

# **Data Wrangling and Data Analysis**

## **Clustering:**

**Daniel Oberski & Erik-Jan van Kesteren**

Department of Methodology & Statistics

Utrecht University



# This week

- Lecture 1: Problems with missing data
- Lecture 2: Solutions to problems with missing data
- **Lecture 3: Clustering**



# Reading materials about clustering (this week & next)

- Selected paragraphs from **Introduction to Statistical Learning (ISLR)** §12.1 and 12.4
- “Mixture models: latent profile and latent class analysis” [Oberski, 2016] §1, §2
- <http://daob.nl/wp-content/papercite-data/pdf/oberski2016mixturemodels.pdf>

Springer Texts in Statistics

Gareth James  
Daniela Witten  
Trevor Hastie  
Robert Tibshirani

## An Introduction to Statistical Learning

with Applications in R

*Second Edition*

 Springer



# Optional, much more in-depth material

*Clustering strategy and method selection (ch. 31),*

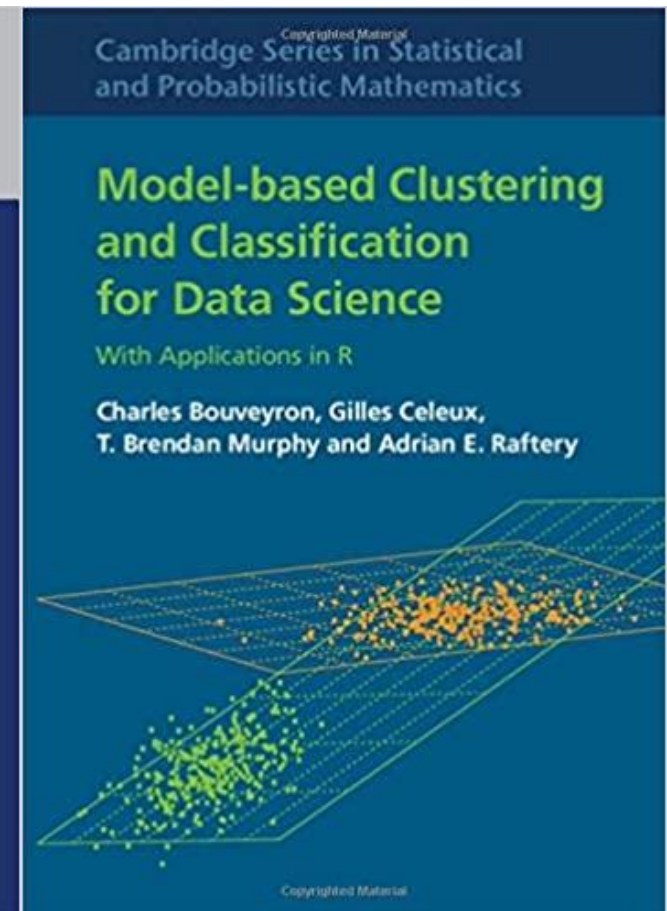
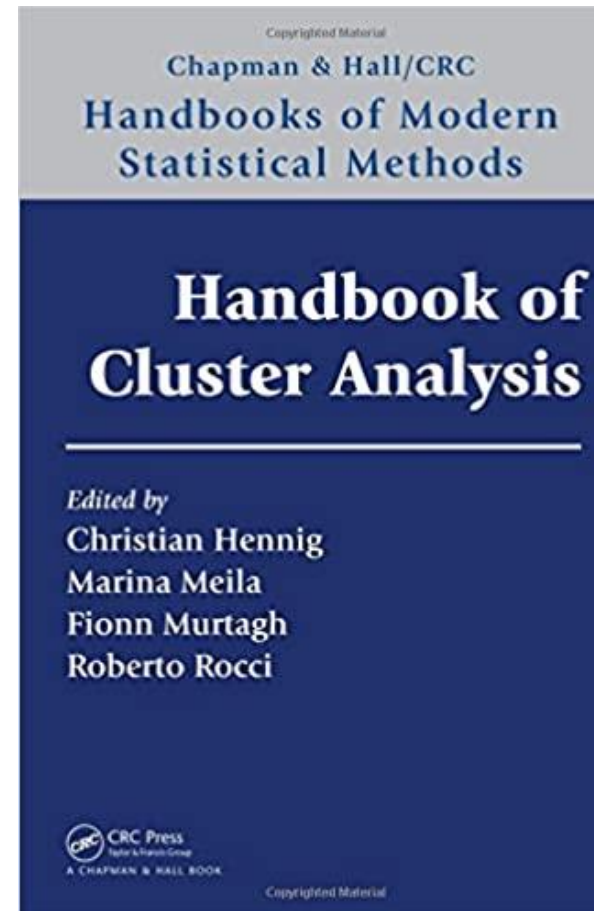
<https://arxiv.org/pdf/1503.02059.pdf>

*Handbook of Cluster Analysis*

Hennig et al. (2016)

*Model-based Clustering and  
Classification for Data Science*

Bouveyron et al. (2018)



# Clustering

Find subgroups (clusters) of similar examples in a database

# Why clustering?

- Unsupervised: expect groups in our data, but were not able to measure them
  - potential new subtypes of cancer tissue
- We want to summarize features into a categorical feature to use in further decisions/analysis
  - subgrouping customers by their spending types



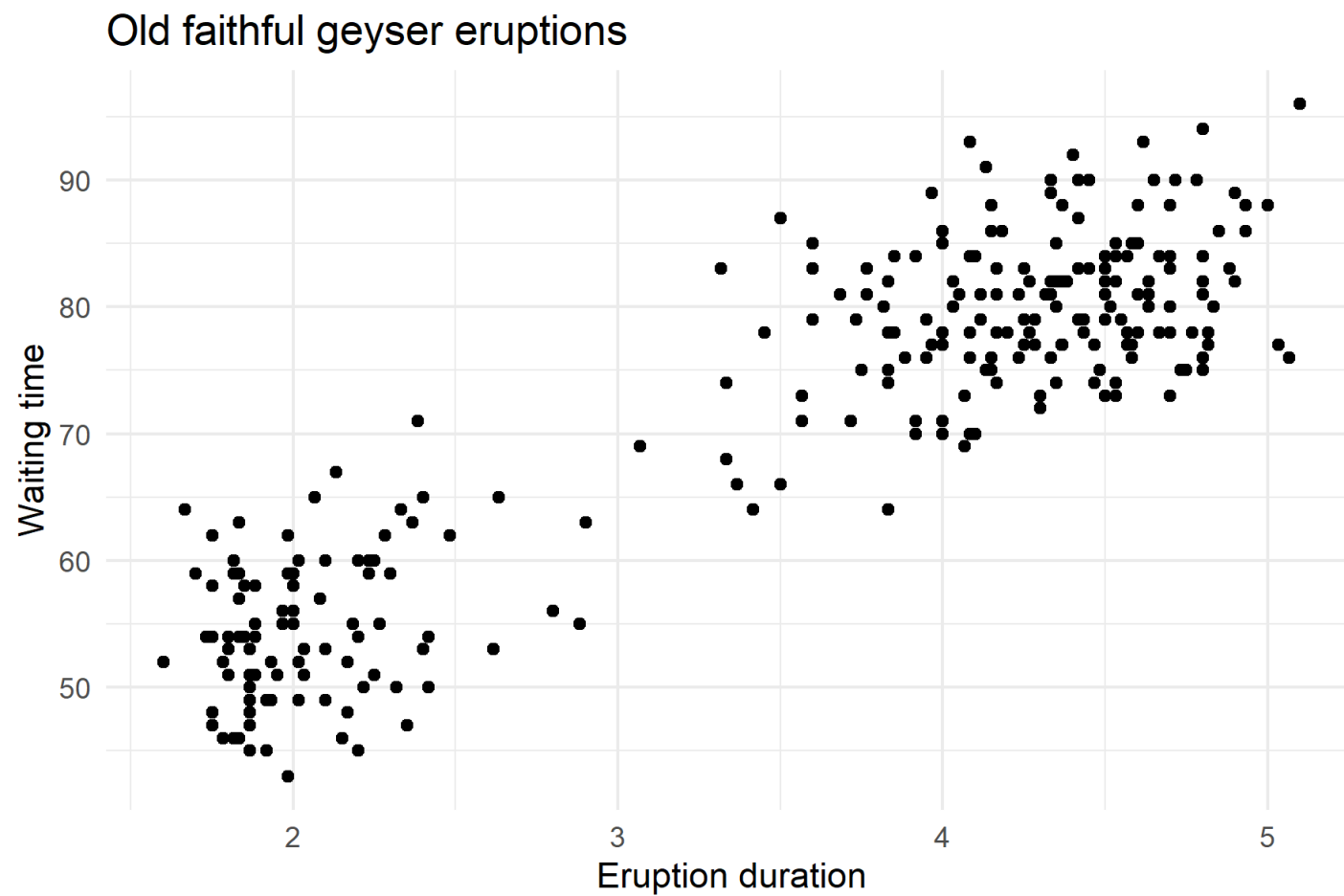
# Some applications of clustering

- Intermediate step for other fundamental data mining problems
- Collaborative filtering
- Customer segmentation
- Data summarization
- Dynamic trend detection
- Multimedia data analysis
- Biological data analysis
- Social network analysis





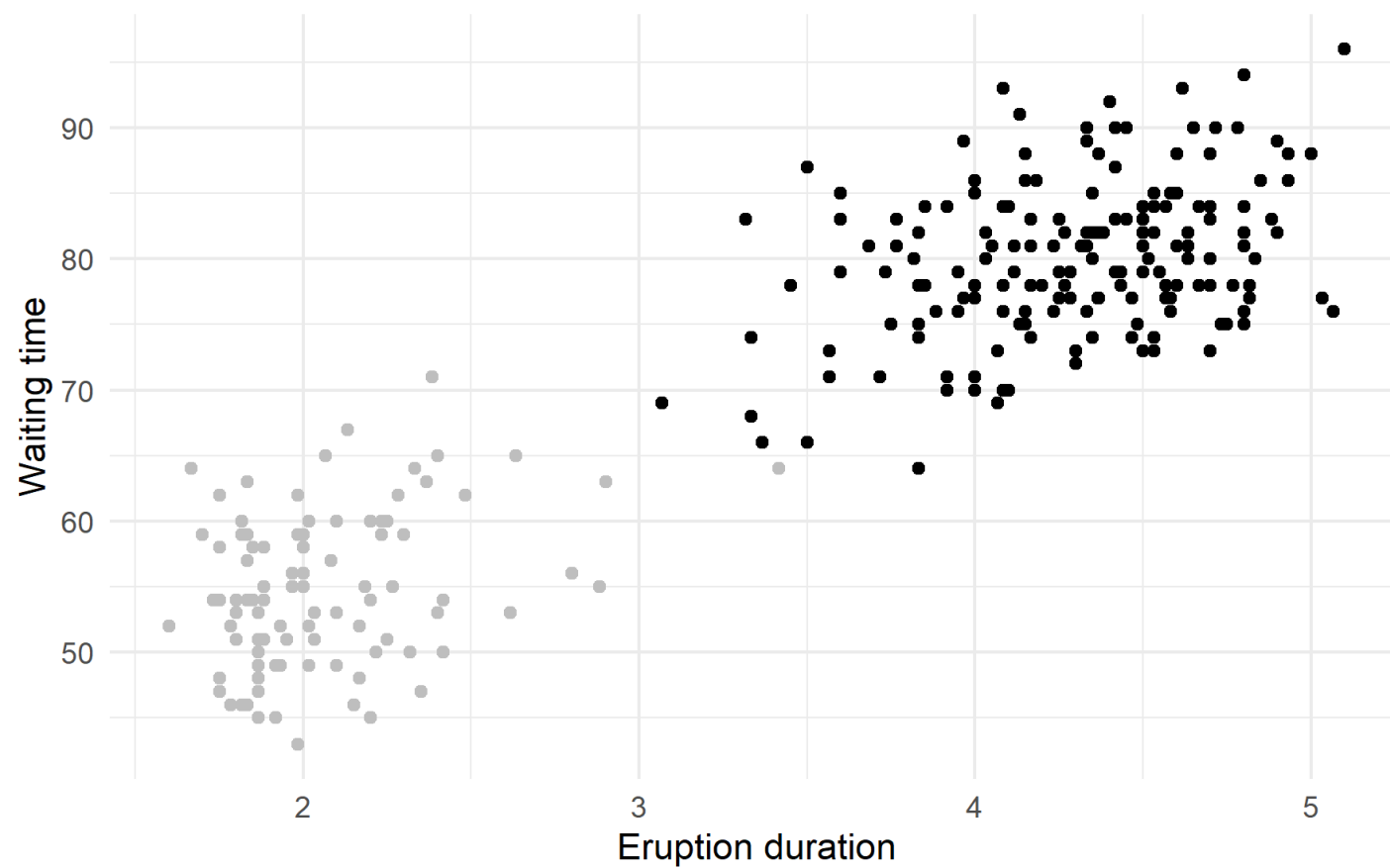
# Old Faithful: two types of eruption?





