

CS 6770 Natural Language Processing

Recurrent Neural Networks Language Models

Yangfeng Ji

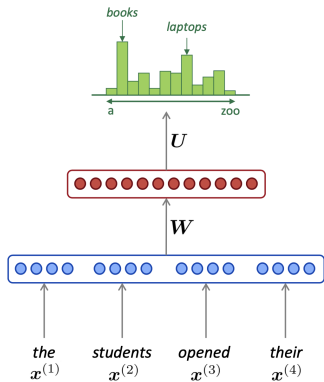
Information and Language Processing Lab
Department of Computer Science
University of Virginia



1. Neural Network Language Models
2. Recurrent Neural Networks
3. Computation Graph
4. RNN Language Modeling
5. Challenge of Training RNNs

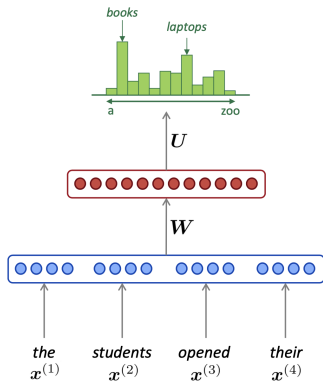
Neural Network Language Models

A Neural Language Model with Fixed Window Size



► Word indices: x_1, x_2, x_3, x_4

A Neural Language Model with Fixed Window Size

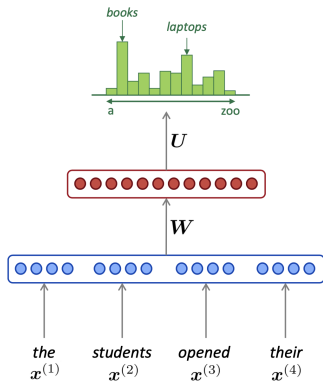


- Concatenated word embeddings

$$v = [v_{x_1}, v_{x_2}, v_{x_3}, v_{x_4}] \quad (3)$$

- Word indices: x_1, x_2, x_3, x_4

A Neural Language Model with Fixed Window Size



- ▶ Hidden layer: $f(\cdot)$ could be any nonlinear activation function

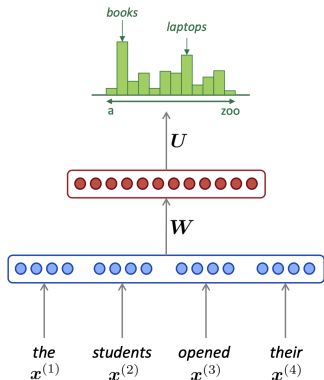
$$h = f(Wv + b_1) \quad (2)$$

- ▶ Concatenated word embeddings

$$v = [v_{x_1}, v_{x_2}, v_{x_3}, v_{x_4}] \quad (3)$$

- ▶ Word indices: x_1, x_2, x_3, x_4

A Neural Language Model with Fixed Window Size



- Output distribution

$$P(X_5 | X_{1:4}) = \text{softmax}(\mathbf{U}\mathbf{h} + \mathbf{b}_2) \in \mathbb{R}^{|\mathcal{V}|} \quad (1)$$

- Hidden layer: $f(\cdot)$ could be any nonlinear activation function

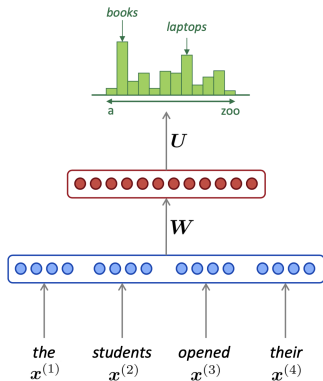
$$\mathbf{h} = f(\mathbf{W}\mathbf{v} + \mathbf{b}_1) \quad (2)$$

- Concatenated word embeddings

$$\mathbf{v} = [\mathbf{v}_{x_1}, \mathbf{v}_{x_2}, \mathbf{v}_{x_3}, \mathbf{v}_{x_4}] \quad (3)$$

- Word indices: x_1, x_2, x_3, x_4

A Neural Language Model with Fixed Window Size



- Output distribution

$$P(X_5 | X_{1:4}) = \text{softmax}(\mathbf{U}\mathbf{h} + \mathbf{b}_2) \in \mathbb{R}^{|\mathcal{V}|} \quad (1)$$

- Hidden layer: $f(\cdot)$ could be any nonlinear activation function

$$\mathbf{h} = f(\mathbf{W}\mathbf{v} + \mathbf{b}_1) \quad (2)$$

- Concatenated word embeddings

$$\mathbf{v} = [v_{x_1}, v_{x_2}, v_{x_3}, v_{x_4}] \quad (3)$$

- Word indices: x_1, x_2, x_3, x_4

This is the very first neural neural language model

[Bengio et al., 2001], which has a similar network architecture as the one discussed text classification.

