



Sequence Modeling: Recurrent and Recursive Nets

Jie Tang

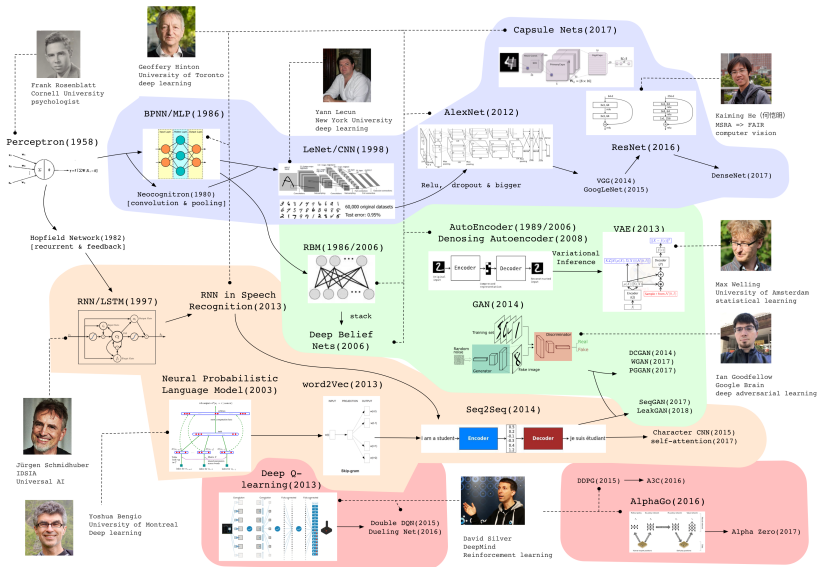
Tsinghua University

Overview

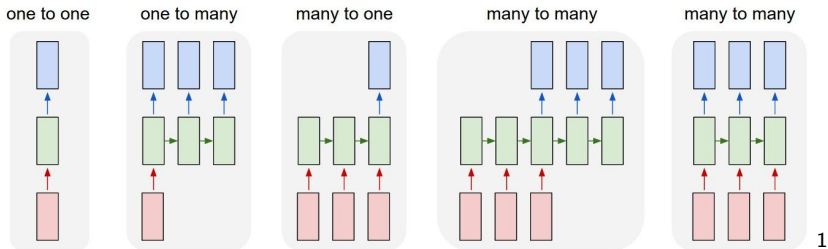
- 1 Introduction
- 2 Bidirectional RNN
- 3 Encoder-Decoder Architectures
- 4 Long-Term Dependencies
- 5 Summary

Introduction

- A recurrent neural network (RNN) is a class of artificial neural network where connections between units form a directed cycle.
- Unlike feedforward neural networks, RNNs can use their internal memory to process arbitrary sequences of inputs.
- The first RNN was invented in 1980s.



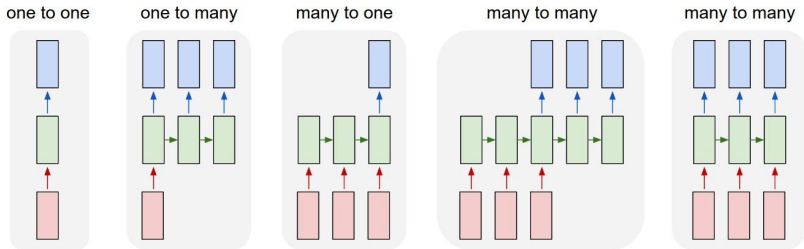
Recurrent Networks offer a lot of flexibility...



Simple neural network

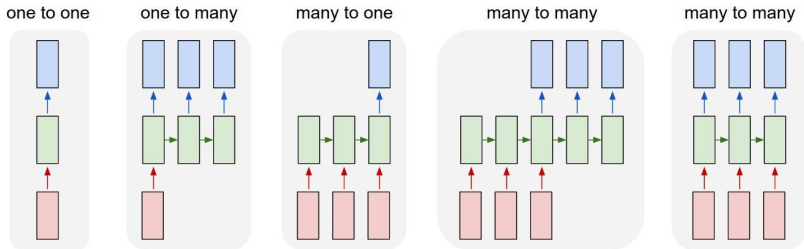
¹The example is from Fei-Fei Li's slides.

Recurrent Networks offer a lot of flexibility...



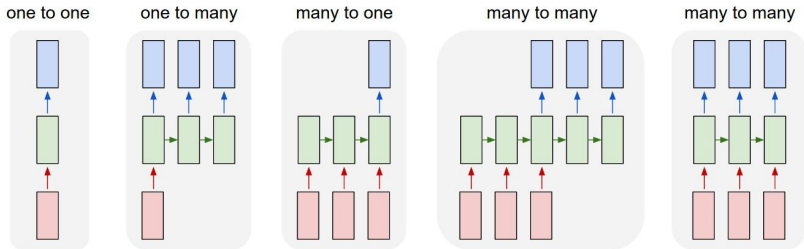
e.g., **Image captioning:** image $\rightarrow$ sequence-of-words

Recurrent Networks offer a lot of flexibility...



