

CS2109S Tutorial 7

Unsupervised Learning

(AY 24/25 Semester 2)

March 27, 2025

(Prepared by Benson)

Contents

Clustering

Q1. K-Means Algorithm

Q2. Kernel K-Means

Principal Component Analysis (PCA)

Q3. Lossy Compression

Bonus. Composing Kernels

Q1. K-Means Algorithm

Intuition:

- ▶ Lines 4-5: Minimize the distortion by reassigning clusters,
keeping the centroids unchanged.
- ▶ Lines 6-7: Minimize the distortion by changing centroids,
keeping the assignment unchanged.

Algorithm K-Means Clustering

```

1: for  $k = 1$  to  $K$  do
2:    $\mu_k \leftarrow$  random location
3: while not converged do
4:   for  $i = 1$  to  $m$  do
5:      $c^{(i)} \leftarrow \operatorname{argmin}_k \|\mathbf{x}^{(i)} - \mu_k\|^2$ 
6:   for  $k = 1$  to  $K$  do
7:      $\mu_k \leftarrow \frac{1}{|\{\mathbf{x}^{(i)} | c^{(i)} = k\}|} \sum_{\mathbf{x} \in \{\mathbf{x}^{(i)} | c^{(i)} = k\}} \mathbf{x}$ 

```

$$\frac{d\mathcal{E}}{d\mu_k} = 0 \quad \Rightarrow \quad \sum_{c^{(i)}=k} 2(\mathbf{x}^{(i)} - \mu_k) = 0 \quad \Rightarrow \quad \mu_k = \frac{1}{|\{\mathbf{x}^{(i)} | c^{(i)} = k\}|} \sum_{c^{(i)}=k} \mathbf{x}^{(i)}$$

Q1. K-Means Algorithm

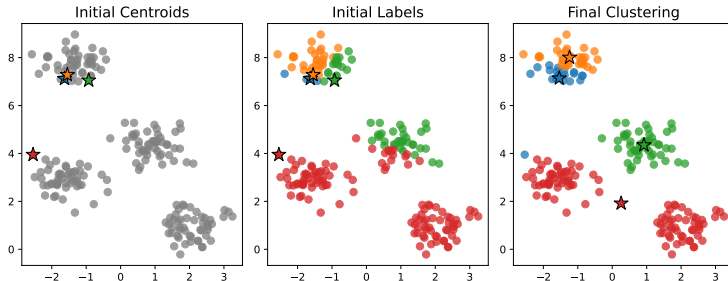
- (a) Given that every iteration produces a partition with a lower distortion, prove that this algorithm always converges. Convergence is when the centroids/medoids do not change after an iteration of the algorithm.
- ▶ There are not more than k^N possible partitions.
 - ▶ Since every iteration of the algorithm produces a partition that has a lower distortion, it eventually terminates.

Algorithm K-Means Clustering

```
1: for  $k = 1$  to  $K$  do
2:    $\mu_k \leftarrow$  random location
3: while not converged do
4:   for  $i = 1$  to  $m$  do
5:      $c^{(i)} \leftarrow \operatorname{argmin}_k \|\mathbf{x}^{(i)} - \mu_k\|^2$ 
6:   for  $k = 1$  to  $K$  do
7:      $\mu_k \leftarrow \frac{1}{|\{\mathbf{x}^{(i)} | c^{(i)} = k\}|} \sum_{\mathbf{x} \in \{\mathbf{x}^{(i)} | c^{(i)} = k\}} \mathbf{x}$ 
```

Q1. K-Means Algorithm

- (b) Although k-means always converges, it may get stuck at a bad local minimum. As mentioned in the lecture, one method of circumventing this is to run the algorithm multiple times and choose the clusters with the minimum distortion. Suggest some other ways to help the algorithm get closer to the global minimum.



- Intuition: Pick centroids that are **as far as possible** during random initialization.

Q1. K-Means Algorithm

Deterministic method:

Algorithm Cluster initialization

- 1: $\mu_1 \leftarrow$ random location
 - 2: **for** $k = 2$ **to** K **do**
 - 3: $\mu_k \leftarrow$ point farthest possible from
 $\mu_1, \dots, \mu_{k-1}$
-

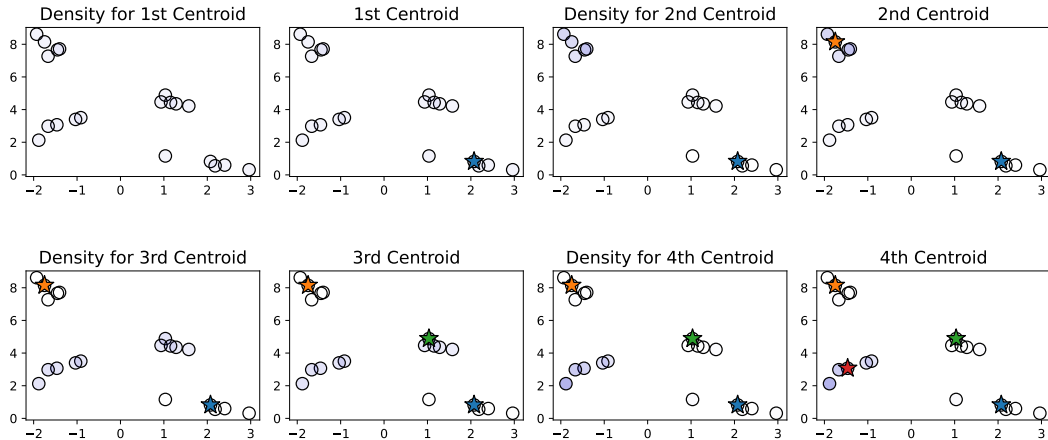
Random method (K-means++):

Algorithm Cluster initialization

- 1: $\mu_1 \leftarrow$ random location
 - 2: **for** $k = 2$ **to** K **do**
 - 3: $D(x) \leftarrow$ distance from x to closest centroid in $\mu_1, \dots, \mu_{k-1}$
 - 4: $\mu_k \leftarrow$ random point with (weighted) probability distribution $D(x)^2$
-

Q1. K-Means Algorithm

K-means++ Initialization



Q1. K-Means Algorithm

- (c) Cluster the 6 points in table 1 into two clusters using the K-means algorithm. The two initial centroids are (0, 1) and (2.5, 2).

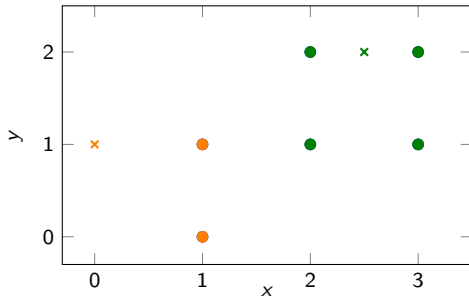
Iteration 1

i	1	2	3	4	5	6
x	1	1	2	2	3	3
y	0	1	1	2	1	2
c	1	1	2	2	2	2

- ▶ $\mu_1 = \frac{1}{2}((1, 0) + (1, 1)) = (1, 0.5)$
- ▶ $\mu_2 = \frac{1}{4}((2, 1) + (2, 2) + (3, 1) + (3, 2)) = (2.5, 1.5)$

Algorithm K-Means Clustering

```
1: for  $k = 1$  to  $K$  do
2:    $\mu_k \leftarrow$  random location
3:   while not converged do
4:     for  $i = 1$  to  $m$  do
5:        $c^{(i)} \leftarrow \operatorname{argmin}_k \|\mathbf{x}^{(i)} - \mu_k\|^2$ 
6:       for  $k = 1$  to  $K$  do
7:          $\mu_k \leftarrow \frac{1}{|\{\mathbf{x}^{(i)} | c^{(i)} = k\}|} \sum_{\mathbf{x} \in \{\mathbf{x}^{(i)} | c^{(i)} = k\}} \mathbf{x}$ 
```



Q1. K-Means Algorithm

- (c) Cluster the 6 points in table 1 into two clusters using the K-means algorithm. The two initial centroids are (0, 1) and (2.5, 2).

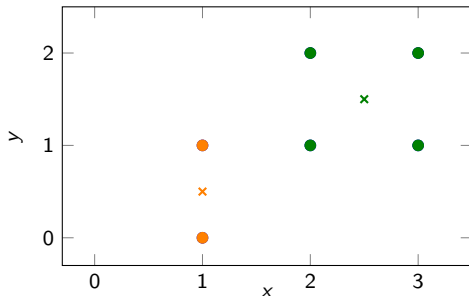
Iteration 2

i	1	2	3	4	5	6
x	1	1	2	2	3	3
y	0	1	1	2	1	2
c	1	1	2	2	2	2

- ▶ $\mu_1 = \frac{1}{2}((1, 0) + (1, 1)) = (1, 0.5)$
- ▶ $\mu_2 = \frac{1}{4}((2, 1) + (2, 2) + (3, 1) + (3, 2)) = (2.5, 1.5)$

Algorithm K-Means Clustering

```
1: for  $k = 1$  to  $K$  do
2:    $\mu_k \leftarrow$  random location
3:   while not converged do
4:     for  $i = 1$  to  $m$  do
5:        $c^{(i)} \leftarrow \operatorname{argmin}_k \|\mathbf{x}^{(i)} - \mu_k\|^2$ 
6:     for  $k = 1$  to  $K$  do
7:        $\mu_k \leftarrow \frac{1}{|\{\mathbf{x}^{(i)} | c^{(i)} = k\}|} \sum_{\mathbf{x} \in \{\mathbf{x}^{(i)} | c^{(i)} = k\}} \mathbf{x}$ 
```



Q1. K-Means Algorithm

- (d) Cluster the 6 points in table 1 into two clusters using the K-medoids algorithm. The initial medoids are point 1 and point 3.

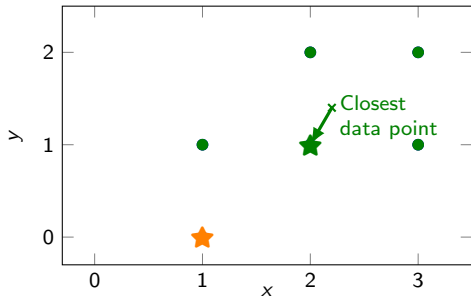
Iteration 1

i	1	2	3	4	5	6
x	1	1	2	2	3	3
y	0	1	1	2	1	2
c	1	2*	2	2	2	2

- ▶ $\mu_1 = (1, 0)$
- ▶ $\mu_2 = \frac{1}{5}((1, 1) + (2, 1) + (2, 2) + (3, 1) + (3, 2)) = (2.2, 1.4) \Rightarrow (2, 1)$

Algorithm K-Medoids Clustering

```
1: for  $k = 1$  to  $K$  do
2:    $\mu_k \leftarrow$  random data point  $\mathbf{x}^{(i)}$ 
3:   while not converged do
4:     for  $i = 1$  to  $m$  do
5:        $c^{(i)} \leftarrow \operatorname{argmin}_k \|\mathbf{x}^{(i)} - \mu_k\|^2$ 
6:       for  $k = 1$  to  $K$  do
7:          $\mu_k \leftarrow \frac{1}{|\{\mathbf{x}^{(i)} | c^{(i)} = k\}|} \sum_{\mathbf{x} \in \{\mathbf{x}^{(i)} | c^{(i)} = k\}} \mathbf{x}$ 
8:       for  $k = 1$  to  $K$  do  $\triangleright$  Closest data point
9:          $\mu_k \leftarrow \operatorname{argmin}_{\mathbf{x} \in \{\mathbf{x}^{(i)} | c^{(i)} = k\}} \|\mathbf{x}^{(i)} - \mu_k\|^2$ 
```



Recap: Kernel Trick

Important Property of a Kernel Function

$$K(\mathbf{u}, \mathbf{v}) = \phi(\mathbf{u}) \cdot \phi(\mathbf{v}) \text{ for some function } \phi$$

If we have a model that **ONLY** relies on the dot products of $\mathbf{x}^{(i)}$:

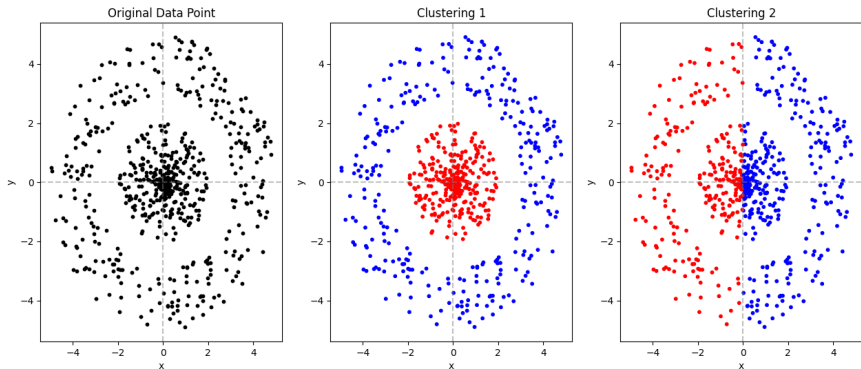
$$\text{SVM}_{\alpha}(\mathbf{x}) = \text{sign} \left(\sum_{i=1}^N \alpha^{(i)} y^{(i)} (\mathbf{x}^{(i)} \cdot \mathbf{x}) \right)$$

Applying the kernel function $\equiv$ applying the transformed features!

$$\text{SVM}_{\alpha}(\mathbf{x}) = \text{sign} \left(\sum_{i=1}^N \alpha^{(i)} y^{(i)} (\phi(\mathbf{x}^{(i)}) \cdot \phi(\mathbf{x})) \right) = \text{sign} \left(\sum_{i=1}^N \alpha^{(i)} y^{(i)} K(\mathbf{x}^{(i)}, \mathbf{x}) \right)$$

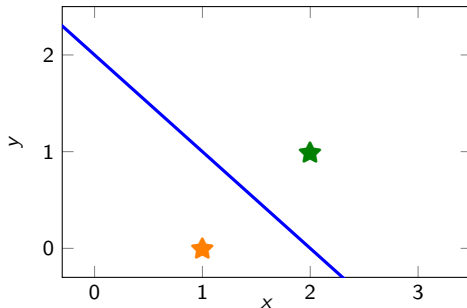
Q2. Kernel K-Means

- (a) Clustering 1 should be considered better than Clustering 2. Can the K-means algorithm achieve the better clustering? Why or why not?



Q2. Kernel K-Means

- (a) Clustering 1 should be considered better than Clustering 2. Can the K-means algorithm achieve the better clustering? Why or why not?
- ▶ No. K-means uses **Euclidean distance**.
 - ▶ Assignment to clusters will follow a **linear boundary**.



Q2. Kernel K-Means

- (b) Write the squared Euclidean distance between two points $\mathbf{p}_i$ and $\mathbf{p}_j$ using vector norms (length) and the dot product. How can we apply the kernel trick?

$$\begin{aligned}\|\mathbf{p}_i - \mathbf{p}_j\|^2 &= (\mathbf{p}_i - \mathbf{p}_j) \cdot (\mathbf{p}_i - \mathbf{p}_j) \\ &= \mathbf{p}_i \cdot \mathbf{p}_i - \mathbf{p}_i \cdot \mathbf{p}_j - \mathbf{p}_j \cdot \mathbf{p}_i + \mathbf{p}_j \cdot \mathbf{p}_j \\ &= \mathbf{p}_i \cdot \mathbf{p}_i - 2(\mathbf{p}_i \cdot \mathbf{p}_j) + \mathbf{p}_j \cdot \mathbf{p}_j\end{aligned}$$

Naively applying transformed features:

$$\phi(\mathbf{p}_i) \cdot \phi(\mathbf{p}_i) - 2(\phi(\mathbf{p}_i) \cdot \phi(\mathbf{p}_j)) + \phi(\mathbf{p}_j) \cdot \phi(\mathbf{p}_j)$$

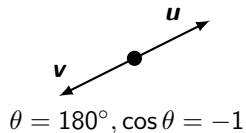
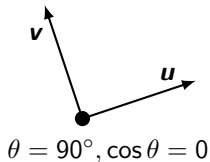
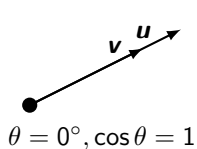
Kernel trick:

$$K(\mathbf{p}_i, \mathbf{p}_i) - 2K(\mathbf{p}_i, \mathbf{p}_j) + K(\mathbf{p}_j, \mathbf{p}_j)$$

Q2. Kernel K-Means

Dot products can be viewed as a **similarity measure**.

$$\mathbf{u} \cdot \mathbf{v} = \|\mathbf{u}\| \|\mathbf{v}\| \cos \theta$$



Q2. Kernel K-Means

(c) Gaussian radial basis function (RBF) kernel:

distance $\Rightarrow$ similarity measure!

$$k_{rbf}(\mathbf{p}_i, \mathbf{p}_j) = \exp\left(-\frac{\|\mathbf{p}_i - \mathbf{p}_j\|^2}{2\sigma^2}\right)$$

Calculate $k_{rbf}(\mathbf{p}_i, \mathbf{p}_j)$ for all pairs of data points. Explain what the value represents.

$\mathbf{p}$	$k(\mathbf{p}_i, \mathbf{p}_1)$	$k(\mathbf{p}_i, \mathbf{p}_2)$	$k(\mathbf{p}_i, \mathbf{p}_3)$	$k(\mathbf{p}_i, \mathbf{p}_4)$	$k(\mathbf{p}_i, \mathbf{p}_5)$
(0, 0)	1	e^{-1}	e^{-1}	e^{-1}	e^{-1}
(4, 4)	e^{-1}	1	e^{-2}	e^{-4}	e^{-2}
(-4, 4)	e^{-1}	e^{-2}	1	e^{-2}	e^{-4}
(-4, -4)	e^{-1}	e^{-4}	e^{-2}	1	e^{-2}
(4, -4)	e^{-1}	e^{-2}	e^{-4}	e^{-2}	1

p	x	y
1	0	0
2	4	4
3	-4	4
4	-4	-4
5	4	-4

Q2. Kernel K-Means

Intuition:

- ▶ Lines 4-5: Minimize the distortion by reassigning clusters,
keeping the centroids unchanged.
- ▶ **Do not compute the centers** (We cannot compute $\phi(\mathbf{x})$...).

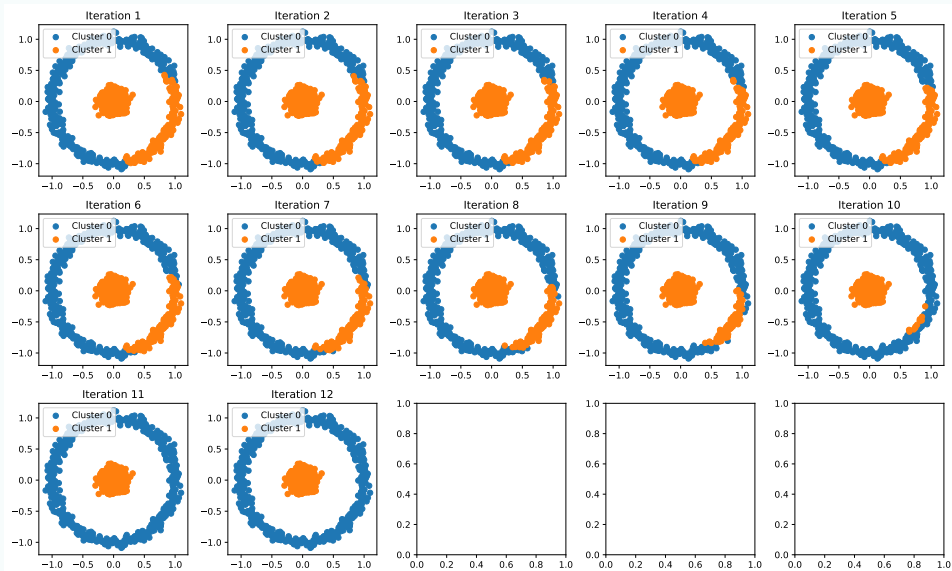
Algorithm Kernel K-Means Clustering

```
1: for  $k = 1$  to  $K$  do  
2:    $\mu_k \leftarrow$  random location  
3: while not converged do  
4:   for  $i = 1$  to  $m$  do  
5:      $c^{(i)} \leftarrow \operatorname{argmin}_k d_{rbf}^2(\mathbf{x}^{(i)}, \mu_k)$ 
```

$$\begin{aligned} & d_{rbf}^2(\mathbf{x}^{(i)}, \mu_c) \\ &= \phi(\mathbf{x}^{(i)}) \cdot \phi(\mathbf{x}^{(i)}) - 2 \cdot \left(\phi(\mathbf{x}^{(i)}) \cdot \frac{1}{\text{size}} \sum_{j \in \mathcal{C}} \phi(\mathbf{x}^{(j)}) \right) + \left(\left(\frac{1}{\text{size}} \sum_{j \in \mathcal{C}} \phi(\mathbf{x}^{(j)}) \right) \cdot \left(\frac{1}{\text{size}} \sum_{k \in \mathcal{C}} \phi(\mathbf{x}^{(k)}) \right) \right) \\ &= k_{rbf}(\mathbf{x}^{(i)}, \mathbf{x}^{(i)}) + \frac{2}{\text{size}} \sum_{j \in \mathcal{C}} k_{rbf}(\mathbf{x}^{(i)}, \mathbf{x}^{(j)}) + \frac{1}{\text{size}^2} \sum_{j, k \in \mathcal{C}} k_{rbf}(\mathbf{x}^{(j)}, \mathbf{x}^{(k)}) \end{aligned}$$

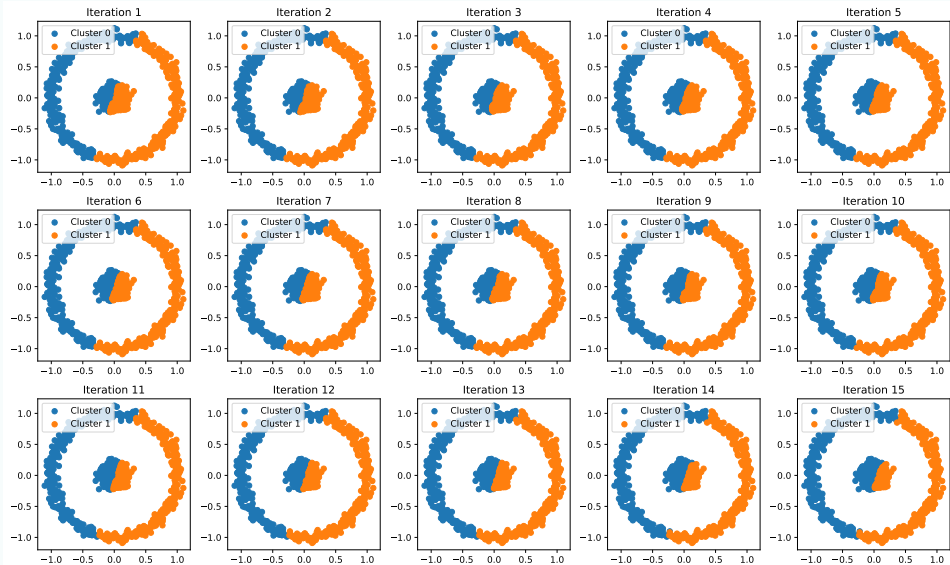
Q2. Kernel K-Means

Extra Slide



Q2. Kernel K-Means

Extra Slide



Recap: Principal Component Analysis (PCA)

$$\mathbf{X} = \begin{bmatrix} 4 & 2 & 5 & 1 \\ 1 & 3 & 4 & 0 \end{bmatrix}$$

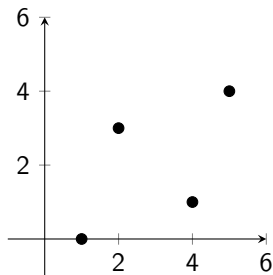
1. **Mean-center** the data (subtract each $\mathbf{x}$ by the mean $\bar{\mathbf{x}}$).

$$\bar{\mathbf{x}} = \frac{1}{4} \begin{bmatrix} 12 \\ 8 \end{bmatrix} = \begin{bmatrix} 3 \\ 2 \end{bmatrix}$$

$$\hat{\mathbf{X}} = \begin{bmatrix} 4 & 2 & 5 & 1 \\ 1 & 3 & 4 & 0 \end{bmatrix} - \begin{bmatrix} 3 & 3 & 3 & 3 \\ 2 & 2 & 2 & 2 \end{bmatrix} = \begin{bmatrix} 1 & -1 & 2 & -2 \\ -1 & 1 & 2 & -2 \end{bmatrix}$$

2. Compute the **covariance matrix** $\text{Cov}(\mathbf{X}) = \frac{1}{m} \hat{\mathbf{X}} \hat{\mathbf{X}}^T$.

$$\text{Cov}(\mathbf{X}) = \frac{1}{4} \begin{bmatrix} 1 & -1 & 2 & -2 \\ -1 & 1 & 2 & -2 \end{bmatrix} \begin{bmatrix} 1 & -1 \\ -1 & 1 \\ 2 & 2 \\ -2 & -2 \end{bmatrix} = \begin{bmatrix} 2.5 & 1.5 \\ 1.5 & 2.5 \end{bmatrix}$$



Recap: Principal Component Analysis (PCA)

$$\mathbf{X} = \begin{bmatrix} 4 & 2 & 5 & 1 \\ 1 & 3 & 4 & 0 \end{bmatrix}, \text{Cov}(\mathbf{X}) = \begin{bmatrix} 2.5 & 1.5 \\ 1.5 & 2.5 \end{bmatrix}$$

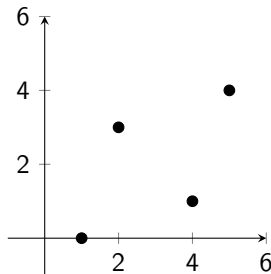
3. Compute the **singular value decomposition** of $\text{Cov}(\mathbf{X})$ ($\text{Cov}(\mathbf{X}) = \mathbf{V}\mathbf{\Sigma}^2\mathbf{V}^\top$).

$$\begin{bmatrix} 2.5 & 1.5 \\ 1.5 & 2.5 \end{bmatrix} = \begin{bmatrix} \frac{1}{\sqrt{2}} & \frac{1}{\sqrt{2}} \\ \frac{1}{\sqrt{2}} & -\frac{1}{\sqrt{2}} \end{bmatrix} \begin{bmatrix} 4 & 0 \\ 0 & 1 \end{bmatrix} \begin{bmatrix} \frac{1}{\sqrt{2}} & \frac{1}{\sqrt{2}} \\ \frac{1}{\sqrt{2}} & -\frac{1}{\sqrt{2}} \end{bmatrix}$$

4. Reduce the basis to k components to obtain the new basis $\tilde{\mathbf{U}}$.

Calculate the ratio of explained variance $\frac{\sum_{i=1}^r \sigma_i^2}{\sum_{i=1}^m \sigma_i^2}$.

$$\tilde{\mathbf{U}} = \begin{bmatrix} \frac{1}{\sqrt{2}} \\ \frac{1}{\sqrt{2}} \end{bmatrix}, \frac{\sum_{i=1}^r \sigma_i^2}{\sum_{i=1}^m \sigma_i^2} = \frac{4}{4+1} = 80\%$$



Recap: Principal Component Analysis (PCA)

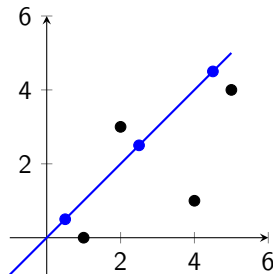
$$\mathbf{X} = \begin{bmatrix} 4 & 2 & 5 & 1 \\ 1 & 3 & 4 & 0 \end{bmatrix}, \tilde{\mathbf{U}} = \begin{bmatrix} \frac{1}{\sqrt{2}} \\ \frac{1}{\sqrt{2}} \end{bmatrix}$$

► **Reduction:** $\mathbf{Z} = \tilde{\mathbf{U}}^\top \mathbf{X}$

$$\mathbf{Z} = \begin{bmatrix} \frac{1}{\sqrt{2}} & \frac{1}{\sqrt{2}} \end{bmatrix} \begin{bmatrix} 4 & 2 & 5 & 1 \\ 1 & 3 & 4 & 0 \end{bmatrix} = \begin{bmatrix} \frac{5}{\sqrt{2}} & \frac{5}{\sqrt{2}} & \frac{9}{\sqrt{2}} & \frac{1}{\sqrt{2}} \end{bmatrix}$$

► **Reconstruction:** $\mathbf{X} \approx \tilde{\mathbf{U}} \mathbf{Z}$

$$\mathbf{X} \approx \begin{bmatrix} \frac{1}{\sqrt{2}} \\ \frac{1}{\sqrt{2}} \end{bmatrix} \begin{bmatrix} \frac{5}{\sqrt{2}} & \frac{5}{\sqrt{2}} & \frac{9}{\sqrt{2}} & \frac{1}{\sqrt{2}} \end{bmatrix} = \begin{bmatrix} 2.5 & 2.5 & 4.5 & 0.5 \\ 2.5 & 2.5 & 4.5 & 0.5 \end{bmatrix}$$

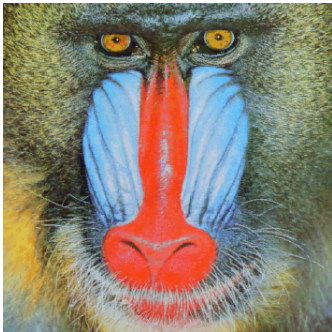


We are only required to store $\tilde{\mathbf{U}}$ and $\mathbf{Z}$ (instead of $\mathbf{X}$)!

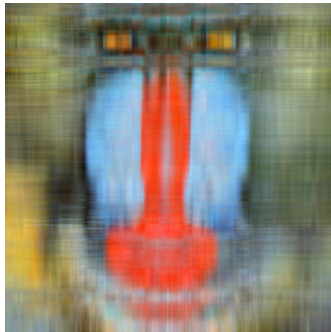
Q3. Lossy Compression

- (a) The current choice of $k = 9$ does not produce a very nice output. What is a good value for k ? Justify your answer.

Before



After $k=9$



Q3. Lossy Compression

1. Compute the **covariance matrix** $\text{Cov}(\mathbf{X}) = \frac{1}{m} \hat{\mathbf{X}} \hat{\mathbf{X}}^\top$ (*already includes mean-centering*).
2. Compute the **singular value decomposition** of $\text{Cov}(\mathbf{X})$ ($\text{Cov}(\mathbf{X}) = \mathbf{V} \mathbf{\Sigma}^2 \mathbf{V}^\top$).

```
1 cov = np.cov(X, bias=True)
2 U, s, VT = np.linalg.svd(cov)
```


Q3. Lossy Compression

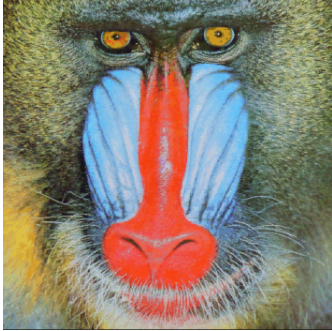
Goal: Retain at least 99% of the explained variance.

```
1 def explained_variance(s, k):
2     """ returns the explained variance when we keep the first k cols """
3     return sum(s[:k]) / sum(s)
4
5 minimum_retained_variance = 0.99
6 k = 0
7 var = 0
8 while var < minimum_retained_variance:
9     k += 1
10    var = explained_variance(s, k)
```

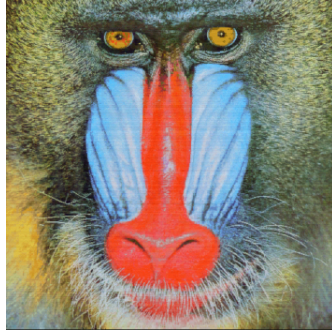
Q3. Lossy Compression

Answer: $k = 286$

Before



After



Q3. Lossy Compression

(b) For the value of k you select in (a), what is the space saved by doing this compression?

- ▶ Original: $\mathbf{X}$ (512×1536)
- ▶ New: $\tilde{\mathbf{U}}$ (512×286) and $\mathbf{Z}$ (286×1536).
- ▶ Compression ratio = $\frac{(512 \times 286) + (286 \times 1536)}{512 \times 1536} = 0.745$.
- ▶ 25.5% space saved!

Reduction: $\mathbf{Z} = \tilde{\mathbf{U}}^\top \mathbf{X}$
Reconstruction: $\mathbf{X} \approx \tilde{\mathbf{U}} \mathbf{Z}$

Q3. Lossy Compression

- (c) What are the drawbacks of this form of compression?
- ▶ This is a lossy compression as we do not regain 100% of the variance.
 - ▶ $\tilde{\mathbf{U}}$ is only used for one image $\Rightarrow$ waste of space (if k is large, we might end up using more space).

Bonus. Composing Kernels

Refer to the template code provided.

1. Let $K^{(1)}(\mathbf{u}, \mathbf{v})$ and $K^{(2)}(\mathbf{u}, \mathbf{v})$ be kernel functions.
Work out the transformed features in the kernel $K^{(1)}(\mathbf{u}, \mathbf{v}) + K^{(2)}(\mathbf{u}, \mathbf{v})$.
2. Let $K^{(1)}(\mathbf{u}, \mathbf{v})$ and $K^{(2)}(\mathbf{u}, \mathbf{v})$ be kernel functions.
Work out the transformed features in the kernel $K^{(1)}(\mathbf{u}, \mathbf{v}) \times K^{(2)}(\mathbf{u}, \mathbf{v})$.
3. Invent your own kernel!

Sample Solution:

1. We have

$$K^{(1)}(\mathbf{u}, \mathbf{v}) = \phi^{(1)}(\mathbf{u}) \cdot \phi_1(\mathbf{v})$$

$$K^{(2)}(\mathbf{u}, \mathbf{v}) = \phi^{(2)}(\mathbf{u}) \cdot \phi_2(\mathbf{v})$$

Then, we get

$$K^{(1)}(\mathbf{u}, \mathbf{v}) + K^{(2)}(\mathbf{u}, \mathbf{v}) = \begin{bmatrix} \phi^{(1)}(\mathbf{u}) \\ \phi^{(2)}(\mathbf{u}) \end{bmatrix} \cdot \begin{bmatrix} \phi^{(1)}(\mathbf{v}) \\ \phi^{(2)}(\mathbf{v}) \end{bmatrix}$$

which means the addition of two kernel functions essentially **concatenates** the transformed features together.

Bonus. Composing Kernels

2. Similarly, we get

$$\begin{aligned} K^{(1)}(\mathbf{u}, \mathbf{v}) \times K^{(2)}(\mathbf{u}, \mathbf{v}) &= (\phi^{(1)}(\mathbf{u}) \cdot \phi^{(1)}(\mathbf{v})) \times (\phi^{(2)}(\mathbf{u}) \cdot \phi^{(2)}(\mathbf{v})) \\ &= \left(\sum_i \phi_i^{(1)}(\mathbf{u}) \phi_i^{(1)}(\mathbf{v}) \right) \times \left(\sum_j \phi_j^{(2)}(\mathbf{u}) \phi_j^{(2)}(\mathbf{v}) \right) \\ &= \sum_{ij} \left(\phi_i^{(1)}(\mathbf{u}) \phi_j^{(2)}(\mathbf{u}) \right) \left(\phi_i^{(1)}(\mathbf{v}) \phi_j^{(2)}(\mathbf{v}) \right) \end{aligned}$$

Define another set of transformed features $\phi^*(\mathbf{u})$ which consists of $\phi_i^{(1)}(\mathbf{u}) \phi_j^{(2)}(\mathbf{u})$ for all pairs (i, j) , then

$$K^{(1)}(\mathbf{u}, \mathbf{v}) \times K^{(2)}(\mathbf{u}, \mathbf{v}) = \phi^*(\mathbf{u}) \cdot \phi^*(\mathbf{v})$$

which means the product of two kernel functions takes the **pairwise product** of between their transformed features.

1. Just a simple concatenation.

```
1 def transform_add(X1, X2):  
2     return np.hstack([X1, X2])
```

2. Broadcasting is your best friend!

```
1 def transform_multiply(X1, X2):  
2     n = X1.shape[0]  
3     pairwise_prod = X1.reshape(n, -1, 1) * X2.reshape(n, 1, -1)  
4     return pairwise_prod.reshape(n, -1)
```