

PURDUE

BME 64600 – 001 and ECE 60146 – 001

Midterm #2, Spring 2025

NAME _____

PUID _____

Exam instructions:

- You have 75 minutes to work the exam.
- This is a closed-book and closed-note exam. You may not use or have access to your book, notes, any supplementary reference, a calculator, or any communication device including a cell-phone or computer.
- You may not communicate with any person other than the official proctor during the exam.

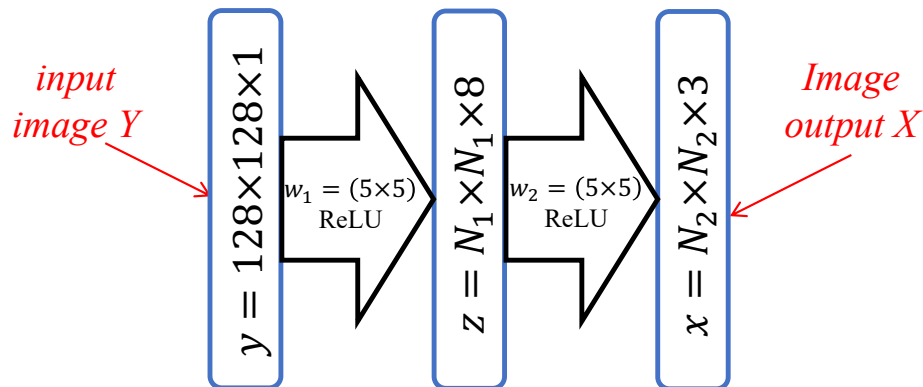
To ensure Gradescope can read your exam:

- Write your full name and PUID above and on the top of every page.
- Answer all questions in the area designated for each problem.
- Write only on the front of the exam pages.
- DO NOT run over to the next question.

Name/PUID: _____

Problem 1. (30pt) Parameters of Convolutional Neural Networks

Consider the following convolutional neural network pictured below with a gray-scale image input and a color image output. Each layer uses a ReLU activation function, a “valid” boundary condition, and denote the convolution kernels by w_1 and w_2 and the associated offsets by b_1 and b_2 .



1a) Calculate the value of N_1 .

Solution: $N_1 = 124$

1b) Calculate the value of N_2 .

Solution: $N_2 = 120$

1c) Calculate the shape of w_1 .

Solution: $5 \times 5 \times 1 \times 8$

1d) Calculate the shape of w_2 .

Solution: $5 \times 5 \times 8 \times 3$

1e) Calculate the shape of b_1 .

Solution: 8

1f) Calculate the total number of parameters in the model.

Solution: Number of parameters in each layer:

- layer1:
 - filter: $5 \times 5 \times 1 \times 8 = 200$
 - offset: 8
- layer2:
 - filter: $5 \times 5 \times 8 \times 3 = 600$
 - offset: 3

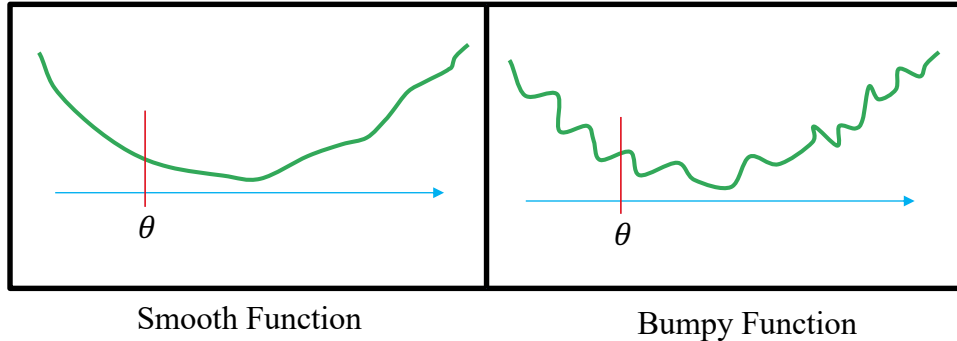
Total number of parameters:

$$200 + 8 + 600 + 3 = 811$$

Name/PUID: _____

Problem 2. (15pt) Batch Size

Consider the following “smooth” function and “bumpy” function illustrated below.



2a) Which function has more local minima?

Solution: The bumpy function has more local minima

2b) For which function is it better to use a **large** batch size? Why?

