

Recurrent Neural Networks

Milan Straka

 April 14, 2020



EUROPEAN UNION
European Structural and Investment Fund
Operational Programme Research,
Development and Education

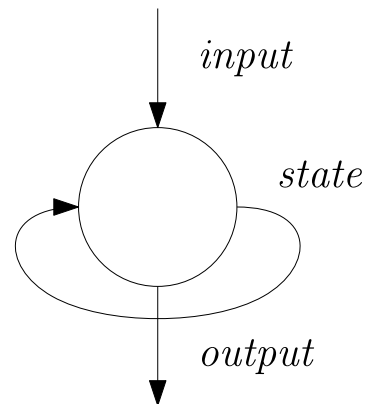
Charles University in Prague
Faculty of Mathematics and Physics
Institute of Formal and Applied Linguistics



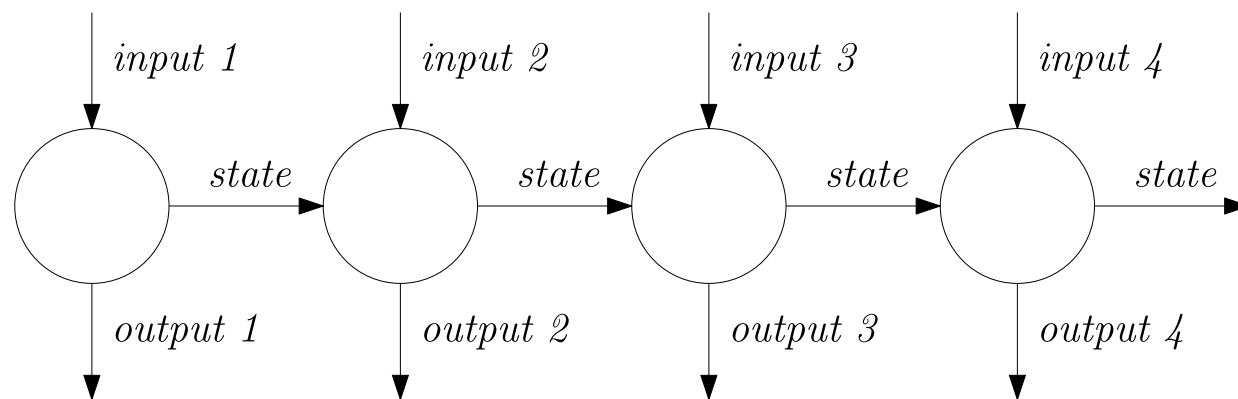
unless otherwise stated

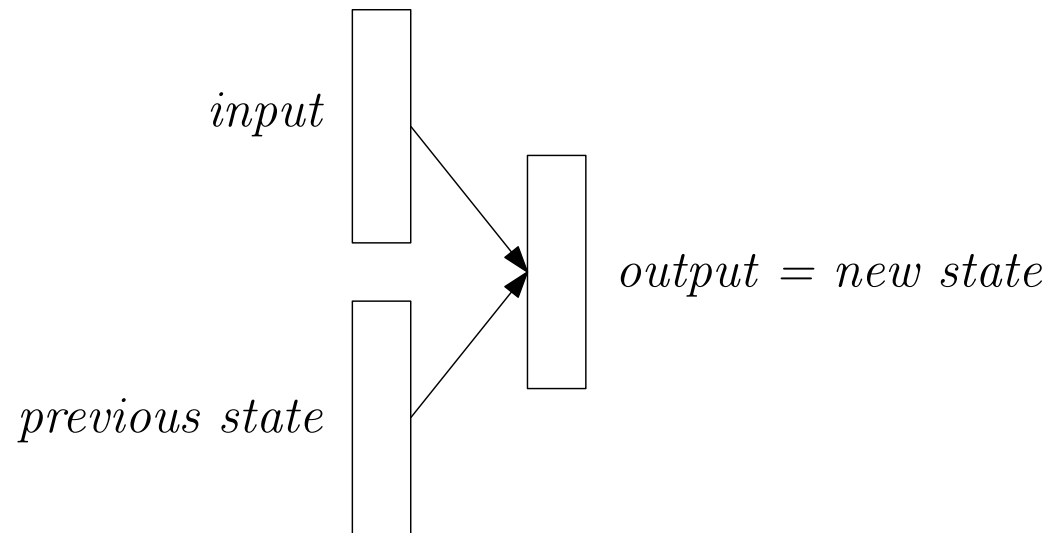
Recurrent Neural Networks

Single RNN cell



Unrolled RNN cells





Given an input $\mathbf{x}^{(t)}$ and previous state $\mathbf{s}^{(t-1)}$, the new state is computed as

$$\mathbf{s}^{(t)} = f(\mathbf{s}^{(t-1)}, \mathbf{x}^{(t)}; \boldsymbol{\theta}).$$

One of the simplest possibilities (called SimpleRNN in TensorFlow) is

$$\mathbf{s}^{(t)} = \tanh(\mathbf{U}\mathbf{s}^{(t-1)} + \mathbf{V}\mathbf{x}^{(t)} + \mathbf{b}).$$

Basic RNN cells suffer a lot from vanishing/exploding gradients (*the challenge of long-term dependencies*).

If we simplify the recurrence of states to

$$\mathbf{s}^{(t)} = \mathbf{U} \mathbf{s}^{(t-1)},$$

we get

$$\mathbf{s}^{(t)} = \mathbf{U}^t \mathbf{s}^{(0)}.$$

If $\mathbf{U}$ has eigenvalue decomposition of $\mathbf{U} = \mathbf{Q} \mathbf{\Lambda} \mathbf{Q}^{-1}$, we get

$$\mathbf{s}^{(t)} = \mathbf{Q} \mathbf{\Lambda}^t \mathbf{Q}^{-1} \mathbf{s}^{(0)}.$$

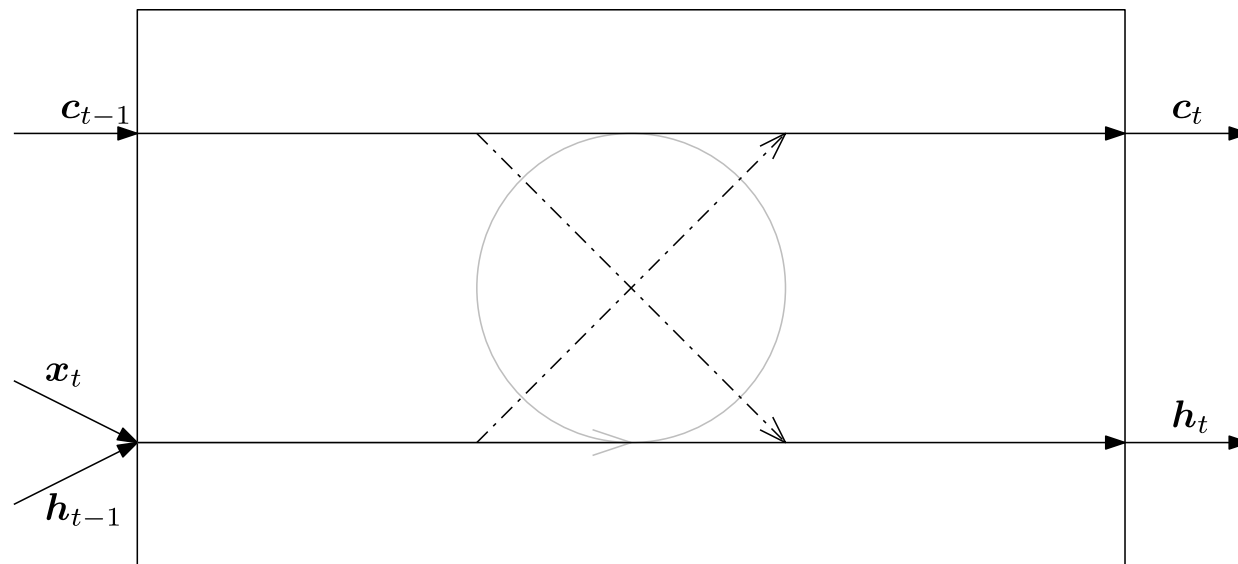
The main problem is that the *same* function is iteratively applied many times.

Several more complex RNN cell variants have been proposed, which alleviate this issue to some degree, namely **LSTM** and **GRU**.

Hochreiter & Schmidhuber (1997) suggested that to enforce *constant error flow*, we would like

$$f' = 1.$$

They propose to achieve that by a *constant error carrousel*.