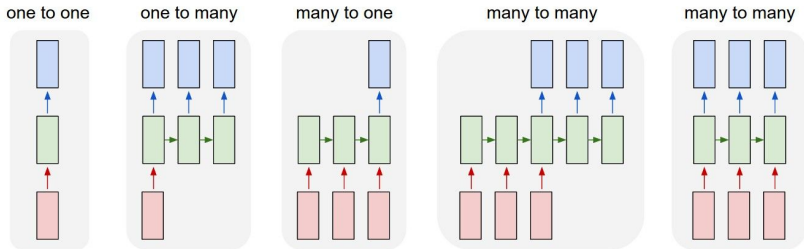
e.g., **Sentiment analysis:** sequence-of-words $\rightarrow$ sentiment

Recurrent Networks offer a lot of flexibility...

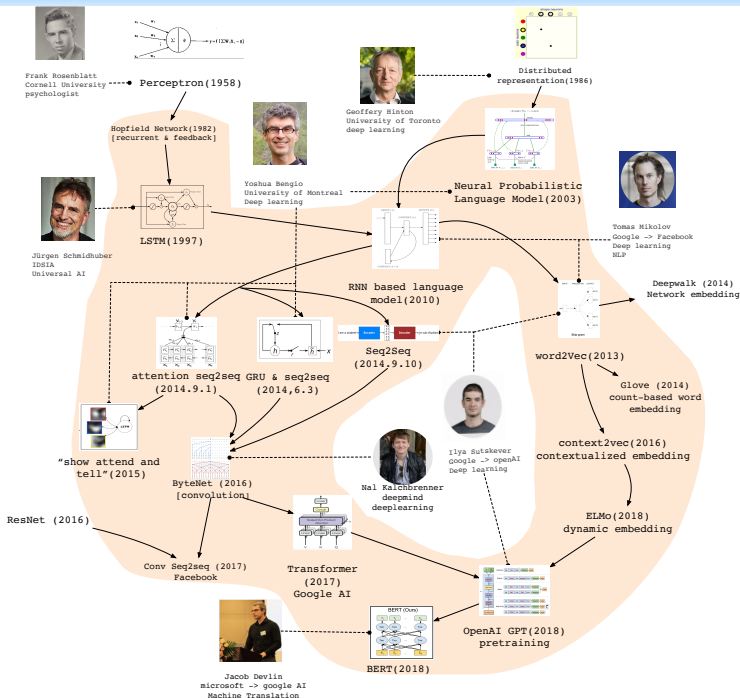


e.g., **Machine translation:** sequence $\rightarrow$ sequence

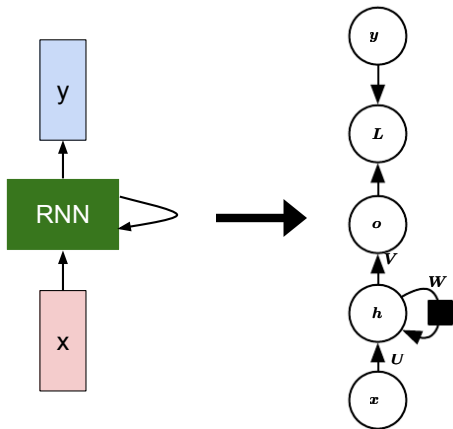
Recurrent Networks offer a lot of flexibility...



e.g., **POS Tagging:** sequence $\rightarrow$ sequence



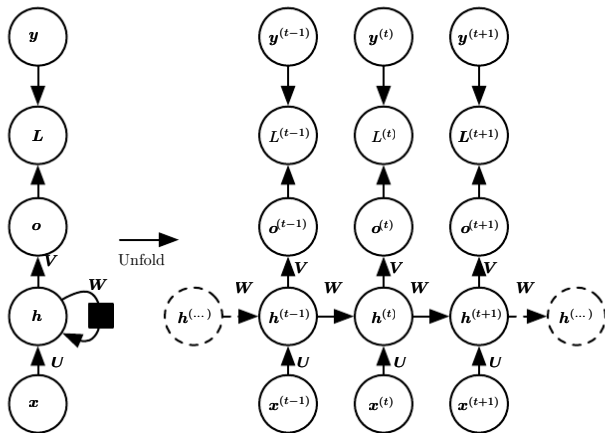
Recurrent Neural Networks — An Example



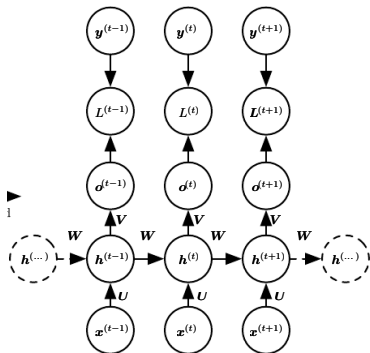
- Input x and output y are from the training data;
- o is the estimated output value;
- L is to evaluate the difference (loss) between the estimated output $\hat{y} = \text{softmax}(o)$ and the true output y ;
- U , W , V are three weight matrices.

