

ECE 661 - Assignment 10

Ali Almualllem, aalmuall@purdue.edu

December 2024

1 Theory Questions

1.1 Overfitting

Overfitting occurs when a model (often a machine learning model) learns the training data very well but fail to generalize to the testing data. This can be due to the scarcity of the training data, or to the capacity of the model. In other words, overfitting can be seen as memorizing the training data but failing to generalize to unseen data. If a model is too complex or too simple for the underlying data, it may fail to generalize. [2]

In machine learning, however, there is some recent research that argue that overparametrization may lead to overcome the overfitting phenomenon in what is known as "*Double Descent*" [3] which is explained below.

1.1.1 Do models overfit in the "modern" overparametrized models?

There is recent research that suggests that as the model complexity gets bigger, the model will initially overfit to the training data (the under-parametrized region in Fig. 1), but as model gets "over-parametrized", it will reach a point where it will actually get better in generalization and avoids overfitting. This hypothesis is known in the literature as "over-parametrization", which as the name suggests, make large models (larger than usually necessary for a given problem) and performs well on the training and testing data. This challenges the usual conventional thinking that very large models will always yield to overfitting. [4]

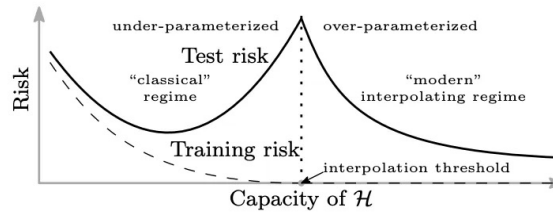


Figure 1: Overparametrization. To the left of the interpolation threshold is the classical thinking of model complexity and overfitting, to the right, the modern hypothesis that suggests that sufficiently large models do not suffer from overfitting. Figure from: [4]

1.2 Reparameterization Trick

The reparameterization trick is a technique used in VAEs and variational inference in general to overcome a backpropagation problem through a random sampling step. The issue arises in VAEs when the latent vector z samples from a distribution, i.e.: $z \sim \mathcal{N}(\mu, \sigma)$, that step is non-differentiable, and therefore the loss cannot be calculated and backpropagated.

The trick therefore solves this issue by defining a latent vector z according to the following equation:

$$z = \mu + \sigma * \epsilon \quad (1)$$

Where $\epsilon \sim \mathcal{N}(0, 1)$. Note however, that drawing a sample from, $z \sim \mathcal{N}(\mu, \sigma)$ or $\epsilon \sim \mathcal{N}(0, 1)$ is non-differentiable. However, employing the trick in the equation above makes z differentiable despite having a non-learnable parameter ϵ .

The role of the new parameter ϵ is said to give a differentiable degree of freedom to the decoder (generator) when it generates a new image from the latent vector z . Others describe the new parameter ϵ as adding noise to the sampled z . [2], [1]

2 Face Recognition using PCA and LDA

The dataset provided has 30 persons faces. Each person has 21 images of their face from slightly different directions. The first four faces can be seen in Fig. 2 *Note*: we resized the images to 32×32 for faster computations.

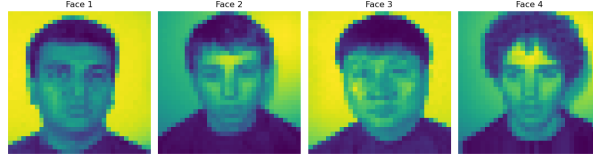


Figure 2: First four persons faces

2.1 Principal Component Analysis: PCA

PCA is a statistical method commonly used for dimensionality reduction tasks. PCA aims to find a lower-dimensional representation of an underlying high-dimensional data. It tries to find directions (principal components) that maximizes the variation in the data so that it can be easily dismantled or identified.

PCA achieves that using the following steps:

- Given N vectorized images $\vec{x}_i$ $i = 1, 2, \dots, N$ where each image is represented as a column in the matrix.
- Compute the mean image vector $\vec{m} = \frac{1}{N} \sum_{i=1}^N \vec{x}_i$
- Subtract the mean from the images: $X = [\vec{x}_1 - \vec{m} \quad \vec{x}_2 - \vec{m} \quad \dots \quad \vec{x}_N - \vec{m}]$
- To avoid variations due to differences in illumination, normalize the each image in $\vec{x}_i$
- Calculate the covariance $C = \frac{1}{N} \sum_{i=1}^N \{(\vec{x}_i - \vec{m})(\vec{x}_i - \vec{m})^T\}$
- The above calculation is expensive since it's XX^T so if X is of size 4096×1000 , the resulting C is of size $4096 \times 4096 = +16\text{million}$ which is huge. So, we employ a computational trick described in [2]. For the interest of conciseness, we refer the reader to the reference for the details.
- We calculate the eigenvectors $W_K = [\vec{w}_1 \quad \vec{w}_2 \quad \dots \quad \vec{w}_K]$ of our covariance matrix.

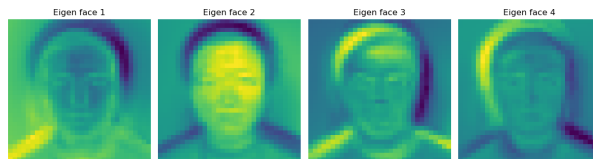


Figure 3: The first 4 eigen faces

- Each vector $\vec{w}_i$ acts as a feature for classification.
- The data can be then projected onto the new space: $\vec{y} = W_K^T(\vec{x} - \vec{m})$

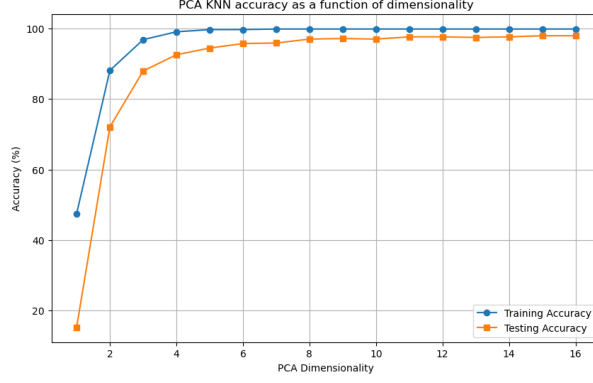


Figure 4: PCA accuracy as a function of dimensionality p .

2.2 UMAP for PCA

The UMAP for PCA shows a good overall separation between classes that increases with the dimensionality (value of P). However, there are still many classes that are very close to each other or cannot be disentangled, at least visually by inspecting the plot. We will see that this is slightly different with LDA.

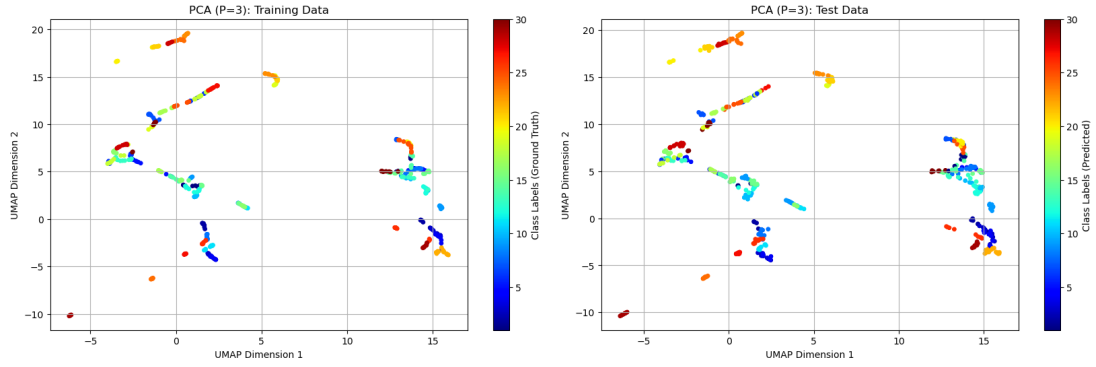


Figure 5: UMAP for PCA with $P=3$. Left: training. Right: Testing

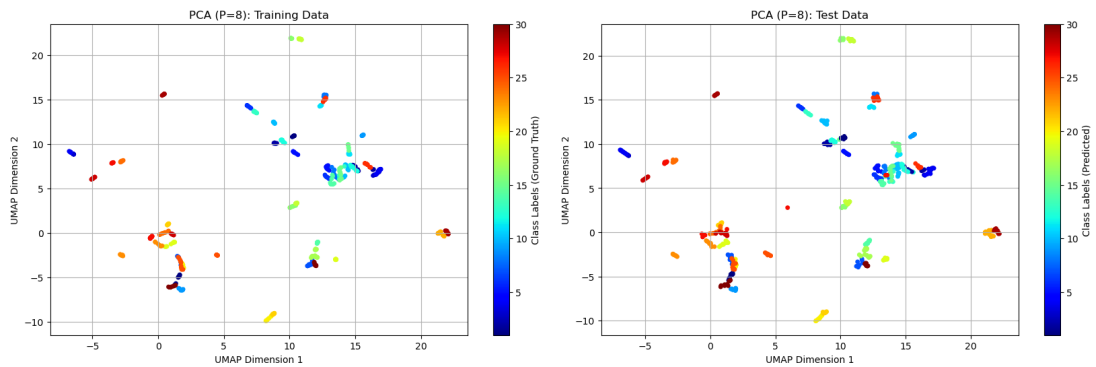


Figure 6: UMAP for PCA with $P=8$. Left: training. Right: Testing

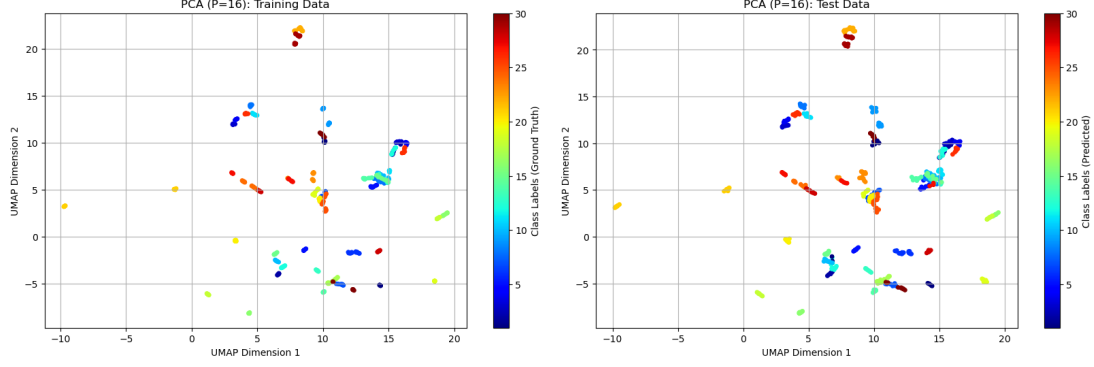


Figure 7: UMAP for PCA with P=16. Left: training. Right: Testing

2.3 Linear Discriminant Analysis: LDA

Just like PCA, LDA is another dimensionality reduction method, but unlike PCA, LDA tries to find the direction of a vector space that maximizes the separability between classes, i.e.: it tries to maximize the discrimination between classes. Instead of maximizing the overall variance, LDA aims to maximize the between-class scatter and minimize the within-class scatter.

The steps are generally as follow:

- Let $\vec{m}$ be the global mean over all images.
- Define the between-classes scatter as: $S_B = \frac{1}{|C|} \sum_{i=1}^{|C|} (\vec{m}_i - \vec{m})(\vec{m}_i - \vec{m})^T$
- Define the within-class scatter as: $S_W = \frac{1}{|C|} \sum_{i=1}^{|C|} \frac{1}{|C_i|} \sum_{k=1}^{|C_i|} (\vec{x}_k^i - \vec{m}_i)(\vec{x}_k^i - \vec{m}_i)^T$
- If $\vec{w}$ is the underlying vector space, then $\vec{w}^T S_B \vec{w}$ is the projection of the between-classes scatter on $\vec{w}$ and $\vec{w}^T S_W \vec{w}$ is the projection of the within-class scatter on $\vec{w}$.
- Use the Fisher Discriminant Function to maximize the between-class and minimize the within-class scatter as follows: $J(\vec{w}) = \frac{\vec{w}^T S_B \vec{w}}{\vec{w}^T S_W \vec{w}}$
- The Fisher function must satisfy $S_B \vec{w} = \lambda S_W \vec{w}$ for some constant λ .
- If the within-class scatter is assumed non-singular, the problem can be transformed into an eigen-decomposition problem as follows: $S_W^{-1} S_B \vec{w} = \lambda \vec{w}$
- Additional steps might be needed if it is singular. For the sake of conciseness, we refer the reader for [2] for more details.

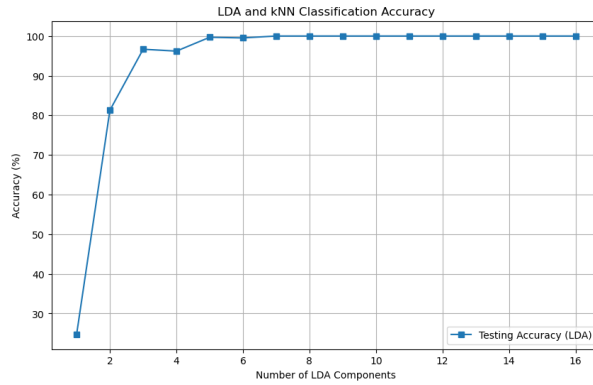


Figure 8: LDA result as a function of dimensionality.

2.4 Overall observation of PCA and LDA

In general, both methods performed similarly with LDA performing a notch better. It could be that PCA needs more dimensionality. It could also be the nature of PCA objective, as it does not work on the variance between and within classes, but rather the overall variance, which makes LDA better in some situations. However, the obtained results still achieved pretty good accuracy with both methods.

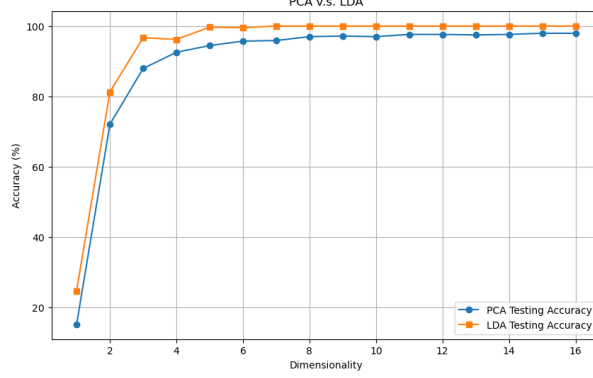


Figure 9: A comparison between LDA and PCA

2.5 UMAP for LDA

The LDA UMAP plots are much more discriminatory than PCA. The separation is "cleaner", i.e.: there is only slight or no overlap between classes. Looking at $P=3$ graphs, we can see only few classes that overlap, and as the dimensionality gets larger $P > 3$, those overlaps disappear. This is supported by the accuracy plots provided earlier, which shows a very high accuracy on the testing dataset as P increases.

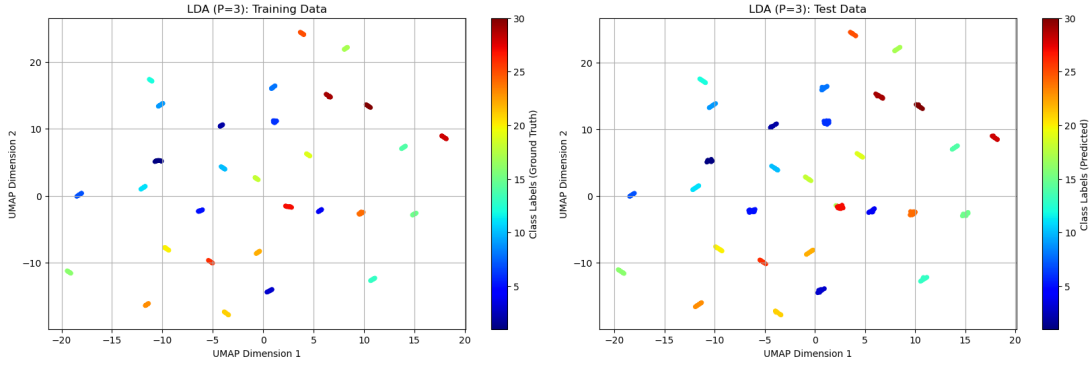


Figure 10: UMAP for LDA with $P=3$. Left: training. Right: Testing

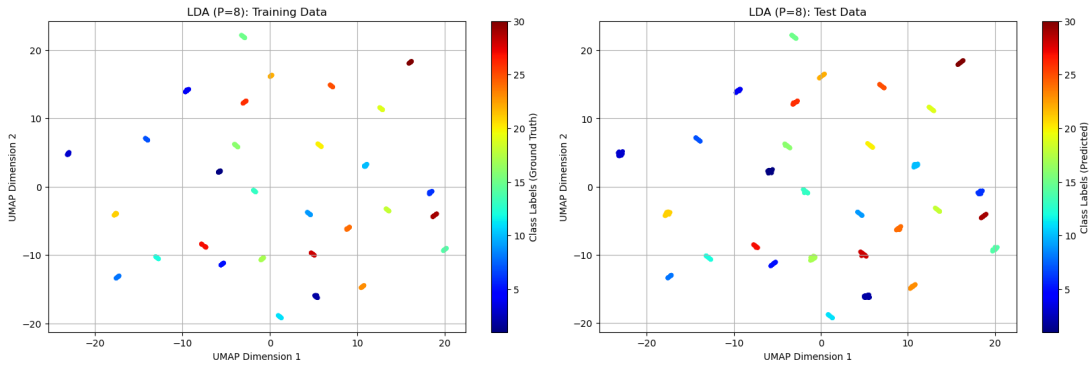


Figure 11: UMAP for LDA with $P=8$. Left: training. Right: Testing

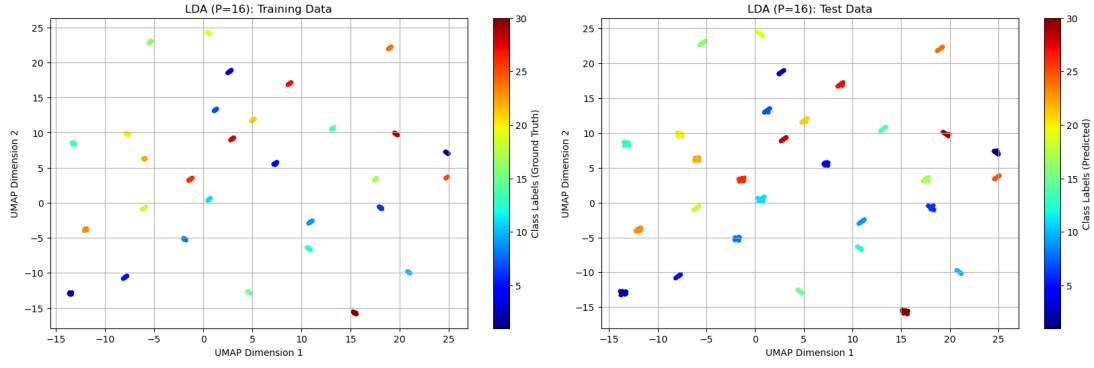


Figure 12: UMAP for LDA with P=16. Left: training. Right: Testing

3 Face Recognition using an Autoencoder

We utilized the provided code and obtain the testing accuracies on different network that were pre-trained on different dimensionalities, namely: $P = 3, 8, 16$. As expected, the accuracy increased as the dimensionality increases, reaching 100% with $P = 16$ as seen in Fig. 13. We will compare the autoencoder results against PCA and LDA in a later subsection, and detail the accuracies in Table. 1

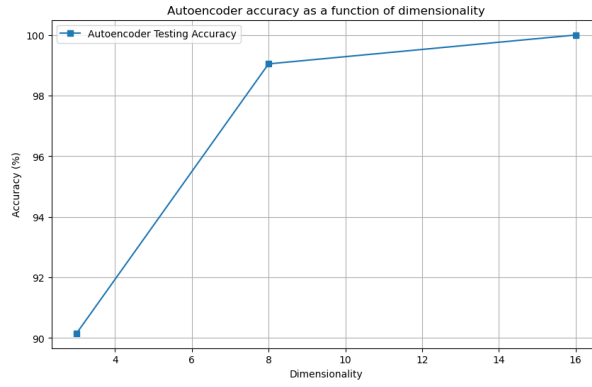


Figure 13: The pre-trained autoencoder result as a function of dimensionality

3.1 UMAP for Autoencoder

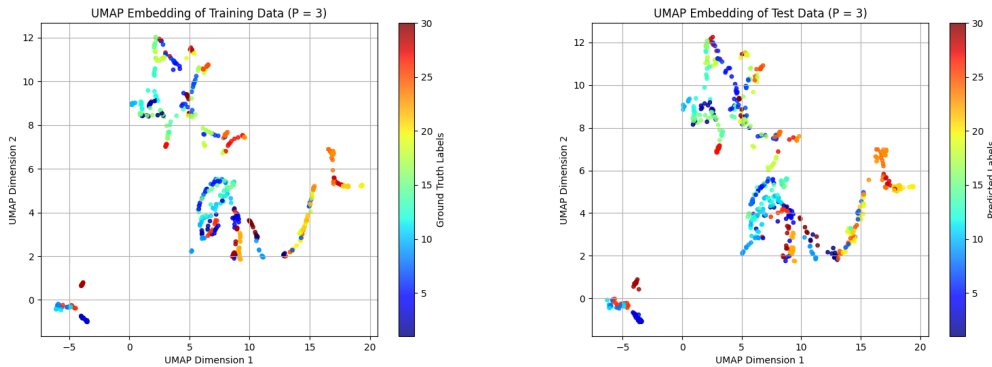


Figure 14: UMAP for Autoencoder with P=3. Left: training. Right: Testing

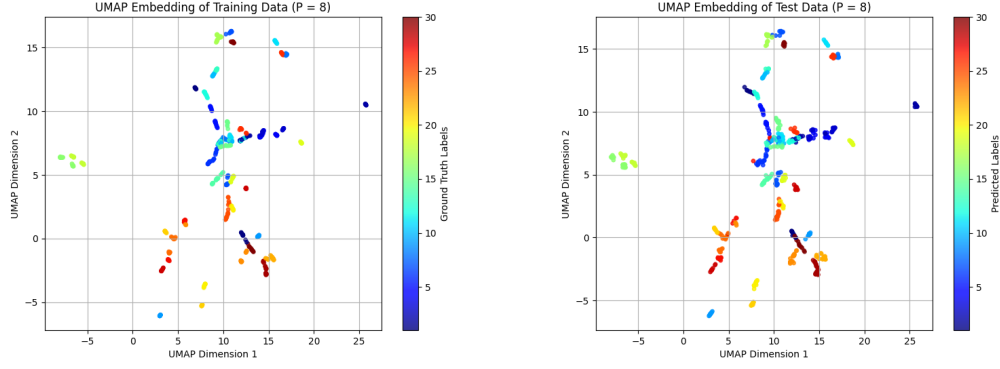


Figure 15: UMAP for Autoencoder with P=8. Left: training. Right: Testing

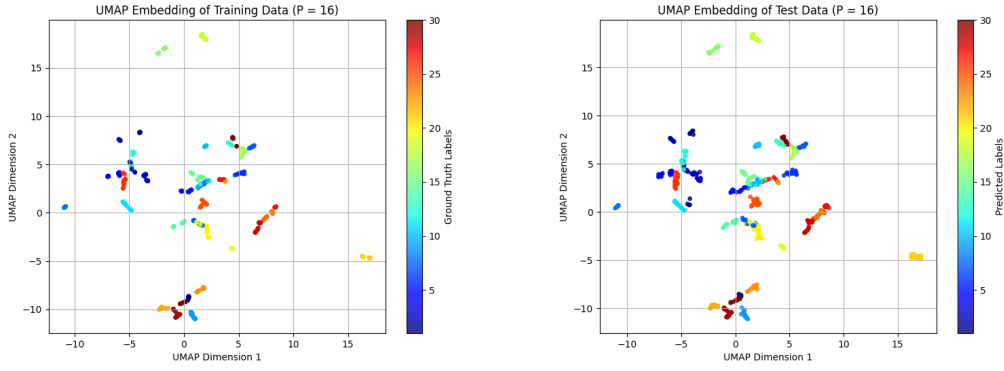


Figure 16: UMAP for Autoencoder with P=16. Left: training. Right: Testing

3.2 Overall comparison: LDA v.s. Autoencoder v.s. PCA

We compare the results of the provided pre-trained autoencoder against our implemented LDA and PCA on the testing data for P values of 3, 8, and 16. As we can see from Fig. 17, the LDA performed the best, especially when the dimensionality was low (P=3), but it has almost the same accuracy as the autoencoder when the dimensionality is sufficiently large (P=16). The PCA in our experiment performed well but slightly lower than the other two methods, which suggests that it could be an issue with the PCA not having the same discrimination power as the LDA and autoencoder, as it does not seek to "discriminate" in its objective function like the LDA does.

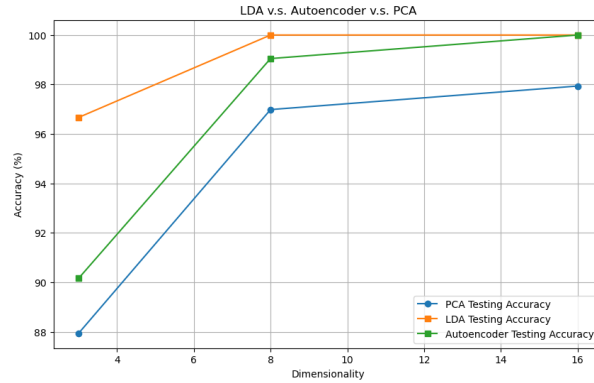


Figure 17: The testing accuracy result for P= 3, 8, 16 for LDA, Autoencoder, and PCA.

	P = 3	P = 8	P = 16
Autoencoder	90.159%	99.048%	100%
PCA	87.937%	96.984%	97.937%
LDA	96.667%	100%	100%

Table 1: Accuracy results with different dimensionality P . Highest values are in **bold**

4 Object Detection using Cascaded AdaBoost Classifiers

AdaBoost, short for Adaptive Boosting, is a learning algorithm that combines multiple "weak" classifiers into a single "strong" classifier. A weak classifier here is one that performs slightly better than random chance, while a strong classifier achieves high accuracy rates and low false positive and false negative rates.

The idea behind the cascaded classifiers relies on the exponential rule, where if we have a true positive rate of $x_{TP} = 0.99$ and a false positive rate of $x_{FP} = 0.3$ as a start, having 10 classifiers would yield: $x_{TP}^{10} = 0.99^{10} \approx 0.90$ and $x_{FP}^{10} = 0.30^{10} \approx 0.000006$.

The general outline for how AdaBoost works is as follows. For detailed steps and intuition, we refer the reader to the comprehensive guide by Prof. Kak [2]:

- Compute image features. There are multiple ways to do that. I choose to implement HAAR like features with window of sizes: 2, 4, 8, 16.
- We initialize a weight matrix uniformly w . They are then updated for each sample to determine whether it should be considered in the next iterations.
- Get a weak classifier that specifically targets the training samples that were misclassified by the previous weak classifier (if any). We'll denote it as h_t
- Apply h_t to all training data. The classifier can be seen as a mapping $h_t : X \mapsto \{-1, 1\}$
- If D_t is the probability distribution, we use the following equation to measure the weighted misclassification rate $Prob\{h_t(x_i) \neq y_i \text{ where } y_i \text{ is the label}$

$$\epsilon_t = \frac{1}{2} \sum_{i=1}^m D_t(x_i) \cdot |h_t(x_i) - y_i| \quad (2)$$

- We calculate the trust factor for the current classifier

$$\alpha_t = \frac{1}{2} \ln \left(\frac{1 - \epsilon_t}{\epsilon_t} \right) \quad (3)$$

- Update the probability distribution for the next iteration

$$D_{t+1}(x_i) = \frac{D_t(x_i) e^{-\alpha_t y_i h_t(x_i)}}{Z_t} \quad (4)$$

Where Z_t is a normalization factor. So it is set such that

$$D_{t+1}(x_i) = \frac{D_t(x_i) e^{-\alpha_t y_i h_t(x_i)}}{Z_t} \quad (5)$$

- Iterate back to the 3rd step by getting another weak classifier.

4.1 AdaBoost Classification performance

The data for the AdaBoost classification tasks consist of images of cars and other images that do not contain cars as seen in Fig. 18. The dataset images are pretty small (40×20), which could have impacted the accuracy of our feature extraction routine. We expand on that in the following subsections.

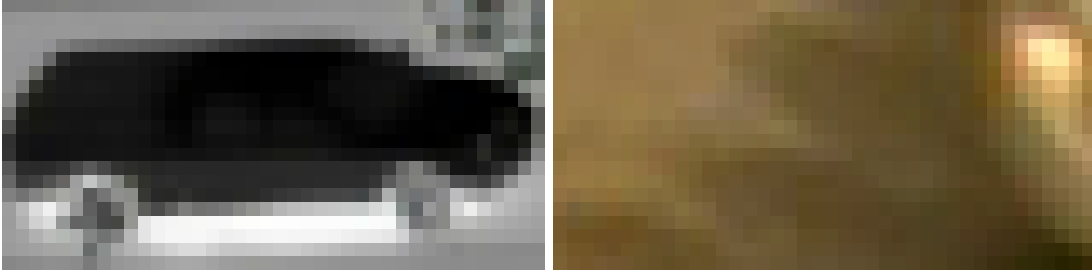


Figure 18: **Left: Positive** sample (car present). **Right: Negative** sample (car not present)

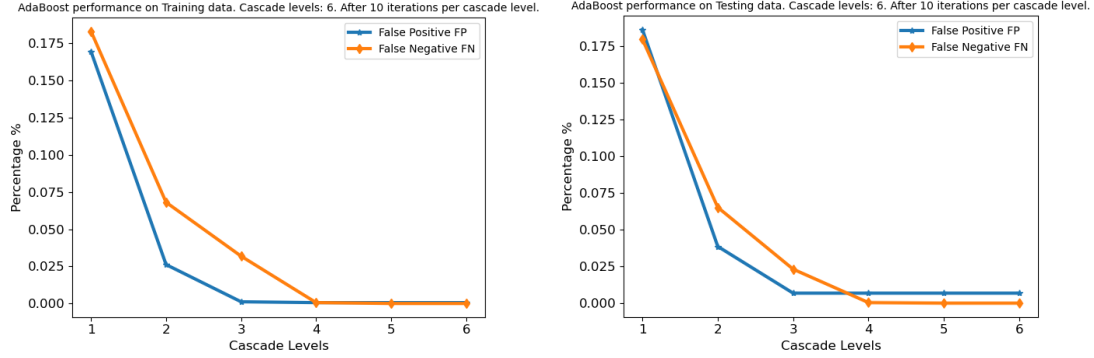


Figure 19: AdaBoost performance. **Left: training** dataset. **Right: Testing** dataset

4.2 Observation on AdaBoost

The plots above in Fig. 19 shows clearly that the classification performance improves substantially as the cascade levels increases. However, it is noticable that the performance plateaus around level four in my implementation. This could be due to a number of reason among which is the dataset which consist of very small low quality images which might have been hard to extract the features from. Another possibility is my feature extraction method which consisted of a HAAR-like windowing technique. The choose of windows sizes might influence the classification performance.

Non-theless, the lowest false positive and false negatives rates achieved are still remarkable as seen in Table.2.

	Training	Testing
FP	0.0005688	0.0068182
FN	0.000000	0.000000

Table 2: Lowest false positive (FP) and false negative (FN) rates achieved on both training and testing datasets.

References

- [1] Gregory Gundersen. The Reparameterization Trick — gregorygundersen.com. <https://gregorygundersen.com/blog/2018/04/29/reparameterization/>, 2018. [Accessed 30-11-2024].
- [2] Avinash Kak. Ece 661: Computer vision. <https://engineering.purdue.edu/kak/computervision/>, 2024.
- [3] Preetum Nakkiran, Gal Kaplun, Yamini Bansal, Tristan Yang, Boaz Barak, and Ilya Sutskever. Deep double descent. *OpenAI Blog*, 2019.
- [4] Jaideep Ray. What are Over-parameterized models? - medium.com. <https://medium.com/better-ml/what-are-over-parameterized-models-40c53c304367>, 2022. Accessed: 30-11-2024.

5 Code

Listing 1: Assignment 10 Code

```
1 import numpy as np
2 import matplotlib.pyplot as plt
3 import cv2
4 import umap
5 import os
6 import pickle
7
8
9
10 def reading_images(images_directory, num_image_per_class, num_classes,
11                   resize = True, resize_to = 32):
12     """Read images from a directory
13     Images are of this format: xx_yy.png, where xx is the class number,
14     and yy is the index of the image per class
15     Input:
16     - images_directory: the folder where images reside
17     - num_image_per_class: number of images per class
18     - num_classes: the number of classes
19
20     returns:
21     - labels: numpy array of labels
22     - images: numpy array of the flatten images. Each image is a column
23       vector
24     """
25     labels = []#np.linspace(1, num_classes, num_classes, dtype= int,
26                       endpoint= True)
27
28     #Classes starts at 1
29     for class_num in range(1, num_classes+1):
30
31         #Images start at 1
32         for image_num in range(1, num_image_per_class+1):
33             current_class_num = str(class_num) if class_num>=10 else
34                 "0"+str(class_num)
35             current_image_num = str(image_num) if image_num>=10 else
36                 "0"+str(image_num)
37             current_path = images_directory + "/" + current_class_num + "_" +
38                 current_image_num + ".png"
39
40             #Opening image
41             current_image = cv2.imread(current_path)
42
43             # Resize the image
44             if (resize):
45                 current_image = cv2.resize(current_image, (resize_to,
46                                                             resize_to),
47                                             interpolation = cv2.
48                                                 INTER_AREA)
49
50             #Convert to grayscale
51             current_image = cv2.cvtColor(current_image, cv2.
52                 COLOR_BGR2GRAY)
53
54             current_image = np.array(current_image).flatten()
```

```

45         current_image = np.expand_dims(current_image, 1)
46
47         #If it's the first class and image, create the numpy array
48         if (class_num==1 and image_num==1):
49             all_images = current_image
50         else:
51             all_images = np.concatenate((all_images, current_image)
52                                         , axis = 1)
53             labels.append(class_num)
54     return all_images, np.array(labels)
55
56
57 # Define the kNN function
58 def knn_classify(trainX, trainY, testX, k):
59     predictions = []
60     for test_point in testX:
61         distances = np.linalg.norm(trainX - test_point, axis=1) #
62         # Compute distances
63         neighbors = np.argsort(distances)[:k] # Find k nearest
64         # neighbors
65         neighbor_labels = trainY[neighbors] # Get their labels
66         prediction = np.bincount(neighbor_labels).argmax() # Majority
67         # vote
68         predictions.append(prediction)
69     return np.array(predictions)
70
71 #Compute PCA
72 def manual_pca(X, num_components):
73     mean_vector = np.mean(X, axis=0) # Compute mean
74     X_centered = X - mean_vector # Center the data
75     covariance_matrix = np.dot(X_centered.T, X_centered) / (X_centered.
76     shape[0] - 1) # Covariance matrix
77     eigenvalues, eigenvectors = np.linalg.eig(covariance_matrix) #
78     # Eigen decomposition
79     idx = np.argsort(-eigenvalues) # Sort eigenvalues in descending
80     # order
81     eigenvectors = eigenvectors[:, idx[:num_components]] # Select top
82     # eigenvectors
83     X_reduced = np.dot(X_centered, eigenvectors) # Project data to the
84     # new basis
85     return X_reduced, mean_vector, eigenvectors
86
87 # Compute accuracy
88 def compute_accuracy(true_labels, predicted_labels):
89     return np.mean(true_labels == predicted_labels) * 100
90
91 #LDA
92 # Define the LDA function
93 def manual_lda(trainX, trainY, num_components):
94     classes = np.unique(trainY)
95     overall_mean = np.mean(trainX, axis=0)
96
97     # Initialize within-class scatter matrix and between-class scatter
98     # matrix
99     Sw = np.zeros((trainX.shape[1], trainX.shape[1]))
100    Sb = np.zeros((trainX.shape[1], trainX.shape[1]))

```

```

93
94 for c in classes:
95     class_samples = trainX[trainY == c]
96     class_mean = np.mean(class_samples, axis=0)
97     Sw += np.dot((class_samples - class_mean).T, (class_samples -
98         class_mean))
99     n_class_samples = class_samples.shape[0]
100     mean_diff = (class_mean - overall_mean).reshape(-1, 1)
101     Sb += n_class_samples * np.dot(mean_diff, mean_diff.T)
102
103 eigvals, eigvecs = np.linalg.eig(np.linalg.pinv(Sw).dot(Sb))
104 idx = np.argsort(-eigvals.real) # Sort eigenvalues in descending
105 order
106 eigvecs = eigvecs[:, idx[:num_components]].real # Select top
107 eigenvectors
108
109 # Project data onto the LDA components
110 trainX_lda = np.dot(trainX, eigvecs)
111 return trainX_lda, eigvecs
112
113 #Plot UMAP
114 def plot_umap_embeddings(trainX_proj, trainY, testX_proj, testY_pred,
115     title_train, title_test):
116     """
117     Function to plot UMAP embeddings for training and test data.
118     """
119     umap_reducer = umap.UMAP(n_components=2, random_state=42)
120
121     # Fit UMAP on training data
122     trainX_umap = umap_reducer.fit_transform(trainX_proj.real)
123     # Transform test data using the same UMAP
124     testX_umap = umap_reducer.transform(testX_proj.real)
125
126     # Plot training data
127     plt.figure(figsize=(10, 6))
128     scatter = plt.scatter(trainX_umap[:, 0], trainX_umap[:, 1], c=
129         trainY, cmap='jet', s=15)
130     plt.title(title_train)
131     plt.xlabel("UMAP Dimension 1")
132     plt.ylabel("UMAP Dimension 2")
133     plt.colorbar(scatter, label="Class Labels (Ground Truth)")
134     plt.grid()
135     plt.show()
136
137     # Plot test data
138     plt.figure(figsize=(10, 6))
139     scatter = plt.scatter(testX_umap[:, 0], testX_umap[:, 1], c=
140         testY_pred, cmap='jet', s=15)
141     plt.title(title_test)
142     plt.xlabel("UMAP Dimension 1")
143     plt.ylabel("UMAP Dimension 2")
144     plt.colorbar(scatter, label="Class Labels (Predicted)")
145     plt.grid()
146     plt.show()
147
148 #Reading data

```

```

145 trainX, trainY = reading_images('FaceRecognition/train', 21, 30, True,
    32)
146 testX, testY = reading_images('FaceRecognition/test', 21, 30, True, 32)
147
148 trainX = trainX.T
149 trainY = trainY.T
150 testX = testX.T
151 testY = testY.T
152
153
154 #plot the first four faces
155 fig, axes = plt.subplots(1, 4, figsize=(12, 4)) # Create a 1x4 grid of
    subplots
156
157 for i in range(4):
158     # Reshape and take the real part of the eigenvector
159     image = trainX[i*21,:].reshape(32, 32).real
160
161     # Plot the image
162     axes[i].imshow(image)
163     axes[i].axis('off') # Turn off axis for a cleaner look
164     axes[i].set_title(f"Face {i+1}")
165
166 plt.tight_layout() # Adjust layout to avoid overlap
167 plt.show()
168
169
170 #PCA calculation
171
172 # Parameters
173 P = 16 # Number of principal components
174 K = 3 # Number of neighbors for KNN
175
176 # Step 1: Perform manual PCA on the training data
177 trainX_pca, train_mean, eigenvectors = manual_pca(trainX, P)
178
179 # Step 2: Project the test data using the same eigenvectors
180 testX_centered = testX - train_mean
181 testX_pca = np.dot(testX_centered, eigenvectors)
182
183 # Step 3: Classify using kNN and compute accuracies
184 train_predictions = knn_classify(trainX_pca, trainY, trainX_pca, K)
185 test_predictions = knn_classify(trainX_pca, trainY, testX_pca, K)
186
187 train_accuracy = compute_accuracy(trainY, train_predictions)
188 test_accuracy = compute_accuracy(testY, test_predictions)
189
190 # Print results
191 print(f"Training Accuracy: {train_accuracy:.2f}%")
192 print(f"Testing Accuracy: {test_accuracy:.2f}%")
193
194 # Step 4: Analyze accuracy for different values of P
195 p_values = range(1, P + 1)
196 train_accuracies = []
197 test_accuracies = []
198
199 for p in p_values:
200     trainX_pca, train_mean, eigenvectors = manual_pca(trainX, p)

```

```

201     testX_pca = np.dot(testX - train_mean, eigenvectors)
202     train_predictions = knn_classify(trainX_pca, trainY, trainX_pca, K)
203     test_predictions = knn_classify(trainX_pca, trainY, testX_pca, K)
204     train_accuracies.append(compute_accuracy(trainY, train_predictions)
205                             )
206     test_accuracies.append(compute_accuracy(testY, test_predictions))
207
208 # Plot PCA results
209 plt.figure(figsize=(10, 6))
210 plt.plot(p_values, train_accuracies, label="Training Accuracy", marker
211         ='o')
212 plt.plot(p_values, test_accuracies, label="Testing Accuracy", marker='s
213         ')
214 plt.xlabel("Number of Principal Components (P)")
215 plt.ylabel("Accuracy (%)")
216 plt.title("PCA and kNN Classification Accuracy")
217 plt.legend()
218 plt.grid()
219 plt.show()
220
221 #plot the first four eigen faces
222 fig, axes = plt.subplots(1, 4, figsize=(12, 4)) # Create a 1x4 grid of
223         subplots
224
225 for i in range(4):
226     # Reshape and take the real part of the eigenvector
227     image = eigenvectors[:, i].reshape(32, 32).real
228
229     # Plot the image
230     axes[i].imshow(image)
231     axes[i].axis('off') # Turn off axis for a cleaner look
232     axes[i].set_title(f"Eigen face {i+1}")
233
234 plt.tight_layout() # Adjust layout to avoid overlap
235 plt.show()
236
237
238 #2. LDA calculation
239
240 # Parameters for LDA
241 L = 2 # Number of LDA components
242
243 # Perform LDA on training data
244 trainX_lda, lda_eigvecs = manual_lda(trainX, trainY, L)
245
246 # Project test data using LDA
247 testX_lda = np.dot(testX, lda_eigvecs)
248
249 # Classify using kNN
250 train_predictions_lda = knn_classify(trainX_lda, trainY, trainX_lda, K)
251 test_predictions_lda = knn_classify(trainX_lda, trainY, testX_lda, K)
252
253 train_accuracy_lda = compute_accuracy(trainY, train_predictions_lda)
254 test_accuracy_lda = compute_accuracy(testY, test_predictions_lda)
255
256 print(f"LDA Training Accuracy: {train_accuracy_lda:.2f}%")

```

```

255 print(f"LDA Testing Accuracy: {test_accuracy_lda:.2f}%")
256
257 # Analyze accuracy for different LDA dimensions
258 lda_dimensions = range(1, 16 + 1)
259 train_accuracies_lda = []
260 test_accuracies_lda = []
261
262 for dim in lda_dimensions:
263     trainX_lda, lda_eigvecs = manual_lda(trainX, trainY, dim)
264     testX_lda = np.dot(testX, lda_eigvecs)
265     train_predictions_lda = knn_classify(trainX_lda, trainY, trainX_lda
266                                         , K)
267     test_predictions_lda = knn_classify(trainX_lda, trainY, testX_lda,
268                                         K)
269     train_accuracies_lda.append(compute_accuracy(trainY,
270                                                  train_predictions_lda))
271     test_accuracies_lda.append(compute_accuracy(testY,
272                                                  test_predictions_lda))
273
274 # Plot LDA results
275 plt.figure(figsize=(10, 6))
276 plt.plot(lda_dimensions, train_accuracies_lda, label="Training Accuracy
277         (LDA)", marker='o')
278 plt.plot(lda_dimensions, test_accuracies_lda, label="Testing Accuracy (
279         LDA)", marker='s')
280 plt.xlabel("Number of LDA Components")
281 plt.ylabel("Accuracy (%)")
282 plt.title("LDA and kNN Classification Accuracy")
283 plt.legend()
284 plt.grid()
285 plt.show()
286
287 # Plot PCA vs LDA testing results
288 plt.figure(figsize=(10, 6))
289 plt.plot(lda_dimensions, test_accuracies, label="PCA Testing Accuracy",
290         marker='o')
291 plt.plot(lda_dimensions, test_accuracies_lda, label="LDA Testing
292         Accuracy", marker='s')
293 plt.xlabel("Dimensionality")
294 plt.ylabel("Accuracy (%)")
295 plt.title("PCA v.s. LDA")
296 plt.legend()
297 plt.grid()
298 plt.show()
299
300 #Already obtained the autoencoder numbers by running the code provided
301     for P=3, 8, 16
302 autoencoder_accuracies = [90.1587, 99.0476, 100.00]
303 PCA_3_8_16_accuracy = [test_accuracies[2], test_accuracies[7],
304                        test_accuracies[-1]]
305 LDA_3_8_16_accuracy = [test_accuracies_lda[2], test_accuracies_lda[7],
306                        test_accuracies_lda[-1]]
307
308 # Plot LDA v.s. PCA v.s. Autoencoder results
309 plt.figure(figsize=(10, 6))

```

```

302 plt.plot([3, 8, 16], PCA_3_8_16_accuracy, label="PCA Testing Accuracy",
303         marker='o')
304 plt.plot([3, 8, 16], LDA_3_8_16_accuracy, label="LDA Testing Accuracy",
305         marker='s')
306 plt.plot([3, 8, 16], autoencoder_accuracies, label="Autoencoder Testing
307         Accuracy", marker='s')
308 plt.xlabel("Dimensionality")
309 plt.ylabel("Accuracy (%)")
310 plt.title("LDA v.s. Autoencoder v.s. PCA")
311 plt.legend()
312 plt.grid()
313 plt.show()
314
315 # Plot Autoencoder results alone
316 plt.figure(figsize=(10, 6))
317 plt.plot([3, 8, 16], autoencoder_accuracies, label="Autoencoder Testing
318         Accuracy", marker='s')
319 plt.xlabel("Dimensionality")
320 plt.ylabel("Accuracy (%)")
321 plt.title("Autoencoder accuracy as a function of dimensionality")
322 plt.legend()
323 plt.grid()
324 plt.show()
325
326 #Plot the UMAP
327
328 # Experiment parameters
329 P_values = [3, 8, 16] # Example values for PCA
330 L_values = [3, 8, 16] # Example values for LDA
331
332 # Generate plots for PCA and LDA
333 for P in P_values:
334     # PCA Projection
335     trainX_pca, _, eigenvectors_pca = manual_pca(trainX, P)
336     testX_pca = np.dot(testX - train_mean, eigenvectors_pca)
337
338     # kNN Predictions
339     testY_pred_pca = knn_classify(trainX_pca, trainY, testX_pca, K)
340
341     # Plot UMAP for PCA embeddings
342     plot_umap_embeddings(
343         trainX_proj=trainX_pca,
344         trainY=trainY,
345         testX_proj=testX_pca,
346         testY_pred=testY_pred_pca,
347         title_train=f"PCA (P={P}): Training Data",
348         title_test=f"PCA (P={P}): Test Data"
349     )
350
351 for L in L_values:
352     # LDA Projection
353     trainX_lda, lda_eigvecs = manual_lda(trainX, trainY, L)
354     testX_lda = np.dot(testX, lda_eigvecs)
355
356     # kNN Predictions

```

```

356 testY_pred_lda = knn_classify(trainX_lda, trainY, testX_lda, K)
357
358 # Plot UMAP for LDA embeddings
359 plot_umap_embeddings(
360     trainX_proj=trainX_lda,
361     trainY=trainY,
362     testX_proj=testX_lda,
363     testY_pred=testY_pred_lda,
364     title_train=f"LDA (P={L}): Training Data",
365     title_test=f"LDA (P={L}): Test Data"
366 )
367
368 #3. AdaBoost
369
370 def calculate_features(img):
371     # to grayscale
372     if len(img.shape) > 2:
373         img = cv2.cvtColor(img, cv2.COLOR_BGR2GRAY)
374
375     # Compute the integral image
376     integral_img = cv2.integral(img)
377
378     window_sizes = [2, 4, 8, 16]
379     features = []
380
381     # Loop over windows
382     for win_size in window_sizes:
383         for y in range(0, img.shape[0] - win_size + 1, win_size):
384             for x in range(0, img.shape[1] - win_size + 1, win_size):
385                 # horizontal feature
386                 mid_x = x + win_size // 2
387                 sum_left = integral_img[y + win_size, mid_x] -
388                     integral_img[y, mid_x] \
389                     - (integral_img[y + win_size, x] -
390                        integral_img[y, x])
391                 sum_right = integral_img[y + win_size, x + win_size] -
392                     integral_img[y, x + win_size] \
393                     - (integral_img[y + win_size, mid_x] -
394                        integral_img[y, mid_x])
395                 two_rect_horizontal = sum_right - sum_left
396                 features.append(two_rect_horizontal)
397
398                 # vertical feature
399                 mid_y = y + win_size // 2
400                 sum_top = integral_img[mid_y, x + win_size] -
401                     integral_img[mid_y, x] \
402                     - (integral_img[y, x + win_size] -
403                        integral_img[y, x])
404                 sum_bottom = integral_img[y + win_size, x + win_size] -
405                     integral_img[y + win_size, x] \
406                     - (integral_img[mid_y, x + win_size] -
407                        integral_img[mid_y, x])
408                 two_rect_vertical = sum_bottom - sum_top
409                 features.append(two_rect_vertical)
410
411     return np.array(features)

```

```

406 def concatenate_feats(pos_feats, neg_feats):
407     return np.concatenate((pos_feats, neg_feats), axis=0)
408
409
410 def get_labels_from_feats(pos_feats, neg_feats):
411     """
412     Return positvie and negative labels, un-concatenated given positive
413     and negative features
414     """
415     positive_labels = np.ones((pos_feats.shape[0]))
416     negative_labels = np.zeros((neg_feats.shape[0]))
417
418     return positive_labels, negative_labels
419
420 def concatenate_labels(pos_labels, neg_labels):
421     combined_labels = np.concatenate((pos_labels, neg_labels), axis =
422     0)
423     combined_labels = combined_labels.astype(np.uint8)
424
425     return combined_labels
426
427 #
428 def get_features(data_path, training = True):
429     """
430     Input:
431     - data_path = "CarDetection"
432     - training = True if the features needed are the training features.
433       False for testing data features
434     Returns:
435     - positive_feat: positive data features
436     - negative_feat: negative data features
437     - positive_labels
438     - negative_labels
439     """
440     positive_feat = []
441     negative_feat = []
442
443     for category in ["positive", "negative"]:
444         print("Getting the features of " + category + " training data"
445             if training
446             else "Getting the features of " + category + " testing
447             data")
448         current_path = data_path
449         current_path += "/train/" + category if training else "/test/"
450         +category
451         print ("Current path: ", current_path)
452
453         for file_name in os.listdir(current_path):
454             #Read the images
455             current_image = cv2.imread(current_path+"/"+file_name)
456             if category == "positive":
457                 positive_feat.append(calculate_features(current_image))
458             else:
459                 negative_feat.append(calculate_features(current_image))
460
461     #Create the labels
462     positive_feat = np.array(positive_feat)

```

```

458     negative_feat = np.array(negative_feat)
459
460     positive_labels, negative_labels = get_labels_from_feats(
461         positive_feat, negative_feat)
462
463     return positive_feat, negative_feat, positive_labels,
464         negative_labels
465
466 #Weak classifier
467 def get_weak_classifier(combined_feats, combined_labels, weights):
468
469     #Init error to large value
470     classifier_err = 1e15
471
472     for feat in range(combined_feats.shape[1]):
473         current_feat = combined_feats[:, feat]
474
475         #Sort the features for Ranking
476         sorted_ind = np.argsort(current_feat)
477
478         #Sort the features, labels, and weights
479         current_feat_sorted = current_feat[sorted_ind]
480         current_labels_sorted = combined_labels[sorted_ind]
481         current_weights_sorted = weights[sorted_ind]
482
483         #Declare pos, neg weights
484         num_feat = combined_feats.shape[0]
485         pos_weights = np.zeros((num_feat, 1))
486         neg_weights = np.zeros((num_feat, 1))
487
488         #Assign pos, neg weights
489
490         pos_weights[current_labels_sorted == 1, 0] =
491             current_weights_sorted[current_labels_sorted == 1]
492         neg_weights[current_labels_sorted == 0, 0] =
493             current_weights_sorted[current_labels_sorted == 0]
494
495         #Error polarity
496         pos_weights_cumsum = np.cumsum(pos_weights)
497         neg_weights_cumsum = np.cumsum(neg_weights)
498
499         pos_weights_sum = np.sum(pos_weights)
500         neg_weights_sum = np.sum(neg_weights)
501
502         error_polarity1 = pos_weights_cumsum + neg_weights_sum -
503             neg_weights_cumsum
504         error_polarity2 = neg_weights_cumsum + pos_weights_sum -
505             pos_weights_cumsum
506
507         #Unsqueeze
508         error_polarity1 = np.reshape(error_polarity1, [-1, 1])#np.
509             expand_dims(error_polarity1, 1)
510         error_polarity2 = np.reshape(error_polarity2, [-1, 1])#np.
511             expand_dims(error_polarity2, 1)
512
513         #Stack errors

```

```

507         total_error = np.concatenate((error_polarity1, error_polarity2)
508                                     , axis = 1)
509
510         #Get the index where error is minimum
511         total_min_error = np.argmin(total_error)
512         error_min_index = np.unravel_index(total_min_error, total_error
513                                     .shape)
514
515         current_min_error = total_error[error_min_index]
516
517         #Check if the new error is lower than the previous classifier
518         error
519         #If yes, assign it as the new classifier error
520         if current_min_error < classifier_err:
521             classifier_err = current_min_error
522
523             new_threshold = current_feat_sorted[error_min_index[0]]
524
525             polarity = 0
526             prediction = current_feat < new_threshold
527             if error_min_index[1]==0:
528                 polarity = 1
529                 prediction = current_feat >= new_threshold
530
531             best_classifier = [feat, new_threshold, polarity,
532                             classifier_err, prediction]
533
534     return best_classifier
535
536 def get_cascaded_AdaBoost(combined_feats, combined_labels, epochs,
537 levels):
538     """
539     Input:
540     - combined_feats = an array of combined features for positive and
541       negative
542     - combined_labels = for positive and negative
543     - epochs = how many epochs or iterations per cascade
544     - levels = how many cascade levels to run
545     """
546
547     #Init trust factor to something very small
548     best_trust = -1e16
549
550     #Initi weights
551     pos_w_total = np.sum(combined_labels == 1)
552     neg_w_total = np.sum(combined_labels == 0)
553     pos_w = np.repeat(1 / pos_w_total, pos_w_total)
554     neg_w = np.repeat(1 / neg_w_total, neg_w_total)
555
556     weights = np.concatenate((pos_w, neg_w), axis = 0)
557
558     for epoch in range(epochs):
559
560         weights_sum = np.sum(weights)
561         weights = weights / weights_sum #Normalizing the weights

```

```

559
560 #Get a weak classifier
561 current_weak_classifier = get_weak_classifier(combined_feats=
    combined_feats, combined_labels= combined_labels, weights=
    weights)
562
563 current_feat_index = current_weak_classifier[0]
564 current_threshold = current_weak_classifier[1]
565 current_polarity = current_weak_classifier[2]
566 current_classifier_err = current_weak_classifier[3]
567 current_prediction = current_weak_classifier[4]
568
569 #Set the negative labels to -1
570 current_prediction = np.where(current_prediction == 0, -1,
    current_prediction)
571
572
573 weights_expanded = np.expand_dims(weights, 0)
574 #Calculate absolute prediction error
575 label_err = np.abs(current_prediction - combined_labels)
576
577
578 current_epsilon = np.matmul(weights_expanded, label_err.T)
579 current_epsilon = np.expand_dims(current_epsilon, 1).squeeze()
580 current_epsilon *= 0.5
581
582 #Update trust factor and weights
583 trust_factor = np.log((1-current_epsilon)/current_epsilon)
584 trust_factor *= 0.5
585
586 #Updating the weights based on the trust factor
587 weights = weights * np.exp(- trust_factor * combined_labels *
    current_prediction)
588
589 #Computing the Falsse Positives FP, and False Negatives FN
590 #Get the false positive and false negative rates
591 total_negatives = np.sum(combined_labels==0)
592 total_positives = np.sum(combined_labels==1)
593
594 misclassified_negatives = np.sum((combined_labels == 0) & (
    current_prediction == 1))
595 misclassified_positives = np.sum((combined_labels == 1) & (
    current_prediction == -1))
596
597 FP = misclassified_negatives / total_negatives
598 FN = misclassified_positives / total_positives
599
600 print ("Training cascase level: %d. Epoch: %d. Trust: %.3f.
    Epsilon: %.3f. FP: %.7f. FN: %.7f." %
601         (levels+1, epoch, trust_factor, current_epsilon, FP, FN)
    )
602
603 #Check the trust factor
604 if (trust_factor > best_trust):
605     best_trust = trust_factor
606     best_prediction = current_prediction
607     best_FP = FP
608     best_FN = FN

```

```

609         best_classifier = [current_feat_index, current_threshold,
610                             current_polarity, current_classifier_err,
611                             current_prediction, FP, FN,
612                             current_epsilon, trust_factor]
613
614     positive_feats = combined_feats[:, np.sum(combined_labels==1),:]
615     negative_feats = combined_feats[np.sum(combined_labels==1):, :]
616     negative_feats_mask = np.where(best_prediction[np.sum(
617         combined_labels==1):]==1)
618
619     negative_feats = negative_feats[negative_feats_mask, :][0]
620
621     combined_feats = concatenate_feats(positive_feats, negative_feats)
622
623     current_pos_labels, current_neg_labels = get_labels_from_feats(
624         positive_feats, negative_feats)
625
626     combined_labels = concatenate_labels(current_pos_labels,
627         current_neg_labels)
628
629     return combined_feats, combined_labels, best_FP, best_FN,
630         best_classifier
631
632 def predict_ada(cascade_classifiers, features, labels):
633     feats_number = features.shape[0]
634     final_labels = np.ones((feats_number, 1))
635     class_indices = np.arange(features.shape[0]).reshape(-1, 1)
636     FPRs, FNRs = [], []
637
638     for level, classifier in enumerate(cascade_classifiers.values()):
639         feature_id, threshold, polarity, trust_factor = classifier[0],
640             classifier[1], classifier[2], classifier[-1]
641
642         feature_column = features[:, feature_id].reshape(-1, 1)
643         predictions = (feature_column >= threshold) if polarity == 1
644             else (feature_column < threshold)
645         predictions = np.where(predictions, 1, -1)
646
647         weighted_preds = trust_factor * predictions
648         final_preds = weighted_preds >= trust_factor
649
650         total_negatives = np.sum(labels == 0)
651         total_positives = np.sum(labels == 1)
652
653         FPR = np.sum(final_preds[total_positives:] == 1) /
654             total_negatives if total_negatives else 0
655         FNR = 1 - np.sum(final_preds[:total_positives] == 1) /
656             total_positives if total_positives else 0
657
658         features = features[final_preds.flatten() == 1]
659         labels = labels[final_preds.flatten() == 1]
660         negative_indices = class_indices[final_preds.flatten() == 0]
661         final_labels[negative_indices] = -1
662         class_indices = class_indices[final_preds.flatten() == 1]

```

```

657         FPRs.append(FPR if level == 0 else FPRs[-1] * FPR)
658         FNRs.append(FNR if level == 0 else FNRs[-1] * FNR)
659
660     return FPRs, FNRs, final_labels
661
662 def plot_adaboost(FPs, FNs, title=None):
663
664     plt.rcParams.update({
665         'font.size': 10,
666         'figure.titlesize': 14,
667         'legend.fontsize': 10,
668         'axes.titlesize': 10,
669         'axes.labelsize': 12,
670         'xtick.labelsize': 12,
671         'ytick.labelsize': 12
672     })
673
674     plt.figure()
675     plt.plot(FPs, '-*', label='False Positive FP', linewidth=3)
676     plt.plot(FNs, '-d', label='False Negative FN', linewidth=3)
677     plt.legend()
678
679     plt.xlabel('Cascade Levels')
680     plt.ylabel('Percentage %')
681     plt.xticks(np.arange(len(FPs)), [str(ix) for ix in range(1, len(FPs)
682         ) + 1)])
683     plt.title(title)
684     plt.tight_layout()
685     plt.show()
686
687 def excute_adaboost(number_of_cascades = 10, iterations = [1, 5, 25],
688     data_path = "CarDetection", train = True):
689
690     """
691     Inputs:
692     - number_of_cascades: how many levels
693     - iterations: iterations per cascade
694     - data_path: parent data directory "CarDetection"
695     - train: if True, excute adaboost on training dataset, else do
696         inference on testing data
697     NOTE: the code has to be run on training mode first to save the
698         classifier that will be used for testing.
699
700     Returns:
701     None. Plots will be shown and accuracies will be printed.
702     """
703     for epoch in iterations:
704         current_FPR = 1
705         current_FNR = 1
706         FPs = []
707         FNs = []
708         tol = 1e-6
709         best_classifier={}
710
711         tolerance_counter = 0
712         if train:
713             print("Performing Adaboost on Training data")

```

```

711     train_positive_feats, train_negative_feats,
       train_positive_labels, train_negative_labels =
       get_features(data_path, train)
712 combined_train_feats = concatenate_feats(
       train_positive_feats, train_negative_feats)
713 combined_train_labels = concatenate_labels(
       train_positive_labels, train_negative_labels)
714
715
716 for level in range(number_of_cascades):
717     current_Ada = get_cascaded_AdaBoost(
       combined_train_feats, combined_train_labels, epoch,
       level)
718
719     combined_train_feats, combined_train_labels, best_FPR,
       best_FNR, best_weak_classifier = current_Ada
720
721     best_classifier[str(level + 1)] = best_weak_classifier
722     current_FPR *= best_FPR
723     current_FNR *= best_FNR
724     FPs.append(current_FPR)
725     FNs.append(current_FNR)
726
727     print("Current cumulative FP: %.7f, cumulative FN %.7F"
       % (current_FPR, current_FNR))
728
729     #Increasing the tolerance
730     tolerance_counter = level + 1
731
732     if current_FPR <= tol:
733         print("We reached FP tolerance level:", tol, 'after
       cascade level =', level + 1)
734
735     if np.sum(combined_train_labels == 0) == 0:
736         print("No images with negative labels remained
       after cascade level: ", level+1)
737         break
738
739     #Plot the FP and FN rates
740     current_title = "AdaBoost performance on Training data.
       Cascade levels: " + str(number_of_cascades) + ". After "
       + str(epoch) + " iterations per cascade level."
741     plot_adaboost(FPs= FPs, FNs= FNs, title = current_title)
742
743     #Save the best classifier
744     with open("Results/AdaboostMine" + '/'
       Adaboost_classifier_level_' + str(epoch) + '_' + str(
       number_of_cascades) + '.pkl', 'wb') as f:
745         pickle.dump(best_classifier, f)
746
747 #Test data
748 else:
749     print("Performing Adaboost on Testing data")
750     #loading the best classifier
751     with open("Results/AdaboostMine" + '/'
       Adaboost_classifier_level_' + str(epoch) + '_' + str(
       number_of_cascades) + '.pkl', 'rb') as f:
752         best_classifier = pickle.load(f)

```

```

753     #Get the test features and labels
754     test_positive_feats, test_negative_feats,
755         test_positive_labels, test_negative_labels =
756         get_features(data_path, False)
757     combined_test_feats = concatenate_feats(test_positive_feats
758         , test_negative_feats)
759     combined_test_labels = concatenate_labels(
760         test_positive_labels, test_negative_labels)
761
762     #Get the FP and FN
763     FPs, FNs, _ = predict_ada(best_classifier,
764         combined_test_feats, combined_test_labels)
765     print("Current Testing cumulative FP: %.7f, cumulative FN
766         %.7F" % (np.array(FPs).min(), np.array(FNs).min()))
767
768     #current_title = "Test Data - Iteration: " + str(epoch) + ".
769         Reached minimum at cascade level: "+str(
770         number_of_cascades)
771     current_title = "AdaBoost performance on Testing data.
772         Cascade levels: " + str(number_of_cascades) + ". After "
773         + str(epoch) + " iterations per cascade level."
774     plot_adaboost(FPs = FPs, FNs = FNs, title = current_title)
775
776 #Training data AdaBoost
777 excute_adaboost(number_of_cascades= 6, iterations=[1, 3, 10], data_path
778     = "CarDetection", train = True)
779
780 #Performing Adaboost on testing data
781 excute_adaboost(number_of_cascades= 6, iterations=[1, 3, 10], data_path
782     = "CarDetection", train = False)

```

Listing 2: Assignment 10 Autoencoder Code

```

1 import os
2
3 import numpy as np
4 import torch
5 from torch import nn, optim
6 from PIL import Image
7 from torch.autograd import Variable
8 from torch.utils.data import Dataset, DataLoader
9 from torchvision import transforms
10 import umap
11 import matplotlib.pyplot as plt
12 from sklearn.neighbors import KNeighborsClassifier
13
14 class DataBuilder(Dataset):
15     def __init__(self, path):
16         self.path = path
17         self.image_list = [f for f in os.listdir(path) if f.endswith('.
18             png')]
19         self.label_list = [int(f.split('_')[0]) for f in self.
20             image_list]
21         self.len = len(self.image_list)
22         self.aug = transforms.Compose([
23             transforms.Resize((64, 64)),

```

```

22         transforms.ToTensor(),
23     ])
24
25     def __getitem__(self, index):
26         fn = os.path.join(self.path, self.image_list[index])
27         x = Image.open(fn).convert('RGB')
28         x = self.aug(x)
29         return {'x': x, 'y': self.label_list[index]}
30
31     def __len__(self):
32         return self.len
33
34
35 class Autoencoder(nn.Module):
36
37     def __init__(self, encoded_space_dim):
38         super().__init__()
39         self.encoded_space_dim = encoded_space_dim
40         ### Convolutional section
41         self.encoder_cnn = nn.Sequential(
42             nn.Conv2d(3, 8, 3, stride=2, padding=1),
43             nn.LeakyReLU(True),
44             nn.Conv2d(8, 16, 3, stride=2, padding=1),
45             nn.LeakyReLU(True),
46             nn.Conv2d(16, 32, 3, stride=2, padding=1),
47             nn.LeakyReLU(True),
48             nn.Conv2d(32, 64, 3, stride=2, padding=1),
49             nn.LeakyReLU(True)
50         )
51         ### Flatten layer
52         self.flatten = nn.Flatten(start_dim=1)
53         ### Linear section
54         self.encoder_lin = nn.Sequential(
55             nn.Linear(4 * 4 * 64, 128),
56             nn.LeakyReLU(True),
57             nn.Linear(128, encoded_space_dim * 2)
58         )
59         self.decoder_lin = nn.Sequential(
60             nn.Linear(encoded_space_dim, 128),
61             nn.LeakyReLU(True),
62             nn.Linear(128, 4 * 4 * 64),
63             nn.LeakyReLU(True)
64         )
65         self.unflatten = nn.Unflatten(dim=1,
66                                       unflattened_size=(64, 4, 4))
67         self.decoder_conv = nn.Sequential(
68             nn.ConvTranspose2d(64, 32, 3, stride=2,
69                               padding=1, output_padding=1),
70             nn.BatchNorm2d(32),
71             nn.LeakyReLU(True),
72             nn.ConvTranspose2d(32, 16, 3, stride=2,
73                               padding=1, output_padding=1),
74             nn.BatchNorm2d(16),
75             nn.LeakyReLU(True),
76             nn.ConvTranspose2d(16, 8, 3, stride=2,
77                               padding=1, output_padding=1),
78             nn.BatchNorm2d(8),
79             nn.LeakyReLU(True),

```

```

80         nn.ConvTranspose2d(8, 3, 3, stride=2,
81                             padding=1, output_padding=1)
82     )
83
84     def encode(self, x):
85         x = self.encoder_cnn(x)
86         x = self.flatten(x)
87         x = self.encoder_lin(x)
88         mu, logvar = x[:, :self.encoded_space_dim], x[:, self.
            encoded_space_dim:]
89         return mu, logvar
90
91     def decode(self, z):
92         x = self.decoder_lin(z)
93         x = self.unflatten(x)
94         x = self.decoder_conv(x)
95         x = torch.sigmoid(x)
96         return x
97
98     @staticmethod
99     def reparameterize(mu, logvar):
100         std = logvar.mul(0.5).exp_()
101         eps = Variable(std.data.new(std.size()).normal_())
102         return eps.mul(std).add_(mu)
103
104
105 class VaeLoss(nn.Module):
106     def __init__(self):
107         super(VaeLoss, self).__init__()
108         self.mse_loss = nn.MSELoss(reduction="sum")
109
110     def forward(self, xhat, x, mu, logvar):
111         loss_MSE = self.mse_loss(xhat, x)
112         loss_KLD = -0.5 * torch.sum(1 + logvar - mu.pow(2) - logvar.exp
            ())
113         return loss_MSE + loss_KLD
114
115
116 def train(epoch):
117     model.train()
118     train_loss = 0
119
120     for batch_idx, data in enumerate(trainloader):
121         optimizer.zero_grad()
122         mu, logvar = model.encode(data['x'])
123         z = model.reparameterize(mu, logvar)
124         xhat = model.decode(z)
125         loss = vae_loss(xhat, data['x'], mu, logvar)
126         loss.backward()
127         train_loss += loss.item()
128         optimizer.step()
129
130     print('====> Epoch: {} Average loss: {:.4f}'.format(
131         epoch, train_loss / len(trainloader.dataset)))
132
133 def calculate_accuracy(true_labels, distances, num_classes, k_neighbors
    =1):
134     temp_predictions = []

```

```

135
136 # Loop through each data point
137 for sample_idx in range(distances.shape[0]):
138     sample_distances = distances[sample_idx, :].copy()
139     class_votes = np.zeros(num_classes)
140     class_total_distances = np.zeros(num_classes)
141
142     # Find the k nearest neighbors
143     for _ in range(k_neighbors):
144         nearest_idx = np.argmin(sample_distances)
145         nearest_class = true_labels[nearest_idx] - 1
146         class_votes[nearest_class] += 1
147         class_total_distances[nearest_class] += sample_distances[
148             nearest_idx]
149         sample_distances[nearest_idx] = np.inf # Exclude this
150         index from future searches
151
152     # Determine the predicted class
153
154     avg_class_distances = class_total_distances / np.max(
155         class_votes)
156     zero_indices = np.where(avg_class_distances == 0)
157     avg_class_distances[zero_indices] = np.inf
158     predicted_class = np.argmin(avg_class_distances) + 1
159     temp_predictions.append(predicted_class)
160
161 # Convert predictions to a numpy array
162 temp_predictions = np.array(temp_predictions, dtype=int).squeeze()
163
164 # Calculate accuracy
165 matches = (true_labels == temp_predictions).astype(int)
166 correct_predictions = np.sum(matches)
167 accuracy_percentage = (correct_predictions / true_labels.shape[0])
168 * 100
169
170 return accuracy_percentage
171 #####
172 # Change these
173 p = 3 # [3, 8, 16]
174 training = False
175 TRAIN_DATA_PATH = 'FaceRecognition/train'
176 EVAL_DATA_PATH = 'FaceRecognition/test'
177 LOAD_PATH = f'weights/model_{p}.pt'
178 OUT_PATH = 'autoencoderOutput'
179 #####
180
181 model = Autoencoder(p)
182
183 if training:
184     epochs = 100
185     log_interval = 1
186     trainloader = DataLoader(
187         dataset=DataBuilder(TRAIN_DATA_PATH),
188         batch_size=12,
189         shuffle=True,
190     )
191     optimizer = optim.Adam(model.parameters(), lr=1e-3)

```

```

189     vae_loss = VaeLoss()
190     for epoch in range(1, epochs + 1):
191         train(epoch)
192     torch.save(model.state_dict(), os.path.join(OUT_PATH, f'model_{p}.
        pt'))
193 else:
194     trainloader = DataLoader(
195         dataset=DataBuilder(TRAIN_DATA_PATH),
196         batch_size=1,
197     )
198     model.load_state_dict(torch.load(Load_PATH))
199     model.eval()
200
201     X_train, y_train = [], []
202     for batch_idx, data in enumerate(trainloader):
203         mu, logvar = model.encode(data['x'])
204         z = mu.detach().cpu().numpy().flatten()
205         X_train.append(z)
206         y_train.append(data['y'].item())
207     X_train = np.stack(X_train)
208     y_train = np.array(y_train)
209
210     testloader = DataLoader(
211         dataset=DataBuilder(EVAL_DATA_PATH),
212         batch_size=1,
213     )
214     X_test, y_test = [], []
215     for batch_idx, data in enumerate(testloader):
216         mu, logvar = model.encode(data['x'])
217         z = mu.detach().cpu().numpy().flatten()
218         X_test.append(z)
219         y_test.append(data['y'].item())
220     X_test = np.stack(X_test)
221     y_test = np.array(y_test)
222
223     #####
224     # Your code starts here
225     samples = X_test.shape[0]
226     current_distance = np.zeros((samples, samples))
227     for i in range(samples):
228         squared_diff = np.square(X_train - X_test[i, :])
229         ith_distance = np.sqrt(np.sum(squared_diff, axis=1)).T
230         current_distance[i, :] = ith_distance
231
232     #Calculate the accuracy
233     #How many images per class
234     imagesPerClass = 21
235     total_accuracy = calculate_accuracy(y_test, current_distance, 30,
        1)
236
237     print ("With P = %d. Accuracy = %.4f" % (p, total_accuracy))
238     #####
239
240
241
242 def plot_umap_embeddings(X_train, y_train, X_test, y_test, y_test_pred,
    p):
243     """

```

```

244     Function to plot UMAP embeddings for training and test data.
245     """
246     # Initialize UMAP reducer
247     reducer = umap.UMAP(n_components=2, random_state=42)
248
249     # Fit UMAP on training data
250     X_train_umap = reducer.fit_transform(X_train)
251
252     # Transform test data
253     X_test_umap = reducer.transform(X_test)
254
255     # Plot training data
256     plt.figure(figsize=(8, 6))
257     scatter = plt.scatter(
258         X_train_umap[:, 0],
259         X_train_umap[:, 1],
260         c=y_train,
261         cmap="jet",
262         s=15,
263         alpha=0.8,
264     )
265     plt.title(f"UMAP Embedding of Training Data (P = {p})")
266     plt.colorbar(scatter, label="Ground Truth Labels")
267     plt.xlabel("UMAP Dimension 1")
268     plt.ylabel("UMAP Dimension 2")
269     plt.grid()
270     plt.show()
271
272     # Plot test data
273     plt.figure(figsize=(8, 6))
274     scatter = plt.scatter(
275         X_test_umap[:, 0],
276         X_test_umap[:, 1],
277         c=y_test_pred,
278         cmap="jet", #"tab10",
279         s=15,
280         alpha=0.8,
281     )
282     plt.title(f"UMAP Embedding of Test Data (P = {p})")
283     plt.colorbar(scatter, label="Predicted Labels")
284     plt.xlabel("UMAP Dimension 1")
285     plt.ylabel("UMAP Dimension 2")
286     plt.grid()
287     plt.show()
288
289
290 if not training:
291     # UMAP and Visualization
292     knn_classifier = KNeighborsClassifier(n_neighbors=1)
293     knn_classifier.fit(X_train, y_train)
294     y_test_pred = knn_classifier.predict(X_test)
295
296     # Generate UMAP plots
297     plot_umap_embeddings(X_train, y_train, X_test, y_test, y_test_pred,
298                          p)

```