



THE UNIVERSITY *of* EDINBURGH
informatics

Applied Machine Learning (AML)

Clustering

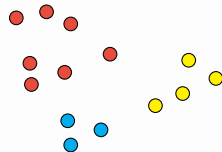
Oisin Mac Aodha • Siddharth N.

Outline

- What is clustering and why is it useful?
- What kinds are there and how are they characterised?
- Explore
 - K-Means
 - Hierarchical Clustering
- How do we evaluate clustering?

Clustering

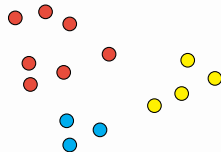
- Discover the underlying structure of data



Clusters in 2D

Clustering

- Discover the underlying structure of data
- What sub-groups exist in the data
 - # clusters, size, ...
 - common properties within sub-group
 - potential for further clustering



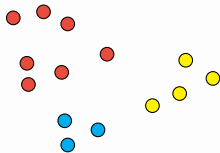
Clusters in 2D

Clustering

- Discover the underlying structure of data
- What sub-groups exist in the data
 - # clusters, size, ...
 - common properties within sub-group
 - potential for further clustering

Applications

- discover classes / structure in an unsupervised manner
 - clustering images of handwritten digits (K=10)
 - finding phylogenetic trees using DNA
- dimensionality reduction: clusters $\leftrightarrow$ “latent factors”
 - use cluster id as representation
 - assume relevant characteristics reflected in cluster membership



Clusters in 2D

Features of Clustering Algorithms

Hard vs. Soft

Hard: objects belong to a single cluster

Soft: objects have soft assignments—distribution over clusters

Features of Clustering Algorithms

Hard vs. Soft

Hard: objects belong to a single cluster

Soft: objects have soft assignments—distribution over clusters

Flat vs. Hierarchical

Flat: single group of clusters

Hierarchical: clusters at different levels

Features of Clustering Algorithms

Hard vs. Soft

Hard: objects belong to a single cluster

Soft: objects have soft assignments—distribution over clusters

Flat vs. Hierarchical

Flat: single group of clusters

Hierarchical: clusters at different levels

Monothetic vs. Polythetic

Monothetic: clustered based on common feature (e.g. hair colour)

Polythetic: clustered based on distance measure(s) over features

Features of Clustering Algorithms

Hard vs. Soft

Hard: objects belong to a single cluster

Soft: objects have soft assignments—distribution over clusters

Flat vs. Hierarchical

Flat: single group of clusters

Hierarchical: clusters at different levels

Monothetic vs. Polythetic

Monothetic: clustered based on common feature (e.g. hair colour)

Polythetic: clustered based on distance measure(s) over features

K-Means



K-Means

Characteristics

Hard: a point belongs to just one cluster

Flat: single level of clustering

Polythetic: distance-based similarity within clusters

K-Means

Characteristics

Hard: a point belongs to just one cluster

Flat: single level of clustering

Polythetic: distance-based similarity within clusters

Idea

Ensure points closest to some special point end up in the same cluster

- Top-down approach
- Produces a partition of the data
- Requires defining a distance metric over points

K-Means Algorithm

Require: $\mathcal{D}, K, \{x_1, \dots, x_N\}$ ▷ # clusters, points

1: $\{c_1, \dots, c_K\} \leftarrow$ **random initialisation** ▷ centroids of clusters

2: **repeat**

3: **for** $x_n \in \{x_1, \dots, x_N\}$ **do**

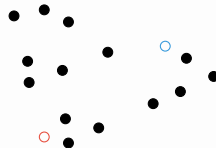
4: $c_k^* = \arg \min_{c_k} \mathcal{D}(x_n, c_k)$ ▷ find nearest centroid id

5: $c_k^* \leftarrow x_n$ ▷ assign point to cluster

6: **for** $c_k \in \{c_1, \dots, c_K\}$ **do**

7: $c_k = \frac{1}{N_k} \sum_{x_n \rightarrow c_k} x_n$ ▷ update cluster centroids

8: **until** cluster assignments do not change



K-Means Algorithm

Require: $\mathcal{D}, K, \{x_1, \dots, x_N\}$ ▷ # clusters, points

1: $\{c_1, \dots, c_K\} \leftarrow$ random initialisation ▷ centroids of clusters

2: **repeat**

3: **for** $x_n \in \{x_1, \dots, x_N\}$ **do**

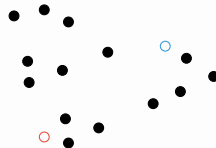
4: $c_k^* = \arg \min_{c_k} \mathcal{D}(x_n, c_k)$ ▷ find nearest centroid id

5: $c_k^* \leftarrow x_n$ ▷ assign point to cluster

6: **for** $c_k \in \{c_1, \dots, c_K\}$ **do**

7: $c_k = \frac{1}{N_k} \sum_{x_n \rightarrow c_k} x_n$ ▷ update cluster centroids

8: **until** cluster assignments do not change



K-Means Algorithm

Require: $\mathcal{D}, K, \{x_1, \dots, x_N\}$ ▷ # clusters, points

1: $\{c_1, \dots, c_K\} \leftarrow$ random initialisation ▷ centroids of clusters

2: **repeat**

3: **for** $x_n \in \{x_1, \dots, x_N\}$ **do**

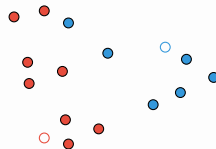
4: $c_k^* = \arg \min_{c_k} \mathcal{D}(x_n, c_k)$ ▷ find nearest centroid id

5: $c_k^* \leftarrow x_n$ ▷ assign point to cluster

6: **for** $c_k \in \{c_1, \dots, c_K\}$ **do**

7: $c_k = \frac{1}{N_k} \sum_{x_n \rightarrow c_k} x_n$ ▷ update cluster centroids

8: **until** cluster assignments do not change



K-Means Algorithm

Require: $\mathcal{D}, K, \{x_1, \dots, x_N\}$ ▷ # clusters, points

1: $\{c_1, \dots, c_K\} \leftarrow$ random initialisation ▷ centroids of clusters

2: **repeat**

3: **for** $x_n \in \{x_1, \dots, x_N\}$ **do**

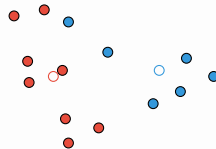
4: $c_k^* = \arg \min_{c_k} \mathcal{D}(x_n, c_k)$ ▷ find nearest centroid id

5: $c_k^* \leftarrow x_n$ ▷ assign point to cluster

6: **for** $c_k \in \{c_1, \dots, c_K\}$ **do**

7: $c_k = \frac{1}{N_k} \sum_{x_n \rightarrow c_k} x_n$ ▷ update cluster centroids

8: **until** cluster assignments do not change



K-Means Algorithm

Require: $\mathcal{D}, K, \{x_1, \dots, x_N\}$ ▷ # clusters, points

1: $\{c_1, \dots, c_K\} \leftarrow$ random initialisation ▷ centroids of clusters

2: **repeat**

3: **for** $x_n \in \{x_1, \dots, x_N\}$ **do**

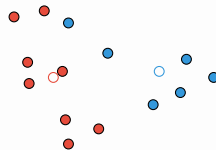
4: $c_k^* = \arg \min_{c_k} \mathcal{D}(x_n, c_k)$ ▷ find nearest centroid id

5: $c_k^* \leftarrow x_n$ ▷ assign point to cluster

6: **for** $c_k \in \{c_1, \dots, c_K\}$ **do**

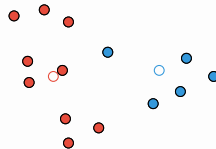
7: $c_k = \frac{1}{N_k} \sum_{x_n \rightarrow c_k} x_n$ ▷ update cluster centroids

8: **until** cluster assignments do not change



K-Means Algorithm

- Require:** $\mathcal{D}, K, \{x_1, \dots, x_N\}$ ▷ # clusters, points
- 1: $\{c_1, \dots, c_K\} \leftarrow$ random initialisation ▷ centroids of clusters
 - 2: **repeat**
 - 3: **for** $x_n \in \{x_1, \dots, x_N\}$ **do**
 - 4: $c_k^* = \arg \min_{c_k} \mathcal{D}(x_n, c_k)$ ▷ find nearest centroid id
 - 5: $c_k^* \leftarrow x_n$ ▷ assign point to cluster
 - 6: **for** $c_k \in \{c_1, \dots, c_K\}$ **do**
 - 7: $c_k = \frac{1}{N_k} \sum_{x_n \rightarrow c_k} x_n$ ▷ update cluster centroids
 - 8: **until** cluster assignments do not change



K-Means Algorithm

Require: $\mathcal{D}, K, \{x_1, \dots, x_N\}$ ▷ # clusters, points

1: $\{c_1, \dots, c_K\} \leftarrow$ random initialisation ▷ centroids of clusters

2: **repeat**

3: **for** $x_n \in \{x_1, \dots, x_N\}$ **do**

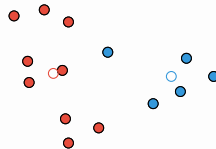
4: $c_k^* = \arg \min_{c_k} \mathcal{D}(x_n, c_k)$ ▷ find nearest centroid id

5: $c_k^* \leftarrow x_n$ ▷ assign point to cluster

6: **for** $c_k \in \{c_1, \dots, c_K\}$ **do**

7: $c_k = \frac{1}{N_k} \sum_{x_n \rightarrow c_k} x_n$ ▷ update cluster centroids

8: **until** cluster assignments do not change



K-Means Properties

- Minimises aggregate intra-cluster distance: $V = \sum_k \sum_{\mathbf{x}_n \rightarrow \mathbf{c}_k} \mathcal{D}(\mathbf{x}_n, \mathbf{c}_k)$
 - if $\mathcal{D}(\mathbf{x}_n, \mathbf{c}_k) = \|\mathbf{x}_n - \mathbf{c}_k\|_2^2$, i.e., Euclidean distance, then V is proportional to variance

K-Means Properties

- Minimises aggregate intra-cluster distance: $V = \sum_k \sum_{\mathbf{x}_n \rightarrow c_k} \mathcal{D}(\mathbf{x}_n, c_k)$
 - if $\mathcal{D}(\mathbf{x}_n, c_k) = \|\mathbf{x}_n - c_k\|_2^2$, i.e., Euclidean distance, then V is proportional to variance
- Converges to *local* minimum
 - *different* initialisations lead to *different* clustering results
 - repeat several random initialisations and pick one with smallest aggregate distance



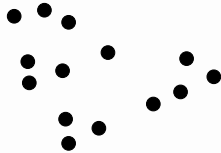
K-Means Properties

- Minimises aggregate intra-cluster distance: $V = \sum_k \sum_{\mathbf{x}_n \rightarrow \mathbf{c}_k} \mathcal{D}(\mathbf{x}_n, \mathbf{c}_k)$
 - if $\mathcal{D}(\mathbf{x}_n, \mathbf{c}_k) = \|\mathbf{x}_n - \mathbf{c}_k\|_2^2$, i.e., Euclidean distance, then V is proportional to variance
- Converges to *local* minimum
 - *different* initialisations lead to *different* clustering results
 - repeat several random initialisations and pick one with smallest aggregate distance



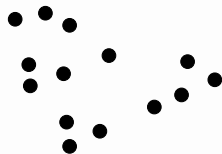
- ‘Adjacent’ points can end up in different clusters

Estimating Number of Clusters

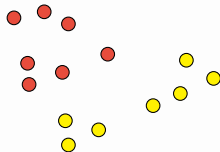


data

Estimating Number of Clusters

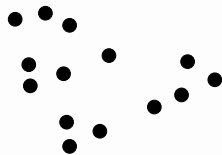


data

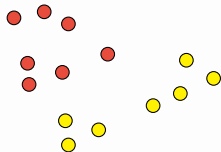


2 clusters

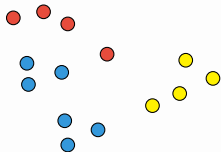
Estimating Number of Clusters



data

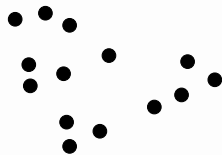


2 clusters

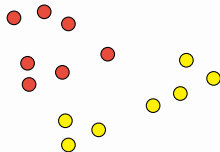


3 clusters

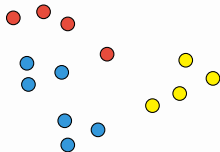
Estimating Number of Clusters



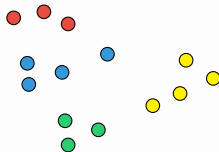
data



2 clusters

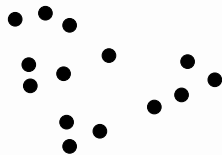


3 clusters

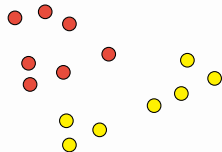


4 clusters

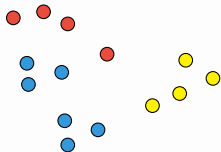
Estimating Number of Clusters



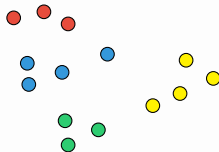
data



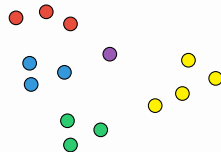
2 clusters



3 clusters



4 clusters



5 clusters

Estimating Number of Clusters

How many clusters does your data have?

- Get (K) from class labels (e.g. digits 0...9)

Estimating Number of Clusters

How many clusters does your data have?

- Get (K) from class labels (e.g. digits 0...9)
- Find an “appropriate” K : optimise for V

Estimating Number of Clusters

How many clusters does your data have?

- Get (K) from class labels (e.g. digits 0...9)
- Find an “appropriate” K : optimise for V
 - Run K-Means for $K = 1, 2, \dots$; choose K with smallest V

Estimating Number of Clusters

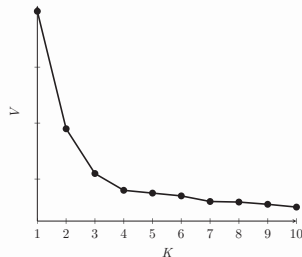
How many clusters does your data have?

- Get (K) from class labels (e.g. digits 0...9)
- Find an “appropriate” K : optimise for V
 - Run K-Means for $K = 1, 2, \dots$; choose K with smallest V
 - **Issue:** What is V when $K = N$?
 - choose best K on *validation* data

Estimating Number of Clusters

How many clusters does your data have?

- Get (K) from class labels (e.g. digits 0...9)
- Find an “appropriate” K : optimise for V
 - Run K-Means for $K = 1, 2, \dots$; choose K with smallest V
 - **Issue:** What is V when $K = N$?
 - choose best K on *validation* data
 - Choose visually from a **elbow** plot
 - point that maximises the 2nd derivative of V



K-Means: Example

Colour Quantisation

- Original Image: 96,615 colours

Original



Figures: Scikit Learn: Colour Quantization using K-Means

K-Means: Example

Colour Quantisation

- Original Image: 96,615 colours
- Quantised Image: 64 colours (K-Means)
 - Replace pixel value x_i with cluster centroid c_k value

$$\mathbf{x}_i \in \mathbb{R}^3$$

(pixel values in RGB)

$$\mathcal{D}(\mathbf{x}_i, \mathbf{x}_j) = \|\mathbf{x}_i - \mathbf{x}_j\|_2^2$$

$$K = 64$$

Original



K-Means Quantised



Figures: Scikit Learn: Colour Quantization using K-Means

K-Means: Example

Colour Quantisation

- Original Image: 96,615 colours
- Quantised Image: 64 colours (K-Means)
 - Replace pixel value x_i with cluster centroid c_k value
- Quantised Image: 64 colours (Random)
 - Select random set of K pixels as “centroids”
 - Replace pixel value x_i with nearest “centroid” value

$$\mathbf{x}_i \in \mathbb{R}^3 \quad (\text{pixel values in RGB})$$

$$\mathcal{D}(\mathbf{x}_i, \mathbf{x}_j) = \|\mathbf{x}_i - \mathbf{x}_j\|_2^2$$

$$K = 64$$

Original



Random Quantised



Figures: Scikit Learn: Colour Quantization using K-Means

K-Means: Example

Clustering Handwritten Digits

- High-dimensional data



$$\mathbf{x} \in \mathbb{R}^{784}$$

K-Means: Example

Clustering Handwritten Digits

- High-dimensional data
- Dimensionality reduction (e.g. PCA)



$$\mathbf{x} \in \mathbb{R}^{784}$$

$$\mathbf{e} \in \mathbb{R}^2 \quad (\text{PCA})$$

K-Means: Example

Clustering Handwritten Digits

- High-dimensional data
- Dimensionality reduction (e.g. PCA)
- K-Means on embeddings

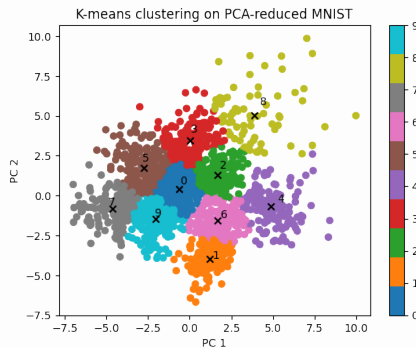
$$\mathbf{x} \in \mathbb{R}^{784}$$

$$\mathbf{e} \in \mathbb{R}^2$$

(PCA)

$$\mathcal{D}(\mathbf{x}_i, \mathbf{x}_j) = \|\mathbf{e}_i - \mathbf{e}_j\|_2^2$$

$$K = 10$$



Hierarchical Clustering

Hierarchical Clustering

Choosing number of clusters

- Depends a lot on *granularity*

Hierarchical Clustering

Choosing number of clusters

- Depends a lot on *granularity*
 - data (e.g. satellite maps—how much does 1 pixel cover?)

Hierarchical Clustering

Choosing number of clusters

- Depends a lot on *granularity*
 - data (e.g. satellite maps—how much does 1 pixel cover?)
 - context—what do we care about? High vs. low level?

Hierarchical Clustering

Choosing number of clusters

- Depends a lot on *granularity*
 - data (e.g. satellite maps—how much does 1 pixel cover?)
 - context—what do we care about? High vs. low level?
- No magical algorithm to give you *correct K*

Hierarchical Clustering

Choosing number of clusters

- Depends a lot on *granularity*
 - data (e.g. satellite maps—how much does 1 pixel cover?)
 - context—what do we care about? High vs. low level?
- No magical algorithm to give you *correct K*

Hierarchical Clustering

Choosing number of clusters

- Depends a lot on *granularity*
 - data (e.g. satellite maps—how much does 1 pixel cover?)
 - context—what do we care about? High vs. low level?
- No magical algorithm to give you *correct* K

Find a hierarchy of structure

- **Upper levels:** coarse groups (e.g. collection of objects; bedroom, kitchen, etc.)
- **Lower levels:** fine-grained (e.g. object parts; chair leg, table top, etc.)
- **Strategies**
 - Top-Down: start with everything in one cluster, then split recursively
 - Bottom-up: start with each item separately, then merge recursively

Hierarchical K-Means

- Top-Down approach

Hierarchical K-Means

- Top-Down approach
 - perform K-Means on data

Hierarchical K-Means

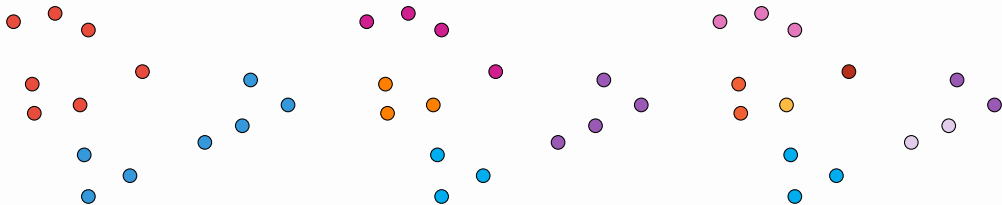
- Top-Down approach
 - perform K-Means on data
 - for each resulting cluster c_i , run K-Means within c_i

Hierarchical K-Means

- Top-Down approach
 - perform K-Means on data
 - for each resulting cluster c_i , run K-Means within c_i
- **Fast:** recursive calls on successively smaller datasets

Hierarchical K-Means

- Top-Down approach
 - perform K-Means on data
 - for each resulting cluster c_i , run K-Means within c_i
- **Fast:** recursive calls on successively smaller datasets
- **Greedy:** once cluster has been determined at top level; cannot change



Agglomerative Clustering

Characteristics

Hard: a point belongs to just one cluster

Hierarchical: multiple levels of clustering

Polythetic: distance-based similarity within clusters

Agglomerative Clustering

Characteristics

Hard: a point belongs to just one cluster

Hierarchical: multiple levels of clustering

Polythetic: distance-based similarity within clusters

Idea

Ensure “nearby” points end up in the same cluster

- Bottom-up approach
- Generates a dendrogram: hierarchical tree of clusters
- Requires defining a distance metric over *clusters*

Agglomerative Clustering: Sketch

$\mathcal{D}(x_l, x_m)$ —distance between *points*

$\mathcal{G}_{\mathcal{D}}(c_i, c_j)$ —distance between *clusters* of points

Require: $\mathcal{G}_{\mathcal{D}}, \{x_1, \dots, x_N\}$ ▷ points

1: $C = \{c_1, \dots, c_N\} = \{\{x_1\}, \dots, \{x_N\}\}$ ▷ initial clusters

2: **repeat**

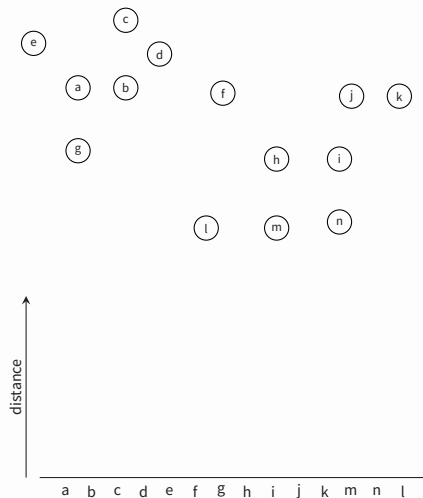
3: $c_i^*, c_j^* = \arg \min_{c_i, c_j} \mathcal{G}_{\mathcal{D}}(c_i, c_j)$ ▷ find closest pair

4: $c_{i,j} \leftarrow c_i^*, c_j^*$ ▷ merge into new cluster

5: $C = C \setminus \{c_i^*, c_j^*\}$ ▷ remove pair of clusters

6: $C = C \cup \{c_{i,j}\}$ ▷ add merged cluster

7: **until** only one cluster remaining



Agglomerative Clustering: Sketch

$\mathcal{D}(x_l, x_m)$ —distance between *points*

$\mathcal{G}_{\mathcal{D}}(c_i, c_j)$ —distance between *clusters* of points

Require: $\mathcal{G}_{\mathcal{D}}, \{x_1, \dots, x_N\}$ ▷ points

1: $C = \{c_1, \dots, c_N\} = \{\{x_1\}, \dots, \{x_N\}\}$ ▷ initial clusters

2: **repeat**

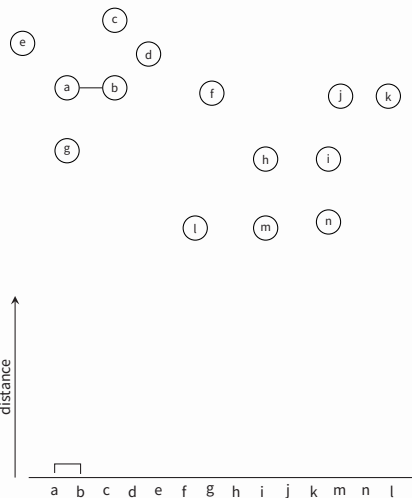
3: $c_i^*, c_j^* = \arg \min_{c_i, c_j} \mathcal{G}_{\mathcal{D}}(c_i, c_j)$ ▷ find closest pair

4: $c_{i,j} \leftarrow c_i^*, c_j^*$ ▷ merge into new cluster

5: $C = C \setminus \{c_i^*, c_j^*\}$ ▷ remove pair of clusters

6: $C = C \cup \{c_{i,j}\}$ ▷ add merged cluster

7: **until** only one cluster remaining



Agglomerative Clustering: Sketch

$\mathcal{D}(x_l, x_m)$ —distance between *points*

$\mathcal{G}_{\mathcal{D}}(c_i, c_j)$ —distance between *clusters* of points

Require: $\mathcal{G}_{\mathcal{D}}, \{x_1, \dots, x_N\}$ ▷ points

1: $C = \{c_1, \dots, c_N\} = \{\{x_1\}, \dots, \{x_N\}\}$ ▷ initial clusters

2: **repeat**

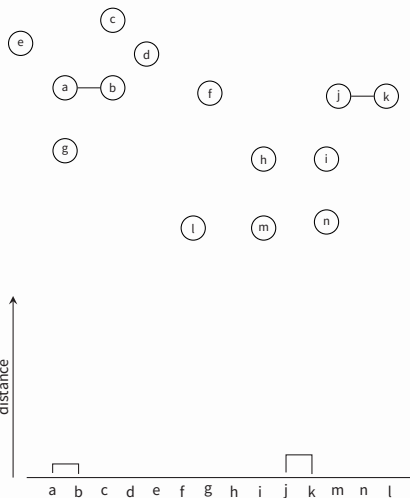
3: $c_i^*, c_j^* = \arg \min_{c_i, c_j} \mathcal{G}_{\mathcal{D}}(c_i, c_j)$ ▷ find closest pair

4: $c_{i,j} \leftarrow c_i^*, c_j^*$ ▷ merge into new cluster

5: $C = C \setminus \{c_i^*, c_j^*\}$ ▷ remove pair of clusters

6: $C = C \cup \{c_{i,j}\}$ ▷ add merged cluster

7: **until** only one cluster remaining



Agglomerative Clustering: Sketch

$\mathcal{D}(x_l, x_m)$ —distance between *points*

$\mathcal{G}_{\mathcal{D}}(c_i, c_j)$ —distance between *clusters* of points

Require: $\mathcal{G}_{\mathcal{D}}, \{x_1, \dots, x_N\}$

‣ points

1: $C = \{c_1, \dots, c_N\} = \{\{x_1\}, \dots, \{x_N\}\}$

‣ initial clusters

2: **repeat**

3: $c_i^*, c_j^* = \arg \min_{c_i, c_j} \mathcal{G}_{\mathcal{D}}(c_i, c_j)$

‣ find closest pair

4: $c_{i,j} \leftarrow c_i^*, c_j^*$

‣ merge into new cluster

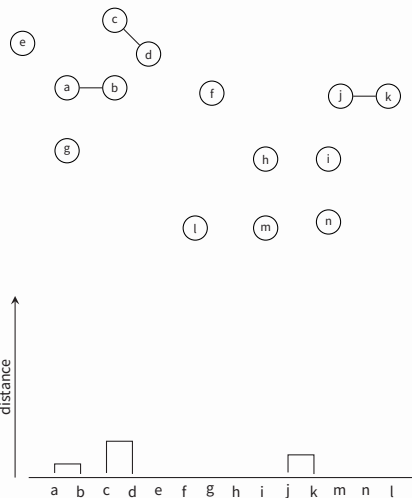
5: $C = C \setminus \{c_i^*, c_j^*\}$

‣ remove pair of clusters

6: $C = C \cup \{c_{i,j}\}$

‣ add merged cluster

7: **until** only one cluster remaining



Agglomerative Clustering: Sketch

$\mathcal{D}(x_l, x_m)$ —distance between *points*

$\mathcal{G}_{\mathcal{D}}(c_i, c_j)$ —distance between *clusters* of points

Require: $\mathcal{G}_{\mathcal{D}}, \{x_1, \dots, x_N\}$

‣ points

1: $C = \{c_1, \dots, c_N\} = \{\{x_1\}, \dots, \{x_N\}\}$

‣ initial clusters

2: **repeat**

3: $c_i^*, c_j^* = \arg \min_{c_i, c_j} \mathcal{G}_{\mathcal{D}}(c_i, c_j)$

‣ find closest pair

4: $c_{i,j} \leftarrow c_i^*, c_j^*$

‣ merge into new cluster

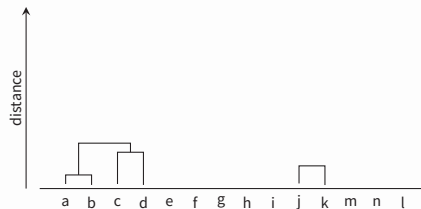
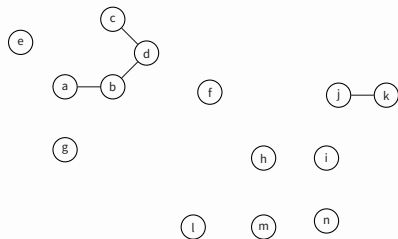
5: $C = C \setminus \{c_i^*, c_j^*\}$

‣ remove pair of clusters

6: $C = C \cup \{c_{i,j}\}$

‣ add merged cluster

7: **until** only one cluster remaining



Agglomerative Clustering: Sketch

$\mathcal{D}(x_l, x_m)$ —distance between *points*

$\mathcal{G}_{\mathcal{D}}(c_i, c_j)$ —distance between *clusters* of points

Require: $\mathcal{G}_{\mathcal{D}}, \{x_1, \dots, x_N\}$

‣ points

1: $C = \{c_1, \dots, c_N\} = \{\{x_1\}, \dots, \{x_N\}\}$

‣ initial clusters

2: **repeat**

3: $c_i^*, c_j^* = \arg \min_{c_i, c_j} \mathcal{G}_{\mathcal{D}}(c_i, c_j)$

‣ find closest pair

4: $c_{i,j} \leftarrow c_i^*, c_j^*$

‣ merge into new cluster

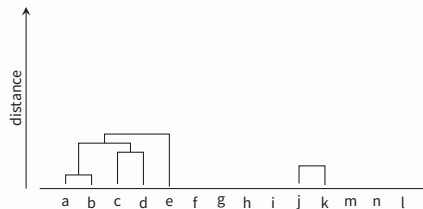
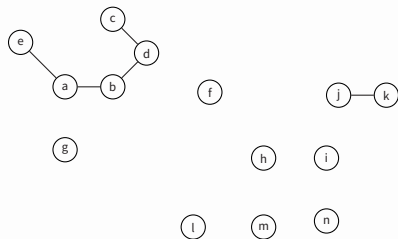
5: $C = C \setminus \{c_i^*, c_j^*\}$

‣ remove pair of clusters

6: $C = C \cup \{c_{i,j}\}$

‣ add merged cluster

7: **until** only one cluster remaining



Agglomerative Clustering: Sketch

$\mathcal{D}(x_l, x_m)$ —distance between *points*

$\mathcal{G}_{\mathcal{D}}(c_i, c_j)$ —distance between *clusters* of points

Require: $\mathcal{G}_{\mathcal{D}}, \{x_1, \dots, x_N\}$

‣ points

1: $C = \{c_1, \dots, c_N\} = \{\{x_1\}, \dots, \{x_N\}\}$

‣ initial clusters

2: **repeat**

3: $c_i^*, c_j^* = \arg \min_{c_i, c_j} \mathcal{G}_{\mathcal{D}}(c_i, c_j)$

‣ find closest pair

4: $c_{i,j} \leftarrow c_i^*, c_j^*$

‣ merge into new cluster

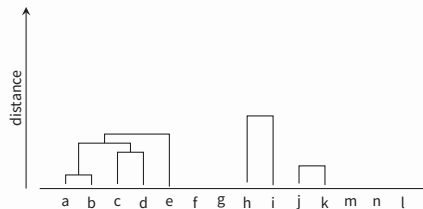
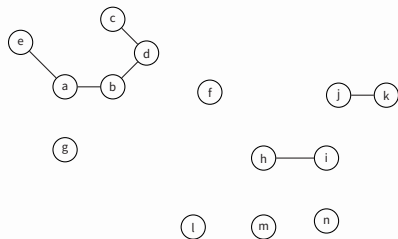
5: $C = C \setminus \{c_i^*, c_j^*\}$

‣ remove pair of clusters

6: $C = C \cup \{c_{i,j}\}$

‣ add merged cluster

7: **until** only one cluster remaining



Agglomerative Clustering: Sketch

$\mathcal{D}(x_l, x_m)$ —distance between *points*

$\mathcal{G}_{\mathcal{D}}(c_i, c_j)$ —distance between *clusters* of points

Require: $\mathcal{G}_{\mathcal{D}}, \{x_1, \dots, x_N\}$

‣ points

1: $C = \{c_1, \dots, c_N\} = \{\{x_1\}, \dots, \{x_N\}\}$

‣ initial clusters

2: **repeat**

3: $c_i^*, c_j^* = \arg \min_{c_i, c_j} \mathcal{G}_{\mathcal{D}}(c_i, c_j)$

‣ find closest pair

4: $c_{i,j} \leftarrow c_i^*, c_j^*$

‣ merge into new cluster

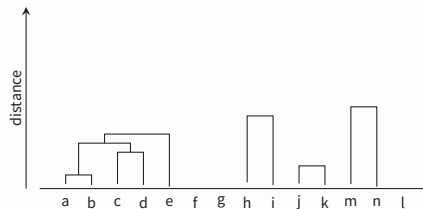
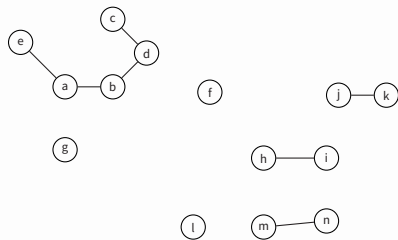
5: $C = C \setminus \{c_i^*, c_j^*\}$

‣ remove pair of clusters

6: $C = C \cup \{c_{i,j}\}$

‣ add merged cluster

7: **until** only one cluster remaining



Agglomerative Clustering: Sketch

$\mathcal{D}(x_l, x_m)$ —distance between *points*

$\mathcal{G}_{\mathcal{D}}(c_i, c_j)$ —distance between *clusters* of points

Require: $\mathcal{G}_{\mathcal{D}}, \{x_1, \dots, x_N\}$

‣ points

1: $C = \{c_1, \dots, c_N\} = \{\{x_1\}, \dots, \{x_N\}\}$

‣ initial clusters

2: **repeat**

3: $c_i^*, c_j^* = \arg \min_{c_i, c_j} \mathcal{G}_{\mathcal{D}}(c_i, c_j)$

‣ find closest pair

4: $c_{i,j} \leftarrow c_i^*, c_j^*$

‣ merge into new cluster

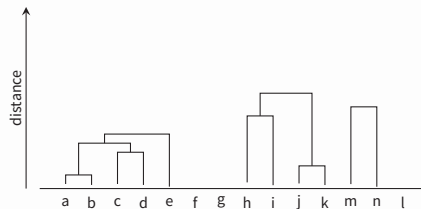
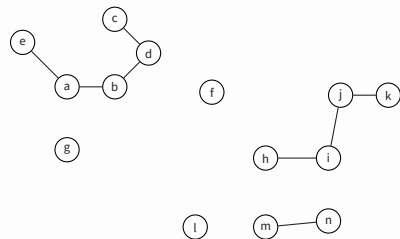
5: $C = C \setminus \{c_i^*, c_j^*\}$

‣ remove pair of clusters

6: $C = C \cup \{c_{i,j}\}$

‣ add merged cluster

7: **until** only one cluster remaining



Agglomerative Clustering: Sketch

$\mathcal{D}(x_l, x_m)$ —distance between *points*

$\mathcal{G}_{\mathcal{D}}(c_i, c_j)$ —distance between *clusters* of points

Require: $\mathcal{G}_{\mathcal{D}}, \{x_1, \dots, x_N\}$

‣ points

1: $C = \{c_1, \dots, c_N\} = \{\{x_1\}, \dots, \{x_N\}\}$

‣ initial clusters

2: **repeat**

3: $c_i^*, c_j^* = \arg \min_{c_i, c_j} \mathcal{G}_{\mathcal{D}}(c_i, c_j)$

‣ find closest pair

4: $c_{i,j} \leftarrow c_i^*, c_j^*$

‣ merge into new cluster

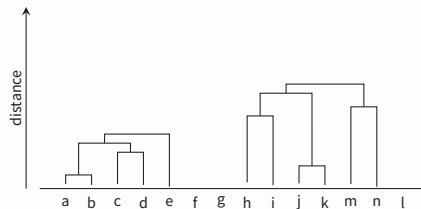
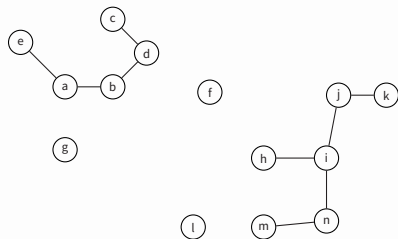
5: $C = C \setminus \{c_i^*, c_j^*\}$

‣ remove pair of clusters

6: $C = C \cup \{c_{i,j}\}$

‣ add merged cluster

7: **until** only one cluster remaining



Agglomerative Clustering: Sketch

$\mathcal{D}(x_l, x_m)$ —distance between *points*

$\mathcal{G}_{\mathcal{D}}(c_i, c_j)$ —distance between *clusters* of points

Require: $\mathcal{G}_{\mathcal{D}}, \{x_1, \dots, x_N\}$

‣ points

1: $C = \{c_1, \dots, c_N\} = \{\{x_1\}, \dots, \{x_N\}\}$

‣ initial clusters

2: **repeat**

3: $c_i^*, c_j^* = \arg \min_{c_i, c_j} \mathcal{G}_{\mathcal{D}}(c_i, c_j)$

‣ find closest pair

4: $c_{i,j} \leftarrow c_i^*, c_j^*$

‣ merge into new cluster

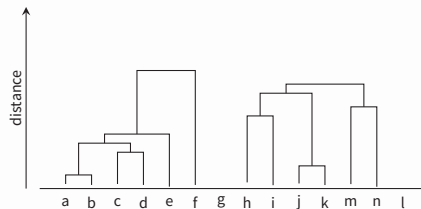
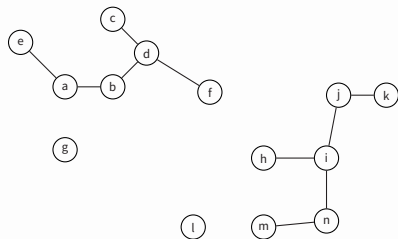
5: $C = C \setminus \{c_i^*, c_j^*\}$

‣ remove pair of clusters

6: $C = C \cup \{c_{i,j}\}$

‣ add merged cluster

7: **until** only one cluster remaining



Agglomerative Clustering: Sketch

$\mathcal{D}(x_l, x_m)$ —distance between *points*

$\mathcal{G}_{\mathcal{D}}(c_i, c_j)$ —distance between *clusters* of points

Require: $\mathcal{G}_{\mathcal{D}}, \{x_1, \dots, x_N\}$

► points

1: $C = \{c_1, \dots, c_N\} = \{\{x_1\}, \dots, \{x_N\}\}$

► initial clusters

2: **repeat**

3: $c_i^*, c_j^* = \arg \min_{c_i, c_j} \mathcal{G}_{\mathcal{D}}(c_i, c_j)$

► find closest pair

4: $c_{i,j} \leftarrow c_i^*, c_j^*$

► merge into new cluster

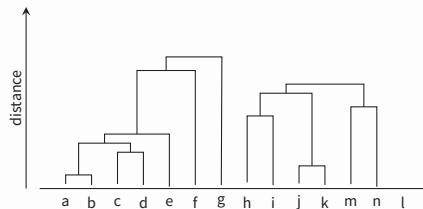
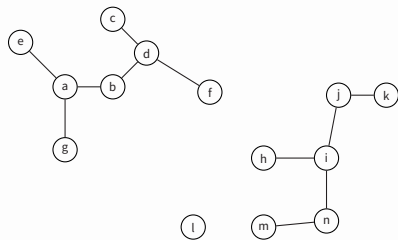
5: $C = C \setminus \{c_i^*, c_j^*\}$

► remove pair of clusters

6: $C = C \cup \{c_{i,j}\}$

► add merged cluster

7: **until** only one cluster remaining



Agglomerative Clustering: Sketch

$\mathcal{D}(x_l, x_m)$ —distance between *points*

$\mathcal{G}_{\mathcal{D}}(c_i, c_j)$ —distance between *clusters* of points

Require: $\mathcal{G}_{\mathcal{D}}, \{x_1, \dots, x_N\}$

▷ points

1: $C = \{c_1, \dots, c_N\} = \{\{x_1\}, \dots, \{x_N\}\}$

▷ initial clusters

2: **repeat**

3: $c_i^*, c_j^* = \arg \min_{c_i, c_j} \mathcal{G}_{\mathcal{D}}(c_i, c_j)$

▷ find closest pair

4: $c_{i,j} \leftarrow c_i^*, c_j^*$

▷ merge into new cluster

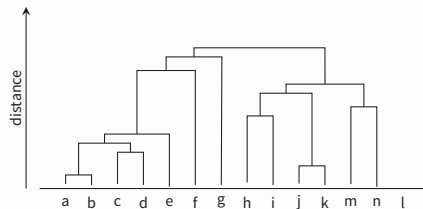
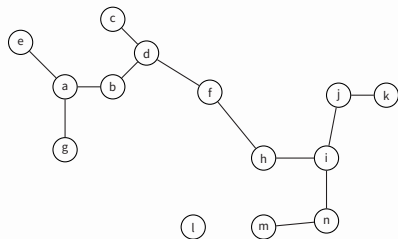
5: $C = C \setminus \{c_i^*, c_j^*\}$

▷ remove pair of clusters

6: $C = C \cup \{c_{i,j}\}$

▷ add merged cluster

7: **until** only one cluster remaining



Agglomerative Clustering: Sketch

$\mathcal{D}(x_l, x_m)$ —distance between *points*

$\mathcal{G}_{\mathcal{D}}(c_i, c_j)$ —distance between *clusters* of points

Require: $\mathcal{G}_{\mathcal{D}}, \{x_1, \dots, x_N\}$

▷ points

1: $C = \{c_1, \dots, c_N\} = \{\{x_1\}, \dots, \{x_N\}\}$

▷ initial clusters

2: **repeat**

3: $c_i^*, c_j^* = \arg \min_{c_i, c_j} \mathcal{G}_{\mathcal{D}}(c_i, c_j)$

▷ find closest pair

4: $c_{i,j} \leftarrow c_i^*, c_j^*$

▷ merge into new cluster

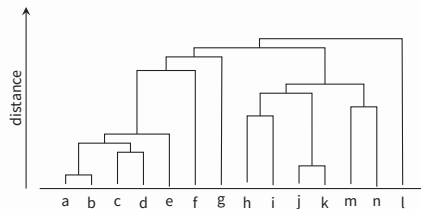
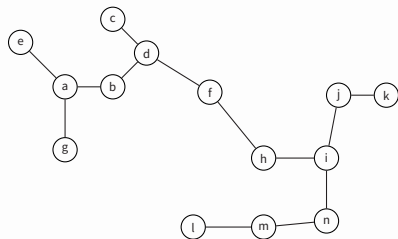
5: $C = C \setminus \{c_i^*, c_j^*\}$

▷ remove pair of clusters

6: $C = C \cup \{c_{i,j}\}$

▷ add merged cluster

7: **until** only one cluster remaining



Agglomerative Clustering: Sketch

$\mathcal{D}(x_l, x_m)$ —distance between *points*

$\mathcal{G}_{\mathcal{D}}(c_i, c_j)$ —distance between *clusters* of points

Require: $\mathcal{G}_{\mathcal{D}}, \{x_1, \dots, x_N\}$

▷ points

1: $C = \{c_1, \dots, c_N\} = \{\{x_1\}, \dots, \{x_N\}\}$

▷ initial clusters

2: **repeat**

3: $c_i^*, c_j^* = \arg \min_{c_i, c_j} \mathcal{G}_{\mathcal{D}}(c_i, c_j)$

▷ find closest pair

4: $c_{i,j} \leftarrow c_i^*, c_j^*$

▷ merge into new cluster

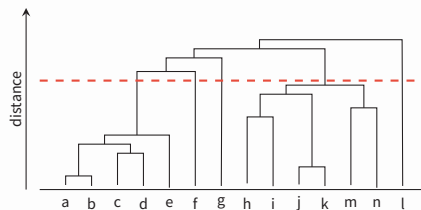
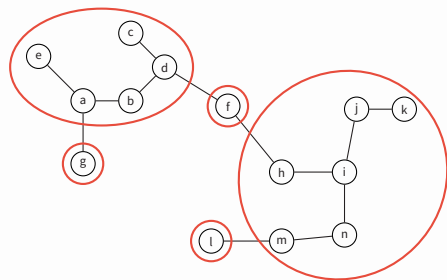
5: $C = C \setminus \{c_i^*, c_j^*\}$

▷ remove pair of clusters

6: $C = C \cup \{c_{i,j}\}$

▷ add merged cluster

7: **until** only one cluster remaining



Cluster Distance Measures

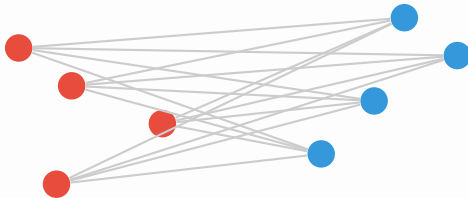
Single Link

$$\mathcal{G}_D(c_i, c_j) = \min_{\substack{\mathbf{x}_{i,l} \in c_i \\ \mathbf{x}_{j,m} \in c_j}} \mathcal{D}(\mathbf{x}_{i,l}, \mathbf{x}_{j,m})$$

Cluster Distance Measures

Single Link

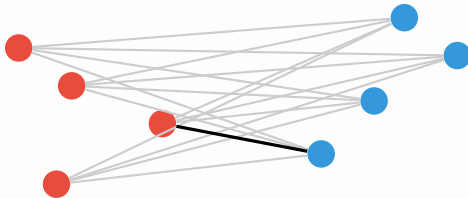
$$\mathcal{G}_D(c_i, c_j) = \min_{\substack{\mathbf{x}_{i,l} \in c_i \\ \mathbf{x}_{j,m} \in c_j}} D(\mathbf{x}_{i,l}, \mathbf{x}_{j,m})$$



Cluster Distance Measures

Single Link

$$\mathcal{G}_D(c_i, c_j) = \min_{\substack{x_{i,l} \in c_i \\ x_{j,m} \in c_j}} D(x_{i,l}, x_{j,m})$$



Cluster Distance Measures

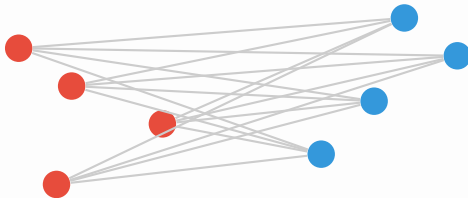
Complete Link

$$\mathcal{G}_D(c_i, c_j) = \max_{\substack{\mathbf{x}_{i,l} \in c_i \\ \mathbf{x}_{j,m} \in c_j}} \mathcal{D}(\mathbf{x}_{i,l}, \mathbf{x}_{j,m})$$

Cluster Distance Measures

Complete Link

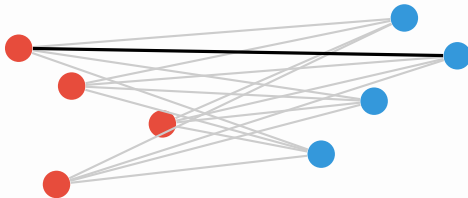
$$\mathcal{G}_D(c_i, c_j) = \max_{\substack{\mathbf{x}_{i,l} \in c_i \\ \mathbf{x}_{j,m} \in c_j}} \mathcal{D}(\mathbf{x}_{i,l}, \mathbf{x}_{j,m})$$



Cluster Distance Measures

Complete Link

$$\mathcal{G}_D(c_i, c_j) = \max_{\substack{\mathbf{x}_{i,l} \in c_i \\ \mathbf{x}_{j,m} \in c_j}} \mathcal{D}(\mathbf{x}_{i,l}, \mathbf{x}_{j,m})$$



Cluster Distance Measures

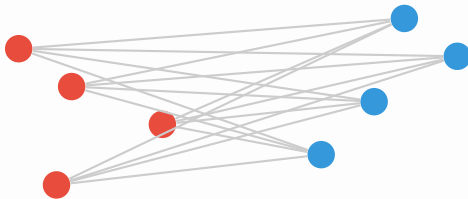
Average Link

$$\mathcal{G}_{\mathcal{D}}(c_i, c_j) = \frac{1}{|c_i| |c_j|} \sum_{\substack{x_{i,l} \in c_i \\ x_{j,m} \in c_j}} \mathcal{D}(x_{i,l}, x_{j,m})$$

Cluster Distance Measures

Average Link

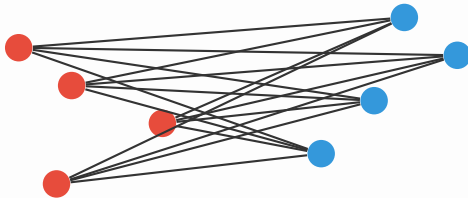
$$\mathcal{G}_{\mathcal{D}}(c_i, c_j) = \frac{1}{|c_i| |c_j|} \sum_{\substack{x_{i,l} \in c_i \\ x_{j,m} \in c_j}} \mathcal{D}(x_{i,l}, x_{j,m})$$



Cluster Distance Measures

Average Link

$$\mathcal{G}_{\mathcal{D}}(c_i, c_j) = \frac{1}{|c_i| |c_j|} \sum_{\substack{x_{i,l} \in c_i \\ x_{j,m} \in c_j}} \mathcal{D}(x_{i,l}, x_{j,m})$$



Cluster Distance Measures

Ward's Method

$$\bar{\mathbf{x}}_{ij} = \frac{1}{|\mathbf{c}_{ij}|} \sum_{\mathbf{x}_l \in \mathbf{c}_{ij}} \mathbf{x}_l \quad (\mathbf{c}_{ij} = \mathbf{c}_i \cup \mathbf{c}_j)$$

$$\mathcal{G}_{\mathcal{D}}(\mathbf{c}_i, \mathbf{c}_j) = \frac{1}{|\mathbf{c}_{ij}|} \sum_{\mathbf{x}_l \in \mathbf{c}_{ij}} \mathcal{D}(\mathbf{x}_l, \bar{\mathbf{x}}_{ij}) = \frac{1}{|\mathbf{c}_{ij}|} \sum_{\mathbf{x}_l \in \mathbf{c}_{ij}} \|\mathbf{x}_l - \bar{\mathbf{x}}_{ij}\|^2$$

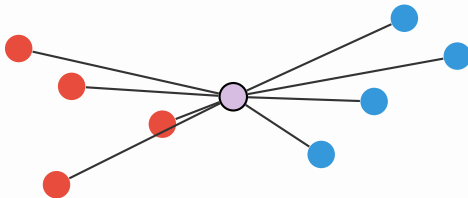
Cluster Distance Measures

Ward's Method

$$\bar{x}_{ij} = \frac{1}{|c_{ij}|} \sum_{x_l \in c_{ij}} x_l$$

$$(c_{ij} = c_i \cup c_j)$$

$$\mathcal{G}_D(c_i, c_j) = \frac{1}{|c_{ij}|} \sum_{x_l \in c_{ij}} \mathcal{D}(x_l, \bar{x}_{ij}) = \frac{1}{|c_{ij}|} \sum_{x_l \in c_{ij}} \|x_l - \bar{x}_{ij}\|^2$$



Unified Formulation

Lance-Williams Algorithm

- When merging two clusters to get $c_{i,j}$
- Need to compute updated distances to all other clusters

Unified Formulation

Lance-Williams Algorithm

- When merging two clusters to get $c_{i,j}$
- Need to compute updated distances to all other clusters

For each remaining cluster c_k , denoting $G_{i,j} = \mathcal{G}_{\mathcal{D}}(c_i, c_j)$

$$G_{k,i,j} = \alpha_i G_{k,i} + \alpha_j G_{k,j} + \beta G_{i,j} + \gamma |G_{k,i} - G_{k,j}|$$

Unified Formulation

Lance-Williams Algorithm

- When merging two clusters to get $c_{i,j}$
- Need to compute updated distances to all other clusters

For each remaining cluster c_k , denoting $G_{i,j} = \mathcal{G}_{\mathcal{D}}(c_i, c_j)$

$$G_{k,i,j} = \alpha_i G_{k,i} + \alpha_j G_{k,j} + \beta G_{i,j} + \gamma |G_{k,i} - G_{k,j}|$$

Method	α_i	α_j	β	γ
Single Link	0.5	0.5	0	-0.5
Complete Link	0.5	0.5	0	0.5
Average Link	$\frac{ c_i }{ c_i + c_j }$	$\frac{ c_j }{ c_i + c_j }$	0	0
Ward's Method	$\frac{ c_i + c_k }{ c_i + c_j + c_k }$	$\frac{ c_j + c_k }{ c_i + c_j + c_k }$	$\frac{- c_k }{ c_i + c_j + c_k }$	0

Evaluation



Evaluation

Extrinsic

Helps solve downstream task

- **Quantisation:** represent data with cluster features
 - colour quantisation—use centroid value
 - feature extraction—use cluster index

Evaluation

Extrinsic

Helps solve downstream task

- **Quantisation:** represent data with cluster features
 - colour quantisation—use centroid value
 - feature extraction—use cluster index
- **Partition:** treat clusters as different datasets
 - train separate classifiers for each sub-group
 - e.g. MNIST 1 vs. not 1; 2 vs. not 2 ...

Evaluation

Extrinsic

Helps solve downstream task

- **Quantisation:** represent data with cluster features
 - colour quantisation—use centroid value
 - feature extraction—use cluster index
- **Partition:** treat clusters as different datasets
 - train separate classifiers for each sub-group
 - e.g. MNIST 1 vs. not 1; 2 vs. not 2 ...
- **Key:** Does it help perform task better?

Evaluation

Intrinsic

Helps understand qualitative makeup of data

- **Unsupervised:** measure how well-separated clusters are
 - compare intra-cluster distances to inter-cluster distances
 - e.g. silhouette scores

Evaluation

Intrinsic

Helps understand qualitative makeup of data

- **Unsupervised:** measure how well-separated clusters are
 - compare intra-cluster distances to inter-cluster distances
 - e.g. silhouette scores
- **Supervised:** measure alignment of clusters to known labels
 - can treat as evaluation of classification
 - reason in terms of pairs belonging to cluster / label
 - **issue:** # cluster $\neq$ # labels

Evaluation

Intrinsic

Helps understand qualitative makeup of data

- **Unsupervised:** measure how well-separated clusters are
 - compare intra-cluster distances to inter-cluster distances
 - e.g. silhouette scores
- **Supervised:** measure alignment of clusters to known labels
 - can treat as evaluation of classification
 - reason in terms of pairs belonging to cluster / label
 - **issue:** # cluster $\neq$ # labels
- **Human:** compare judgements to humans on exemplars
 - ask human if pair x_i, x_j belong together
 - compute match between human judgements and predictions: F1-score, κ , etc.

Intrinsic Evaluation: Unsupervised

In the absence of labels, or any other external measure of utility, can compute a generic measure of how well-clustered the data is.

Silhouette Score

Let data point $x_l \in c_i$ be denoted $x_{i,l}$, then

Intrinsic Evaluation: Unsupervised

In the absence of labels, or any other external measure of utility, can compute a generic measure of how well-clustered the data is.

Silhouette Score

Let data point $x_l \in c_i$ be denoted $x_{i,l}$, then

$$a_l = \frac{1}{|c_i| - 1} \sum_{\substack{x_{i,m} \in c_i \\ m \neq l}} \mathcal{D}(x_{i,l}, x_{i,m})$$

mean distance within cluster

Intrinsic Evaluation: Unsupervised

In the absence of labels, or any other external measure of utility, can compute a generic measure of how well-clustered the data is.

Silhouette Score

Let data point $x_l \in c_i$ be denoted $x_{i,l}$, then

$$a_l = \frac{1}{|c_i| - 1} \sum_{\substack{x_{i,m} \in c_i \\ m \neq l}} \mathcal{D}(x_{i,l}, x_{i,m})$$

mean distance within cluster

$$b_l = \min_{j \neq i} \frac{1}{|c_j|} \sum_{x_{j,m} \in c_j} \mathcal{D}(x_{i,l}, x_{j,m})$$

mean distance with *nearest* cluster

Intrinsic Evaluation: Unsupervised

In the absence of labels, or any other external measure of utility, can compute a generic measure of how well-clustered the data is.

Silhouette Score

Let data point $x_l \in c_i$ be denoted $x_{i,l}$, then

$$a_l = \frac{1}{|c_i| - 1} \sum_{\substack{x_{i,m} \in c_i \\ m \neq l}} \mathcal{D}(x_{i,l}, x_{i,m})$$

mean distance within cluster

$$b_l = \min_{j \neq i} \frac{1}{|c_j|} \sum_{x_{j,m} \in c_j} \mathcal{D}(x_{i,l}, x_{j,m})$$

mean distance with *nearest* cluster

$$s_l = \frac{b_l - a_l}{\max\{a_l, b_l\}} \quad |c_i| > 1$$

Intrinsic Evaluation: Unsupervised

In the absence of labels, or any other external measure of utility, can compute a generic measure of how well-clustered the data is.

Silhouette Score

Let data point $x_l \in c_i$ be denoted $x_{i,l}$, then

$$a_l = \frac{1}{|c_i| - 1} \sum_{\substack{x_{i,m} \in c_i \\ m \neq l}} \mathcal{D}(x_{i,l}, x_{i,m})$$

mean distance within cluster

$$b_l = \min_{j \neq i} \frac{1}{|c_j|} \sum_{x_{j,m} \in c_j} \mathcal{D}(x_{i,l}, x_{j,m})$$

mean distance with *nearest* cluster

$$s_l = \frac{b_l - a_l}{\max\{a_l, b_l\}} \quad |c_i| > 1$$

$$s = \frac{1}{N} \sum_{l=1}^N s_l \quad -1 \leq s \leq 1$$

Intrinsic Evaluation: Supervised

Issue: Alignment

Clustering produces clusters $C = \{c_1, \dots, c_U\}$

Intrinsic Evaluation: Supervised

Issue: Alignment

Clustering produces clusters $C = \{c_1, \dots, c_U\}$

Labels induce *reference* clusters $\mathcal{R} = \{r_1, \dots, r_V\}$

Intrinsic Evaluation: Supervised

Issue: Alignment

Clustering produces clusters $C = \{c_1, \dots, c_U\}$

Labels induce *reference* clusters $\mathcal{R} = \{r_1, \dots, r_V\}$

- if $U = V$

Intrinsic Evaluation: Supervised

Issue: Alignment

Clustering produces clusters $C = \{c_1, \dots, c_U\}$

Labels induce *reference* clusters $\mathcal{R} = \{r_1, \dots, r_V\}$

- if $U = V$
 - still cannot compare directly—permutation unknown!

Intrinsic Evaluation: Supervised

Issue: Alignment

Clustering produces clusters $C = \{c_1, \dots, c_U\}$

Labels induce *reference* clusters $\mathcal{R} = \{r_1, \dots, r_V\}$

- if $U = V$
 - still cannot compare directly—permutation unknown!
 - which u corresponds to which v ?

Intrinsic Evaluation: Supervised

Issue: Alignment

Clustering produces clusters $C = \{c_1, \dots, c_U\}$

Labels induce *reference* clusters $\mathcal{R} = \{r_1, \dots, r_V\}$

- if $U = V$
 - still cannot compare directly—permutation unknown!
 - which u corresponds to which v ?
 - if $u \leftrightarrow v$ matching known
- standard measures: accuracy, F1-score, etc.

Intrinsic Evaluation: Supervised

Issue: Alignment

Clustering produces clusters $C = \{c_1, \dots, c_U\}$

Labels induce *reference* clusters $\mathcal{R} = \{r_1, \dots, r_V\}$

- if $U = V$
 - still cannot compare directly—permutation unknown!
 - which u corresponds to which v ?
 - if $u \leftrightarrow v$ matching known
 - standard measures: accuracy, F1-score, etc.
- if $U \neq V$

Intrinsic Evaluation: Supervised

Issue: Alignment

Clustering produces clusters $C = \{c_1, \dots, c_U\}$

Labels induce *reference* clusters $\mathcal{R} = \{r_1, \dots, r_V\}$

- if $U = V$
 - still cannot compare directly—permutation unknown!
 - which u corresponds to which v ?
 - if $u \leftrightarrow v$ matching known
standard measures: accuracy, F1-score, etc.
- if $U \neq V$
 - need to *also* find best alignment

Intrinsic Evaluation: Supervised

Issue: Alignment

Clustering produces clusters $C = \{c_1, \dots, c_U\}$

Labels induce *reference* clusters $\mathcal{R} = \{r_1, \dots, r_V\}$

- if $U = V$
 - still cannot compare directly—permutation unknown!
 - which u corresponds to which v ?
 - if $u \leftrightarrow v$ matching known
 - standard measures: accuracy, F1-score, etc.
- if $U \neq V$
 - need to *also* find best alignment
 - can have multiple $c_u \rightarrow$ *same* r_v

Intrinsic Evaluation: Supervised

Issue: Alignment

Clustering produces clusters $C = \{c_1, \dots, c_U\}$

Labels induce *reference* clusters $\mathcal{R} = \{r_1, \dots, r_V\}$

- if $U = V$
 - still cannot compare directly—permutation unknown!
 - which u corresponds to which v ?
 - if $u \leftrightarrow v$ matching known
standard measures: accuracy, F1-score, etc.
- if $U \neq V$
 - need to *also* find best alignment
 - can have multiple $c_u \rightarrow$ *same* r_v
 - can have multiple $r_v \rightarrow$ *same* c_u

Intrinsic Evaluation: Supervised

Key Idea: Evaluate relationship between *pairs* of data points x_l, x_m

Rand Index (RI)

- $+$: x_l, x_m are in the same cluster
- $-$: x_l, x_m are in different clusters

Intrinsic Evaluation: Supervised

Key Idea: Evaluate relationship between *pairs* of data points $\mathbf{x}_l, \mathbf{x}_m$

Rand Index (RI)

- $+$: $\mathbf{x}_l, \mathbf{x}_m$ are in the same cluster
- $-$: $\mathbf{x}_l, \mathbf{x}_m$ are in different clusters

		Predicted (C)	
		+	-
True (R)	+	TP	FN
	-	FP	TN

$$\text{RI} = \frac{TP + TN}{TP + TN + FP + FN}$$

= Accuracy!

Intrinsic Evaluation: Supervised

Issue: Expected value of RI of two *random* partitions $\neq 0$ (or any constant)

Intrinsic Evaluation: Supervised

Issue: Expected value of RI of two *random* partitions $\neq 0$ (or any constant)

Adjusted Rand Index (ARI)

	c_1	c_2	$\dots$	c_U	sum
r_1	N_{11}	N_{12}	$\dots$	N_{1U}	a_1
r_2	N_{21}	N_{22}	$\dots$	N_{2U}	a_2
$\vdots$	$\vdots$	$\vdots$	$\ddots$	$\vdots$	$\vdots$
r_V	N_{V1}	N_{V2}	$\dots$	N_{VU}	a_V
sum	b_1	b_2	$\dots$	b_U	N

$$N_{ij} = |\mathbf{r}_i \cap \mathbf{c}_j| \quad \binom{N}{2} = \frac{N(N-1)}{2}$$

Intrinsic Evaluation: Supervised

Issue: Expected value of RI of two *random* partitions $\neq 0$ (or any constant)

Adjusted Rand Index (ARI)

	c_1	c_2	$\dots$	c_U	sum
r_1	N_{11}	N_{12}	$\dots$	N_{1U}	a_1
r_2	N_{21}	N_{22}	$\dots$	N_{2U}	a_2
$\vdots$	$\vdots$	$\vdots$	$\ddots$	$\vdots$	$\vdots$
r_V	N_{V1}	N_{V2}	$\dots$	N_{VU}	a_V
sum	b_1	b_2	$\dots$	b_U	N

$$N_{ij} = |\mathbf{r}_i \cap \mathbf{c}_j| \quad \binom{N}{2} = \frac{N(N-1)}{2}$$

$$\text{TP} = \sum_{ij} \binom{N_{ij}}{2}$$

$$\text{Expected RI} = \frac{1}{\binom{N}{2}} \left[\sum_v \binom{a_v}{2} \cdot \sum_u \binom{b_u}{2} \right]$$

$$\text{Max RI} = \frac{1}{2} \left[\sum_v \binom{a_v}{2} + \sum_u \binom{b_u}{2} \right]$$

$$\text{ARI} = \frac{\text{TP} - \text{Expected RI}}{\text{Max RI} - \text{Expected RI}}$$

Summary

- **Clustering:** Means of discovering structure / sub-groups in data

Summary

- **Clustering:** Means of discovering structure / sub-groups in data
- K-Means
 - Hard; Flat; Polythetic
 - Requires knowing K; search for best K
 - Fast; Iterative; Local Minima

Summary

- **Clustering:** Means of discovering structure / sub-groups in data
- K-Means
 - Hard; Flat; Polythetic
 - Requires knowing K; search for best K
 - Fast; Iterative; Local Minima
- Hierarchical Clustering
 - Hard; Hierarchical; Polythetic
 - Top-Down: Hierarchical K-Means
 - Bottom-Up: Agglomerative Clustering
 - multiple variants: single, complete, etc.

Summary

- **Clustering:** Means of discovering structure / sub-groups in data
- K-Means
 - Hard; Flat; Polythetic
 - Requires knowing K; search for best K
 - Fast; Iterative; Local Minima
- Hierarchical Clustering
 - Hard; Hierarchical; Polythetic
 - Top-Down: Hierarchical K-Means
 - Bottom-Up: Agglomerative Clustering
 - multiple variants: single, complete, etc.
- Evaluation
 - Unsupervised, Supervised, and Human-judgement driven