Recurrent Neural Networks — An Example

- Let us unfold the graphical representation.



Recurrent Neural Networks — An Example



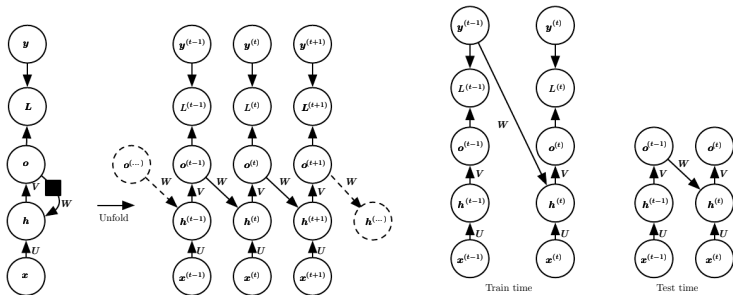
Update Equations

- $h^{(t)} = \tanh(\mathbf{b} + \mathbf{W}h^{(t-1)} + \mathbf{U}x^{(t)})$
- $\mathbf{o}^{(t)} = \mathbf{c} + \mathbf{V}h^{(t)}$
- $\hat{y}^{(t)} = \text{softmax}(\mathbf{o}^{(t)})$
- **Objective function:**
$$L = \sum_t L^{(t)} = \sum_t \log P(y^{(t)} | \{\mathbf{x}^{(1)}, \dots, \mathbf{x}^{(t)}\})$$

Recurrent Neural Networks

- There are many different types of recurrent neural networks
 - Recurrent connections between hidden units
 - Recurrent connections from the output at one time step to the hidden units at the next time step
 - Recurrent connections between hidden units but produces only one output
 - ...

Recurrent Neural Networks – Output Recurrence



(a) Network with Output Recurrence

(b) Teacher Forcing

Figure 1: Other Types of RNN. Output Recurrence: the only recurrence is from the output to the hidden layer; Teacher Forcing: training technique for RNN.

Conditional likelihood of Teacher Forcing:

$$\log p(y^{(1)}, y^{(2)} | x^{(1)}, x^{(2)}) = \log p(y^{(2)} | y^{(1)}, x^{(1)}, x^{(2)}) + \log p(y^{(1)} | x^{(1)}, x^{(2)})$$

Discussion on Output Recurrence

Advantage

- all the time steps are decoupled, training can be in parallel

Disadvantage

- lack hidden-to-hidden recurrence, strictly less powerful

Solution

- randomly choose to use generated values or actual values as input

Gradient Computation in RNN

The gradient can be computed using a generalized back-propagation algorithm, specifically called **back-propagation through time (BPTT)**:

$$\begin{aligned}(\nabla_{\mathbf{o}^{(t)}} L)_i &= \frac{\partial L}{\partial o_i^{(t)}} = \frac{\partial L}{\partial L^{(t)}} \frac{\partial L^{(t)}}{\partial o_i^{(t)}} = \hat{y}_i^{(t)} - 1_{i, y^{(t)}} \\ \nabla_{\mathbf{h}^{(t)}} L &= \mathbf{W}^\top (\nabla_{\mathbf{h}^{(t+1)}} L) \text{diag} \left(1 - (\mathbf{h}^{(t+1)})^2 \right) + \mathbf{V}^\top (\nabla_{\mathbf{o}^{(t)}} L) \\ \nabla_{\mathbf{c}} L &= \sum_t \nabla_{\mathbf{o}^{(t)}} L \\ \nabla_{\mathbf{b}} L &= \sum_t \text{diag} \left(1 - (\mathbf{h}^{(t)})^2 \right) \nabla_{\mathbf{h}^{(t)}} L \\ \nabla_{\mathbf{v}} L &= \sum_t (\nabla_{\mathbf{o}^{(t)}} L) \mathbf{h}^{(t)\top} \\ \nabla_{\mathbf{w}} L &= \sum_t \text{diag} \left(1 - (\mathbf{h}^{(t)})^2 \right) (\nabla_{\mathbf{h}^{(t)}} L) \mathbf{h}^{(t-1)\top} \\ \nabla_{\mathbf{u}} L &= \sum_t \text{diag} \left(1 - (\mathbf{h}^{(t)})^2 \right) (\nabla_{\mathbf{h}^{(t)}} L) \mathbf{x}^{(t)\top}\end{aligned} \tag{1}$$

RNN as Directed Graphical Models

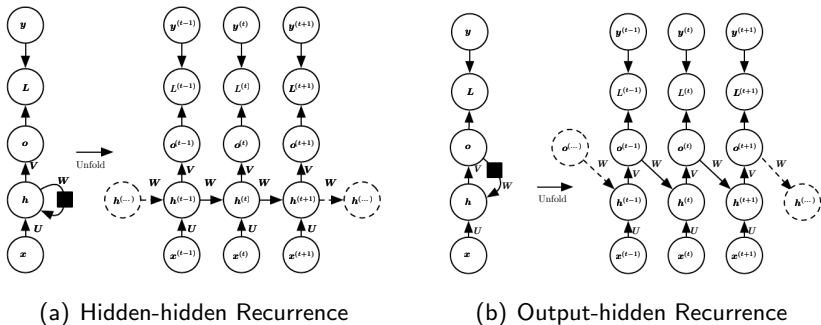
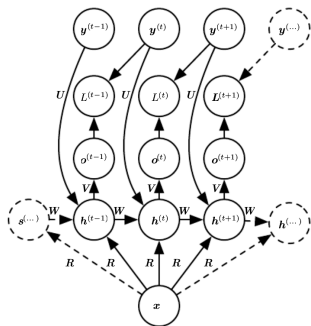
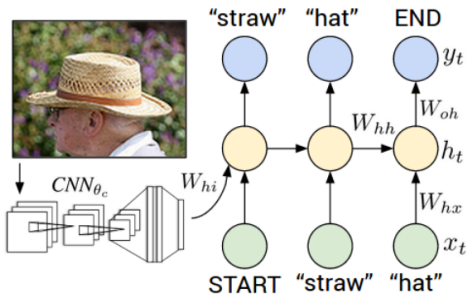


Figure 2: Two Types of RNN we have discussed. Left: maximize the log-likelihood $\log P(\mathbf{y}^{(t)} | \mathbf{x}^{(1)}, \dots, \mathbf{x}^{(t)})$; Right: maximize the log-likelihood $\log P(\mathbf{y}^{(t)} | \mathbf{x}^{(1)}, \dots, \mathbf{x}^{(t)}, \mathbf{y}^{(1)}, \dots, \mathbf{y}^{(t-1)})$

Modeling Sequences Conditioned on Context with RNNs



(a)



(b)

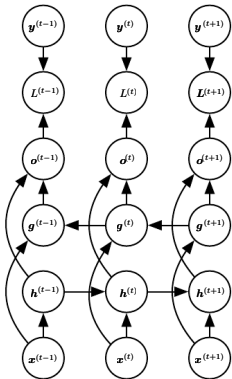
Figure 3: RNN for tasks like image captioning Left: a single image as input, the model produces a sequence of words describing the image. Right: a real example from Feifei Li's CVPR'15 paper. However, the RNN is only conditioned on the image information at the first time step

Overview

- 1 Introduction
- 2 Bidirectional RNN**
- 3 Encoder-Decoder Architectures
- 4 Long-Term Dependencies
- 5 Summary

Bidirectional RNNs

For classification you want to incorporate information from both preceding and following.



- Two hidden layers, one for the left-to-right propagation and another for the right-to-left propagation

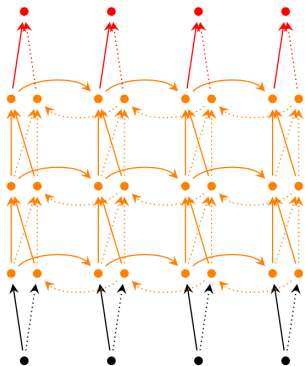
$$\vec{h}_t = f(\vec{W}x_t + \vec{V}\vec{h}_{t-1} + \vec{b})$$

$$\bar{h}_t = f(\bar{W}x_t + \bar{V}\bar{h}_{t-1} + \bar{b})$$

$$\hat{y} = g(Uh_t + c) = g(U[\vec{h}_t, \bar{h}_t] + c) \quad (2)$$

Bidirectional RNNs

Deep Bidirectional RNN



$$\vec{h}_t^i = f(\vec{W}^i h_t^{i-1} + \vec{V}^i \vec{h}_{t-1}^i + \vec{b}^i)$$

$$\vec{h}_t^i = f(\vec{W}^i h_t^{i-1} + \vec{V}^i \vec{h}_{t-1}^i + \vec{b}^i)$$

$$\hat{y} = g(Uh_t + c) = g(U[\vec{h}_t^L, \vec{h}_t^L] + c) \quad (3)$$

A deep bidirectional RNN
with three RNN layers.

Overview

- 1 Introduction
- 2 Bidirectional RNN
- 3 Encoder-Decoder Architectures**
- 4 Long-Term Dependencies
- 5 Summary

Encoder-Decoder Sequence-to-Sequence Architectures

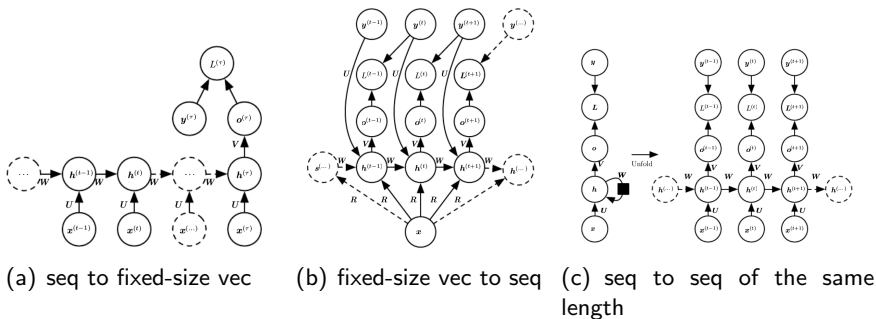
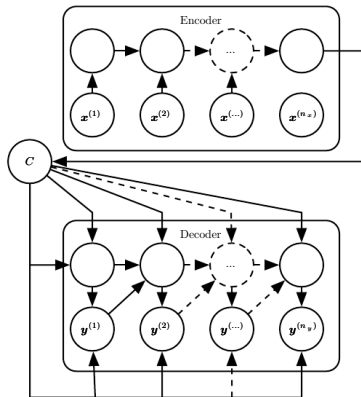


Figure 4: RNNs that we have discussed so far.

Can we train a RNN model to map an input sequence to an output sequence which is **not necessarily of the same length**?

This is important in many applications: speech recognition, machine translation, or question answering.

Encoder-Decoder Sequence-to-Sequence Architectures



When C is a fixed-size vector, this architecture could be reviewed as a **sequence to fixed-size vector** RNN and a **fixed-size vector to sequence** RNN.

Figure 5: Encoder-decoder or sequence-to-sequence RNN architecture.

Encoder-Decoder Sequence-to-Sequence Architectures

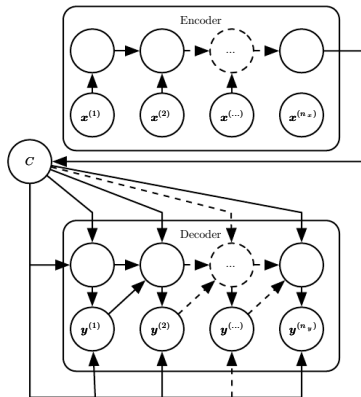


Figure 5: Encoder-decoder or sequence-to-sequence RNN architecture.

When C is a fixed-size vector, this architecture could be reviewed as a **sequence to fixed-size vector** RNN and a **fixed-size vector to sequence** RNN.

Limitation: C may be too small to summarize a long sequence (Bahdanau et al. (2015)).

Encoder-Decoder Architectures with Attention Mechanism

C may be too small to summarize a long sequence (Bahdanau et al. (2015)). They proposed to make C a **variable-length** sequence. Additionally, they introduced the **attention mechanism**.

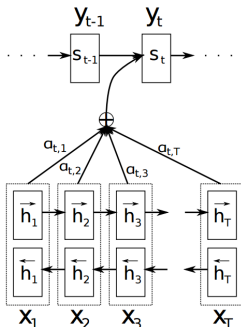
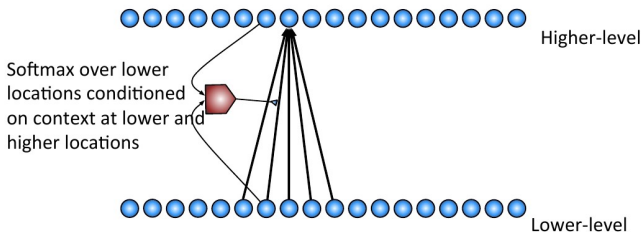


Figure 6: Formulation: $a(\cdot)$ is an alignment model which scores how well the inputs around position j and the output at position i match, which is parameterized as a feedforward neural network

Encoder-Decoder Architectures with Attention Mechanism

- Consider an input (or intermediate) sequence
- Consider an upper level representation, which can choose “where to look” when producing some task (e.g., machine translation) at each position



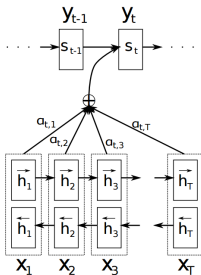
Encoder-Decoder Architectures with Attention Mechanism

- Decoder: is often trained to predict the next word y_t given the context vector c and all the previous words $\{y_1, y_2, \dots, y_{t-1}\}$

$$p(y_t | y_1, \dots, y_{t-1}, x) = g(y_{t-1}, s_t, c_t) \quad (4)$$

where s_t is an RNN hidden state for position t , computed by

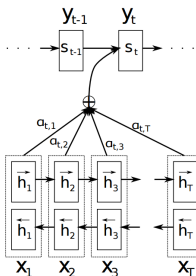
$$s_i = f(s_{i-1}, y_{i-1}, c_i) \quad (5)$$



Attention Mechanism

- The context vector c_t depends on a sequence of *annotations* ($h_1, \dots, h_T$) to which an encoder maps the input sentence, computed by

$$c_t = \sum_{j=1}^T \alpha_{tj} h_j \quad (6)$$



Attention Mechanism

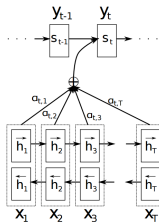
- The weight α is so called “attention matrix” (used to choose “where to look”) and is computed by

$$\alpha_{tj} = \frac{\exp\{e_{tj}\}}{\sum_{k=1}^T e_{tk}} \quad (7)$$

where

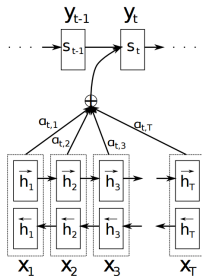
$$e_{tj} = h(s_{t-1}, h_j) \quad (8)$$

is an attention model which scores how well the inputs around position j and the output at position t match.



Attention Mechanism

- The attention model can be parametrized as a feedforward neural network (or multilayer perceptron, etc.) which is jointly trained with all the other components of the system.



Attention Mechanism with Seq-to-Seq Architectures

To summarize, the attention mechanism with sequence-to-sequence architectures have the following formalization:

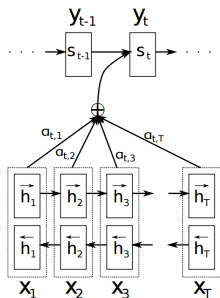


Figure 7: Formulatin: $a(\cdot)$ is an alignment model which scores how well the inputs around position j and the output at position i match, which is parameterized as a feedforward neural network

$$\begin{aligned}
 p(y_i | y_1, \dots, y_{i-1}, \mathbf{x}) &= g(y_{i-1}, s_i, c_i) \\
 s_i &= f(s_{i-1}, y_{i-1}, c_i) \\
 c_i &= \sum_{j=1}^{T_x} \alpha_{ij} h_j \\
 \alpha_{ij} &= \frac{\exp(e_{ij})}{\sum_{j=1}^{T_x} \exp(e_{ij})} \\
 e_{ij} &= a(s_{i-1}, h_j)
 \end{aligned}$$

Overview

- 1 Introduction
- 2 Bidirectional RNN
- 3 Encoder-Decoder Architectures
- 4 Long-Term Dependencies**
- 5 Summary

Let us review the vanishing and exploding gradient problem

Let's first review what is long-term dependencies mentioned in Lecture "Optimization for Training Deep Models".

Pure Linear Case

- Consider a RNN with no nonlinearity

$$h^{(t)} = Wh^{(t-1)} + Vx^{(t)} + b$$

$$\frac{\partial h^{(t)}}{\partial h^{(k)}} = \prod_{j=k+1}^t \frac{\partial h^{(j)}}{\partial h^{(j-1)}} = \prod_{j=k+1}^t W^T = (W^T)^{t-k} \quad (9)$$

- Suppose that W has an eigendecomposition $W^T = V\text{diag}(\lambda)V^{-1}$, then

$$(W^T)^{t-k} = V\text{diag}(\lambda)^{t-k}V^{-1} \quad (10)$$

- Vanishing and exploding gradient problem:** Any eigenvalues λ_i that are not near an absolute value of 1 will either explode (if $|\lambda_i| > 1$) or vanish (if $|\lambda_i| < 1$).

Vanishing and exploding gradient problem

Non-linear Case

Let's consider the Jacobi Matrix and its norm:

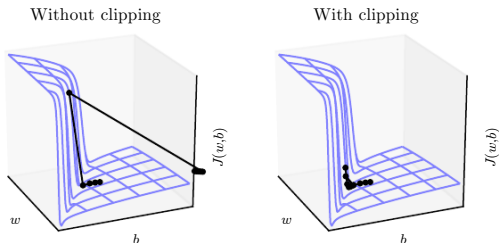
$$\begin{aligned}h^{(t)} &= f(Wh^{(t-1)} + Vx^{(t)} + b) \\ \frac{\partial h^{(t)}}{\partial h^{(k)}} &= \prod_{j=k+1}^t \frac{\partial h^{(j)}}{\partial h^{(j-1)}} = \prod_{j=k+1}^t W^\top \text{diag} \left(f'(h^{(j-1)}) \right) \\ \left\| \frac{\partial h^{(t)}}{\partial h^{(j-1)}} \right\| &\leq \|W^\top\| \times \left\| \text{diag} \left(f'(h^{(j-1)}) \right) \right\|\end{aligned} \quad (11)$$

Non-linear activation function may

- bound recurrent dynamics, alleviating the exploding gradient problem.
- make vanishing gradient problem severer.

The remaining of this slides discuss various approaches to **reduce the difficulty of learning long-term dependencies**.

Trick for exploding gradient: Clipping Gradients



The solution first introduced by Mikolov is to clip gradients...

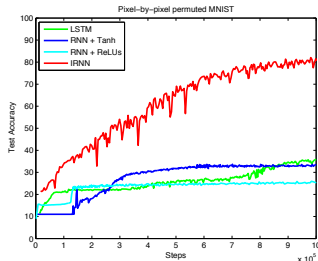
Algorithm 1 Pseudo-code for norm clipping

```
 $\hat{\mathbf{g}} \leftarrow \frac{\partial \mathcal{E}}{\partial \theta}$   
if  $\|\hat{\mathbf{g}}\| \geq threshold$  then  
     $\hat{\mathbf{g}} \leftarrow \frac{threshold}{\|\hat{\mathbf{g}}\|} \hat{\mathbf{g}}$   
end if
```

It makes big difference in RNN!

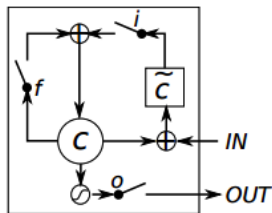
For vanishing gradients: Initialization+ReLU

- Initialize $W(\star)$'s to identity matrix I and ReLU $f(z) = \max(0, z)$
- It makes huge difference!
- Several ideas about initialization for DL Initialization idea first introduced in Parsing with Compositional Vector Grammars, Socher et al. 2013
A Simple Way to Initialize Recurrent Networks of Rectified

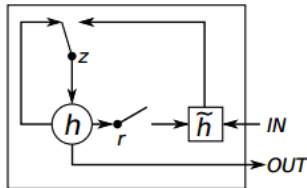


Gated RNNs and Long Short-Term Memory

- The **Long-Short-Term-Memory (LSTM)** algorithm was proposed in 1997 (Hochreiter and Schmidhuber, 1997).
- Recent work on **gated RNNs**, **Gated Recurrent Units (GRU)** was proposed in 2014
- The LSTM has been found extremely **successful** in many applications, such as unconstrained handwriting recognition, speech recognition, handwriting generation, machine translation, image captioning and parsing.



(a) Long Short-Term Memory



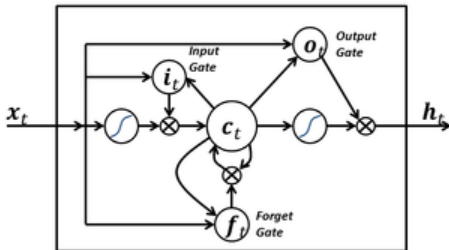
(b) Gated Recurrent Unit

LSTM

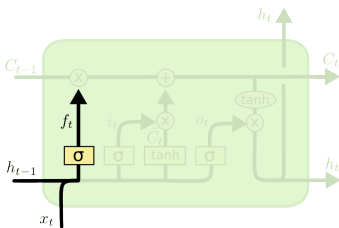
- Forget Gate: $\mathbf{f}^{(t)} = \text{sigmoid}(\mathbf{W}^f \mathbf{h}^{(t-1)} + \mathbf{U}^f \mathbf{x}^{(t)} + \mathbf{b}^f)$
- Input Gate: $\mathbf{i}^{(t)} = \text{sigmoid}(\mathbf{W}^i \mathbf{h}^{(t-1)} + \mathbf{U}^i \mathbf{x}^{(t)} + \mathbf{b}^i)$
- Output Gate: $\mathbf{o}^{(t)} = \text{sigmoid}(\mathbf{W}^o \mathbf{h}^{(t-1)} + \mathbf{U}^o \mathbf{x}^{(t)} + \mathbf{b}^o)$
- Cell State Vector:

$$\mathbf{c}^{(t)} = \mathbf{f}^{(t)} \circ \mathbf{c}^{(t-1)} + \mathbf{i}^{(t)} \circ \tanh(\mathbf{W}^c \mathbf{h}^{(t-1)} + \mathbf{U}^c \mathbf{x}^{(t)} + \mathbf{b}^c)$$

- Final Hidden State: $\mathbf{h}^{(t)} = \mathbf{o}^{(t)} \circ \tanh(\mathbf{c}^{(t)})$



LSTM: Forget Gate

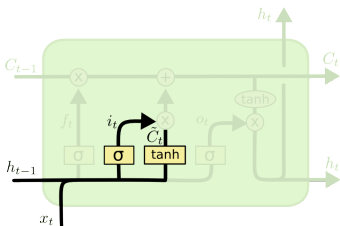


$$f_t = \sigma(W_f \cdot [h_{t-1}, x_t] + b_f)$$

Forget gate looks at h_{t-1} and x_t , and outputs a number between 0 and 1 for each number in the cell state C_{t-1} .

