

Neural networks 2018 Fall
paper-based closed room test
REPLACEMENT

PPCU-ITK
3rd December, 2018

Name: Ekarf CYBA

NEPTUN ID: 7409140

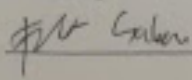
Question	Points	Student's points
1	12	12
2	12	8
3	12	12
4	30	30
Total	66	62

94%

Instructions:

1. This examination contains 6 pages, including this page.
2. You have **ninety (90) minutes** to complete the examination. As a courtesy to your classmates, we ask that you not leave during the last fifteen minutes.
3. You may not use any external resources, including lecture notes, books, other students or other engineers.
4. You may use a calculator. You may not share a calculator with anyone.
5. Please sign the below statement.

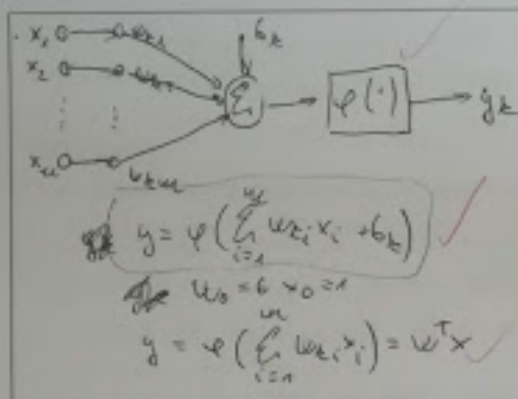
I hereby certify that I will neither give nor receive unpermitted aid on this examination.

Signature: 

Question 1: The artificial neuron

Overall: [12 pts]

- (a) [5 pts] Please write down the working steps of an artificial neuron! (What happens when the perceptron receives input(s) through its synapse?)



0. receives input vector on its synapse ✓
1. multiply it with the corresponding weights ✓
2. sum the values ✓
3. add the bias (for non-linearity) ✓
4. apply activation function on the whole (also for non-linearity) ✓
5. output ✓

5

- (b) [3 pts] State the perceptron convergence theorem! (No proof required.)

The perceptron convergence theorem states that the Perceptron Learning Algorithm a.k.a. Rosenblatt algorithm converges in finite number of steps for a linearly separable dataset.

3

- (c) [4 pts] Write down the output equation of a multi-layer perceptron (a.k.a. fully connected neural network)!

$$\text{Net}(x, w) = \varphi^{(L)}(w^{(L)} \varphi^{(L-1)}(w^{(L-1)} (\varphi^{(L-2)}(\dots \varphi^{(2)}(w^{(2)} (\varphi^{(1)}(w^{(1)} x)) \dots)))$$

4

12

Question 2: Regularization techniques

Overall: [12 pts]

- (a) [4 pts] Please describe the dropout normalization technique in details!

Uses mini-batch approach, and for each minibatch select a random subset of neurons which will be deactivated in the training. Selection & deactivation probability: p .
When testing multiply the outputs with p to account the deactivated neurons.

4

- (b) [5 pts] What is local response normalization? Please give your answer in detail.

~~I don't know what is local response normalization, but I know it's used for batch-normal, input vector normalization, & for the weights (right now)~~
We want to normalize to avoid overfitting & make the network less sensitive to errors.
Then we learned about input vector norm, batch norm. and weight norm.

In normalization processes we want to squeeze the range, where these values exist.

Input vector norm:

x : input vector
 $\bar{x}$: mean
 σ : deviation

$$x_{\text{norm}} = \frac{x - \bar{x}}{\sigma}$$

L_1 & L_2 norm

$$C_0 = \frac{1}{2n} \sum_i \|y - a\|^2 \quad C_2 = \frac{1}{2n} \sum_i \|y - a\|^2 + \frac{\lambda}{2n} \sum_i w^2$$

$$C_1 = \frac{1}{2n} \sum_i \|y - a\|^1 + \frac{\lambda}{2n} \sum_i |w|$$

Batch norm

$$\hat{x} = \frac{x - E[x]}{\sqrt{\text{Var}[x]}}$$

YES
THIS IS
WHY
WE USE
NORMALIZATION
E.g. input
variability

- (c) [3 pts] Write 3 common loss functions!

