

Lecture slides for this course
have been prepared by Dr. Le Minh Huy,
EEE, Phenikaa University



Deep Learning

Chapter 2 Building Neural Network from Scratch

Dr. Van-Toi NGUYEN
EEE, Phenikaa University

1



Chapter 2: Building Neural Network from Scratch

1. Shallow neural network
2. Deep neural network
3. Building neural network: step-by-step (modulation)
4. Regularization
5. Dropout
6. Batch Normalization
7. Optimizers
8. Hyper-parameters
9. Practice

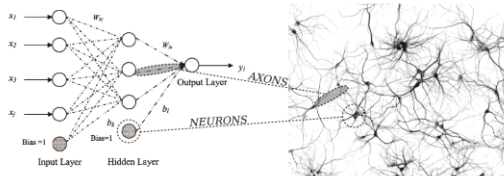
2

Previous Lecture Overview

Biologically inspired (akin to the neurons in a brain)



NEURAL NETWORK MAPPING



These slides are provided by Minh Huy Le, ICNLab, Phenikaa Univ.

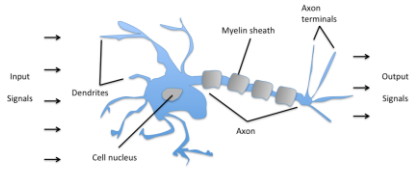
3

3

Previous Lecture Overview



Artificial Neurons and the McCulloch-Pitts Model (1943)



Schematic of a biological neuron.

W. S. McCulloch and W. Pitts. A logical calculus of the ideas immanent in nervous activity. The bulletin of mathematical biophysics, 5(4):115-133, 1943.

These slides are provided by Mubhayee Le, ICSTLab, Phenikaa UTM.

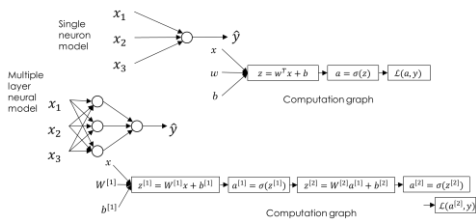
4

4

Previous Lecture Overview



What is a Shallow Neural Network?



These slides are provided by Mubhayee Le, ICSTLab, Phenikaa UTM.

5

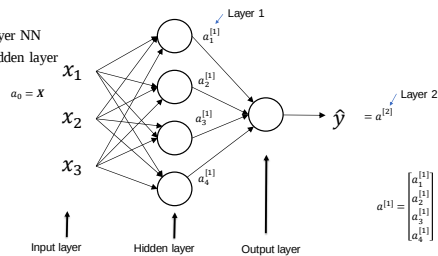
5

Previous Lecture Overview

One hidden layer Neural Network



- 2-layer NN
- 1 hidden layer



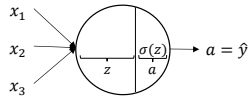
These slides are provided by Mubhayee Le, ICSTLab, Phenikaa UTM.

6

6

Previous Lecture Overview

Computing NN's Output



$$z = w^T x + b$$

$$a = \sigma(z)$$

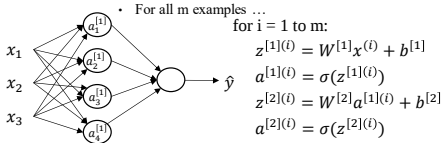
These slides are provided by Minhuy Le, ICSLab, Phenikaa Utd.

7

7

Previous Lecture Overview

Vectorizing across multiple examples



These slides are provided by Minhuy Le, ICSLab, Phenikaa Utd.

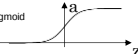
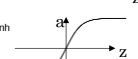
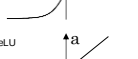
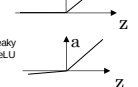
8

8

Previous Lecture Overview

Activation functions



Activation Function	Formula (g(z))	Derivative (g'(z))	
sigmoid	$a = \frac{1}{1 + e^{-z}}$	$a(1 - a)$	sigmoid 
tanh	$a = \frac{e^z - e^{-z}}{e^z + e^{-z}}$	$1 - a^2$	tanh 
ReLU	$\max(0, z)$	0 if $z < 0$ 1 if $z \geq 0$	ReLU 
Leaky ReLU	$\max(0.01z, z)$	0.01 if $z < 0$ 1 if $z \geq 0$	Leaky ReLU 

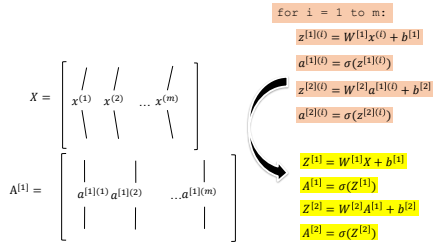
These slides are provided by Minhuy Le, ICSLab, Phenikaa Utd.

9

9

Previous Lecture Overview

Vectorizing across multiple examples



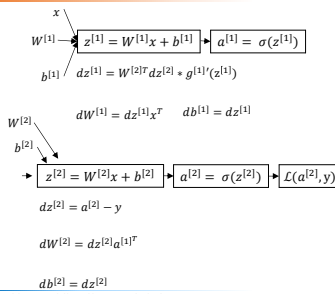
These slides are provided by Minhuy Le, ICSLab, Phenikaa Utd.

10

10

Previous Lecture Overview

Gradient descent for one hidden layer



These slides are provided by Minhuy Le, ICSLab, Phenikaa Utd.

11

11

Previous Lecture Overview

Vectorizing Gradient Descent



$$dz^{[2]} = a^{[2]} - y$$

$$dW^{[2]} = dz^{[2]} a^{[1]T}$$

$$db^{[2]} = dz^{[2]}$$

$$dz^{[1]} = W^{[2]T} dz^{[2]} + g^{[1]'}(z^{[1]})$$

$$dW^{[1]} = dz^{[1]} x^T$$

$$db^{[1]} = dz^{[1]}$$

$$dZ^{[2]} = A^{[2]} - Y$$

$$dW^{[2]} = \frac{1}{m} dz^{[2]} A^{[1]T}$$

$$db^{[2]} = \frac{1}{m} \text{np.sum}(dz^{[2]}, \text{axis}=1, \text{keepdims}=True)$$

$$dZ^{[1]} = W^{[2]T} dz^{[2]} + g^{[1]'}(Z^{[1]})$$

$$dW^{[1]} = \frac{1}{m} dZ^{[1]} x^T$$

$$db^{[1]} = \frac{1}{m} \text{np.sum}(dZ^{[1]}, \text{axis}=1, \text{keepdims}=True)$$

These slides are provided by Minhuy Le, ICSLab, Phenikaa Utd.

12

12

2. Deep neural network

Deep Learning???



These slides are provided by Minhuy Le, ICSLab, Phenikaa Uin.

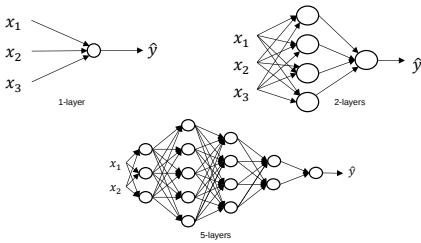


13

13

2. Deep neural network

Multiple layers



These slides are provided by Minhuy Le, ICSLab, Phenikaa Uin.



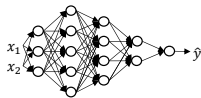
14

14

2. Deep neural network

Notations in Deep Neural Networks

- #Layers = $L = 5$
- $n^{[l]}$ = #units in layer l
 - $n^{[1]} = 3, n^{[2]} = 5, n^{[3]} = 4, n^{[4]} = 2, n^{[5]} = n^{[L]} = 1$
 - $n^{[0]} = n^{[L]} = 2$
- $a^{[l]} = g^{[l]}(z^{[l]})$
- $W^{[l]}$ = weights for computing $z^{[l]}$
- $b^{[l]}$ = bias for computing $z^{[l]}$
- $a^{[0]} = X$ (input)
- $a^{[L]} = \hat{y}$ (prediction output)



These slides are provided by Minhuy Le, ICSLab, Phenikaa Uin.



15

15

2. Deep neural network

Dimensions of vectorized implementations



- For one single training example:
 - $z^{[l]} = W^{[l]}a^{[l-1]} + b^{[l]}$
 - $(n^{[l]}, 1) = (n^{[l]}, n^{[l-1]}) \times (n^{[l-1]}, 1) + (n^{[l]}, 1)$
- For a vectorized implementation over m examples
 - $Z^{[l]} = W^{[l]}A^{[l-1]} + b^{[l]}$
 - $(n^{[l]}, m) = (n^{[l]}, n^{[l-1]}) \times (n^{[l-1]}, m) + (n^{[l]}, m)$

Matrix	Dimensions
$Z^{[l]}, A^{[l]}, b^{[l]}, dZ^{[l]}, dA^{[l]}, db^{[l]}$	$(n^{[l]}, m)$
$W^{[l]}, dW^{[l]}$	$(n^{[l]}, n^{[l-1]})$

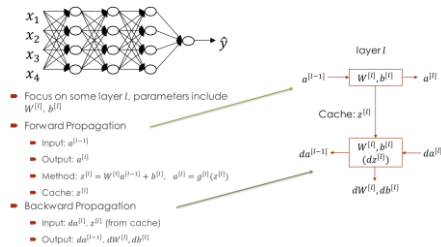
These slides are provided by Mithun Le, ICSLab, Phenikaa Utd.

16

16

2. Deep neural network

Building Blocks of Deep Neural Networks



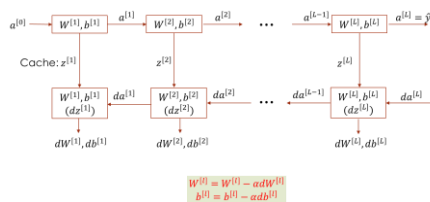
These slides are provided by Mithun Le, ICSLab, Phenikaa Utd.

17

17

2. Deep neural network

Forward and backward functions



These slides are provided by Mithun Le, ICSLab, Phenikaa Utd.

18

18

2. Deep neural network

Forward Propagation for DNN



- Input: $a^{[l-1]}$
- Forward Propagation:
 - For single example:

$$z^{[l]} = W^{[l]}a^{[l-1]} + b^{[l]}, \quad a^{[l]} = g^{[l]}(z^{[l]})$$
 - Vectorized version:

$$Z^{[l]} = W^{[l]}A^{[l-1]} + b^{[l]}, \quad A^{[l]} = g^{[l]}(Z^{[l]})$$
- Output: $a^{[l]}$
- Cache: $z^{[l]}, W^{[l]}, b^{[l]}$

These slides are provided by Minhuy Le, ICSLab, Phenikaa Utd.

19

19

2. Deep neural network

Backward Propagation for DNN



- Inputs: $da^{[l]}, z^{[l]}, W^{[l]}, b^{[l]}$ (last 3 from cache)
- Outputs: $dA^{[l-1]}, dW^{[l]}, db^{[l]}$
- Backward Propagation:
 - For **single example**

$$dz^{[l]} = da^{[l]} * g^{[l]'}(z^{[l]})$$

$$dW^{[l]} = dz^{[l]} a^{[l-1]}$$

$$db^{[l]} = dz^{[l]}$$

$$da^{[l-1]} = W^{[l]T} dz^{[l]}$$

$$dz^{[l]} = W^{[l+1]T} dz^{[l+1]} * g^{[l]'}(z^{[l]})$$

Vectorized implementation

$$dZ^{[l]} = dA^{[l]} * g^{[l]'}(Z^{[l]})$$

$$dW^{[l]} = \frac{1}{m} dz^{[l]} A^{[l-1]T}$$

$$db^{[l]} = \frac{1}{m} np.sum(dZ^{[l]}, axis = 1, keepdims = True)$$

$$dA^{[l-1]} = W^{[l]T} dz^{[l]}$$

These slides are provided by Minhuy Le, ICSLab, Phenikaa Utd.

20

20

2. Deep neural network

Parameters vs. Hyperparameters



- Parameters: $W^{[1]}, b^{[1]}, W^{[2]}, b^{[2]}, \dots$
- Hyperparameters
 - Learning rate: α
 - Number of iterations
 - Number of hidden layers or L
 - Number of hidden units in each layer: $n^{[1]}, n^{[2]}, \dots$
 - Choice of activation function: sigmoid, ReLU, tanh, etc.
 - Momentum, mini-batch size, regularization parameters, ... (in the next Chapter)

These slides are provided by Minhuy Le, ICSLab, Phenikaa Utd.

21

21

2. Deep neural network

Summary of Forward/Backward Computations



$Z^{[1]} = W^{[1]}X + b^{[1]}$	$dZ^{[L]} = A^{[L]} - Y$
$A^{[1]} = g^{[1]}(Z^{[1]})$	$dW^{[L]} = \frac{1}{m} dZ^{[L]} A^{[L]T}$
$Z^{[2]} = W^{[2]}A^{[1]} + b^{[2]}$	$db^{[L]} = \frac{1}{m} np.sum(dZ^{[L]}, axis = 1, keepdims = True)$
$A^{[2]} = g^{[2]}(Z^{[2]})$	$dZ^{[L-1]} = dW^{[L]T} dZ^{[L]} g'^{[L]}(Z^{[L-1]})$
$\vdots$	$\vdots$
$A^{[L]} = g^{[L]}(Z^{[L]}) = \hat{y}$	$dZ^{[1]} = dW^{[L]T} dZ^{[2]} g'^{[1]}(Z^{[1]})$
	$dW^{[1]} = \frac{1}{m} dZ^{[1]} A^{[1]T}$
	$db^{[1]} = \frac{1}{m} np.sum(dZ^{[1]}, axis = 1, keepdims = True)$

Forward Propagation

Backward Propagation

These slides are provided by Minhuy Le, ICSLab, Phenikaa Utd.

22

22

3. Building neural network: step-by-step (modulation)



On Assignments
Coding time !!!

These slides are provided by Minhuy Le, ICSLab, Phenikaa Utd.

23

23