



Language
Technologies
Institute

Carnegie
Mellon
University

Multimodal Machine Learning

Lecture 3.2: Recurrent Networks

Louis-Philippe Morency

* Original version co-developed with Tadas Baltrusaitis

Lecture Objectives

- Sequential modeling with convolutional networks
- Word representations
 - Distributional hypothesis
 - Learning neural representations
- Language models and sequence modeling tasks
- Recurrent neural networks
 - Gated recurrent neural networks
 - Long Short-Term Memory (LSTM) model
 - Multi-view LSTM
 - Backpropagation through time

Sequential Modeling with Convolutional Networks



Modeling Temporal and Sequential Data

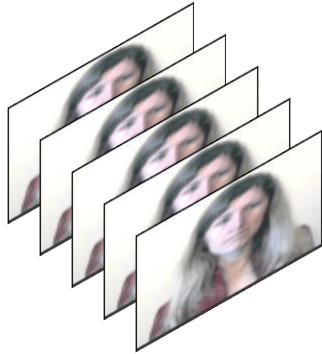


How to represent a video sequence?

One option: Recurrent Neural Networks
(more about this on Thursday)

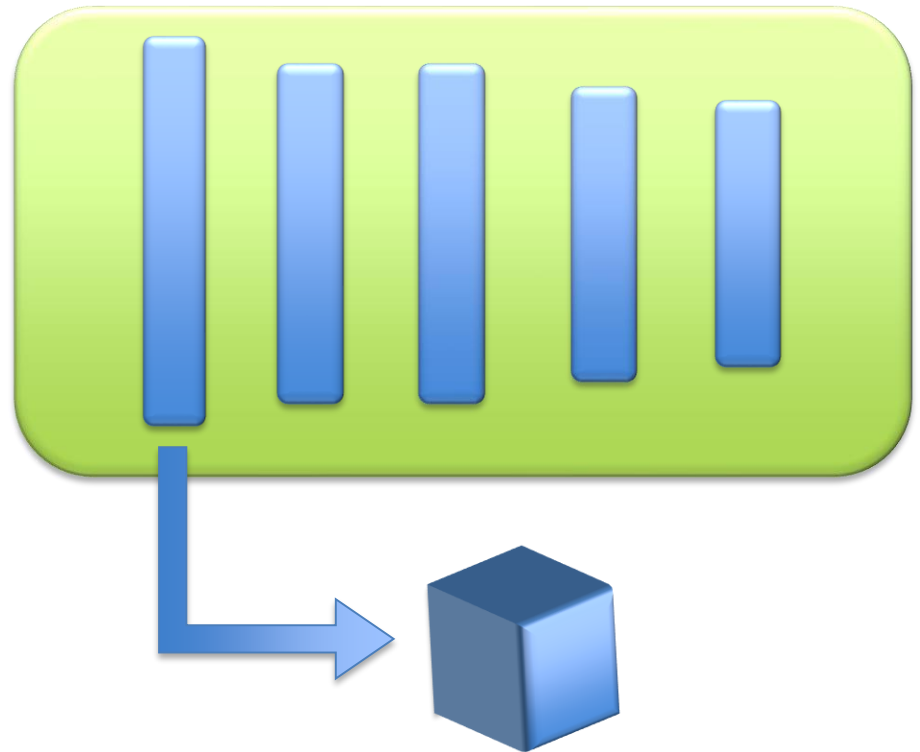


3D CNN



Input as a 3D tensor
(stacking video images)

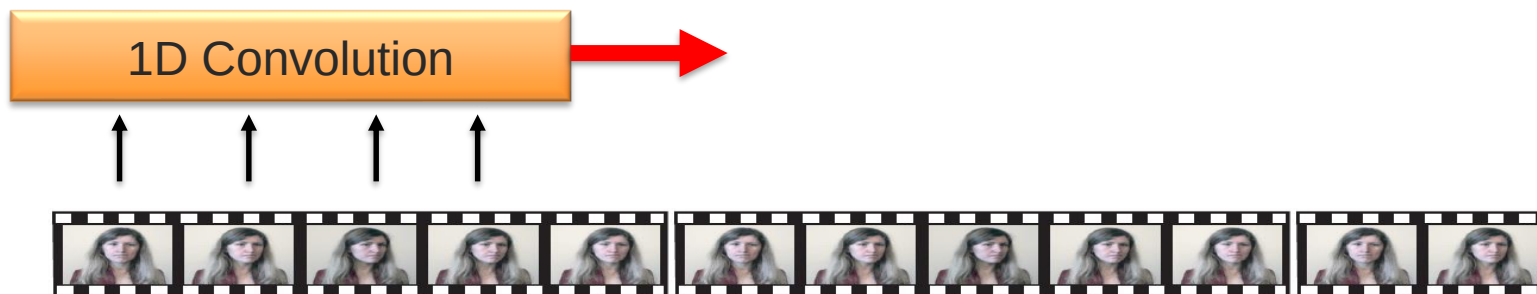
3D CNN



First layer with 3D kernels



Time-Delay Neural Network

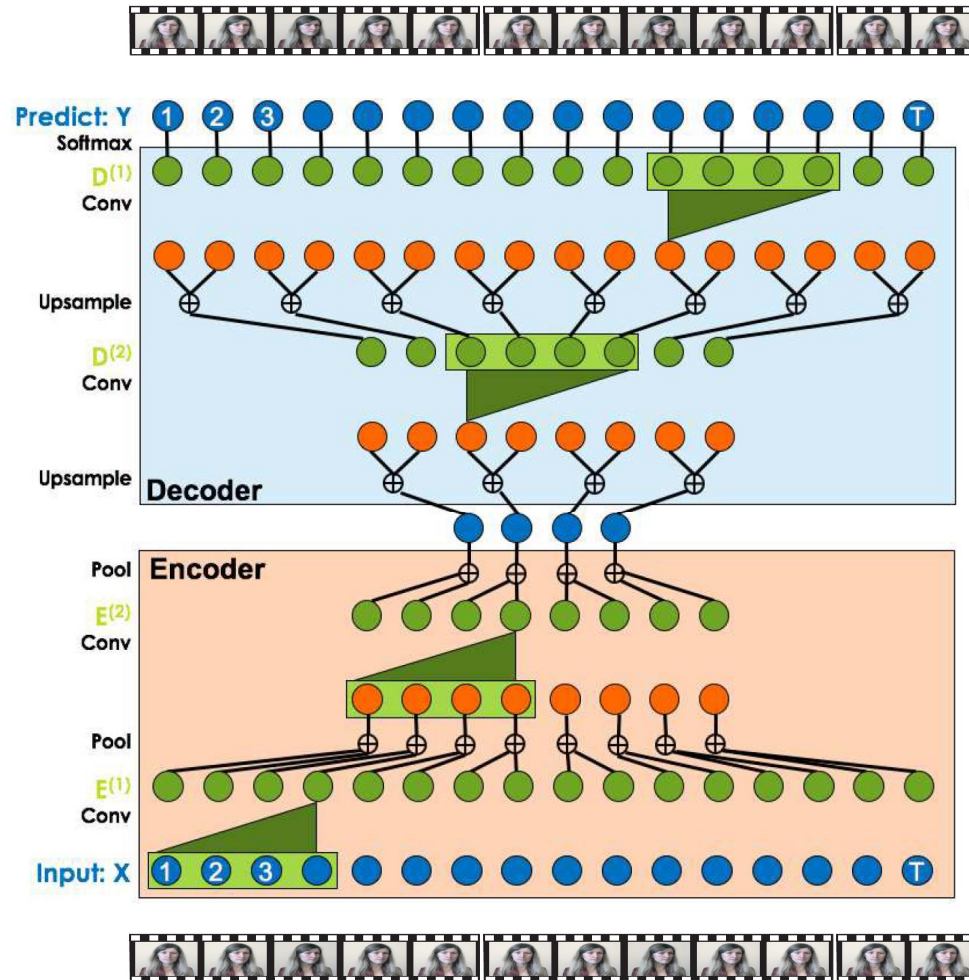


Alexander Waibel, Phoneme Recognition Using Time-Delay Neural Networks, SP87-100, Meeting of the Institute of Electrical, Information and Communication Engineers (IEICE), December, 1987, Tokyo, Japan.

Temporal Convolution Network (TCN) [Lea et al., CVPR 2017]

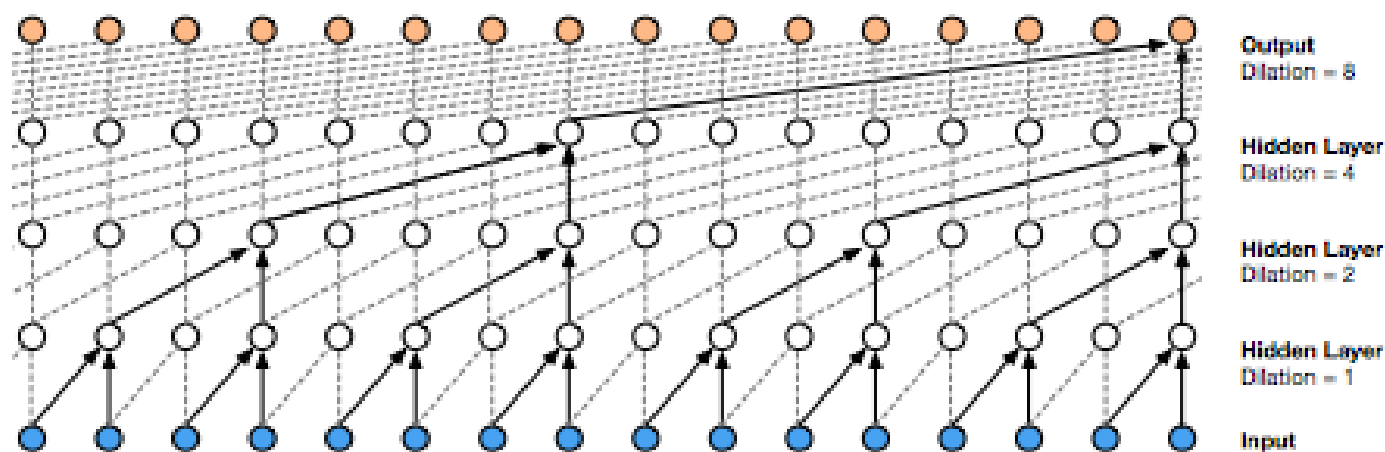
Decoder

Encoder



Dilated TCN Model [Lea et al., CVPR 2017]

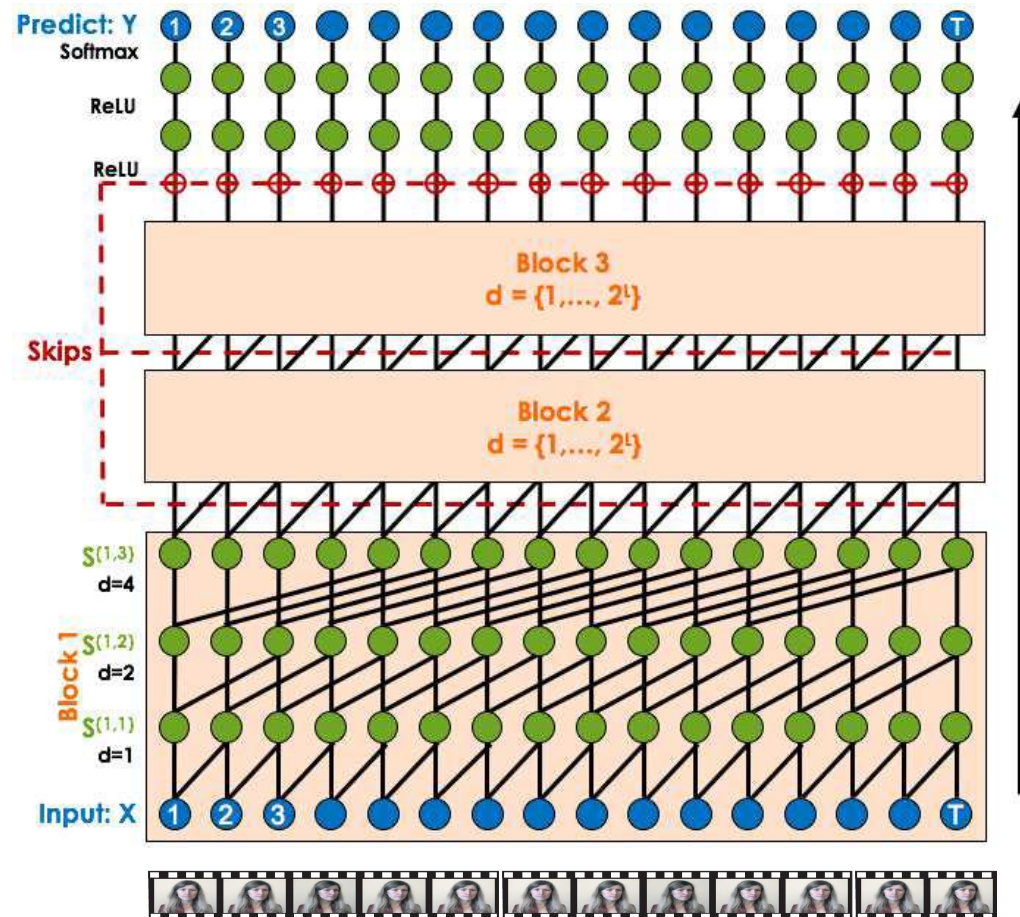
Dilated Convolutions



Dilation of 4: Step size of 4 when convoluting

+ Skip connections to help with deep modeling

Dilated TCN Models [Lea et al., CVPR 2017]



Representing Words: Distributed Semantics



What is the meaning of “bardiwac”?

- He handed her glass of **bardiwac**.
 - Beef dishes are made to complement the **bardiwacs**.
 - Nigel staggered to his feet, face flushed from too much **bardiwac**.
 - Malbec, one of the lesser-known **bardiwac** grapes, responds well to Australia’s sunshine.
 - I dined off bread and cheese and this excellent **bardiwac**.
 - The drinks were delicious: blood-red **bardiwac** as well as light, sweet Rhenish.
- ⇒ **bardiwac** is a heavy red alcoholic beverage made from grapes

The Distributional Hypothesis

- Distribution Hypothesis (DH) [Lenci 2008]
 - At least certain aspects of the meaning of lexical expressions depend on their distributional properties in the linguistic contexts
 - The degree of semantic similarity between two linguistic expressions α and β is a function of the similarity of the linguistic contexts in which α and β can appear
- Weak and strong DH
 - Weak view as a quantitative method for semantic analysis and lexical resource induction
 - Strong view as a cognitive hypothesis about the form and origin of semantic representations; assuming that word distributions in context play a specific *causal role* in forming meaning representations.