Quadratic-loss function:

$$E(w) = \frac{1}{K} \sum_{k=1}^K (d_k - \text{Net}(x_k, w))^2$$

Negative-log likelihood

$$L(w) = \frac{1}{K} \sum_{k=1}^K -\log(P(y_k | x_k, w))$$

Cross entropy:

$$C = -\frac{1}{n} \sum_x [y \ln a + (1-y) \ln (1-a)]$$

$$C(w) = -\frac{1}{K} \sum_{k=1}^K [d_k \ln P(y_k | x_k, w) + (1-d_k) \ln (1 - P(y_k | x_k, w))]$$

$$= -\log(P(y_k | x_k, w)) - \log(1 - P(y_k | x_k, w))$$

3

Question 3: Optimization techniques

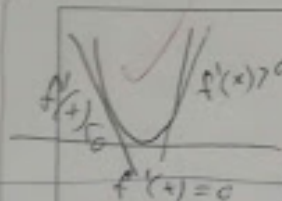
Overall: [12pts]

- (a) [3pts] What is the mathematical definition of the Newton method (first order gradient based optimization method)?

$$\Delta x = -H(f(x(u)))^{-1} \nabla f(x(u)) \quad x(u+1) = x(u) - \eta H(f(x(u)))^{-1} \nabla f(x(u))$$

3

- (b) [5pts] Write down the general idea of gradient descent in optimization problems!



Goal: find local minima
 → Start from a random point and calculate the gradient of the function.
 → always follow the ~~descending~~ descending gradient.
 → ~~follows~~ follows the chain-rule

100%
 BATCHES
 LEARNING
 RATE

- (c) [4pts] Explain the minibatch approach in gradient-based optimization algorithms!

3 ~~approach~~ approaches: instant update, batch, minibatch
 If we use the minibatch approach, it means, that we select a random set from the input vectors, calculate the updates, then take their average and apply.

4

→ update weights!

If ~~the~~ $f'(x) < 0$ we should increase x to get closer to the min.
 if $f'(x) > 0$ we should decrease x to get closer to the min.

m_b : size of minibatch usually ~ 100

$$L = \frac{1}{m_b} \sum_{k=1}^{m_b} [d_k - \text{Net}(x, w)]^2$$

72

Question 4: Mathematical methods in practice

Overall: [30pts]

- (a) [12pts] Apply the Rosenblatt perceptron training algorithm to learn the implication logical function $f(x,y) = \neg x \vee y$ (NOT x or Y) with a single perceptron! Validate your results graphically! Please show your work.

x_1	x_2	$f(x,y)$
-1	-1	0
-1	1	1
1	-1	1
1	1	1

$\phi^k = \text{sign}(\cdot)$
 $\eta = \frac{1}{2}$
 $u^k(0) = 0$
 $y^k(k) = \phi(u^T x(k))$
 $b = 1$

$\oplus: u^k(k+1) = u^k(k) + \eta x(k)$
 $\ominus: u^k(k+1) = u^k(k) - \eta x(k)$

$y(0) = \text{sign}([0 \ 0 \ 0] \begin{bmatrix} 1 \\ -1 \\ -1 \end{bmatrix}) = 0$
 $E(0) = d_0 - y(0) = 1 - 0 = 1 \Rightarrow \oplus \rightarrow u(1) = \begin{bmatrix} 0 \\ 0 \\ 0 \end{bmatrix} + \frac{1}{2} \begin{bmatrix} 1 \\ -1 \\ -1 \end{bmatrix} = \begin{bmatrix} 1/2 \\ -1/2 \\ -1/2 \end{bmatrix}$

$y(1) = \text{sign}(\begin{bmatrix} 1/2 & -1/2 & -1/2 \end{bmatrix} \begin{bmatrix} 1 \\ -1 \\ -1 \end{bmatrix}) = 1$
 $E(1) = d_1 - y(1) = 1 - 1 = 0 \Rightarrow OK$

$y(2) = \text{sign}(\begin{bmatrix} 1/2 & -1/2 & -1/2 \end{bmatrix} \begin{bmatrix} 1 \\ 1 \\ -1 \end{bmatrix}) = 1$
 $E(2) = d_2 - y(2) = 1 - 1 = 0 \Rightarrow OK$

$y(3) = \text{sign}(\begin{bmatrix} 1/2 & -1/2 & -1/2 \end{bmatrix} \begin{bmatrix} 1 \\ 1 \\ 1 \end{bmatrix}) = 1$
 $E(3) = d_3 - y(3) = 1 - 1 = 0 \Rightarrow OK$

$y(4) = \text{sign}(\begin{bmatrix} 1/2 & -1/2 & -1/2 \end{bmatrix} \begin{bmatrix} -1 \\ -1 \\ -1 \end{bmatrix}) = 0$
 $E(4) = d_4 - y(4) = 1 - 0 = 1 \Rightarrow \oplus \rightarrow u(5) = \begin{bmatrix} 1/2 \\ -1/2 \\ -1/2 \end{bmatrix} + \frac{1}{2} \begin{bmatrix} -1 \\ -1 \\ -1 \end{bmatrix} = \begin{bmatrix} 0 \\ -1 \\ -1 \end{bmatrix}$

$y(5) = \text{sign}(\begin{bmatrix} 0 & -1 & -1 \end{bmatrix} \begin{bmatrix} 1 \\ -1 \\ -1 \end{bmatrix}) = 1$
 $E(5) = d_5 - y(5) = 1 - 1 = 0 \Rightarrow OK$

$y(6) = \text{sign}(\begin{bmatrix} 0 & -1 & -1 \end{bmatrix} \begin{bmatrix} 1 \\ 1 \\ -1 \end{bmatrix}) = 1$
 $E(6) = d_6 - y(6) = 1 - 1 = 0 \Rightarrow OK$

$E(6) = 0 \Rightarrow OK$

- (b) [2pts] If we are able to learn this logical function using a single perceptron, what property of the learned function does this imply?

Linear separability

$\text{sign}(\begin{bmatrix} 0 & -1 & -1 \end{bmatrix} \begin{bmatrix} 1 \\ 1 \\ 1 \end{bmatrix}) = 1$
 $E(7) = d_7 - y(7) = 1 - 1 = 0 \Rightarrow OK$

$y(8) = \text{sign}(\begin{bmatrix} 0 & -1 & -1 \end{bmatrix} \begin{bmatrix} -1 \\ -1 \\ -1 \end{bmatrix}) = -1$
 $E(8) = d_8 - y(8) = 1 - (-1) = 2 \Rightarrow \oplus \rightarrow u(9) = \begin{bmatrix} 0 \\ -1 \\ -1 \end{bmatrix} + \frac{1}{2} \begin{bmatrix} -1 \\ -1 \\ -1 \end{bmatrix} = \begin{bmatrix} -1/2 \\ -3/2 \\ -3/2 \end{bmatrix}$

$y(9) = \text{sign}(\begin{bmatrix} -1/2 & -3/2 & -3/2 \end{bmatrix} \begin{bmatrix} 1 \\ -1 \\ -1 \end{bmatrix}) = 1$
 $E(9) = d_9 - y(9) = 1 - 1 = 0 \Rightarrow OK$

$y(10) = \text{sign}(\begin{bmatrix} -1/2 & -3/2 & -3/2 \end{bmatrix} \begin{bmatrix} 1 \\ 1 \\ -1 \end{bmatrix}) = 1$
 $E(10) = d_{10} - y(10) = 1 - 1 = 0 \Rightarrow OK$

$y(11) = \text{sign}(\begin{bmatrix} -1/2 & -3/2 & -3/2 \end{bmatrix} \begin{bmatrix} 1 \\ 1 \\ 1 \end{bmatrix}) = -1$
 $E(11) = d_{11} - y(11) = 1 - (-1) = 2 \Rightarrow \oplus \rightarrow u(12) = \begin{bmatrix} -1/2 \\ -3/2 \\ -3/2 \end{bmatrix} + \frac{1}{2} \begin{bmatrix} 1 \\ 1 \\ 1 \end{bmatrix} = \begin{bmatrix} 0 \\ -1 \\ -1 \end{bmatrix}$

$y(12) = \text{sign}(\begin{bmatrix} 0 & -1 & -1 \end{bmatrix} \begin{bmatrix} 1 \\ -1 \\ -1 \end{bmatrix}) = 1$
 $E(12) = d_{12} - y(12) = 1 - 1 = 0 \Rightarrow OK$

