

CS 378 Lecture 16

Announcements

Assignments: A4 only coding (but simplified), A5 no coding,
FP in progress (w/partner possible, no pres.)

Extra slip days

Zoom protocol: type Qs in chat, polls on Canvas, breakouts

Recap

Language modeling: $P(\bar{w})$ prob of sentence
or document

n-gram model: $P(\bar{w}) = \prod_{i=1}^n P(w_i | w_{i-n+1}, \dots, w_{i-1})$

3-gram prob of "I want to go":

$P(I | \langle s \rangle \langle s \rangle) P(\text{want} | \langle s \rangle I) P(\text{to} | I \text{ want})$
 $P(\text{go} | \text{want to}) P(\text{STOP} | \text{to go})$

Estimate counts from large corpora, smooth

Goals RECURRENT neural networks (RNN)

1. key RNN abstraction
2. Key properties of LSTM
(long short-term memory)

This lecture: LM, RNN definition

Next lecture: training + implementation

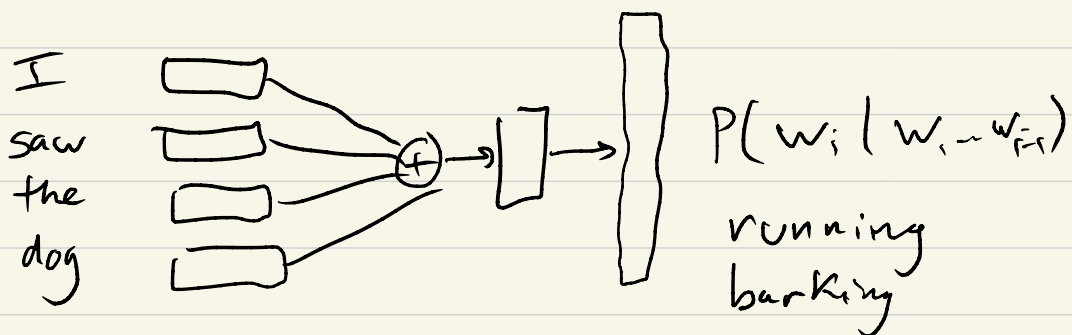
This sets the stage for machine translation, summarization, dialogue

Neural Language Models

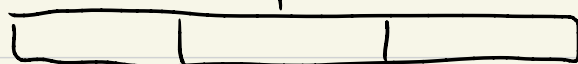
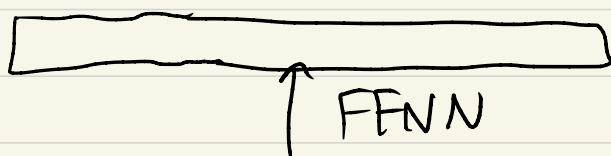
Discrim model for $P(w_i | \underbrace{w_1 \dots w_{i-1}}_{\text{all past words}})$

What we want: neural net to look at $w_1 \dots w_{i-1}$, place distribution over w_i

DAN: + fast — ignores order



Feedforward n-gram model | + uses order
4-gram model: — not much context



saw the dog

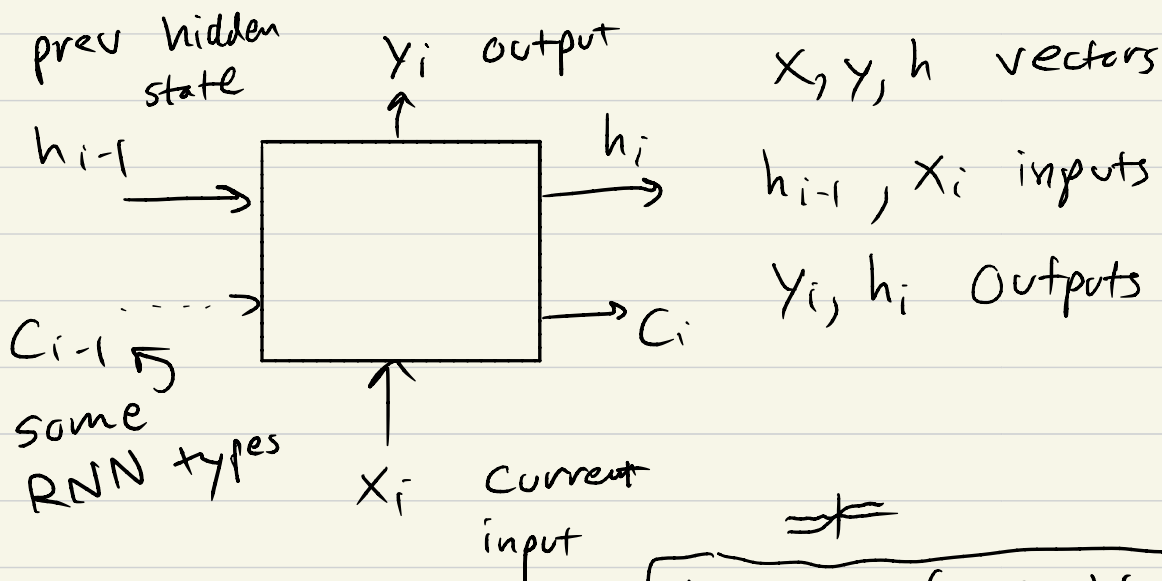
$P(w_i | w_{i-3}, w_{i-2}, w_{i-1})$

past 3 words

input: $| \text{word emb} | \cdot (n-1)$ $3 \Rightarrow 10$: more
need large hidden states (1000+) params in NN

params: $| \text{word emb} | \cdot (n-1) \cdot 1000 + \dots$

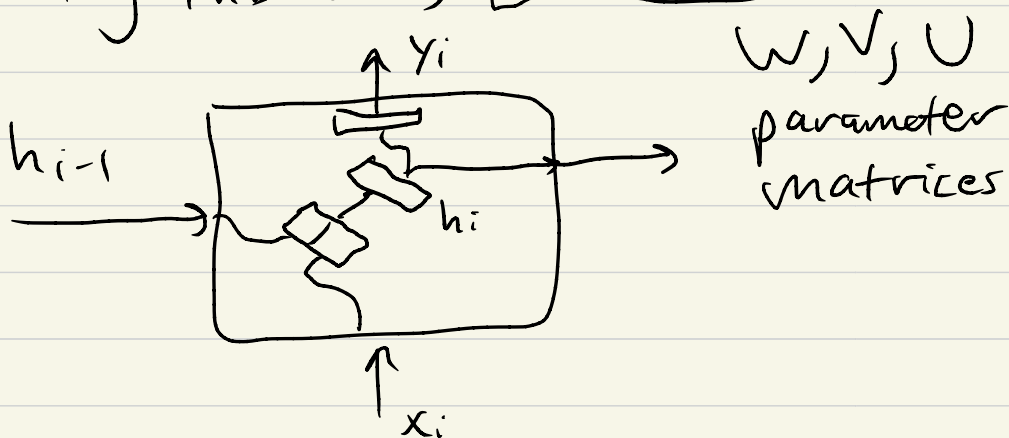
RNN: neural model for encoding seqs of arbitrary length

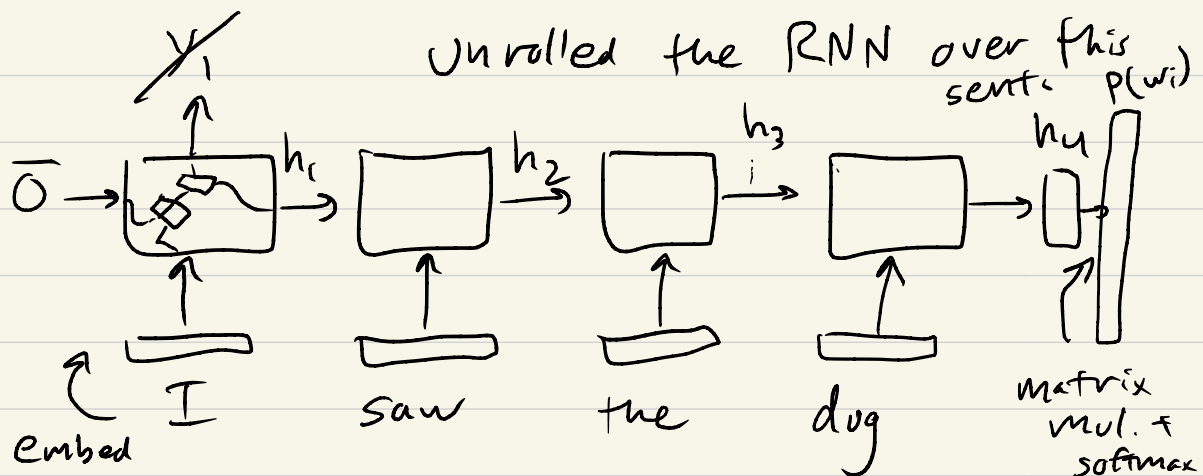


Elman network
(UT ling PhD 1977)

$$h_i = \tanh(Wx_i + Vh_{i-1})$$

$$y_i = \tanh(Uh_i)$$





`forward()` in PyTorch has a for loop, but otherwise this is just a special feedforward network

At each step, h_i captures the model's "picture" of the sentence so far

$$P(w_i | w_1, \dots, w_{i-1}) = \text{softmax}(M h_{i-1})$$

$$h_{i-1} = \text{RNN}(w_1, \dots, w_{i-1})$$

W, U, V, M , are our only params!

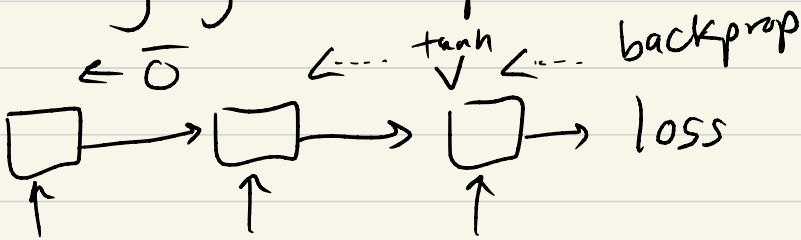
$|V| - |h|$

Training RNNs

"Backpropagation through time" = backpropagation
Move next time

Long short-term memory network

Vanishing gradient problem



Elman networks (not LSTMs) : error term
"dies out"

model can't learn to feed info
long distance

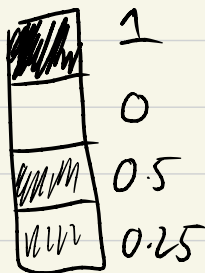
Key idea: gates

Elman: $h_i = \tanh(Wx_i + Vh_{i-1})$

Gated: $h_i = h_{i-1} \odot f + \text{func}(x_i) \odot i$

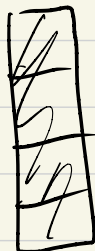
$\odot$ = element-wise mul

f : vector $\in [0, 1]^d$



f

$\odot$



$h_{i-1} \in \mathbb{R}^d$

=

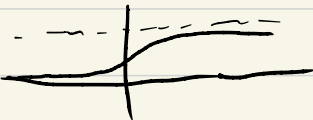


$f = \mathbf{I}$ preserves h_i : $h_i = \sum_{j=0}^{i-1} \text{func}(x_j) \odot i$

otherwise zeroes out some parts

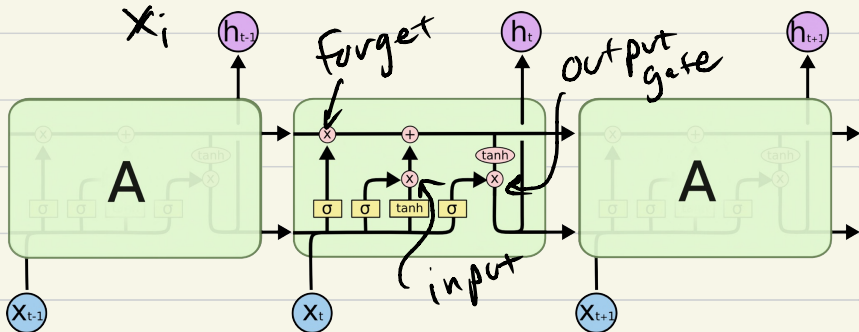
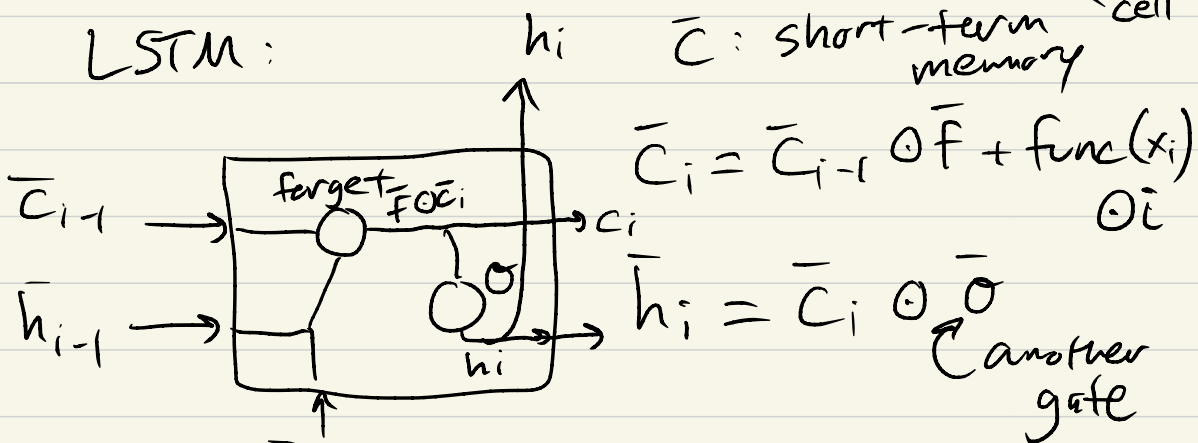
$\bar{}$ elementwise

$$\bar{f} = \text{sigmoid}(W^1 \bar{x}_i + W^2 \bar{h}_{i-1}) + \text{bias}$$

$$\sigma = \frac{e^x}{1+e^x}$$


$$\bar{i} = \text{sigmoid}(W^3 \bar{x}_i + W^4 \bar{h}_{i-1})$$

LSTM: $\bar{c}$: short-term memory 'cell'



Source: Chris Olah blog post (link on course website)

Key properties:

RNN: seq of words (vectors) as
input
encodes into state h_i

LSTM: 2 hidden states (h, c)

h vs. c distinction is not
better at remembering info for long
sequences by using gates rather
than tanh + matrix mul. critical

Resources: Chris Olah blog
PyTorch example lstm-lecture.py
Floydhub tutorial