Geometric interpretation

- row vector $\mathbf{x}_{\text{dog}}$ describes usage of word *dog* in the corpus
- can be seen as coordinates of point in n -dimensional Euclidean space $\mathbb{R}^n$

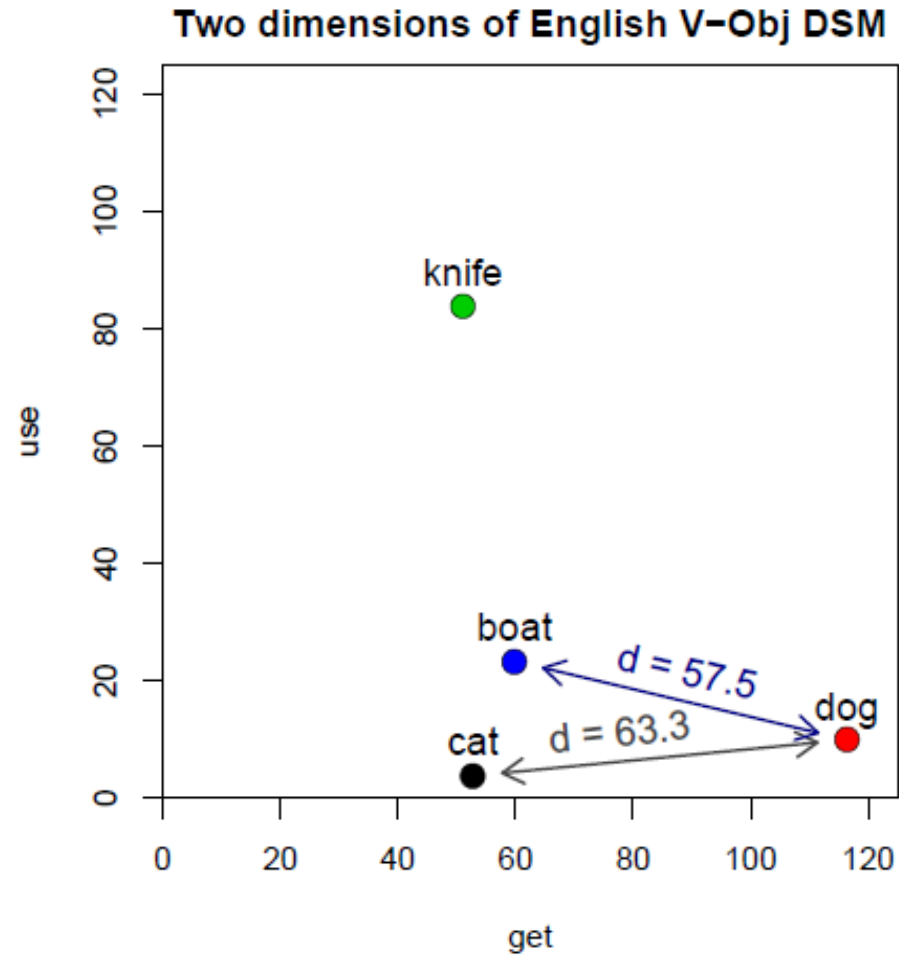


	get	see	use	hear	eat	kill
knife	51	20	84	0	3	0
cat	52	58	4	4	6	26
dog	115	83	10	42	33	17
boat	59	39	23	4	0	0
cup	98	14	6	2	1	0
pig	12	17	3	2	9	27
banana	11	2	2	0	18	0

co-occurrence matrix $\mathbf{M}$

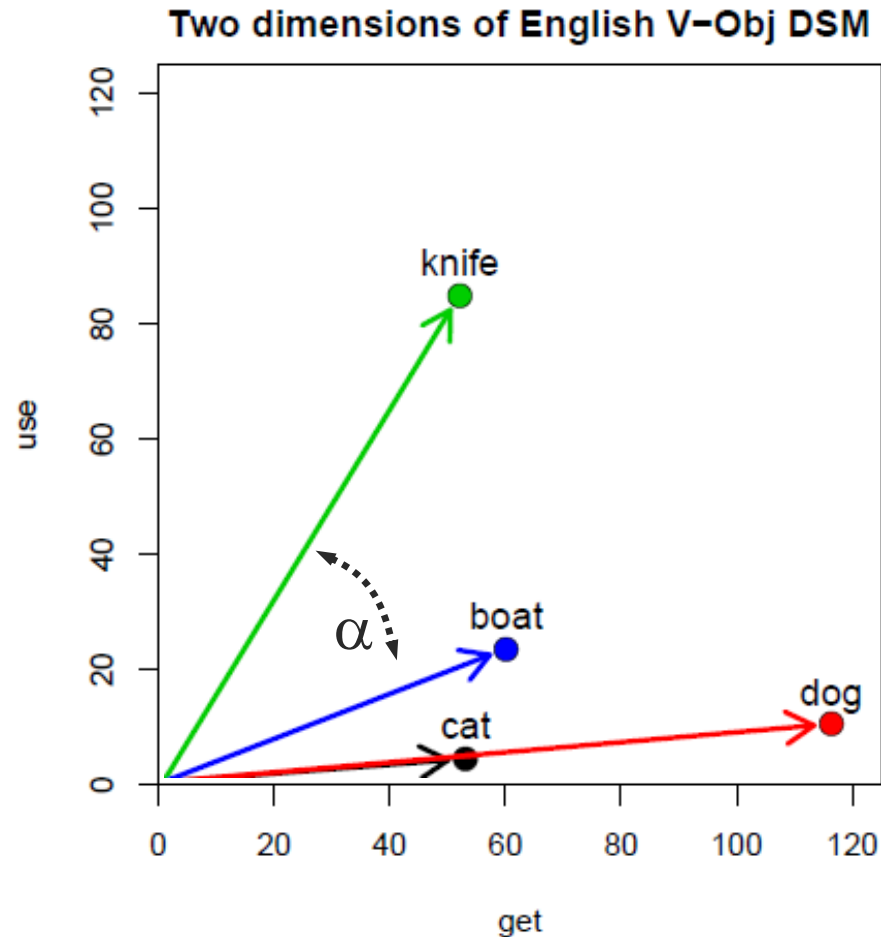
Distance and similarity

- illustrated for two dimensions: *get* and *use*: $\mathbf{x}_{\text{dog}} = (115, 10)$
- similarity = spatial proximity (Euclidean distance)
- location depends on frequency of noun ($f_{\text{dog}} \approx 2.7 \cdot f_{\text{cat}}$)

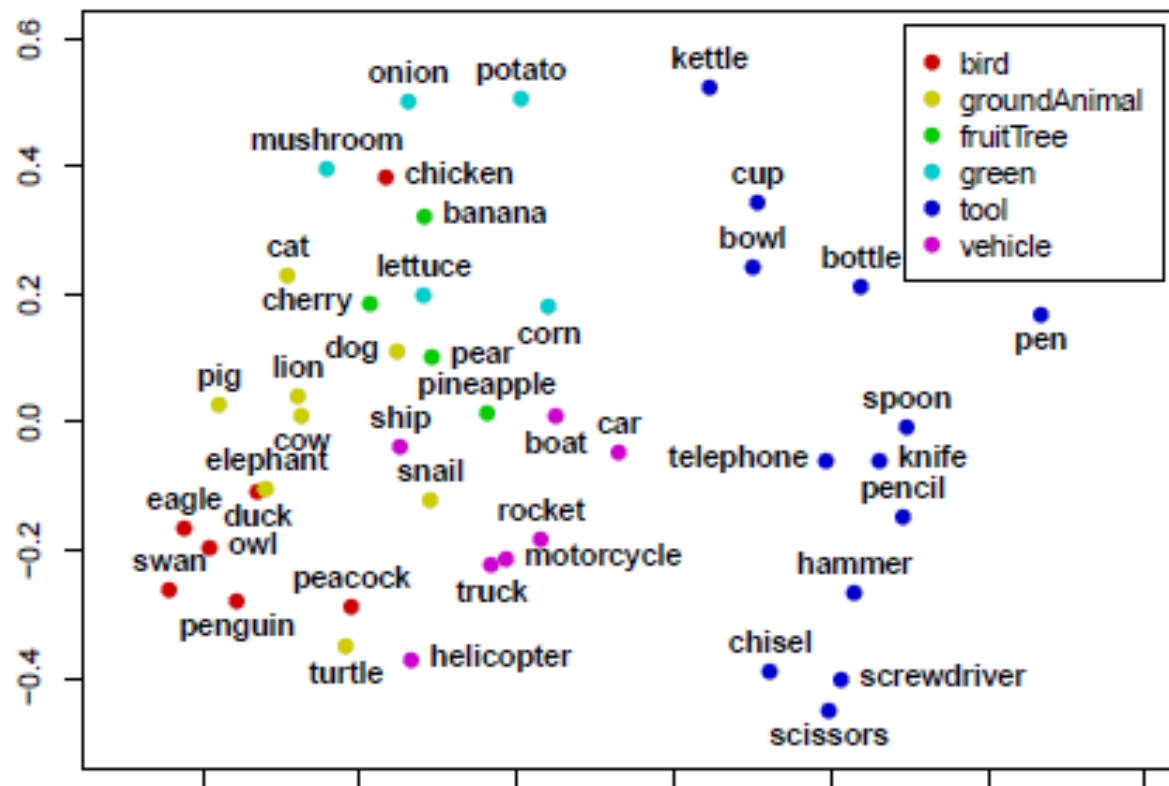


Angle and similarity

- direction more important than location
- normalise “length”
 $\|\mathbf{x}_{\text{dog}}\|$ of vector
- or use angle α as distance measure



Semantic maps



Learning Neural Word Representations



How to learn neural word representations?

- ➔ **Distribution hypothesis:** Approximate the word meaning by its surrounding words
- ➔ Words used in a similar context will lie close together

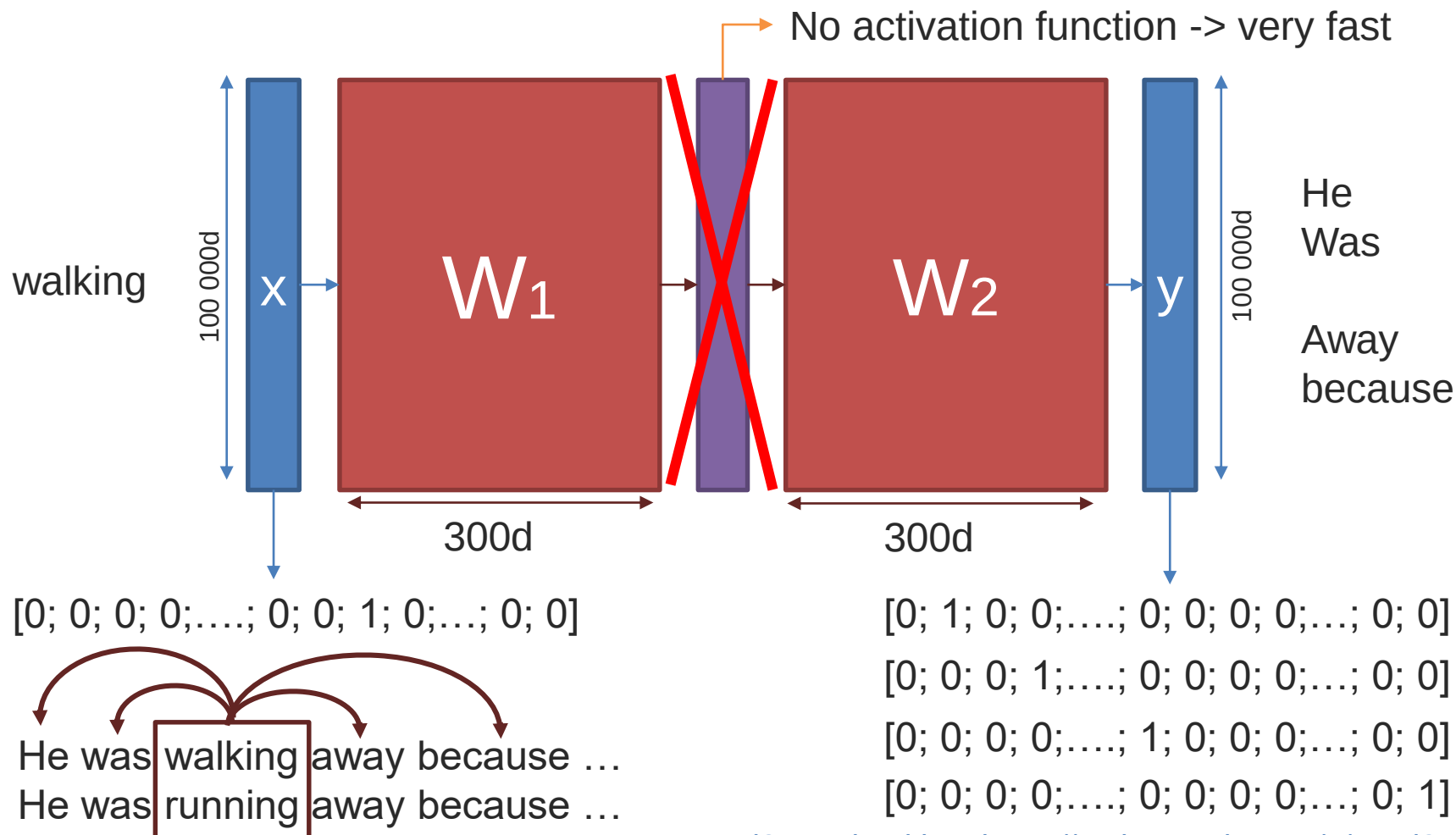


- ➔ **Instead of capturing co-occurrence counts directly, predict surrounding words of every word**

$$\frac{1}{T} \sum_{t=1}^T \sum_{-c \leq j \leq c, j \neq 0} \log p(w_{t+j} | w_t)$$



How to learn neural word representations?



Word2vec algorithm: <https://code.google.com/p/word2vec/>

How to use these word representations

If we would have a vocabulary of 100 000 words:

Classic NLP: $\xleftarrow{\text{100 000 dimensional vector}}$

Walking: $[0; 0; 0; 0; \dots; 0; 0; 1; 0; \dots; 0; 0]$

Running: $[0; 0; 0; 0; \dots; 0; 0; 0; 0; \dots; 1; 0]$

➡ Similarity = 0.0

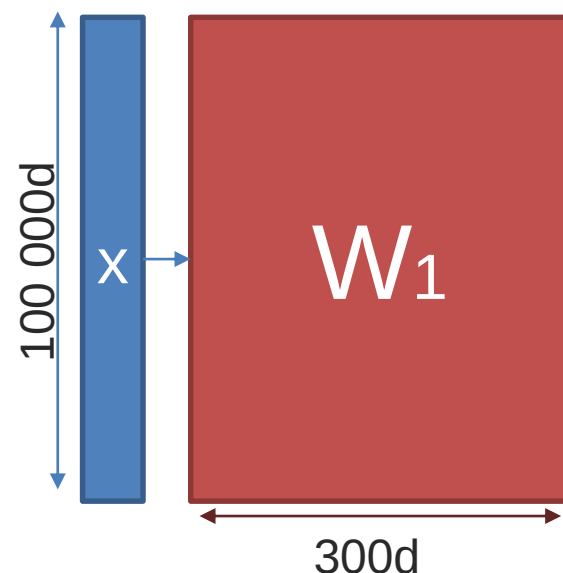
↓ Transform: $x' = x * W$

Goal: $\xleftarrow{\text{300 dimensional vector}}$

Walking: $[0,1; 0,0003; 0; \dots; 0,02; 0,08; 0,05]$

Running: $[0,1; 0,0004; 0; \dots; 0,01; 0,09; 0,05]$

➡ Similarity = 0.9



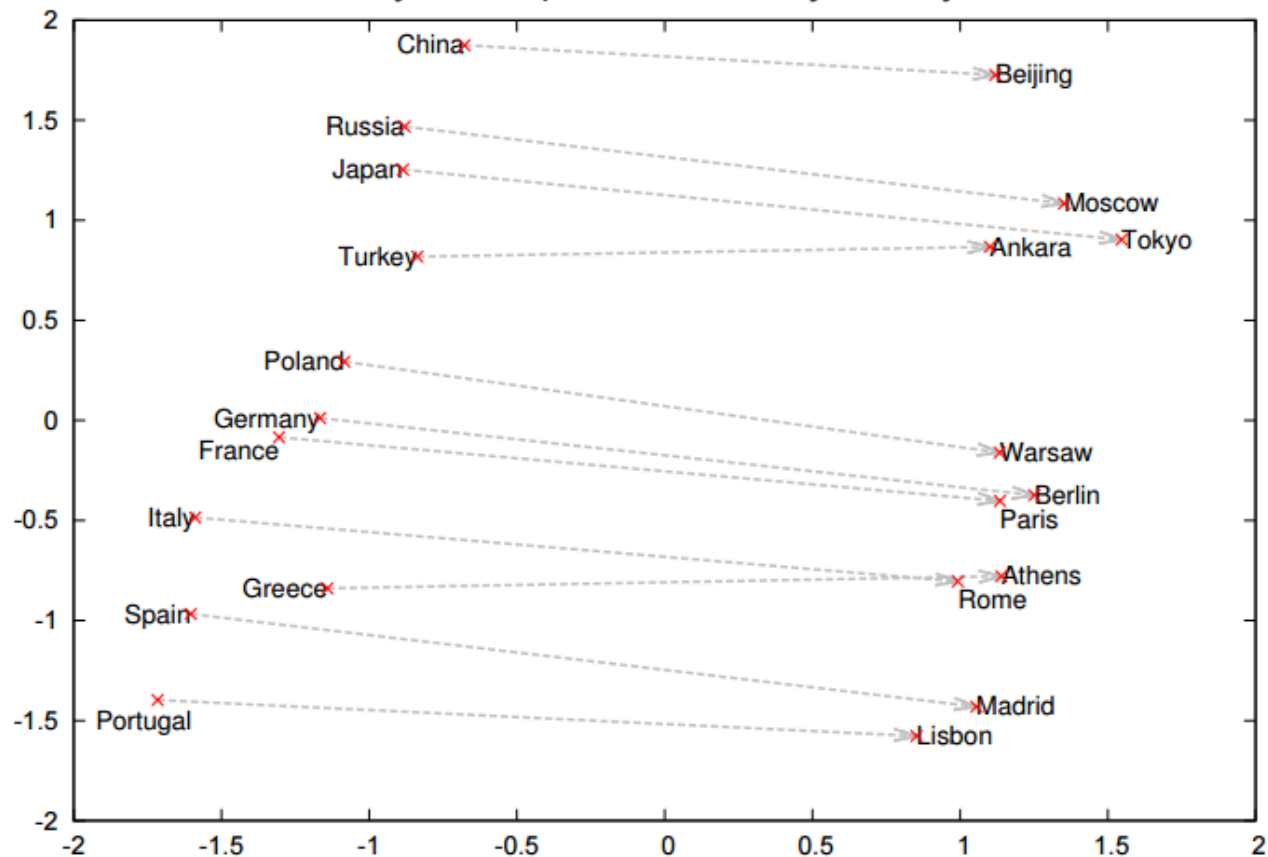
Vector space models of words

- ➡ While learning these word representations, we are actually building a vector space in which all words reside with certain relationships between them
- ➡ Encodes both syntactic and semantic relationships
- ➡ This vector space allows for algebraic operations:

$$\text{Vec}(\text{king}) - \text{vec}(\text{man}) + \text{vec}(\text{woman}) \approx \text{vec}(\text{queen})$$

Why linear algebra is working?