A Neural Probabilistic Language Model

The first paragraph of the paper *A Neural Probabilistic Language Model* [Bengio et al., 2001]

1 Introduction

A fundamental problem that makes language modeling and other learning problems difficult is the *curse of dimensionality*. It is particularly obvious in the case when one wants to model the joint distribution between many discrete random variables (such as words in a sentence, or discrete attributes in a data-mining task). For example, if one wants to model the joint distribution of 10 consecutive words in a natural language with a vocabulary V of size 100,000, there are potentially $100\,000^{10} - 1 = 10^{50} - 1$ free parameters.

¹For more precise description, please refer to [Shalev-Shwartz and Ben-David, 2014]

A Neural Probabilistic Language Model

The first paragraph of the paper *A Neural Probabilistic Language Model* [Bengio et al., 2001]

1 Introduction

A fundamental problem that makes language modeling and other learning problems difficult is the *curse of dimensionality*. It is particularly obvious in the case when one wants to model the joint distribution between many discrete random variables (such as words in a sentence, or discrete attributes in a data-mining task). For example, if one wants to model the joint distribution of 10 consecutive words in a natural language with a vocabulary V of size 100,000, there are potentially $100\,000^{10} - 1 = 10^{50} - 1$ free parameters.

Curse of dimensionality

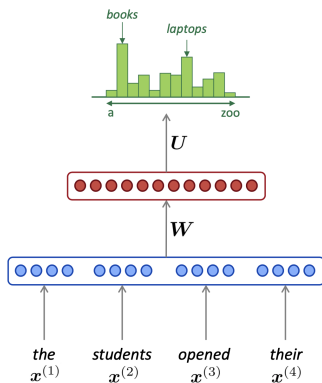
The sample complexity is an exponential function of the dimensionality of data¹

¹For more precise description, please refer to [Shalev-Shwartz and Ben-David, 2014]

A Neural Language Model with Fixed Window Size

Improvement over n -gram language models

- ▶ Less parameters (with large n 's)
- ▶ No sparsity problem
- ▶ No smoothing is needed



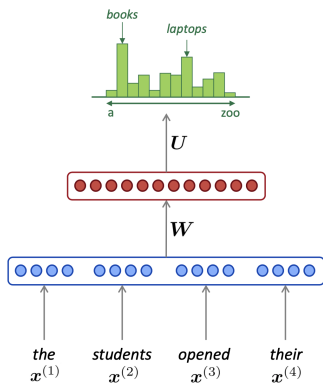
A Neural Language Model with Fixed Window Size

Improvement over n -gram language models

- ▶ Less parameters (with large n 's)
- ▶ No sparsity problem
- ▶ No smoothing is needed

Remaining issues

- ▶ Fixed window size – similar to n -gram models
- ▶ No explicit modeling of word order beyond the window size
- ▶ Same word will be computed k times along the sliding window, where k is the window size



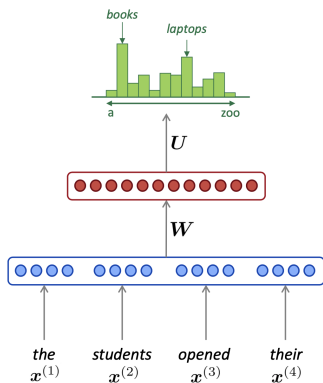
A Neural Language Model with Fixed Window Size

Improvement over n -gram language models

- ▶ Less parameters (with large n 's)
- ▶ No sparsity problem
- ▶ No smoothing is needed

Remaining issues

- ▶ Fixed window size – similar to n -gram models
- ▶ No explicit modeling of word order beyond the window size
- ▶ Same word will be computed k times along the sliding window, where k is the window size



We need a new neural network **architecture** that can read words continuously along predictions

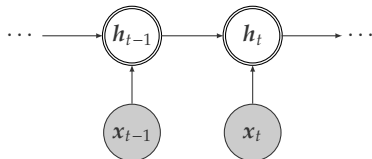
Recurrent Neural Networks

Recurrent Neural Networks (RNNs)