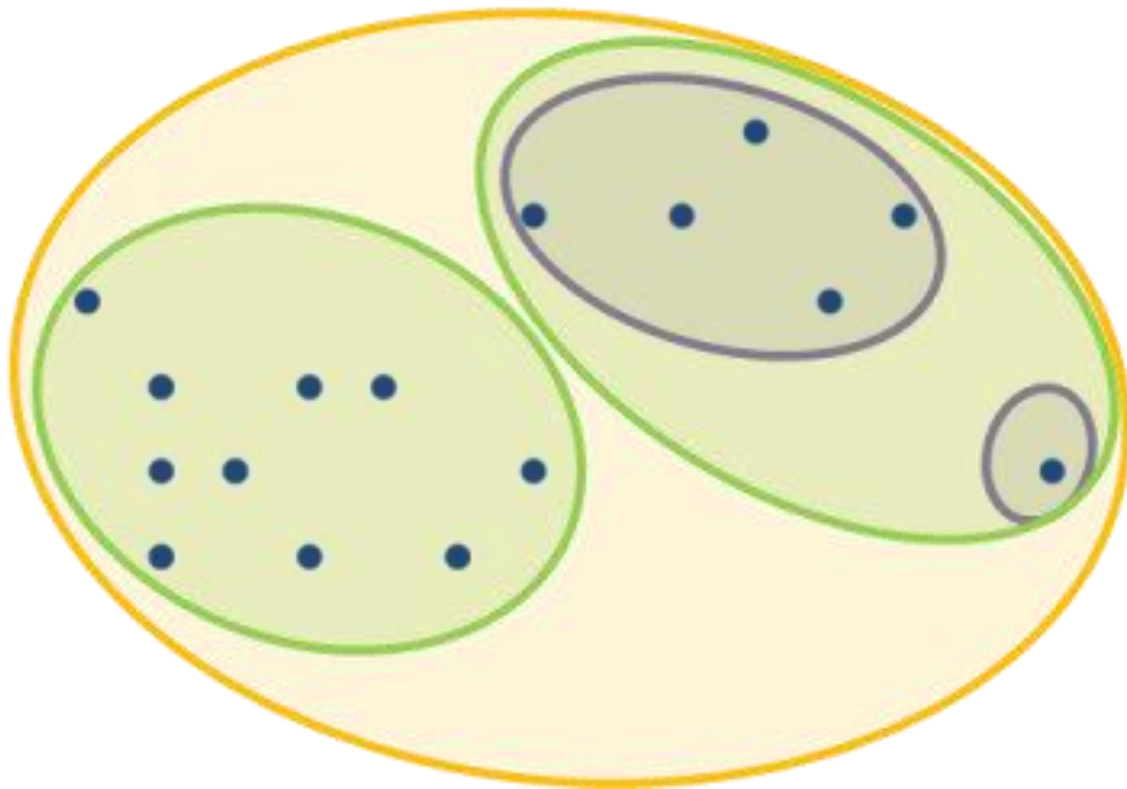
# Old Faithful: two types of eruption?

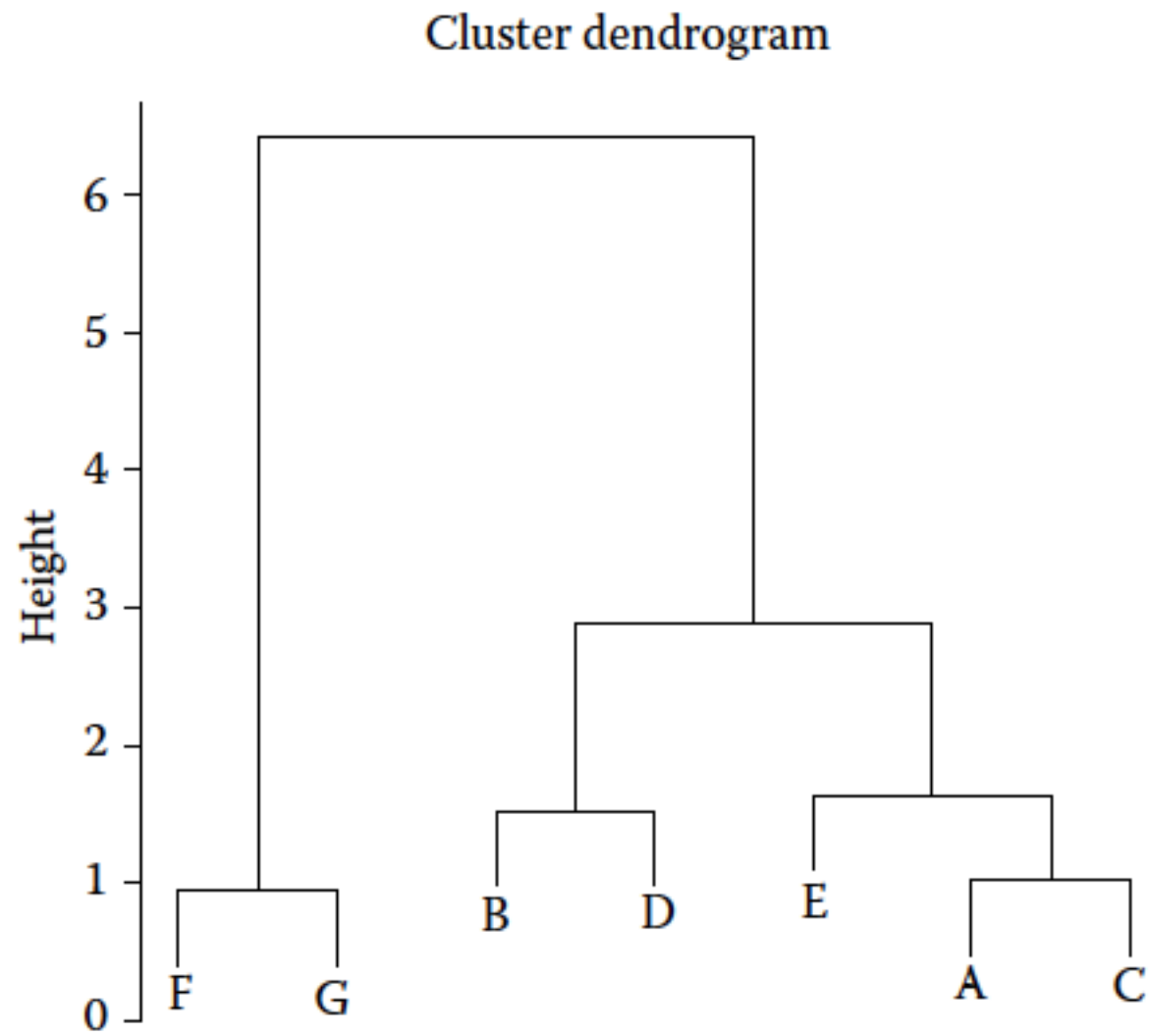
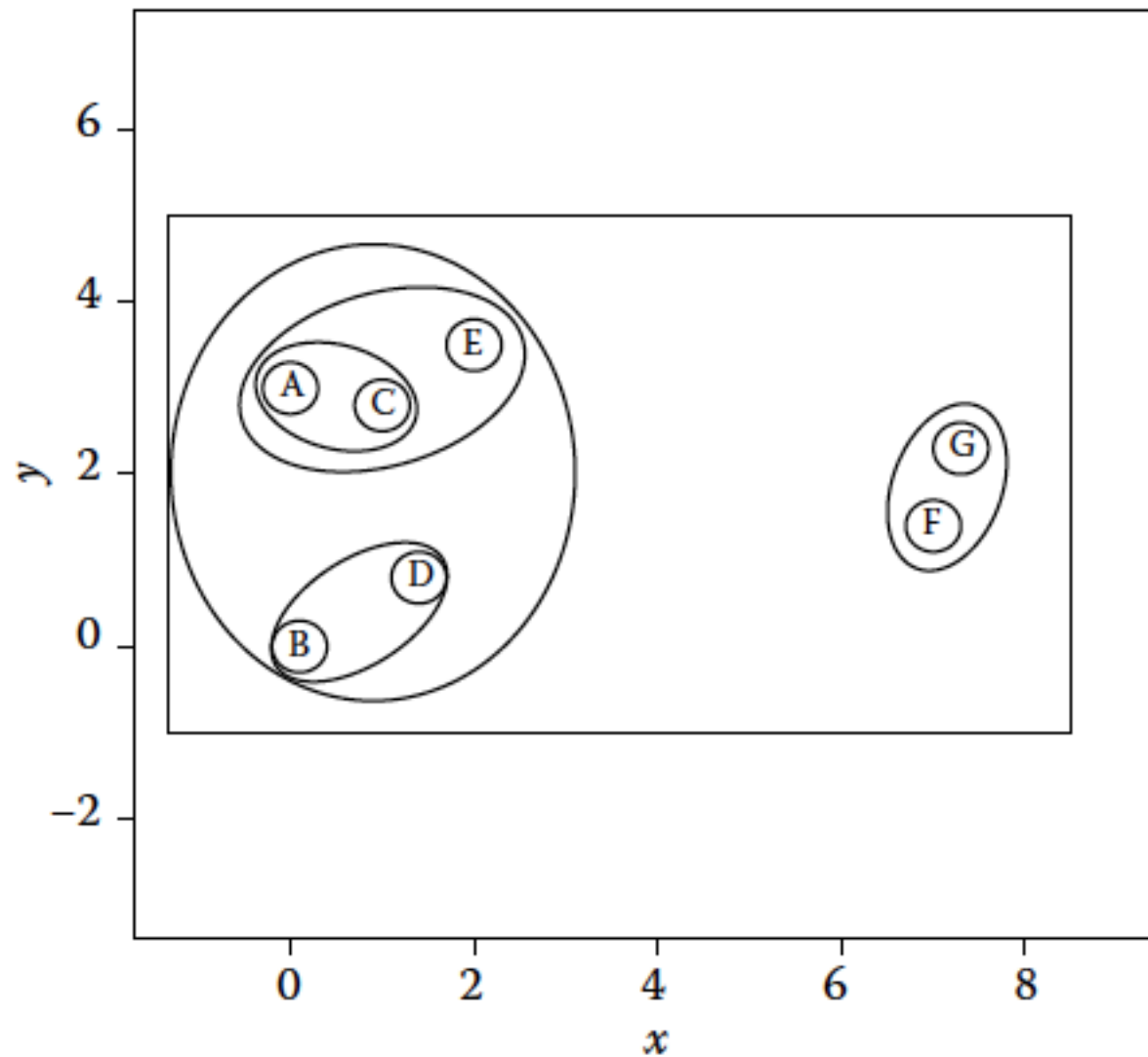
Old faithful geyser eruptions (clustered)



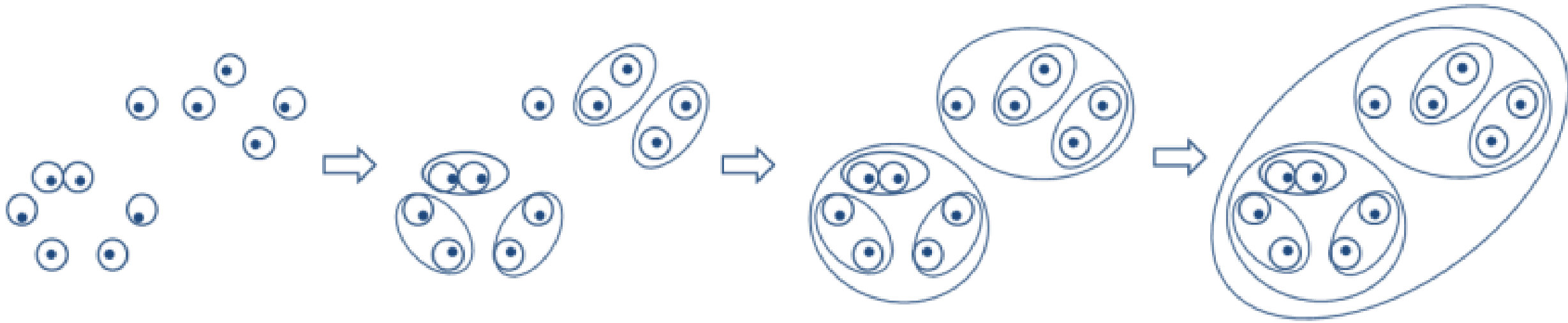
# Clustering types

## Hierarchical Clustering

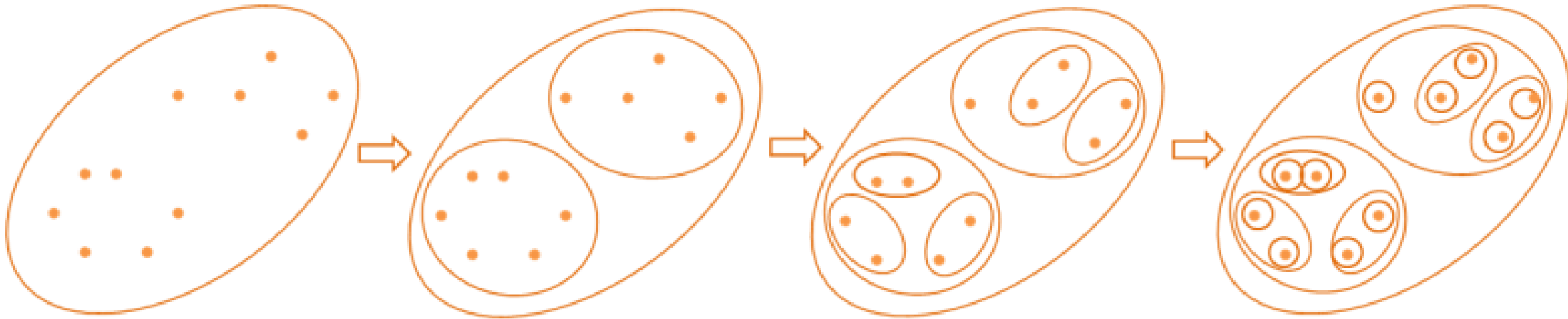




## Agglomerative Hierarchical Clustering



## Divisive Hierarchical Clustering



# Hierarchical clustering

## *Bottom-up agglomerative clustering*

- For each observation, compute the *distance* to all other observations
- Assign all examples to their individual cluster
- Combine *most similar* clusters
- Keep combining clusters until there is only one cluster left
- Select number of clusters for the final solution