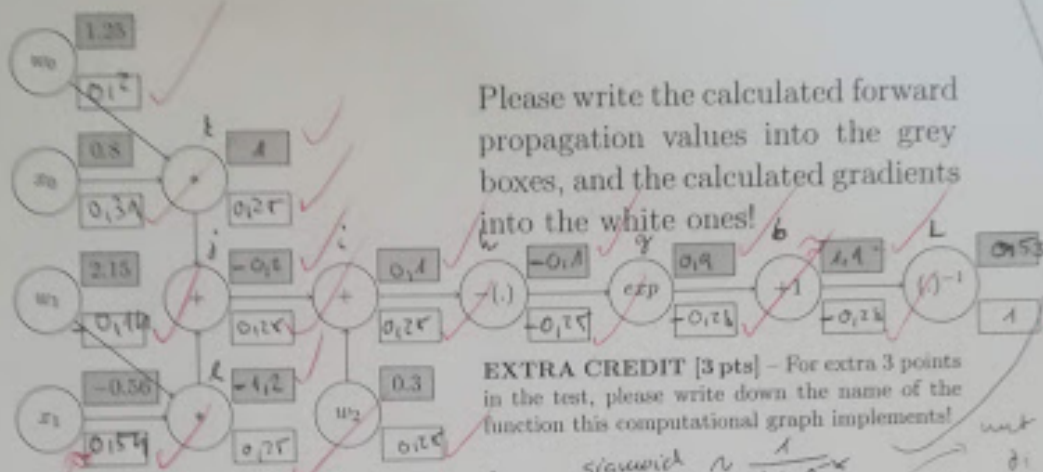
$W^T x = 0$
 $0 = \begin{bmatrix} 1/2 \\ -1/2 \\ 1/2 \end{bmatrix}^T \begin{bmatrix} x_1 \\ x_2 \\ x_3 \end{bmatrix} = \frac{1}{2}x_1 - \frac{1}{2}x_2 + \frac{1}{2}x_3$
 $-\frac{1}{2}x_1 + \frac{1}{2}x_2 + \frac{1}{2} = 0$

$$L = (w_0 x_0 + w_1 x_1 + w_2) \cdot A$$

(e) [16 (13+3) pts] Calculate the gradient for each node of the computational graph below! Please show your work! Also, here is some help with derivation:

$$f(x) = e^x \rightarrow \frac{df}{dx} = e^x \quad f(x) = ax \rightarrow \frac{df}{dx} = a$$

$$f(x) = \frac{1}{x} \rightarrow \frac{df}{dx} = -\frac{1}{x^2} \quad f(x) = c + x \rightarrow \frac{df}{dx} = 1$$



EXTRA CREDIT [3 pts] - For extra 3 points in the test, please write down the name of the function this computational graph implements!

Answer: sigmoid $\sim \frac{1}{1+e^x}$

$$\frac{\partial L}{\partial L} = 1$$

$$\frac{\partial L}{\partial b} = \frac{\partial L}{\partial L} \cdot \frac{\partial L}{\partial b} = 1 \cdot \left(-\frac{1}{(1.9)^2}\right) = -0.28$$

$$\frac{\partial L}{\partial y} = \frac{\partial L}{\partial L} \cdot \frac{\partial L}{\partial y} = (-0.28) \cdot 1 = -0.28$$

$$\frac{\partial L}{\partial u} = \frac{\partial L}{\partial L} \cdot \frac{\partial L}{\partial u} = (-0.28) \cdot e^{-0.1} = -0.25$$

$$\frac{\partial L}{\partial i} = \frac{\partial L}{\partial L} \cdot \frac{\partial L}{\partial i} = (-0.28) \cdot (-1) = 0.28$$

$$\frac{\partial L}{\partial w_0} = \frac{\partial L}{\partial L} \cdot \frac{\partial L}{\partial w_0} = 0.28 \cdot 1.25 = 0.35$$

$$\frac{\partial L}{\partial w_1} = \frac{\partial L}{\partial L} \cdot \frac{\partial L}{\partial w_1} = 0.28 \cdot 2.15 = 0.60$$

$$\frac{\partial L}{\partial w_2} = \frac{\partial L}{\partial L} \cdot \frac{\partial L}{\partial w_2} = 0.28 \cdot (-0.56) = -0.16$$

$$\frac{\partial L}{\partial w_0} = \frac{\partial L}{\partial L} \cdot \frac{\partial L}{\partial w_0} = 0.28 \cdot (-0.56) = -0.16$$

$$\frac{\partial L}{\partial w_1} = \frac{\partial L}{\partial L} \cdot \frac{\partial L}{\partial w_1} = 0.28 \cdot 1.25 = 0.35$$

$$\frac{\partial L}{\partial w_2} = \frac{\partial L}{\partial L} \cdot \frac{\partial L}{\partial w_2} = 0.28 \cdot 2.15 = 0.60$$

$$\frac{\partial L}{\partial w_0} = \frac{\partial L}{\partial L} \cdot \frac{\partial L}{\partial w_0} = 0.28 \cdot 2.15 = 0.60$$

$$\frac{\partial L}{\partial w_1} = \frac{\partial L}{\partial L} \cdot \frac{\partial L}{\partial w_1} = 0.28 \cdot (-0.56) = -0.16$$

$$\frac{\partial L}{\partial w_2} = \frac{\partial L}{\partial L} \cdot \frac{\partial L}{\partial w_2} = 0.28 \cdot 1.25 = 0.35$$

16