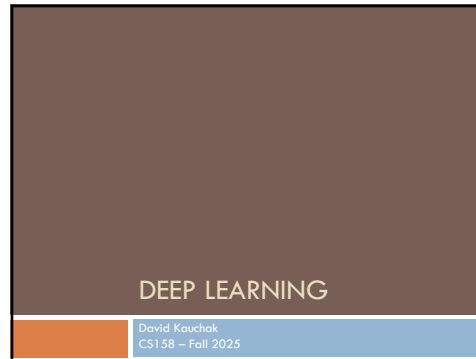
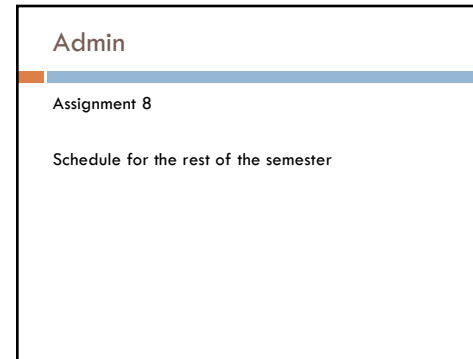
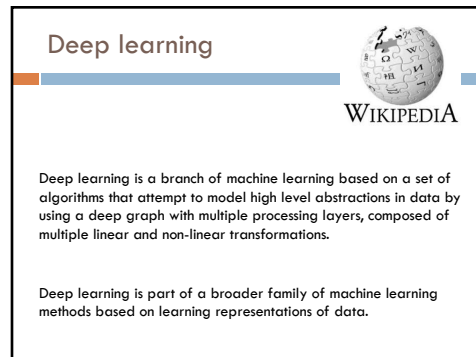




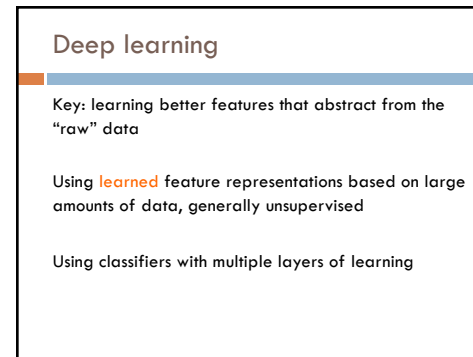
1



2



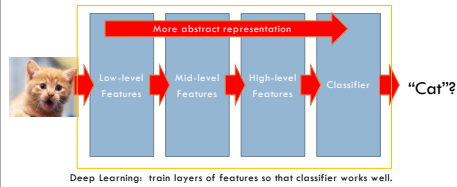
3



4

Deep learning

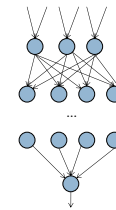
- Train *multiple layers* of features/abstractions from data.
- Try to discover *representation* that makes decisions easy.



Slide adapted from: Adam Coates

5

Deep learning for neural networks



Traditional NN models: 1-2 hidden layers

Deep learning NN models: 3+ hidden layers

6

Importance of features

Feature quality is critical to the performance of ML methods

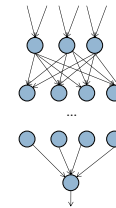
Normal process = hand-crafted features

Deep learning: find algorithms to automatically discover features from the data

7

Challenges

What makes "deep learning" hard for NNs?



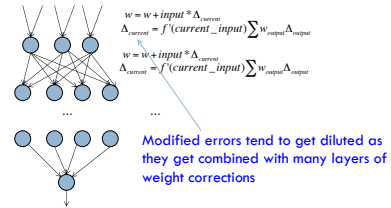
$$w = w + \text{input} * \Delta_{\text{output}}$$

$$\Delta_{\text{output}} = f'(\text{current_input}) \sum w_{\text{output}} \Delta_{\text{output}}$$

8

Challenges

What makes “deep learning” hard for NNs?



10

Deep learning

Growing field

Driven by:

- Increase in data availability
- Increase in computational power
- Parallelizability of many of the algorithms

Involves more than just neural networks (though, they're a very popular model)

11

word2vec

How many people have heard of it?

What is it?

12

Word representations

Wine data uses word occurrences as a feature

What does this miss?

13

Word representations

Wine data uses word occurrences as a feature

What does this miss?

"The wine had a dark red color"	Zinfandel
"The wine was a deep crimson color"	label?
"The wine was a deep yellow color"	label?

Would like to recognize that words have similar meaning even though they aren't lexically the same

14

Word representations

Key idea: project words into a multi-dimensional "meaning" space

word $\Rightarrow [x_1, x_2, \dots, x_d]$

15

Word representations

Key idea: project words into a multi-dimensional "meaning" space

word $\Rightarrow [x_1, x_2, \dots, x_d]$

red	$\Rightarrow [r_1, r_2, \dots, r_d]$	$[y_1, y_2, \dots, y_d]$
crimson	$\Rightarrow [c_1, c_2, \dots, c_d]$	$[r_1, r_2, \dots, r_d]$
yellow	$\Rightarrow [y_1, y_2, \dots, y_d]$	$[c_1, c_2, \dots, c_d]$

16

Word representations

Key idea: project words into a multi-dimensional "meaning" space

word $\Rightarrow [x_1, x_2, \dots, x_d]$

The idea of word representations is not new:

- Co-occurrence matrices
- Latent Semantic Analysis (LSA)

New idea: learn word representation using a task-driven approach

17

A prediction problem

I like to eat bananas with cream cheese

Given a **context** of words

Predict what **words** are likely to occur in that context

18

A prediction problem

Given text, can generate lots of examples:

I like to eat bananas with cream cheese

input

___ like to eat

I ___ to eat bananas

I like ___ eat bananas with

I like to ___ bananas with cream

...

prediction

I

like

to

eat

...

19

A prediction problem

Use data like this to learn a distribution:

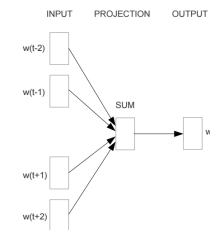
$$p(\text{word} | \text{context})$$

$$p(w_i | w_{i-2} w_{i-1} w_{i+1} w_{i+2})$$

words before
words after

20

Train a neural network on this problem



<https://arxiv.org/pdf/1301.3781v3.pdf>

21

Encoding words

How can we input a "word" into a network?

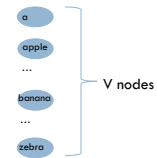


22

"One-hot" encoding

For a vocabulary of V words, have V input nodes

All inputs are 0 except the one corresponding to the word

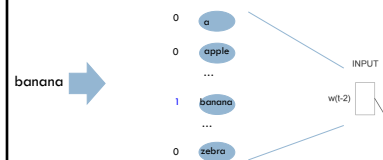


23

"One-hot" encoding

For a vocabulary of V words, have V input nodes

All inputs are 0 except the one corresponding to the word

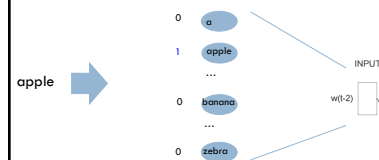


24

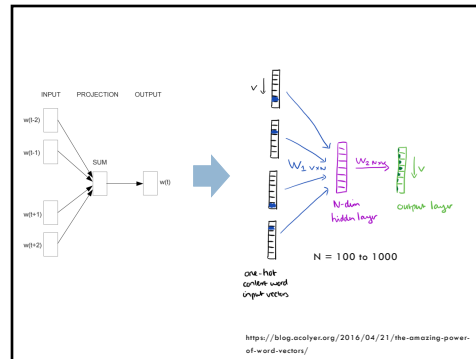
"One-hot" encoding

For a vocabulary of V words, have V input nodes

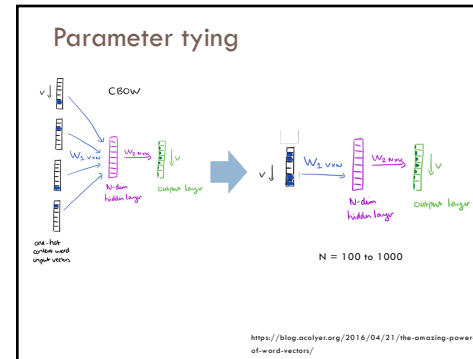
All inputs are 0 except the one corresponding to the word



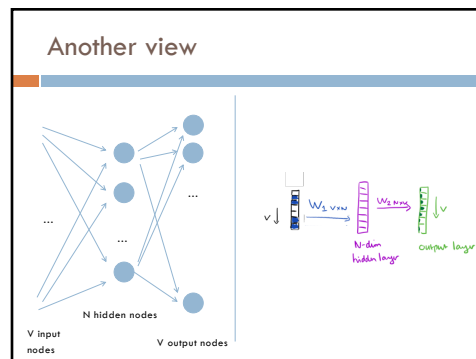
25



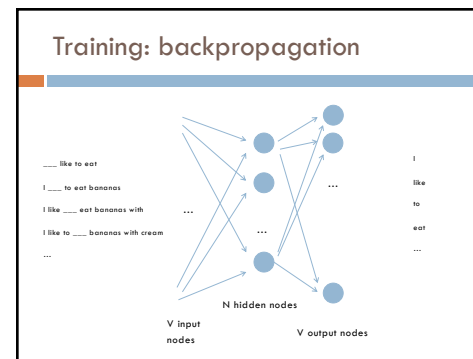
26



27

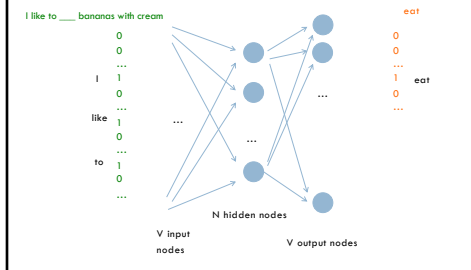


28



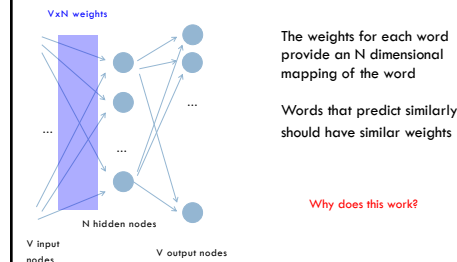
29

Training: backpropagation



30

Word representation



31

Results

$$\text{vector}(\text{word1}) - \text{vector}(\text{word2}) = \text{vector}(\text{word3}) - X$$

word1 is to word2 as word3 is to X

Type of relationship	Word Pair 1		Word Pair 2	
Common capital city	Athens	Greece	Oslo	Norway
All capital cities	Astana	Kazakhstan	Harare	Zimbabwe
Currency	Angola	kwana	Iran	rial
City-in-state	Chicago	Illinois	Stockton	California
Man-Woman	brother	sister	grandson	granddaughter

32

Results

$$\text{vector}(\text{word1}) - \text{vector}(\text{word2}) = \text{vector}(\text{word3}) - X$$

word1 is to word2 as word3 is to X

Type of relationship	Word Pair 1		Word Pair 2	
Adjective to adverb	apparent	apparently	rapid	rapidly
Opposite	possibly	impossibly	ethical	unethical
Comparative	great	greater	tough	tougher
Superlative	easy	easiest	lucky	luckiest
Present Participle	think	thinking	read	reading
Nationality adjective	Switzerland	Swiss	Cambodia	Cambodian
Past tense	walking	walked	swimming	swam
Plural nouns	mouse	mice	dollar	dollars
Plural verbs	work	works	speak	speaks

33

Results

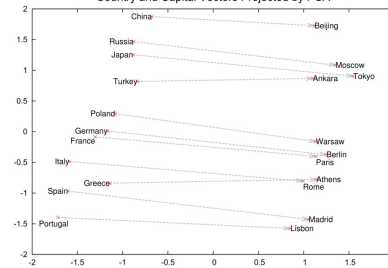
$$\text{vector}(\text{word1}) - \text{vector}(\text{word2}) = \text{vector}(\text{word3}) - X$$

word1 is to word2 as word3 is to X

Newspapers			
New York Times	Baltimore Sun		
San Jose Mercury News	Cincinnati Enquirer		
NHL Teams			
Boston Bruins	Montreal Canadiens		
Phoenix Coyotes	Nashville Predators		
NBA Teams			
Detroit Pistons	Toronto Raptors		
Golden State Warriors	Memphis Grizzlies		
Airlines			
Austrian Airlines	Spain		
Brussels Airlines	Greece		
Company executives			
Steve Ballmer	Microsoft	Larry Page	Google
Samuel J. Palmisano	IBM	Werner Vogels	Amazon

34

Country and Capital Vectors Projected by PCA



2-Dimensional projection of the N-dimensional space

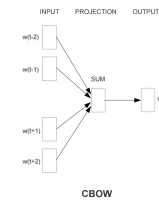
35

Visualized

<https://projector.tensorflow.org/>

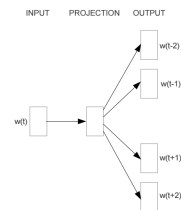
36

Continuous Bag Of Words



37

Other models: skip-gram



38

word2vec

A model for learning word representations from large amounts of data

Has become a popular pre-processing step for learning a more robust feature representation

Models like word2vec have also been incorporated into other learning approaches (e.g., translation tasks)

39

word2vec resources

<https://blog.acolyer.org/2016/04/21/the-amazing-power-of-word-vectors/>

<https://code.google.com/archive/p/word2vec/>

<https://deeplearning4j.org/word2vec>

<https://arxiv.org/pdf/1301.3781v3.pdf>

40

Language modeling

What does natural language look like?

$p(\text{ sentence })$

- $p(\text{"I like to eat pizza"})$
- $p(\text{"pizza like I eat"})$

Often is posed as: $p(\text{ word } | \text{ previous words })$ – or some other notion of context

- $p(\text{"pizza"} | \text{"I like to eat"})$
- $p(\text{"garbage"} | \text{"I like to eat"})$
- $p(\text{"run"} | \text{"I like to eat"})$

41

Ideas?

p("I like to eat pizza")

p("pizza like I eat")

p("pizza" | "I like to eat")

p("garbage" | "I like to eat")

p("run" | "I like to eat")

42

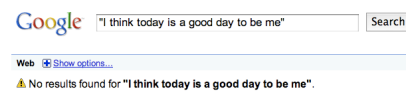
Look at a corpus



43

Language modeling

I think today is a good day to be me



Language modeling is about dealing with data sparsity!

44

Language modeling

Many approaches:

- n-gram language modeling
 - Start at the beginning of the sentence
 - Generate one word at a time based on the previous words
- syntax-based language modeling
 - Construct the syntactic tree from the top down
 - e.g. context free grammar
 - eventually at the leaves, generate the words
- Neural language models
 - Predict the likelihood of the word based on the context
 - Often allows for generalization beyond the lexical strings

45

n-gram language modeling

I think today is a good day to be me

Google "I think" Search

Web Show options... Results 1 - 10 of about 564,000,000 for "I think". (0.28 seconds)

Google "today is a good day" Search

Web Show options... Results 1 - 10 of about 10,100,000 for "today is a good day".

Google "to be me" Search

Web Show options... Results 1 - 10 of about 70,200,000 for "to be me".

46

Our friend the chain rule

Step 1: decompose the probability

$$P(\text{I think today is a good day to be me}) =$$

$$P(\text{I} \mid \langle \text{start} \rangle) \times$$

$$P(\text{think} \mid \text{I}) \times$$

$$P(\text{today} \mid \text{I think}) \times$$

$$P(\text{is} \mid \text{I think today}) \times$$

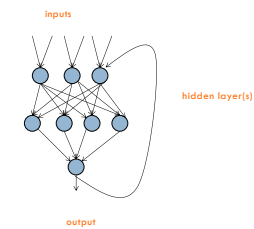
$$P(\text{a} \mid \text{I think today is}) \times$$

$$P(\text{good} \mid \text{I think today is a}) \times$$

...

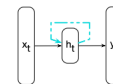
47

Recurrent neural networks



48

Recurrent neural nets



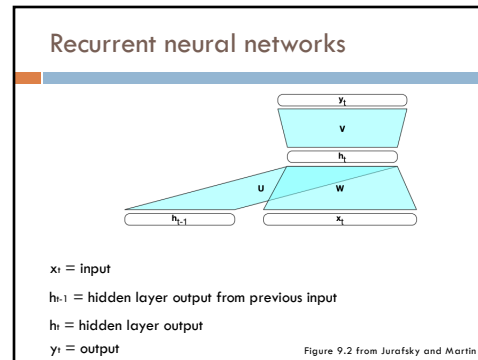
x_t = input

h_t = hidden layer output

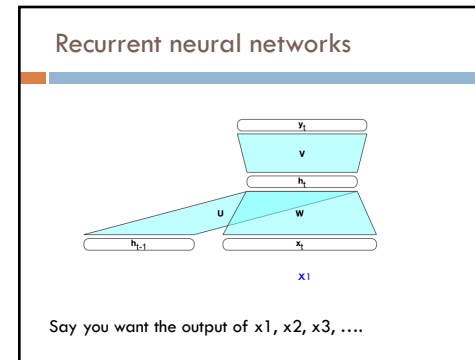
y_t = output

Figure 9.1 from Jurafsky and Martin

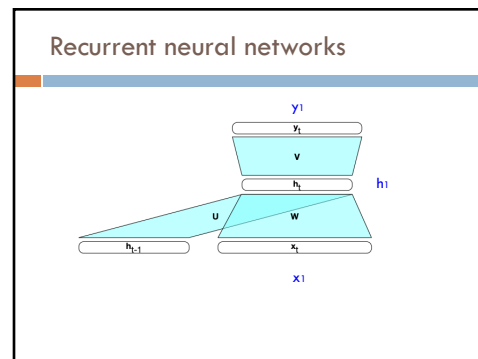
49



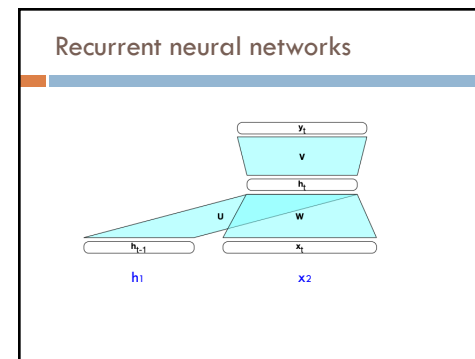
50



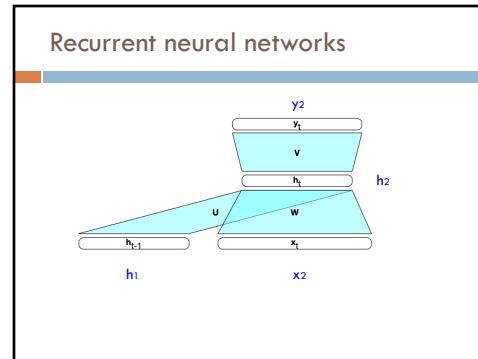
51



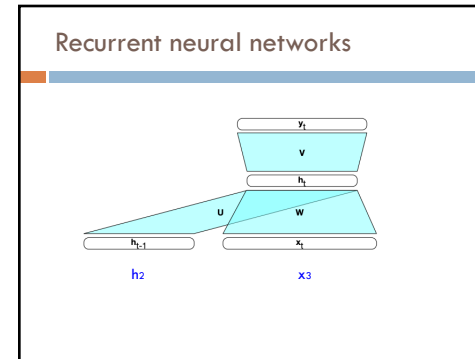
52



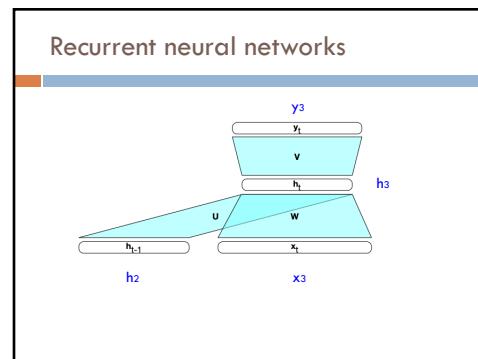
53



54



55



56

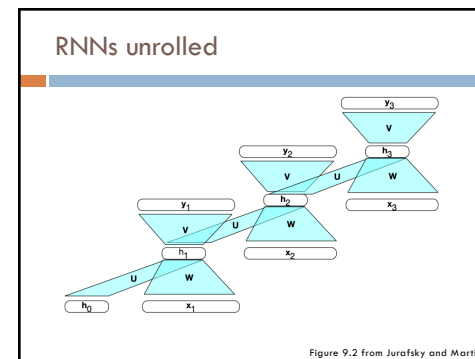
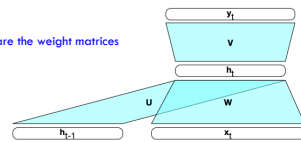


Figure 9.2 from Jurafsky and Martin

57

Still just a single neural network

U , W and V are the weight matrices



x_t = input

h_{t-1} = hidden layer output from previous input

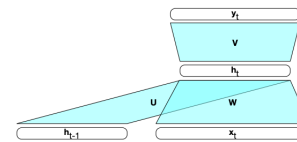
h_t = hidden layer output

y_t = output

Figure 9.2 from Jurafsky and Martin

58

RNN language models



How can we use RNNs as language models $p(w_1, w_2, \dots, w_n)$?

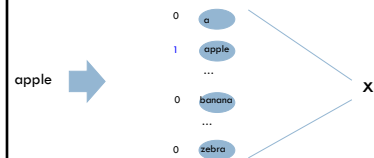
How do we input a word into a NN?

59

"One-hot" encoding

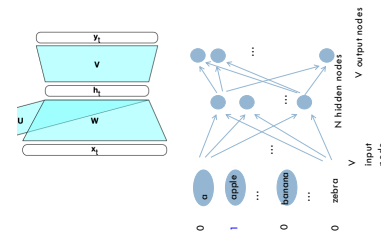
For a vocabulary of V words, have V input nodes