*(Divisive: start with **one** cluster and keep splitting most **different**)*



# Hierarchical clustering

In R:

```
distances <- dist(faithful, method = "euclidean")  
result    <- hclust(distances, method = "average")
```

Then we can plot the dendrogram with `plot()` or `ggdendrogram`

```
library(ggdendro)  
ggdendrogram(result)
```

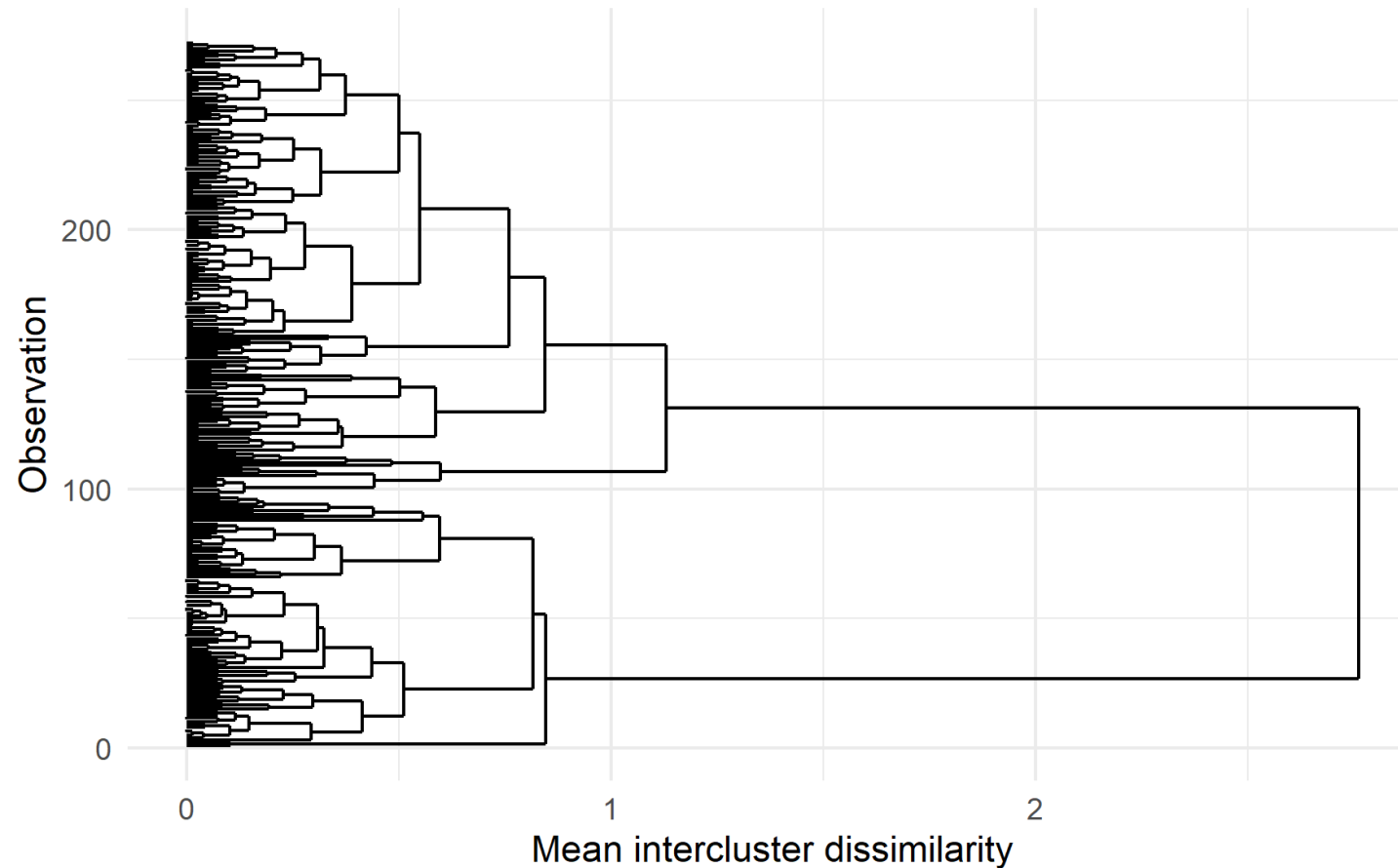
Then, select the number of clusters using a cutoff point

```
cutree(result, h = 2)
```

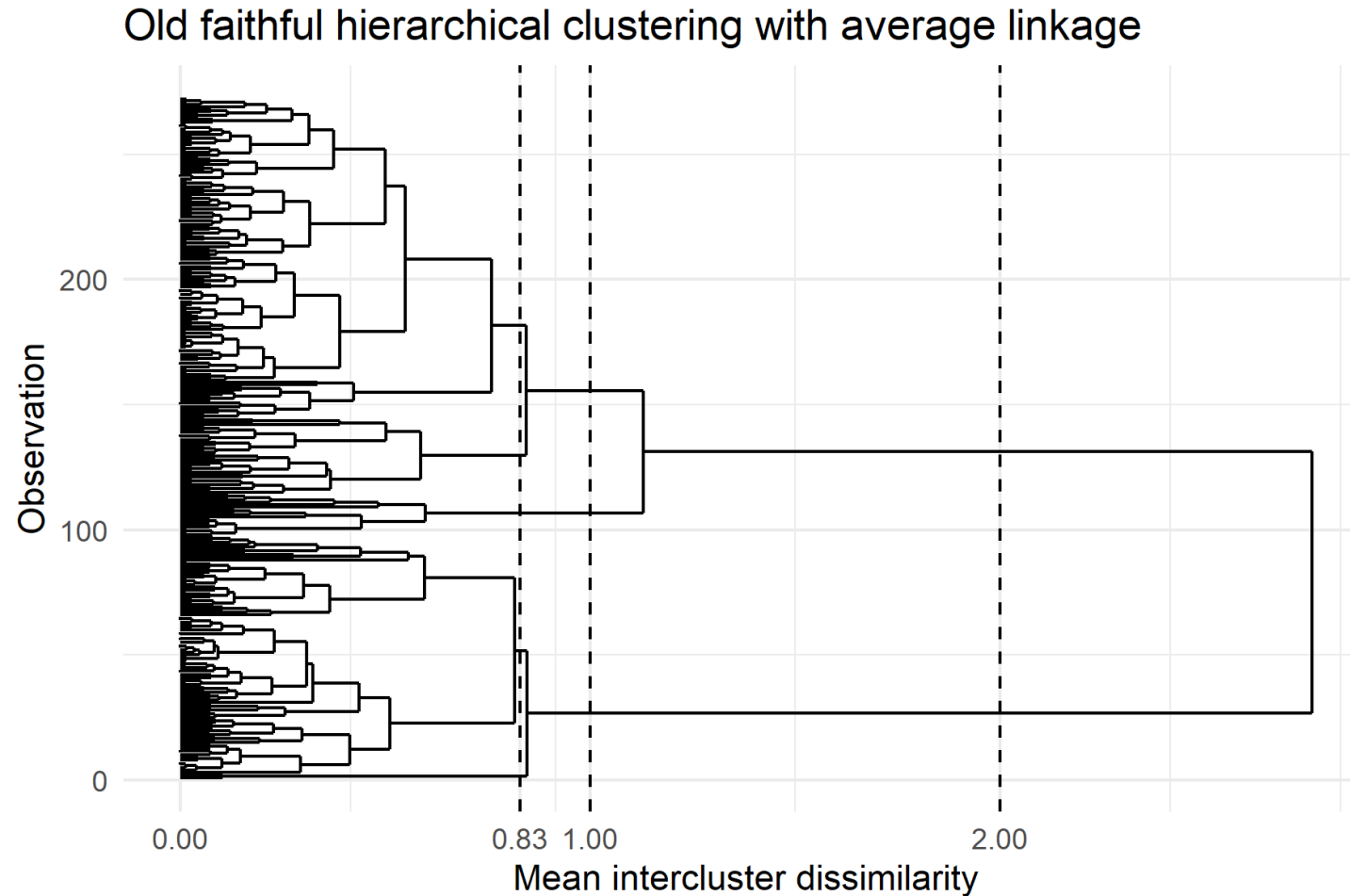


# Hierarchical clustering

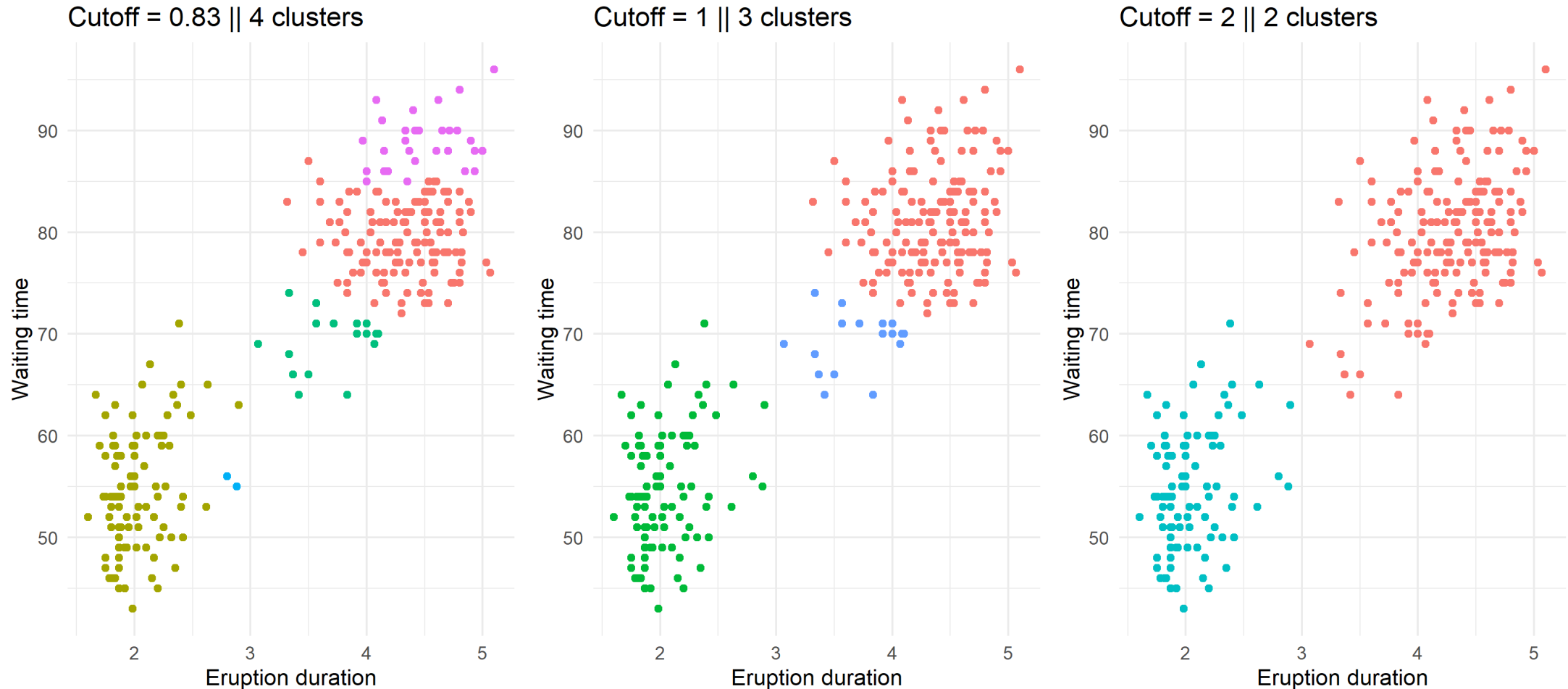
Old faithful hierarchical clustering with average linkage