Solution: The large batch size is better for the smooth function because it reduces that random variation in the gradient. This in turn allow the global minimum to be more precisely computed in the final “exploitation” stage of optimization.

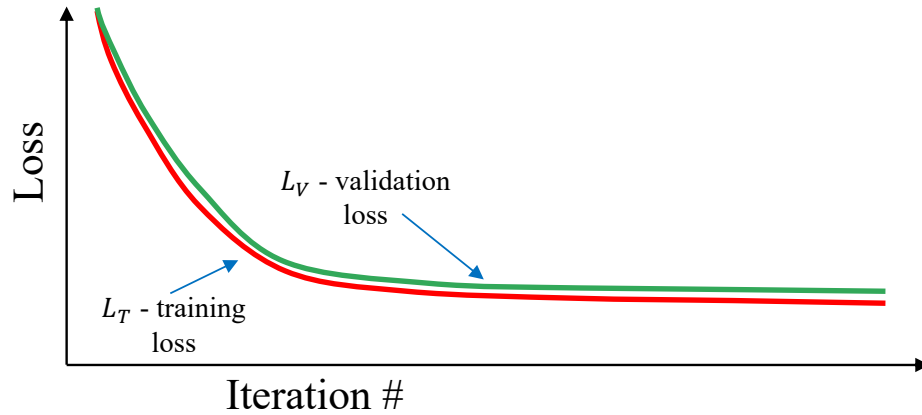
2c) For which function is it better to use a **small** batch size? Why?

Solution: The small batch size is better for the bumpy function because for two reasons: a) It generates more noise in the gradient that helps jump out of local minima. This improves the “exploration” portion of the optimization in the early stages of optimization; b) It speeds evaluation of the gradient, so the algorithm can perform more gradient steps per unit time.

Name/PUID: _____

Problem 3. (10pt) Training Convergence

Consider the following convergence functions.



3a) What is the data telling you?

Solution: This is telling you that even after many iterations, the training loss is not significantly better than the validation loss. This means two things: a) There is sufficient training data for the capacity of the model; b) There model may not have sufficient capacity to best solve the ML problem.

3b) What actions should you take to improve your results? Why?

Solution: Based on this it would be good to increase the capacity of the model by, for example, increasing the number of layers, increasing the size of the convolution kernel, adding blocks to the networks that make it more expressive.

Name/PUID: _____

Problem 4. (35pt) Back Propagation

A neural network can be represented by

$$x = f(y) = Ay + b ,$$

where $y \in \mathbb{R}^N$ is the input, $x \in \mathbb{R}^N$, $A \in \mathbb{R}^{N \times N}$, $b \in \mathbb{R}^N$, and $N = 4$. Furthermore, assume that

$$A_{i,j} = w_{i-j} ,$$

for some function w_i . Also define the loss function

$$L(y) = \frac{1}{2K} \sum_{k=0}^{K-1} \|x_k - f(y)\|^2 ,$$

where $\{x_k\}_{k=0}^{K-1}$ are K training samples.

4a) Write out the matrix A in terms of w_i .

Solution:

$$A = \begin{bmatrix} w_0 & w_{-1} & w_{-2} & w_{-3} \\ w_1 & w_0 & w_{-1} & w_{-2} \\ w_2 & w_1 & w_0 & w_{-1} \\ w_3 & w_2 & w_1 & w_0 \end{bmatrix}$$

4b) Describe in words the meaning of the operation Ay .

Solution: The operation Ay corresponds to convolution of

$$y = [y_0, y_1, y_2, y_3] ,$$

with the kernel

$$w = [w_{-3}, w_{-2}, w_{-1}, w_0, w_1, w_2, w_3] ,$$

using circular convolution with the “same” boundary condition.

4c) Write out an explicit expression for the adjoint gradient of the function $[\nabla_y f(y)]^t$ in terms of A .

Solution:

$$[\nabla_y f(y)]^t = [A]^t = A^t$$

4d) Write out the matrix A^t in terms of w_i .

Solution:

$$A = \begin{bmatrix} w_0 & w_1 & w_2 & w_3 \\ w_{-1} & w_0 & w_1 & w_2 \\ w_{-2} & w_{-1} & w_0 & w_1 \\ w_{-3} & w_{-2} & w_{-1} & w_0 \end{bmatrix}$$

4e) Describe in words the meaning of the operation $[\nabla_y f(y)]^t \epsilon$.

