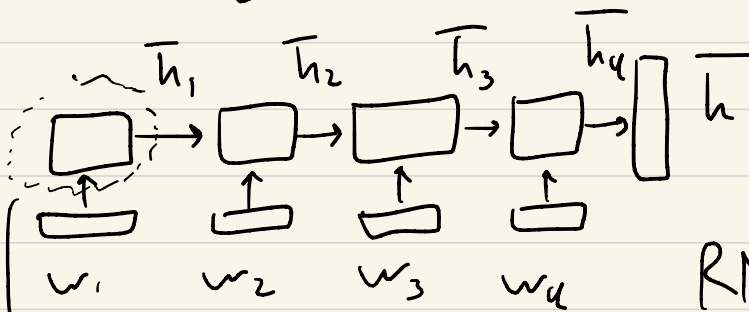
FFNN: $n-1$ words

5 gram



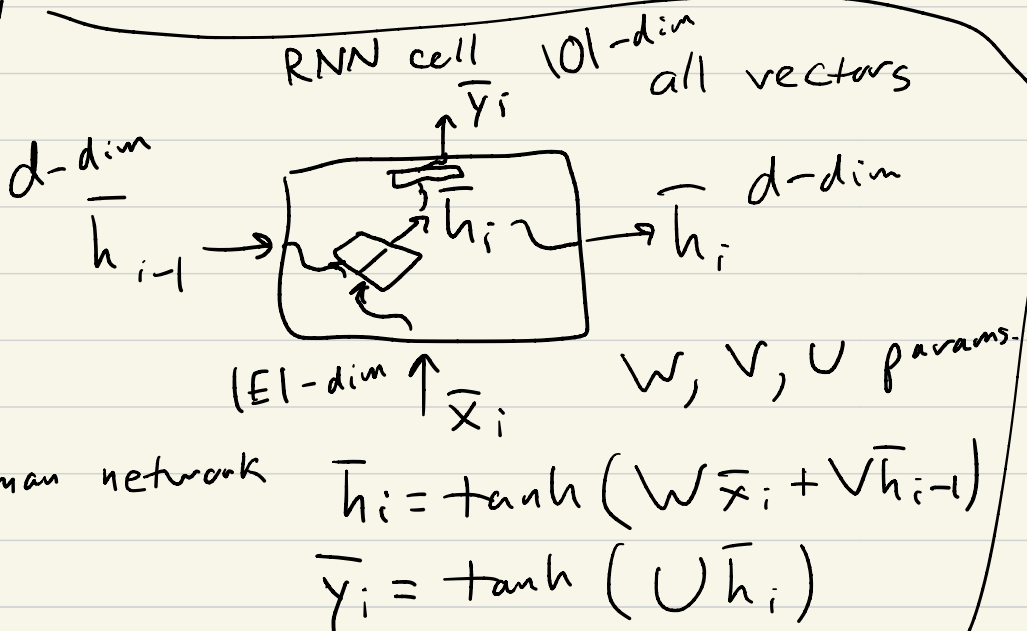
doesn't scale to large n

RNN: encodes a sequence of arbitrary length into a single vector



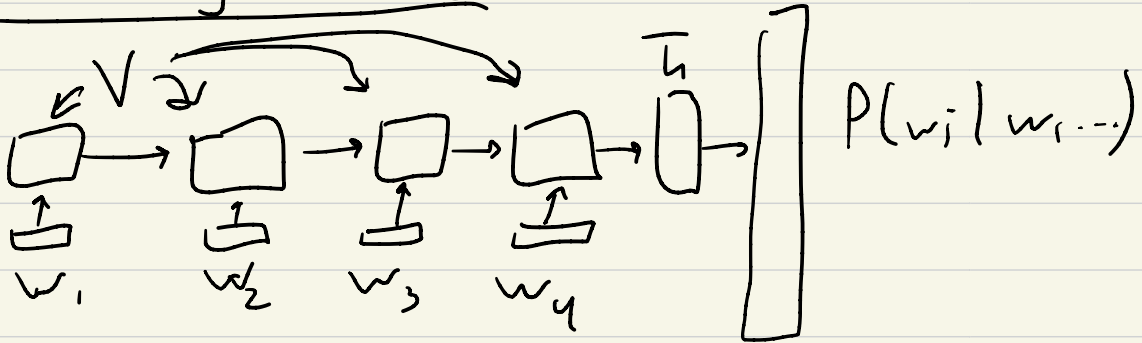
Zoom in

RNN outputs:
 $\bar{y}_i, \bar{h}_i$ at each step



Elman network

Training RNNs



This is a differentiable computation!

Some layers share params