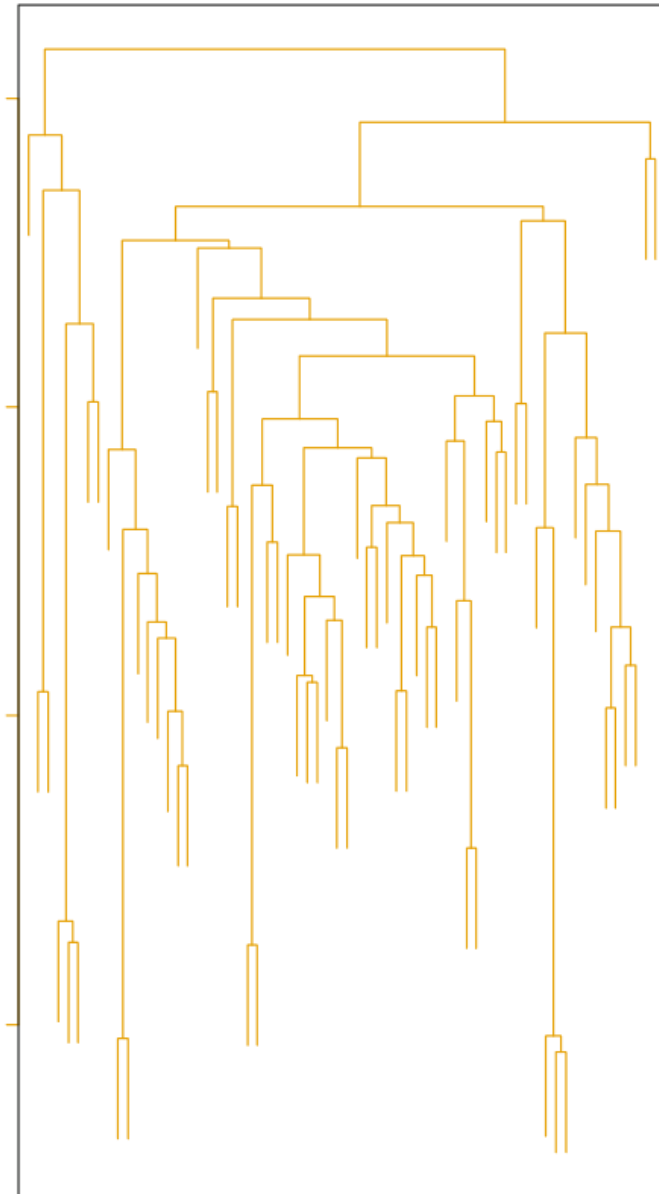
# Hierarchical clustering



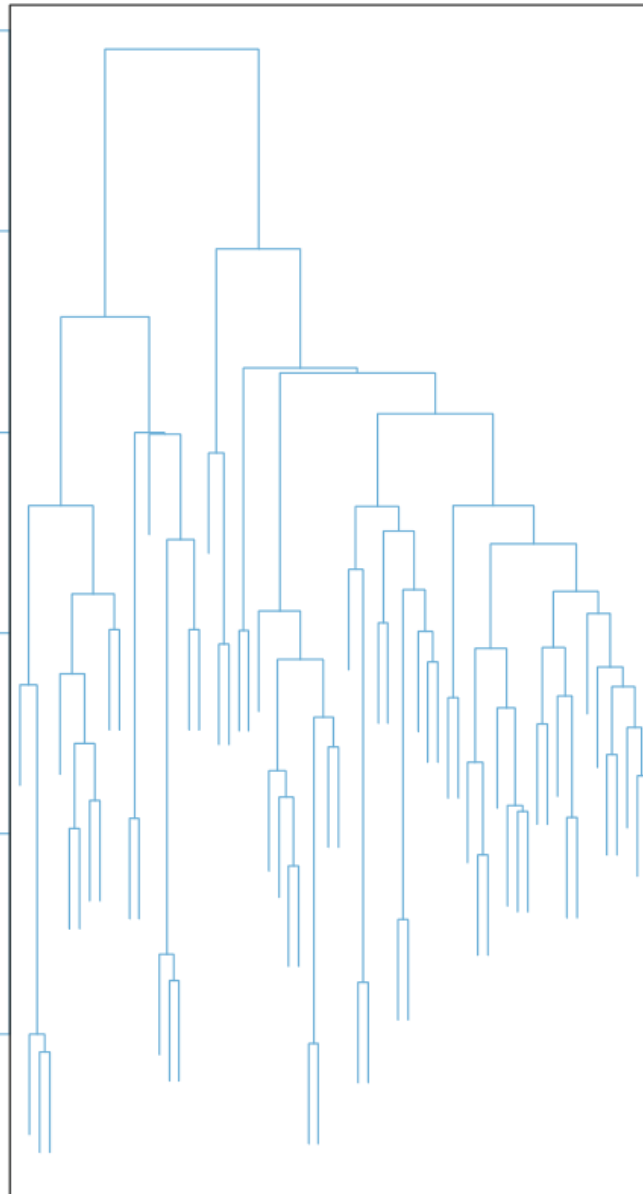
# Hierarchical clustering



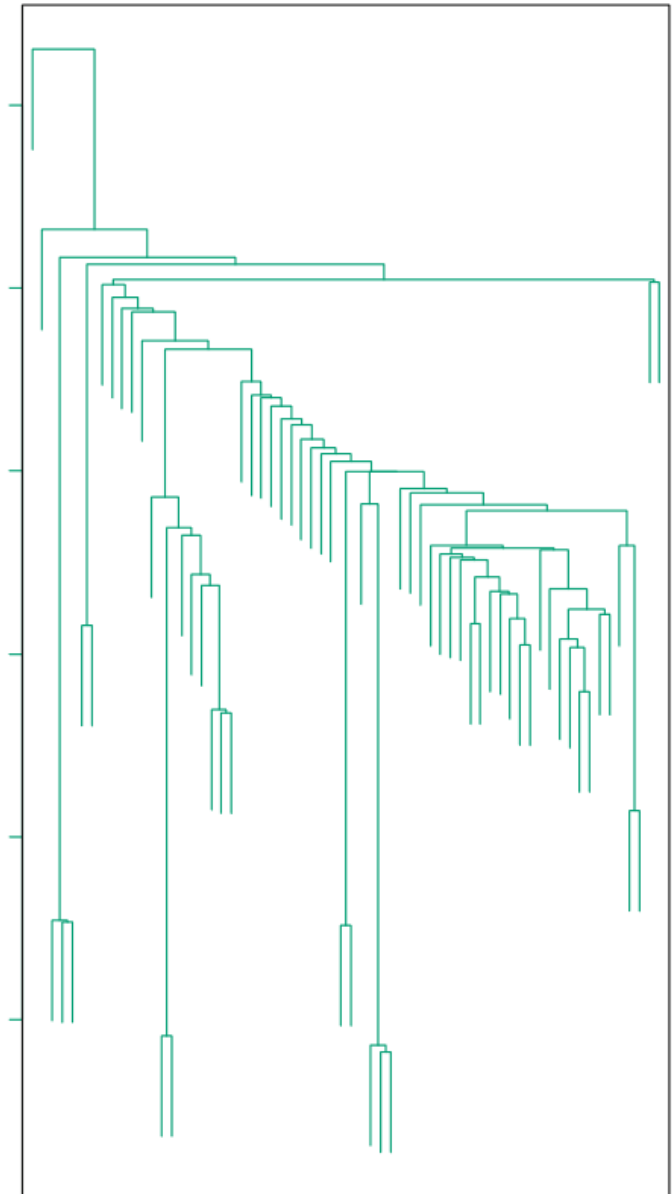
Average Linkage



Complete Linkage



Single Linkage

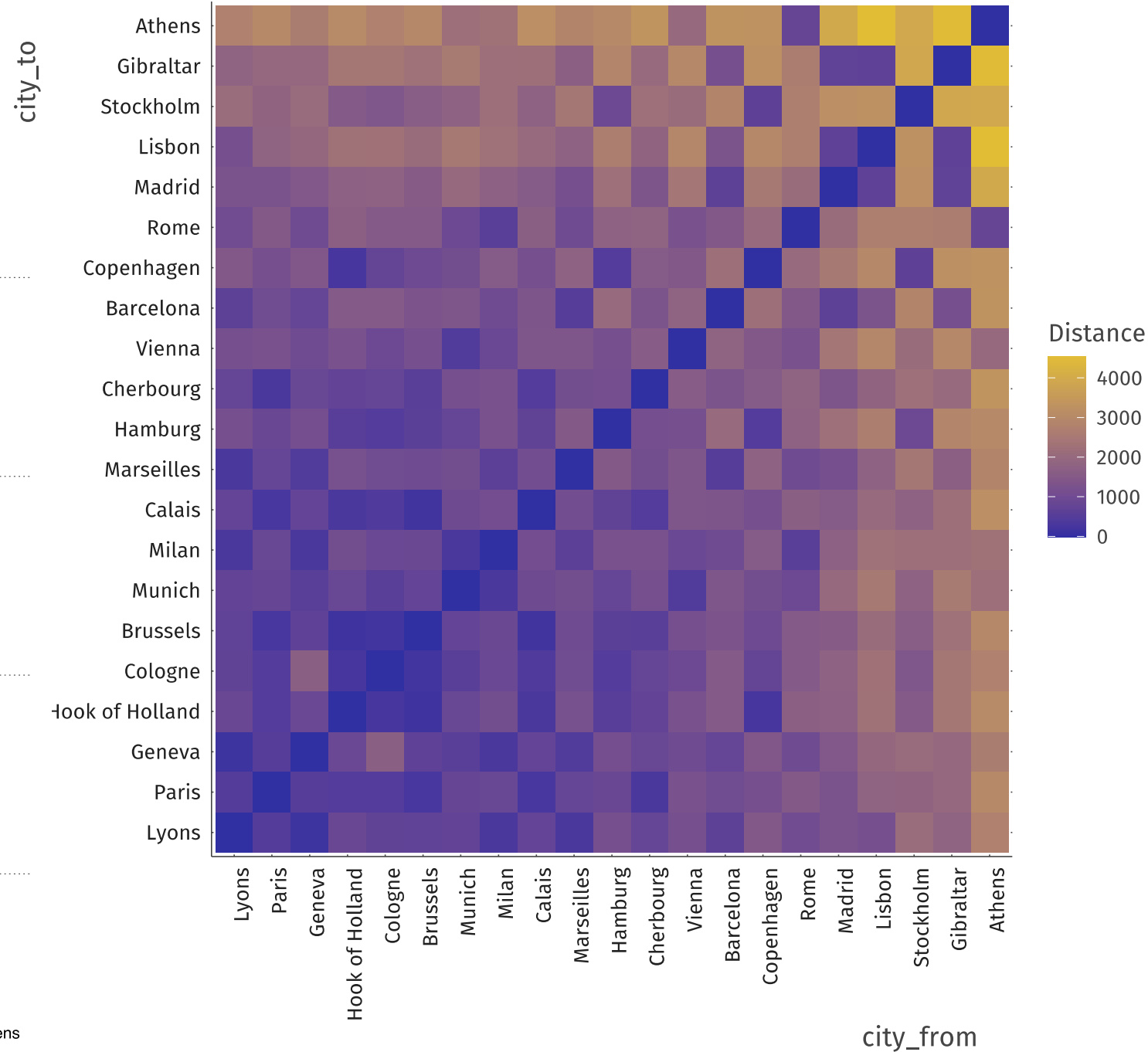


# Note: scaling

- It is generally a good idea to measure your features in the same scale before entering them into a clustering algorithm
- Otherwise, height in **cm** will be more important in the distance computation than width in **m**
- Generally, you want the features to have a similar scale
- This can be done by **standardization**, or z-transformation: subtract the mean from each feature and divide by its observed standard deviation
- Changes the interpretation of the values, but not their association

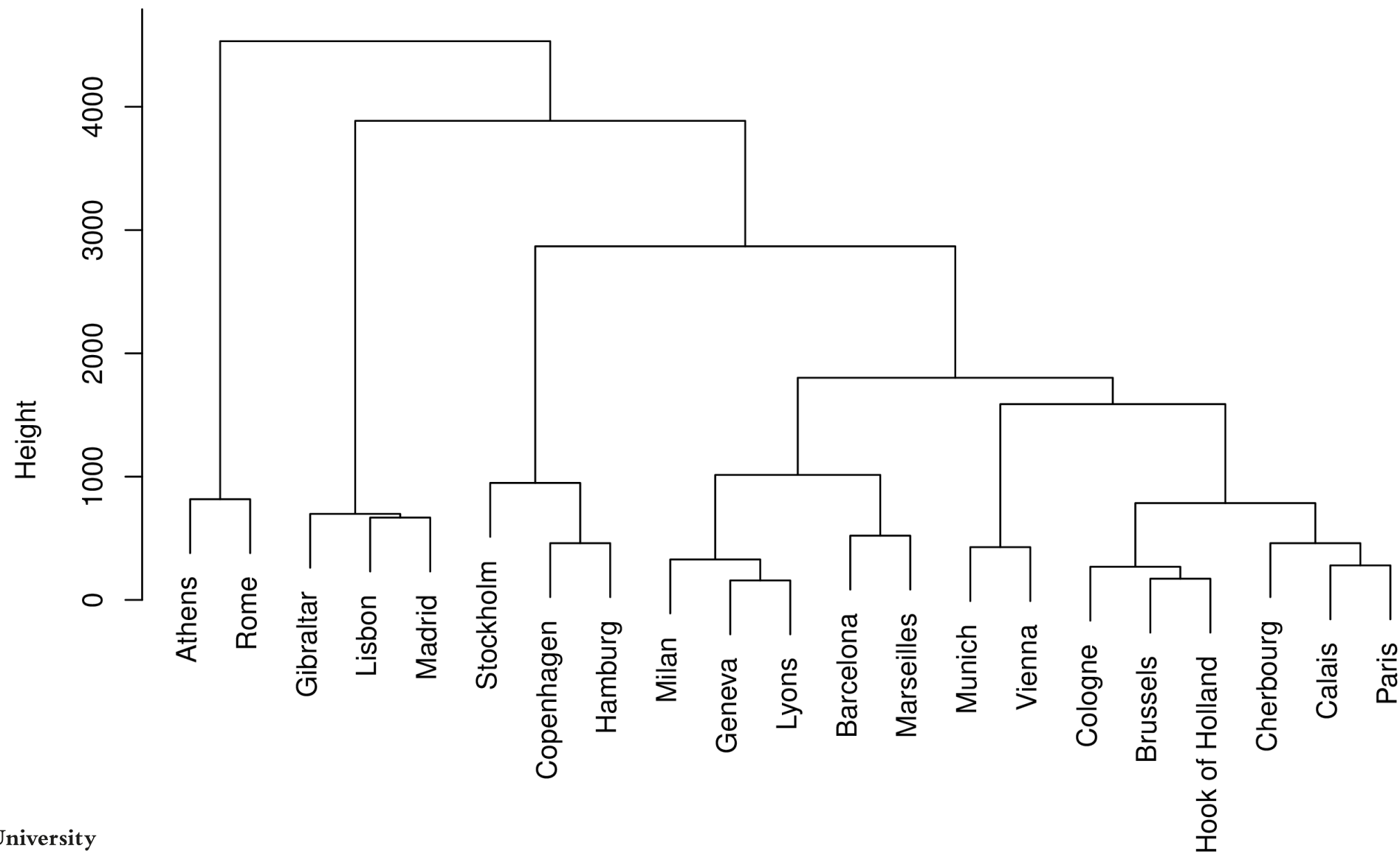


# “Distance” matrix



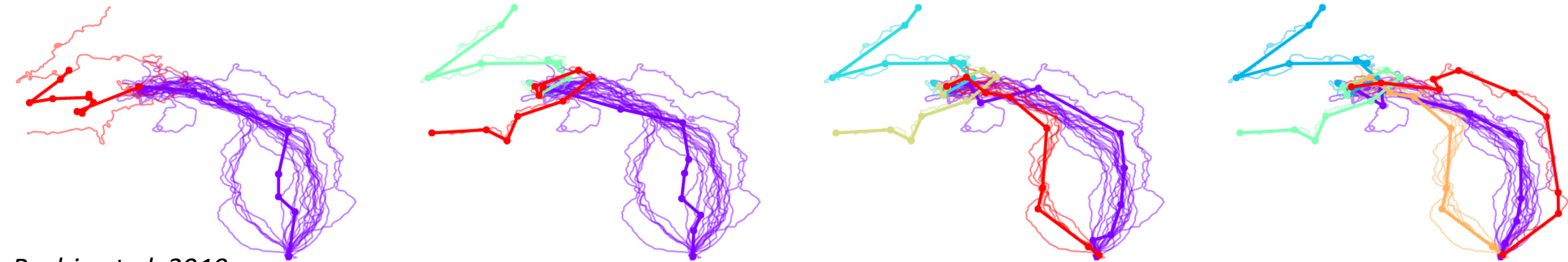
```
plot(hclust(eurodist))
```

Cluster Dendrogram



# “Distance” can mean lots of things

- **Continuous:** Euclidean, maximum, Manhattan, Minkowski, ...
- Time series: “Dynamic time warping”, Fréchet, cross-corr., wavelets, ...
- Networks: Modularity, Shortest path, ...
- Text/DNA: Edit distance, Hamming distance, TF-IDF distance, ..
- Images: “Structural similarity”, GAN loss...



*Buchin et al. 2019*

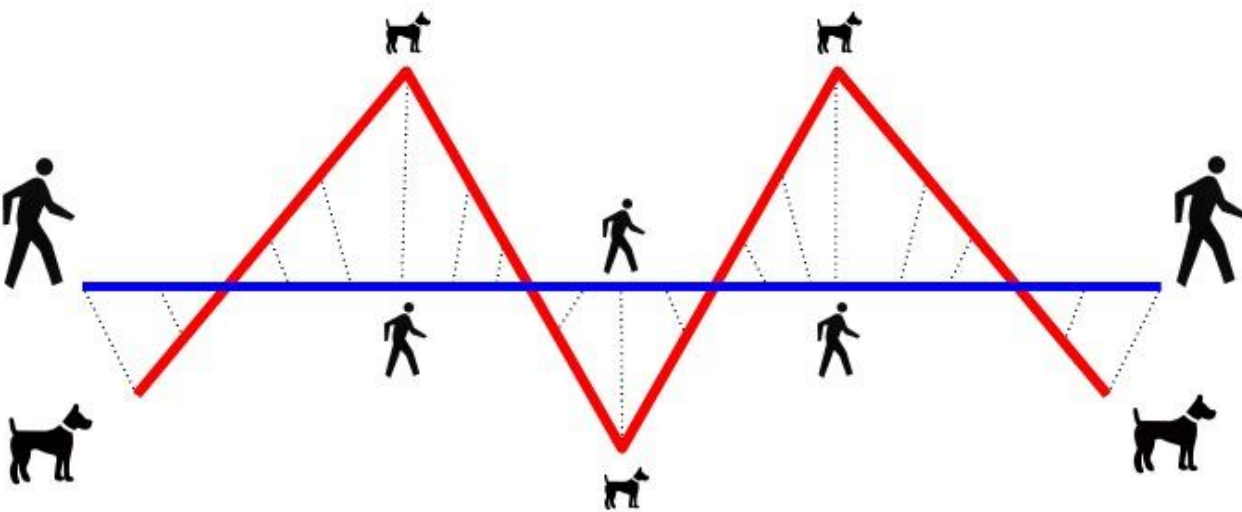
**Figure 1: Example of a  $(k, \ell)$ -clustering for the flight paths of a pigeon with the number of clusters  $k$  increasing from 2 (left) until 5 (right) and the complexity of the clusters being  $\ell = 10$ . Trajectories belonging to the same cluster are shown in the same color. For each cluster, a center trajectory generated by the algorithm is shown using thick lines of the same color.**

Some example distances that aren't Euclidean but useful for specific data types

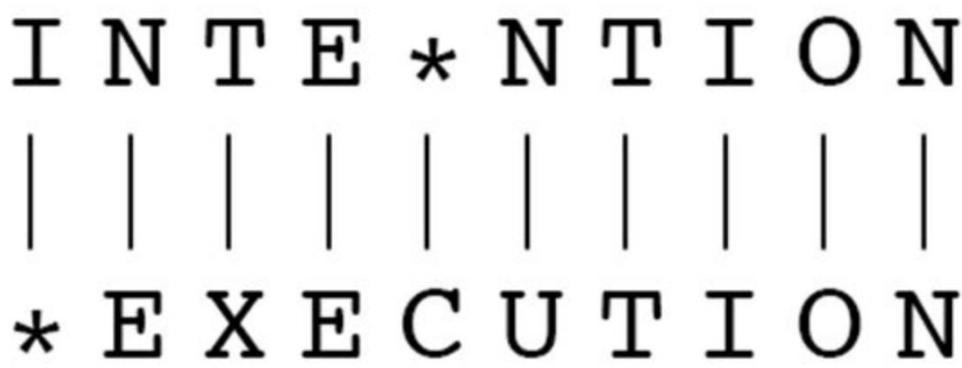
Fréchet distance for time series/curves

Minimum edit distance for text/DNA

Walking your dog



The **Fréchet distance** between the curves is the minimum leash length that permits such a walk



(edit distance = 5)

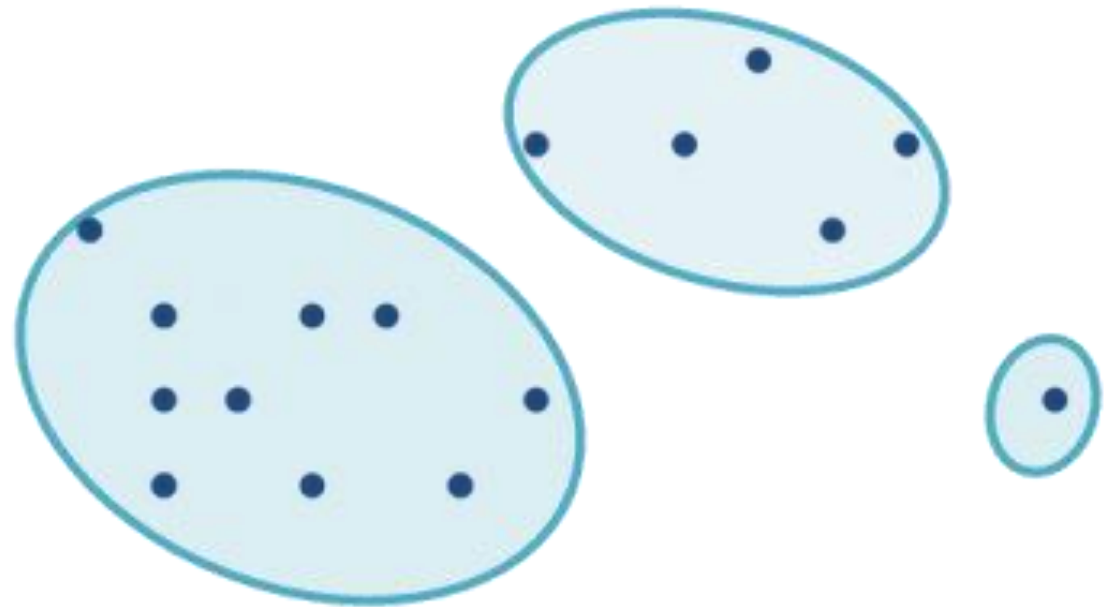
# Hierarchical clustering: conclusion

- Tree-based representation - dendrogram
- Determine number of clusters afterwards
- Different distance metrics possible
- Different agglomeration methods (“linkage” rules) possible
- Taking distance matrix as input → very flexible technique
- Distance matrix  $N \times N$  → can be tricky when  $N$  is large (esp. divisive)



# Clustering types

## Partitional Clustering



Source: T. Fuertes

<https://quantdare.com/hierarchical-clustering/>

# K-means clustering algorithm

1. Randomly assign examples to  $K$  clusters
- 2a. Calculate the **centroid** (per-feature mean) for each cluster

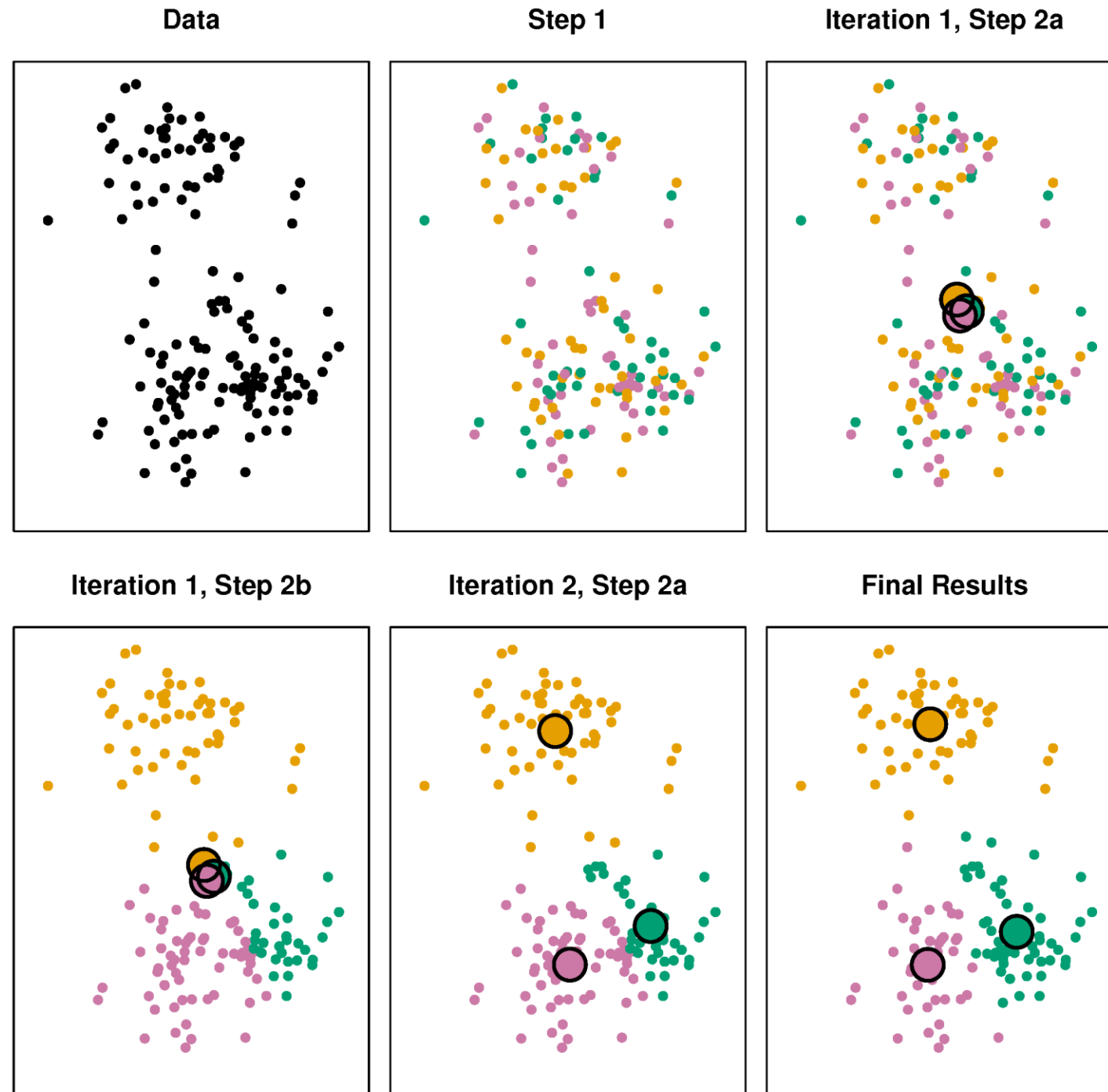
```
faithful %>% group_by(cluster) %>% summarize_all(mean)

#>   cluster eruptions waiting
#>   <int>      <dbl>    <dbl>
#> 1      1        3.69     73.4
#> 2      2        3.30     68.6
```