Long Short-Term Memory

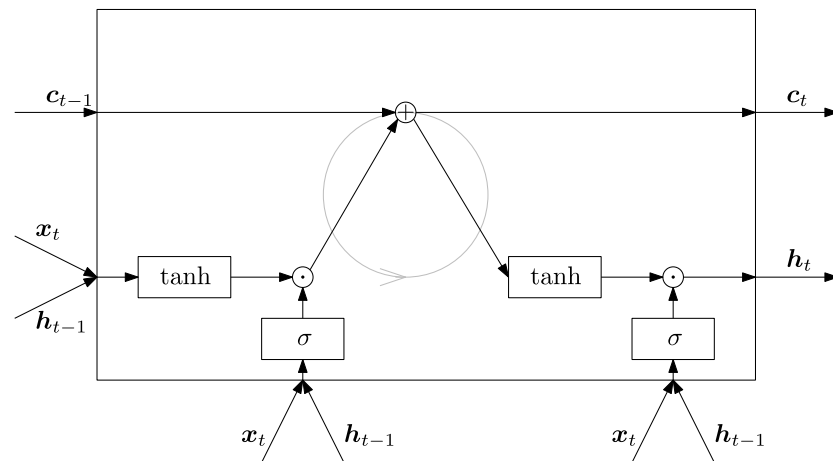
They also propose an *input* and *output* gates which control the flow of information into and out of the carrousel (*memory cell* $\mathbf{c}_t$).

$$\mathbf{i}_t \leftarrow \sigma(\mathbf{W}^i \mathbf{x}_t + \mathbf{V}^i \mathbf{h}_{t-1} + \mathbf{b}^i)$$

$$\mathbf{o}_t \leftarrow \sigma(\mathbf{W}^o \mathbf{x}_t + \mathbf{V}^o \mathbf{h}_{t-1} + \mathbf{b}^o)$$

$$\mathbf{c}_t \leftarrow \mathbf{c}_{t-1} + \mathbf{i}_t \cdot \tanh(\mathbf{W}^y \mathbf{x}_t + \mathbf{V}^y \mathbf{h}_{t-1} + \mathbf{b}^y)$$

$$\mathbf{h}_t \leftarrow \mathbf{o}_t \cdot \tanh(\mathbf{c}_t)$$



Long Short-Term Memory

Later in Gers, Schmidhuber & Cummins (1999) a possibility to *forget* information from memory cell c_t was added.

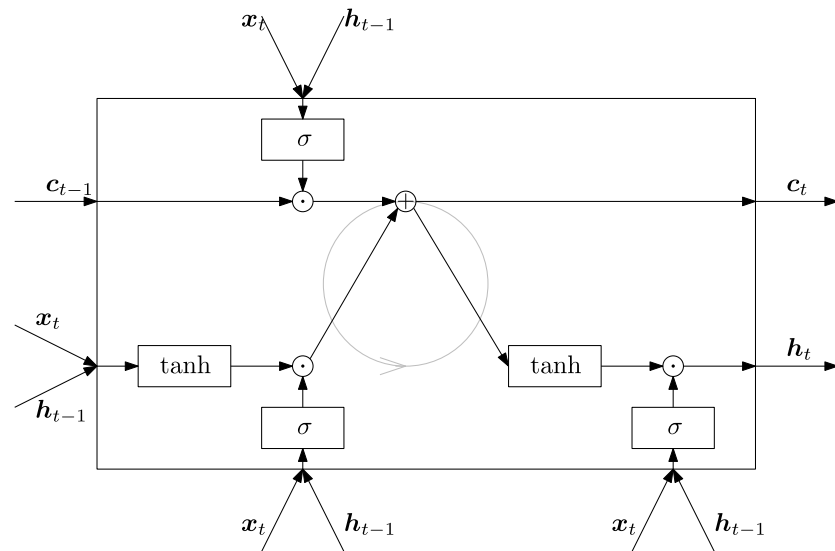
$$i_t \leftarrow \sigma(W^i x_t + V^i h_{t-1} + b^i)$$

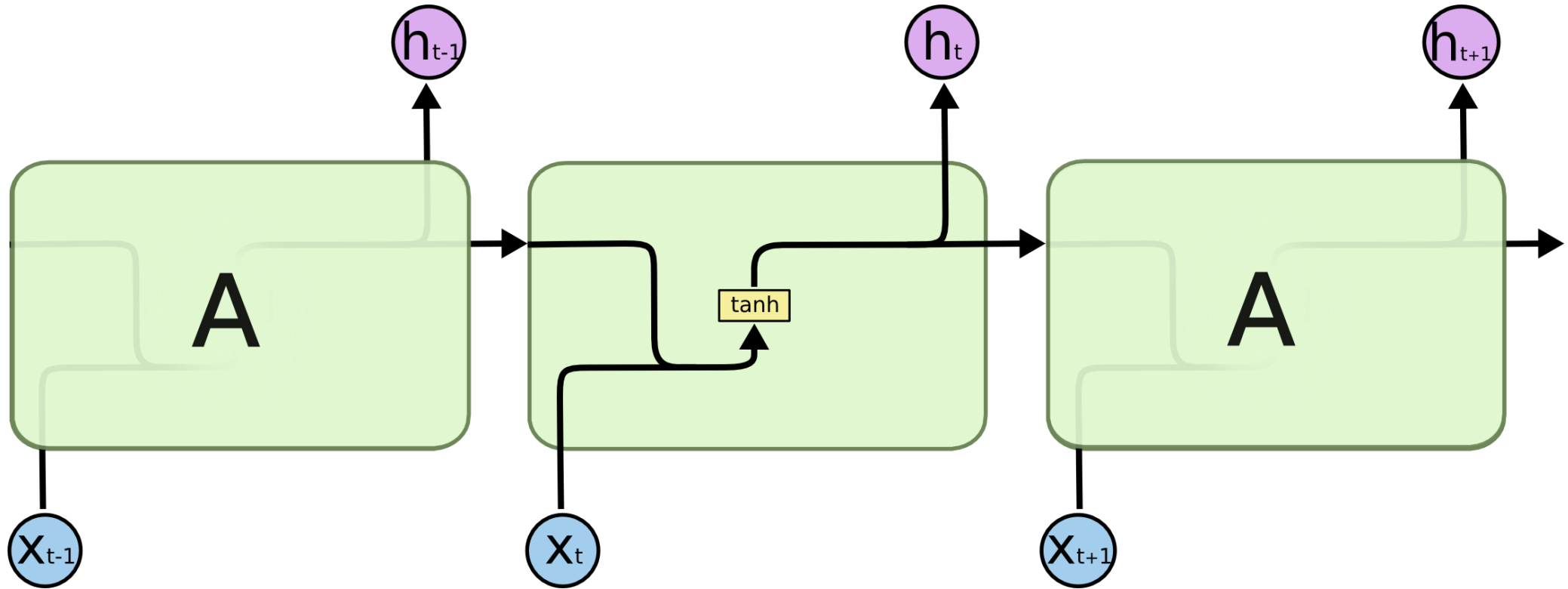
$$f_t \leftarrow \sigma(W^f x_t + V^f h_{t-1} + b^f)$$

$$o_t \leftarrow \sigma(W^o x_t + V^o h_{t-1} + b^o)$$

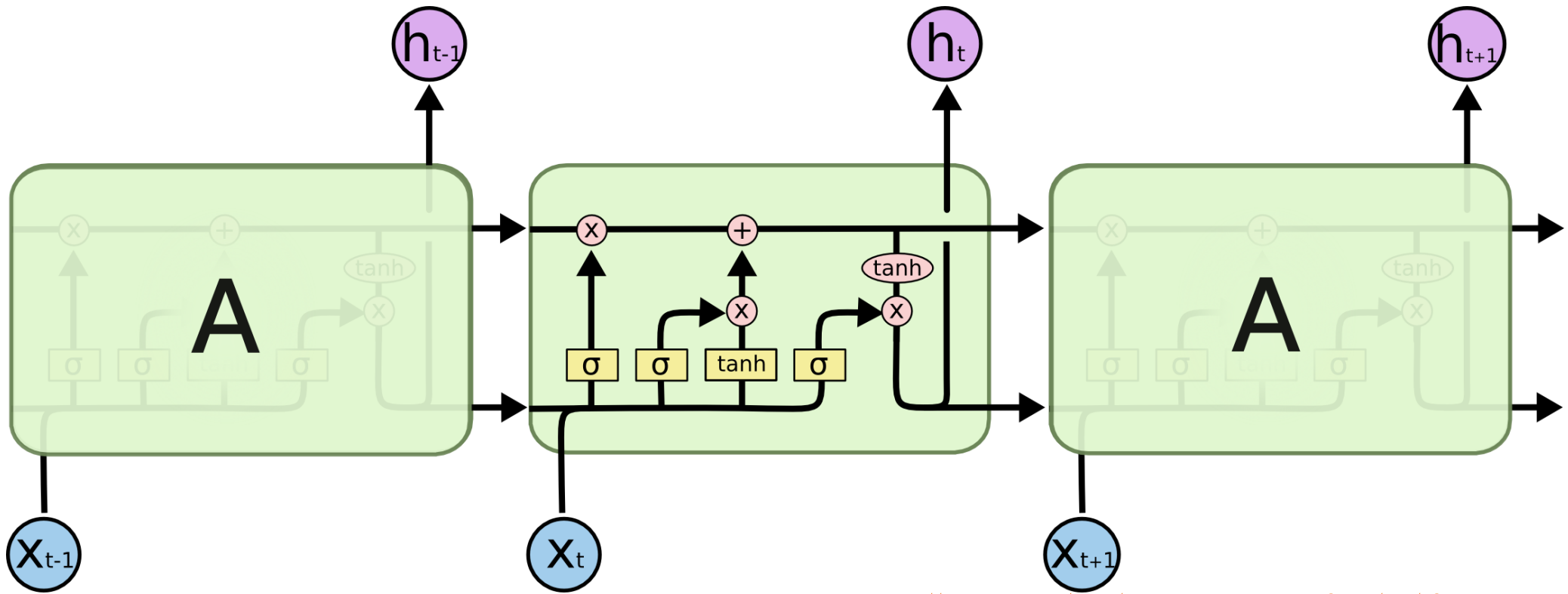
$$c_t \leftarrow f_t \cdot c_{t-1} + i_t \cdot \tanh(W^y x_t + V^y h_{t-1} + b^y)$$