A simple RNN is defined by the following recursive function

$$h_t = f(x_t, h_{t-1}) \quad (4)$$

and depicted as



where

- ▶ h_{t-1} : hidden state at time step $t - 1$
- ▶ x_t : input at time step t
- ▶ h_t : hidden state at time step t

A Simple Transition Function

In the simplest case, the transition function f is defined with an **element-wise** Sigmoid function and a linear transformation of x_t and h_{t-1}

$$h_t = f(x_t, h_{t-1}) = \sigma(W_h h_{t-1} + W_i x_t + b) \quad (5)$$

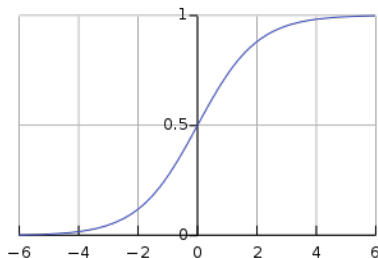
where

- ▶ x_t : input word embedding
- ▶ h_{t-1} : hidden state from previous time step
- ▶ W_h : parameter matrix for **hidden states**
- ▶ W_i : parameter matrix for **inputs**
- ▶ b : bias term (also a parameter)

Sigmoid Function

A Sigmoid function with one-dimensional input $x \in (-\infty, \infty)$

$$\sigma(x) = \frac{1}{1 + e^{-x}}$$



The potential numeric issue caused by the Sigmoid function

- ▶ $\sigma(x) \rightarrow 1$ with $x \gg 6$
- ▶ $\sigma(x) \rightarrow 0$, $x \ll -6$

The output of the Sigmoid function will approximate a constant, when the input value is beyond certain ranges

Unfolding RNNs

We can unfold this recursive definition of a RNN

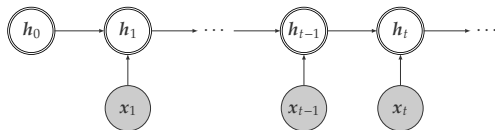
$$h_t = f(x_t, h_{t-1}) \tag{6}$$

Unfolding RNNs

We can unfold this recursive definition of a RNN

$$h_t = f(x_t, h_{t-1}) \quad (6)$$

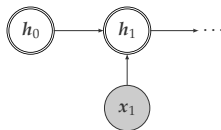
as



$$\begin{aligned} h_t &= f(x_t, f(x_{t-1}, h_{t-2})) \\ &= f(x_t, f(x_{t-1}, f(x_{t-2}, h_{t-3}))) \\ &= \dots \\ &= f(x_t, f(x_{t-1}, f(x_{t-2}, \dots f(x_1, h_0) \dots))) \end{aligned} \quad (7)$$

Base Condition

Base condition defines the starting point of the recursive computation



$$h_t = f(x_t, f(x_{t-1}, f(x_{t-2}, \dots f(\mathbf{x}_1, \mathbf{h}_0) \dots))) \quad (8)$$

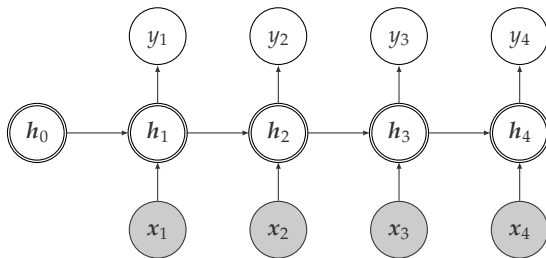
- ▶ h_0 : zero vector or parameter
- ▶ x_1 : input at time $t = 1$

RNN for Sequential Prediction

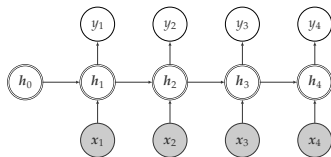
In general, RNNs can be used for any sequential modeling tasks

$$h_t = f(x_t, h_{t-1}) \quad (9)$$

$$P(Y_t; h_t) = \text{softmax}(W_o h_t + b_o) \quad (10)$$



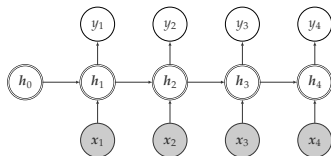
Sequential Modeling as Classification



- Prediction at each time step t

$$\hat{y}_t = \operatorname{argmax}_y P(Y_t = y; h_t) \quad (11)$$

Sequential Modeling as Classification



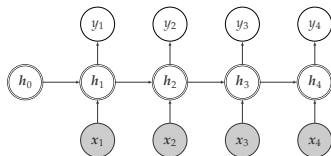
- Prediction at each time step t

$$\hat{y}_t = \operatorname{argmax}_y P(Y_t = y; h_t) \quad (11)$$

- Loss at single time step t

$$L_t(y_t, \hat{y}_t) = -\log P(y_t; h_t) \quad (12)$$

Sequential Modeling as Classification



- Prediction at each time step t

$$\hat{y}_t = \operatorname{argmax}_y P(Y_t = y; \mathbf{h}_t) \quad (11)$$

- Loss at single time step t

$$L_t(y_t, \hat{y}_t) = -\log P(y_t; \mathbf{h}_t) \quad (12)$$

- The total loss

$$\ell = \sum_{t=1}^T L_t(y_t, \hat{y}_t) \quad (13)$$

Computation Graph

For simplicity, consider the example of a two-layer neural network

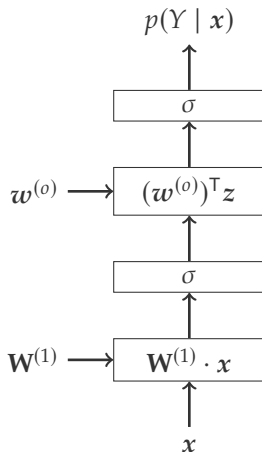
$$P(Y = +1 \mid \mathbf{x}) = \sigma \left((\mathbf{w}^{(o)})^\top \sigma(\mathbf{W}^{(1)} \mathbf{x}) \right) \quad (14)$$

A neural network is a composition of some basic functions and operations. For example

- ▶ the Sigmoid function $\sigma(\cdot)$
- ▶ matrix transpose $(\mathbf{w}^{(o)})^\top$
- ▶ matrix-vector multiplication $\mathbf{W}^{(1)} \mathbf{x}$

Forward Graph

The computation graph of the two-layer neural network²



²For simplicity, the *transpose* operation blocks are ignored from the graph

Backward Operations

Similarly, the gradient of neural network parameters are computed with a series of backward operations associated with the derivative of some basic function. For example

- ▶ $\frac{dbx}{dx} = b$

- ▶ $\frac{d \log(x)}{dx} = \frac{1}{x}$

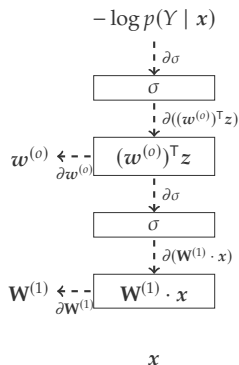
- ▶ $\frac{d\sigma(x)}{dx} = \sigma(x)(1 - \sigma(x))$

- ▶ $\frac{da^T x}{dx} = a$

- ▶ $\frac{dWx}{dW} = \begin{bmatrix} x^T \\ \vdots \\ x^T \end{bmatrix}$

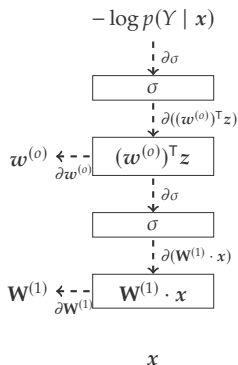
Backward Graph

- ▶ The gradient of each operations can be computed individually based on the input and output
- ▶ With the chain rule, gradient of the loss function with respect to any parameter is a sequence of multiplications of individual gradients



Backward Graph

- ▶ The gradient of each operations can be computed individually based on the input and output
- ▶ With the chain rule, gradient of the loss function with respect to any parameter is a sequence of multiplications of individual gradients



For example,

$$\frac{\partial L(\theta)}{\partial w^{(0)}} = - \frac{\partial \log \sigma(\cdot)}{\partial \sigma(\cdot)} \cdot \frac{\partial \sigma((w^{(0)})^T \sigma(W^{(1)} x))}{\partial (w^{(0)})^T \sigma(W^{(1)} x)} \cdot \frac{\partial (w^{(0)})^T \sigma(W^{(1)} x)}{\partial w^{(0)}}$$

Example: MiniTorch

A screenshot from the MINITORCH library³

```
class Add(Function):
    @staticmethod
    def forward(ctx, t1, t2):
        return add_zip(t1, t2)

    @staticmethod
    def backward(ctx, grad_output):
        return grad_output, grad_output

class Mul(Function):
    @staticmethod
    def forward(ctx, a, b):
        ctx.save_for_backward(a, b)
        return mul_zip(a, b)

    @staticmethod
    def backward(ctx, grad_output):
        a, b = ctx.saved_values
        return mul_zip(b, grad_output), mul_zip(a, grad_output)

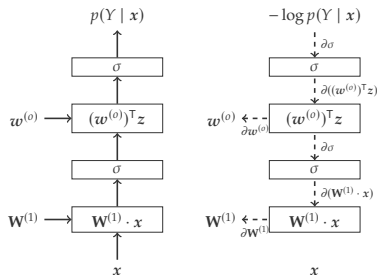
class Sigmoid(Function):
    @staticmethod
    def forward(ctx, a):
        out = sigmoid_map(a)
        ctx.save_for_backward(out)
        return out

    @staticmethod
    def backward(ctx, grad_output):
        sigma = ctx.saved_values
        return sigma * (-sigma + 1.0) * grad_output
```

³A teaching library of PyTorch developed by Sasha Rush and Ge Gao

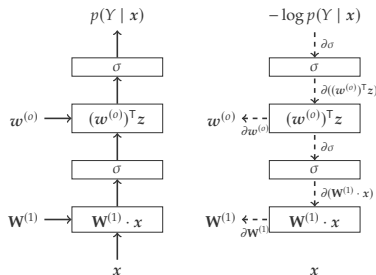
Computation Graph

Perform the forward/backward step with a graph of basic operations (e.g., PyTorch, Tensorflow)



Computation Graph

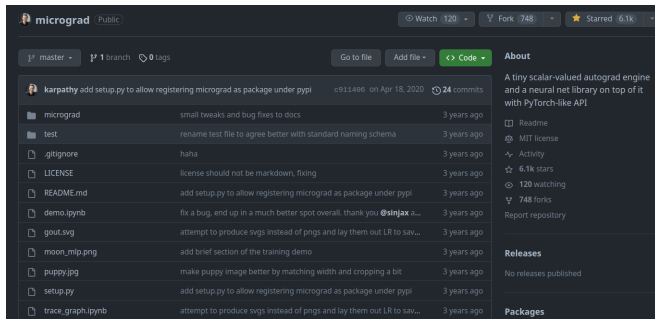
Perform the forward/backward step with a graph of basic operations (e.g., PyTorch, Tensorflow)



- ▶ Modular implementation: implement each module with its forward/backward operations together
- ▶ Automatic differentiation: automatically run with the backward step

Example: Micrograd

MICROGRAD is an extremely simple example to illustrate the idea of using autograd to implementing the backpropagation algorithm.



The screenshot shows the GitHub repository page for 'micrograd' by user 'karpthy'. The repository is public and has 120 watches, 748 forks, and 6.1k stars. It is currently on the 'master' branch with 1 branch and 0 tags. The repository description states: 'A tiny scalar-valued autograd engine and a neural net library on top of it with PyTorch-like API'. The file list includes: 'micrograd' (small tweaks and bug fixes to docs, 3 years ago), 'test' (rename test file to agree better with standard naming schema, 3 years ago), '.gitignore' (haha, 3 years ago), 'LICENSE' (license should not be markdown, fixing, 3 years ago), 'README.md' (add setup.py to allow registering micrograd as package under pypi, 3 years ago), 'demo.ipynb' (fix a bug, end up in a much better spot overall, thank you @sinjx a..., 3 years ago), 'gout.svg' (attempt to produce svgs instead of pngs and lay them out LR to sav..., 3 years ago), 'moon_mlp.png' (add brief section of the training demo, 3 years ago), 'puppy.jpg' (make puppy image better by matching width and cropping a bit, 3 years ago), 'setup.py' (add setup.py to allow registering micrograd as package under pypi, 3 years ago), and 'trace_graph.ipynb' (attempt to produce svgs instead of pngs and lay them out LR to sav..., 3 years ago). The right sidebar shows the 'About' section with the repository description, 'Readme', 'MIT license', 'Activity', '6.1k stars', '120 watching', '748 forks', and 'Report repository'. The 'Releases' section shows 'No releases published'. The 'Packages' section is also visible.

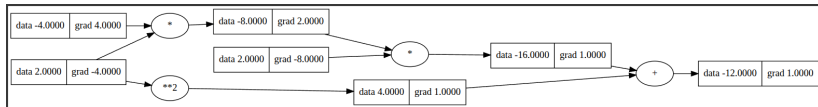
Demo

Example: Gradient Computation

Function

$$f(a, b) = 2ab + b^2$$

The computation graph on forward and backward operations



RNN Language Modeling

A language model defines the probability of x_t given $\mathbf{x} = (x_1, x_2, \dots, x_{t-1})$ as

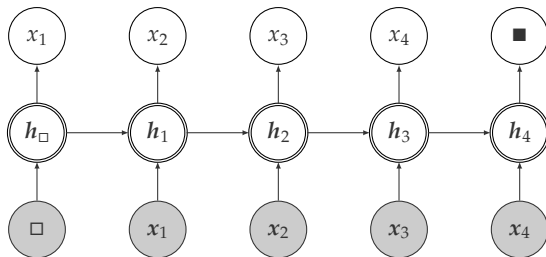
$$P(x_t \mid x_1, \dots, x_{t-1}) \tag{15}$$

and the joint probability as

$$\begin{aligned} P(\mathbf{x}_{1:T}) &= P(x_1) \cdot P(x_2 \mid x_1) \\ &\quad \cdot \dots \cdot \\ &\quad \cdot P(x_T \mid x_1, x_2, \dots, x_{T-1}) \end{aligned}$$

Language Modeling with RNNs

Using RNNs for language modeling

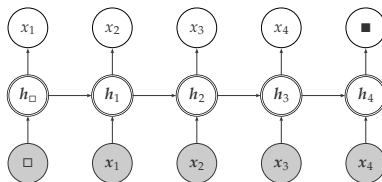


with two special tokens

$$\{\square, x_1, \dots, x_T, \blacksquare\}$$

RNN Language Models

For a given sentence $\{x_1, \dots, x_t\}$, the input at time t is **word embedding** x_t



The probability distribution of next word X_t

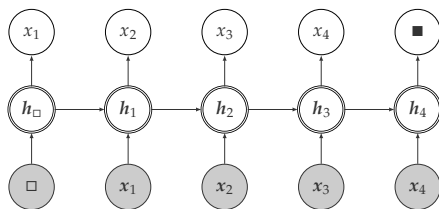
$$P(X_t = x \mid \mathbf{x}_{1:t-1}) = \frac{\exp(\mathbf{w}_{0,x}^\top \mathbf{h}_{t-1})}{\sum_{x' \in \mathcal{V}} \exp(\mathbf{w}_{0,x'}^\top \mathbf{h}_{t-1})} \quad (16)$$

where

- ▶ $\mathbf{w}_{0,x}$ is the output weight vector (parameter) associated with word x
- ▶ $\mathcal{V}$ is the word vocabulary

Special Cases

Similar to statistical language modeling, there are also two special cases that we need to consider



$$\{\square, x_1, \dots, x_T, \blacksquare\}$$

The corresponding prediction functions are defined as

- ▶ At time $t = 1$

$$P(X_1 = x) \propto \exp(w_{o,x}^\top h_\square) \quad (17)$$

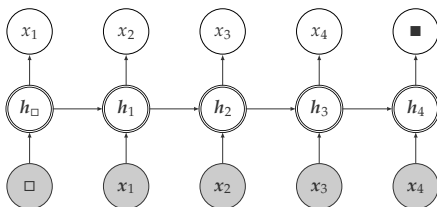
- ▶ At time $t = T$

$$P(X_T = \blacksquare \mid x_{1:T-1}) \propto \exp(w_{o,x}^\top h_{T-1}) \quad (18)$$

Challenge of Training RNNs

Objective

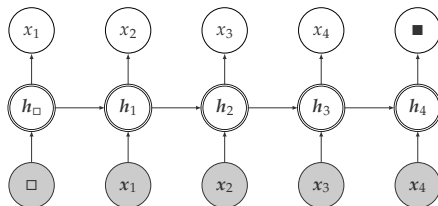
The training objective for each timestep is to predict the next token in the text



- ▶ Prediction at step t , $P(X_t = x \mid \mathbf{x}_{1:t-1}) = \frac{\exp(w_{o,x}^T h_{t-1})}{\sum_{x' \in \mathcal{V}} \exp(w_{o,x'}^T h_{t-1})}$
- ▶ Loss at step t , $L_t = -\log P(X_t = x \mid \mathbf{x}_{1:t-1})$

Let θ denote **all** model parameters

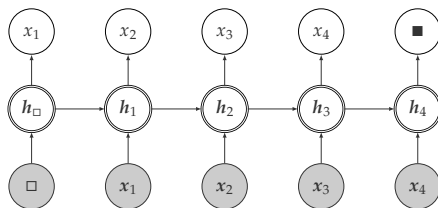
$$\frac{\partial \ell}{\partial \theta} = \sum_{t=1}^T \frac{\partial L_t}{\partial \theta} \quad (19)$$



Backpropagation Through Time [Rumelhart et al., 1985, BPTT]

Model Parameters

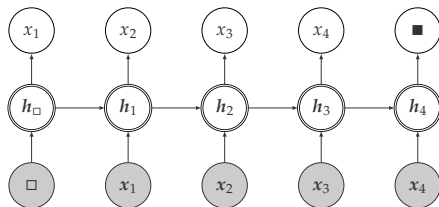
Before computing the gradient of each L_t with respect to model parameters, let us count how many parameters that we need consider



- Output parameter matrix $W_o = (w_{o,1}, \dots, w_{o,V})$

Model Parameters

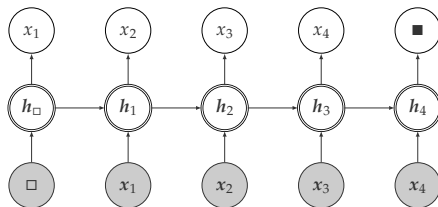
Before computing the gradient of each L_t with respect to model parameters, let us count how many parameters that we need consider



- ▶ Output parameter matrix $W_o = (w_{o,1}, \dots, w_{o,V})$
- ▶ Input word embedding matrix $X = (x_1, \dots, x_V)$

Model Parameters

Before computing the gradient of each L_t with respect to model parameters, let us count how many parameters that we need consider



- ▶ Output parameter matrix $W_o = (w_{o,1}, \dots, w_{o,V})$
- ▶ Input word embedding matrix $X = (x_1, \dots, x_V)$
- ▶ Neural network parameters W_h, W_i, b

Backpropagation Through Time

Take time step t as an example, we can take a look the gradient computation of some specific parameters

- ▶ Output model parameter $\frac{\partial L_t}{\partial w_o}$.

Backpropagation Through Time

Take time step t as an example, we can take a look the gradient computation of some specific parameters

- ▶ Output model parameter $\frac{\partial L_t}{\partial w_o}$
- ▶ Neural network parameters, for example W_h

$$\frac{\partial L_t}{\partial W_h} = \sum_{i=1}^t \left\{ \frac{\partial L_t}{\partial h_t} \cdot \left(\prod_{j=i}^{t-1} \frac{\partial h_{j+1}}{\partial h_j} \right) \cdot \frac{\partial h_i}{\partial W_h} \right\} \quad (20)$$

Similar patterns for the other two neural network parameters W_i and b

Backpropagation Through Time

Take time step t as an example, we can take a look the gradient computation of some specific parameters

- ▶ Output model parameter $\frac{\partial L_t}{\partial w_o}$.
- ▶ Neural network parameters, for example W_h

$$\frac{\partial L_t}{\partial W_h} = \sum_{i=1}^t \left\{ \frac{\partial L_t}{\partial h_t} \cdot \left(\prod_{j=i}^{t-1} \frac{\partial h_{j+1}}{\partial h_j} \right) \cdot \frac{\partial h_i}{\partial W_h} \right\} \quad (20)$$

Similar patterns for the other two neural network parameters W_i and b

- ▶ Word embedding $\frac{\partial L_t}{\partial x_{t'}}$
 - ▶ E.g., word embedding $x_{t'}$ is the input of h_t if $t' \leq t$, so ...

Challenges

For each timestep, we need to compute the gradient using the chain rule:

$$\frac{\partial L_t}{\partial \mathbf{W}_h} = \sum_{i=1}^t \left\{ \frac{\partial L_t}{\partial \mathbf{h}_t} \cdot \left(\prod_{j=i}^{t-1} \frac{\partial \mathbf{h}_{j+1}}{\partial \mathbf{h}_j} \right) \cdot \frac{\partial \mathbf{h}_i}{\partial \mathbf{W}_h} \right\} \quad (21)$$

The chain rule of gradient will cause two potential problems in training RNNs

- ▶ vanishing gradients: $\frac{\partial L_t}{\partial \theta} \rightarrow 0$
- ▶ exploding gradients: $\frac{\partial L_t}{\partial \theta} \geq M$

[Pascanu et al., 2013]

Exploding Gradients

Solution: **norm clipping** [Pascanu et al., 2013].

Consider the gradient $g = \frac{\partial \ell}{\partial \theta}$,

$$\hat{g} \leftarrow \tau \cdot \frac{g}{\|g\|} \quad (22)$$

when $\|g\| > \tau$.

- ▶ Usually, $\tau = 3$ or 5 in practice.
- ▶ Smaller gradient will cause slower learning progress

Solution:

- ▶ initialize parameters carefully
- ▶ replace hidden state transition function $\sigma(\cdot)$ with other options

$$f(x_t, h_{t-1}) = \sigma(\mathbf{W}_h h_{t-1} + \mathbf{W}_i x_t + \mathbf{b}) \quad (23)$$

- ▶ LSTM [Hochreiter and Schmidhuber, 1997]
- ▶ GRU [Cho et al., 2014]

From the first page of the original paper proposing LSTM [Hochreiter and Schmidhuber, 1997]

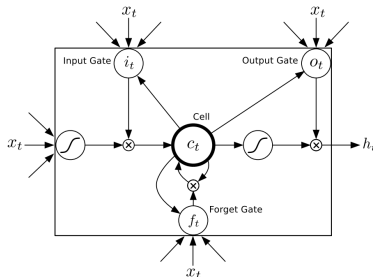
The problem. With conventional “Back-Propagation Through Time” (BPTT, e.g., Williams and Zipser 1992, Werbos 1988) or “Real-Time Recurrent Learning” (RTRL, e.g., Robinson and Fallside 1987), error signals “flowing backwards in time” tend to either (1) blow up or (2) vanish: the temporal evolution of the backpropagated error exponentially depends on the size of the weights (Hochreiter 1991). Case (1) may lead to oscillating weights, while in case (2) learning to bridge long time lags takes a prohibitive amount of time, or does not work at all (see section 3).

The remedy. This paper presents “*Long Short-Term Memory*” (LSTM), a novel recurrent network architecture in conjunction with an appropriate gradient-based learning algorithm. LSTM is designed to overcome these error back-flow problems. It can learn to bridge time intervals in excess of 1000 steps even in case of noisy, incompressible input sequences, without loss of short time lag capabilities. This is achieved by an efficient, gradient-based algorithm for an architecture

Long Short-Term Memory

Rather than directly taking input and hidden state as simple transition function, LSTM relies on three gates to control *how much* information it should take from input and hidden state before combining them together

$$\begin{aligned}i_t &= \sigma(\mathbf{W}_{xi}\mathbf{x}_t + \mathbf{W}_{hi}\mathbf{h}_{t-1} + \mathbf{W}_{ci}\mathbf{c}_{t-1} + \mathbf{b}_i) \\f_t &= \sigma(\mathbf{W}_{xf}\mathbf{x}_t + \mathbf{W}_{hf}\mathbf{h}_{t-1} + \mathbf{W}_{cf}\mathbf{c}_{t-1} + \mathbf{b}_f) \\c_t &= f_t \circ \mathbf{c}_{t-1} + i_t \circ \tanh(\mathbf{W}_{xc}\mathbf{x}_t + \mathbf{W}_{hc}\mathbf{h}_{t-1} + \mathbf{b}_c) \\o_t &= \sigma(\mathbf{W}_{xo}\mathbf{x}_t + \mathbf{W}_{ho}\mathbf{h}_{t-1} + \mathbf{W}_{co}\mathbf{c}_t + \mathbf{b}_o) \\h_t &= o_t \circ \tanh(c_t)\end{aligned}$$



where $\circ$ is the element-wise multiplication, $\{\mathbf{W}\}$ and $\{\mathbf{b}\}$ are parameters. [Graves, 2013]

Reference



Bengio, Y., Ducharme, R., and Vincent, P. (2001).
A neural probabilistic language model.
In *NIPS*.



Cho, K., Van Merriënboer, B., Bahdanau, D., and Bengio, Y. (2014).
On the properties of neural machine translation: Encoder-decoder approaches.
arXiv preprint arXiv:1409.1259.



Graves, A. (2013).
Generating sequences with recurrent neural networks.
arXiv preprint arXiv:1308.0850.



Hochreiter, S. and Schmidhuber, J. (1997).
Long short-term memory.
Neural computation, 9(8):1735–1780.



Pascanu, R., Mikolov, T., and Bengio, Y. (2013).
On the difficulty of training recurrent neural networks.
In *International Conference on Machine Learning*, pages 1310–1318.



Rumelhart, D. E., Hinton, G. E., and Williams, R. J. (1985).
Learning internal representations by error propagation.
Technical report, California Univ San Diego La Jolla Inst for Cognitive Science.



Shalev-Shwartz, S. and Ben-David, S. (2014).
Understanding machine learning: From theory to algorithms.
Cambridge university press.