

HOMework 9: LEARNING PARADIGMS *

10-301 / 10-601 INTRODUCTION TO MACHINE LEARNING (SPRING 2026)

<http://www.cs.cmu.edu/~mgormley/courses/10601/>

OUT: Thursday, April 16

DUE: Tuesday, April 21

TAs: Claire, Muskan, Soham, Jerry, Markov, Neural

Summary This is the final homework assignment. This assignment covers **Ensemble Methods**, *k*-Means, **PCA**, and **Recommender Systems**.

START HERE: Instructions

- **Collaboration Policy:** Please read the collaboration policy here: <http://mlcourse.org/index.html#7-collaboration-and-academic-integrity-policies>
- **Late Submission Policy:** See the late submission policy here: <http://mlcourse.org/index.html#6-general-policies>
- **Submitting your work:** You will use Gradescope to submit answers to all questions.
 - **Written:** For written problems such as short answer, multiple choice, derivations, proofs, or plots, please use the provided template. Submissions can be handwritten onto the template, but should be labeled and clearly legible. If your writing is not legible, you will not be awarded marks. Alternatively, submissions can be written in $\text{\LaTeX}$. Each derivation/proof should be completed in the boxes provided. You are responsible for ensuring that your submission contains exactly the same number of pages and the same alignment as our PDF template. If you do not follow the template, your assignment may not be graded correctly by our AI assisted grader and there will be a **2% penalty** (e.g., if the homework is out of 100 points, 2 points will be deducted from your final score).

Instructions for Specific Problem Types

For “Select One” questions, please fill in the appropriate bubble completely:

Select One: Who taught this course?

- ☒ Matt Gormley
☐ Henry Chai
☐ Noam Chomsky

If you need to change your answer, you may cross out the previous answer and bubble in the new answer:

Select One: Who taught this course?

- ☒ Matt Gormley
☐ Henry Chai
☒ Noam Chomsky

For “Select all that apply” questions, please fill in all appropriate squares completely:

Select all that apply: Which are instructors for this course?

- ☒ Matt Gormley
☒ Pat Virtue
☐ Henry Chai
☐ I don't know

Again, if you need to change your answer, you may cross out the previous answer(s) and bubble in the new answer(s):

Select all that apply: Which are the instructors for this course?

- ☒ Matt Gormley
☒ Pat Virtue
☒ Henry Chai
☒ I don't know

For questions where you must fill in a blank, please make sure your final answer is fully included in the given space. You may cross out answers or parts of answers, but the final answer must still be within the given space.

Fill in the blank: What is the course number?

10-601

10-~~6~~301

Written Questions (38 points)

1 $\text{\LaTeX}$ Point and Template Alignment (1 points)

1.1. (1 point) **Select one:** Did you use $\text{\LaTeX}$ for the entire written portion of this homework?

☐ Yes

☐ No

1.2. (0 points) **Select one:** I have ensured that my final submission is aligned with the original template given to me in the handout file and that I haven't deleted or resized any items or made any other modifications which will result in a misaligned template. I understand that incorrectly responding yes to this question will result in a penalty equivalent to 2% of the points on this assignment.

Note: Failing to answer this question will not exempt you from the 2% misalignment penalty.

☐ Yes

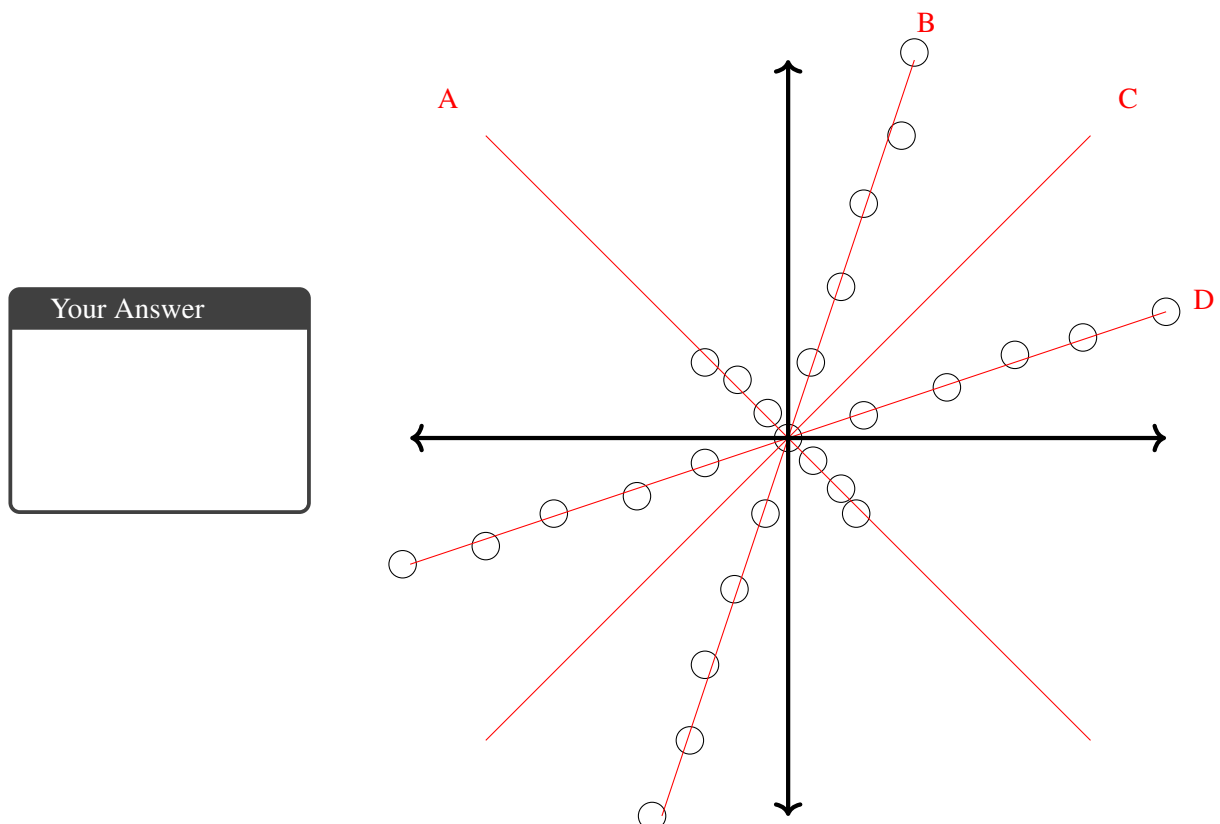
1.3. (0 points) **Select one:** Did you fill out the [Exit Poll](#) for the previous HW? Completing the exit poll will count towards your participation grade.

☐ Yes

2 PCA (7 points)

Some PCA Theory

- 2.1. (1 point) **Select one:** Assume we apply PCA to a matrix $\mathbf{X} \in \mathbb{R}^{n \times 2}$ and obtain two sets of PCA feature scores, $\mathbf{z}_1, \mathbf{z}_2 \in \mathbb{R}^n$, where $\mathbf{z}_1$ corresponds to the first principal component and $\mathbf{z}_2$ corresponds to the second principal component. Which is more common in the training data:
- ☐ a point with large feature values in $\mathbf{z}_1$ and small feature values in $\mathbf{z}_2$
 - ☐ a point with small feature values in $\mathbf{z}_1$ and large feature values in $\mathbf{z}_2$
- 2.2. (2 points) Based on the dataset shown below, list only the principal components in order from first to last. You might not list all four principle components. Be sure to consider all data points in the graph, not just those that lie on the indicated lines.



- 2.3. (1 point) **Select one:** If we wanted to perform dimensionality reduction to have just two features and train a classifier on them, we could represent our data by EITHER (a) picking any two features from the dataset, OR (b) using the first 2 principal components we obtained from PCA. Which one should we prefer and why?
- ☐ We prefer (a), because it ensures randomness in the selection and have a better chance of representing the data well.
 - ☐ We prefer (a), because PCA introduces artificial bias and does not reflect the original features.

- ☐ We prefer (b), because PCA usually preserves higher variance than two original features.
- ☐ We prefer (b), because PCA ensures that variance is evenly distributed across the features.

PCA in Practice

For this section, refer to the PCA demo linked [here](#). In this demonstration, we have performed PCA for you on a [simple four-feature dataset](#). Run the code in the notebook, then answer the questions below based on the results. The questions have also been copied to the relevant code cells in the Colab notebook for ease.

