

## Recurrent Neural Network

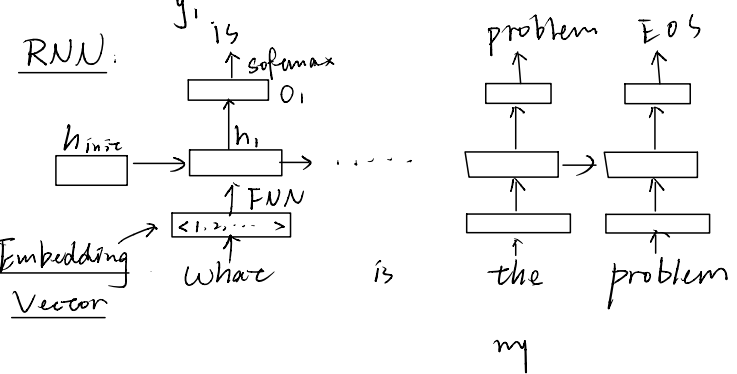
### Previous models' prediction

$\{x_1, y_1, \dots, x_n, y_n\} \rightarrow P(y_i | x_i) \Rightarrow 1 \text{ input} \rightarrow 1 \text{ output}$

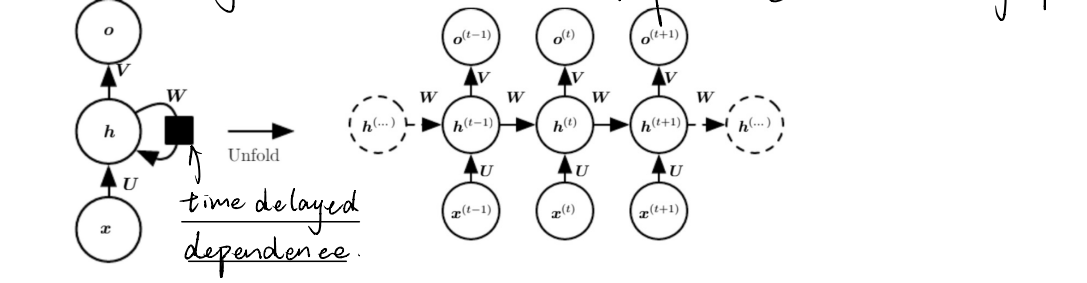
### RNN → Sequential Data

$\{(x_1^{(1)} \dots x_1^{(n)}), (y_1^{(1)} \dots y_1^{(n)})\}_{n=1}^N \rightarrow P(y_i^{(n)} | (x_1^{(n)} \dots x_{i-1}^{(n)}))$

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

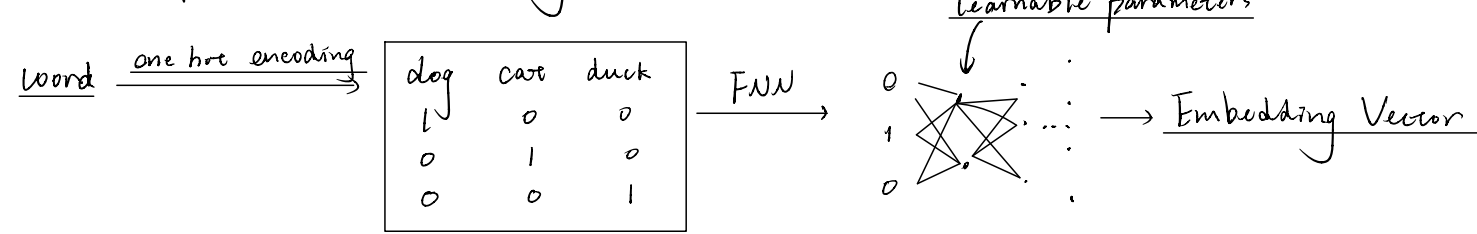


### "circuit diagram"



$$\begin{cases} h_t = \tanh(W h_{t-1} + U x_t + b) \\ o_t = \text{softmax}(V h_t + c) \end{cases}$$