Solution: Since

$$[\nabla_y f(y)]^t \epsilon = A^t \epsilon ,$$

we see that this represents convolution of

$$\epsilon = [\epsilon_0, \epsilon_1, \epsilon_2, \epsilon_3] ,$$

with the kernel

$$w = [w_3, w_2, w_1, w_0, w_{-1}, w_{-2}, w_{-3}] ,$$

using circular convolution with the “same” boundary condition.

4f) Write out an explicit expression for the adjoint gradient of the loss function $[\nabla_y L(y)]^t$ in terms of A , x_k , and $f(y)$.

(Hint, this should result in a column vector.)

Solution:

$$\begin{aligned} [\nabla_y L(y)]^t &= \left[\nabla_y \frac{1}{2K} \sum_{k=0}^{K-1} \|x_k - f(y)\|^2 \right]^t \\ &= \left[\frac{1}{2K} \sum_{k=0}^{K-1} \nabla_y \|x_k - f(y)\|^2 \right]^t \\ &= \left[\frac{1}{K} \sum_{k=0}^{K-1} [x_k - f(y)]^t A \right]^t \\ &= \frac{1}{K} \sum_{k=0}^{K-1} A^t [x_k - f(y)] \\ &= A^t \frac{1}{K} \sum_{k=0}^{K-1} [x_k - f(y)] \end{aligned}$$

4g) Describe in words how you evaluate the function, $[\nabla_y L(y)]^t$.

Solution: First you average all the prediction errors $[x_k - f(y)]$ for each training sample. Then you perform a convolution with the time reverse kernel w_{-i} using “same” boundary conditions.

Name/PUID: _____

Problem 5. (20pt) Bayes Classifier

Let $Y \in \mathbb{R}^N$ be a random vector whose distribution depends on the class, C . Moreover, C is a binary random variable such that $P\{C = 0\} = \pi_0$ and $P\{C = 1\} = 1 - \pi_0$.

Then when $C = 0$, the distribution is given by $Y \sim p_0(y)$, and when $C = 1$, the distribution is given by $Y \sim p_1(y)$.

Also define the likelihood ratio

$$R(y) = \frac{p_1(y)}{p_0(y)}.$$

5a) Calculate the conditional distribution of Y given C denoted by $p(y|c)$.

Solution:

$$p(y|c) = \begin{cases} p_0(y) & \text{if } c = 0 \\ p_1(y) & \text{if } c = 1 \end{cases}$$

5b) Calculate $p(y)$, the marginal distribution of Y .

Solution:

$$p(y) = p_0(y)\pi_0 + p_1(y)(1 - \pi_0)$$

5c) Calculate the conditional probability of $P\{C = 0|Y = y\} = f(y)$ in terms of the likelihood ratio $R(y)$.

Solution:

$$\begin{aligned} f(y) &= P\{C = 0|Y = y\} \\ &= \frac{p_0(y)\pi_0}{p_0(y)\pi_0 + p_1(y)\pi_1} \\ &= \frac{\pi_0}{\pi_0 + R(y)\pi_1} \end{aligned}$$

5d) Specify an expression for the Bayes classifier, $B(y)$, that classifies Y as either $C = 0$ or $C = 1$ with the minimum probability of classification error in terms of $f(y)$.

Solution: The Bayes classifier is given by

$$B(y) = \begin{cases} 0 & \text{if } f(y) \geq 1/2 \\ 1 & \text{if } f(y) < 1/2 \end{cases}$$

Name/PUID: _____

Problem 6. (25pt) Nash Equilibrium

Consider the two loss functions for the generator and discriminator in a GAN defined below

$$\begin{aligned} g(\theta_g, \theta_d) \\ d(\theta_g, \theta_d) , \end{aligned}$$

where g and d are loss functions, and θ_g and θ_d are the parameters of the generator and discriminator, respectively. Next define the following function

$$D(\theta_g) = \arg \min_{\theta_d} d(\theta_g, \theta_d) .$$

6a) Write the equations for the Nash equilibrium using the two loss functions g and d .

Solution:

$$\begin{aligned} \theta_g^* &= \arg \min_{\theta_g} g(\theta_g, \theta_d^*) \\ \theta_d^* &= \arg \min_{\theta_d} d(\theta_g^*, \theta_d) \end{aligned}$$

6b) Give an interpretation for the function $D(\theta_g)$.