- 2b. Assign each example to the cluster belonging to its **closest** centroid
3. If the assignments changed, go to step 2a, else stop



# K-means clustering



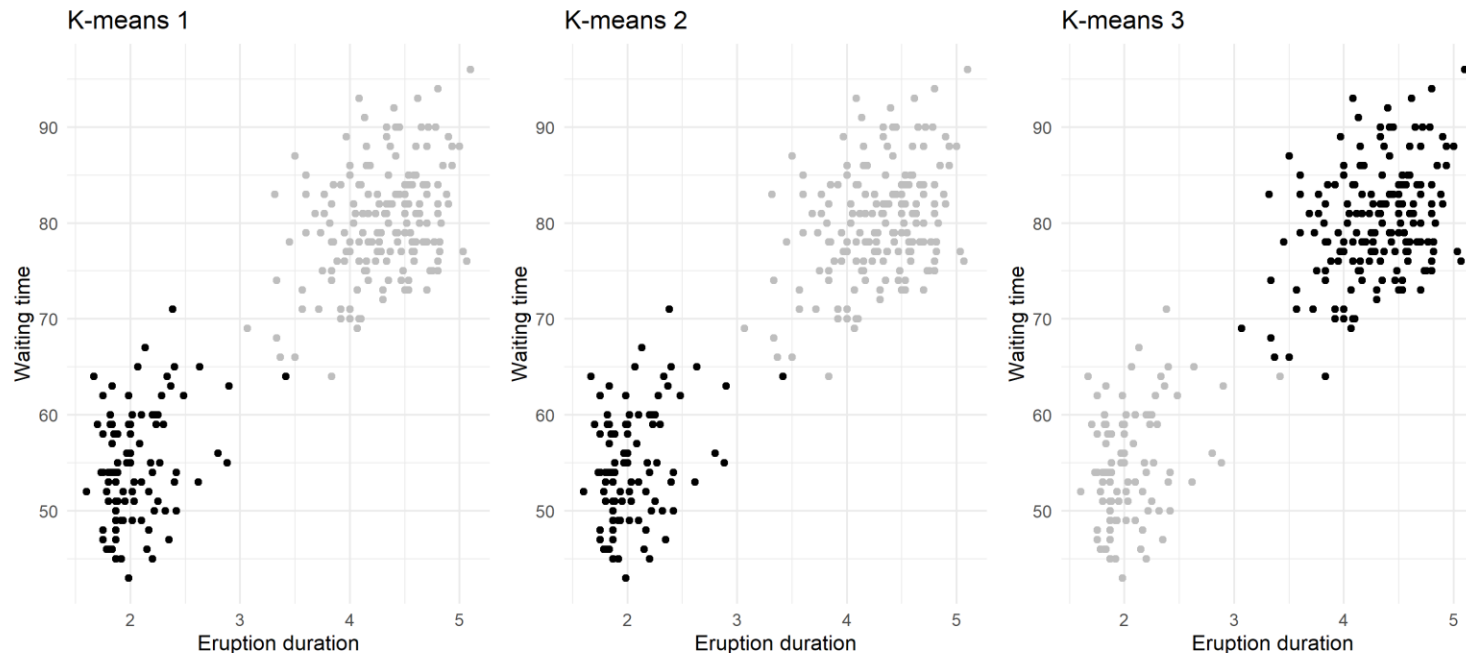
# K-means clustering

- $K$  (and, in theory, the distance metric) are hyperparameters (ISLR Sec 12.4.1)  
*(in practice, distance is almost always Euclidean (sum of squares criterion))*
- $K$  is determined in advance by the analyst
- Could be based on knowledge about the data or the goal of the analysis
  - Perhaps there are generally 2 types of geyser eruption because of physics
  - We may have resources to approach customers in at most 3 different ways
- Other criteria discussed later.



# K-means clustering

- Because the initialization is random, the result is random
  - Label switching: cluster 1, 2, 3 may end up in each other's locations
  - Some examples at the boundary may end up in different clusters altogether
  - Use **multiple starts** to obtain the best solution



# How to **evaluate** clustering results

1. Use of external information
2. Visual exploration
3. Stability assessment / sensitivity analysis
4. Internal validation indexes
5. (Testing for clustering structure)

*Much more info & helpful advice: Clustering strategy & method selection (ch 31 of Handbook of clustering), <https://arxiv.org/pdf/1503.02059.pdf>*



# 1. External validation

Are the clusters associated with *external* feature  $Y$ ?

“Making unsupervised supervised”

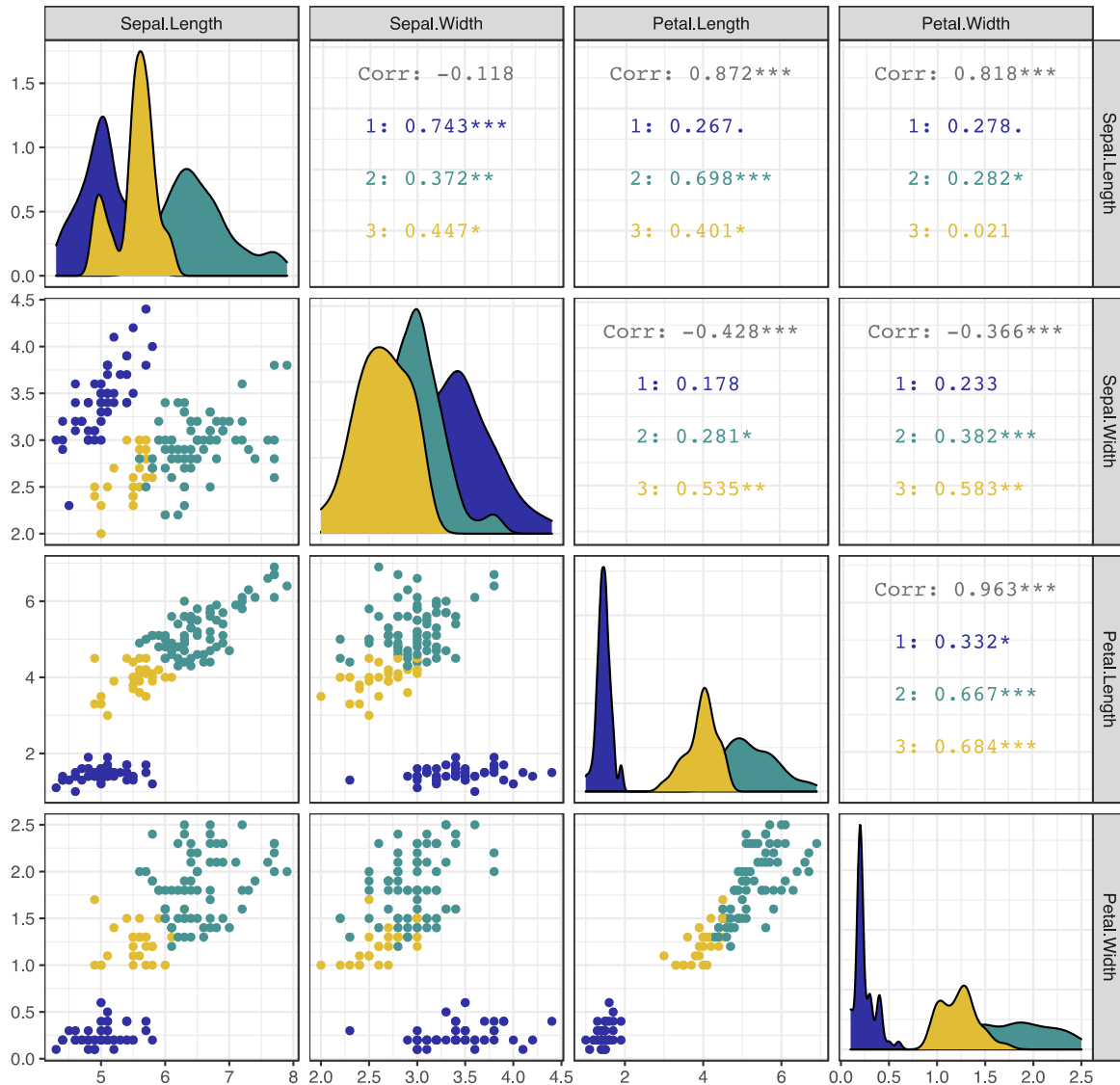
*Examples:*

- Are my customer segments based on spending associated with the demographics of the customers?
- Are the geyser eruption types strongly correlated with water pressure or temperature?
- Can I recognize the person in the vector quantized picture?



## 2. Visual exploration

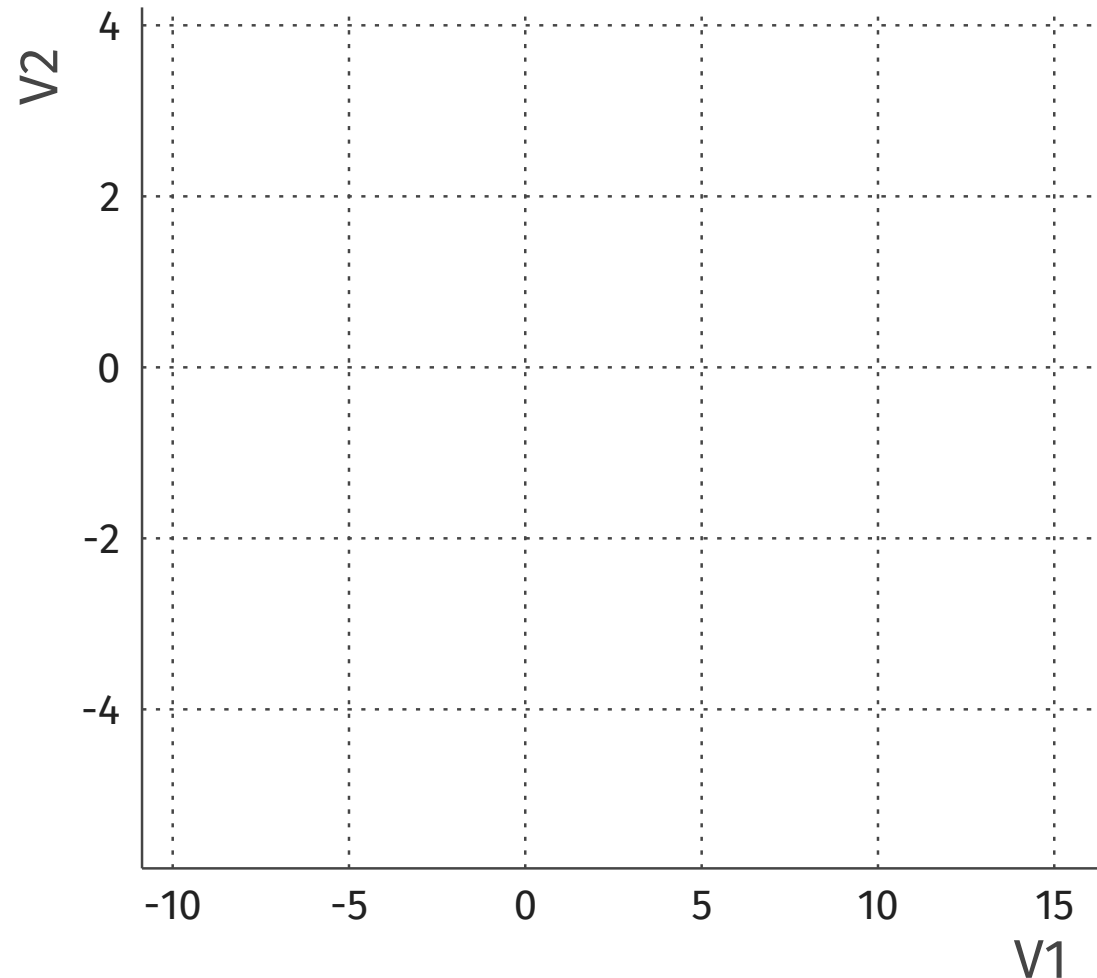
Iris: Scatterplot pairs



- **Problem:** Kind of hard to see already...
- Wait till you get 1000 variables!
- **New idea:** Reduce variables into 2D “manifold” for visualization
- Popular techniques: UMAP, t-SNE, MDS, Discriminant Coordinates, (PCA)

## 2. Visual exploration (using “manifold”)

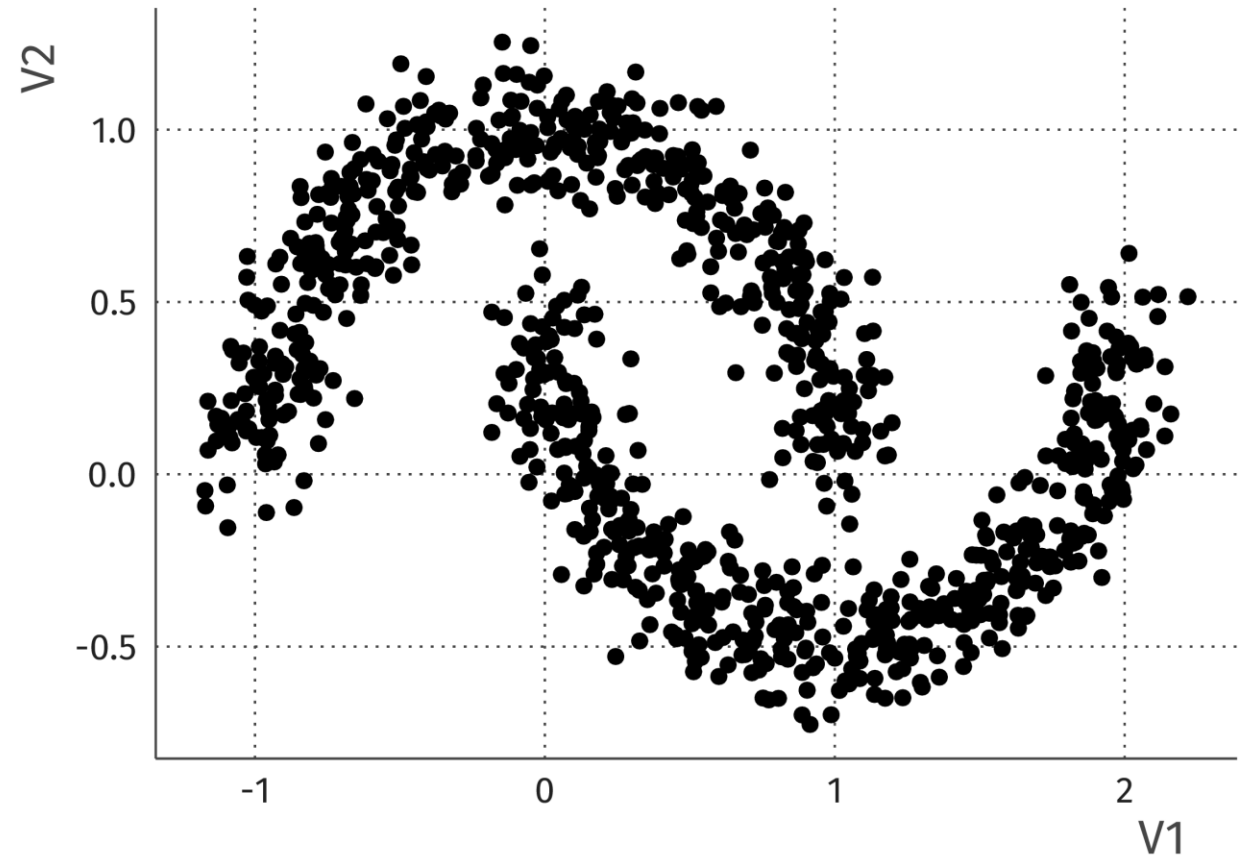
Iris: UMAP representation



# 3. Stability assessment

A.k.a.: Clustering can be fiddly

Toy dataset 'moons'





# 3. Cluster “stability”

Three “stabilities”. How much does clustering change when:

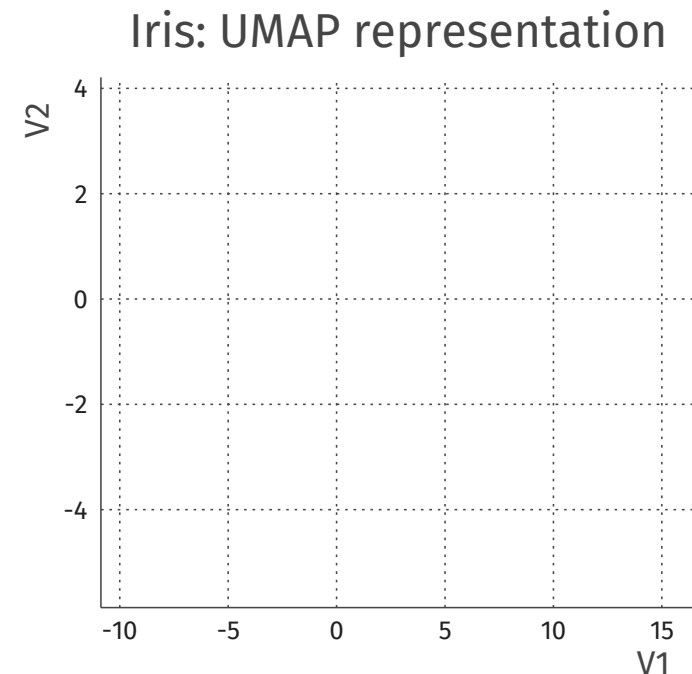
1. Changing some **hyperparameters** (distance metric, linkage, K, ...)
2. Changing some **observations** (bootstrapping, [Hennig, 2007](#))
3. Changing some **features**

Check if observations are classified into same cluster across choices  
e.g. using Jaccard index, Rand index



### 3. Clusterwise stability in R

```
library(fpc)
data(iris)
clusterboot(iris[, 1:4],
             clustermethod = hclustCBI,
             method = "complete", k = 3)
```



Clusterwise Jaccard bootstrap (omitting multiple points) mean:

```
[1] 0.891 0.459 0.719
```

## 4. Internal validation indices

- Only look at “unsupervised” bit: data and clustering
- Quantify how “successful” the clustering is in some sense
- Popular **measures**:
  - Average silhouette width (**ASW**)      (*how close are points to other clusters*)
  - “Gap statistic”      (*Tibshirani et al. 2001*)
  - (measures from model-based clustering → tomorrow)

*Disadvantage:* don't take account of the clustering **aim**!



# Silhouette analysis in R

```
distmat_faithful <- dist(faithful)
hclust_faithful <- hclust(distmat_faithful)

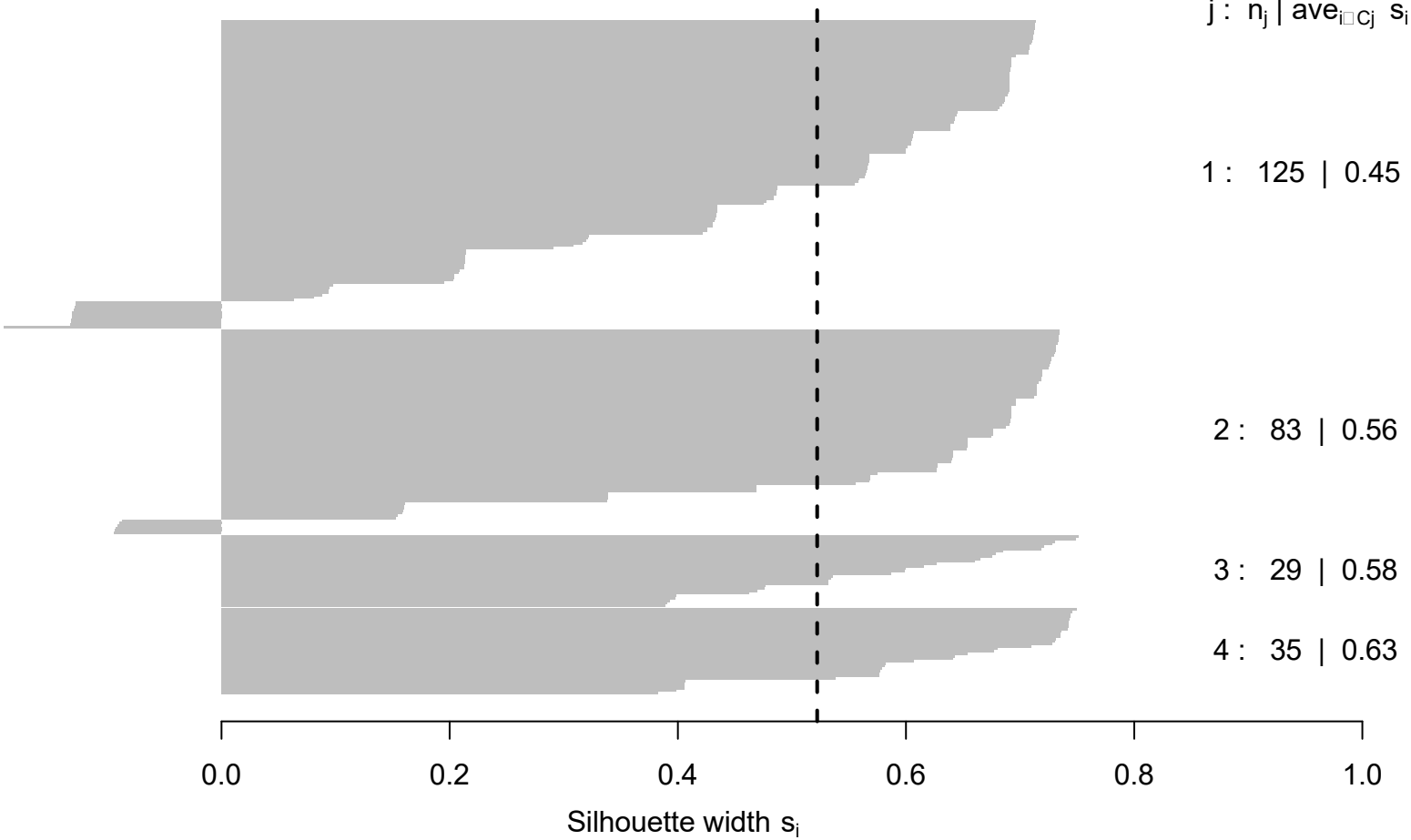
clustering_faithful <- cutree(hclust_faithful, 2)
silhouette_scores <- silhouette(clustering_faithful, distmat_faithful)

plot(silhouette_scores)
```

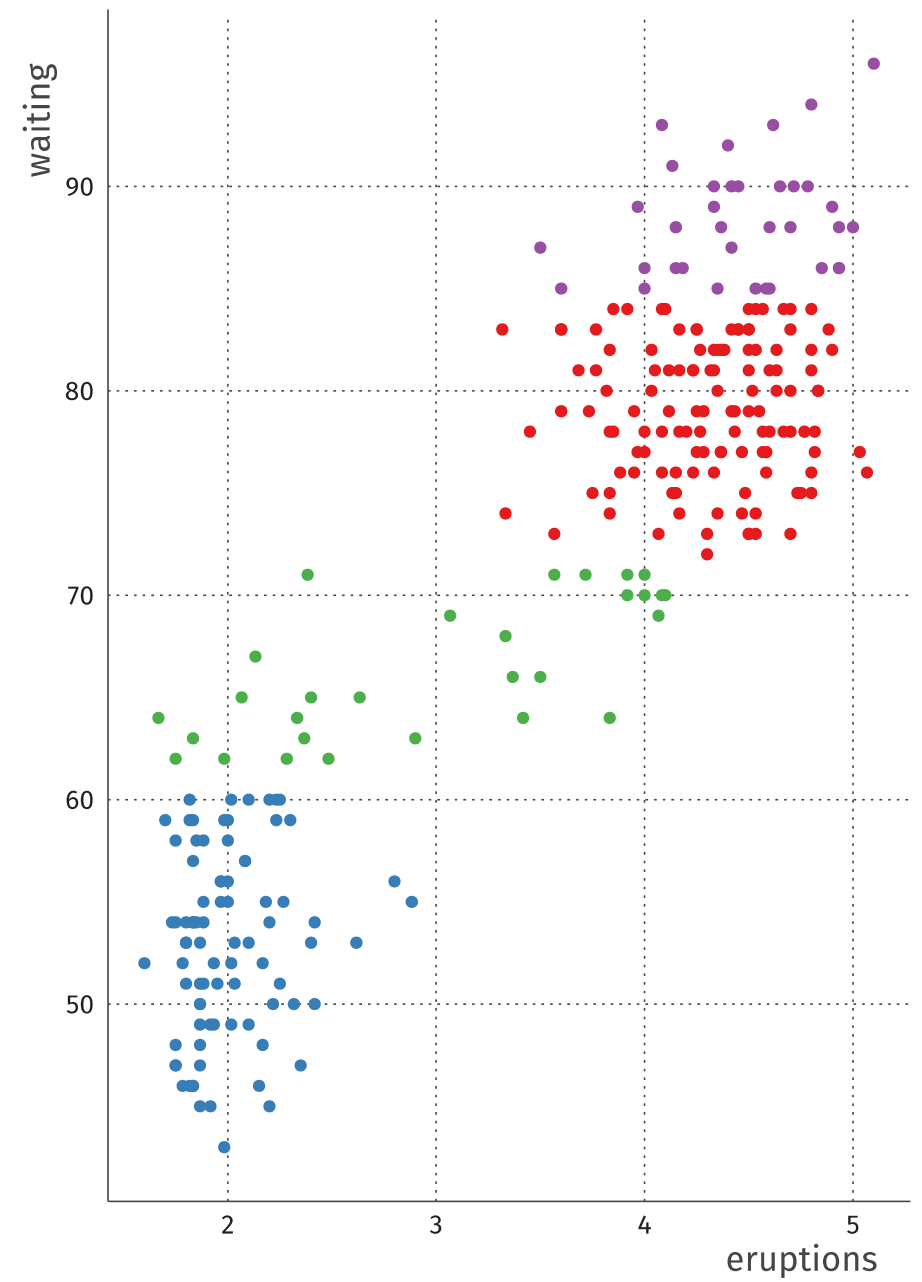
*Note that this works for any type of clustering!*