- A one represents “completely keep this”.
- A zero represents “completely get rid of this”.

LSTM: Input Gate

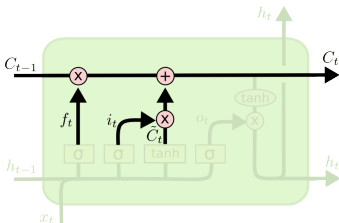


$$i_t = \sigma(W_i \cdot [h_{t-1}, x_t] + b_i)$$
$$\tilde{C}_t = \tanh(W_C \cdot [h_{t-1}, x_t] + b_C)$$

Decide what new information we're going to store in the cell state.

- a sigmoid layer called the “input gate layer” decides which values we'll update.
- a tanh layer creates a vector of new candidate values of cell state.

LSTM: Update Cell State

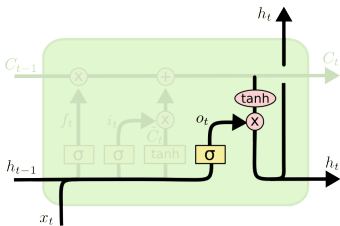


$$C_t = f_t * C_{t-1} + i_t * \tilde{C}_t$$

Update the old cell state

- multiply the old state by f_t (results from forget gate), forgetting the things we decided to forget.
- add the new candidate values, scaled by i_t — how much we decided to update each state value.

LSTM: Output Gate



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

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

Output final results.

- a sigmoid layer which decides what parts of the cell state we're going to output.
- put the cell state through $\tanh$ and multiply it by the output of the sigmoid gate.

Gated Recurrent Units (GRU, 2014)

Let's denote element-wise product (Hadamard product) as $\circ$

- Update Gate: $\mathbf{z}^{(t)} = \text{sigmoid}(\mathbf{W}^z \mathbf{h}^{(t-1)} + \mathbf{U}^z \mathbf{x}^{(t)} + \mathbf{b}^z)$
- Reset Gate: $\mathbf{r}^{(t)} = \text{sigmoid}(\mathbf{W}^r \mathbf{h}^{(t-1)} + \mathbf{U}^r \mathbf{x}^{(t)} + \mathbf{b}^r)$
- New Memory Content:
 $\bar{\mathbf{h}}^{(t)} = \tanh(\mathbf{W}(\mathbf{r}^{(t)} \circ \mathbf{h}^{(t-1)}) + \mathbf{U}\mathbf{x}^{(t)} + \mathbf{b})$
- Final Memory: $\mathbf{h}^{(t)} = \mathbf{z}^{(t)} \circ \mathbf{h}^{(t-1)} + (1 - \mathbf{z}^{(t)}) \circ \bar{\mathbf{h}}^{(t)}$

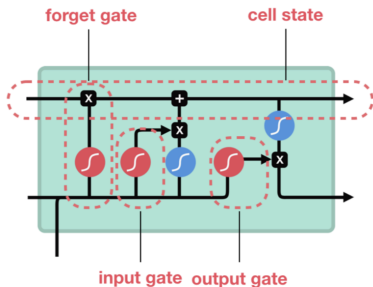
Remark

- The reset gate $\mathbf{r}^{(t)}$ close to 0 makes RNN ignore previous hidden state — allow model to drop information that is irrelevant
- The Update gate $\mathbf{z}^{(t)}$ control how much of past state should matter now.

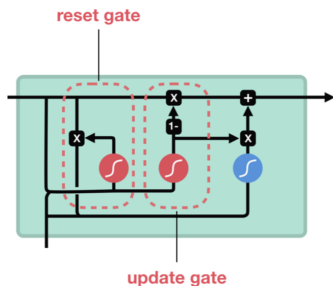
LSTM vs GRU

- GRUs achieve similar performance in a variety of tasks compared to LSTM.
- GRUs train faster than LSTMs on less training data.
- LSTM outperform GRUs in tasks requiring modeling long-distance relations.

LSTM



GRU



Overview

- 1 Introduction
- 2 Bidirectional RNN
- 3 Encoder-Decoder Architectures
- 4 Long-Term Dependencies
- 5 Summary**

Summary

- RNN is able to model sequential data with lots of flexibilities
- RNN can be also extended to model bidirectional dependencies and work with attention mechanism.
- LSTM is a special RNN and can be used to model long-term dependencies. GRU is a simplification of LSTM that runs faster and requires less training data.

Thanks.

HP: <http://keg.cs.tsinghua.edu.cn/jietang/>
Email: jietang@tsinghua.edu.cn