- 2.4. (1 point) **Select one:** Do you see any special relationships between any of the features? In particular, take a look at the `petal_length` feature. How would you describe its association with each of the **other features**? Select the correct statement *with appropriate justification*.
- ☐ The features are highly correlated: we observe linearly proportional relationships where increases in `petal_length` often correspond to increases in another feature
 - ☐ The features are highly correlated: we observe that the color classes can be separated with decision boundaries along the `petal_length` axis.
 - ☐ The features are uncorrelated: we observe random noise as if the features were generated from independent distributions
 - ☐ The features are uncorrelated: we observe the “default $y = x$ ” relationship between features
- 2.5. (1 point) If we wanted to find k principal components such that we preserve **at least** 95% of the variance in the data, what would be the value of k ? Hint: it is helpful here to look at the cumulative variance in the first k components, which we have calculated for you.

k

- 2.6. (1 point) See the visualization below of the first two principal components for the Fisher Iris Dataset (i.e. the projection of the Fisher Iris Dataset down to two dimensions), where the points are colored by label. From this graph, do you think PCA is an appropriate dimensionality reduction method for this dataset? Why or why not? Explain in your own words.



Figure 1: Fisher Iris Dataset in 2D

Explanation

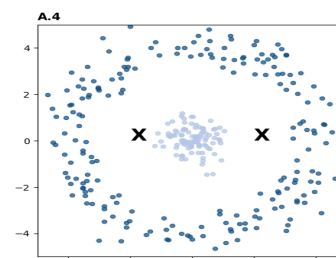
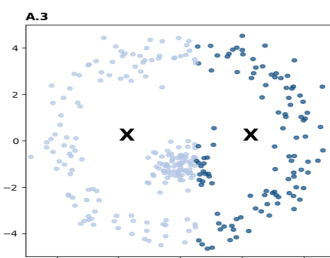
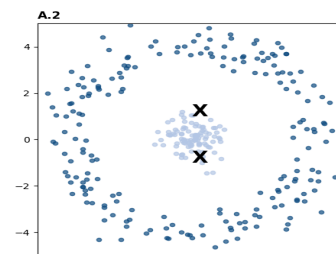
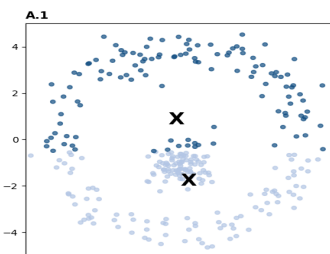
3 k -Means (10 points)

3.1. Consider the 2 datasets A and B. Each dataset is classified into k clusters, with centers marked X and cluster membership represented by different colors in the figure. For each dataset, select each image with a clustering that **could** have been generated by k -means with Euclidean distance.

3.1.a. (1 point) Dataset A

Select all that apply:

- ☐ A.1
- ☐ A.2
- ☐ A.3
- ☐ A.4
- ☐ None of the above.

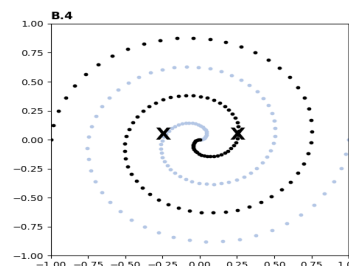
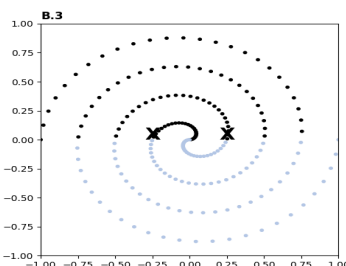
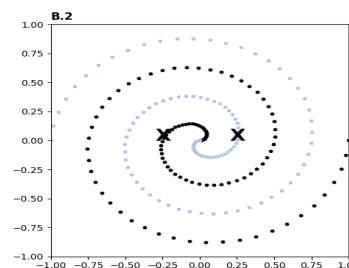
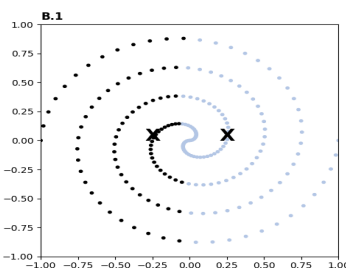


3.1.b. (1 point) Dataset B

Select all that apply:

X

- ☐ B.1
- ☐ B.2
- ☐ B.3
- ☐ B.4
- ☐ None of the above



- 3.2. Consider a dataset $\mathcal{D}$ with 5 points as shown below. Perform a k -means clustering on this dataset with $k = 2$ using the Euclidean distance as the distance function. Remember that in the k -means algorithm, one iteration consists of following two steps: first, we assign each data point to its nearest cluster center; second, we recompute each center as the average of the data points assigned to it. Initially, the 2 cluster centers are chosen randomly as $\mu_0 = (5.3, 3.5)$, $\mu_1 = (5.1, 4.2)$. Parts (a) through (c) refer only to the first iteration of k -means clustering performed on $\mathcal{D}$.

$$\mathcal{D} = \begin{bmatrix} 5.5 & 3.1 \\ 5.1 & 4.8 \\ 6.6 & 3.0 \\ 5.5 & 4.6 \\ 6.8 & 3.8 \end{bmatrix}$$

- 3.2.a. (1 point) **Select one:** Which of the following points will be the new center for cluster 0?

- ☐ (5.7, 4.1)
☐ (5.6, 4.8)
☐ (6.3, 3.3)
☐ (6.7, 3.4)

- 3.2.b. (1 point) **Select one:** Which of the following points will be the new center for cluster 1?

- ☐ (6.1, 3.8)
☐ (5.5, 4.6)
☐ (5.4, 4.7)
☐ (5.3, 4.7)

- 3.2.c. (1 point) How many points will belong to cluster 0, using the new centers?

Answer

- 3.3. Recall that in k -means clustering we attempt to find k cluster centers $\mathbf{c}_1, \dots, \mathbf{c}_k$ such that the total distance between each point and the nearest cluster center is minimized. We thus solve

$$\min \sum_{i=1}^N \min_{j \in \{1, \dots, k\}} \|\mathbf{x}^{(i)} - \mathbf{c}_j\|_2^2$$

where N is the number of data points. Instead of holding the number of clusters k fixed, your friend John tries to also minimize the objective over k , solving

$$\operatorname{argmin}_k \operatorname{argmin}_{\mathbf{c}_1, \dots, \mathbf{c}_k} \sum_{i=1}^N \min_{j \in \{1, \dots, k\}} \|\mathbf{x}^{(i)} - \mathbf{c}_j\|_2^2$$

You found this idea to be a bad one.

- 3.3.a. (1 point) What is the minimum possible value of the objective function when minimizing over k ?

Answer

- 3.3.b. (1 point) What is a value of k for which we achieve the minimum possible value of the objective function when $N = 100$?

Answer

- 3.4. (1 point) Consider the following brute-force algorithm for minimizing the k -means objective: Iterate through each possible assignment of the points to k clusters, $\mathbf{z} = [z^{(1)}, \dots, z^{(N)}]$. For each assignment $\mathbf{z} \in \{1, \dots, k\}^N$, you evaluate the following objective function:

$$J(\mathbf{z}) = \operatorname{argmin}_{\mathbf{c}_1, \dots, \mathbf{c}_k} \sum_{i=1}^N \|\mathbf{x}^{(i)} - \mathbf{c}_{z^{(i)}}\|_2^2$$

At the end, you pick the assignment $\mathbf{z}$ that had lowest $J(\mathbf{z})$.

Suppose we have N points and k clusters. For how many possible assignments $\mathbf{z}$ does the brute force algorithm have to evaluate $J(\mathbf{z})$?

Answer

- 3.5. Initializing the centers has a big impact on the performance of the k -means clustering algorithm. Usually, we randomly initialize k cluster centers. However, there are other methods, namely, furthest point initialization and k -means++ initialization.

- 3.5.a. (1 point) **Select one:** Clustering at convergence generated by furthest point initialization is sensitive to outliers. Which of the following statements is correct about this phenomenon?

- ☐ Although outliers will not be selected in the first several iterations, they will temporarily be chosen as centers during training.
- ☐ Outliers will slow convergence, but will never be centers at convergence time.
- ☐ Outliers are likely to be selected as centers in the first few iterations.

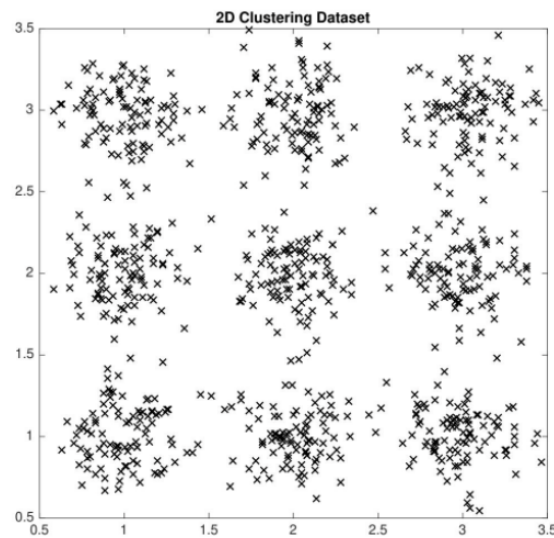


Figure 2: 2D Dataset

3.5.b. (1 point) **Select one:** Given that $k \leq 9$, using the dataset in Figure 2 above, compared to random initialization, using k -means++ initialization is _____.

- ☐ more likely to choose one sample from each cluster because centers are chosen with probability proportional to squared distance from existing centers.
- ☐ less likely to choose one sample from each cluster because the formula does not account for the number of clusters which may be found and thus won't be calibrated to correctly choose one point from each cluster.
- ☐ equally likely to choose one sample from each cluster because as the number of points grows large, k -means++ asymptotically approaches random initialization.

4 Ensemble Methods (14 points)

4.1. (1 point) **True or False:** In a random forest, it is generally better for the trees to be highly correlated, as this reduces ensemble variance.

☐ True

☐ False

4.2. (2 points) **Select all that apply:** Which of the following is true about OOB error?

☐ OOB error is calculated on a held-out dataset separate from the dataset used to generate bootstrap samples

☐ OOB error is the aggregated value of the errors of subsets of the ensemble on samples those subsets were not trained on

☐ Cross-validation error is the same as OOB error

☐ OOB error can be used instead of validation error when hyperparameter tuning a random forest model

☐ None of the above

4.3. (2 points) **Select all that apply:** Which of the following are hyperparameters that can be tuned in a random forest?

☐ Number of trees trained

☐ Number of points used to train each decision tree

☐ Size of feature subsets used to train each decision tree

☐ Which features are used for splits in each decision tree

☐ None of the above

4.4. In this question, we will consider the behavior of an error bound for random forests in the case of **binary classification**. Given a random forest of B trees $\{h_i(x)\}_{i=1}^B$ and a sample (x, y) drawn from some data distribution $\mathcal{D}$, define the classification margin as:

$$m(x, y) = \frac{1}{B} \left(\sum_{i=1}^B \mathbb{I}[h_i(x) = y] - \sum_{i=1}^B \mathbb{I}[h_i(x) \neq y] \right)$$

In words, the margin $m(x, y)$ is the difference between the average vote for the correct label and the average vote for the incorrect label.

4.4.a. (1 point) **Fill in the blank:** For any example (x, y) , the example is classified incorrectly if and only if $m(x, y) \leq \underline{\hspace{1cm}}$. Assume majority vote ties are classified incorrectly.

Answer

- 4.4.b. (1 point) **Select one:** Using [Chebyshev's inequality](#), it is possible to bound the generalization error of a random forest as $P(m(x, y) < c) \leq \frac{\text{Var}(m(x, y))}{s^2}$, where c is your answer to part (a) and s is the *strength* of the ensemble

$$s = \mathbb{E}_{(x, y) \sim \mathcal{D}}[m(x, y)].$$

Through some additional manipulation, it is further possible to show that $\text{Var}(m(x, y)) \leq \bar{\rho}(1 - s^2)$, where $\bar{\rho}$ is the mean correlation between trees in the ensemble. Substitute this into the given bound. Which of the following describes how the error bound is affected by s and $\bar{\rho}$?

- ☐ The error bound gets smaller as $\bar{\rho}$ increases and $|s|$ increases.
- ☐ The error bound gets smaller as $\bar{\rho}$ increases and $|s|$ decreases.
- ☐ The error bound gets smaller as $\bar{\rho}$ decreases and $|s|$ increases.
- ☐ The error bound gets smaller as $\bar{\rho}$ decreases and $|s|$ decreases.

- 4.5. (1 point) **True or False:** Consider some training point $(x^{(i)}, y^{(i)})$ to the AdaBoost algorithm. If for all t , the weak learner h_t learned during training at time t correctly classifies $h_t(x^{(i)}) = y^{(i)}$, there will eventually be a finite time t such that the weight assigned to $x^{(i)}$ in the training distribution $\mathcal{D}_t$ reaches exactly 0.

☐ True

☐ False

- 4.6. Assume we use a deterministic training procedure for weak learners. Suppose for some iteration t' of AdaBoost we find that the weak classifier learned by the algorithm at time t' has error $\epsilon_{t'} = 0.5$ of the weak learner $h_{t'}$ on the training distribution weighted by $\mathcal{D}_{t'}$.

- 4.6.a. (1 point) What weight $\alpha_{t'}$ will AdaBoost assign to the classifier $h_{t'}$ from above?

$\alpha_{t'}$

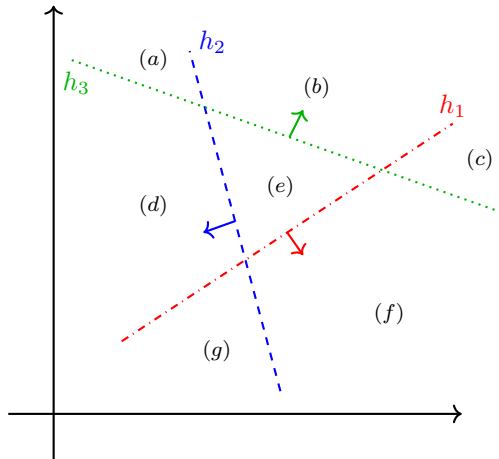
- 4.6.b. (1 point) **Select all that apply:** Which of the following scenarios are possible?

- ☐ $\mathcal{D}_{t'+1}(i) > \mathcal{D}_{t'}(i)$ and $h_{t'}$ correctly classifies $x^{(i)}$
- ☐ $\mathcal{D}_{t'+1}(i) > \mathcal{D}_{t'}(i)$ and $h_{t'}$ incorrectly classifies $x^{(i)}$
- ☐ $\mathcal{D}_{t'+1}(i) < \mathcal{D}_{t'}(i)$ and $h_{t'}$ correctly classifies $x^{(i)}$
- ☐ $\mathcal{D}_{t'+1}(i) < \mathcal{D}_{t'}(i)$ and $h_{t'}$ incorrectly classifies $x^{(i)}$
- ☐ None of the above.

- 4.6.c. (1 point) **Select all that apply:** Which of the following are true about the next iteration of the AdaBoost algorithm?

- ☐ The errors $\epsilon_{t'+1}$ and $\epsilon_{t'}$ are equivalent
- ☐ The weak learners $h_{t'+1}$ and $h_{t'}$ will be equivalent (i.e., they will have the same output for every input)
- ☐ None of the above

- 4.7. The figure below shows three binary classifiers h_1, h_2, h_3 in a 2D space. We intend to use AdaBoost to combine them into a strong classifier. The data set consists of 7 points (a) - (g), and the weight initializations for these points are uniform (i.e. weights are equal and add up to 1). The true label of (g) is negative. Use this information to answer the following:



- 4.7.a. (1 point) **Select all that apply:** Using $\alpha_1 = \alpha_2 = \alpha_3 = 1$ as weights of each classifier, which data points among (a), (b), (c), (d), (e), (f), and (g) would be classified as +1 by weighted majority vote?

- ☐ (a)
☐ (b)
☐ (c)
☐ (d)
☐ (e)
☐ (f)
☐ (g)
☐ None of the above

- 4.7.b. (1 point) **Fill in the blank:** Use $\alpha_1 = 0, \alpha_2 = 1, \alpha_3 = 0$ as weights of each classifier, and assume that h_1, h_2 , and h_3 update the sample distribution in this order. Fill in the weight of the datapoint (g) after each weak classifier is applied. No weight normalization is needed.

w_1	w_2	w_3

4.7.c. (1 point) **True or False:** Neural argues that no classifier can ever receive weight $\alpha = 0$, because the dataset contains 7 (odd number) datapoints, the weighted error ϵ_t of any weak classifier can never be exactly 0.5. **Select one and briefly justify your answer.**

☐ True

☐ False

5 Recommender Systems (6 points)

5.1. (2 points) **Select all that apply:** In which of the following situations will a collaborative filtering system be a more appropriate learning algorithm than a linear or logistic regression model?

- ☐ You manage an online bookstore, and you have book ratings and sales data from many users. For each user, you want to recommend other books she will enjoy, based on her own ratings and the ratings of other users.
- ☐ You manage an online bookstore, and you have book ratings and sales data from many users. You want to learn to predict the expected sales volume (number of books sold) as a function of the average rating of a book.
- ☐ You run an online news aggregator, and for every user, you know some subset of articles that the user likes and some different subset that the user dislikes. You want to use this to find other articles that a given user likes.
- ☐ You've written a piece of software that downloads news articles from many news websites. In your system, you also keep track of which articles you personally like and which ones you dislike, and the system also stores away features of these articles (e.g., word counts, name of author). Using this information, you want to build a system to try to find additional new articles that you personally will like.
- ☐ None of the above

5.2. Neural the Narwhal decides to set up a friend-recommendation system for all the students in 10-301/601, of which there are $N = 10,301,601$. Ideally, Neural would store the full $N \times N$ matrix M , where M_{ij} is 1 if student i and j are friends, 1 if $i = j$, 0 if student i and j are nemeses, or null if students i and j have never met. Assume that these are the only possible relationships between 2 people and that all relationships are symmetric (so it cannot be the case that student i thinks student j is their friend while j thinks i is their nemesis). Unfortunately, storing M in its entirety would take over 10 TB of storage, which Neural cannot afford on a TA salary. Neural instead uses the following procedure to approximate M as UU^T for some $N \times d$ matrix U , where each row of U corresponds to a student.

```

1: Given learning rate  $\eta$ , ground truth relationships  $M$ 
2: Randomly initialize user embedding matrix  $U \in \mathbb{R}^{N \times d}$ 
3: while not converged do
4:   Sample  $i \sim \text{Uniform}(1, N)$ ,  $j \sim \text{Uniform}(1, N)$ 
5:    $\hat{M}_{ij} \leftarrow \sigma(\vec{u}_i^T \vec{u}_j)$  //  $\sigma$  is sigmoid function
6:    $\mathcal{L}(\vec{u}_i, \vec{u}_j) \leftarrow -M_{ij} \log(\hat{M}_{ij}) - (1 - M_{ij}) \log(1 - \hat{M}_{ij})$  // Compute logistic loss
7:    $\vec{g}_i \leftarrow \nabla_{\vec{u}_i} \mathcal{L}(\vec{u}_i, \vec{u}_j)$  // Compute and perform gradient updates
8:    $\vec{g}_j \leftarrow \nabla_{\vec{u}_j} \mathcal{L}(\vec{u}_i, \vec{u}_j)$ 
9:    $\vec{u}_i \leftarrow \vec{u}_i - \eta \cdot \vec{g}_i$ 
10:   $\vec{u}_j \leftarrow \vec{u}_j - \eta \cdot \vec{g}_j$ 

```

- 5.2.a. (1 point) Explain why line 4 of the algorithm is incorrect. You do not need to provide a line of code unless you feel that it would help you in your explanation.

Explanation

- 5.2.b. (1 point) **Select one:** Based on the loss function given on line 6, derive an expression for $\nabla_{\vec{u}_j} \mathcal{L}(\vec{u}_i, \vec{u}_j)$ in terms of $\vec{u}_i$, $\vec{u}_j$, and M_{ij} . Note that σ denotes the sigmoid function and $\log$ is natural log.

- ☐ $\vec{u}_i M_{ij} \sigma(\vec{u}_i^T \vec{u}_j)$
☐ $\vec{u}_i (M_{ij} + \sigma(\vec{u}_i^T \vec{u}_j))$
☐ $\vec{u}_i (M_{ij} - \sigma(\vec{u}_i^T \vec{u}_j))$
☐ $\vec{u}_i (-M_{ij} + \sigma(\vec{u}_i^T \vec{u}_j))$

- 5.2.c. (2 points) **Select all that apply:** Why is it appropriate here to factorize M as the product of a single matrix by itself UU^T rather than as the product of two distinct matrices VW^T ?

- ☐ UU^T enforces that our approximation is symmetric, while VW^T is not necessarily symmetric.
☐ SGD is guaranteed to converge faster and to a better optimum for U because we have fewer parameters to learn.
☐ We can parallelize training for U , whereas we could not parallelize training for V and W .
☐ We wish to model relationships between objects of a single type, not between objects of two different types.
☐ None of the above.

6 Collaboration Questions

After you have completed all other components of this assignment, report your answers to these questions regarding the collaboration policy. Details of the policy can be found [here](#).

1. Did you receive any help whatsoever from anyone in solving this assignment? If so, include full details.
2. Did you give any help whatsoever to anyone in solving this assignment? If so, include full details.
3. Did you find or come across code that implements any part of this assignment? If so, include full details.

Your Answer