$x_t$ : input token (Embedding Vectors)



$$L = E[-\sum_i \log P(y_i | x_1, \dots, x_n)]$$

Objective: minimize  $L(W, U, V, b, c, \text{Embedding}) \rightarrow \text{BPTT}$

BPTT: Back Propagation Through Time

### Learning Long Sequences

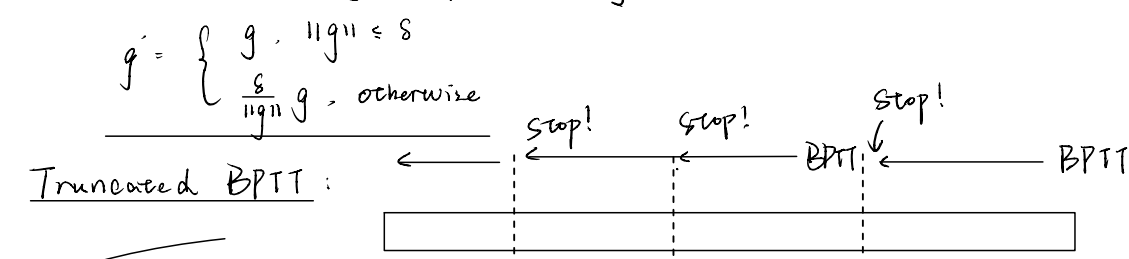
Problem:

① Vanishing Gradient  $\rightarrow$  Change Activation Function / Residual Connection / Batch Normalization

② Exploding Gradient

More serious: weight sharing (W) across different time steps

Gradient Clipping: if exceeding threshold  $\rightarrow$  clips the gradients.



Truncates a long sequence into multiple short subsequence

$\Rightarrow$  Later context cannot have attention to previous context

$\Rightarrow$  Better Approach: replace simple recurrent units by more powerful memory cell

### LSTM: Long Short-Term Memory

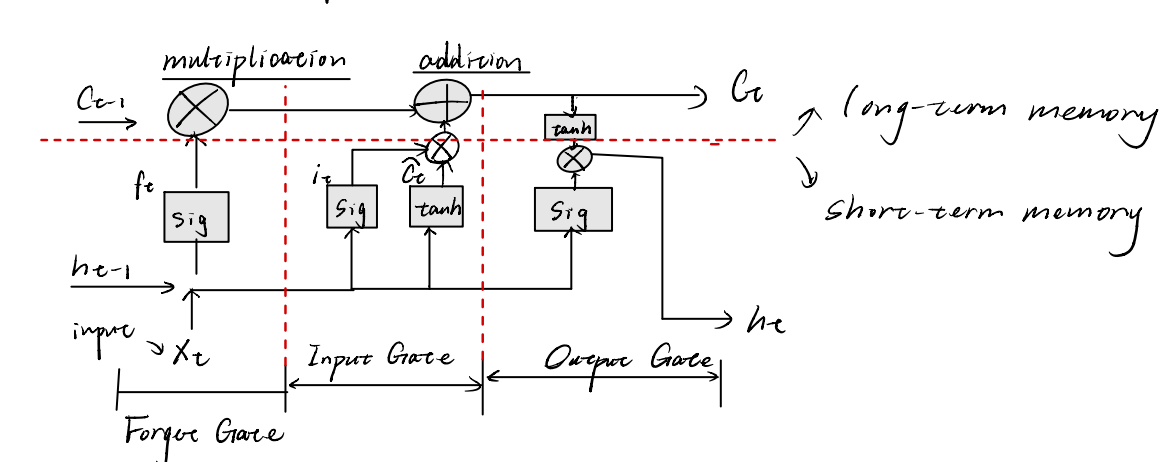
if truncated BPTT, we can never know what is "it"

I bought Harry Potter yesterday. Today I want to read it.

LSTM Cell:

$$C_t = f_t \odot C_{t-1} + i_t \odot \tanh(a_t) \quad h_t = \tanh(b_t)$$

### LSTM



Forget Gate:  $f_t = \sigma(W_f \cdot [h_{t-1}, x_t] + b_f)$

Input Gate:  $i_t = \sigma(W_i \cdot [h_{t-1}, x_t] + b_i)$

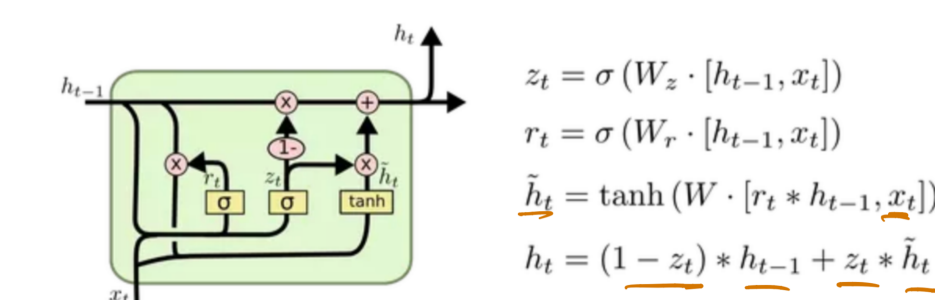
$\Rightarrow C_t = f_t \odot C_{t-1} + i_t \odot x_t$   
How much to forget How much to input (add)

Output Gate:  $o_t = \sigma(W_o \cdot [h_{t-1}, x_t] + b_o)$

$h_t = o_t \odot \tanh(C_t)$

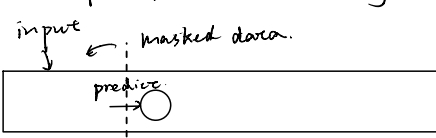
### GRU: Gated Recurrent Unit

similar to LSTM, fewer parameters, more efficient.



Performance similar to LSTM, except Better on small datasets.

### Self-Supervised Learning



"Training Data is automatically labelled"

in RNN: "Masking"

$y_i$ : the next word after  $x_{i-1}$

$x_i \rightarrow o_i \leftrightarrow y_i \rightarrow \text{loss, BP} \dots$

### RNN Architecture

#### ① Deep RNNs

More layers often usually implies better performance.

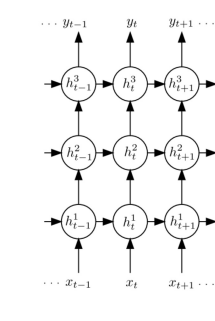
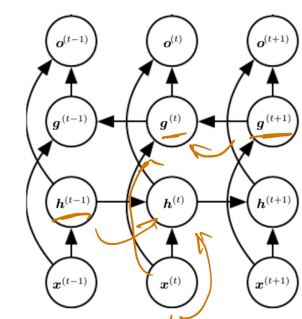


Fig. 3: Deep Recurrent Neural Network

#### ② Bidirectional RNNs

"I \_\_\_\_\_ a horse", "I \_\_\_\_\_ coffee."

We need information from BOTH past and future



#### ③ The Encoder-Decoder Structure

"Seq2Seq" Architecture.

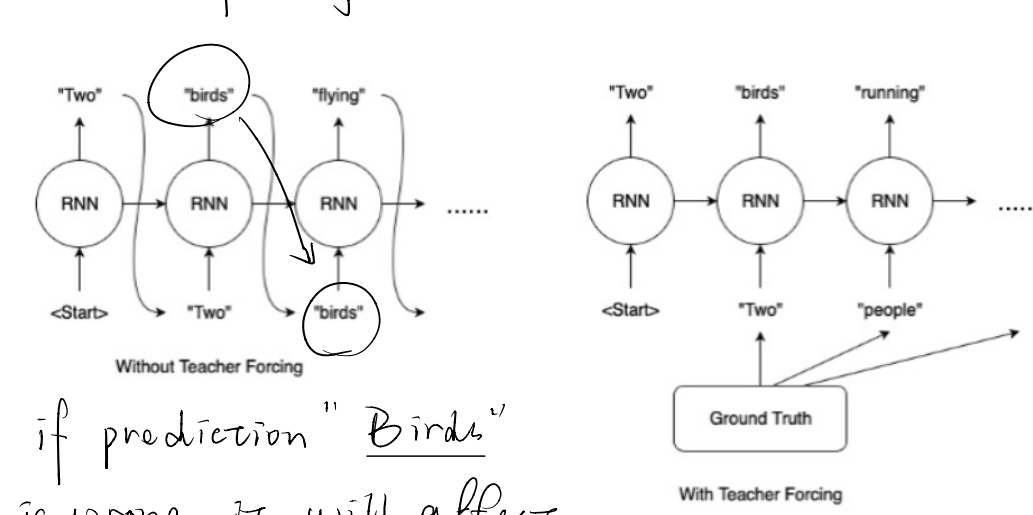
$(x^{(1)} \dots x^{(n)}) \rightarrow C \rightarrow (y^{(1)} \dots y^{(n)})$

$n_x \neq n_y$ , vanilla RNN can only output sequence which has same length

C: context variable  $\Rightarrow$  "Summary", with input

Used in dialogue system / machine translation, ("Transformer")

### Teacher forcing in Decoder



if prediction "Birds" is wrong, it will affect next prediction

Pros: makes training faster

if we don't use it, errors will accumulate

Cons: Discrepancy Between training and inference

"Since when inference, this is no ground truth!"

"Exposure Bias"