All inputs are 0 except the one corresponding to the word

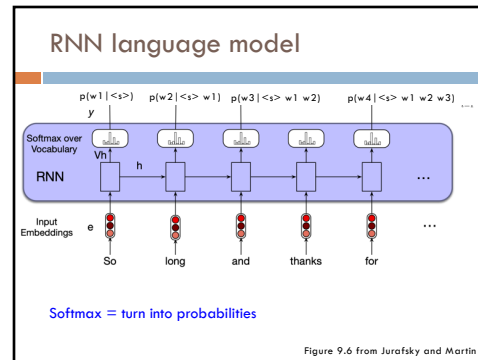


60

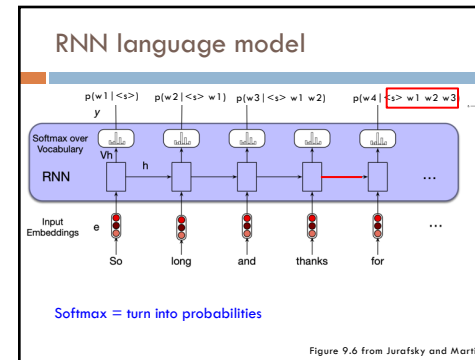
RNN language model



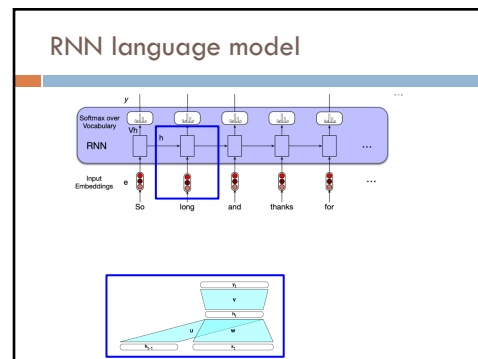
61



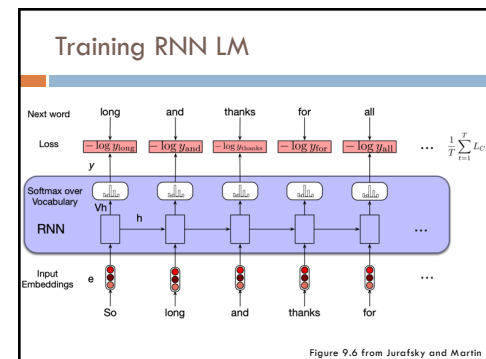
62



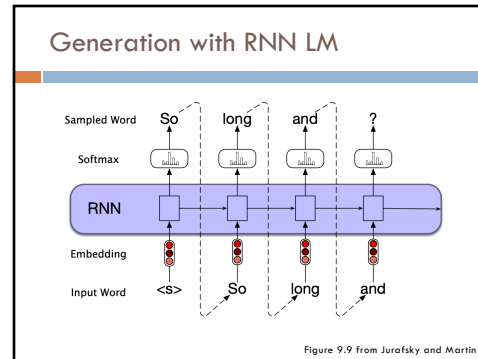
63



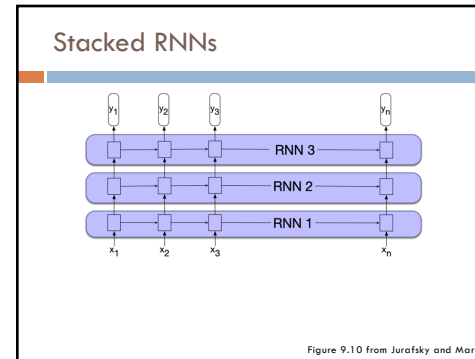
64



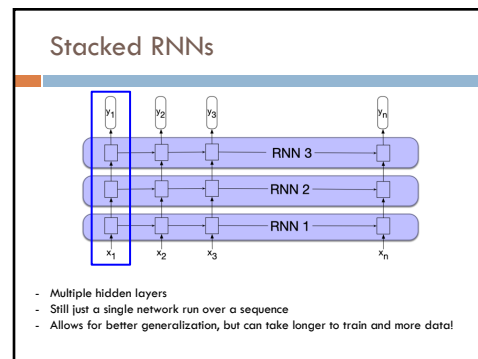
65



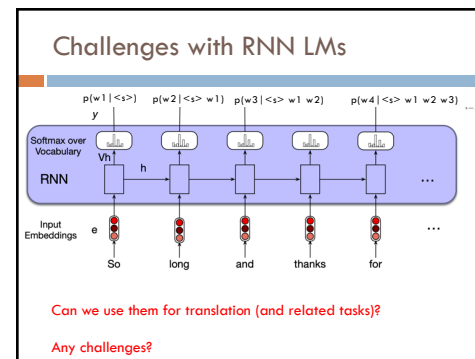
66



67

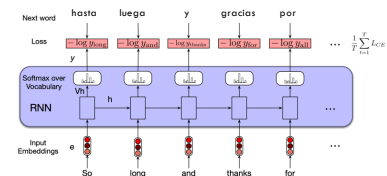


68



75

Challenges with RNN LMs

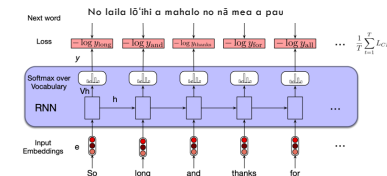


Can we use them for translation (and related tasks)?

Any challenges?

76

Challenges with RNN LMs

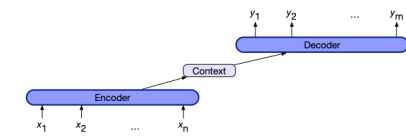


Translation isn't word-to-word

Worse for other tasks like summarization

77

Encoder-decoder models



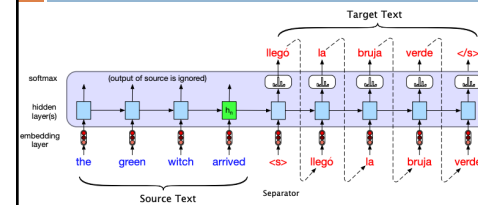
Idea:

- Process the input sentence (e.g., sentence to be translated) with a network
- Represent the sentence as some function of the hidden states (encoding)
- Use this context to generate the output

Figure 9.16 from Jurafsky and Martin

78

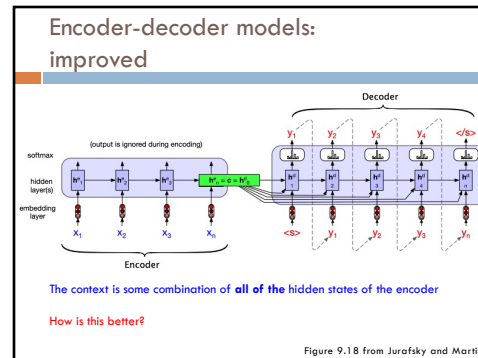
Encoder-decoder models: simple version



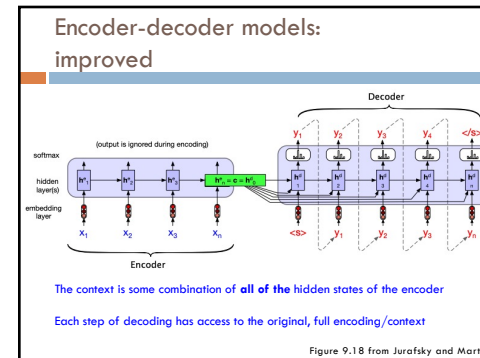
The context is the final hidden state of the encoder and is provided as input to the first step of the decoder

Figure 9.17 from Jurafsky and Martin

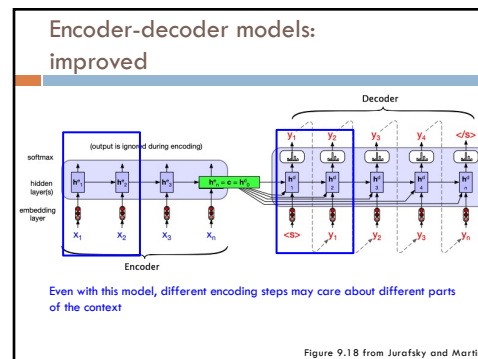
79



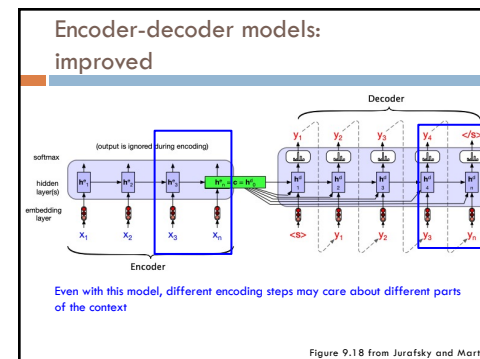
80



81

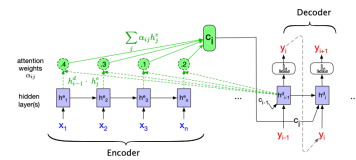


82



83

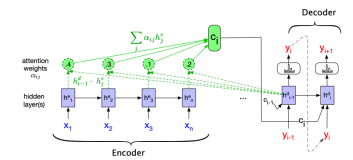
Attention



Context is dependent on where we are in decoding step and the relationship between encoder and decoder hidden states

84

Attention

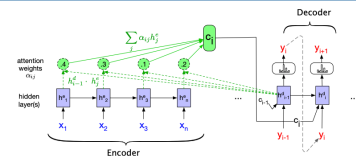


Simple version attention is static, but can learn attention mechanism (i.e., relationship between encoder and decoder hidden states)

Figure 9.23 from Jurafsky and Martin

85

Attention

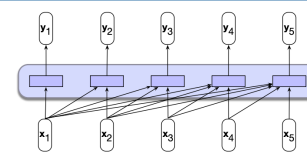


Key RNN challenge: computation is sequential
 - This prevents parallelization
 - Harder to model contextual dependencies

Figure 9.23 from Jurafsky and Martin

86

Another model

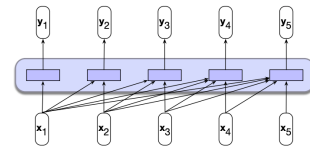


How is this setup different from the RNN?

Figure 10.1 from Jurafsky and Martin

87

Another model

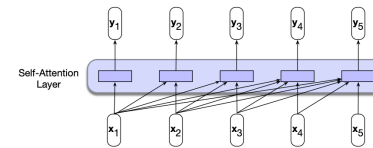


- Do not rely on the hidden states for context information
- Parallel: computation can all happen at once

Figure 10.1 from Jurafsky and Martin

88

Self-attention



- Self-attention:
 - Input is some context (for LMs, the previous words)
 - Learn what parts of the context are important based

Figure 10.1 from Jurafsky and Martin

89

Self-attention

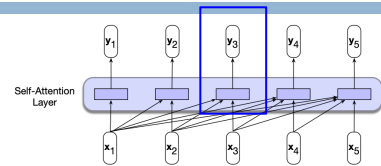


Figure 10.1 from Jurafsky and Martin

90

Transformer block

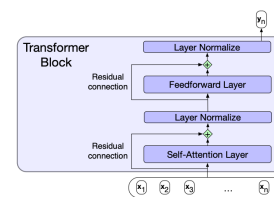
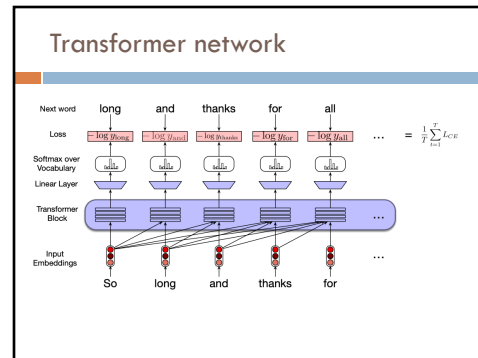
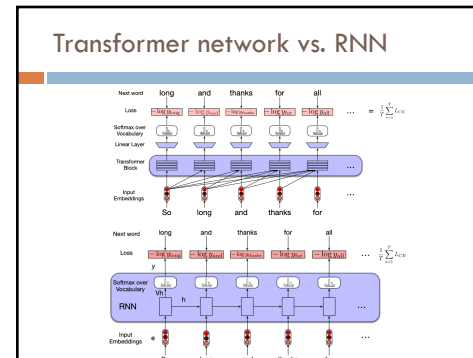


Figure 10.4 from Jurafsky and Martin

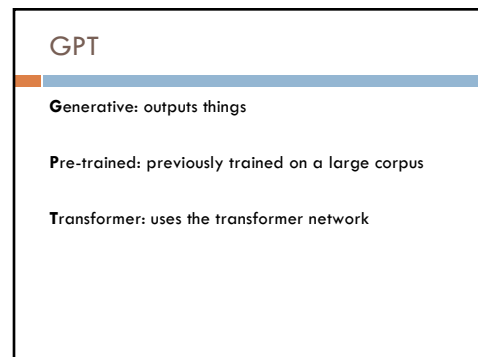
91



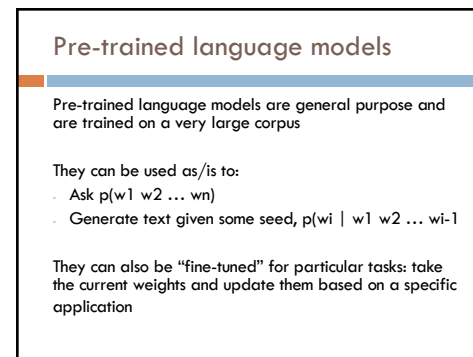
92



93



94



95

ChatGPT

ChatGPT is based on particular [GPT foundation models](#), namely [GPT-4](#), [GPT-4o](#) and [GPT-4o mini](#), that were [fine-tuned](#) to target conversational usage.^[17] The fine-tuning process leveraged [supervised learning](#) and [reinforcement learning from human feedback](#) (RLHF).^{[18][19]} Both approaches employed human trainers to improve model performance. In the case of supervised learning, the trainers played both sides: the user and the AI assistant. In the reinforcement learning stage, human trainers first ranked responses that the model had created in a previous conversation.^[20] These rankings were used to create "reward models" that were used to fine-tune the model further by using several iterations of [proximal policy optimization](#).^{[18][21]}

96

Another example

https://adamharley.com/nn_vis/

97