Vector space models of words: semantic relationships



Trained on the Google news corpus with over 300 billion words

Language Sequence Modeling Tasks



Sequence Modeling: Sequence Label Prediction



Masterful!

By Antony Witheyman - January 12, 2006

Ideal for anyone with an interest in
disguises who likes to see the subject
tackled in a humorous manner.

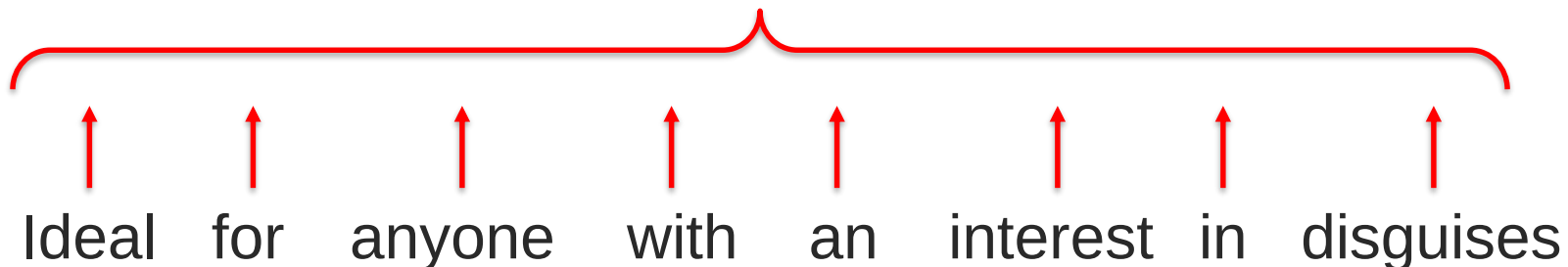
0 of 4 people found this review helpful

Prediction



Sentiment ?
(positive or negative)

Sentiment label?



Sequence Modeling: Sequence Prediction

★★★★★ **Masterful!**

By Antony Witheyman - January 12, 2006

Ideal for anyone with an interest in
disguises who likes to see the subject
tackled in a humorous manner.

0 of 4 people found this review helpful

Prediction



Part-of-speech ?
(noun, verb,...)

POS?

POS?

POS?

POS?

POS?

POS?

POS?

POS?



Ideal

for

anyone

with

an

interest

in

disguises



Sequence Modeling: Sequence Representation

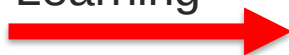
★★★★★ **Masterful!**

By Antony Witheyman - January 12, 2006

Ideal for anyone with an interest in disguises who likes to see the subject tackled in a humorous manner.

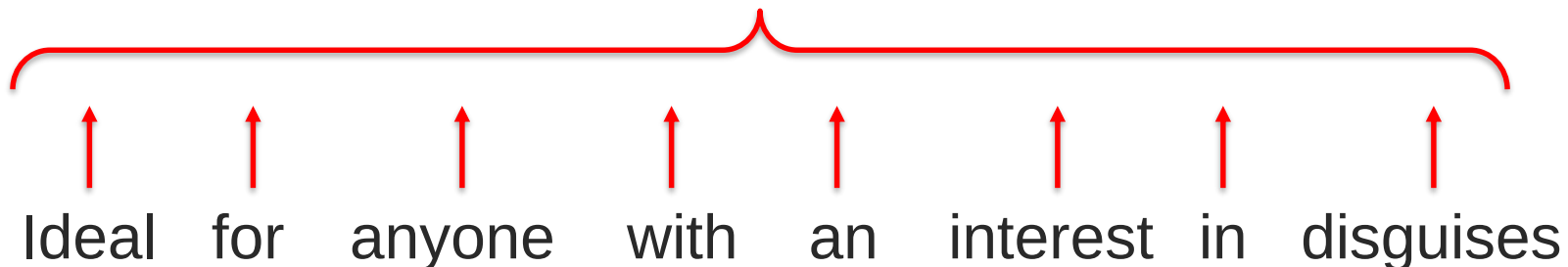
0 of 4 people found this review helpful

Learning



Sequence representation

[0,1; 0,0004; 0;...; 0,01; 0.09; 0,05]



Sequence Modeling: Language Model

★★★★★ **Masterful!**

By Antony Witheyman - January 12, 2006

Ideal for anyone with an interest in disguises who likes to see the subject tackled in a humorous manner.

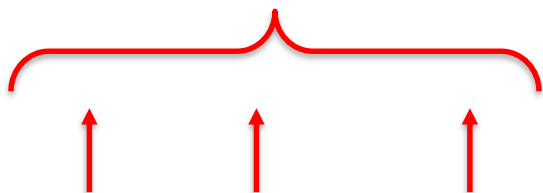
0 of 4 people found this review helpful

Prediction



Language Model

Next word?



Ideal for anyone

with an interest in disguises



Application: Speech Recognition

$$\arg \max_{wordsequence} P(wordsequence | acoustics) =$$

$$\arg \max_{wordsequence} \frac{P(acoustics | wordsequence) \times P(wordsequence)}{P(acoustics)}$$

$$\arg \max_{wordsequence} P(acoustics | wordsequence) \times P(wordsequence)$$



Language model



Application: Language Generation

Embedding

[0,1;
0,0004;
....;
0.09;
0,05]

Generation



Ideal for anyone with an interest in
disguises who likes to see the subject
tackled in a humourous manner.

Example: Image captioning



[0,1;
0,0004;
....;
0.09;
0,05]



The man at bat readies to swing at the
pitch while the umpire looks on.

N-Gram Language Model Formulations

- Word sequences

$$w_1^n = w_1 \dots w_n$$

- Chain rule of probability

$$P(w_1^n) = P(w_1)P(w_2 | w_1)P(w_3 | w_1^2) \dots P(w_n | w_1^{n-1}) = \prod_{k=1}^n P(w_k | w_1^{k-1})$$

- Bigram approximation

$$P(w_1^n) = \prod_{k=1}^n P(w_k | w_{k-1})$$

- N-gram approximation

$$P(w_1^n) = \prod_{k=1}^n P(w_k | w_{k-N+1}^{k-1})$$



Evaluating Language Model: Perplexity

The best language model is one that best predicts an unseen test set

- Gives the highest $P(\text{sentence})$

Perplexity is the inverse probability of the test set, normalized by the number of words:

Chain rule:

For bigrams:

$$PP(W) = P(w_1 w_2 \dots w_N)^{-\frac{1}{N}}$$

$$= \sqrt[N]{\frac{1}{P(w_1 w_2 \dots w_N)}}$$

$$PP(W) = \sqrt[N]{\prod_{i=1}^N \frac{1}{P(w_i | w_1 \dots w_{i-1})}}$$

$$PP(W) = \sqrt[N]{\prod_{i=1}^N \frac{1}{P(w_i | w_{i-1})}}$$

Challenges in Sequence Modeling

★★★★★ **Masterful!**

By Antony Witheyman - January 12, 2006

Ideal for anyone with an interest in disguises who likes to see the subject tackled in a humorous manner.

0 of 4 people found this review helpful

Model

- Part-of-speech ?
(noun, verb,...)
- Sentiment ?
(positive or negative)
- Language Model
- Sequence representation

Main Challenges:

- Sequences of variable lengths (e.g., sentences)
- Keep the number of parameters at a minimum
- Take advantage of possible redundancy

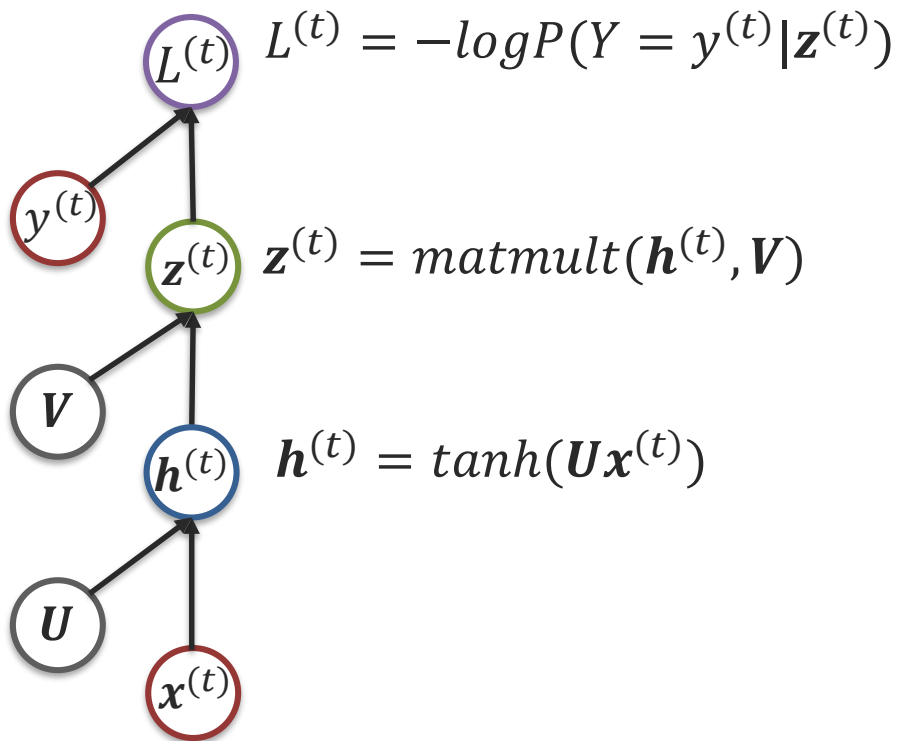


Recurrent Neural Networks



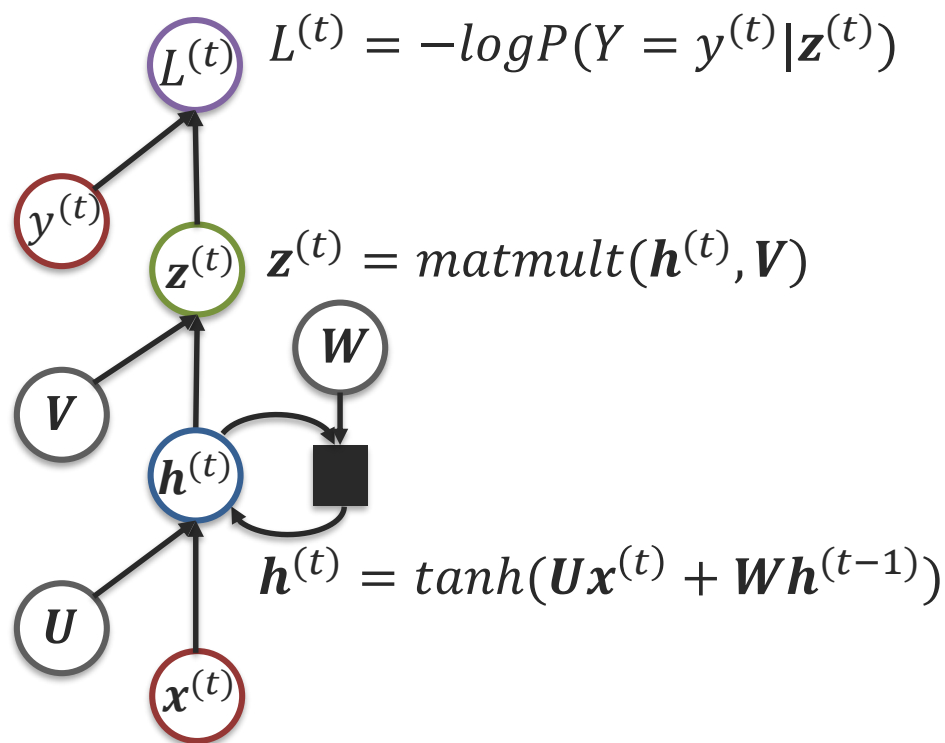
Recurrent Neural Network

Feedforward Neural Network



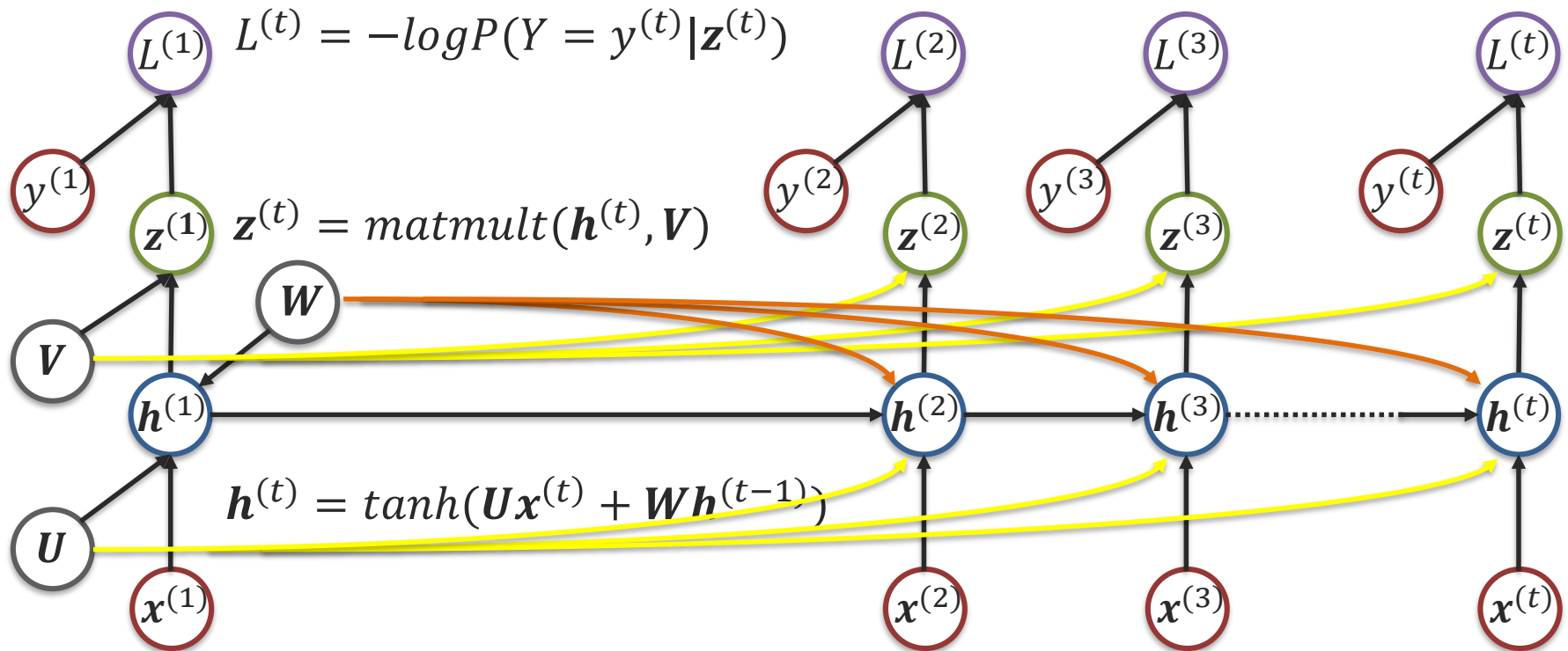
Recurrent Neural Networks

$$L = \sum_t L^{(t)}$$



Recurrent Neural Networks - Unrolling

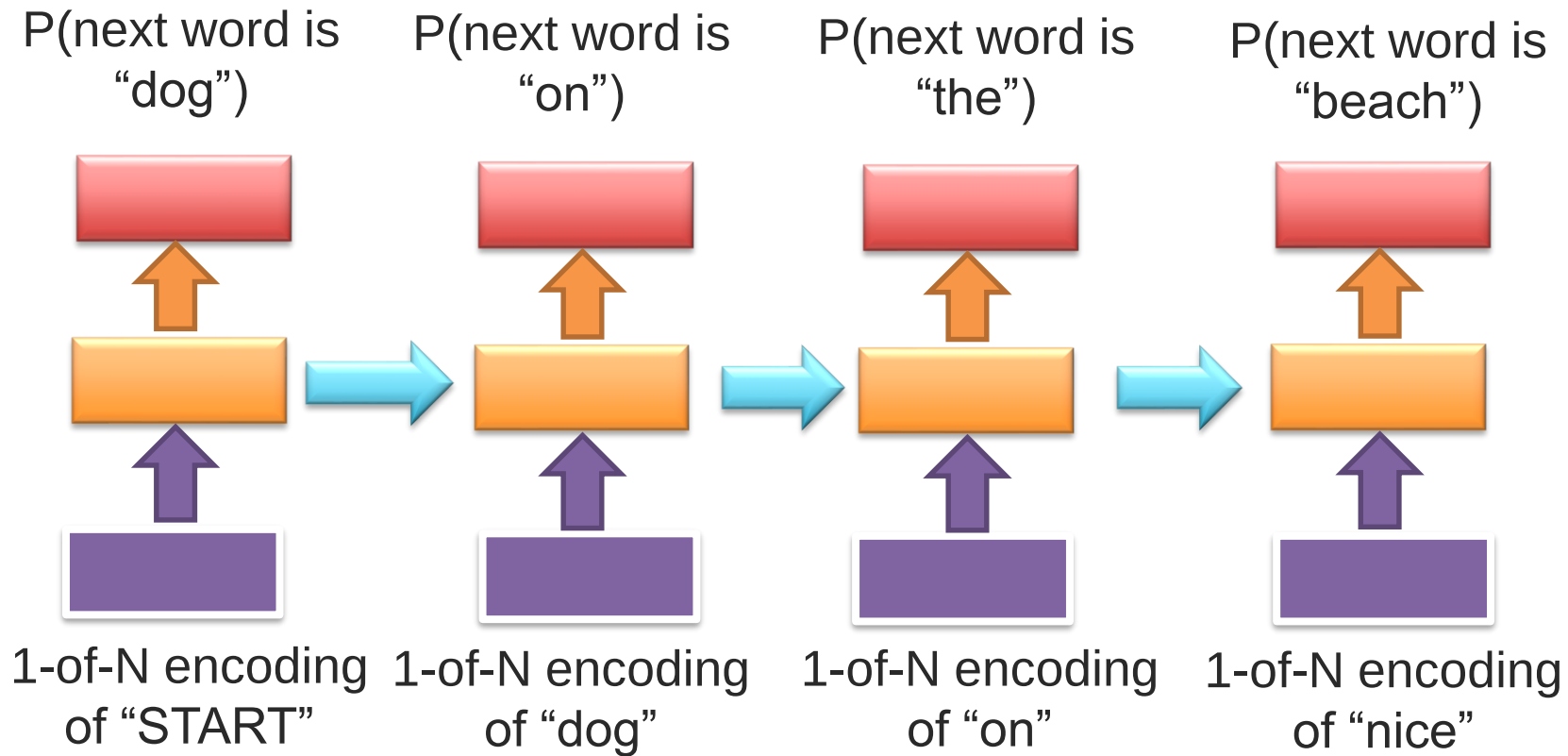
$$L = \sum_t L^{(t)}$$



Same model parameters are used for all time parts.

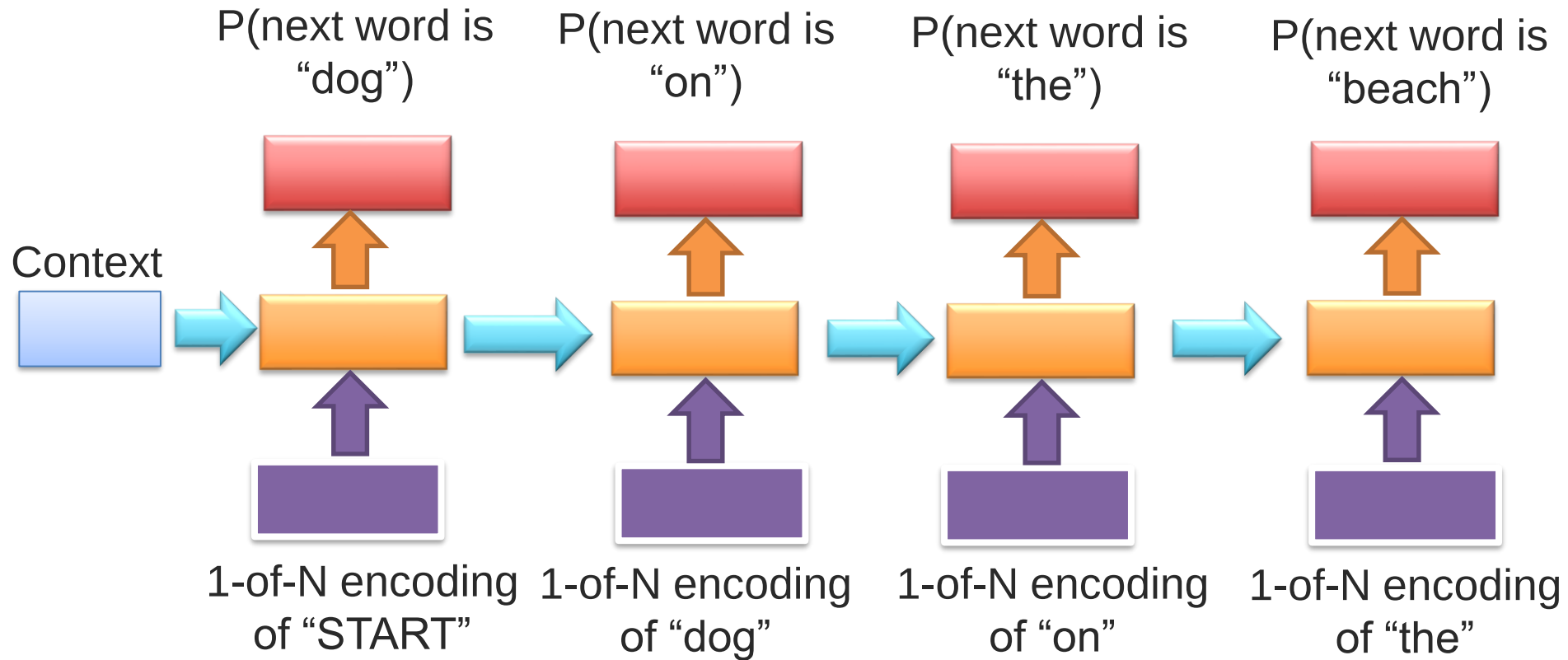


RNN-based Language Model



➤ Models long-term information

RNN-based Sentence Generation (Decoder)



➤ Models long-term information

Sequence Modeling: Sequence Prediction

★★★★★ **Masterful!**

By Antony Witheyman - January 12, 2006

Ideal for anyone with an interest in
disguises who likes to see the subject
tackled in a humorous manner.

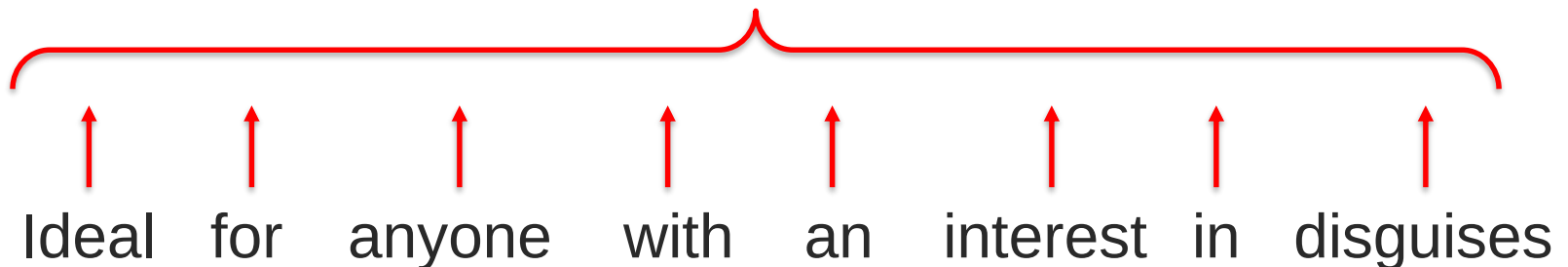
0 of 4 people found this review helpful

Prediction

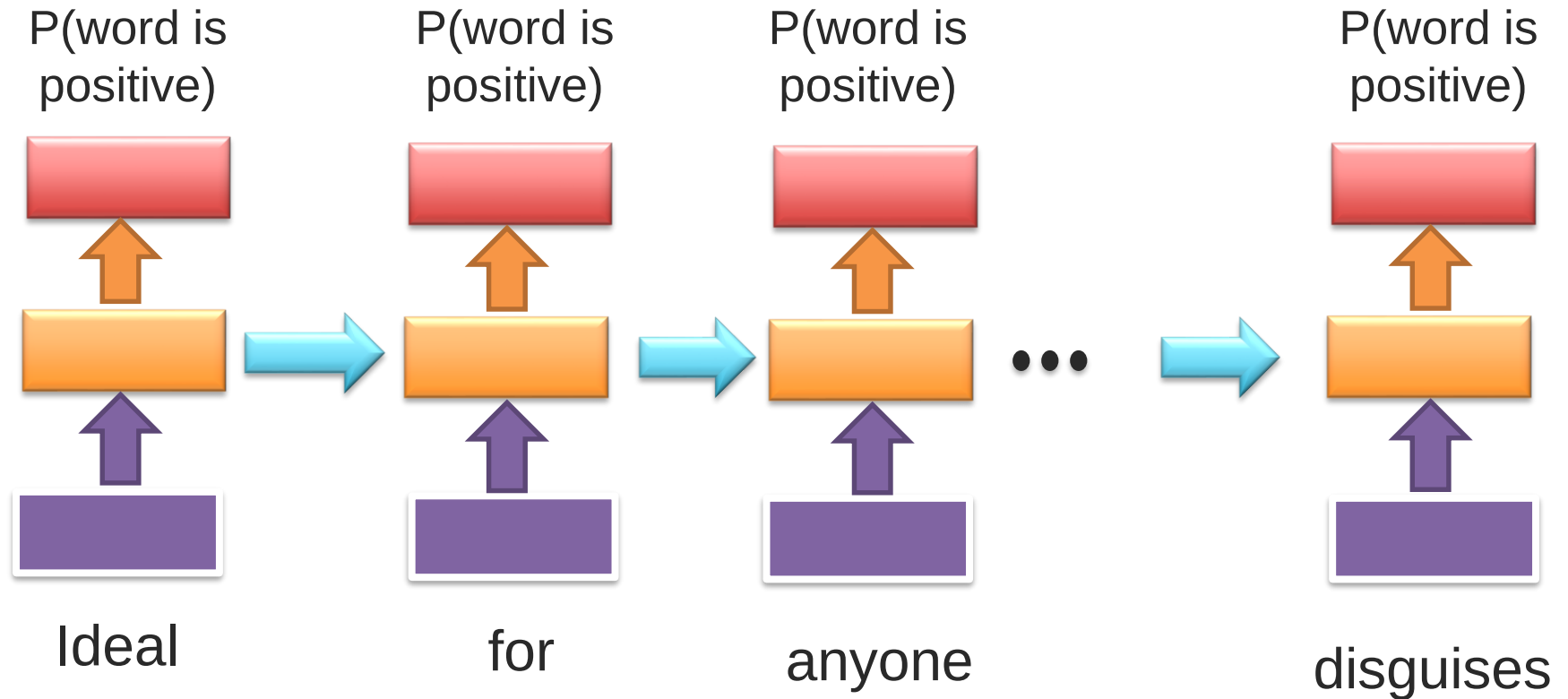


Sentiment ?
(positive or negative)

Sentiment label?

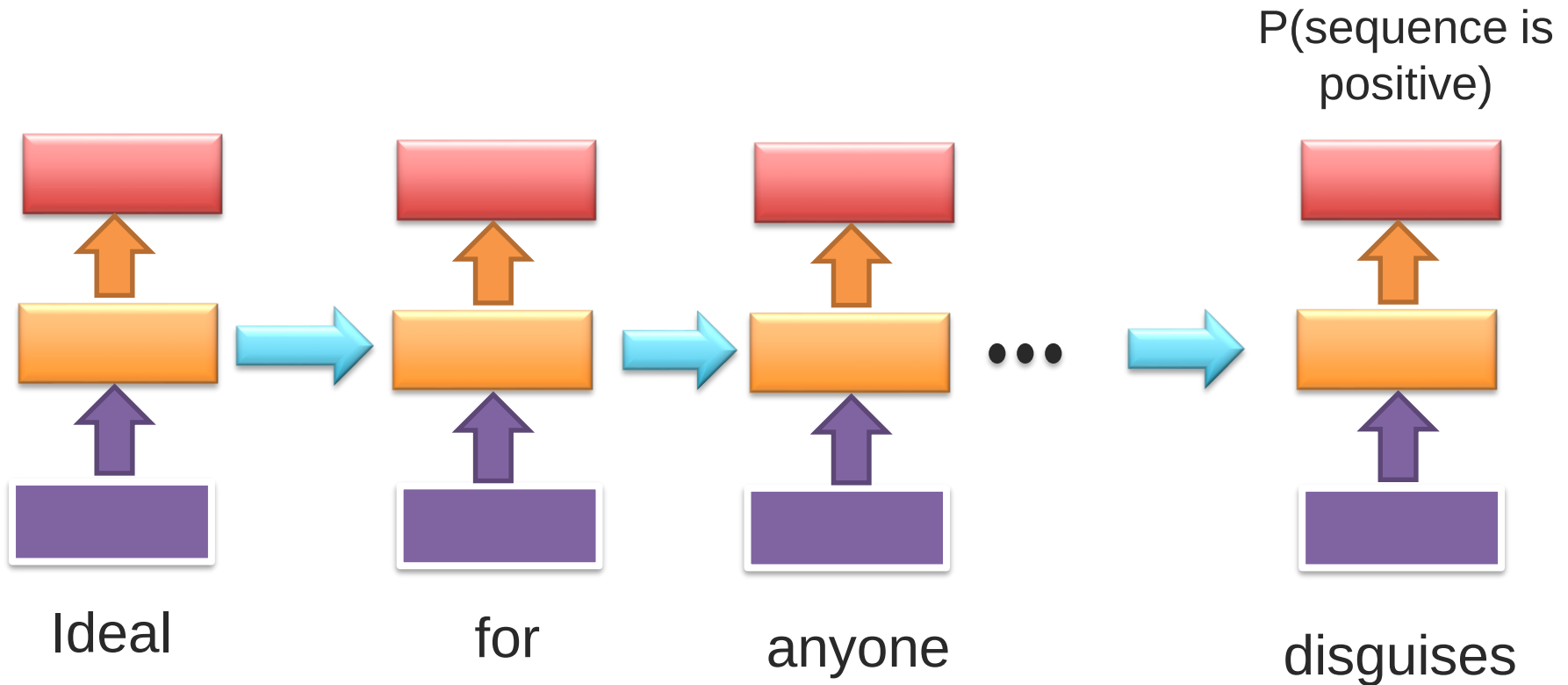


RNN for Sequence Prediction



$$L = \frac{1}{N} \sum_t L^{(t)} = \frac{1}{N} \sum_t -\log P(Y = y^{(t)} | \mathbf{z}^{(t)})$$

RNN for Sequence Prediction



$$L = L^{(N)} = -\log P(Y = y^{(N)} | \mathbf{z}^{(N)})$$

Sequence Modeling: Sequence Representation

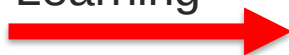
★★★★★ **Masterful!**

By Antony Witheyman - January 12, 2006

Ideal for anyone with an interest in disguises who likes to see the subject tackled in a humorous manner.

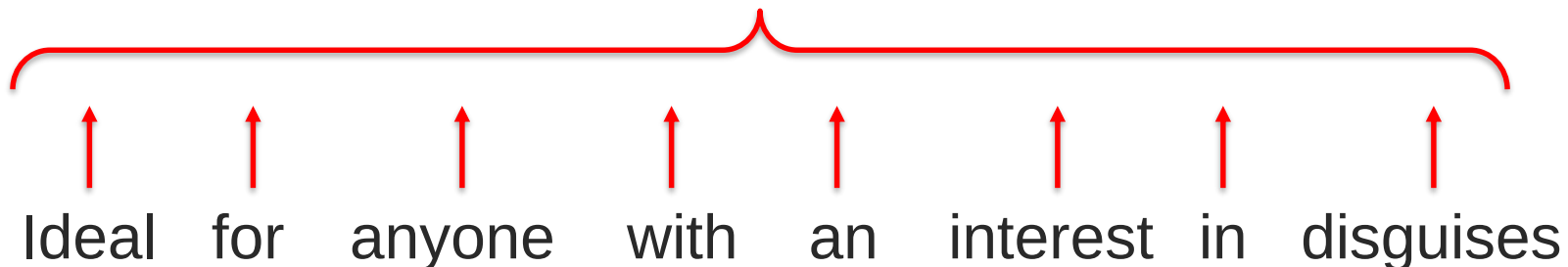
0 of 4 people found this review helpful

Learning

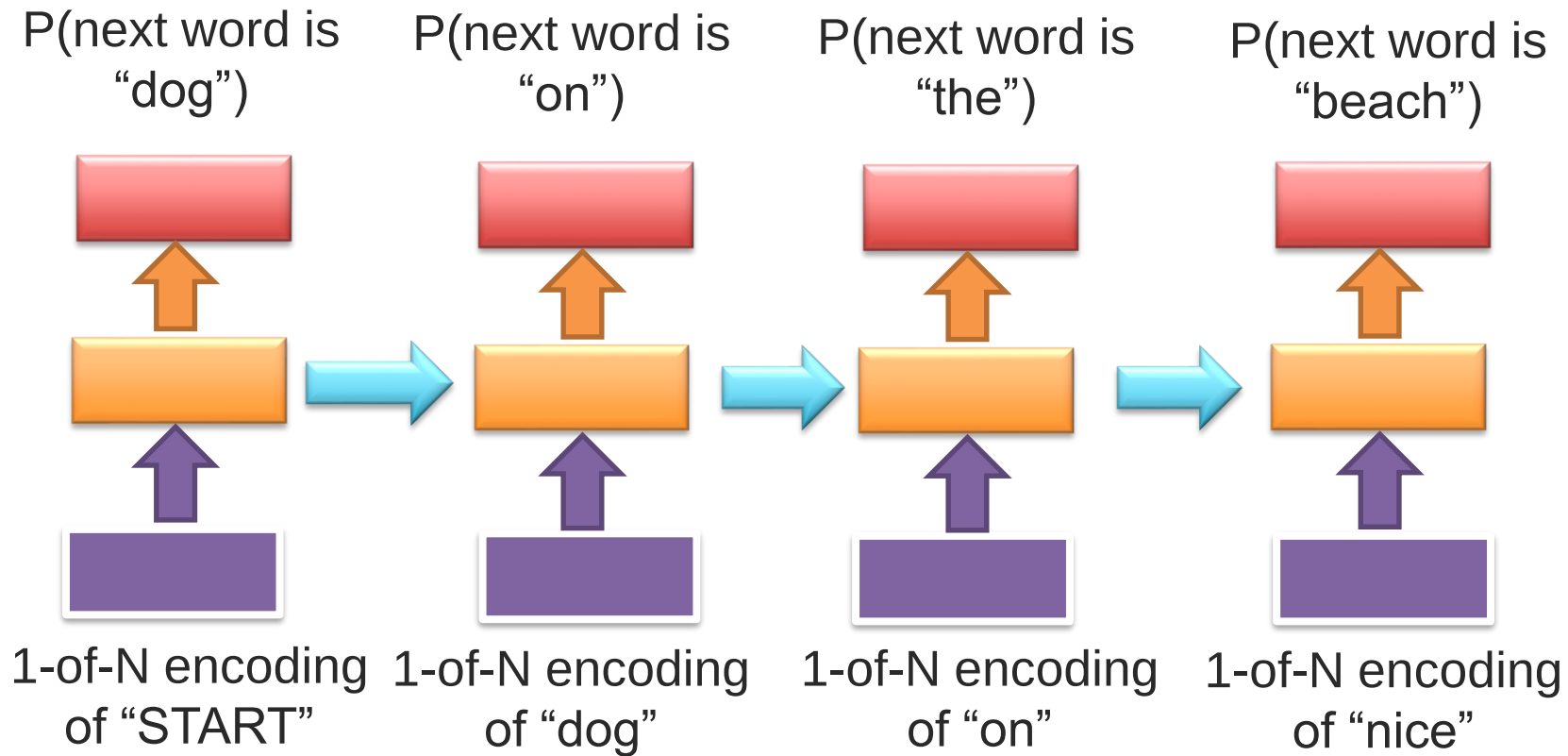


Sequence representation

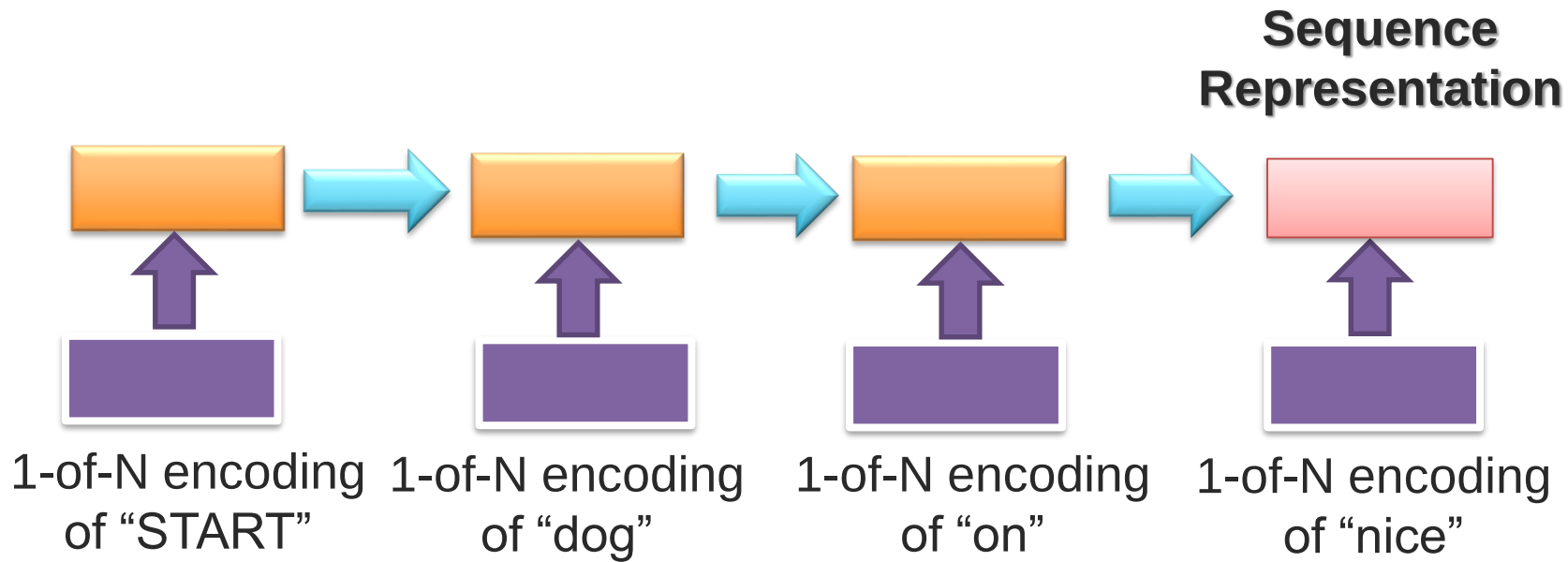
[0,1; 0,0004; 0;...; 0,01; 0.09; 0,05]



RNN for Sequence Representation

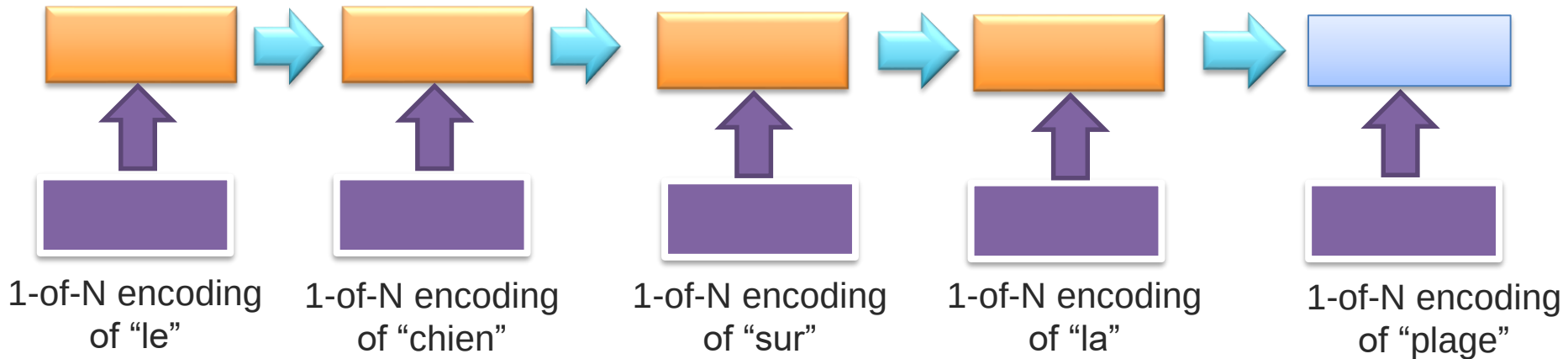


RNN for Sequence Representation (Encoder)

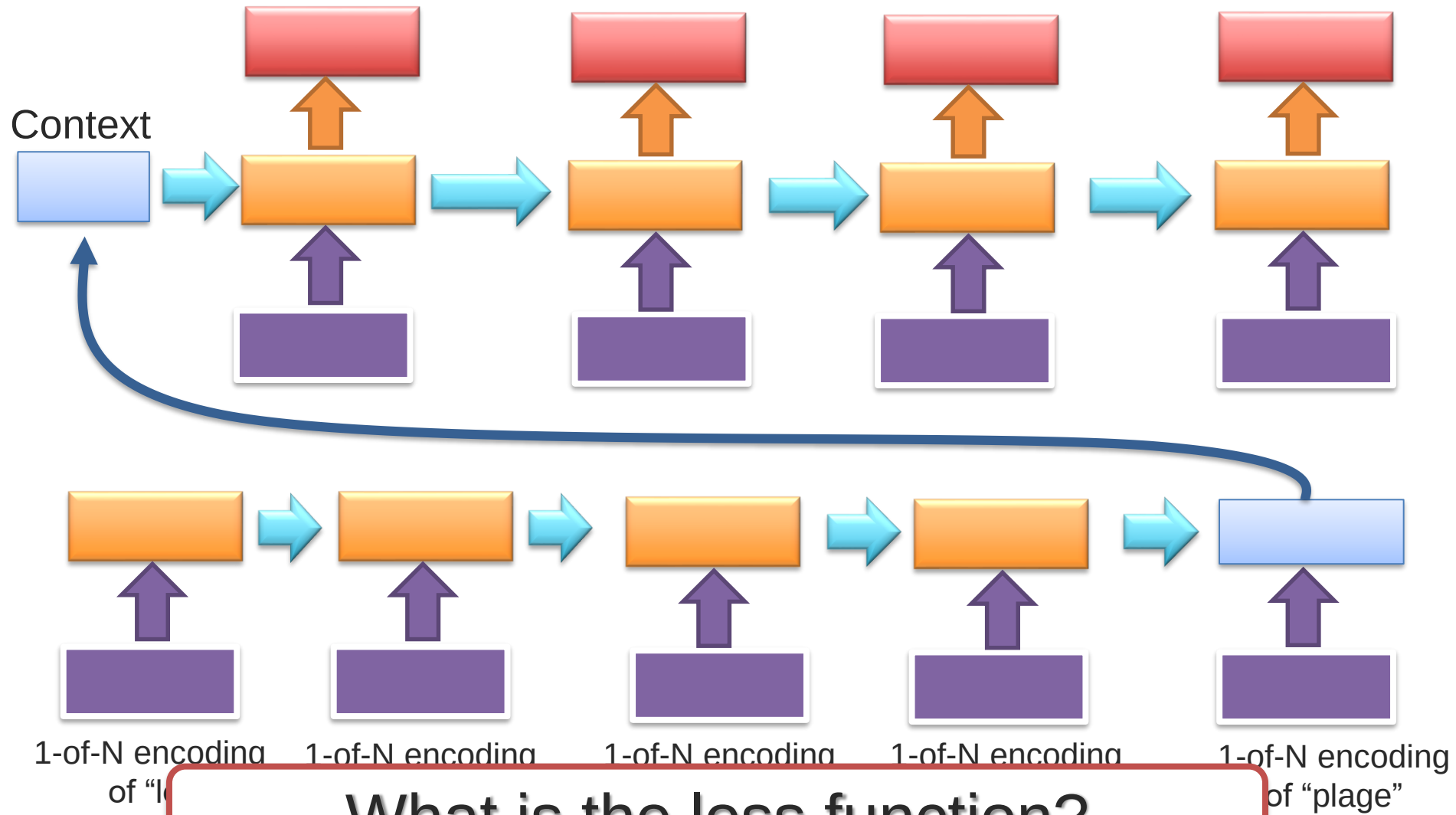


RNN-based for Machine Translation

Le chien sur la plage → The dog on the beach



Encoder-Decoder Architecture



What is the loss function?

Related Topics

- Character-level “language models”
 - Xiang Zhang, Junbo Zhao and Yann LeCun, Character-level Convolutional Networks for Text Classification, NIPS 2015
<http://arxiv.org/pdf/1509.01626v2.pdf>
- Skip-thought: embedding at the sentence level
 - Ryan Kiros, Yukun Zhu, Ruslan Salakhutdinov, Richard S. Zemel, Antonio Torralba, Raquel Urtasun, Sanja Fidler. Skip-Thought Vectors, NIPS 2015
<http://arxiv.org/pdf/1506.06726v1.pdf>

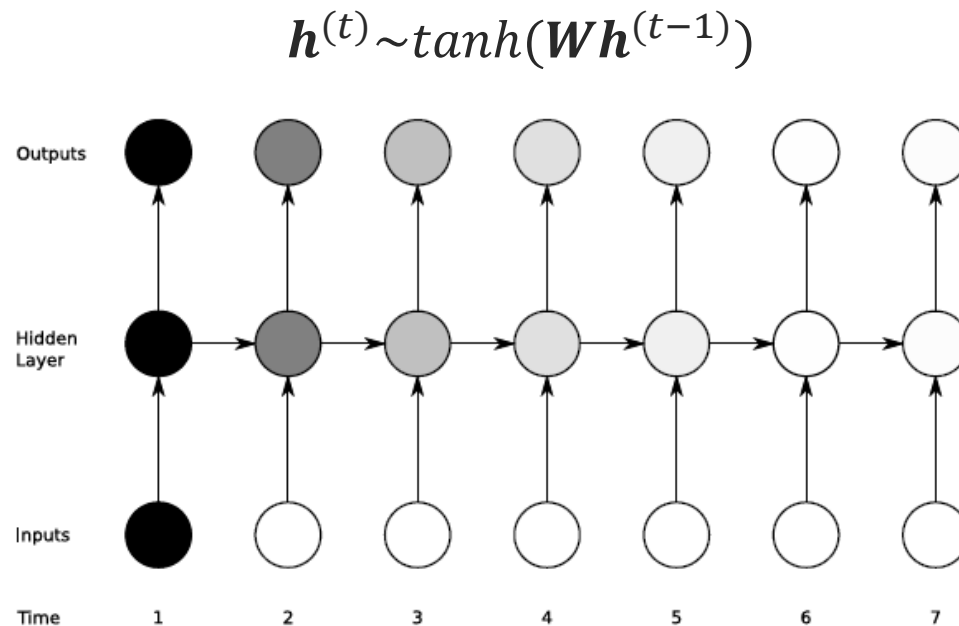


Gated Recurrent Neural Networks



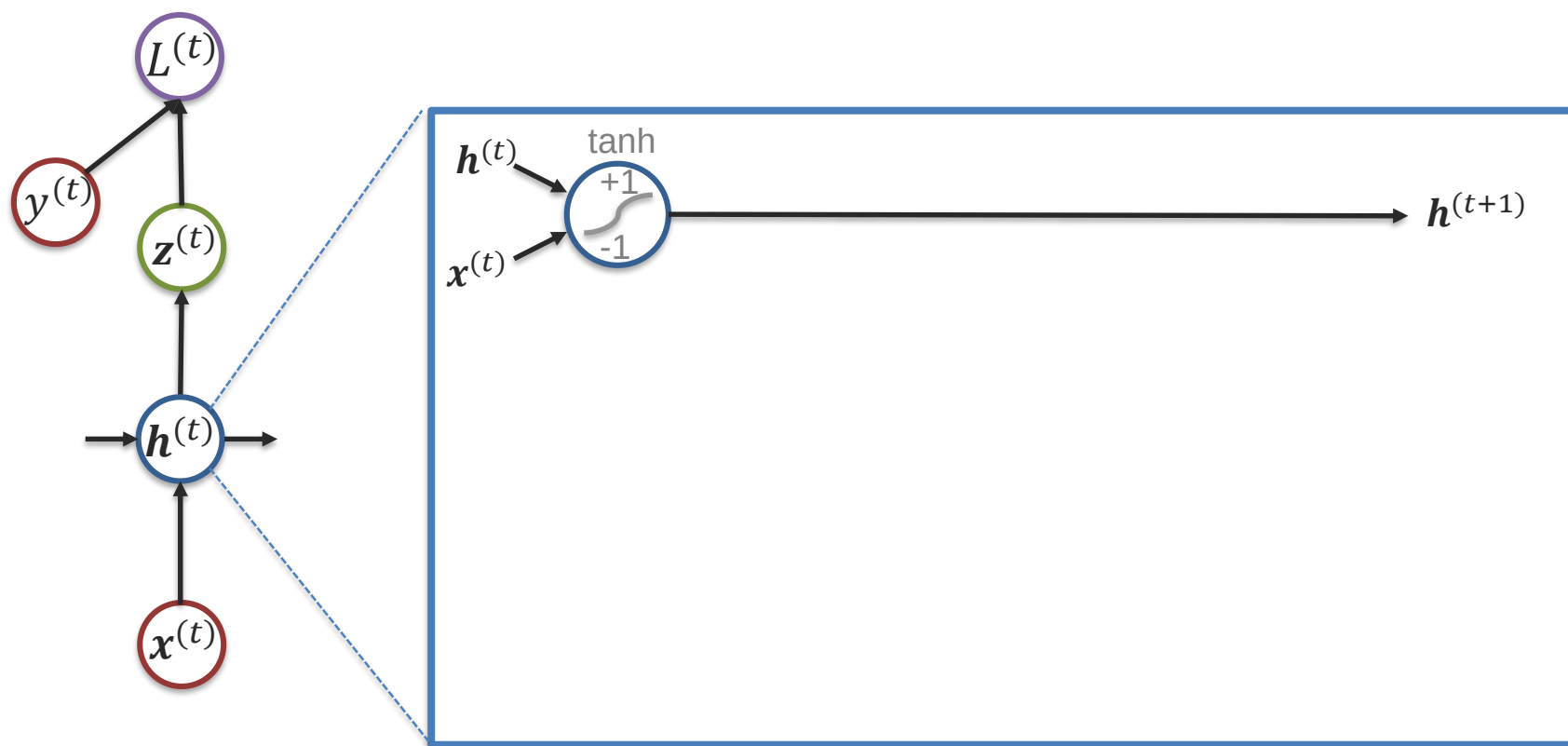
Long-term Dependencies

Vanishing gradient problem for RNNs:



- The influence of a given input on the hidden layer, and therefore on the network output, either decays or blows up exponentially as it cycles around the network's recurrent connections.

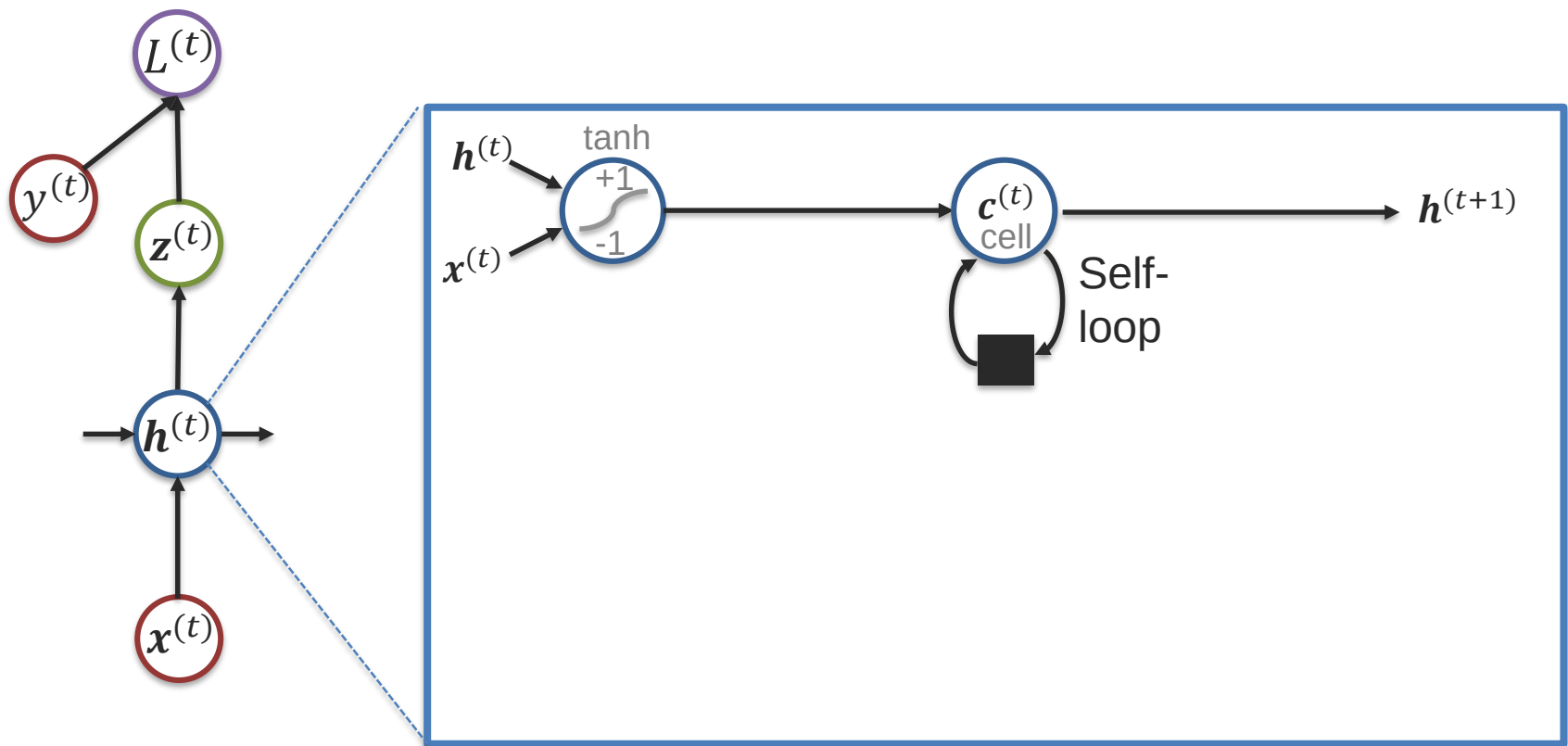
Recurrent Neural Networks



LSTM ideas: (1) “Memory” Cell and Self Loop

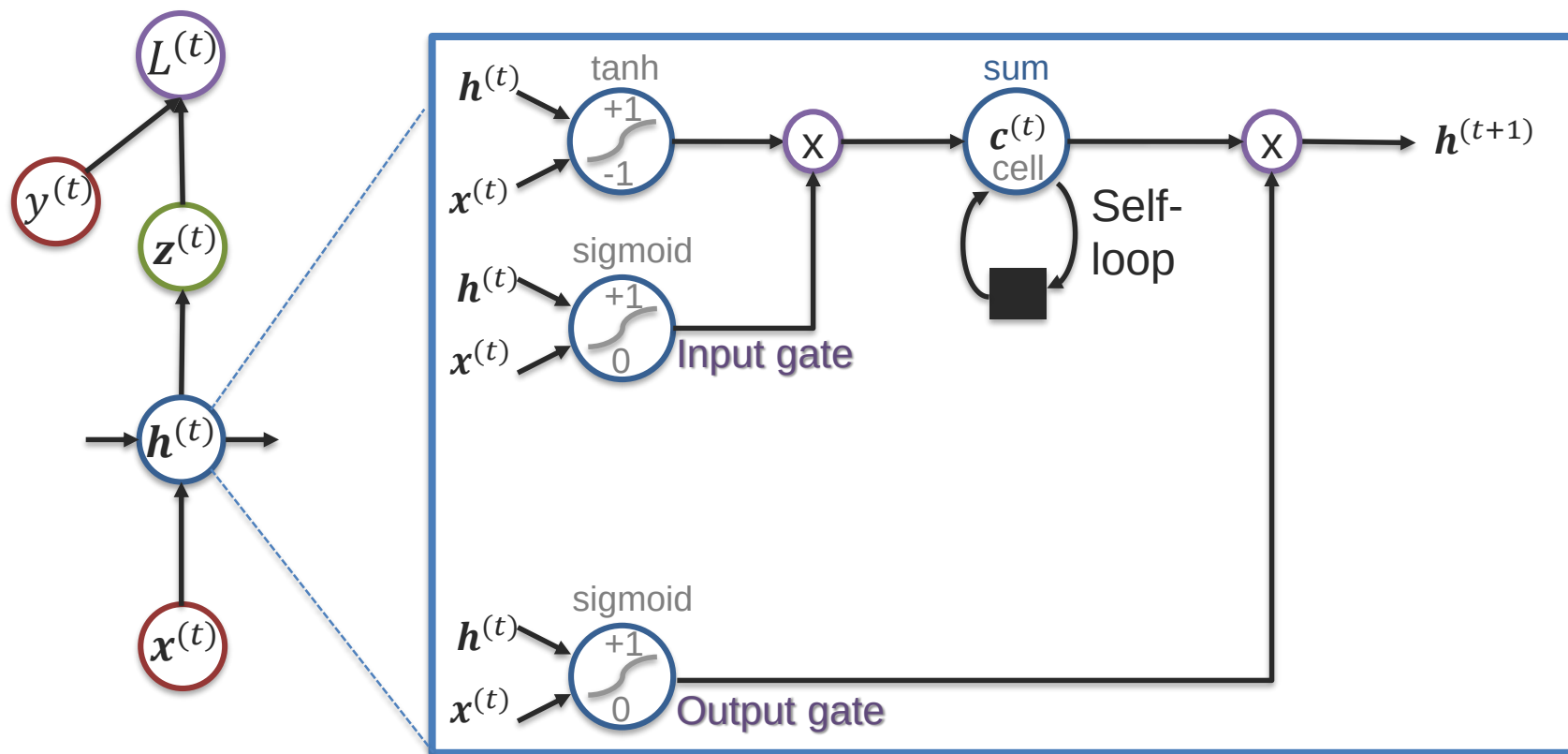
[Hochreiter and Schmidhuber, 1997]

Long Short-Term Memory (LSTM)



LSTM Ideas: (2) Input and Output Gates

[Hochreiter and Schmidhuber, 1997]

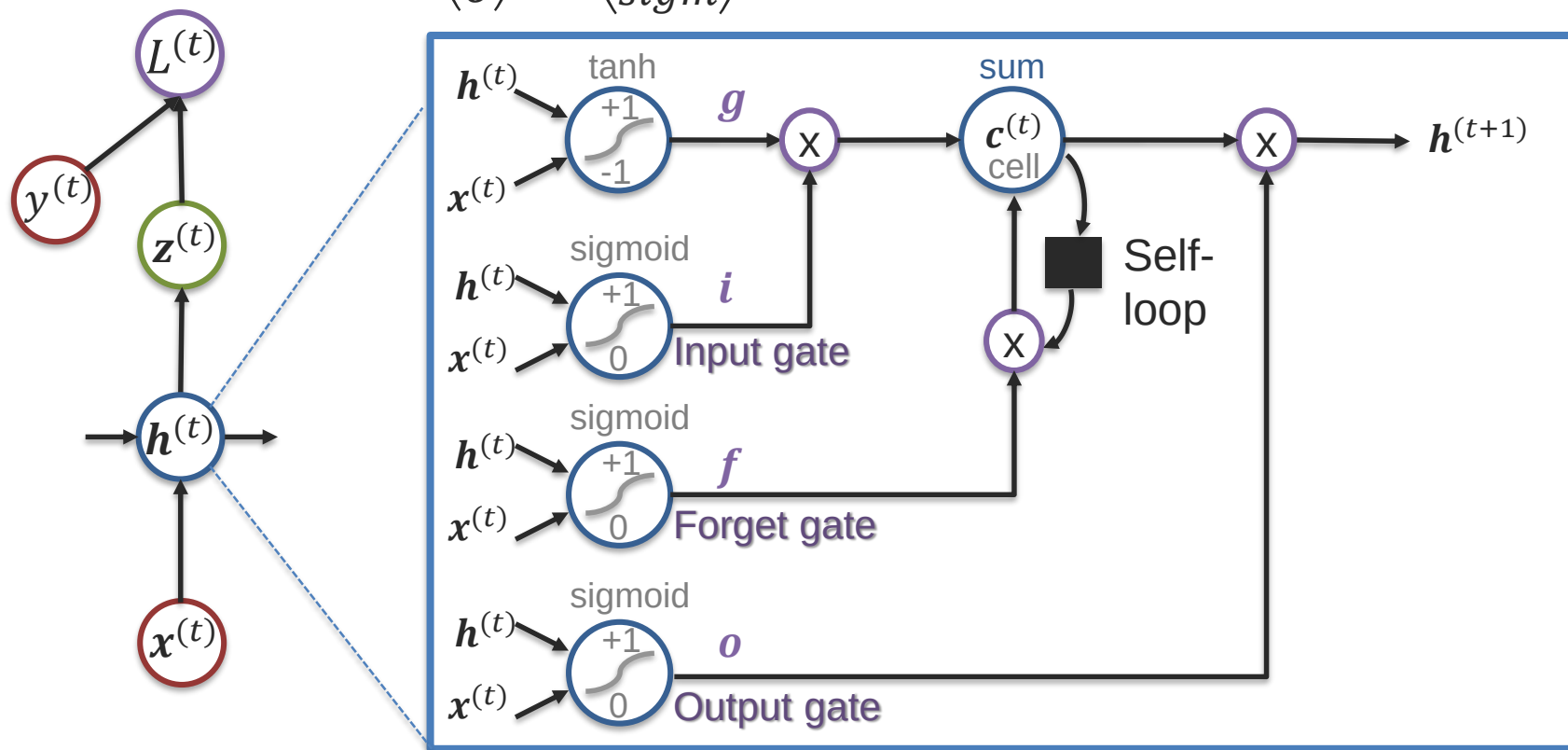


LSTM Ideas: (3) Forget Gate [Gers et al., 2000]

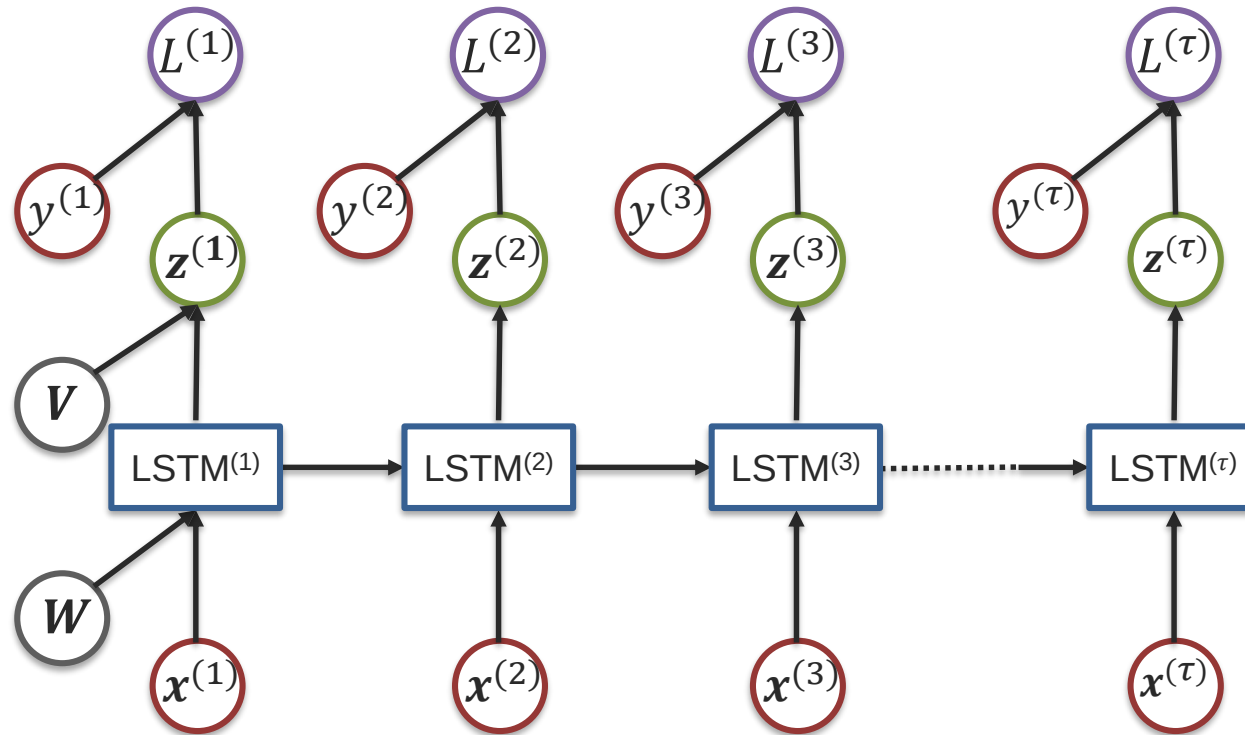
$$\begin{pmatrix} g \\ i \\ f \\ o \end{pmatrix} = \begin{pmatrix} \tanh \\ \text{sigm} \\ \text{sigm} \\ \text{sigm} \end{pmatrix} W \begin{pmatrix} h^{(t)} \\ x^{(t)} \end{pmatrix}$$

$$c^{(t)} = f \odot c^{(t-1)} + i \odot g$$

$$h^{(t)} = o \odot \tanh(c^{(t)})$$

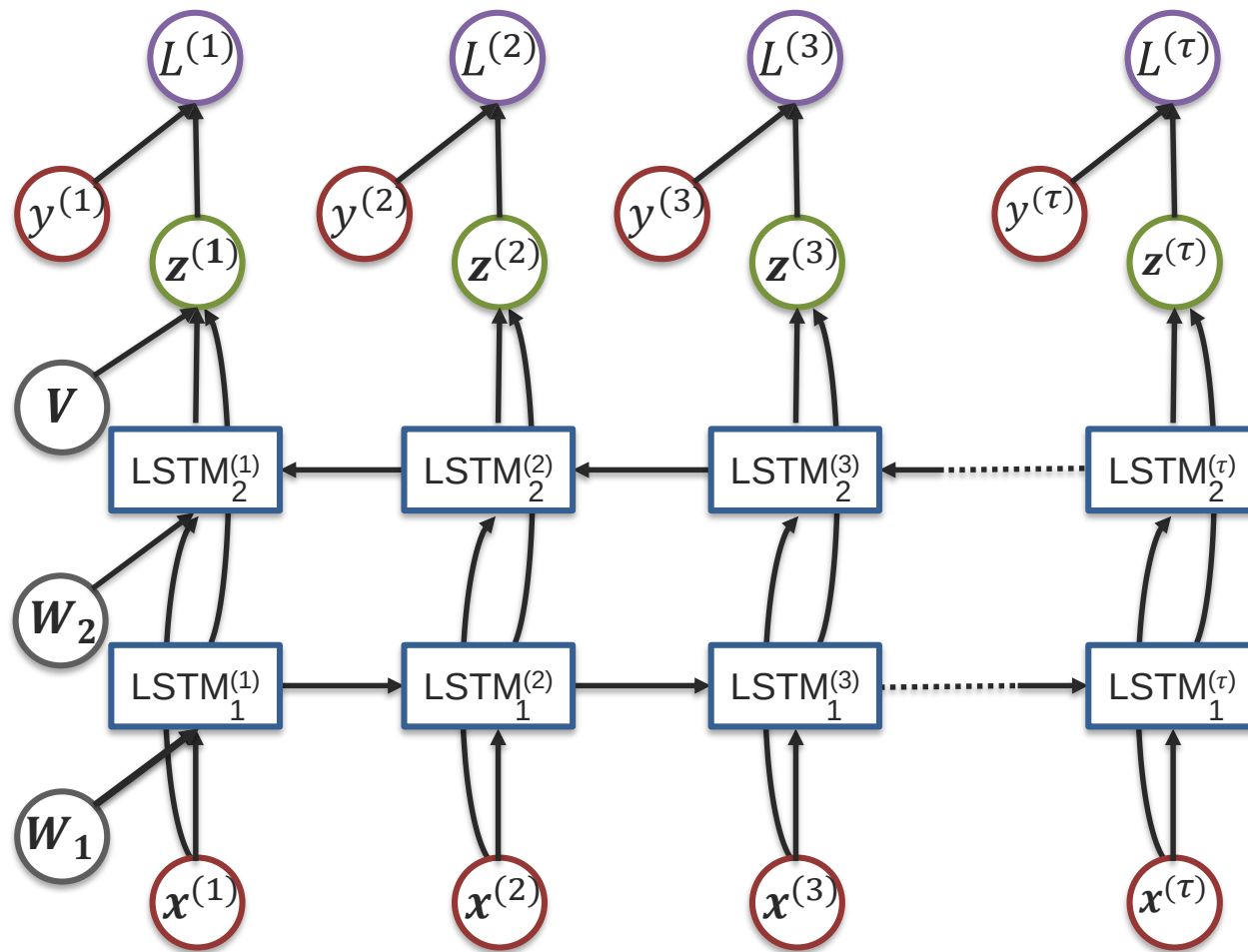


Recurrent Neural Network using LSTM Units

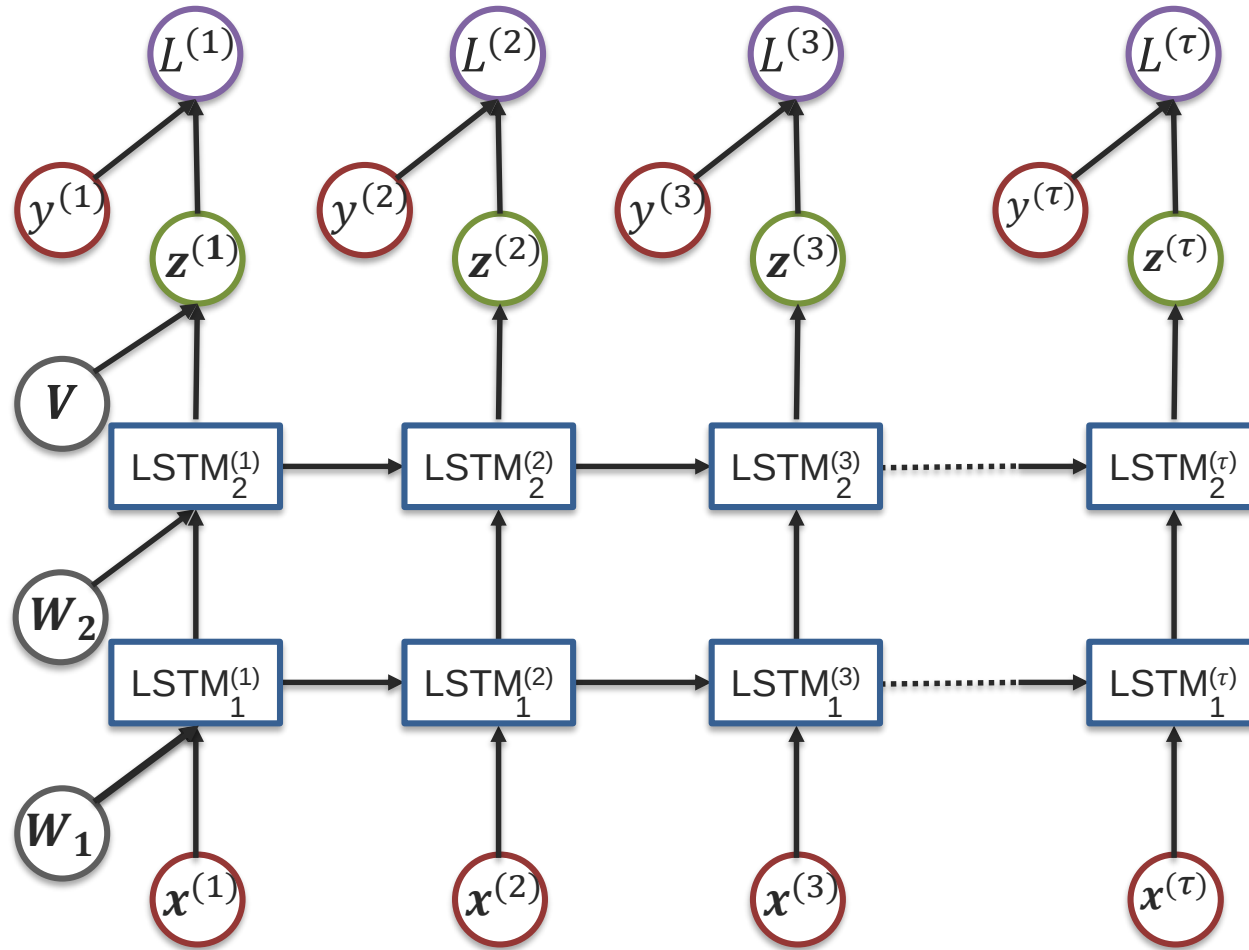


Gradient can still be computed using backpropagation!

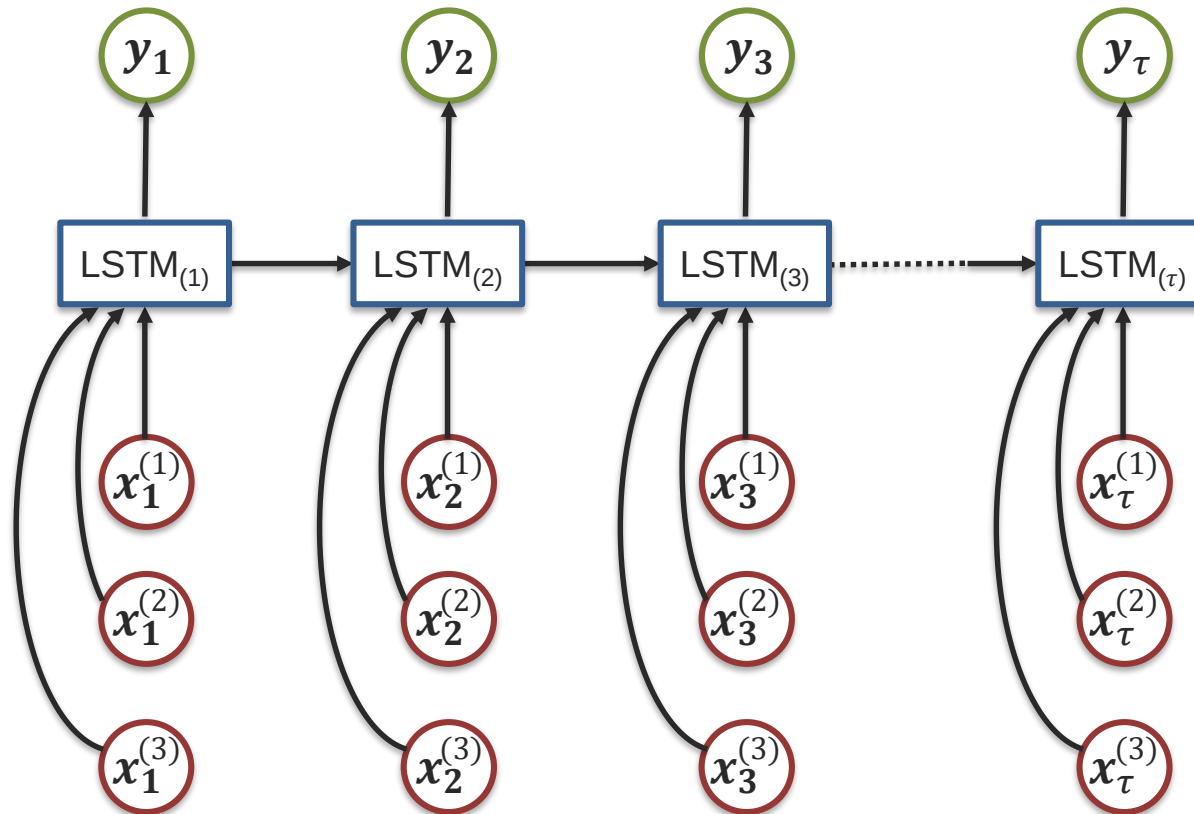
Bi-directional LSTM Network



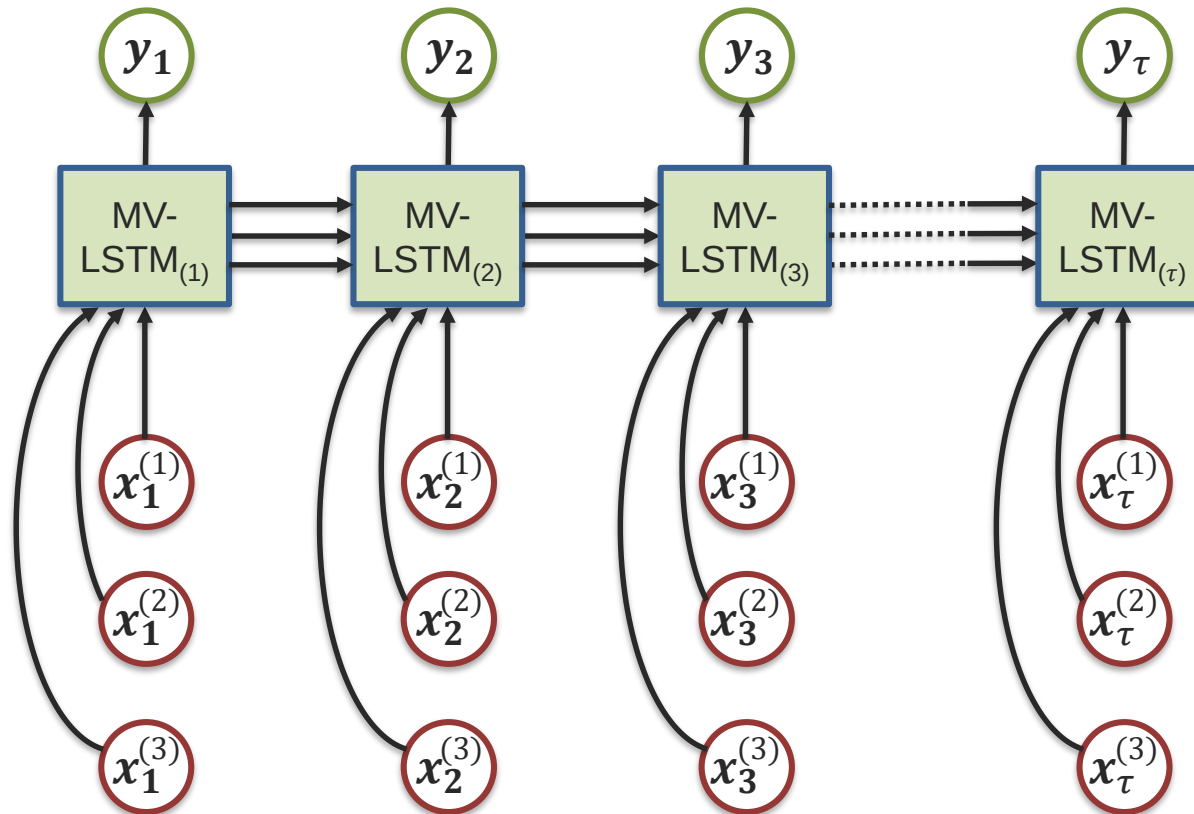
Deep LSTM Network



Multimodal Sequence Modeling – Early Fusion



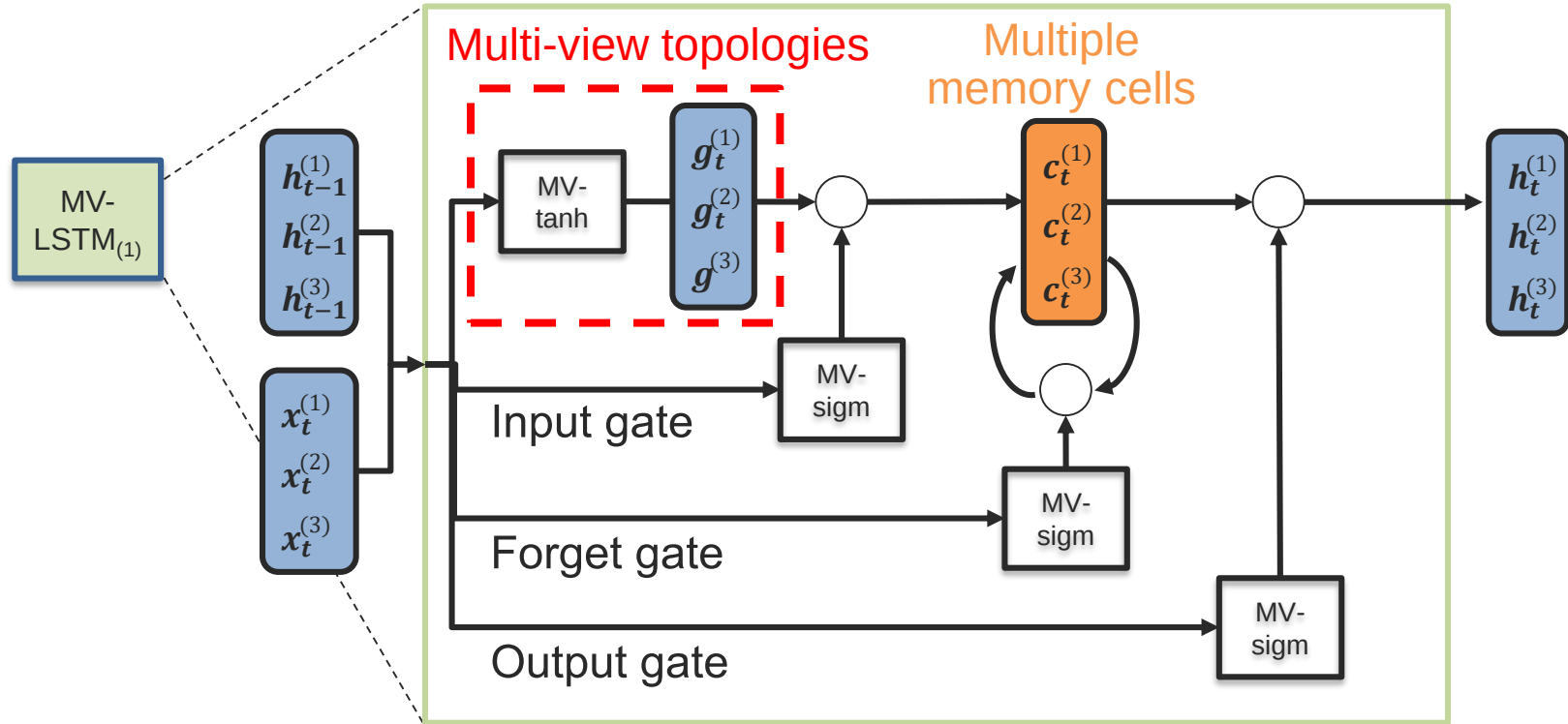
Multi-View Long Short-Term Memory (MV-LSTM)



[Shyam, Morency, et al. Extending Long Short-Term Memory for Multi-View Structured Learning, ECCV, 2016]



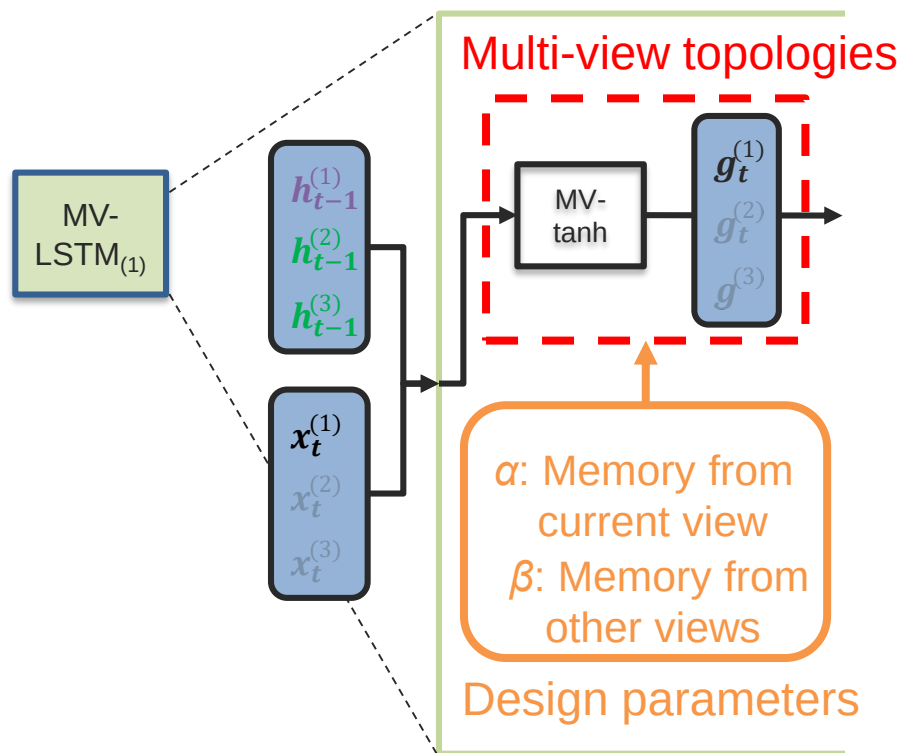
Multi-View Long Short-Term Memory



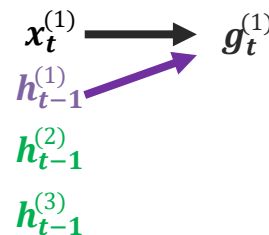
[Shyam, Morency, et al. Extending Long Short-Term Memory for Multi-View Structured Learning, ECCV, 2016]



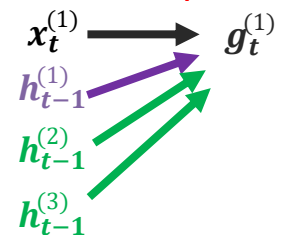
Topologies for Multi-View LSTM



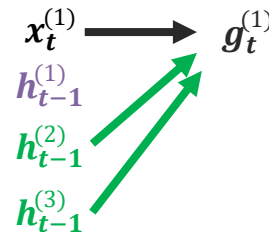
View-specific
 $\alpha=1, \beta=0$



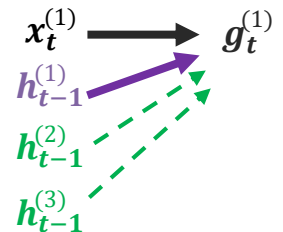
Fully-connected
 $\alpha=1, \beta=1$



Coupled
 $\alpha=0, \beta=1$



Hybrid
 $\alpha=2/3, \beta=1/3$



[Shyam, Morency, et al. Extending Long Short-Term Memory for Multi-View Structured Learning, ECCV, 2016]



Multi-View Long Short-Term Memory (MV-LSTM)

Multimodal prediction of children engagement

Class labels	Model	Precision	Recall	F1
Easy to engage	LSTM (Early fusion)	0.75	0.81	0.78
	MV-LSTM Full	0.81	0.81	0.81
	MV-LSTM Coupled	0.79	0.81	0.80
	MV-LSTM Hybrid	0.80	0.86	0.83
Difficult to engage	LSTM (Early fusion)	0.63	0.55	0.59
	MV-LSTM Full	0.68	0.68	0.68
	MV-LSTM Coupled	0.67	0.64	0.65
	MV-LSTM Hybrid	0.74	0.64	0.68

[Shyam, Morency, et al. Extending Long Short-Term Memory for Multi-View Structured Learning, *ECCV*, 2016]

Backpropagation Through Time



Optimization: Gradient Computation

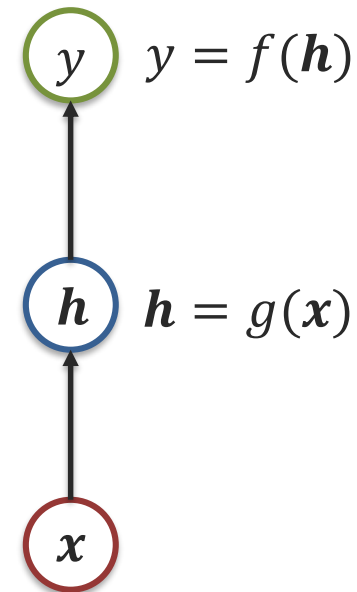
Vector representation:

Gradient $\nabla_x y = \left[\frac{\partial y}{\partial x_1}, \frac{\partial y}{\partial x_2}, \frac{\partial y}{\partial x_3} \right]$

$\nabla_x y = \left(\frac{\partial \mathbf{h}}{\partial \mathbf{x}} \right)^T \nabla_h y$

“local” Jacobian (matrix of size $|h| \times |x|$ computed using partial derivatives)

“backprop” Gradient



Backpropagation Algorithm

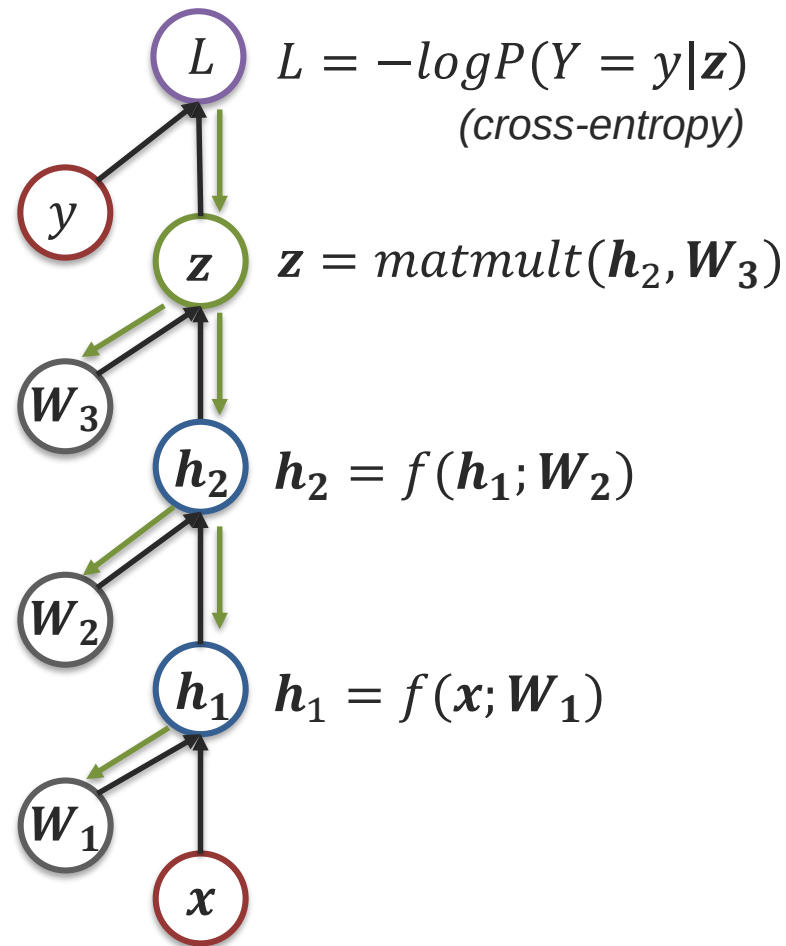
Forward pass

- Following the graph topology, compute value of each unit

Backpropagation pass

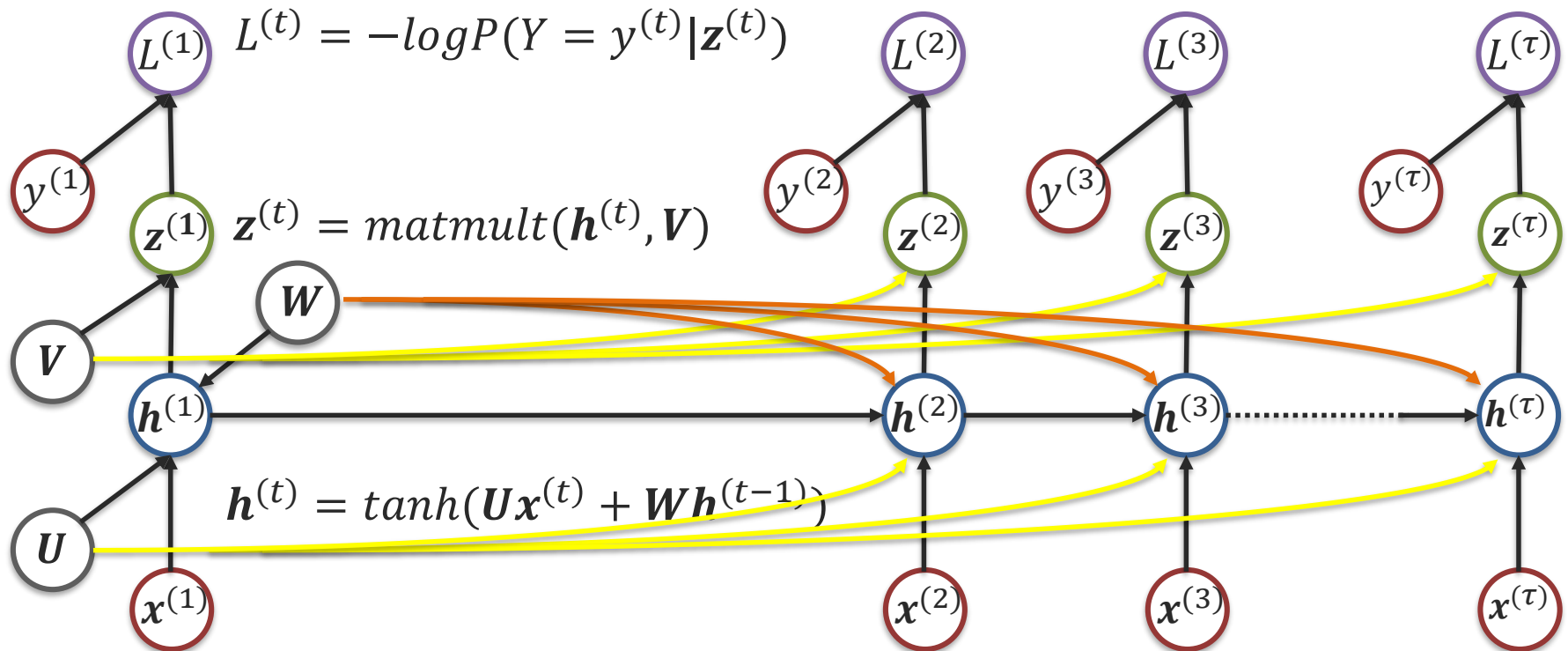
- Initialize output gradient = 1
- Compute “local” Jacobian matrix using values from forward pass
- Use the chain rule:

Gradient = “local” Jacobian $\times$
“backprop” gradient



Recurrent Neural Networks

$$L = \sum_t L^{(t)}$$



Backpropagation Through Time

$$L = \sum_t L^{(t)} = - \sum_t \log P(Y = y^{(t)} | \mathbf{z}^{(t)})$$

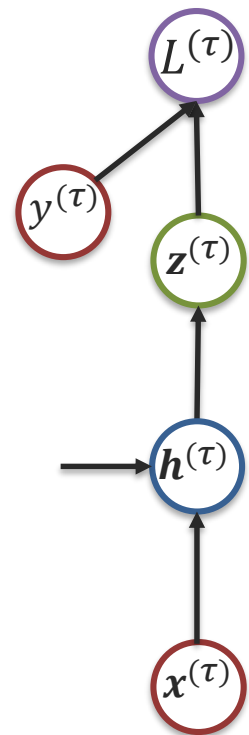
Gradient = “backprop” gradient
x “local” Jacobian

$$L^{(\tau)} \text{ or } L^{(t)} \quad \frac{\partial L}{\partial L^{(t)}} = 1$$

$$\mathbf{z}^{(\tau)} \text{ or } \mathbf{z}^{(t)} \quad (\nabla_{\mathbf{z}^{(t)}} L)_i = \frac{\partial L}{\partial z_i^{(t)}} = \frac{\partial L}{\partial L^{(t)}} \frac{\partial L^{(t)}}{\partial z_i^{(t)}} = \text{sigmoid}(z_i^t) - \mathbf{1}_{i,y^{(t)}}$$

$$\mathbf{h}^{(\tau)} \quad \nabla_{\mathbf{h}^{(\tau)}} L = \nabla_{\mathbf{z}^{(\tau)}} L \frac{\partial \mathbf{z}^{(\tau)}}{\partial \mathbf{h}^{(\tau)}} = \nabla_{\mathbf{z}^{(\tau)}} L \mathbf{V}$$

$$\mathbf{h}^{(t)} \rightarrow \mathbf{h}^{(t+1)} \quad \nabla_{\mathbf{h}^{(t)}} L = \nabla_{\mathbf{z}^{(t)}} L \frac{\partial \mathbf{o}^{(t)}}{\partial \mathbf{h}^{(t)}} + \nabla_{\mathbf{z}^{(t+1)}} L \frac{\partial \mathbf{h}^{(t+1)}}{\partial \mathbf{h}^{(t)}}$$



Backpropagation Through Time

$$L = \sum_t L^{(t)} = - \sum_t \log P(Y = y^{(t)} | \mathbf{z}^{(t)})$$

Gradient = “backprop” gradient
x “local” Jacobian

$$\textcircled{V} \quad \nabla_V L = \sum_t (\nabla_{\mathbf{z}^{(t)}} L) \frac{\partial \mathbf{z}^{(t)}}{\partial V}$$

$$\textcircled{W} \quad \nabla_W L = \sum_t (\nabla_{\mathbf{h}^{(t)}} L) \frac{\partial \mathbf{h}^{(t)}}{\partial W}$$

$$\textcircled{U} \quad \nabla_U L = \sum_t (\nabla_{\mathbf{h}^{(t)}} L) \frac{\partial \mathbf{h}^{(t)}}{\partial U}$$