Recent work has shown exposure bias is not a serious issue

### Decoding:

Method 1: Greedy  $\Rightarrow$  Choose word with max prob

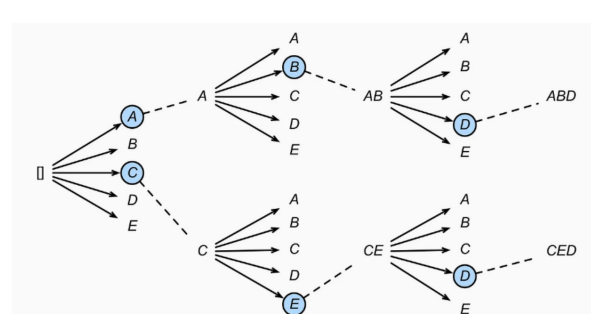
Method 2: Sampling

1. Top-k: Sample among top-k words.

2. Top-p (nucleus sampling)

among words with highest prob  $\rightarrow$  till  $\sum(\text{prob}) \geq p$

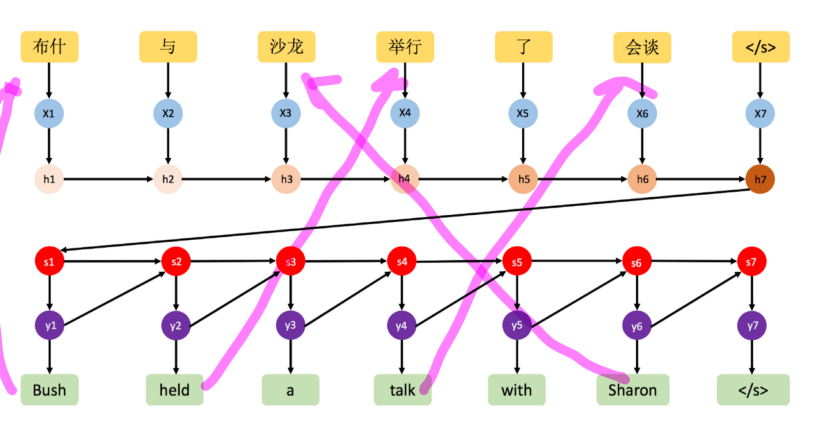
Method 3: Beam Search



### 5. Attention

A mechanism to allow model focus on different parts

of the input sequence when generating each output token

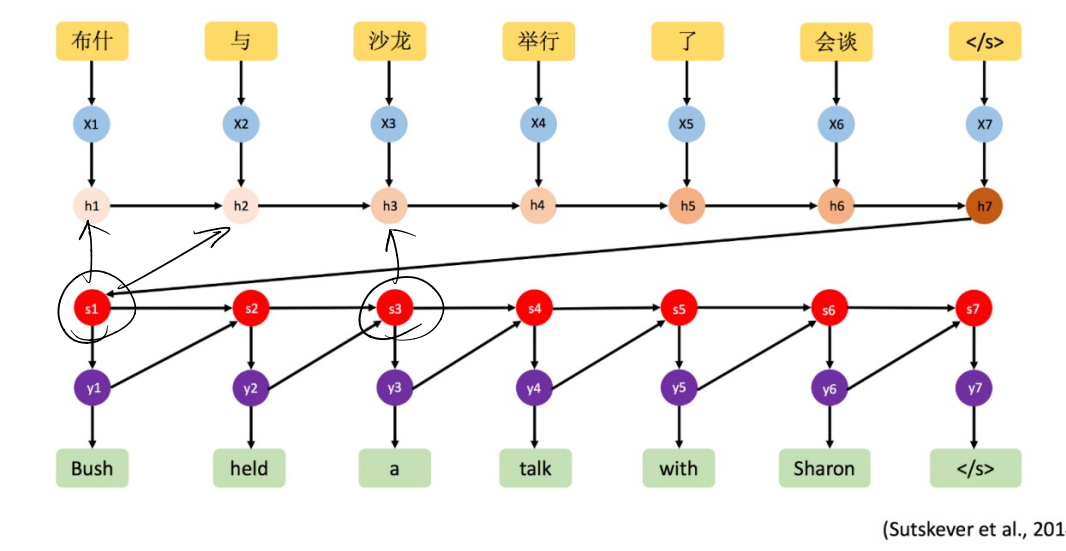


In vanilla Seq2Seq: context variable  $C \rightarrow$  token 1

$\rightarrow$  token 2

"Write using same Summary when generating at All time step"

$\Rightarrow$  Not desirable

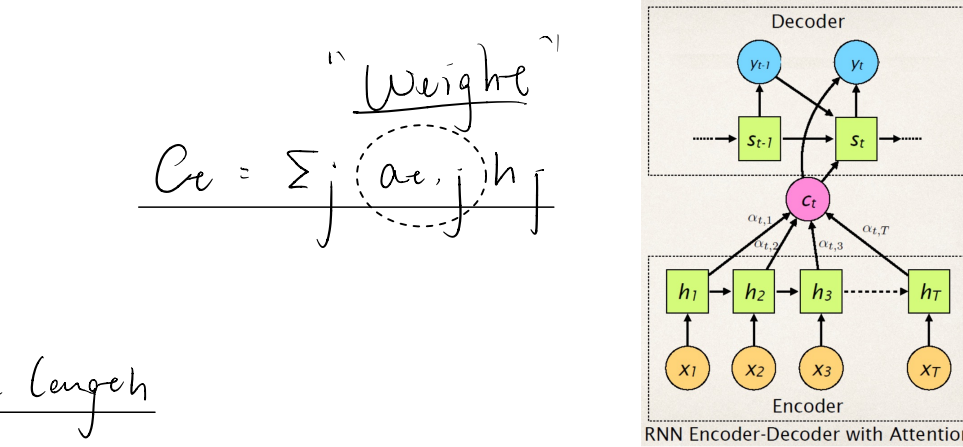


in fact, We want to focus on different parts of input at

different time

- Focus at  $h_1$  at time step 1,
- Focus at  $h_4$  at time step 2,
- Focus at  $h_6$  at time step 4,
- ...

$\therefore C_t$ : context variable at time step t



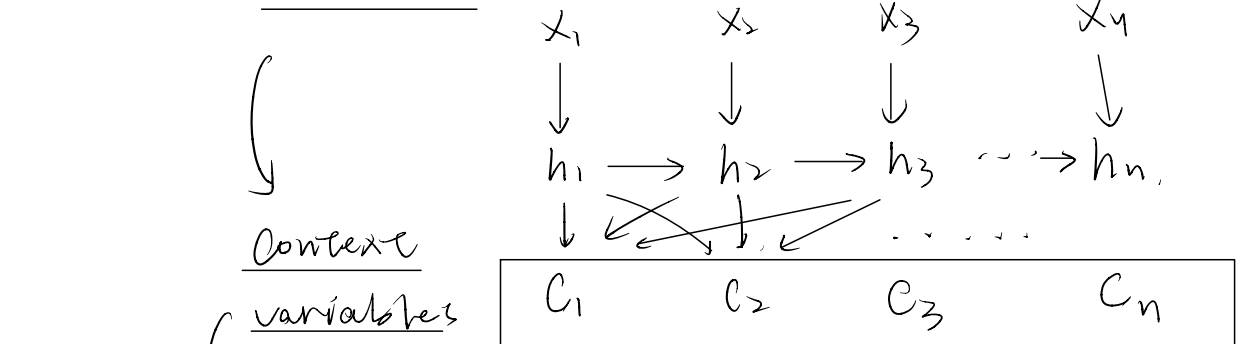
At different step, we use different value of  $\alpha$

"What we want to see / What we don't want to see"

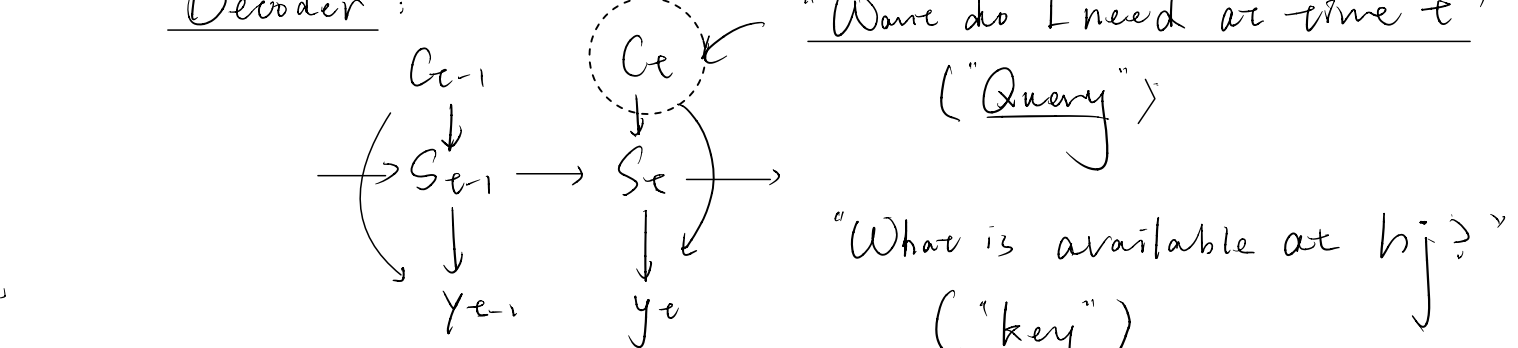
$$C_t = \alpha_1 h_1 + \alpha_2 h_2 + \dots + \alpha_n h_n$$

What we want to see What we don't want to see

### Encoder



Decoder: "What do I need at time t?" ("Query")



"What is available at  $h_j$ ?" ("key")

$$C_t = \sum_j \alpha_{t,j} h_j$$

attention paid to  $h_j$  at time t

$$\alpha_{t,j} = \frac{e^{w_{t,j}}}{\sum_j e^{w_{t,j}}} \quad (\text{softmax})$$

where  $w_{t,j} = C_{t-1}^T \cdot h_j$  ("Query" \* "key")

Looking for  $(C_{t-1}) \rightarrow$  "How much we want it"

what we have

$h_j$

"dynamic Weight"  $\Rightarrow$  Depends on input

### Attention Example:

