

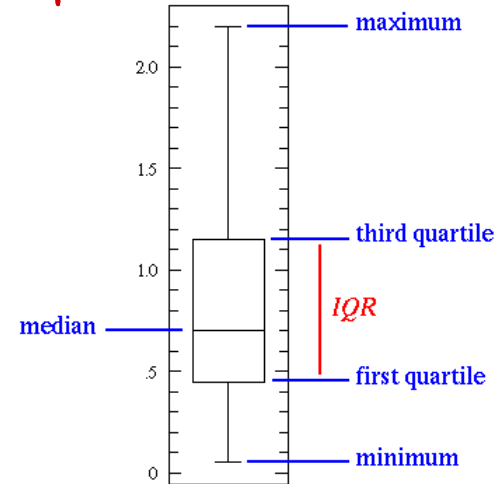
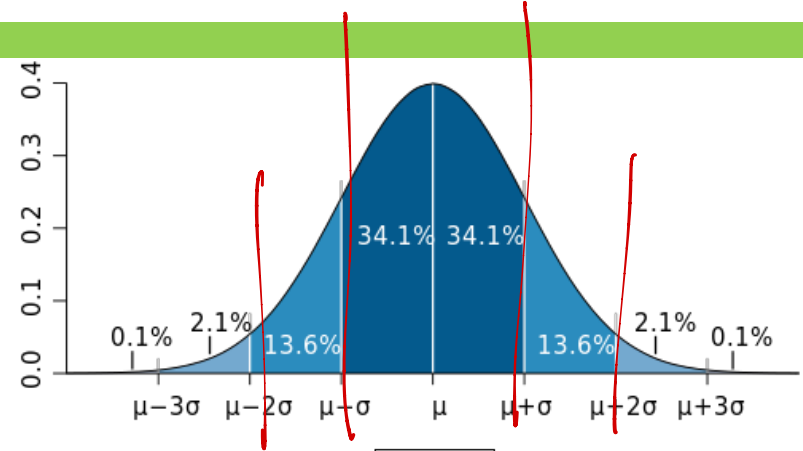
UNSUPERVISED LEARNING

OUTLIERS

APPLICATIONS

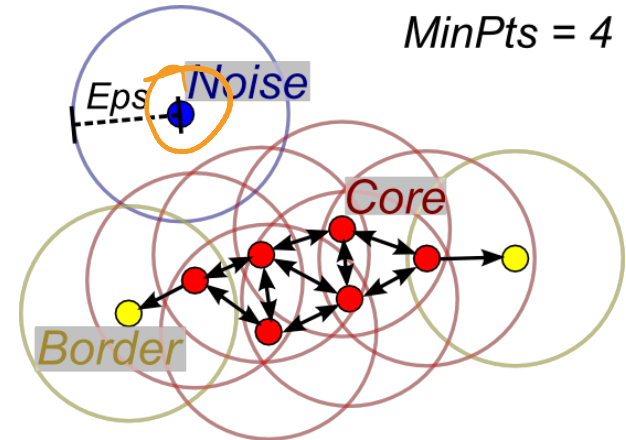
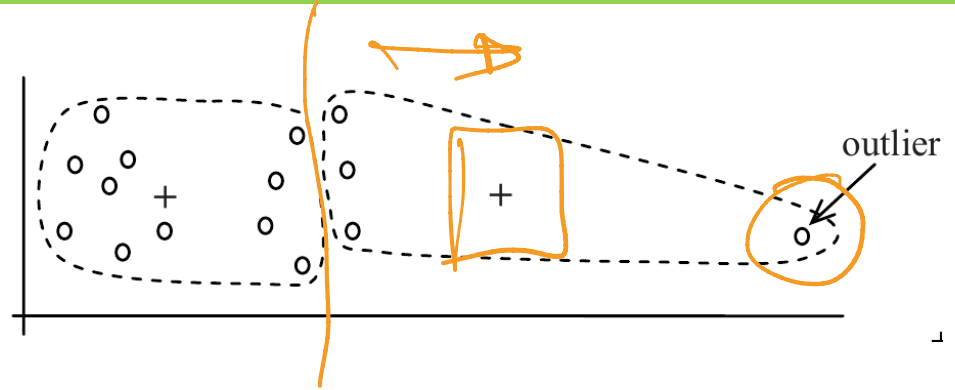
Outliers

- Values outside the normal range
- Statistically, define in terms of deviation from mean or median
- Gaussian distribution – number of standard deviations from mean
- Box and whisker plots – outer and inner fences based on median, interquartile (IQR) value



Outliers and clustering

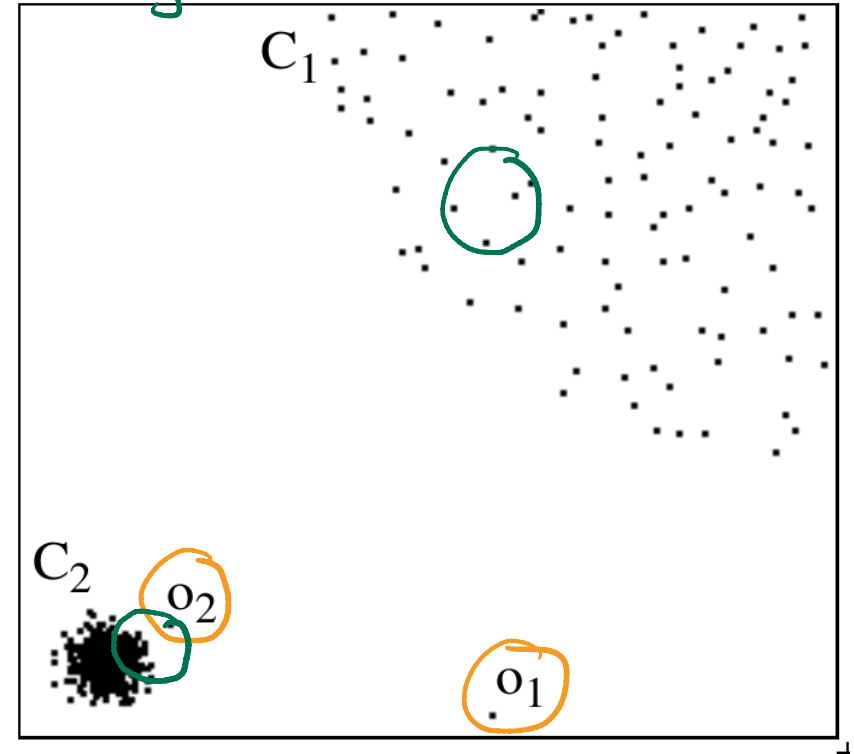
- Outliers are points that lie outside natural clusters
- K Means – far away from all centroids
 - But outliers can distort the clustering process
- Density based clustering – not connected to any core point
 - But density is applied uniformly



Outliers and density

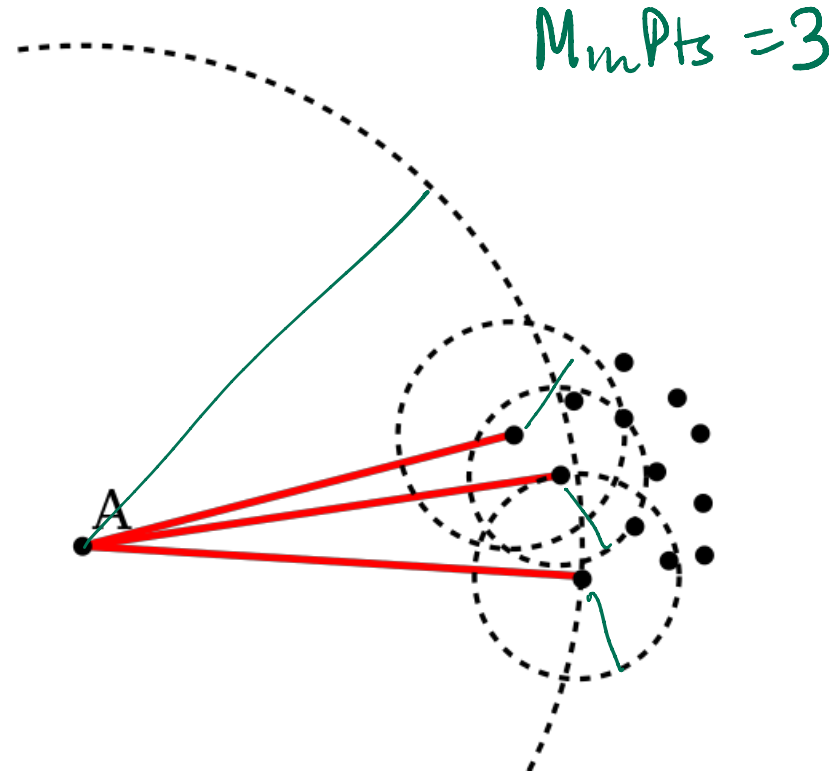
Is the density of a point similar to the density of its neighbours?

- An outlier is less dense than its nearest neighbours
- But difference in density may be local
- A distance metric to eliminate o_2 could make all of C_1 outliers
- C_1 has 400 points, C_2 has 100 points
- Larger distance would make all of C_2 outliers with respect to C_1



Outliers and density

- For clustering, we defined a radius Eps and looked for $MinPts$ neighbours within that ball
- Instead, fix $MinPts$ and find smallest ball with that many neighbours
- Compare $radius(p)$ with radius of its neighbours
- A is an outlier because its radius is much more than that of its neighbours

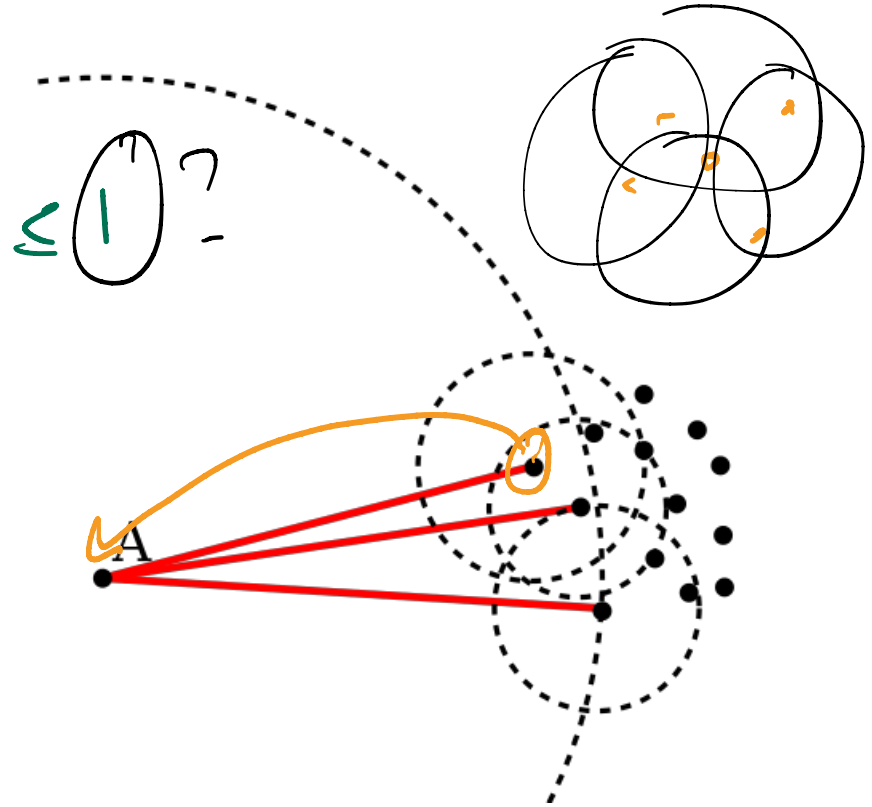


Outliers and density

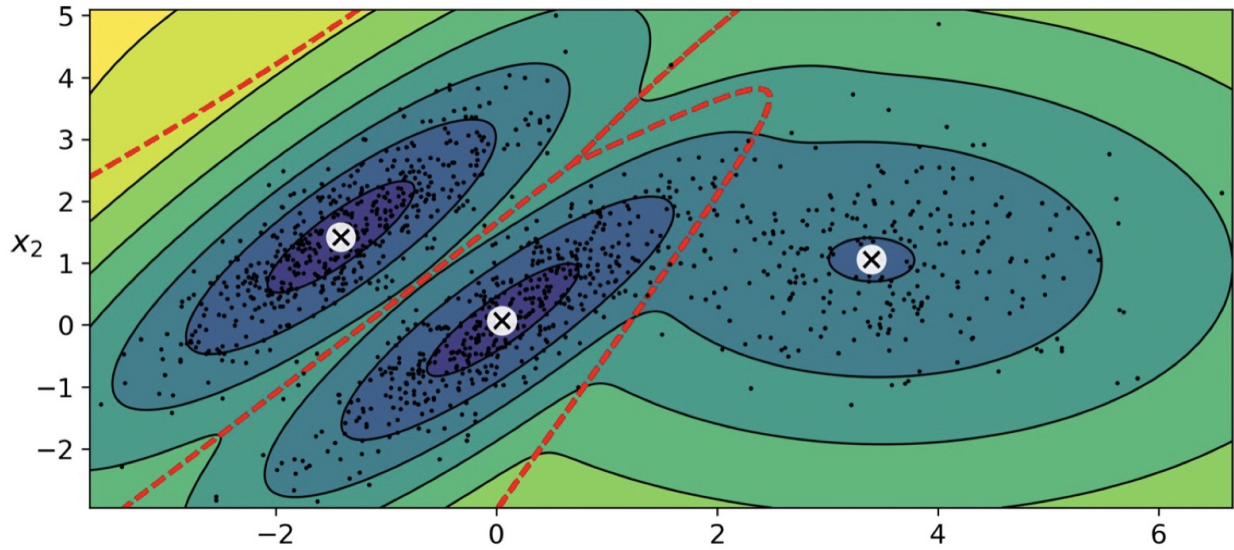
- Local outlier factor $LOF(p)$

$$0 \leq \frac{\text{Mean radius of } MinPts\text{-neighbours}(p)}{radius(p)}$$

- The smaller this ratio, the more likely that p is an outlier
- Comparison is local to neighbourhood, so this can deal with different densities across range of data



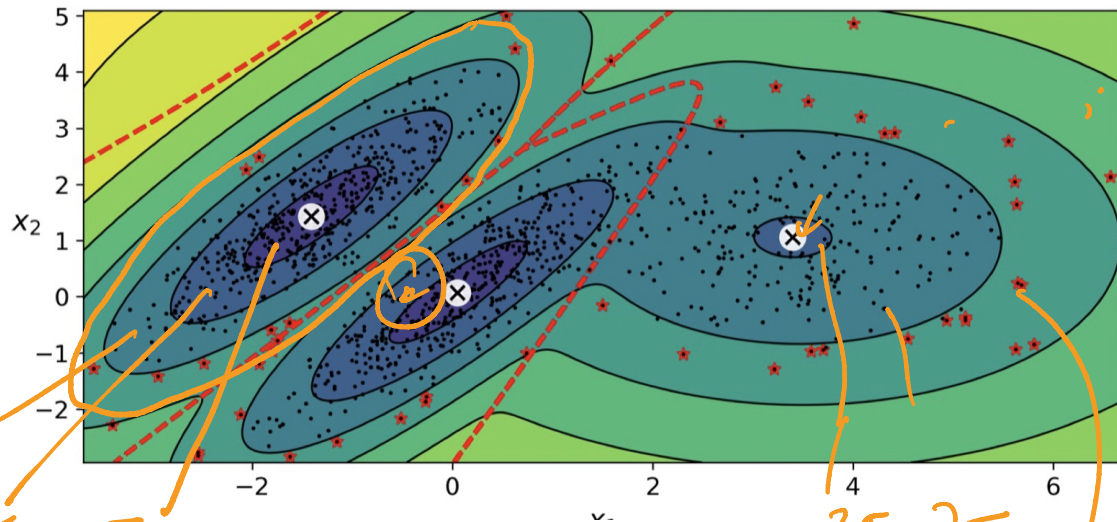
Clustering using Mixture of Gaussians



Examples : a Iris dataset

- Tshirt sizes - shoulder, chest, length
XS, S, M, L, XL

Outlier detection
using mixture of
Gaussians



Numeric data $\rightarrow (\mu_1, \sigma_1), (\mu_2, \sigma_2), (\mu_3, \sigma_3)$

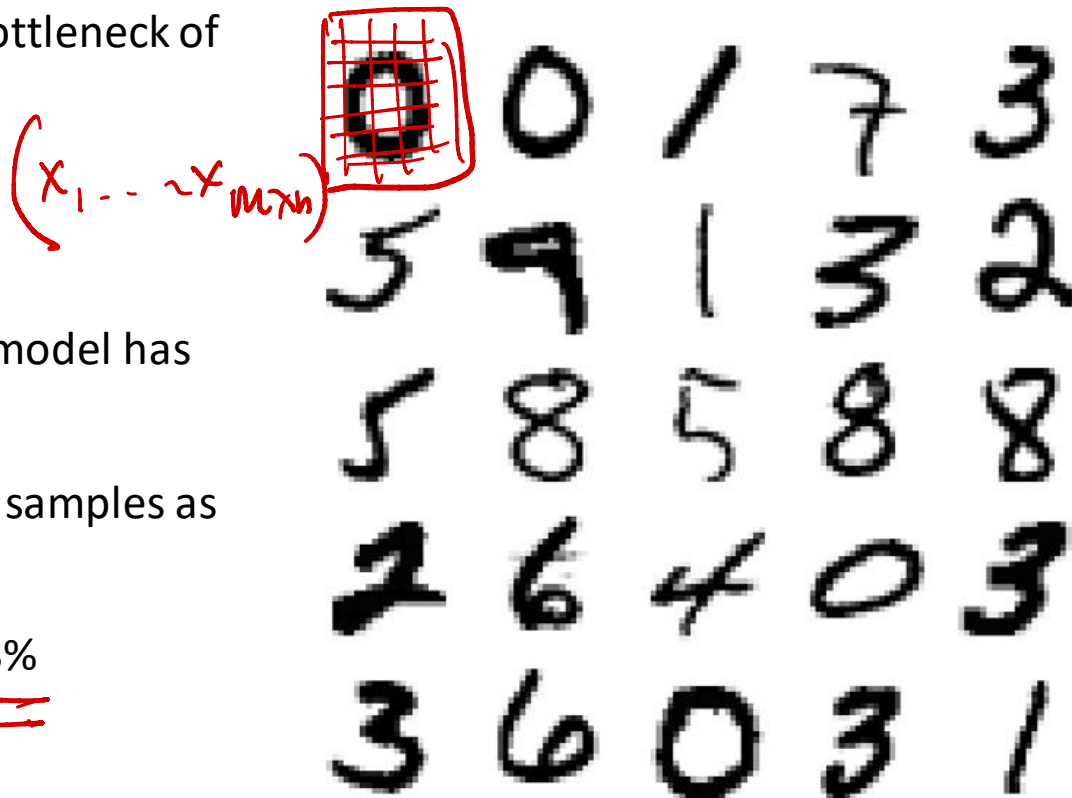
$> 3\sigma$ for all
Gaussians

Outliers

Semi-supervised learning

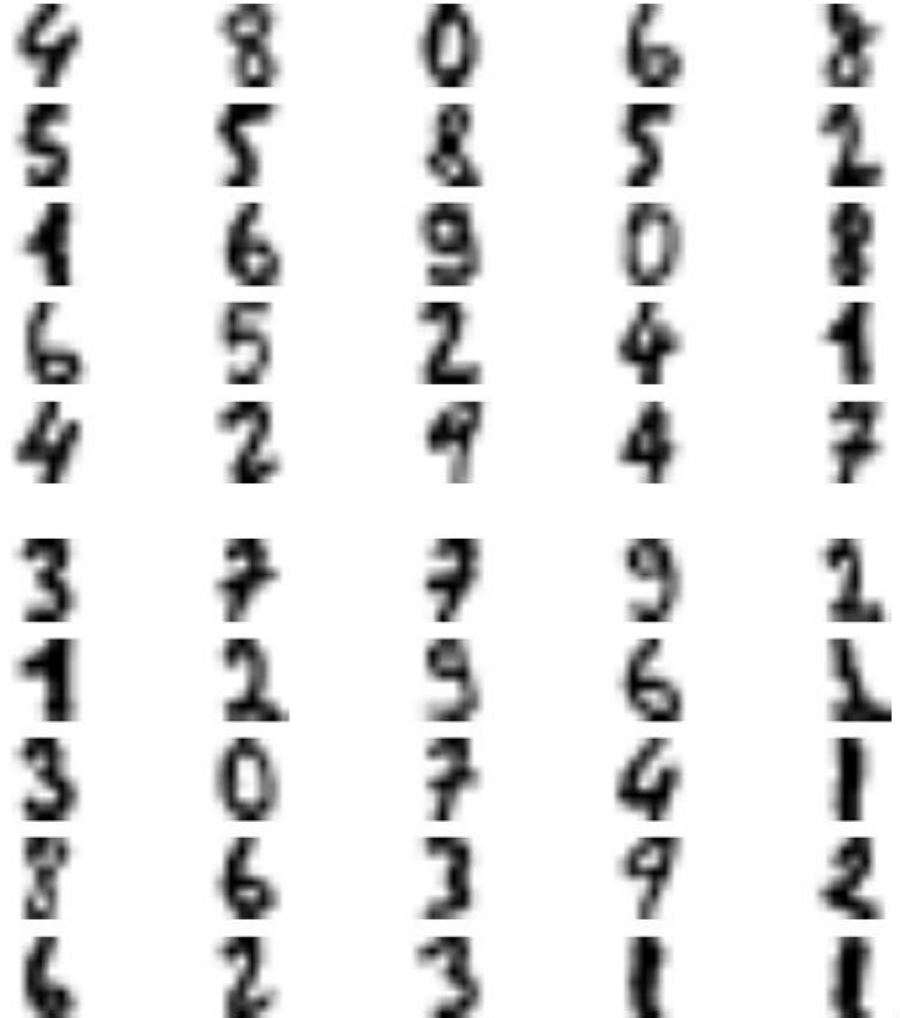
MNIST

- Labelling training data is a bottleneck of supervised learning
- Handwritten digits 0,1,...,9
 - 1797 images
- Standard logistic regression model has 96.9% accuracy
- Suppose we take 50 random samples as training set
- Logistic regression gives 83.3%



Semi-supervised learning

- Instead of 50 random samples, 50 clusters using K means
- Use image nearest to each centroid as trainingset
 - 50 *representative images*
- Logistic regression accuracy jumps to 92.2%



Semi-supervised learning



- Propagate representative image label to entire cluster
- Logistic regression improves to 93.3%
- Propagate representative image label to only 20% items closest to centroid
- Logistic regression improves to 94%
- Only 50 actual labels used, about 5 per class!

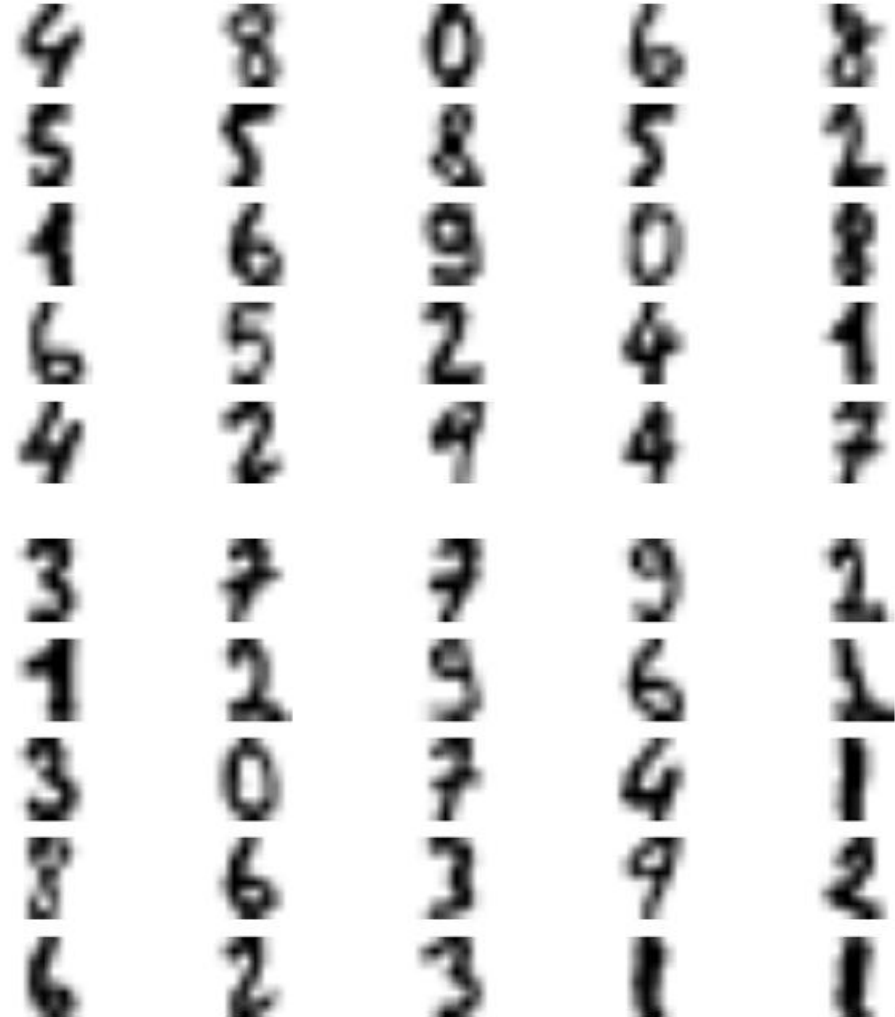


Image segmentation

- An image is a matrix of pixels
- Each pixel has (R,G,B) values
- K means clustering on these values merges colours



Image segmentation

- An image is a matrix of pixels
- Each pixel has (R,G,B) values
- K means clustering on these values merges colours
- With 10 clusters, not much change

10 colors



Image segmentation

- An image is a matrix of pixels
- Each pixel has (R,G,B) values
- K means clustering on these values merges colours
- With 10 clusters, not much change
- Same with 8

8 colors



Image segmentation

- An image is a matrix of pixels
- Each pixel has (R,G,B) values
- K means clustering on these values merges colours
- With 10 clusters, not much change
- Same with 8
- At 6 colours, ladybug red goes

6 colors



Image segmentation

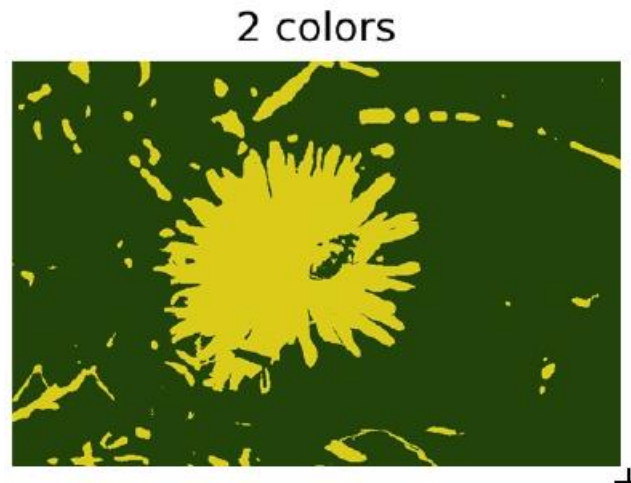
- An image is a matrix of pixels
- Each pixel has (R,G,B) values
- K means clustering on these values merges colours
- With 10 clusters, not much change
- Same with 8
- At 6 colours, ladybug red goes
- 4 colours

4 colors



Image segmentation

- An image is a matrix of pixels
- Each pixel has (R,G,B) values
- K means clustering on these values merges colours
- With 10 clusters, not much change
- Same with 8
- At 6 colours, ladybug red goes
- 4 colours
- Finally 2 colours, flower and rest



Dimensionality reduction

Madhavan Mukund

<https://www.cmi.ac.in/~madhavan>

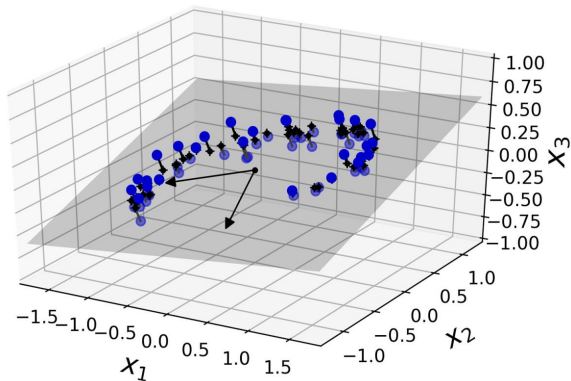
Data Mining and Machine Learning
August–December 2020

Dimensionality reduction

- **Principal Component Analysis** — transform d -dimensional input to k -dimensional input, preserving essential features

Dimensionality reduction

- **Principal Component Analysis** — transform d -dimensional input to k -dimensional input, preserving essential features
- Example: PCA projection of blue points in 3D to black points in 2D

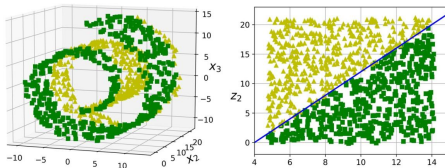


Dimensionality reduction ...

- Unsupervised preprocessing technique — may make later steps easier, like simplifying classification boundaries

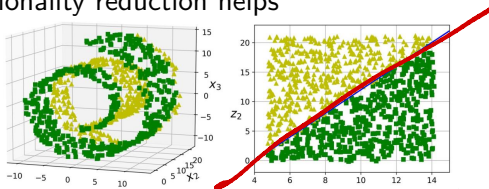
Dimensionality reduction ...

- Unsupervised preprocessing technique — may make later steps easier, like simplifying classification boundaries
- Swiss roll dataset: dimensionality reduction helps

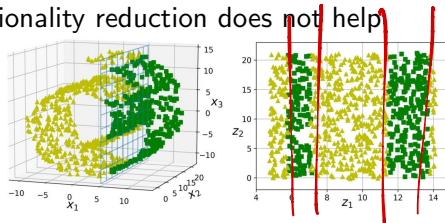


Dimensionality reduction ...

- Unsupervised preprocessing technique — may make later steps easier, like simplifying classification boundaries
- Swiss roll dataset: dimensionality reduction helps



- Swiss roll dataset: dimensionality reduction does not help



Singular Value Decomposition (SVD)

- Input matrix M , dimensions $n \times d$
 - Rows are items, columns are features

$$\begin{matrix} x_1 & x_1^1, x_1^2, \dots, x_1^d \\ x_2 & x_2^1, x_2^2, \dots, x_2^d \\ \vdots & \vdots \\ x_n & x_n^1, x_n^2, \dots, x_n^d \end{matrix}$$

Singular Value Decomposition (SVD)

- Input matrix M , dimensions $n \times d$
 - Rows are items, columns are features

$\rightsquigarrow n \times k$

- Decompose M as UDV^T

- D is a $k \times k$ diagonal matrix, positive real entries

- U is $n \times k$, V is $d \times k$

- Columns of U , V are **orthonormal** — unit vectors, mutually orthogonal

$n \times k$ $\begin{bmatrix} 1 & & \\ & 0 & \\ & & \ddots \end{bmatrix}$ $k \times k$ $k \times d$ V^T V $d \times k$

Singular Value Decomposition (SVD)

- Input matrix M , dimensions $n \times d$
 - Rows are items, columns are features
- Decompose M as UDV^T
 - D is a $k \times k$ diagonal matrix, positive real entries
 - U is $n \times k$, V is $d \times k$
 - Columns of U , V are **orthonormal** — unit vectors, mutually orthogonal
- Interpretation
 - Columns of V correspond to new abstract features

$$U D V^T$$

Singular Value Decomposition (SVD)

- Input matrix M , dimensions $n \times d$

- Rows are items, columns are features

$$n \begin{bmatrix} | & | & | & | & | & | \end{bmatrix}^d$$

$$k \ll d$$

- Decompose M as UDV^T

- D is a $k \times k$ diagonal matrix, positive real entries

- U is $n \times k$, V is $d \times k$

- Columns of U , V are **orthonormal** — unit vectors, mutually orthogonal

$$V^d \begin{bmatrix} | & | & | & | & | \end{bmatrix}^k$$

- Interpretation

- Columns of V correspond to new abstract features

- Rows of U describe decomposition of ~~terms~~^{inputs} across features

$$U^n \begin{bmatrix} | & | & | \end{bmatrix}^k$$

Singular Value Decomposition (SVD)

- Input matrix M , dimensions $n \times d$
 - Rows are items, columns are features
- Decompose M as UDV^T
 - D is a $k \times k$ diagonal matrix, positive real entries
 - U is $n \times k$, V is $d \times k$
 - Columns of U , V are **orthonormal** — unit vectors, mutually orthogonal
- Interpretation
 - Columns of V correspond to new abstract features
 - Rows of U describe decomposition of terms across features
 - For columns u_i of U and v_i of V , $u_i \cdot v_i^T$ is an $n \times d$ matrix, like M

$$M \approx \sum u_i v_i^T$$

Singular Value Decomposition (SVD)

- Input matrix M , dimensions $n \times d$
 - Rows are items, columns are features
- Decompose M as UDV^T
 - D is a $k \times k$ diagonal matrix, positive real entries
 - U is $n \times k$, V is $d \times k$
 - Columns of U , V are **orthonormal** — unit vectors, mutually orthogonal
- Interpretation
 - Columns of V correspond to new abstract features
 - Rows of U describe decomposition of terms across features
 - For columns u_i of U and v_i of V , $u_i \cdot v_i^T$ is an $n \times d$ matrix, like M
 - $u_i \cdot v_i^T$ describes components of rows of M along direction v_i

Singular vectors

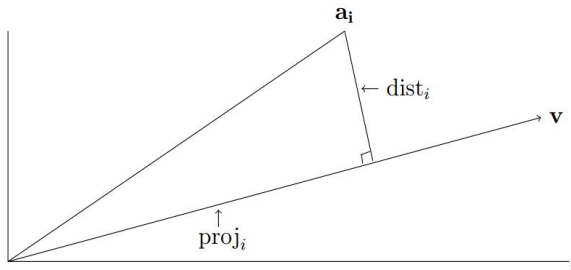
- Unit vectors passing through the origin

Singular vectors

- Unit vectors passing through the origin
- Want to find “best” k singular vectors to represent feature space

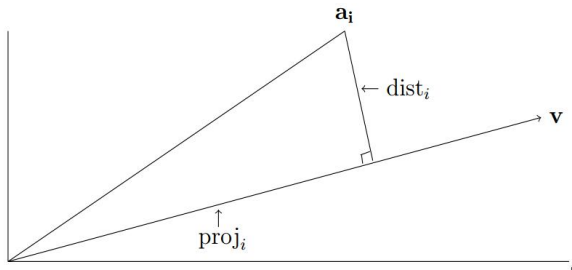
Singular vectors

- Unit vectors passing through the origin
- Want to find “best” k singular vectors to represent feature space
- Suppose we project $\mathbf{a}_i = (a_{i1}, a_{i2}, \dots, a_{id})$ onto $\mathbf{v}$ through origin



Singular vectors

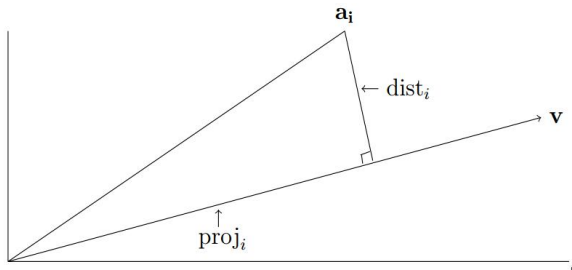
- Unit vectors passing through the origin
- Want to find “best” k singular vectors to represent feature space
- Suppose we project $\mathbf{a}_i = (a_{i1}, a_{i2}, \dots, a_{id})$ onto $\mathbf{v}$ through origin



- Minimizing distance of $\mathbf{a}_i$ from $\mathbf{v}$ is equivalent to maximizing the projection of $\mathbf{a}_i$ onto $\mathbf{v}$

Singular vectors

- Unit vectors passing through the origin
- Want to find “best” k singular vectors to represent feature space
- Suppose we project $\mathbf{a}_i = (a_{i1}, a_{i2}, \dots, a_{id})$ onto $\mathbf{v}$ through origin



- Minimizing distance of $\mathbf{a}_i$ from $\mathbf{v}$ is equivalent to maximizing the projection of $\mathbf{a}_i$ onto $\mathbf{v}$
- Length of the projection is $\mathbf{a}_i \cdot \mathbf{v}$

Singular vectors ...

- Sum of squares of lengths of projections of all rows in M onto $\mathbf{v}$ — $|M\mathbf{v}|^2$

Singular vectors ...

- Sum of squares of lengths of projections of all rows in M onto $\mathbf{v}$ — $|M\mathbf{v}|^2$
- First singular vector — unit vector through origin that maximizes the sum of projections of all rows in M

$$\mathbf{v}_1 = \arg \max_{|\mathbf{v}|=1} |M\mathbf{v}|$$

Singular vectors ...

- Sum of squares of lengths of projections of all rows in M onto $\mathbf{v}$ — $|M\mathbf{v}|^2$
- First singular vector — unit vector through origin that maximizes the sum of projections of all rows in M

$$\mathbf{v}_1 = \arg \max_{|\mathbf{v}|=1} |M\mathbf{v}|$$

- Second singular vector — unit vector through origin, perpendicular to $\mathbf{v}_1$, that maximizes the sum of projections of all rows in M

$$\mathbf{v}_2 = \arg \max_{\mathbf{v} \perp \mathbf{v}_1; |\mathbf{v}|=1} |M\mathbf{v}|$$

Singular vectors ...

- Sum of squares of lengths of projections of all rows in M onto $\mathbf{v}$ — $|M\mathbf{v}|^2$
- First singular vector — unit vector through origin that maximizes the sum of projections of all rows in M

$$\mathbf{v}_1 = \arg \max_{|\mathbf{v}|=1} |M\mathbf{v}|$$

$$|M\mathbf{v}_1| > |M\mathbf{v}_2| > |M\mathbf{v}_3|$$

- Second singular vector — unit vector through origin, perpendicular to $\mathbf{v}_1$, that maximizes the sum of projections of all rows in M

$$\mathbf{v}_2 = \arg \max_{\mathbf{v} \perp \mathbf{v}_1; |\mathbf{v}|=1} |M\mathbf{v}|$$

- Third singular vector — unit vector through origin, perpendicular to $\mathbf{v}_1$, $\mathbf{v}_2$, that maximizes the sum of projections of all rows in M

$$\mathbf{v}_3 = \arg \max_{\mathbf{v} \perp \mathbf{v}_1, \mathbf{v}_2; |\mathbf{v}|=1} |M\mathbf{v}|$$

Singular vectors ...

- With each singular vector $\mathbf{v}_j$, associated singular value is $\sigma_j = |M\mathbf{v}_j|$

Singular vectors ...

- With each singular vector $\mathbf{v}_j$, associated singular value is $\sigma_j = |M\mathbf{v}_j|$
- Repeat r times till $\max_{\mathbf{v} \perp \mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_r; |\mathbf{v}|=1} |M\mathbf{v}| = 0$
 - r turns out to be the rank of M
 - Vectors $\{\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_r\}$ are orthonormal right singular vectors

Singular vectors ...

- With each singular vector $\mathbf{v}_j$, associated singular value is $\sigma_j = |M\mathbf{v}_j|$
- Repeat r times till $\max_{\mathbf{v} \perp \mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_r; |\mathbf{v}|=1} |M\mathbf{v}| = 0$
 - r turns out to be the rank of M
 - Vectors $\{\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_r\}$ are orthonormal right singular vectors
- Our greedy strategy provably produces “best-fit” dimension r subspace for M
 - Dimension r subspace that maximizes content of M projected onto it

Singular vectors ...

- With each singular vector $\mathbf{v}_j$, associated singular value is $\sigma_j = |M\mathbf{v}_j|$
- Repeat r times till $\max_{\mathbf{v} \perp \mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_r; |\mathbf{v}|=1} |M\mathbf{v}| = 0$
 - r turns out to be the rank of M
 - Vectors $\{\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_r\}$ are orthonormal **right singular vectors**
- Our greedy strategy provably produces “best-fit” dimension r subspace for M
 - Dimension r subspace that maximizes content of M projected onto it
- Corresponding **left singular vectors** are given by $\mathbf{u}_i = \frac{1}{\sigma_i} M\mathbf{v}_i$
- Can show that $\{\mathbf{u}_1, \mathbf{u}_2, \dots, \mathbf{u}_r\}$ are also orthonormal

Singular Value Decomposition

- M , dimension $n \times d$, of rank r uniquely decomposes as $M = UDV^T$
 - $V = [\mathbf{v}_1 \ \mathbf{v}_2 \ \cdots \ \mathbf{v}_r]$ are the right singular vectors
 - D is a diagonal matrix with $D[i, i] = \sigma_i$, the singular values
 - $U = [\mathbf{u}_1 \ \mathbf{u}_2 \ \cdots \ \mathbf{u}_r]$ are the left singular vectors

The diagram illustrates the SVD equation $M = UDV^T$ with matrix dimensions:

- M is $n \times d$
- U is $n \times r$
- D is $r \times r$ (indicated by a green diagonal line)
- V^T is $r \times d$

A handwritten green formula below the D matrix states: $\sigma_i = |Mv_i|$

Rank- k approximation

- M has rank r , SVD gives rank r decomposition

Rank- k approximation

- M has rank r , SVD gives rank r decomposition
- Singular values are non-increasing — $\sigma_1 \geq \sigma_2 \geq \dots \geq \sigma_r$

Rank- k approximation

- M has rank r , SVD gives rank r decomposition
- Singular values are non-increasing — $\sigma_1 \geq \sigma_2 \geq \dots \geq \sigma_r$
- Suppose we retain only k largest ones
- We have
 - Matrix of first k right singular vectors $V_k = [\mathbf{v}_1 \ \mathbf{v}_2 \ \dots \ \mathbf{v}_k]$,
 - Corresponding singular values $\sigma_1, \sigma_2, \dots, \sigma_k$
 - Matrix of k left singular vectors $U_k = [\mathbf{u}_1 \ \mathbf{u}_2 \ \dots \ \mathbf{u}_k]$

Rank- k approximation

- M has rank r , SVD gives rank r decomposition
- Singular values are non-increasing — $\sigma_1 \geq \sigma_2 \geq \dots \geq \sigma_r$
- Suppose we retain only k largest ones
- We have
 - Matrix of first k right singular vectors $V_k = [\mathbf{v}_1 \ \mathbf{v}_2 \ \dots \ \mathbf{v}_k]$,
 - Corresponding singular values $\sigma_1, \sigma_2, \dots, \sigma_k$
 - Matrix of k left singular vectors $U_k = [\mathbf{u}_1 \ \mathbf{u}_2 \ \dots \ \mathbf{u}_k]$
- Let D_k be the $k \times k$ diagonal matrix with entries $\sigma_1, \sigma_2, \dots, \sigma_k$
- Then $U_k D_k V_k^T$ is the best fit rank- k approximation of M

Rank- k approximation

- M has rank r , SVD gives rank r decomposition
- Singular values are non-increasing — $\sigma_1 \geq \sigma_2 \geq \dots \geq \sigma_r$
- Suppose we retain only k largest ones
- We have
 - Matrix of first k right singular vectors $V_k = [\mathbf{v}_1 \ \mathbf{v}_2 \ \dots \ \mathbf{v}_k]$,
 - Corresponding singular values $\sigma_1, \sigma_2, \dots, \sigma_k$
 - Matrix of k left singular vectors $U_k = [\mathbf{u}_1 \ \mathbf{u}_2 \ \dots \ \mathbf{u}_k]$
- Let D_k be the $k \times k$ diagonal matrix with entries $\sigma_1, \sigma_2, \dots, \sigma_k$
- Then $U_k D_k V_k^T$ is the best fit rank- k approximation of M
- In other words, by truncating the SVD, we can focus on k most significant features implicit in M

Summary

- Singular Value Decomposition (SVD) finds best fit k -dimensional subspace for any matrix M
- Principal Component Analysis uses SVD for dimensionality reduction
- Unsupervised technique — often helps simplify the problem, but may not
- SVD/PCA can only compress features that have a linear relationship
- More general techniques based on neural networks — **autoencoders**