Solution: The function $D(\theta_g)$ determines the best discriminator parameter, θ_d^* , for each choice of generator parameter, θ_g .

6c) Substitute the function $D(\theta_g)$ into the generator loss function and write a single equation for the Nash equilibrium solution for the generator parameter.

Solution:

$$\theta_g^* = \arg \min_{\theta_g} g(\theta_g, D(\theta_g))$$

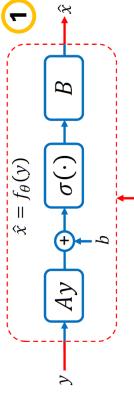
6d) What do we call the special case when $g(\theta_g, \theta_d) = -d(\theta_g, \theta_d)$?

Solution: Zero-sum game

6e) Give an example of an activity people engage in that can be modeled as a Nash equilibrium with $g(\theta_g, \theta_d) = -d(\theta_g, \theta_d)$.

Solution: Any game, such as chess or checkers or go, in which only one player can win is a zero-sum game.

Single Layer NN Graphically:



Where: $f_\theta(y) = B\sigma(Ay + b)$

M-dimensional simplex:

$$S^M = \left\{ x \in \mathbb{R}^M : \forall i, x_i \geq 0 \text{ and } 1 = \sum_{i=0}^{M-1} x_i \right\}$$

Softmax: $\sigma_i(z) = \frac{e^{z_i}}{\sum_j e^{z_j}}$

Gradient of the Softmax:

$$[\nabla \sigma(z)]_{i,j} = \frac{1}{\sum_k e^{z_k}} \left(e^{z_i} \delta_{i-j} - \frac{e^{z_i} e^{z_j}}{\sum_k e^{z_k}} \right)$$

MSE Loss:

$$L_{MSE}(\theta) = \frac{1}{K} \sum_{k=0}^{K-1} \|x_k - f_\theta(y_k)\|^2$$

Gradient Descent and line search:

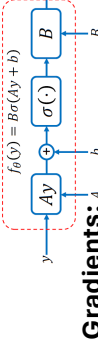
Repeat until converged {
 $d \leftarrow -\nabla L(\theta)$
 $\alpha^* \leftarrow \arg \min_t L(\theta + \alpha d^t)$
 $\theta \leftarrow \theta + \alpha^* d^t$
}

Loss gradient (MSE):

$$-\nabla_\theta L_{MSE}(\theta) = \frac{2}{K} \sum_{k=0}^{K-1} (x_k - f_\theta(y_k))^T \nabla_\theta f_\theta(y_k)$$

Einstein Notation: Delta Functions:

3 $\delta^i_j = \begin{cases} 1 & \text{if } i = j \\ 0 & \text{if } i \neq j \end{cases}$
 $G^k_j = G^k_i \delta^i_j$



Gradients: $A \leftarrow A + \alpha d^t$ with b or B

Update Directions:

Gradient step: $\rightarrow$ Replace A with b or B

For A: $d_{i,j,l_2} = -\nabla_{A,l_2} L(\theta) = \frac{2}{K} \sum_{k=0}^{K-1} \epsilon_{k,i} B^{k,l_2} [\nabla \sigma_k]^{l_2}_{j_1} y_{k,j_2}$

For b: $d_{j_1} = -\nabla_b L(\theta) = \frac{2}{K} \sum_{k=0}^{K-1} \epsilon_{k,i_1} B^{k,l_2} [\nabla \sigma_k]^{l_2}_{j_1}$

For B: $d_{i,j,l_2} = -\nabla_B L(\theta) = \frac{2}{K} \sum_{k=0}^{K-1} \epsilon_{k,i,j} \sigma_k$

Adjoint of the Loss Gradient:

4 $[\nabla L_{MSE}(\theta)]^T = -\frac{2}{K} \sum_{k=0}^{K-1} [\nabla f_\theta(y_k)]^T (x_k - f_\theta(y_k))$



Back

Propagation: $g_{m,k} = \nabla_{\theta_m} \|x_k - f_\theta(y_k)\|^2 = -2(x_k - f_\theta(y_k))^T \nabla_\theta f_\theta(y_k)$

Function: $\nabla_{\theta_m} L(\theta) = g_m = \frac{1}{K} \sum_{k=0}^{K-1} g_{m,k}$