$$h_t \leftarrow o_t \cdot \tanh(c_t)$$

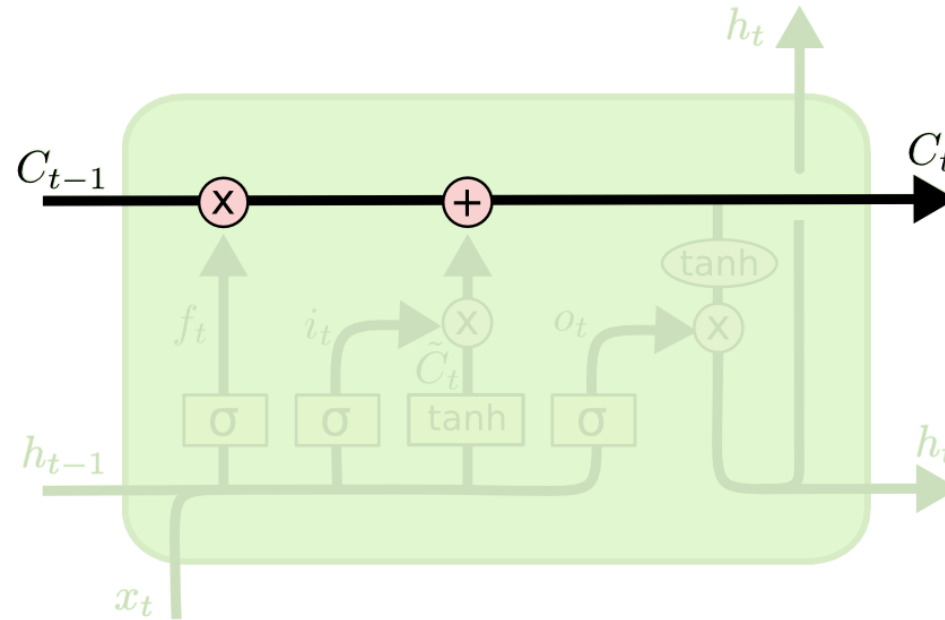




<http://colah.github.io/posts/2015-08-Understanding-LSTMs/img/LSTM3-SimpleRNN.png>

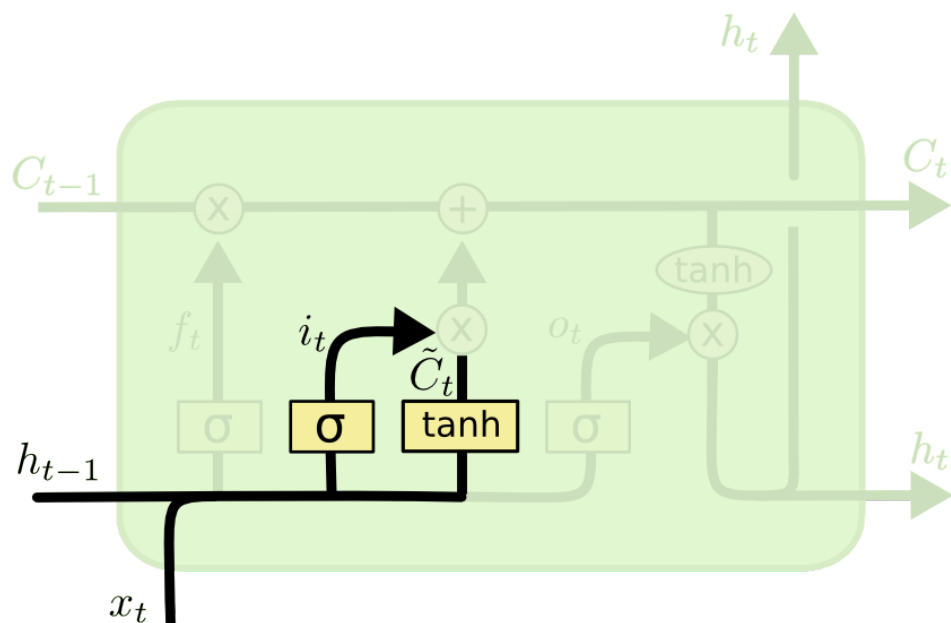


<http://colah.github.io/posts/2015-08-Understanding-LSTMs/img/LSTM3-chain.png>



<http://colah.github.io/posts/2015-08-Understanding-LSTMs/img/LSTM3-C-line.png>

Long Short-Term Memory

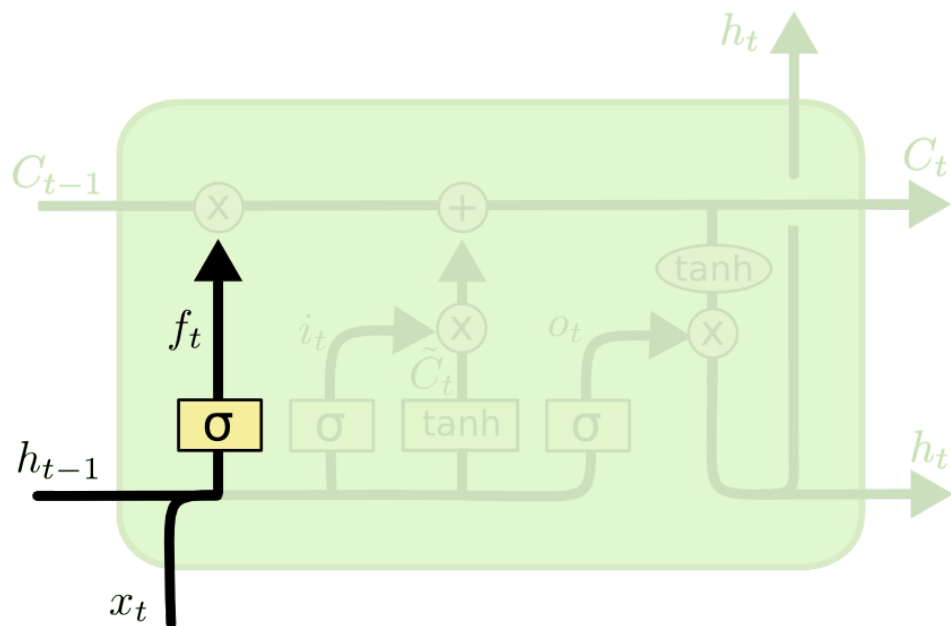


$$i_t = \sigma (W_i \cdot [h_{t-1}, x_t] + b_i)$$

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

<http://colah.github.io/posts/2015-08-Understanding-LSTMs/img/LSTM3-focus-i.png>

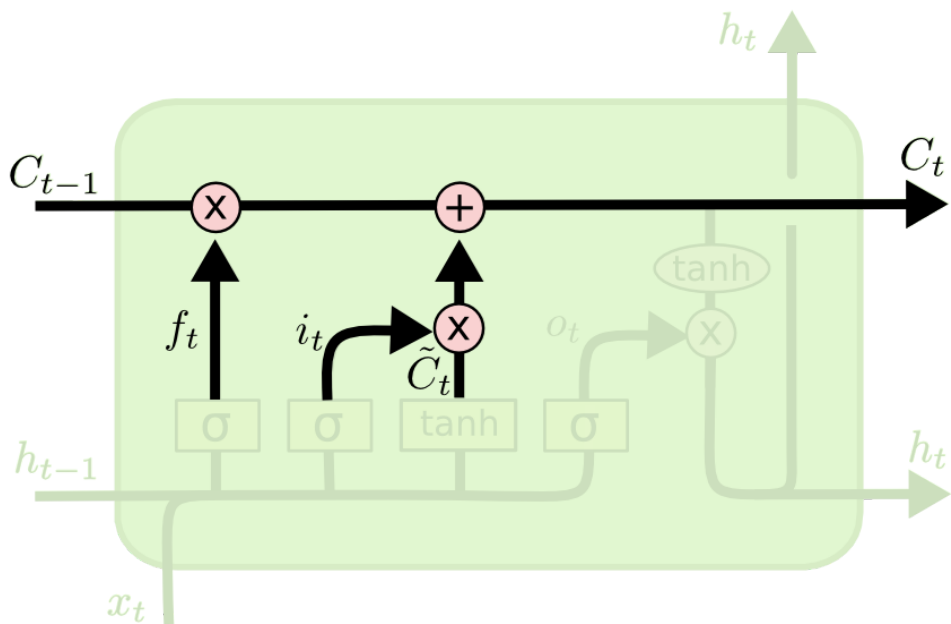
Long Short-Term Memory



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

<http://colah.github.io/posts/2015-08-Understanding-LSTMs/img/LSTM3-focus-f.png>

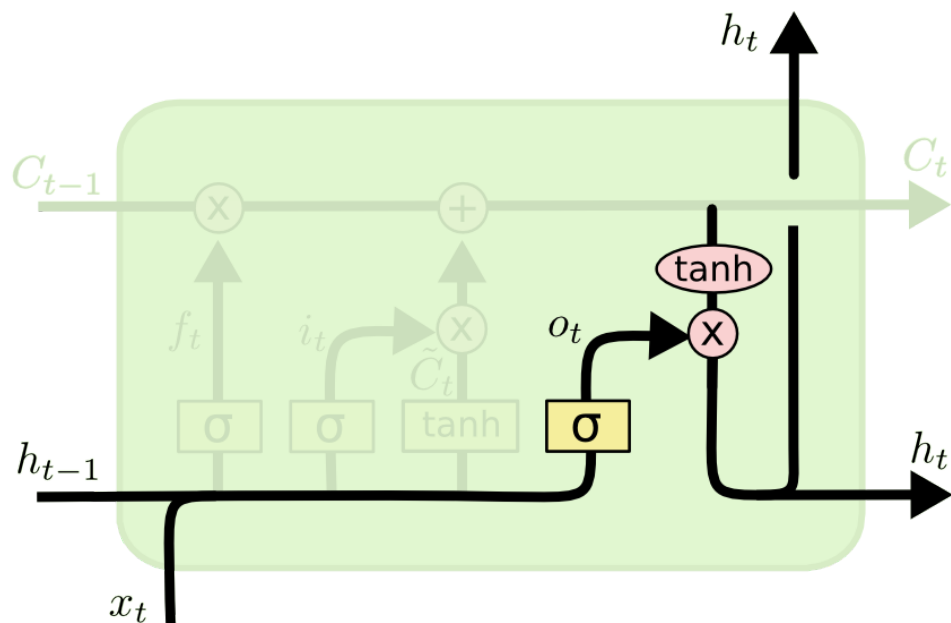
Long Short-Term Memory



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

<http://colah.github.io/posts/2015-08-Understanding-LSTMs/img/LSTM3-focus-C.png>

Long Short-Term Memory



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

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

<http://colah.github.io/posts/2015-08-Understanding-LSTMs/img/LSTM3-focus-o.png>

Gated Recurrent Unit

Gated recurrent unit (GRU) was proposed by Cho et al. (2014) as a simplification of LSTM. The main differences are

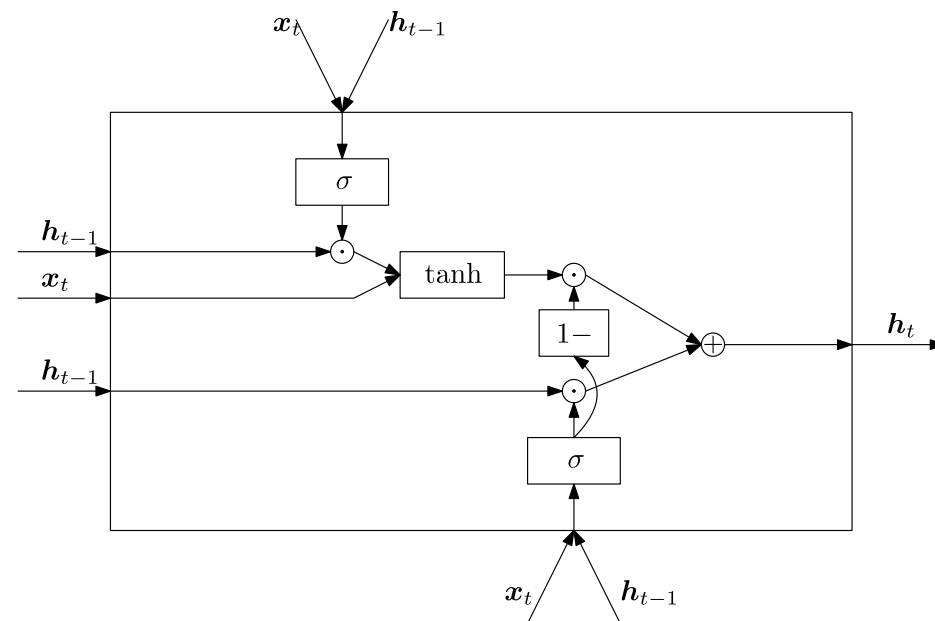
- no memory cell
- forgetting and updating tied together

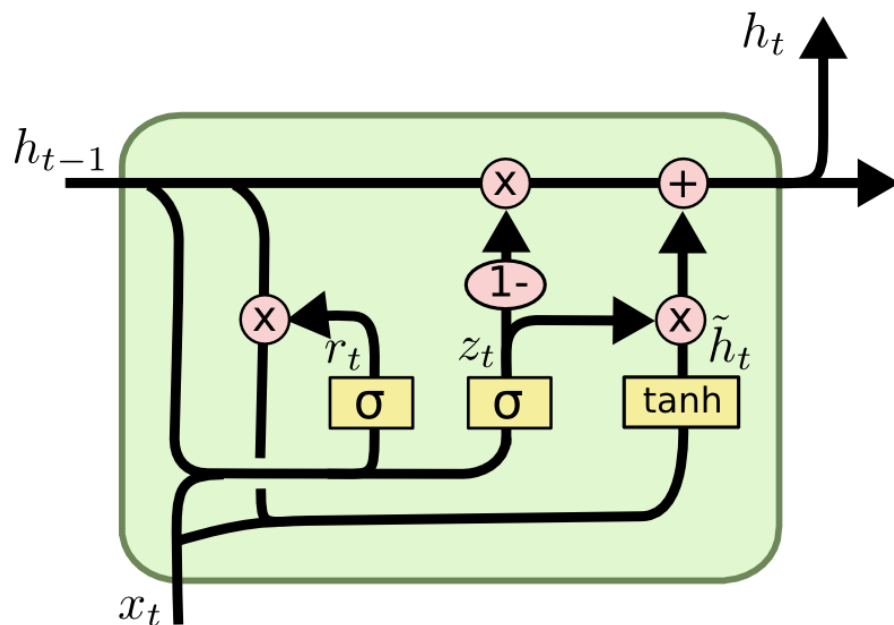
$$\mathbf{r}_t \leftarrow \sigma(\mathbf{W}^r \mathbf{x}_t + \mathbf{V}^r \mathbf{h}_{t-1} + \mathbf{b}^r)$$

$$\mathbf{u}_t \leftarrow \sigma(\mathbf{W}^u \mathbf{x}_t + \mathbf{V}^u \mathbf{h}_{t-1} + \mathbf{b}^u)$$

$$\hat{\mathbf{h}}_t \leftarrow \tanh(\mathbf{W}^h \mathbf{x}_t + \mathbf{V}^h (\mathbf{r}_t \cdot \mathbf{h}_{t-1}) + \mathbf{b}^h)$$

$$\mathbf{h}_t \leftarrow \mathbf{u}_t \cdot \mathbf{h}_{t-1} + (1 - \mathbf{u}_t) \cdot \hat{\mathbf{h}}_t$$





$$z_t = \sigma(W_z \cdot [h_{t-1}, x_t])$$

$$r_t = \sigma(W_r \cdot [h_{t-1}, x_t])$$

$$\tilde{h}_t = \tanh(W \cdot [r_t * h_{t-1}, x_t])$$

$$h_t = (1 - z_t) * h_{t-1} + z_t * \tilde{h}_t$$

<http://colah.github.io/posts/2015-08-Understanding-LSTMs/img/LSTM3-var-GRU.png>

Highway Networks

For input $\mathbf{x}$, fully connected layer computes

$$\mathbf{y} \leftarrow H(\mathbf{x}, \mathbf{W}_H).$$

Highway networks add residual connection with gating:

$$\mathbf{y} \leftarrow H(\mathbf{x}, \mathbf{W}_H) \cdot T(\mathbf{x}, \mathbf{W}_T) + \mathbf{x} \cdot (1 - T(\mathbf{x}, \mathbf{W}_T)).$$

Usually, the gating is defined as

$$T(\mathbf{x}, \mathbf{W}_T) \leftarrow \sigma(\mathbf{W}_T \mathbf{x} + \mathbf{b}_T).$$

Note that the resulting update is very similar to a GRU cell with $\mathbf{h}_t$ removed; for a fully connected layer $H(\mathbf{x}, \mathbf{W}_H) = \tanh(\mathbf{W}_H \mathbf{x} + \mathbf{b}_H)$ it is exactly it.

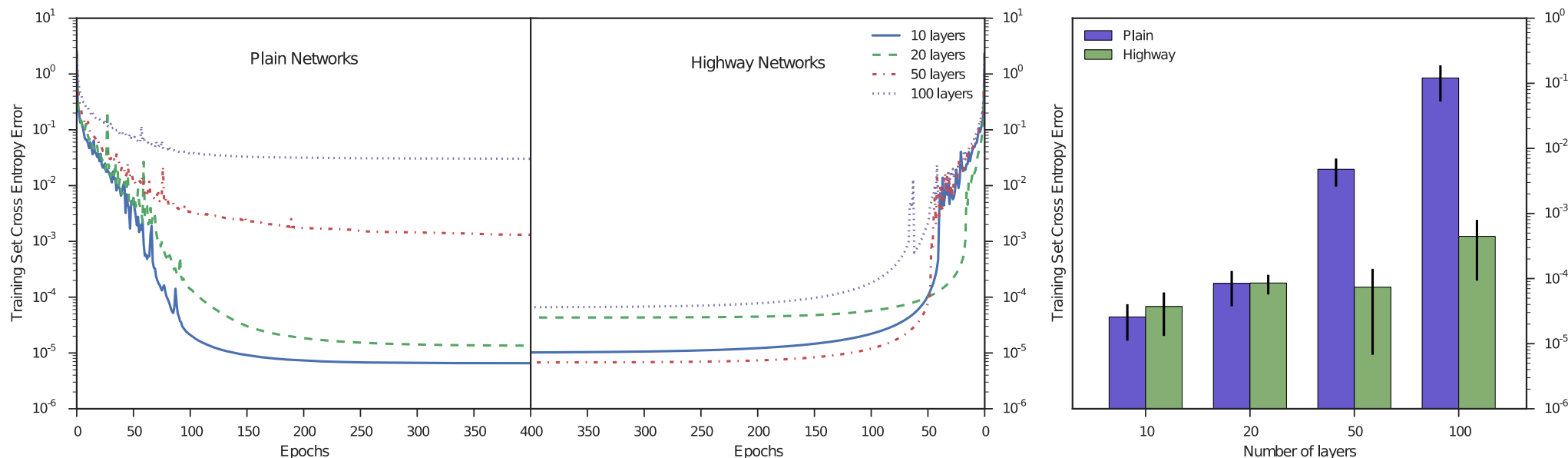


Figure 1: Comparison of optimization of plain networks and highway networks of various depths. *Left:* The training curves for the best hyperparameter settings obtained for each network depth. *Right:* Mean performance of top 10 (out of 100) hyperparameter settings. Plain networks become much harder to optimize with increasing depth, while highway networks with up to 100 layers can still be optimized well. Best viewed on screen (larger version included in Supplementary Material).

Figure 1 of paper "Training Very Deep Networks", <https://arxiv.org/abs/1507.06228>.

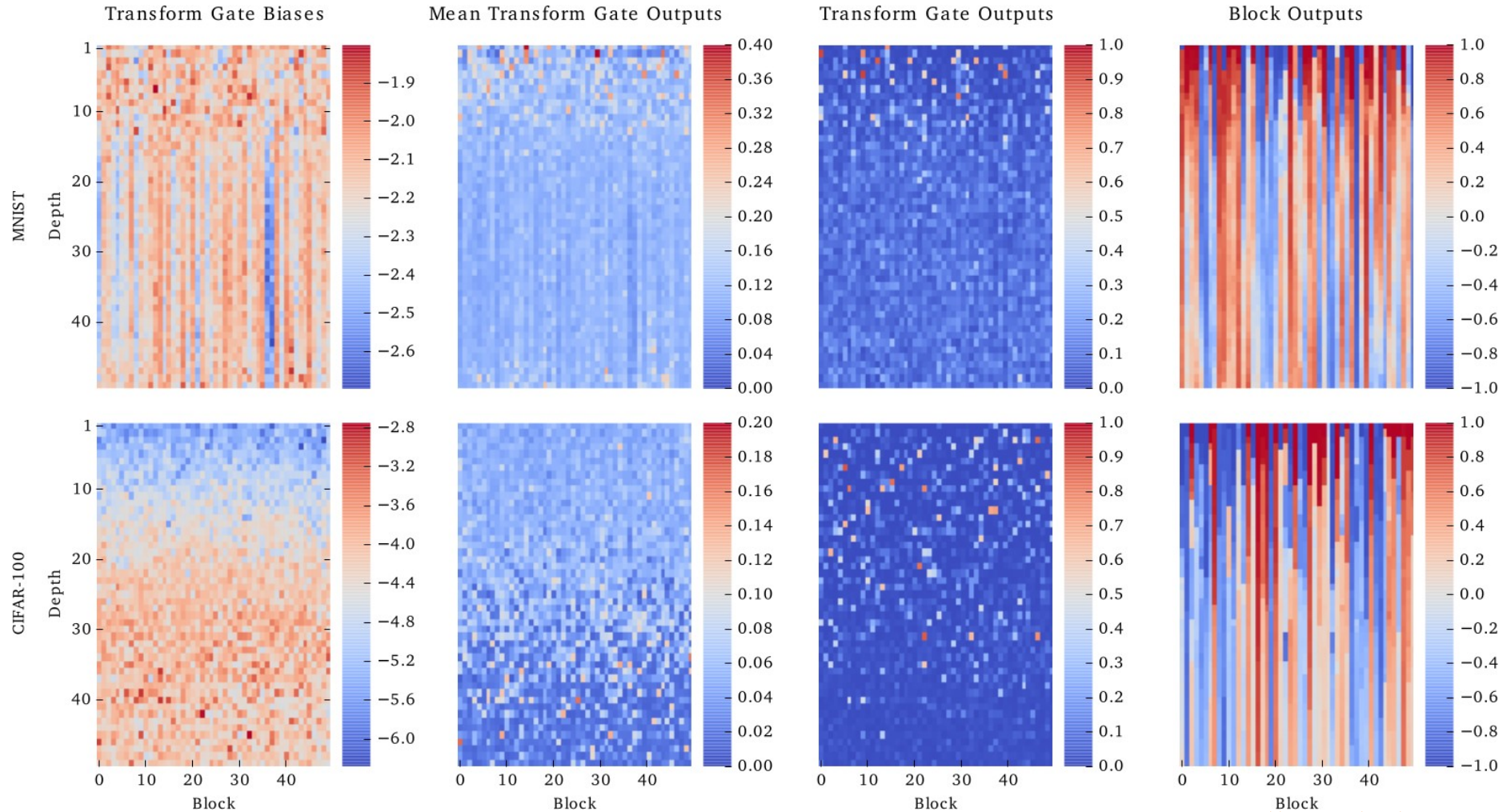


Figure 2 of paper "Training Very Deep Networks", <https://arxiv.org/abs/1507.06228>.

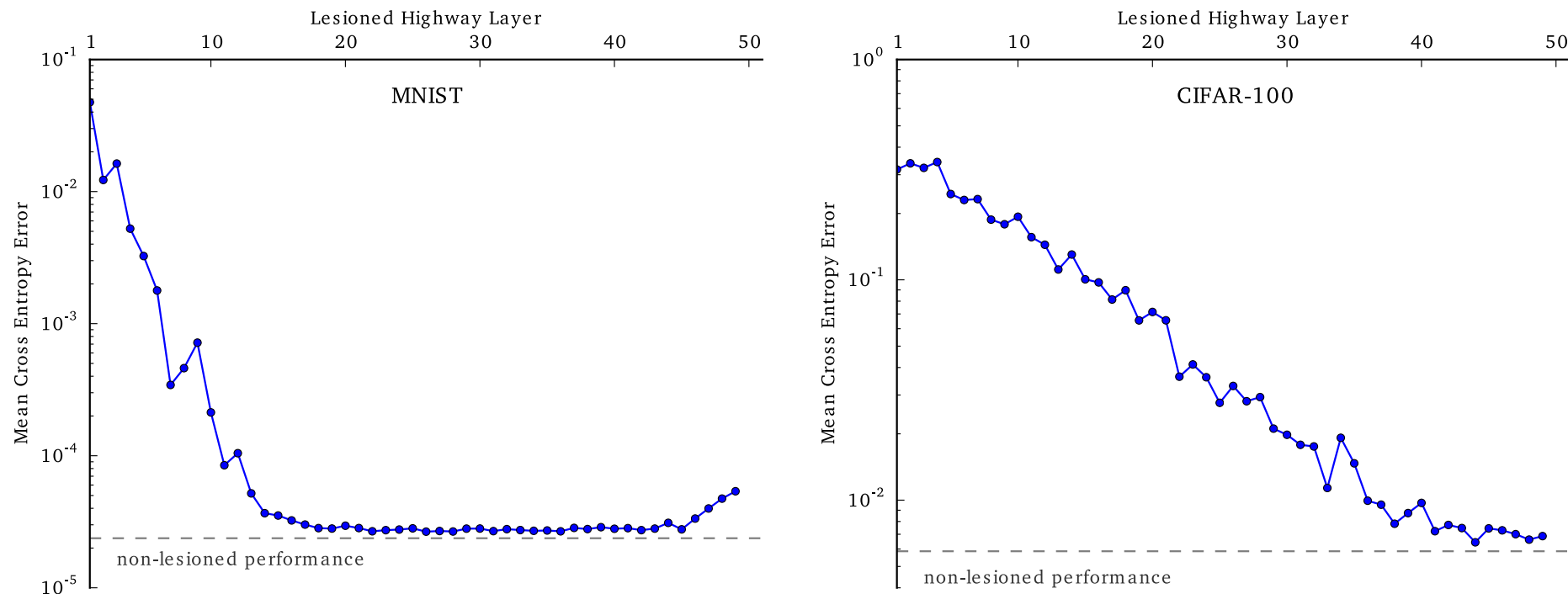


Figure 4: Lesioned training set performance (y-axis) of the best 50-layer highway networks on MNIST (left) and CIFAR-100 (right), as a function of the lesioned layer (x-axis). Evaluated on the full training set while forcefully closing all the transform gates of a single layer at a time. The non-lesioned performance is indicated as a dashed line at the bottom.

Figure 4 of paper "Training Very Deep Networks", <https://arxiv.org/abs/1507.06228>.

Dropout

- Using dropout on hidden states interferes with long-term dependencies.
- However, using dropout on the inputs and outputs works well and is used frequently.
 - In case residual connections are present, the output dropout needs to be applied before adding the residual connection.
- Several techniques were designed to allow using dropout on hidden states.
 - Variational Dropout
 - Recurrent Dropout
 - Zoneout

Variational Dropout

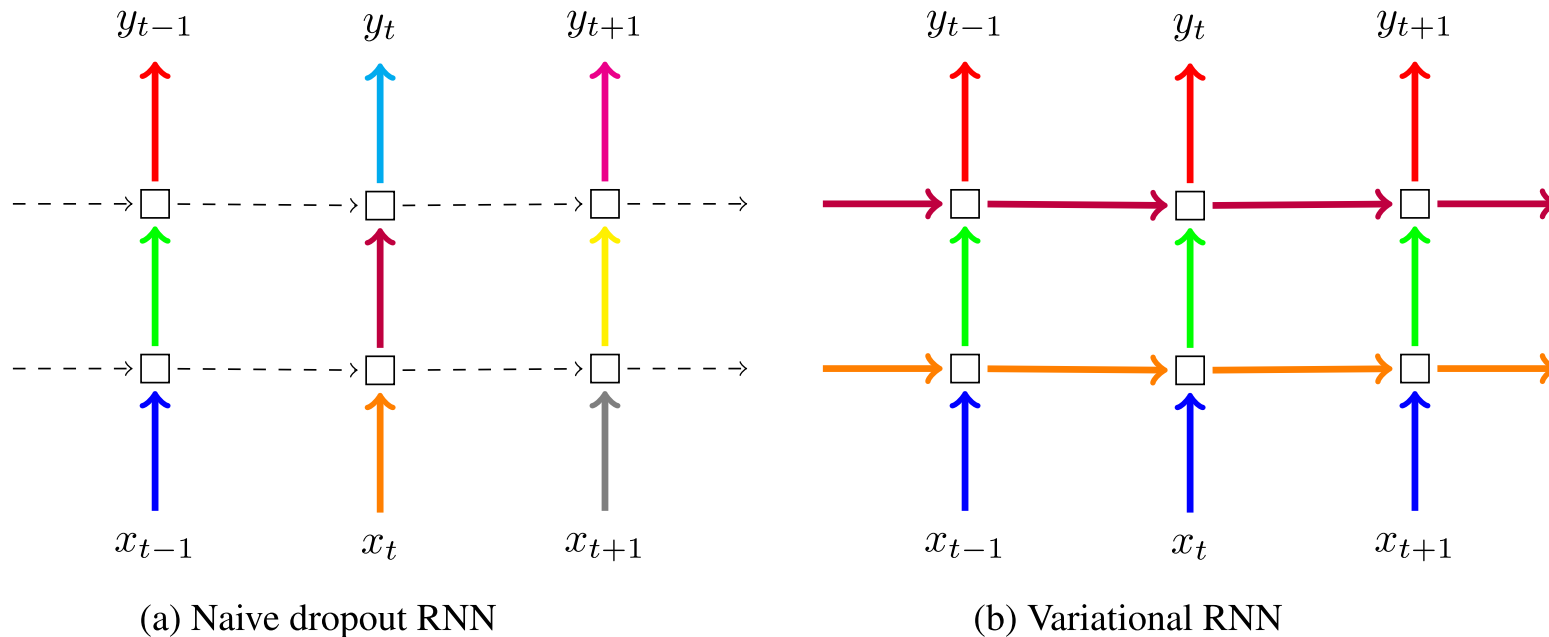


Figure 1 of paper "A Theoretically Grounded Application of Dropout in Recurrent Neural Networks", <https://arxiv.org/abs/1512.05287.pdf>

To implement variational dropout on inputs in TensorFlow, use `noise_shape` of `tf.keras.layers.Dropout` to force the same mask across time-steps. The variational dropout on the hidden states can be implemented using `recurrent_dropout` argument of `tf.keras.layers.{LSTM,GRU,SimpleRNN}{,Cell}`.

Recurrent Dropout

Dropout only candidate states (i.e., values added to the memory cell in LSTM and previous state in GRU).

Zoneout

Randomly preserve hidden activations instead of dropping them.

Batch Normalization

Very fragile and sensitive to proper initialization – there were papers with negative results (*Dario Amodei et al, 2015: Deep Speech 2* or *Cesar Laurent et al, 2016: Batch Normalized Recurrent Neural Networks*) until people managed to make it work (*Tim Cooijmans et al, 2016: Recurrent Batch Normalization*; specifically, initializing $\gamma = 0.1$ did the trick).

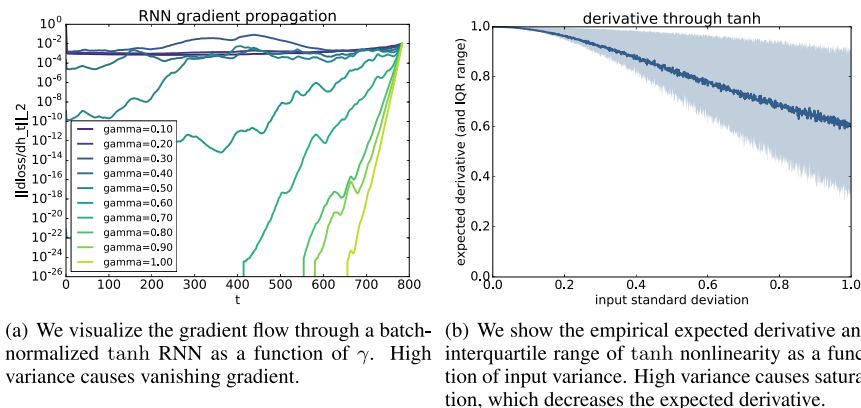


Figure 1 of paper "Recurrent Batch Normalization", <https://arxiv.org/abs/1603.09025>.

Layer Normalization

Much more stable than batch normalization.

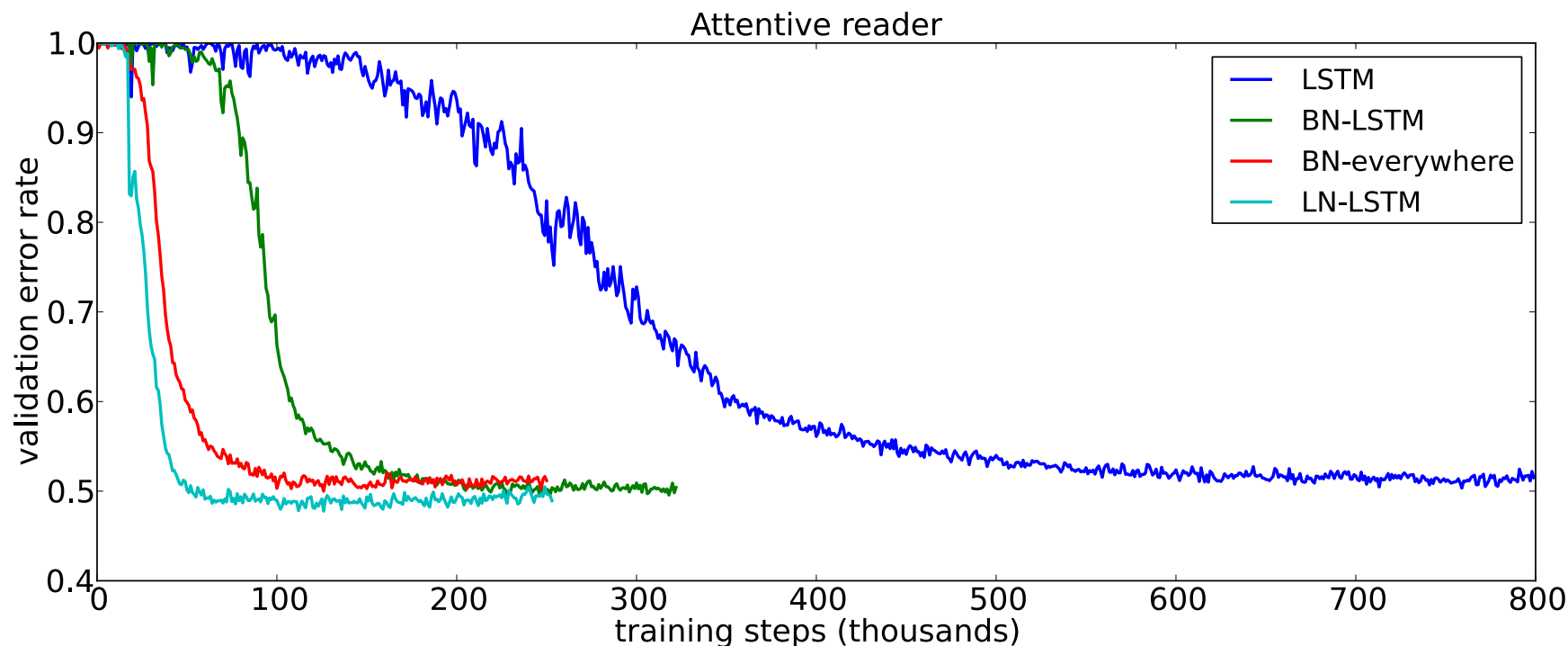
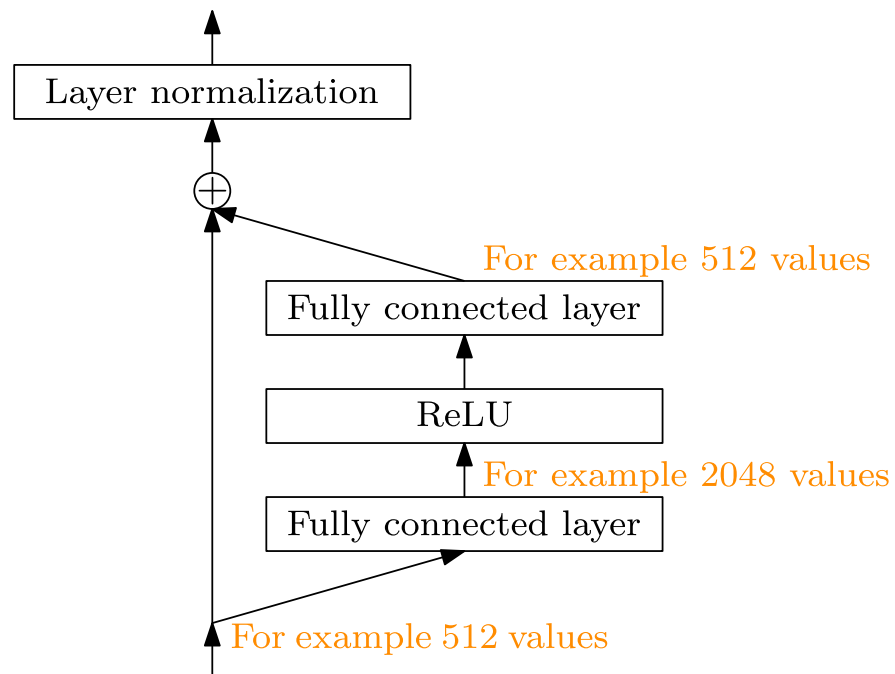


Figure 2: Validation curves for the attentive reader model. BN results are taken from [Cooijmans et al., 2016].

Figure 2 of paper "Layer Normalization", <https://arxiv.org/abs/1607.06450>.

Layer Normalization

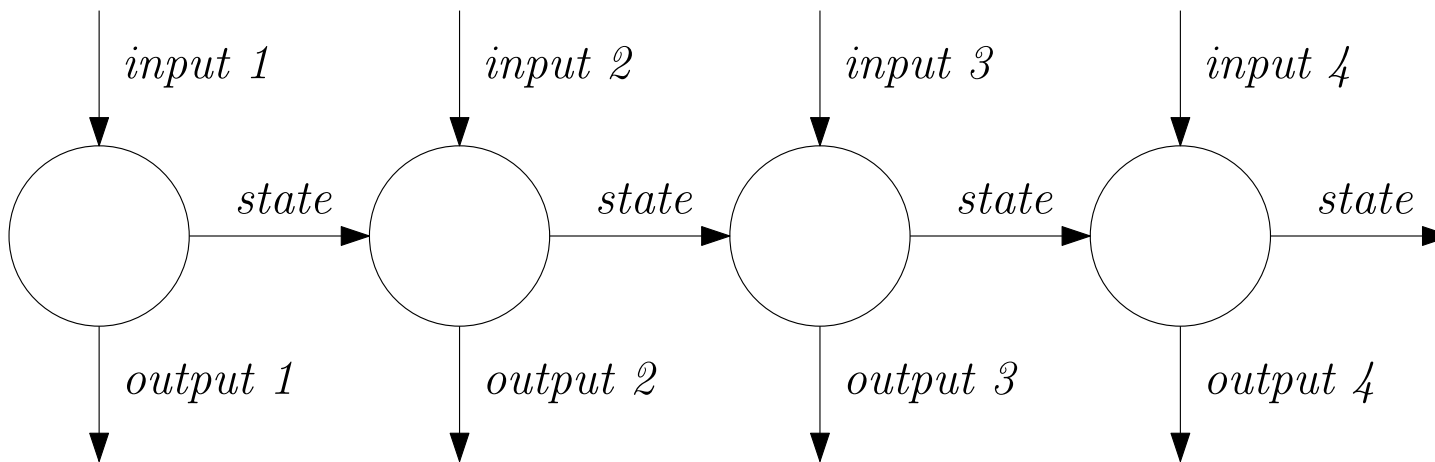
In an important recent architecture (namely Transformer), many fully connected layers are used, with a residual connection and a layer normalization.



This could be considered an alternative to highway networks, i.e., a suitable residual connection for fully connected layers. Note the architecture can be considered as a variant of a mobile inverted bottleneck 1×1 convolution block.

Sequence Element Representation

Create output for individual elements, for example for classification of the individual elements.

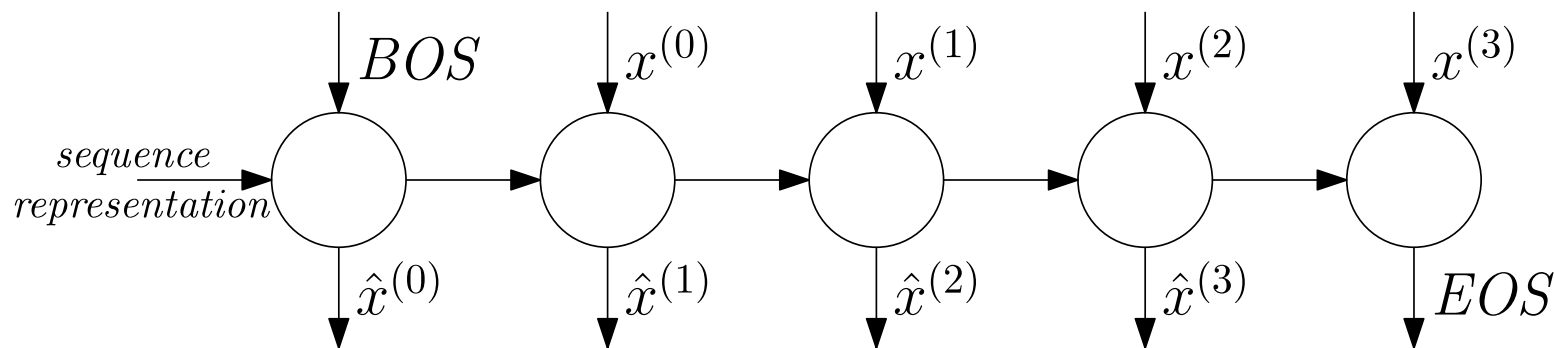


Sequence Representation

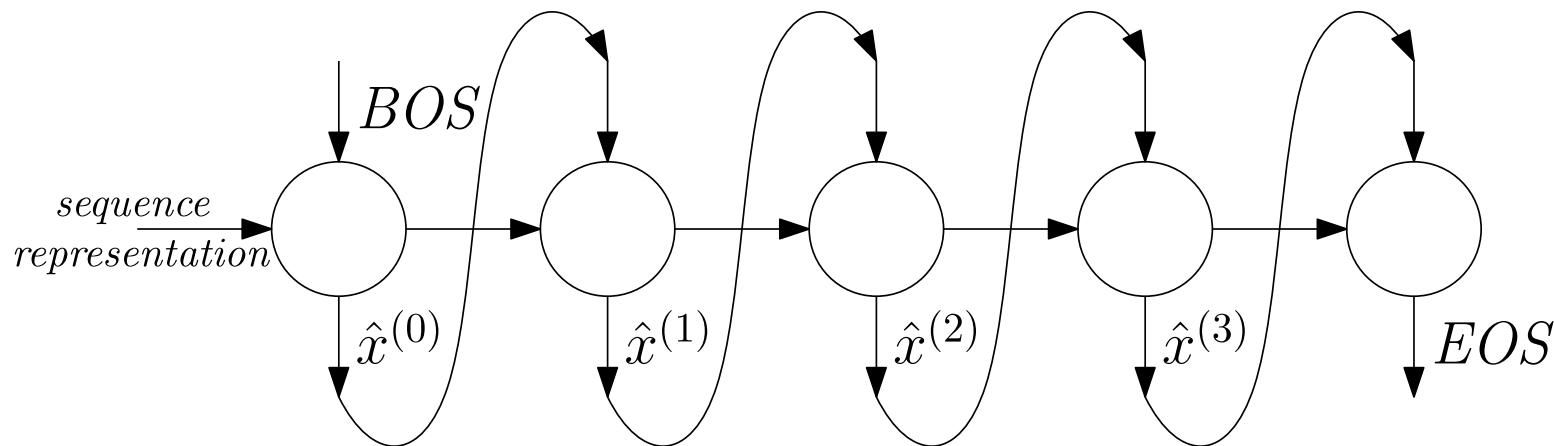
Generate a single output for the whole sequence (either the last output of the last state).

Sequence Prediction

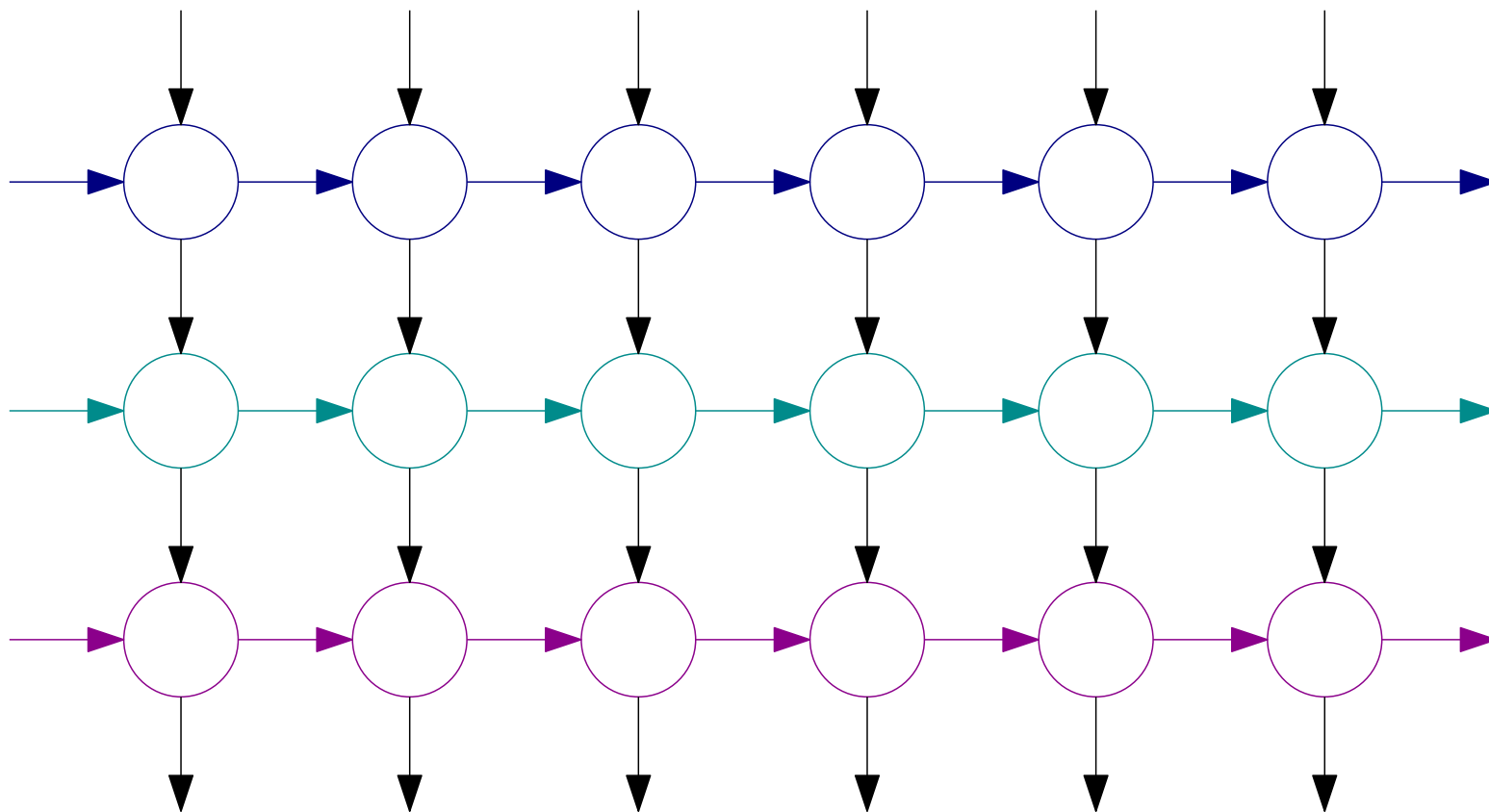
During training, predict next sequence element.



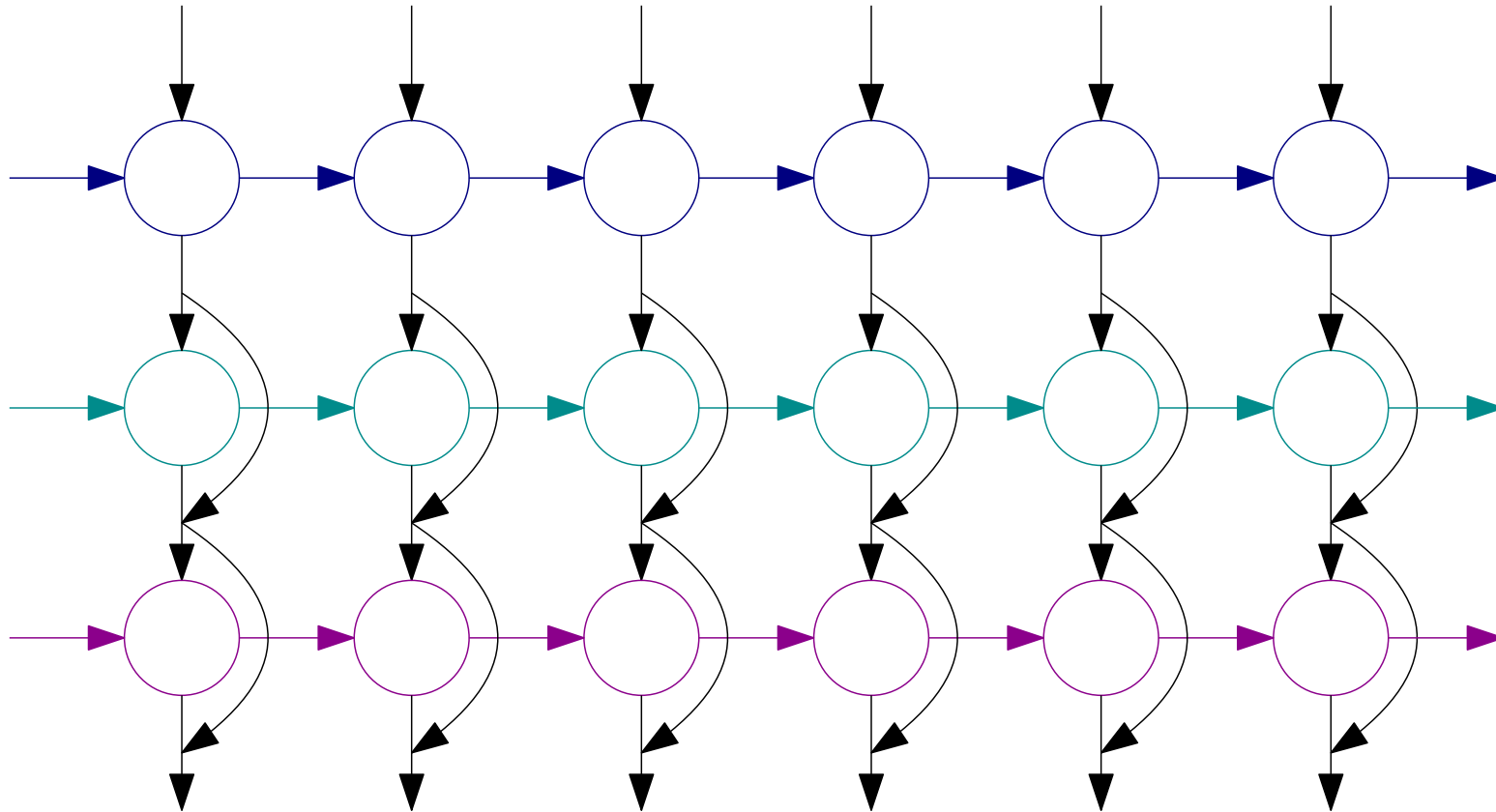
During inference, use predicted elements as further inputs.



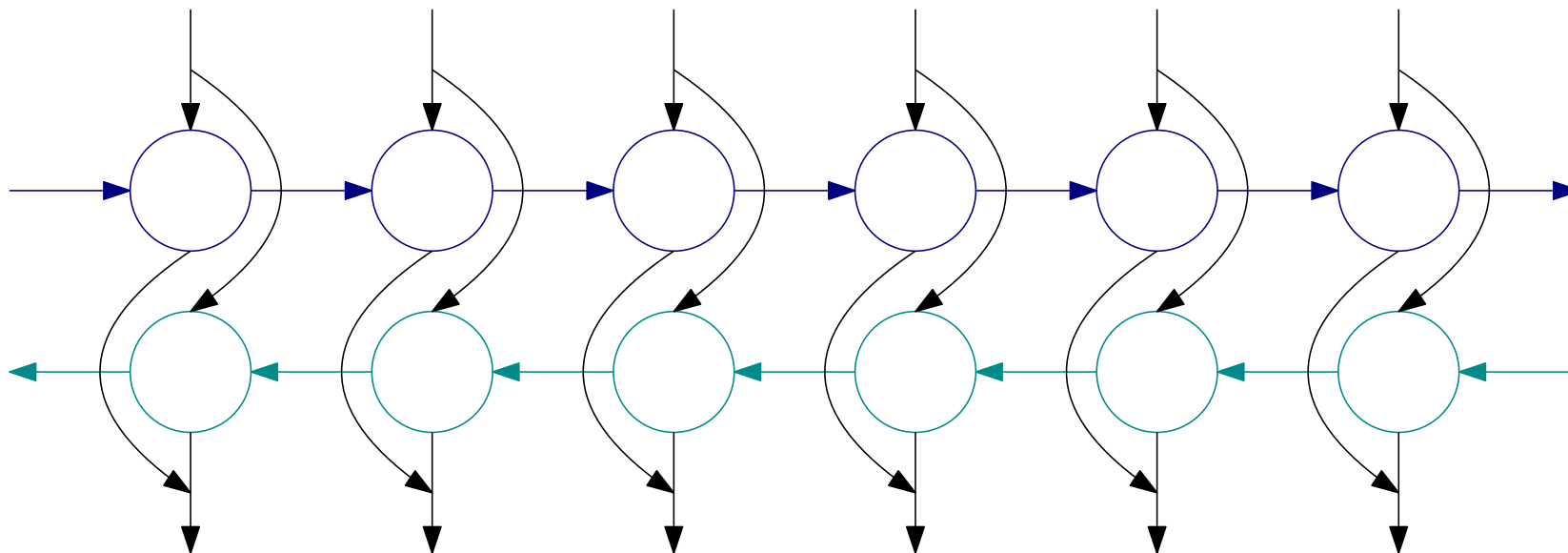
We might stack several layers of recurrent neural networks. Usually using two or three layers gives better results than just one.



In case of multiple layers, residual connections usually improve results. Because dimensionality has to be the same, they are usually applied from the second layer.



To consider both the left and right contexts, a *bidirectional* RNN can be used, which consists of parallel application of a *forward* RNN and a *backward* RNN.



The outputs of both directions can be either *added* or *concatenated*. Even if adding them does not seem very intuitive, it does not increase dimensionality and therefore allows residual connections to be used in case of multilayer bidirectional RNN.

We might represent *words* using one-hot encoding, considering all words to be independent of each other.

However, words are not independent – some are more similar than others.

Ideally, we would like some kind of similarity in the space of the word representations.

Distributed Representation

The idea behind distributed representation is that objects can be represented using a set of common underlying factors.

We therefore represent words as fixed-size *embeddings* into $\mathbb{R}^d$ space, with the vector elements playing role of the common underlying factors.

These embeddings are initialized randomly and trained together with the rest of the network.

The word embedding layer is in fact just a fully connected layer on top of one-hot encoding. However, it is not implemented in that way.

Instead, a so-called *embedding* layer is used, which is much more efficient. When a matrix is multiplied by an one-hot encoded vector (all but one zeros and exactly one 1), the row corresponding to that 1 is selected, so the embedding layer can be implemented only as a simple lookup.

In TensorFlow, the embedding layer is available as

```
tf.keras.layers.Embedding(input_dim, output_dim)
```

Even if the embedding layer is just a fully connected layer on top of one-hot encoding, it is important that this layer is *shared* across the whole network.

