

CS 6770 Natural Language Processing

Text Classification (II): Neural Classifiers

Yangfeng Ji

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



1. From Logistic Regression to Neural Networks
2. Learning Neural Networks
3. Neural Text Classifiers
4. Neural Networks: Tricks of the Trade

From Logistic Regression to Neural Networks

Logistic Regression

A quick review of logistic regression with $\mathbf{x} = [x_1, \dots, x_d] \in \mathbb{R}^d$

- ▶ An basic form for binary classification $y \in \{-1, +1\}$

$$P(Y = +1 \mid \mathbf{x}) = \frac{1}{1 + \exp(-\langle \mathbf{w}, \mathbf{x} \rangle)} \quad (1)$$

Logistic Regression

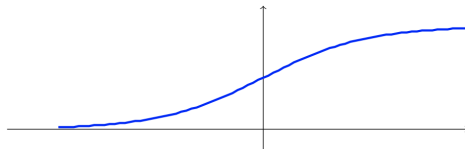
A quick review of logistic regression with $x = [x_1, \dots, x_d] \in \mathbb{R}^d$

- ▶ An basic form for binary classification $y \in \{-1, +1\}$

$$P(Y = +1 \mid x) = \frac{1}{1 + \exp(-\langle w, x \rangle)} \quad (1)$$

- ▶ The sigmoid function $\sigma(a)$ with $a \in \mathbb{R}$, which is a monotonic nonlinear function

$$\sigma(a) = \frac{1}{1 + \exp(-a)} \quad (2)$$



Graphical Representation

- ▶ A specific example of LR with four input features

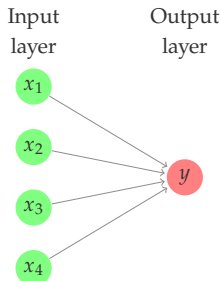
$$P(Y = 1 \mid \mathbf{x}) = \sigma\left(\sum_{j=1}^4 w_j x_j\right) \quad (3)$$

Graphical Representation

- ▶ A specific example of LR with four input features

$$P(Y = 1 \mid \mathbf{x}) = \sigma\left(\sum_{j=1}^4 w_j x_j\right) \quad (3)$$

- ▶ The graphical representation of this LR model is



From LR to a Simple Neural Network

Build upon logistic regression, a simple neural network with K hidden units $\{z_k\}_{k=1}^K$ can be constructed as

$$z_k = \sigma\left(\sum_{j=1}^d w_{k,j}^{(1)} x_j\right) \quad k \in \{1, \dots, K\} \quad (4)$$

(5)

- ▶ $\mathbf{x} \in \mathbb{R}^d$: d -dimensional input
- ▶ K is the number of **hidden** units, each of them has the same form as a LR.

From LR to a Simple Neural Network

Build upon logistic regression, a simple neural network with K hidden units $\{z_k\}_{k=1}^K$ can be constructed as

$$z_k = \sigma\left(\sum_{j=1}^d w_{k,j}^{(1)} x_j\right) \quad k \in \{1, \dots, K\} \quad (4)$$

$$P(y = +1 \mid \mathbf{x}) = \sigma\left(\sum_{k=1}^K w_k^{(o)} z_k\right) \quad (5)$$

- ▶ $\mathbf{x} \in \mathbb{R}^d$: d -dimensional input
- ▶ K is the number of **hidden** units, each of them has the same form as a LR.
- ▶ $y \in \{-1, +1\}$ (binary classification problem)

From LR to a Simple Neural Network

Build upon logistic regression, a simple neural network with K hidden units $\{z_k\}_{k=1}^K$ can be constructed as

$$z_k = \sigma\left(\sum_{j=1}^d w_{k,j}^{(1)} x_j\right) \quad k \in \{1, \dots, K\} \quad (4)$$

$$P(y = +1 \mid \mathbf{x}) = \sigma\left(\sum_{k=1}^K w_k^{(o)} z_k\right) \quad (5)$$

- ▶ $\mathbf{x} \in \mathbb{R}^d$: d -dimensional input
- ▶ K is the number of **hidden** units, each of them has the same form as a LR.
- ▶ $y \in \{-1, +1\}$ (binary classification problem)
- ▶ $\{w_{k,i}^{(1)}\}$ and $\{w_k^{(o)}\}$ are two sets of the parameters, and
- ▶ $\sigma(\cdot)$ is called the **activation function**

Mathematical Formulation

With the notations of matrix-vector multiplication, we can rewrite these equations into a more concise form

► Element-wise formulation

$$z_k = \sigma\left(\sum_{j=1}^d w_{k,j}^{(1)} x_j\right) \quad k \in [K] \quad (6)$$

$$P(y = +1 \mid \mathbf{x}) = \sigma\left(\sum_{k=1}^K w_k^{(o)} z_k\right) \quad (7)$$

Mathematical Formulation

With the notations of matrix-vector multiplication, we can rewrite these equations into a more concise form

- ▶ Element-wise formulation

$$z_k = \sigma\left(\sum_{j=1}^d w_{k,j}^{(1)} x_j\right) \quad k \in [K] \quad (6)$$

$$P(y = +1 \mid \mathbf{x}) = \sigma\left(\sum_{k=1}^K w_k^{(o)} z_k\right) \quad (7)$$

- ▶ Matrix-vector formulation

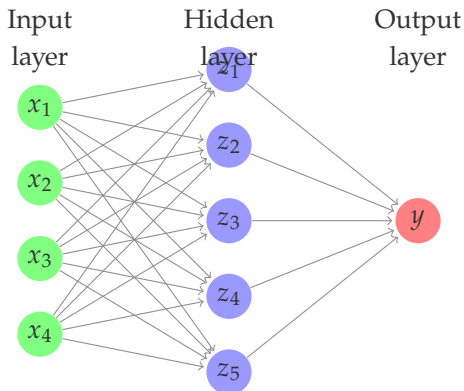
$$\mathbf{z} = \sigma(\mathbf{W}^{(1)} \mathbf{x}) \quad (8)$$

$$P(y = +1 \mid \mathbf{x}) = \sigma((\mathbf{w}^{(o)})^\top \mathbf{z}) \quad (9)$$

where $\mathbf{W}^{(1)} \in \mathbb{R}^{K \times d}$ and $\mathbf{w}^{(o)} \in \mathbb{R}^K$

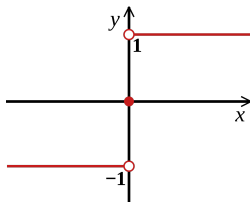
Graphical Representation

Assume the input dimension $d = 4$ and the number of hidden units (hidden dimension) $K = 5$, we can represent this neural network model with a similar graphical representation



Other Activation Functions

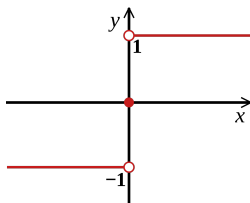
In deep learning, we have a few options about activation functions:



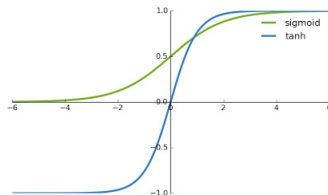
(a) Sign function

Other Activation Functions

In deep learning, we have a few options about activation functions:



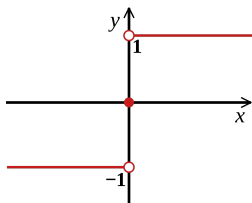
(a) Sign function



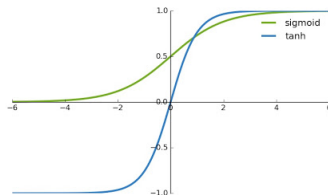
(b) Sigmoid and Tanh function

Other Activation Functions

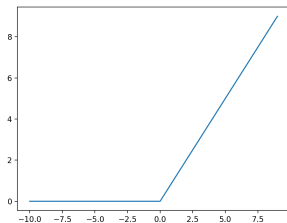
In deep learning, we have a few options about activation functions:



(a) Sign function



(b) Sigmoid and Tanh function



(c) ReLU function

[Jarrett et al., 2009]

Learning Neural Networks

Neural Network Predictions

Consider the previously defined neural network model for binary classification $\mathcal{Y} = \{-1, +1\}$,

- ▶ Re-write the model definition into a compact form

$$P(Y = +1 \mid \mathbf{x}) = \sigma \left((\mathbf{w}^{(0)})^\top \sigma(\mathbf{W}^{(1)} \mathbf{x}) \right) \quad (10)$$

where $\{\mathbf{w}^{(0)}, \mathbf{W}^{(1)}\}$ are the parameters

Neural Network Predictions

Consider the previously defined neural network model for binary classification $\mathcal{Y} = \{-1, +1\}$,

- ▶ Re-write the model definition into a compact form

$$P(Y = +1 \mid \mathbf{x}) = \sigma \left((\mathbf{w}^{(0)})^\top \sigma(\mathbf{W}^{(1)} \mathbf{x}) \right) \quad (10)$$

where $\{\mathbf{w}^{(0)}, \mathbf{W}^{(1)}\}$ are the parameters

- ▶ Assume the ground-truth label is y , let's introduce an **empirical** distribution

$$Q(Y = y' \mid \mathbf{x}) = \delta(y', y) = \begin{cases} 1 & y' = y \\ 0 & y' \neq y \end{cases} \quad (11)$$

Cross Entropy

Given one data point, The loss function of a neural network is usually defined as the **cross entropy** of the prediction distribution p and the empirical distribution q

$$\begin{aligned} H(Q, P) = & -Q(Y = +1 | x) \log P(Y = +1 | x) \\ & -Q(Y = -1 | x) \log P(Y = -1 | x) \end{aligned} \quad (12)$$

Cross Entropy

Given one data point, The loss function of a neural network is usually defined as the **cross entropy** of the prediction distribution p and the empirical distribution q

$$\begin{aligned} H(Q, P) = & -Q(Y = +1 | x) \log P(Y = +1 | x) \\ & -Q(Y = -1 | x) \log P(Y = -1 | x) \end{aligned} \quad (12)$$

Since q is defined with a Delta function, depending on y , we have

$$H(Q, P) = \begin{cases} -\log P(Y = +1 | x) & Y = +1 \\ -\log P(Y = -1 | x) & Y = -1 \end{cases} \quad (13)$$

- ▶ Given a set of training example $S = \{(\mathbf{x}^{(i)}, y^{(i)})\}_{i=1}^m$, the loss function is defined as

$$L(\boldsymbol{\theta}) = - \sum_{i=1}^m \log p(y^{(i)} | \mathbf{x}^{(i)}) \quad (14)$$

where $\boldsymbol{\theta}$ indicates all the parameters in a network.

- ▶ Given a set of training example $S = \{(\mathbf{x}^{(i)}, y^{(i)})\}_{i=1}^m$, the loss function is defined as

$$L(\boldsymbol{\theta}) = - \sum_{i=1}^m \log p(y^{(i)} | \mathbf{x}^{(i)}) \quad (14)$$

where $\boldsymbol{\theta}$ indicates all the parameters in a network.

- ▶ For example, $\boldsymbol{\theta} = \{\mathbf{w}^{(o)}, \mathbf{W}^{(1)}\}$, for the previously defined two-layer neural network

- ▶ Given a set of training example $S = \{(\mathbf{x}^{(i)}, y^{(i)})\}_{i=1}^m$, the loss function is defined as

$$L(\boldsymbol{\theta}) = - \sum_{i=1}^m \log p(y^{(i)} | \mathbf{x}^{(i)}) \quad (14)$$

where $\boldsymbol{\theta}$ indicates all the parameters in a network.

- ▶ For example, $\boldsymbol{\theta} = \{\mathbf{w}^{(o)}, \mathbf{W}^{(1)}\}$, for the previously defined two-layer neural network
- ▶ Just like learning a LR, we can use the **gradient-based** learning

Gradient-based Learning

A simple scratch of gradient-based learning

1. Compute the gradient of θ , $\frac{\partial L(\theta)}{\partial \theta}$

Gradient-based Learning

A simple scratch of gradient-based learning

1. Compute the gradient of θ , $\frac{\partial L(\theta)}{\partial \theta}$
2. Update the parameter with the gradient

$$\theta^{(\text{new})} \leftarrow \theta^{(\text{old})} - \eta \cdot \frac{\partial L(\theta)}{\partial \theta} \Big|_{\theta=\theta^{(\text{old})}} \quad (15)$$

where η is the learning rate

Gradient-based Learning

A simple scratch of gradient-based learning

1. Compute the gradient of θ , $\frac{\partial L(\theta)}{\partial \theta}$
2. Update the parameter with the gradient

$$\theta^{(\text{new})} \leftarrow \theta^{(\text{old})} - \eta \cdot \frac{\partial L(\theta)}{\partial \theta} \Big|_{\theta=\theta^{(\text{old})}} \quad (15)$$

where η is the learning rate

3. Go back step 1 until it converges

Gradient Computation

Consider the two-layer neural network with **one** training example $(\mathbf{x}, y)$, to further simplify the computation, we assume $y = +1$

$$\log p(y | \mathbf{x}) = \log \sigma \left((\mathbf{w}^{(o)})^\top \sigma(\mathbf{W}^{(1)} \mathbf{x}) \right) \quad (16)$$

Gradient Computation

Consider the two-layer neural network with **one** training example $(\mathbf{x}, y)$, to further simplify the computation, we assume **$y = +1$**

$$\log p(y | \mathbf{x}) = \log \sigma \left((\mathbf{w}^{(o)})^\top \sigma(\mathbf{W}^{(1)} \mathbf{x}) \right) \quad (16)$$

The gradient with respect to $\mathbf{w}^{(o)}$ is

$$\frac{\partial L(\theta)}{\partial \mathbf{w}^{(o)}} = - \frac{\partial \log \sigma(\cdot)}{\partial \sigma(\cdot)} \quad (17)$$

Gradient Computation

Consider the two-layer neural network with **one** training example $(\mathbf{x}, y)$, to further simplify the computation, we assume $y = +1$

$$\log p(y | \mathbf{x}) = \log \sigma \left((\mathbf{w}^{(o)})^\top \sigma(\mathbf{W}^{(1)} \mathbf{x}) \right) \quad (16)$$

The gradient with respect to $\mathbf{w}^{(o)}$ is

$$\frac{\partial L(\theta)}{\partial \mathbf{w}^{(o)}} = - \frac{\partial \log \sigma(\cdot)}{\partial \sigma(\cdot)} \cdot \frac{\partial \sigma \left((\mathbf{w}^{(o)})^\top \sigma(\mathbf{W}^{(1)} \mathbf{x}) \right)}{\partial (\mathbf{w}^{(o)})^\top \sigma(\mathbf{W}^{(1)} \mathbf{x})} \quad (17)$$

Gradient Computation

Consider the two-layer neural network with **one** training example $(\mathbf{x}, y)$, to further simplify the computation, we assume **$y = +1$**

$$\log p(y | \mathbf{x}) = \log \sigma \left((\mathbf{w}^{(o)})^\top \sigma(\mathbf{W}^{(1)} \mathbf{x}) \right) \quad (16)$$

The gradient with respect to $\mathbf{w}^{(o)}$ is

$$\frac{\partial L(\theta)}{\partial \mathbf{w}^{(o)}} = - \frac{\partial \log \sigma(\cdot)}{\partial \sigma(\cdot)} \cdot \frac{\partial \sigma \left((\mathbf{w}^{(o)})^\top \sigma(\mathbf{W}^{(1)} \mathbf{x}) \right)}{\partial (\mathbf{w}^{(o)})^\top \sigma(\mathbf{W}^{(1)} \mathbf{x})} \cdot \frac{\partial (\mathbf{w}^{(o)})^\top \sigma(\mathbf{W}^{(1)} \mathbf{x})}{\partial \mathbf{w}^{(o)}} \quad (17)$$

Gradient Computation

Consider the two-layer neural network with **one** training example $(\mathbf{x}, y)$, to further simplify the computation, we assume $y = +1$

$$\log p(y | \mathbf{x}) = \log \sigma \left((\mathbf{w}^{(o)})^\top \sigma(\mathbf{W}^{(1)} \mathbf{x}) \right) \quad (16)$$

The gradient with respect to $\mathbf{w}^{(o)}$ is

$$\begin{aligned} \frac{\partial L(\theta)}{\partial \mathbf{w}^{(o)}} &= - \frac{\partial \log \sigma(\cdot)}{\partial \sigma(\cdot)} \cdot \frac{\partial \sigma \left((\mathbf{w}^{(o)})^\top \sigma(\mathbf{W}^{(1)} \mathbf{x}) \right)}{\partial (\mathbf{w}^{(o)})^\top \sigma(\mathbf{W}^{(1)} \mathbf{x})} \cdot \frac{\partial (\mathbf{w}^{(o)})^\top \sigma(\mathbf{W}^{(1)} \mathbf{x})}{\partial \mathbf{w}^{(o)}} \\ &= - \left\{ 1 - \sigma \left((\mathbf{w}^{(o)})^\top \sigma(\mathbf{W}^{(1)} \mathbf{x}) \right) \right\} \cdot \sigma(\mathbf{W}^{(1)} \mathbf{x}) \end{aligned} \quad (17)$$

Gradient Computation (II)

The gradient with respect to $W^{(1)}$ is

$$\begin{aligned} \frac{\partial L(\theta)}{\partial \mathbf{W}^{(1)}} &= - \frac{\partial \log \sigma(\cdot)}{\partial \sigma(\cdot)} \cdot \frac{\partial \sigma\left((\mathbf{w}^{(o)})^\top \sigma(\mathbf{W}^{(1)} \mathbf{x})\right)}{\partial (\mathbf{w}^{(o)})^\top \sigma(\mathbf{W}^{(1)} \mathbf{x})} \\ &\quad \cdot \frac{\partial (\mathbf{w}^{(o)})^\top \sigma(\mathbf{W}^{(1)} \mathbf{x})}{\partial \sigma(\mathbf{W}^{(1)} \mathbf{x})} \cdot \frac{\partial \sigma(\mathbf{W}^{(1)} \mathbf{x})}{\partial \mathbf{W}^{(1)} \mathbf{x}} \cdot \frac{\partial \mathbf{W}^{(1)} \mathbf{x}}{\partial \mathbf{W}^{(1)}} \end{aligned} \quad (18)$$

Gradient Computation (II)

The gradient with respect to $W^{(1)}$ is

$$\begin{aligned} \frac{\partial L(\theta)}{\partial \mathbf{W}^{(1)}} &= - \frac{\partial \log \sigma(\cdot)}{\partial \sigma(\cdot)} \cdot \frac{\partial \sigma((\mathbf{w}^{(o)})^\top \sigma(\mathbf{W}^{(1)} \mathbf{x}))}{\partial (\mathbf{w}^{(o)})^\top \sigma(\mathbf{W}^{(1)} \mathbf{x})} \\ &\quad \cdot \frac{\partial (\mathbf{w}^{(o)})^\top \sigma(\mathbf{W}^{(1)} \mathbf{x})}{\partial \sigma(\mathbf{W}^{(1)} \mathbf{x})} \cdot \frac{\partial \sigma(\mathbf{W}^{(1)} \mathbf{x})}{\partial \mathbf{W}^{(1)} \mathbf{x}} \cdot \frac{\partial \mathbf{W}^{(1)} \mathbf{x}}{\partial \mathbf{W}^{(1)}} \end{aligned} \quad (18)$$

- ▶ Both of them are the applications of the chain rule in calculus plus some derivatives of basic functions
- ▶ In the literature of neural networks, it is called the back-propagation algorithm [Rumelhart et al., 1986]

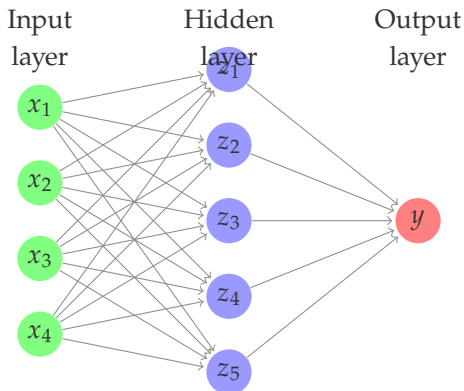
Demo

Neural Text Classifiers

Network Architecture

We are going to build a simple neural network for text classification. It includes three layers as the previous example

- ▶ Input layer
- ▶ Hidden layer
- ▶ Output layer



Example

Consider the following special case, where we have a 4-dimensional BoW representation $x \in \mathbb{R}^4$ and a weight matrix $W \in \mathbb{R}^{5 \times 4}$

$$Wx = \begin{bmatrix} 0.1 & 0.3 & 0.7 & 0.9 \\ 0.2 & 0.8 & 0.3 & 0.5 \\ 0.4 & 0.8 & 0.6 & 0.1 \\ 0.7 & 0.2 & 0.9 & 0.2 \\ 0.4 & 0.5 & 0.8 & 0.9 \end{bmatrix} \cdot \begin{bmatrix} 0 \\ 1 \\ 0 \\ 1 \end{bmatrix} \quad (19)$$

(20)

Example

Consider the following special case, where we have a 4-dimensional BoW representation $x \in \mathbb{R}^4$ and a weight matrix $W \in \mathbb{R}^{5 \times 4}$

$$Wx = \begin{bmatrix} 0.1 & 0.3 & 0.7 & 0.9 \\ 0.2 & 0.8 & 0.3 & 0.5 \\ 0.4 & 0.8 & 0.6 & 0.1 \\ 0.7 & 0.2 & 0.9 & 0.2 \\ 0.4 & 0.5 & 0.8 & 0.9 \end{bmatrix} \cdot \begin{bmatrix} 0 \\ 1 \\ 0 \\ 1 \end{bmatrix} \quad (19)$$

$$= \begin{bmatrix} 0.3 \\ 0.8 \\ 0.8 \\ 0.2 \\ 0.5 \end{bmatrix} + \begin{bmatrix} 0.9 \\ 0.5 \\ 0.1 \\ 0.2 \\ 0.9 \end{bmatrix} \quad (20)$$

Example

Consider the following special case, where we have a 4-dimensional BoW representation $x \in \mathbb{R}^4$ and a weight matrix $W \in \mathbb{R}^{5 \times 4}$

$$Wx = \begin{bmatrix} 0.1 & 0.3 & 0.7 & 0.9 \\ 0.2 & 0.8 & 0.3 & 0.5 \\ 0.4 & 0.8 & 0.6 & 0.1 \\ 0.7 & 0.2 & 0.9 & 0.2 \\ 0.4 & 0.5 & 0.8 & 0.9 \end{bmatrix} \cdot \begin{bmatrix} 0 \\ 1 \\ 0 \\ 1 \end{bmatrix} \quad (19)$$

$$= \begin{bmatrix} 0.3 \\ 0.8 \\ 0.8 \\ 0.2 \\ 0.5 \end{bmatrix} + \begin{bmatrix} 0.9 \\ 0.5 \\ 0.1 \\ 0.2 \\ 0.9 \end{bmatrix} \quad (20)$$

- ▶ Each column vector in W corresponds one word in the BoW representation
- ▶ The column vectors can be considered as representations of words, in other words, *word embeddings*

Similar to building a logistic regression classifier, we need the following steps in data processing

- ▶ Tokenize texts
- ▶ Build a vocab
- ▶ Convert a text into a collection of word indices (instead of using BoW representations)

Similar to building a logistic regression classifier, we need the following steps in data processing

- ▶ Tokenize texts
- ▶ Build a vocab
- ▶ Convert a text into a collection of word indices (instead of using BoW representations)

In addition, we also need to divide the whole set of texts into small batches

- ▶ Create mini-batches

Input Layer

In the mini-batch setting, the model usually takes a subset of texts instead of one single text. The input X represents a word index matrix, where each column vector x_i is the word indices from a corresponding text with paddings

$$X = \begin{bmatrix} 12 & 31 & \dots & 11 & 2 \\ 5 & 16 & \dots & 15 & 8 \\ 9 & \mathbf{1} & \dots & 21 & 10 \\ \mathbf{1} & \mathbf{1} & \dots & 7 & \mathbf{1} \end{bmatrix} \in \mathbb{R}^{L \times B} \quad (21)$$

- ▶ This is a slightly different representation, compared to BoW

Input Layer

In the mini-batch setting, the model usually takes a subset of texts instead of one single text. The input X represents a word index matrix, where each column vector x_i is the word indices from a corresponding text with paddings

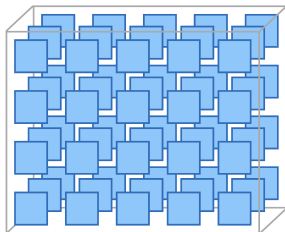
$$X = \begin{bmatrix} 12 & 31 & \cdots & 11 & 2 \\ 5 & 16 & \cdots & 15 & 8 \\ 9 & \mathbf{1} & \cdots & 21 & 10 \\ \mathbf{1} & \mathbf{1} & \cdots & 7 & \mathbf{1} \end{bmatrix} \in \mathbb{R}^{L \times B} \quad (21)$$

- ▶ This is a slightly different representation, compared to BoW
- ▶ B : the mini batch size
- ▶ L : the biggest length of a text in the mini batch
- ▶ Add extra padding tokens will make sentences in the mini batch have the same length

Hidden Layer

With the word index matrix, the hidden layer of a neural network compute the representation of texts with the following steps

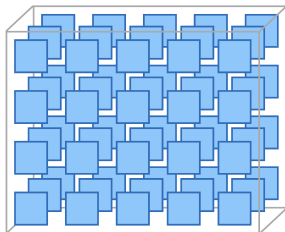
1. Create a mode-3 tensor $\mathcal{X} \in \mathbb{R}^{L \times B \times E}$ by representing each word with a vector (word embedding)



Hidden Layer

With the word index matrix, the hidden layer of a neural network compute the representation of texts with the following steps

1. Create a mode-3 tensor $\mathcal{X} \in \mathbb{R}^{L \times B \times E}$ by representing each word with a vector (word embedding)

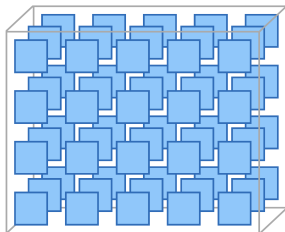


2. Create text representations by summing over the first dimension, which gives a matrix $\mathbf{H}$ with size $B \times E$

Hidden Layer

With the word index matrix, the hidden layer of a neural network compute the representation of texts with the following steps

1. Create a mode-3 tensor $\mathcal{X} \in \mathbb{R}^{L \times B \times E}$ by representing each word with a vector (word embedding)



2. Create text representations by summing over the first dimension, which gives a matrix $\mathbf{H}$ with size $B \times E$
3. Go through an element-wise activation function, e.g., the Sigmoid function.

With the representations of texts $\mathbf{H} \in \mathbb{R}^{B \times E}$, the output layer is nothing different from the logistic regression model.

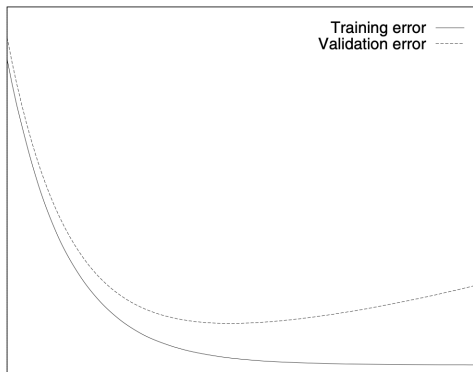
With the representations of texts $\mathbf{H} \in \mathbb{R}^{B \times E}$, the output layer is nothing different from the logistic regression model.

Checkout the demo code!

Neural Networks: Tricks of the Trade

Early Stopping

The idealized training and validation error curves: Vertical axis – errors, horizontal axis – time

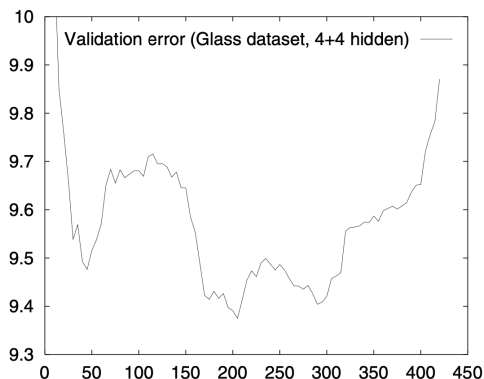


Stop training as soon as the validation error increases

[Montavon et al., 2012]

Early Stopping (II)

A real validation error during training



- ▶ Use check points and always the best model according to the validation performance
- ▶ Stop training if the validation error keep increasing for a *while* (?)

Regularization: Weight Decay

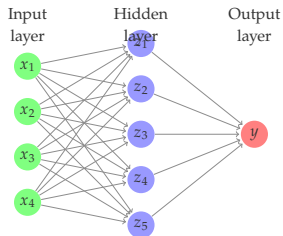
Add a small modification to the gradient of θ , such that

$$\theta^{(\text{new})} \leftarrow \theta^{(\text{old})} - \eta \cdot \left(\frac{\partial L(\theta)}{\partial \theta} \Big|_{\theta=\theta^{(\text{old})}} + \alpha \theta^{(\text{old})} \right) \quad (22)$$

- ▶ It is equivalent to the ℓ_2 regularization
- ▶ Proposed in [Plaut et al., 1986] based on a different intuition

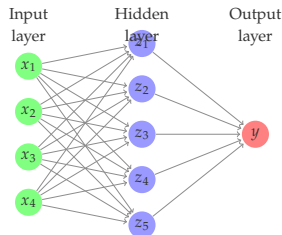
Regularization: Dropout

Randomly remove some components from the input layer during training



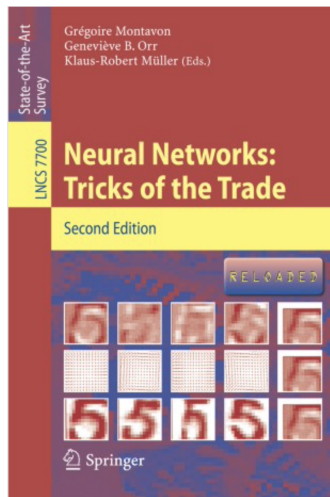
Regularization: Dropout

Randomly remove some components from the input layer during training



- ▶ Proposed in [Hinton et al., 2012] and later published as [Srivastava et al., 2014]
- ▶ Provides a regularization effect that reduce the correlation with noisy features
- ▶ Can also be applied to hidden layers

Neural Networks: Tricks of the Trade



Speeding Learning

Preface	7
1. Efficient BackProp	9
<i>Yann LeCun, Leon Bottou, Genevieve B. Orr, and Klaus-Robert Müller</i>	

Regularization Techniques to Improve Generalization

Preface	49
2. Early Stopping — But When?	53
<i>Lutz Prechelt</i>	
3. A Simple Trick for Estimating the Weight Decay Parameter	69
<i>Thorsteinn S. Rognvaldsson</i>	
4. Controlling the Hyperparameter Search in MacKay's Bayesian Neural Network Framework	91
<i>Tony Plate</i>	
5. Adaptive Regularization in Neural Network Modeling	111
<i>Jan Larsen, Claus Svarer, Lars Nonboe Andersen, and Lars Kai Hansen</i>	
6. Large Ensemble Averaging	131
<i>David Horn, Ury Naftaly, and Nathan Intrator</i>	

1. From Logistic Regression to Neural Networks
2. Learning Neural Networks
3. Neural Text Classifiers
4. Neural Networks: Tricks of the Trade

Reference



Hinton, G. E., Srivastava, N., Krizhevsky, A., Sutskever, I., and Salakhutdinov, R. R. (2012).
Improving neural networks by preventing co-adaptation of feature detectors.
arXiv preprint arXiv:1207.0580.



Jarrett, K., Kavukcuoglu, K., Ranzato, M., and LeCun, Y. (2009).
What is the best multi-stage architecture for object recognition?
In *Proceedings of the 12th International Conference on Computer Vision*, pages 2146–2153. IEEE.



Montavon, G., Orr, G., and Müller, K.-R. (2012).
Neural networks-tricks of the trade second edition.
Springer.



Plaut, D. C., Nowlan, S., and Hinton, G. (1986).
Experiments on learning by back propagation.



Rumelhart, D. E., Hinton, G. E., and Williams, R. J. (1986).
Learning representations by back-propagating errors.
Nature, 323(6088):533–536.



Srivastava, N., Hinton, G., Krizhevsky, A., Sutskever, I., and Salakhutdinov, R. (2014).
Dropout: a simple way to prevent neural networks from overfitting.
The journal of machine learning research, 15(1):1929–1958.