MAE: $\|a - b\|_1 = \sum_i |a_i - b_i|$

MC Cross Entropy: $-\sum_i a_i \log b_i$

Back Prop for Loss Functions:

MSE: $\epsilon_k^L = -2(x_k - f_\theta(y_k))$

MAE: $\epsilon_k^L = -\text{sign}(x_k - f_\theta(y_k))$

CE: $\epsilon_k^L = \frac{x_k}{f_\theta(y_k)}$

Convolution "Same" Boundary:

Convolution "Valid" Boundary:

2D Convolution w/ vector input: $D=2, D_o=1$

2D Convolution w/ vector input/output: $D=2, D_o=2$

CNN Average Pooling: $d=\text{stride}$

CNN Max Pooling:

Frequentist View:

Density: $p\{X \in A\} = \sum_{i \in A} p_i$

Expected, Mean, Variance:

Regression: $\hat{\theta} = \arg \min_{\theta} \{ -\log p_\theta(X, Y) \}$

Classification (w/ 1-hot encoding): $\hat{\theta} = \arg \min_{\theta} \{ -\log p_\theta(X, y) \}$

Multinomial Distro: $\hat{\theta} = \arg \min_{\theta} \left\{ -\sum_{i=1}^M \frac{N_i}{N} \log \theta_i \right\} = \frac{N_i}{N}$

Bayesian View:

Bayes Law: $p_{Y|X}(x|y) = \frac{p_Y(y)}{\sum_{k=0}^{K-1} p_{Y|X}(x|y_k)} p_X(x)$

MAP Estimate: $\hat{x} = \arg \min_x \{ -\log p_{Y|X}(y|x) - \log p_X(x) \}$

Bias and Variance:

MAP Estimate: $\hat{x} = \arg \min_x \{ -\log p_{Y|X}(y|x) - \log p_X(x) \}$

MAP Estimate: $\hat{x} = \arg \min_x \{ -\log p_{Y|X}(y|x) - \log p_X(x) \}$

MAP Estimate: $\hat{x} = \arg \min_x \{ -\log p_{Y|X}(y|x) - \log p_X(x) \}$

MAP Estimate: $\hat{x} = \arg \min_x \{ -\log p_{Y|X}(y|x) - \log p_X(x) \}$

MAP Estimate: $\hat{x} = \arg \min_x \{ -\log p_{Y|X}(y|x) - \log p_X(x) \}$

MAP Estimate: $\hat{x} = \arg \min_x \{ -\log p_{Y|X}(y|x) - \log p_X(x) \}$

MAP Estimate: $\hat{x} = \arg \min_x \{ -\log p_{Y|X}(y|x) - \log p_X(x) \}$

MAP Estimate: $\hat{x} = \arg \min_x \{ -\log p_{Y|X}(y|x) - \log p_X(x) \}$

MAP Estimate: $\hat{x} = \arg \min_x \{ -\log p_{Y|X}(y|x) - \log p_X(x) \}$

MAP Estimate: $\hat{x} = \arg \min_x \{ -\log p_{Y|X}(y|x) - \log p_X(x) \}$

MAP Estimate: $\hat{x} = \arg \min_x \{ -\log p_{Y|X}(y|x) - \log p_X(x) \}$

MAP Estimate: $\hat{x} = \arg \min_x \{ -\log p_{Y|X}(y|x) - \log p_X(x) \}$

MAP Estimate: $\hat{x} = \arg \min_x \{ -\log p_{Y|X}(y|x) - \log p_X(x) \}$

MAP Estimate: $\hat{x} = \arg \min_x \{ -\log p_{Y|X}(y|x) - \log p_X(x) \}$

Stochastic Gradient Descent:

8 $\hat{\theta} = g + \frac{w}{\sqrt{k_0}}$

Batch Parameters:

Batch Parameters:

Batch Parameters:

Batch Parameters:

Batch Parameters:

Batch Parameters:

Batch Parameters:

Batch Parameters:

Batch Parameters:

Batch Parameters:

Batch Parameters:

Batch Parameters:

Batch Parameters:

Batch Parameters:

Batch Parameters:

Batch Parameters:

Batch Parameters:

Batch Parameters:

Batch Parameters:

Batch Parameters:

Batch Parameters:

Batch Parameters: