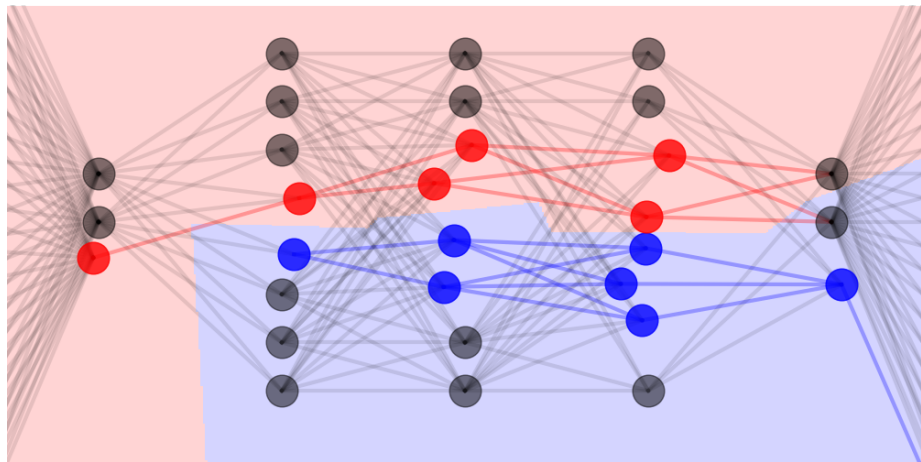


TU-Delft Machine and Deep Learning



RNNs

Lecture 6

Lecturer: Jan van Gemert

Topic: Recurrent networks

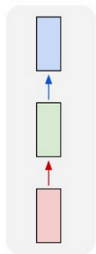
- RNNs
- Gradients in RNNs
- GRU and LSTM

Deep Learning book chapter 10;
D2L-book chapters: 9, 10

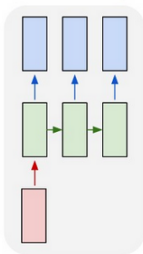
Types of sequential data

<http://karpathy.github.io/2015/05/21/rnn-effectiveness/>

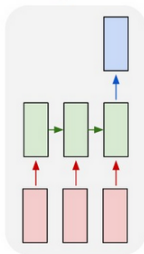
one to one



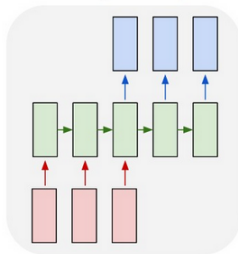
one to many



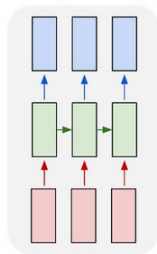
many to one



many to many



many to many



Sequences in the input, in the hidden layers, or in the output. ¹

¹Andrej Karpathy's blog: "The Unreasonable Effectiveness of Recurrent Neural Networks"

Sequence data

Lets deep learn a language translation system.

Sequence data

Lets deep learn a language translation system.

Consider an example sentence, where each word is an input:

The student, after reading the book, was well prepared.

Sequence data

Lets deep learn a language translation system.

Consider an example sentence, where each word is an input:

The student, after reading the book, was well prepared.

Q: What are reasons to not use a feed forward MLP network?

Sequence data

Lets deep learn a language translation system.

Consider an example sentence, where each word is an input:

The student, after reading the book, was well prepared.

Q: What are reasons to not use a feed forward MLP network?

A: Sequences can have variable length.

Sequence data

Lets deep learn a language translation system.

Consider an example sentence, where each word is an input:

The student, after reading the book, was well prepared.

Q: What are reasons to not use a feed forward MLP network?

A: Sequences can have variable length.

Q: So what? Can't I just pad with zeros?

Sequence data

Lets deep learn a language translation system.

Consider an example sentence, where each word is an input:

The student, after reading the book, was well prepared.

Q: What are reasons to not use a feed forward MLP network?

A: Sequences can have variable length.

Q: So what? Can't I just pad with zeros?

A: No; Feed forward will learn exact word locations;

Sequence data

Lets deep learn a language translation system.

Consider an example sentence, where each word is an input:

The student, after reading the book, was well prepared.

Q: What are reasons to not use a feed forward MLP network?

A: Sequences can have variable length.

Q: So what? Can't I just pad with zeros?

A: No; Feed forward will learn exact word locations; Q: Why is that a problem?

Sequence data

Lets deep learn a language translation system.

Consider an example sentence, where each word is an input:

The student, after reading the book, was well prepared.

Q: What are reasons to not use a feed forward MLP network?

A: Sequences can have variable length.

Q: So what? Can't I just pad with zeros?

A: No; Feed forward will learn exact word locations; Q: Why is that a problem?

A: Need to see all possible word/sentence options during training (impossible).

Example: Predict the next character

DL book, Chapter 10; and: <http://karpathy.github.io/2015/05/21/rnn-effectiveness/>



Example: Predict the next character

DL book, Chapter 10; and: <http://karpathy.github.io/2015/05/21/rnn-effectiveness/>



Q: Explain the layout?

Example: Predict the next character

DL book, Chapter 10; and: <http://karpathy.github.io/2015/05/21/rnn-effectiveness/>



Q: Explain the layout? A: A 2-layer network for letter prediction.

Example: Predict the next character

DL book, Chapter 10; and: <http://karpathy.github.io/2015/05/21/rnn-effectiveness/>



Q: Explain the layout? A: A 2-layer network for letter prediction.

Q: How to represent input letters?

Example: Predict the next character

DL book, Chapter 10; and: <http://karpathy.github.io/2015/05/21/rnn-effectiveness/>

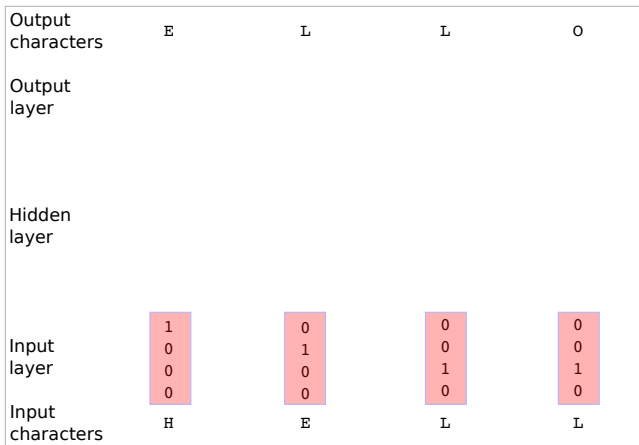


Q: Explain the layout? A: A 2-layer network for letter prediction.

Q: How to represent input letters? A: One-Hot encoding.

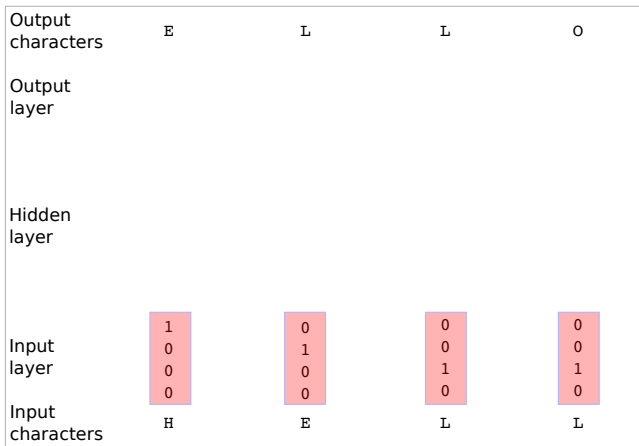
Example: Predict the next character

DL book, Chapter 10; and: <http://karpathy.github.io/2015/05/21/rnn-effectiveness/>



Example: Predict the next character

DL book, Chapter 10; and: <http://karpathy.github.io/2015/05/21/rnn-effectiveness/>



Q: How to represent output?

Example: Predict the next character

DL book, Chapter 10; and: <http://karpathy.github.io/2015/05/21/rnn-effectiveness/>

Output characters	E	L	L	O
Output layer	1.0 2.2 -3 4.1	0.5 0.3 -1 1.2	0.1 0.5 1.9 -1.1	0.2 -1.5 -0.1 2.2
Hidden layer				
Input layer	1 0 0 0	0 1 0 0	0 0 1 0	0 0 1 0
Input characters	H	E	L	L

Example: Predict the next character

DL book, Chapter 10; and: <http://karpathy.github.io/2015/05/21/rnn-effectiveness/>

Output characters	E	L	L	O
	1.0	0.5	0.1	0.2
Output layer	2.2	0.3	0.5	-1.5
	-3	-1	1.9	-0.1
	4.1	1.2	-1.1	2.2
Hidden layer				
Input layer	1	0	0	0
	0	1	0	0
	0	0	1	1
	0	0	0	0
Input characters	H	E	L	L

Q: What do the bold numbers mean?

Example: Predict the next character

DL book, Chapter 10; and: <http://karpathy.github.io/2015/05/21/rnn-effectiveness/>

Output characters	E	L	L	O
	1.0	0.5	0.1	0.2
Output layer	2.2	0.3	0.5	-1.5
	-3	-1	1.9	-0.1
	4.1	1.2	-1.1	2.2
Hidden layer				
Input layer	1	0	0	0
	0	1	0	0
	0	0	1	1
	0	0	0	0
Input characters	H	E	L	L

Q: What do the bold numbers mean? A: Target (make high)

Example: Predict the next character

DL book, Chapter 10; and: <http://karpathy.github.io/2015/05/21/rnn-effectiveness/>

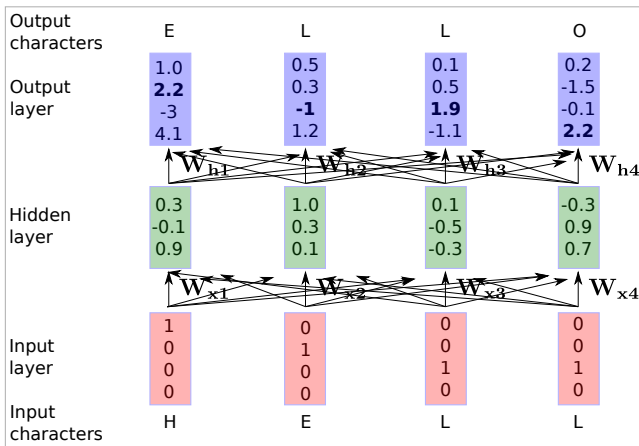
Output characters	E	L	L	O
Output layer	1.0 2.2 -3 4.1	0.5 0.3 -1 1.2	0.1 0.5 1.9 -1.1	0.2 -1.5 -0.1 2.2
Hidden layer				
Input layer	1 0 0 0	0 1 0 0	0 0 1 0	0 0 1 0
Input characters	H	E	L	L

Q: What do the bold numbers mean? A: Target (make high)

Q: What would a fully connected network look like?

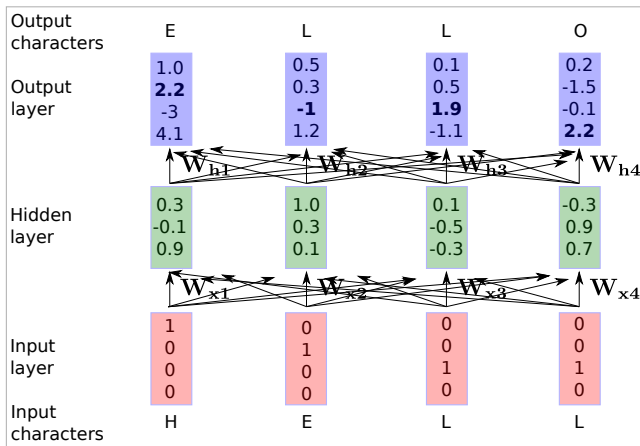
Example: Predict the next character

DL book, Chapter 10; and: <http://karpathy.github.io/2015/05/21/rnn-effectiveness/>



Example: Predict the next character

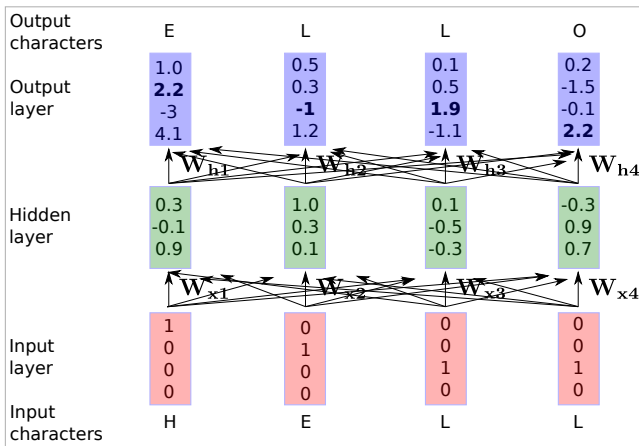
DL book, Chapter 10; and: <http://karpathy.github.io/2015/05/21/rnn-effectiveness/>



Q: Easy solution to remove position-specific weights?

Example: Predict the next character

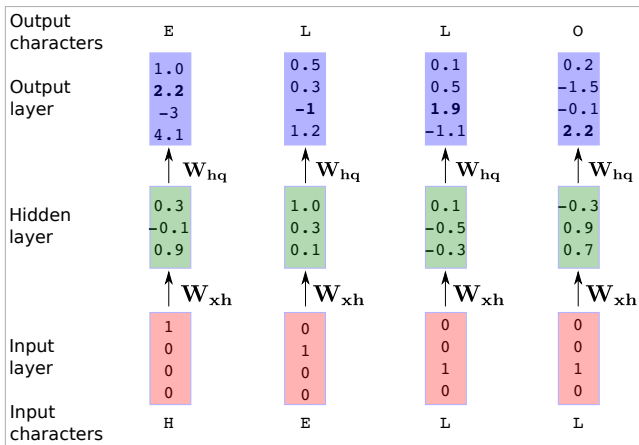
DL book, Chapter 10; and: <http://karpathy.github.io/2015/05/21/rnn-effectiveness/>



Q: Easy solution to remove position-specific weights? A: Share all weights (convolution).

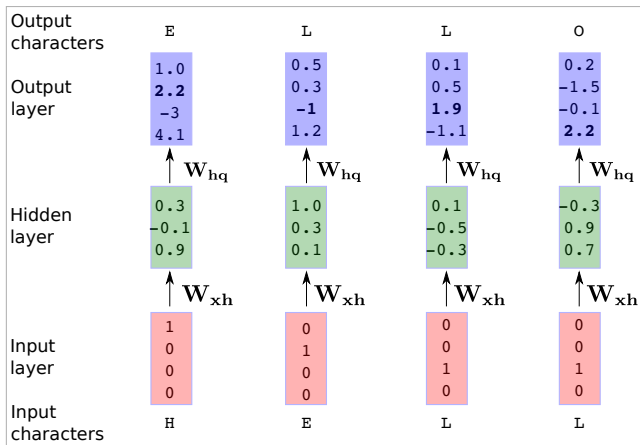
Example: Predict the next character

DL book, Chapter 10; and: <http://karpathy.github.io/2015/05/21/rnn-effectiveness/>



Example: Predict the next character

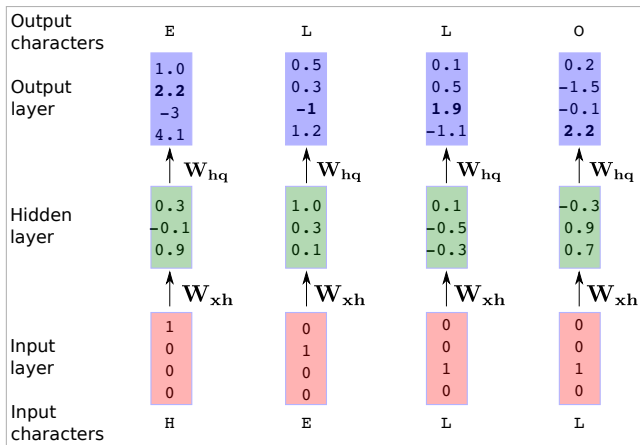
DL book, Chapter 10; and: <http://karpathy.github.io/2015/05/21/rnn-effectiveness/>



Q: What is the problem for 'L'?

Example: Predict the next character

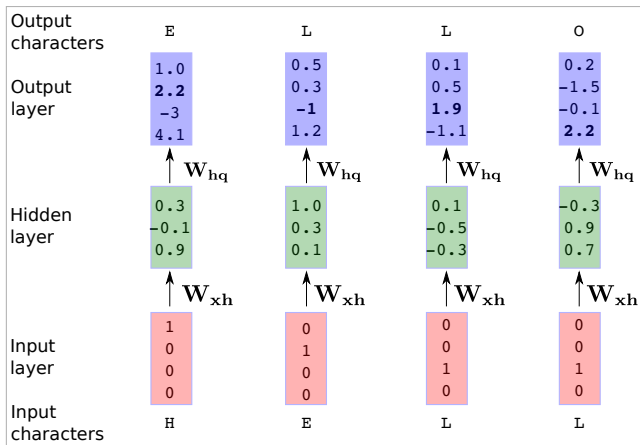
DL book, Chapter 10; and: <http://karpathy.github.io/2015/05/21/rnn-effectiveness/>



Q: What is the problem for 'L'? A: Individual predictions. No past/context.

Example: Predict the next character

DL book, Chapter 10; and: <http://karpathy.github.io/2015/05/21/rnn-effectiveness/>

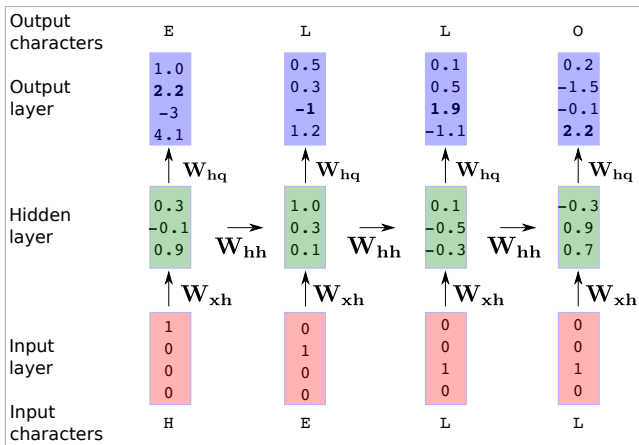


Q: What is the problem for 'L'? A: Individual predictions. No past/context.

Q: How can the past be included?

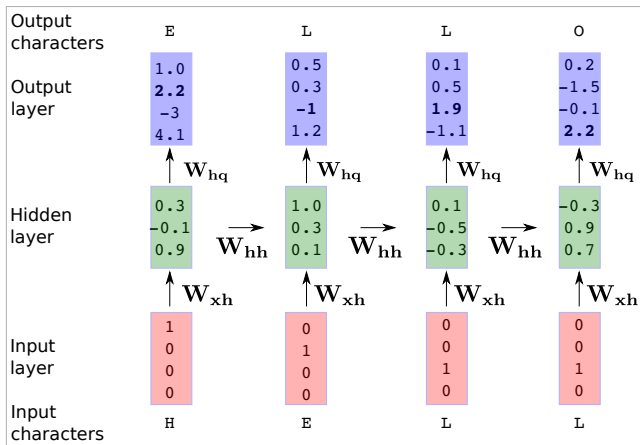
Example: Predict the next character

DL book, Chapter 10; and: <http://karpathy.github.io/2015/05/21/rnn-effectiveness/>



Example: Predict the next character

DL book, Chapter 10; and: <http://karpathy.github.io/2015/05/21/rnn-effectiveness/>



Recurrent Neural Network:

Process a sequence in order and share the same parameters at each time step.

RNN

https://d2l.ai/chapter_recurrent-neural-networks/rnn.html#recurrent-networks-with-hidden-states

Recurrent Neural Network:

Process a sequence in order and share the same parameters at each time step.

RNN

https://d2l.ai/chapter_recurrent-neural-networks/rnn.html#recurrent-networks-with-hidden-states

Recurrent Neural Network:

Process a sequence in order and share the same parameters at each time step.

For a sequence $\mathbf{X}_t$ for $t = 1, 2, 3, \dots, T$

RNN

https://d2l.ai/chapter_recurrent-neural-networks/rnn.html#recurrent-networks-with-hidden-states

Recurrent Neural Network:

Process a sequence in order and share the same parameters at each time step.

For a sequence $\mathbf{X}_t$ for $t = 1, 2, 3, \dots, T$

$$\mathbf{H}_t = \phi(\mathbf{X}_t \mathbf{W}_{\text{xh}} + \mathbf{H}_{t-1} \mathbf{W}_{\text{hh}} + \mathbf{b}_h),$$

RNN

https://d2l.ai/chapter_recurrent-neural-networks/rnn.html#recurrent-networks-with-hidden-states

Recurrent Neural Network:

Process a sequence in order and share the same parameters at each time step.

For a sequence $\mathbf{X}_t$ for $t = 1, 2, 3, \dots, T$

$$\mathbf{H}_t = \phi(\mathbf{X}_t \mathbf{W}_{xh} + \mathbf{H}_{t-1} \mathbf{W}_{hh} + \mathbf{b}_h), \quad \text{output } \mathbf{O}_t = \mathbf{H}_t \mathbf{W}_{hq} + \mathbf{b}_q$$

RNN

https://d2l.ai/chapter_recurrent-neural-networks/rnn.html#recurrent-networks-with-hidden-states

Recurrent Neural Network:

Process a sequence in order and share the same parameters at each time step.

For a sequence $\mathbf{X}_t$ for $t = 1, 2, 3, \dots, T$

$$\mathbf{H}_t = \phi(\mathbf{X}_t \mathbf{W}_{xh} + \mathbf{H}_{t-1} \mathbf{W}_{hh} + \mathbf{b}_h), \quad \text{output } \mathbf{O}_t = \mathbf{H}_t \mathbf{W}_{hq} + \mathbf{b}_q$$

Q: Explain the symbols?

RNN

https://d2l.ai/chapter_recurrent-neural-networks/rnn.html#recurrent-networks-with-hidden-states

Recurrent Neural Network:

Process a sequence in order and share the same parameters at each time step.

For a sequence $\mathbf{X}_t$ for $t = 1, 2, 3, \dots, T$

$$\mathbf{H}_t = \phi(\mathbf{X}_t \mathbf{W}_{xh} + \mathbf{H}_{t-1} \mathbf{W}_{hh} + \mathbf{b}_h), \quad \text{output } \mathbf{O}_t = \mathbf{H}_t \mathbf{W}_{hq} + \mathbf{b}_q$$

Q: Explain the symbols?

- $\mathbf{H}_t$: Activation at current time step t
- ϕ : non-linearity, (often a $\tanh$, range in $(-1,1)$)
- $\mathbf{X}_t$: input at time t
- $\mathbf{W}_{xh}$: Learned weights for input at time t
- $\mathbf{H}_{t-1}$: Activation at previous time step $t - 1$
- $\mathbf{W}_{hh}$: Learned weights: how to use the previous information at $t - 1$
- $\mathbf{W}_{hq}$: Learned weights for output at time t
- $\mathbf{b}_h, \mathbf{b}_q$: Learned bias terms

RNN: Nr parameters doesn't grow with sequence length

https://d2l.ai/chapter_recurrent-neural-networks/rnn.html#recurrent-networks-with-hidden-states

Recurrent Neural Network:

Process a sequence in order and share the same parameters at each time step.

For a sequence $\mathbf{X}_t$ for $t = 1, 2, 3, \dots, T$

$$\mathbf{H}_t = \phi(\mathbf{X}_t \mathbf{W}_{xh} + \mathbf{H}_{t-1} \mathbf{W}_{hh} + \mathbf{b}_h), \quad \text{output } \mathbf{O}_t = \mathbf{H}_t \mathbf{W}_{hq} + \mathbf{b}_q$$

RNN: Nr parameters doesn't grow with sequence length

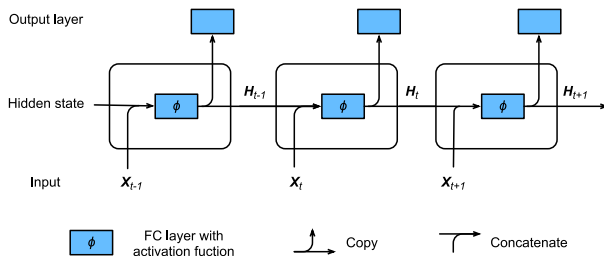
https://d2l.ai/chapter_recurrent-neural-networks/rnn.html#recurrent-networks-with-hidden-states

Recurrent Neural Network:

Process a sequence in order and share the same parameters at each time step.

For a sequence $\mathbf{X}_t$ for $t = 1, 2, 3, \dots, T$

$$\mathbf{H}_t = \phi(\mathbf{X}_t \mathbf{W}_{xh} + \mathbf{H}_{t-1} \mathbf{W}_{hh} + \mathbf{b}_h), \quad \text{output } \mathbf{O}_t = \mathbf{H}_t \mathbf{W}_{hq} + \mathbf{b}_q$$



RNN: Nr parameters doesn't grow with sequence length

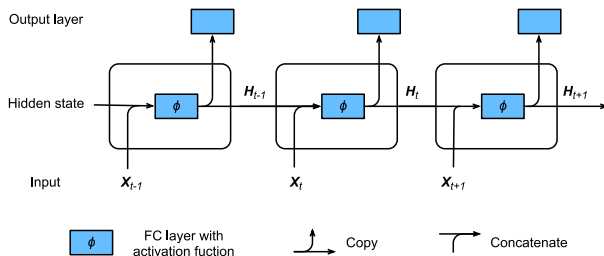
https://d2l.ai/chapter_recurrent-neural-networks/rnn.html#recurrent-networks-with-hidden-states

Recurrent Neural Network:

Process a sequence in order and share the same parameters at each time step.

For a sequence $\mathbf{X}_t$ for $t = 1, 2, 3, \dots, T$

$$\mathbf{H}_t = \phi(\mathbf{X}_t \mathbf{W}_{xh} + \mathbf{H}_{t-1} \mathbf{W}_{hh} + \mathbf{b}_h), \quad \text{output } \mathbf{O}_t = \mathbf{H}_t \mathbf{W}_{hq} + \mathbf{b}_q$$



Q: How does the "concatenate" operator relate to the equation?

RNN: Nr parameters doesn't grow with sequence length

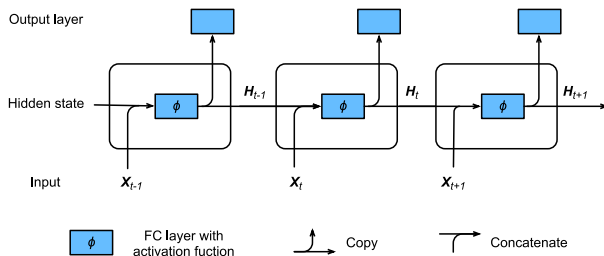
https://d2l.ai/chapter_recurrent-neural-networks/rnn.html#recurrent-networks-with-hidden-states

Recurrent Neural Network:

Process a sequence in order and share the same parameters at each time step.

For a sequence $\mathbf{X}_t$ for $t = 1, 2, 3, \dots, T$

$$\mathbf{H}_t = \phi(\mathbf{X}_t \mathbf{W}_{xh} + \mathbf{H}_{t-1} \mathbf{W}_{hh} + \mathbf{b}_h), \quad \text{output } \mathbf{O}_t = \mathbf{H}_t \mathbf{W}_{hq} + \mathbf{b}_q$$



Q: How does the "concatenate" operator relate to the equation?

A: It's the same. $A \cdot B + C \cdot D = A|C \cdot \frac{B}{D}$ (where hor/ver lines are concatenations), see also last slide of this lecture.

Questions?

Gradients in RNNs

https://d2l.ai/chapter_recurrent-neural-networks/bptt.html

RNNs apply the same (shared weights) function at each time step:

$$z = f_3(\mathbf{x}_3, f_2(\mathbf{x}_2, f_1(\mathbf{x}_1, w), w), w)$$

Gradients in RNNs

https://d2l.ai/chapter_recurrent-neural-networks/bptt.html

RNNs apply the same (shared weights) function at each time step:

$$z = f_3(\mathbf{x}_3, f_2(\mathbf{x}_2, f_1(\mathbf{x}_1, w), w), w)$$

Length 3 recurrence derivative $\frac{dz}{dw}$ includes: $\frac{dz}{dw} = \frac{df_3}{df_2} \frac{df_2}{df_1} \frac{df_1}{dw}$

Gradients in RNNs

https://d2l.ai/chapter_recurrent-neural-networks/bptt.html

RNNs apply the same (shared weights) function at each time step:

$$z = f_3(\mathbf{x}_3, f_2(\mathbf{x}_2, f_1(\mathbf{x}_1, w), w), w)$$

Length 3 recurrence derivative $\frac{dz}{dw}$ includes: $\frac{dz}{dw} = \frac{df_3}{df_2} \frac{df_2}{df_1} \frac{df_1}{dw}$

Q: How about a recurrence of length T ?

A:

Gradients in RNNs

https://d2l.ai/chapter_recurrent-neural-networks/bptt.html

RNNs apply the same (shared weights) function at each time step:

$$z = f_3(\mathbf{x}_3, f_2(\mathbf{x}_2, f_1(\mathbf{x}_1, w), w), w)$$

Length 3 recurrence derivative $\frac{dz}{dw}$ includes:

$$\frac{dz}{dw} = \frac{df_3}{df_2} \frac{df_2}{df_1} \frac{df_1}{dw}$$

Q: How about a recurrence of length T ?

$$\text{A: } \frac{dz}{dw} = \frac{df_1}{dw} \prod_{t=2}^T \frac{df_t}{df_{t-1}}$$

Gradients in RNNs

https://d2l.ai/chapter_recurrent-neural-networks/bptt.html

RNNs apply the same (shared weights) function at each time step:

$$z = f_3(\mathbf{x}_3, f_2(\mathbf{x}_2, f_1(\mathbf{x}_1, w), w), w)$$

Length 3 recurrence derivative $\frac{dz}{dw}$ includes: $\frac{dz}{dw} = \frac{df_3}{df_2} \frac{df_2}{df_1} \frac{df_1}{dw}$

Q: How about a recurrence of length T ? A: $\frac{dz}{dw} = \frac{df_1}{dw} \prod_{t=2}^T \frac{df_t}{df_{t-1}}$

Q: If $|\frac{df_t}{df_{t-1}}| < 1$, what happens for $T = 1000$? A:

Gradients in RNNs

https://d2l.ai/chapter_recurrent-neural-networks/bptt.html

RNNs apply the same (shared weights) function at each time step:

$$z = f_3(\mathbf{x}_3, f_2(\mathbf{x}_2, f_1(\mathbf{x}_1, w), w), w)$$

Length 3 recurrence derivative $\frac{dz}{dw}$ includes: $\frac{dz}{dw} = \frac{df_3}{df_2} \frac{df_2}{df_1} \frac{df_1}{dw}$

Q: How about a recurrence of length T ? A: $\frac{dz}{dw} = \frac{df_1}{dw} \prod_{t=2}^T \frac{df_t}{df_{t-1}}$

Q: If $|\frac{df_t}{df_{t-1}}| < 1$, what happens for $T = 1000$? A: $\frac{dz}{dw} \approx 0$ (vanishes)

Gradients in RNNs

https://d2l.ai/chapter_recurrent-neural-networks/bptt.html

RNNs apply the same (shared weights) function at each time step:

$$z = f_3(\mathbf{x}_3, f_2(\mathbf{x}_2, f_1(\mathbf{x}_1, w), w), w)$$

Length 3 recurrence derivative $\frac{dz}{dw}$ includes: $\frac{dz}{dw} = \frac{df_3}{df_2} \frac{df_2}{df_1} \frac{df_1}{dw}$

Q: How about a recurrence of length T ? A: $\frac{dz}{dw} = \frac{df_1}{dw} \prod_{t=2}^T \frac{df_t}{df_{t-1}}$

Q: If $|\frac{df_t}{df_{t-1}}| < 1$, what happens for $T = 1000$? A: $\frac{dz}{dw} \approx 0$ (vanishes)

Q: If $|\frac{df_t}{df_{t-1}}| > 1$, what happens for $T = 1000$? A:

Gradients in RNNs

https://d2l.ai/chapter_recurrent-neural-networks/bptt.html

RNNs apply the same (shared weights) function at each time step:

$$z = f_3(\mathbf{x}_3, f_2(\mathbf{x}_2, f_1(\mathbf{x}_1, w), w), w)$$

Length 3 recurrence derivative $\frac{dz}{dw}$ includes: $\frac{dz}{dw} = \frac{df_3}{df_2} \frac{df_2}{df_1} \frac{df_1}{dw}$

Q: How about a recurrence of length T ? A: $\frac{dz}{dw} = \frac{df_1}{dw} \prod_{t=2}^T \frac{df_t}{df_{t-1}}$

Q: If $|\frac{df_t}{df_{t-1}}| < 1$, what happens for $T = 1000$? A: $\frac{dz}{dw} \approx 0$ (vanishes)

Q: If $|\frac{df_t}{df_{t-1}}| > 1$, what happens for $T = 1000$? A: $\frac{dz}{dw} \approx \infty$ (explodes)

Gradients in RNNs

https://d2l.ai/chapter_recurrent-neural-networks/bptt.html

RNNs apply the same (shared weights) function at each time step:

$$z = f_3(\mathbf{x}_3, f_2(\mathbf{x}_2, f_1(\mathbf{x}_1, w), w), w)$$

Length 3 recurrence derivative $\frac{dz}{dw}$ includes: $\frac{dz}{dw} = \frac{df_3}{df_2} \frac{df_2}{df_1} \frac{df_1}{dw}$

Q: How about a recurrence of length T ? A: $\frac{dz}{dw} = \frac{df_1}{dw} \prod_{t=2}^T \frac{df_t}{df_{t-1}}$

Q: If $|\frac{df_t}{df_{t-1}}| < 1$, what happens for $T = 1000$? A: $\frac{dz}{dw} \approx 0$ (vanishes)

Q: If $|\frac{df_t}{df_{t-1}}| > 1$, what happens for $T = 1000$? A: $\frac{dz}{dw} \approx \infty$ (explodes)

The student[s], after reading the paper.. book.. blog.. was [were] well prepared.

Gradients in RNNs

https://d2l.ai/chapter_recurrent-neural-networks/bptt.html

RNNs apply the same (shared weights) function at each time step:

$$z = f_3(\mathbf{x}_3, f_2(\mathbf{x}_2, f_1(\mathbf{x}_1, w), w), w)$$

Length 3 recurrence derivative $\frac{dz}{dw}$ includes: $\frac{dz}{dw} = \frac{df_3}{df_2} \frac{df_2}{df_1} \frac{df_1}{dw}$

Q: How about a recurrence of length T ? A: $\frac{dz}{dw} = \frac{df_1}{dw} \prod_{t=2}^T \frac{df_t}{df_{t-1}}$

Q: If $|\frac{df_t}{df_{t-1}}| < 1$, what happens for $T = 1000$? A: $\frac{dz}{dw} \approx 0$ (vanishes)

Q: If $|\frac{df_t}{df_{t-1}}| > 1$, what happens for $T = 1000$? A: $\frac{dz}{dw} \approx \infty$ (explodes)

The student[s], after reading the paper.. book.. blog.. was [were] well prepared.

Q: What is the problem for such a sequence?

Gradients in RNNs

https://d2l.ai/chapter_recurrent-neural-networks/bptt.html

RNNs apply the same (shared weights) function at each time step:

$$z = f_3(\mathbf{x}_3, f_2(\mathbf{x}_2, f_1(\mathbf{x}_1, w), w), w)$$

Length 3 recurrence derivative $\frac{dz}{dw}$ includes: $\frac{dz}{dw} = \frac{df_3}{df_2} \frac{df_2}{df_1} \frac{df_1}{dw}$

Q: How about a recurrence of length T ? A: $\frac{dz}{dw} = \frac{df_1}{dw} \prod_{t=2}^T \frac{df_t}{df_{t-1}}$

Q: If $|\frac{df_t}{df_{t-1}}| < 1$, what happens for $T = 1000$? A: $\frac{dz}{dw} \approx 0$ (vanishes)

Q: If $|\frac{df_t}{df_{t-1}}| > 1$, what happens for $T = 1000$? A: $\frac{dz}{dw} \approx \infty$ (explodes)

The student[s], after reading the paper.. book.. blog.. was [were] well prepared.

Q: What is the problem for such a sequence?

A: RNNs have difficulty with backpropping long-range relations

GRU: Gated Recurrent Unit

https://d2l.ai/chapter_recurrent-modern/gru.html

Improve RNNs hidden state $\mathbf{H}_t$ for long-range relationship learning.

GRU: Gated Recurrent Unit

https://d2l.ai/chapter_recurrent-modern/gru.html

Improve RNNs hidden state $\mathbf{H}_t$ for long-range relationship learning.

- Reset for different logical chunks in input (e.g.: chapters in a book)

GRU: Gated Recurrent Unit

https://d2l.ai/chapter_recurrent-modern/gru.html

Improve RNNs hidden state $\mathbf{H}_t$ for long-range relationship learning.

- Reset for different logical chunks in input (e.g.: chapters in a book)
- Do not update for uninformative input (e.g.: html codings)

GRU: Gated Recurrent Unit

https://d2l.ai/chapter_recurrent-modern/gru.html

Improve RNNs hidden state $\mathbf{H}_t$ for long-range relationship learning.

- Reset for different logical chunks in input (e.g.: chapters in a book)
- Do not update for uninformative input (e.g.: html codings)

Gate: Vector with entries in $(0, 1)$ Inspired by logical gates

GRU: Gated Recurrent Unit

https://d2l.ai/chapter_recurrent-modern/gru.html

Improve RNNs hidden state $\mathbf{H}_t$ for long-range relationship learning.

- Reset for different logical chunks in input (e.g.: chapters in a book)
- Do not update for uninformative input (e.g.: html codings)

Gate: Vector with entries in $(0, 1)$ Inspired by logical gates

Reset gate $\mathbf{R}_t$:

- Q: How can a gate $\mathbf{R}_t$ cause to ignore previous memories $\mathbf{H}_{t-1}$?

GRU: Gated Recurrent Unit

https://d2l.ai/chapter_recurrent-modern/gru.html

Improve RNNs hidden state $\mathbf{H}_t$ for long-range relationship learning.

- Reset for different logical chunks in input (e.g.: chapters in a book)
- Do not update for uninformative input (e.g.: html codings)

Gate: Vector with entries in $(0, 1)$ Inspired by logical gates

Reset gate $\mathbf{R}_t$:

- Q: How can a gate $\mathbf{R}_t$ cause to ignore previous memories $\mathbf{H}_{t-1}$?
- A: Multiply the gate $\mathbf{R}_t$ element-wise $\odot$ with $\mathbf{H}_{t-1}$.

GRU: Gated Recurrent Unit

https://d2l.ai/chapter_recurrent-modern/gru.html

Improve RNNs hidden state $\mathbf{H}_t$ for long-range relationship learning.

- Reset for different logical chunks in input (e.g.: chapters in a book)
- Do not update for uninformative input (e.g.: html codings)

Gate: Vector with entries in $(0, 1)$ Inspired by logical gates

Reset gate $\mathbf{R}_t$:

- Q: How can a gate $\mathbf{R}_t$ cause to ignore previous memories $\mathbf{H}_{t-1}$?
- A: Multiply the gate $\mathbf{R}_t$ element-wise $\odot$ with $\mathbf{H}_{t-1}$.

Update gate $\mathbf{Z}_t$:

- Q: For uninformative input $\mathbf{X}_t$, how can a gate $\mathbf{Z}_t$ weigh between:
(the previous hidden state $\mathbf{H}_{t-1}$) vs (an already updated hidden state $\tilde{\mathbf{H}}_t$)?

GRU: Gated Recurrent Unit

https://d2l.ai/chapter_recurrent-modern/gru.html

Improve RNNs hidden state $\mathbf{H}_t$ for long-range relationship learning.

- Reset for different logical chunks in input (e.g.: chapters in a book)
- Do not update for uninformative input (e.g.: html codings)

Gate: Vector with entries in $(0, 1)$ Inspired by logical gates

Reset gate $\mathbf{R}_t$:

- Q: How can a gate $\mathbf{R}_t$ cause to ignore previous memories $\mathbf{H}_{t-1}$?
- A: Multiply the gate $\mathbf{R}_t$ element-wise $\odot$ with $\mathbf{H}_{t-1}$.

Update gate $\mathbf{Z}_t$:

- Q: For uninformative input $\mathbf{X}_t$, how can a gate $\mathbf{Z}_t$ weigh between: (the previous hidden state $\mathbf{H}_{t-1}$) vs (an already updated hidden state $\tilde{\mathbf{H}}_t$)?
- A: $\mathbf{H}_{t-1} \odot \mathbf{Z}_t + (1 - \mathbf{Z}_t) \odot \tilde{\mathbf{H}}_t$, where $\tilde{\mathbf{H}}_t$ is the already updated state.

GRU: Gated Recurrent Unit

https://d2l.ai/chapter_recurrent-modern/gru.html

Improve RNNs hidden state $\mathbf{H}_t$ for long-range relationship learning.

- Reset for different logical chunks in input (e.g.: chapters in a book)
- Do not update for uninformative input (e.g.: html codings)

Gate: Vector with entries in $(0, 1)$ Inspired by logical gates

Reset gate $\mathbf{R}_t$:

- Q: How can a gate $\mathbf{R}_t$ cause to ignore previous memories $\mathbf{H}_{t-1}$?
- A: Multiply the gate $\mathbf{R}_t$ element-wise $\odot$ with $\mathbf{H}_{t-1}$.

Update gate $\mathbf{Z}_t$:

- Q: For uninformative input $\mathbf{X}_t$, how can a gate $\mathbf{Z}_t$ weigh between: (the previous hidden state $\mathbf{H}_{t-1}$) vs (an already updated hidden state $\tilde{\mathbf{H}}_t$)?
 - A: $\mathbf{H}_{t-1} \odot \mathbf{Z}_t + (1 - \mathbf{Z}_t) \odot \tilde{\mathbf{H}}_t$, where $\tilde{\mathbf{H}}_t$ is the already updated state.
- Q: What happens for a $\mathbf{Z}_t$ as a vector of ones: $(1, 1, \dots, 1)^\top$?

GRU: Gated Recurrent Unit

https://d2l.ai/chapter_recurrent-modern/gru.html

Improve RNNs hidden state $\mathbf{H}_t$ for long-range relationship learning.

- Reset for different logical chunks in input (e.g.: chapters in a book)
- Do not update for uninformative input (e.g.: html codings)

Gate: Vector with entries in $(0, 1)$ Inspired by logical gates

Reset gate $\mathbf{R}_t$:

- Q: How can a gate $\mathbf{R}_t$ cause to ignore previous memories $\mathbf{H}_{t-1}$?
- A: Multiply the gate $\mathbf{R}_t$ element-wise $\odot$ with $\mathbf{H}_{t-1}$.

Update gate $\mathbf{Z}_t$:

- Q: For uninformative input $\mathbf{X}_t$, how can a gate $\mathbf{Z}_t$ weigh between: (the previous hidden state $\mathbf{H}_{t-1}$) vs (an already updated hidden state $\tilde{\mathbf{H}}_t$)?
- A: $\mathbf{H}_{t-1} \odot \mathbf{Z}_t + (1 - \mathbf{Z}_t) \odot \tilde{\mathbf{H}}_t$, where $\tilde{\mathbf{H}}_t$ is the already updated state.
Q: What happens for a $\mathbf{Z}_t$ as a vector of ones: $(1, 1, \dots, 1)^\top$?
A: Input $\mathbf{X}_t$ has no effect: only previous state $\mathbf{H}_{t-1}$ used (ie: no update).

GRU: Gated Recurrent Unit

https://d2l.ai/chapter_recurrent-modern/gru.html

Gates allow to selectively keep or remove dimensions from the hidden state.

GRU: Gated Recurrent Unit

https://d2l.ai/chapter_recurrent-modern/gru.html

Gates allow to selectively keep or remove dimensions from the hidden state.

Q: How to force gates $\mathbf{R}_t$ and $\mathbf{Z}_t$ to be vectors in $(0, 1)$?

GRU: Gated Recurrent Unit

https://d2l.ai/chapter_recurrent-modern/gru.html

Gates allow to selectively keep or remove dimensions from the hidden state.

Q: How to force gates $\mathbf{R}_t$ and $\mathbf{Z}_t$ to be vectors in $(0, 1)$?

A: Use sigmoid activation σ .

GRU: Gated Recurrent Unit

https://d2l.ai/chapter_recurrent-modern/gru.html

Gates allow to selectively keep or remove dimensions from the hidden state.

Q: How to force gates $\mathbf{R}_t$ and $\mathbf{Z}_t$ to be vectors in $(0, 1)$?

A: Use sigmoid activation σ .

Q: Based on what information would you learn the gates? (what info is available?)

GRU: Gated Recurrent Unit

https://d2l.ai/chapter_recurrent-modern/gru.html

Gates allow to selectively keep or remove dimensions from the hidden state.

Q: How to force gates $\mathbf{R}_t$ and $\mathbf{Z}_t$ to be vectors in $(0, 1)$?

A: Use sigmoid activation σ .

Q: Based on what information would you learn the gates? (what info is available?)

A: The current input $\mathbf{X}_t$ and previous state $\mathbf{H}_{t-1}$:

GRU: Gated Recurrent Unit

https://d2l.ai/chapter_recurrent-modern/gru.html

Gates allow to selectively keep or remove dimensions from the hidden state.

Q: How to force gates $\mathbf{R}_t$ and $\mathbf{Z}_t$ to be vectors in $(0, 1)$?

A: Use sigmoid activation σ .

Q: Based on what information would you learn the gates? (what info is available?)

A: The current input $\mathbf{X}_t$ and previous state $\mathbf{H}_{t-1}$:

$$\mathbf{R}_t = \sigma(\mathbf{X}_t \mathbf{W}_{xr} + \mathbf{H}_{t-1} \mathbf{W}_{hr} + \mathbf{b}_r)$$

$$\mathbf{Z}_t = \sigma(\mathbf{X}_t \mathbf{W}_{xz} + \mathbf{H}_{t-1} \mathbf{W}_{hz} + \mathbf{b}_z)$$

GRU: Gated Recurrent Unit

https://d2l.ai/chapter_recurrent-modern/gru.html

Gates allow to selectively keep or remove dimensions from the hidden state.

Q: How to force gates $\mathbf{R}_t$ and $\mathbf{Z}_t$ to be vectors in $(0, 1)$?

A: Use sigmoid activation σ .

Q: Based on what information would you learn the gates? (what info is available?)

A: The current input $\mathbf{X}_t$ and previous state $\mathbf{H}_{t-1}$:

$$\mathbf{R}_t = \sigma(\mathbf{X}_t \mathbf{W}_{xr} + \mathbf{H}_{t-1} \mathbf{W}_{hr} + \mathbf{b}_r)$$

$$\mathbf{Z}_t = \sigma(\mathbf{X}_t \mathbf{W}_{xz} + \mathbf{H}_{t-1} \mathbf{W}_{hz} + \mathbf{b}_z)$$

Q: Looks familiar?

GRU: Gated Recurrent Unit

https://d2l.ai/chapter_recurrent-modern/gru.html

Gates allow to selectively keep or remove dimensions from the hidden state.

Q: How to force gates $\mathbf{R}_t$ and $\mathbf{Z}_t$ to be vectors in $(0, 1)$?

A: Use sigmoid activation σ .

Q: Based on what information would you learn the gates? (what info is available?)

A: The current input $\mathbf{X}_t$ and previous state $\mathbf{H}_{t-1}$:

$$\mathbf{R}_t = \sigma(\mathbf{X}_t \mathbf{W}_{xr} + \mathbf{H}_{t-1} \mathbf{W}_{hr} + \mathbf{b}_r)$$

$$\mathbf{Z}_t = \sigma(\mathbf{X}_t \mathbf{W}_{xz} + \mathbf{H}_{t-1} \mathbf{W}_{hz} + \mathbf{b}_z)$$

Q: Looks familiar? A: Standard RNN: $\mathbf{H}_t = \phi(\mathbf{X}_t \mathbf{W}_{xh} + \mathbf{H}_{t-1} \mathbf{W}_{hh} + \mathbf{b}_h)$

GRU: Gated Recurrent Unit

https://d2l.ai/chapter_recurrent-modern/gru.html

Gates allow to selectively keep or remove dimensions from the hidden state.

Q: How to force gates $\mathbf{R}_t$ and $\mathbf{Z}_t$ to be vectors in $(0, 1)$?

A: Use sigmoid activation σ .

Q: Based on what information would you learn the gates? (what info is available?)

A: The current input $\mathbf{X}_t$ and previous state $\mathbf{H}_{t-1}$:

$$\mathbf{R}_t = \sigma(\mathbf{X}_t \mathbf{W}_{xr} + \mathbf{H}_{t-1} \mathbf{W}_{hr} + \mathbf{b}_r)$$

$$\mathbf{Z}_t = \sigma(\mathbf{X}_t \mathbf{W}_{xz} + \mathbf{H}_{t-1} \mathbf{W}_{hz} + \mathbf{b}_z)$$

Q: Looks familiar? A: Standard RNN: $\mathbf{H}_t = \phi(\mathbf{X}_t \mathbf{W}_{xh} + \mathbf{H}_{t-1} \mathbf{W}_{hh} + \mathbf{b}_h)$

How a GRU differs from a standard RNN:

GRU adds a reset gate:

- Candidate state $\tilde{\mathbf{H}}_t = \tanh(\mathbf{X}_t \mathbf{W}_{xh} + (\mathbf{R}_t \odot \mathbf{H}_{t-1}) \mathbf{W}_{hh} + \mathbf{b}_h)$

GRU: Gated Recurrent Unit

https://d2l.ai/chapter_recurrent-modern/gru.html

Gates allow to selectively keep or remove dimensions from the hidden state.

Q: How to force gates $\mathbf{R}_t$ and $\mathbf{Z}_t$ to be vectors in $(0, 1)$?

A: Use sigmoid activation σ .

Q: Based on what information would you learn the gates? (what info is available?)

A: The current input $\mathbf{X}_t$ and previous state $\mathbf{H}_{t-1}$:

$$\mathbf{R}_t = \sigma(\mathbf{X}_t \mathbf{W}_{xr} + \mathbf{H}_{t-1} \mathbf{W}_{hr} + \mathbf{b}_r)$$

$$\mathbf{Z}_t = \sigma(\mathbf{X}_t \mathbf{W}_{xz} + \mathbf{H}_{t-1} \mathbf{W}_{hz} + \mathbf{b}_z)$$

Q: Looks familiar? A: Standard RNN: $\mathbf{H}_t = \phi(\mathbf{X}_t \mathbf{W}_{xh} + \mathbf{H}_{t-1} \mathbf{W}_{hh} + \mathbf{b}_h)$

How a GRU differs from a standard RNN:

GRU adds a reset gate:

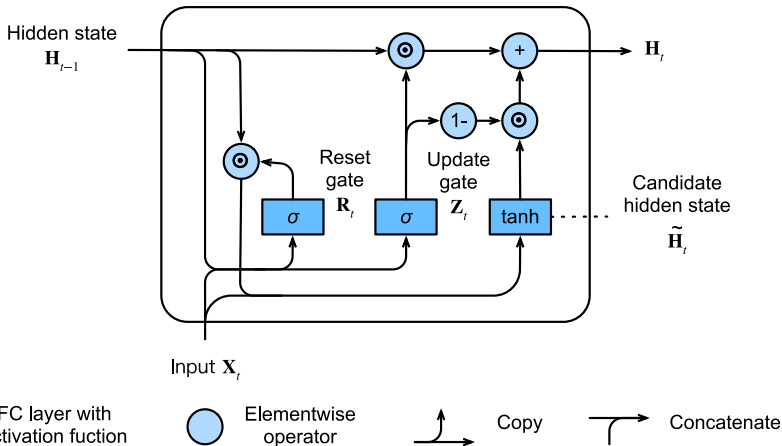
- Candidate state $\tilde{\mathbf{H}}_t = \tanh(\mathbf{X}_t \mathbf{W}_{xh} + (\mathbf{R}_t \odot \mathbf{H}_{t-1}) \mathbf{W}_{hh} + \mathbf{b}_h)$

GRU adds an update gate:

- Current state $\mathbf{H}_t = \mathbf{H}_{t-1} \odot \mathbf{Z}_t + (1 - \mathbf{Z}_t) \odot \tilde{\mathbf{H}}_t$, where $\tilde{\mathbf{H}}_t$ is candidate state.

GRU: Gated Recurrent Unit

https://d2l.ai/chapter_recurrent-modern/gru.html



Questions?

Questions?

Q: Why does gating allow for long-term relations? (ie: what about vanishing/exploding gradients?)

Questions?

Q: Why does gating allow for long-term relations? (ie: what about vanishing/exploding gradients?)

A: Relevant long term information can pass through unchanged.

LSTM: Long Short Term Memory networks

https://d2l.ai/chapter_recurrent-modern/lstm.html

See also: <https://colah.github.io/posts/2015-08-Understanding-LSTMs>

Input gate, Forget gate, Output gate, and an additional 'memory cell' state C_t

LSTM: Long Short Term Memory networks

https://d2l.ai/chapter_recurrent-modern/lstm.html

See also: <https://colah.github.io/posts/2015-08-Understanding-LSTMs>

Input gate, Forget gate, Output gate, and an additional 'memory cell' state C_t

- LSTM linearly adds/removes information to 'conveyor belt' cell C_t

LSTM: Long Short Term Memory networks

https://d2l.ai/chapter_recurrent-modern/lstm.html

See also: <https://colah.github.io/posts/2015-08-Understanding-LSTMs>

Input gate, Forget gate, Output gate, and an additional 'memory cell' state C_t

- LSTM linearly adds/removes information to 'conveyor belt' cell C_t
- Gates and candidate Cell depend on input and previous hidden state.

LSTM: Long Short Term Memory networks

https://d2l.ai/chapter_recurrent-modern/lstm.html

See also: <https://colah.github.io/posts/2015-08-Understanding-LSTMs>

Input gate, Forget gate, Output gate, and an additional 'memory cell' state C_t

- LSTM linearly adds/removes information to 'conveyor belt' cell C_t
- Gates and candidate Cell depend on input and previous hidden state.
- Current Memory cell depends on previous Cell and candidate Cell

LSTM: Long Short Term Memory networks

https://d2l.ai/chapter_recurrent-modern/lstm.html

See also: <https://colah.github.io/posts/2015-08-Understanding-LSTMs>

Input gate, Forget gate, Output gate, and an additional 'memory cell' state C_t

- LSTM linearly adds/removes information to 'conveyor belt' cell C_t
- Gates and candidate Cell depend on input and previous hidden state.
- Current Memory cell depends on previous Cell and candidate Cell
- Current hidden state depends on current Cell

LSTM: Long Short Term Memory networks

https://d2l.ai/chapter_recurrent-modern/lstm.html

See also: <https://colah.github.io/posts/2015-08-Understanding-LSTMs>

Input gate, Forget gate, Output gate, and an additional 'memory cell' state C_t

- LSTM linearly adds/removes information to 'conveyor belt' cell C_t
- Gates and candidate Cell depend on input and previous hidden state.
- Current Memory cell depends on previous Cell and candidate Cell
- Current hidden state depends on current Cell

Gates computed like in a GRU:

LSTM: Long Short Term Memory networks

https://d2l.ai/chapter_recurrent-modern/lstm.html

See also: <https://colah.github.io/posts/2015-08-Understanding-LSTMs>

Input gate, Forget gate, Output gate, and an additional 'memory cell' state C_t

- LSTM linearly adds/removes information to 'conveyor belt' cell C_t
- Gates and candidate Cell depend on input and previous hidden state.
- Current Memory cell depends on previous Cell and candidate Cell
- Current hidden state depends on current Cell

Gates computed like in a GRU:

Input gate	$I_t = \sigma(X_t W_{xi} + H_{t-1} W_{hi} + b_i)$
Forget gate	$F_t = \sigma(X_t W_{xf} + H_{t-1} W_{hf} + b_f)$
Output gate	$O_t = \sigma(X_t W_{xo} + H_{t-1} W_{ho} + b_o)$

LSTM: Long Short Term Memory networks

https://d2l.ai/chapter_recurrent-modern/lstm.html

See also: <https://colah.github.io/posts/2015-08-Understanding-LSTMs>

Input gate, Forget gate, Output gate, and an additional 'memory cell' state C_t

- LSTM linearly adds/removes information to 'conveyor belt' cell C_t
- Gates and candidate Cell depend on input and previous hidden state.
- Current Memory cell depends on previous Cell and candidate Cell
- Current hidden state depends on current Cell

Gates computed like in a GRU:

$$\begin{array}{ll}\text{Input gate} & \mathbf{I}_t = \sigma(\mathbf{X}_t \mathbf{W}_{xi} + \mathbf{H}_{t-1} \mathbf{W}_{hi} + \mathbf{b}_i) \\ \text{Forget gate} & \mathbf{F}_t = \sigma(\mathbf{X}_t \mathbf{W}_{xf} + \mathbf{H}_{t-1} \mathbf{W}_{hf} + \mathbf{b}_f) \\ \text{Output gate} & \mathbf{O}_t = \sigma(\mathbf{X}_t \mathbf{W}_{xo} + \mathbf{H}_{t-1} \mathbf{W}_{ho} + \mathbf{b}_o)\end{array}$$

$$\text{Candidate memory cell } \tilde{C}_t = \tanh(\mathbf{X}_t \mathbf{W}_{xc} + \mathbf{H}_{t-1} \mathbf{W}_{hc} + \mathbf{b}_c)$$

LSTM: Long Short Term Memory networks

https://d2l.ai/chapter_recurrent-modern/lstm.html

See also: <https://colah.github.io/posts/2015-08-Understanding-LSTMs>

Input gate, Forget gate, Output gate, and an additional 'memory cell' state C_t

- LSTM linearly adds/removes information to 'conveyor belt' cell C_t
- Gates and candidate Cell depend on input and previous hidden state.
- Current Memory cell depends on previous Cell and candidate Cell
- Current hidden state depends on current Cell

Gates computed like in a GRU:

$$\begin{array}{ll}\text{Input gate} & \mathbf{I}_t = \sigma(\mathbf{X}_t \mathbf{W}_{xi} + \mathbf{H}_{t-1} \mathbf{W}_{hi} + \mathbf{b}_i) \\ \text{Forget gate} & \mathbf{F}_t = \sigma(\mathbf{X}_t \mathbf{W}_{xf} + \mathbf{H}_{t-1} \mathbf{W}_{hf} + \mathbf{b}_f) \\ \text{Output gate} & \mathbf{O}_t = \sigma(\mathbf{X}_t \mathbf{W}_{xo} + \mathbf{H}_{t-1} \mathbf{W}_{ho} + \mathbf{b}_o)\end{array}$$

Candidate memory cell $\tilde{C}_t = \tanh(\mathbf{X}_t \mathbf{W}_{xc} + \mathbf{H}_{t-1} \mathbf{W}_{hc} + \mathbf{b}_c)$

Where GRU has a single update gate: $\mathbf{H}_t = \mathbf{H}_{t-1} \odot \mathbf{Z}_t + (1 - \mathbf{Z}_t) \odot \tilde{\mathbf{H}}_t$

LSTM has 2 gates: Forget gate $\mathbf{F}_t$; input gate $\mathbf{I}_t$: $\mathbf{C}_t = \mathbf{F}_t \odot \mathbf{C}_{t-1} + \mathbf{I}_t \odot \tilde{C}_t$

LSTM: Long Short Term Memory networks

https://d2l.ai/chapter_recurrent-modern/lstm.html

See also: <https://colah.github.io/posts/2015-08-Understanding-LSTMs>

Input gate, Forget gate, Output gate, and an additional 'memory cell' state C_t

- LSTM linearly adds/removes information to 'conveyor belt' cell C_t
- Gates and candidate Cell depend on input and previous hidden state.
- Current Memory cell depends on previous Cell and candidate Cell
- Current hidden state depends on current Cell

Gates computed like in a GRU:

$$\begin{array}{ll}\text{Input gate} & \mathbf{I}_t = \sigma(\mathbf{X}_t \mathbf{W}_{xi} + \mathbf{H}_{t-1} \mathbf{W}_{hi} + \mathbf{b}_i) \\ \text{Forget gate} & \mathbf{F}_t = \sigma(\mathbf{X}_t \mathbf{W}_{xf} + \mathbf{H}_{t-1} \mathbf{W}_{hf} + \mathbf{b}_f) \\ \text{Output gate} & \mathbf{O}_t = \sigma(\mathbf{X}_t \mathbf{W}_{xo} + \mathbf{H}_{t-1} \mathbf{W}_{ho} + \mathbf{b}_o)\end{array}$$

Candidate memory cell $\tilde{C}_t = \tanh(\mathbf{X}_t \mathbf{W}_{xc} + \mathbf{H}_{t-1} \mathbf{W}_{hc} + \mathbf{b}_c)$

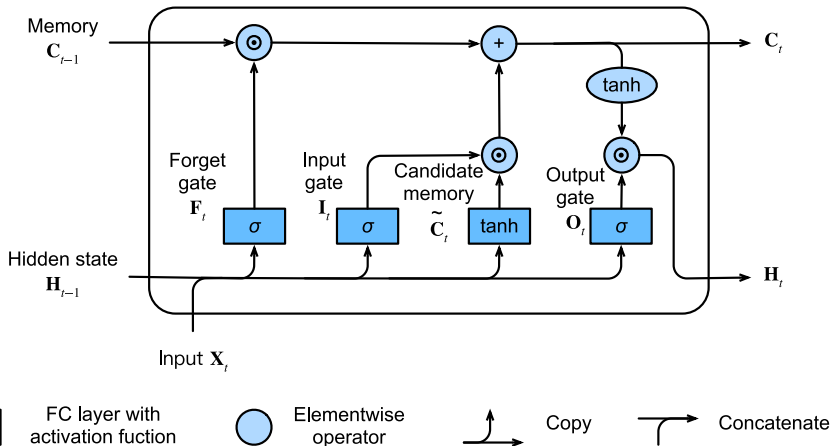
Where GRU has a single update gate: $\mathbf{H}_t = \mathbf{H}_{t-1} \odot \mathbf{Z}_t + (1 - \mathbf{Z}_t) \odot \tilde{\mathbf{H}}_t$

LSTM has 2 gates: Forget gate $\mathbf{F}_t$; input gate $\mathbf{I}_t$: $\mathbf{C}_t = \mathbf{F}_t \odot \mathbf{C}_{t-1} + \mathbf{I}_t \odot \tilde{C}_t$

Current hidden state uses output gate: $\mathbf{H}_t = \mathbf{O}_t \odot \tanh(\mathbf{C}_t)$

LSTM: Long Short Term Memory

D2L-book chapter 9.2



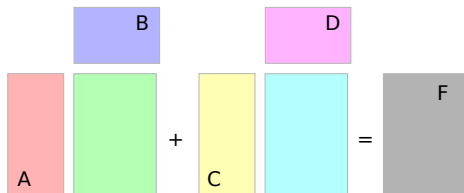
Questions?

Summary

- RNNs
- Vanishing/exploding gradients in RNNs
- GRU and LSTM

Appendix: Concatenating matrices

Adding



Concatenation