Silhouette plot of (x = clustering\_faithful, dist = distmat\_faithful)

n = 272

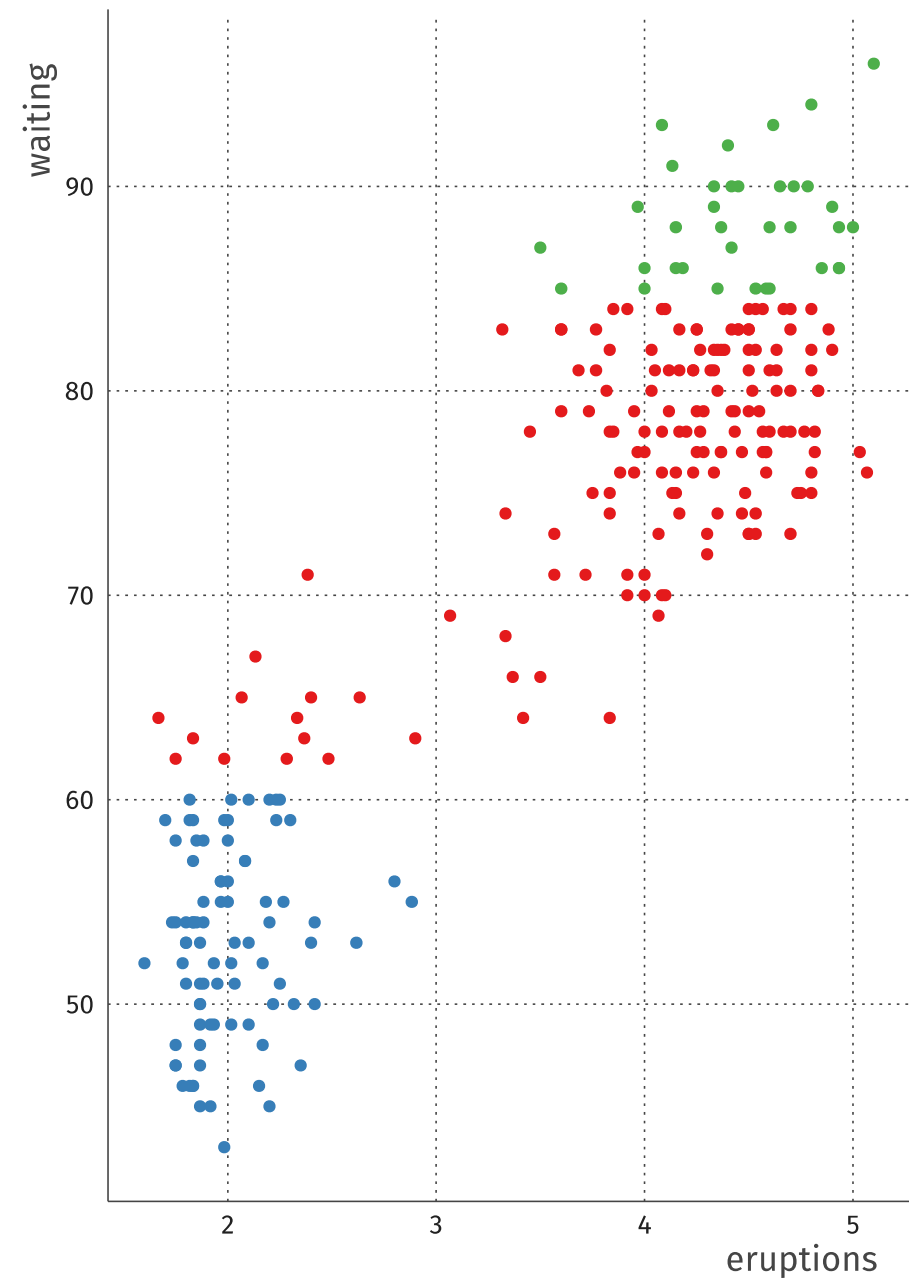
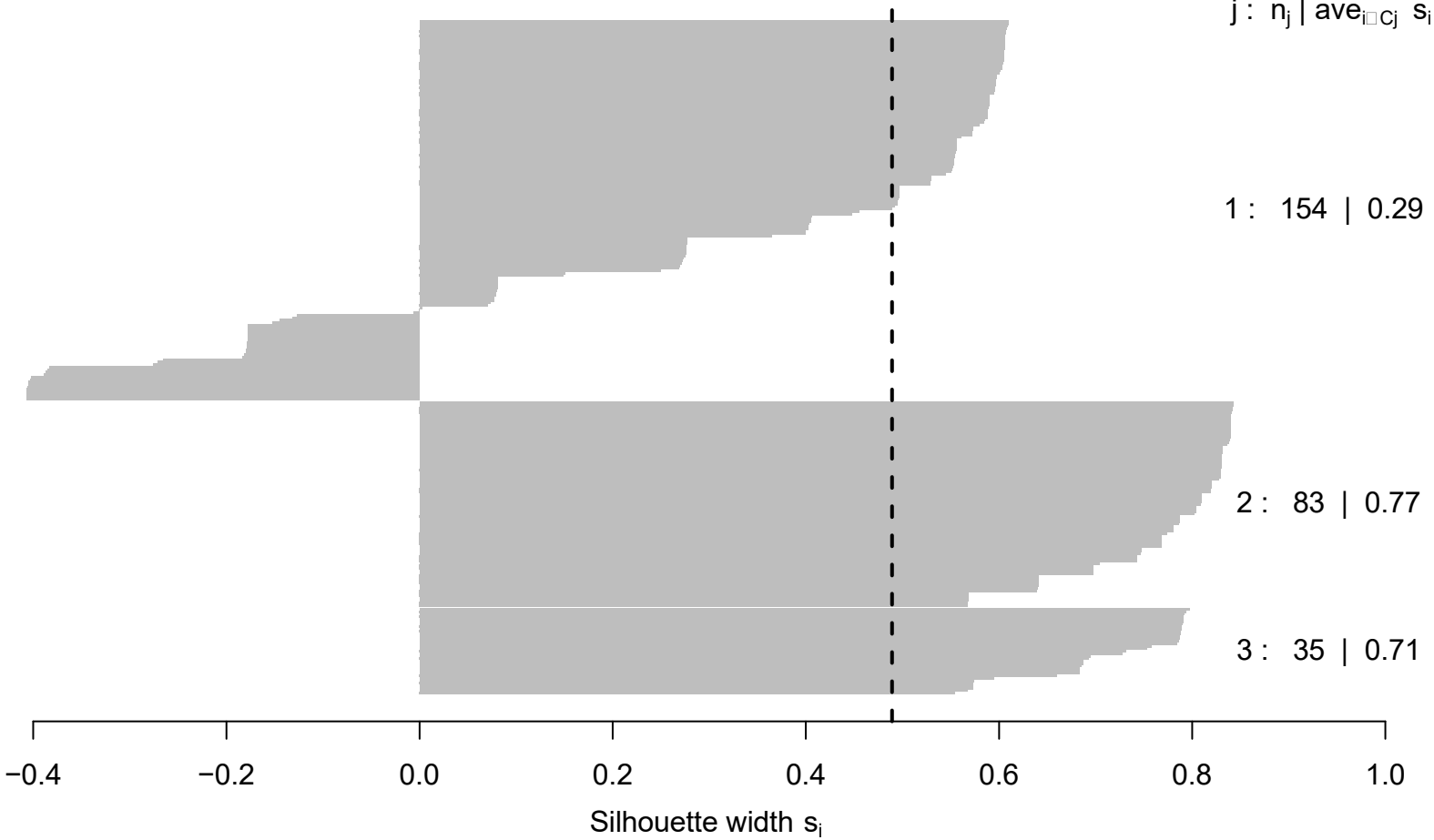


Average silhouette width : 0.52



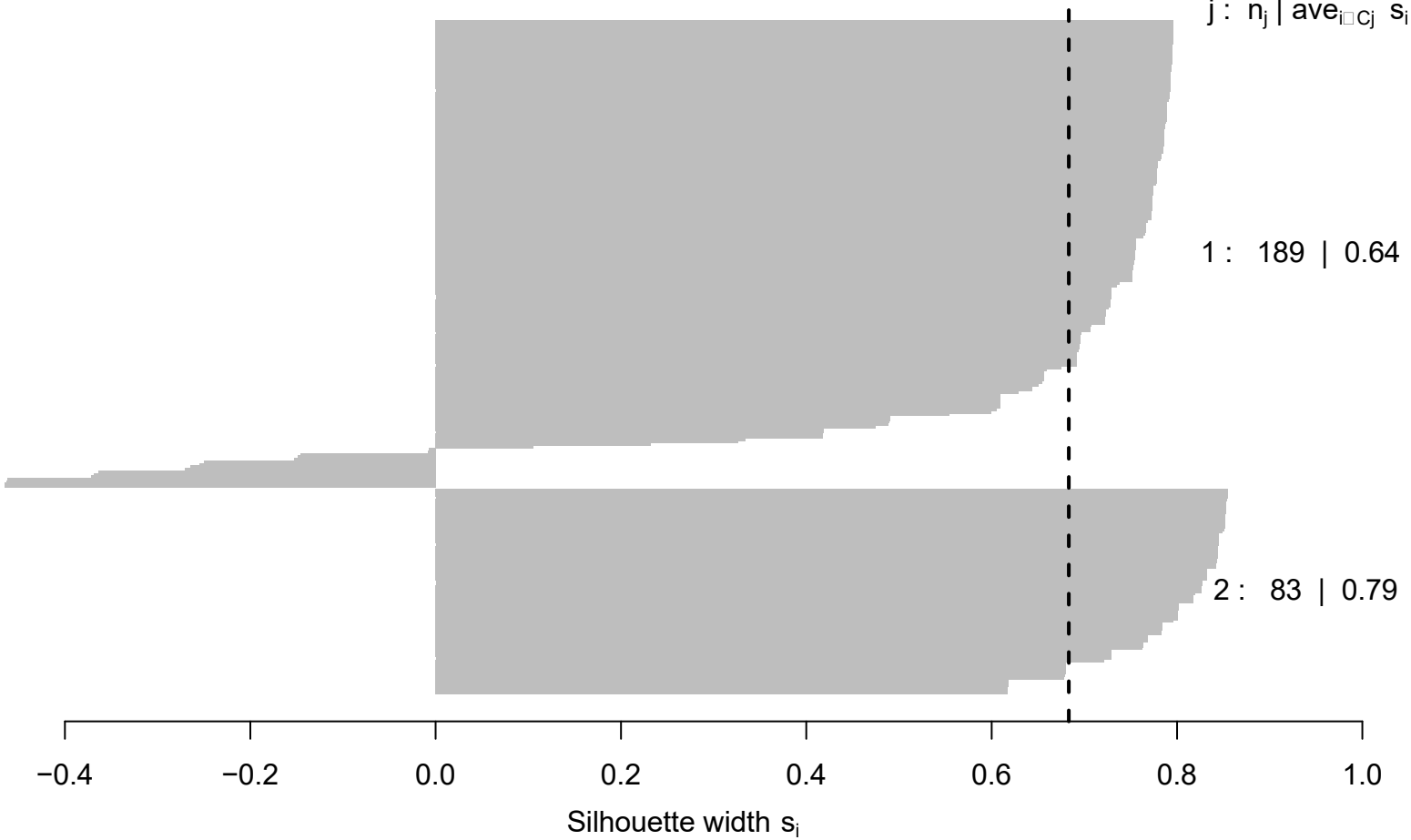
Silhouette plot of (x = clustering\_faithful, dist = distmat\_faithful)

n = 272

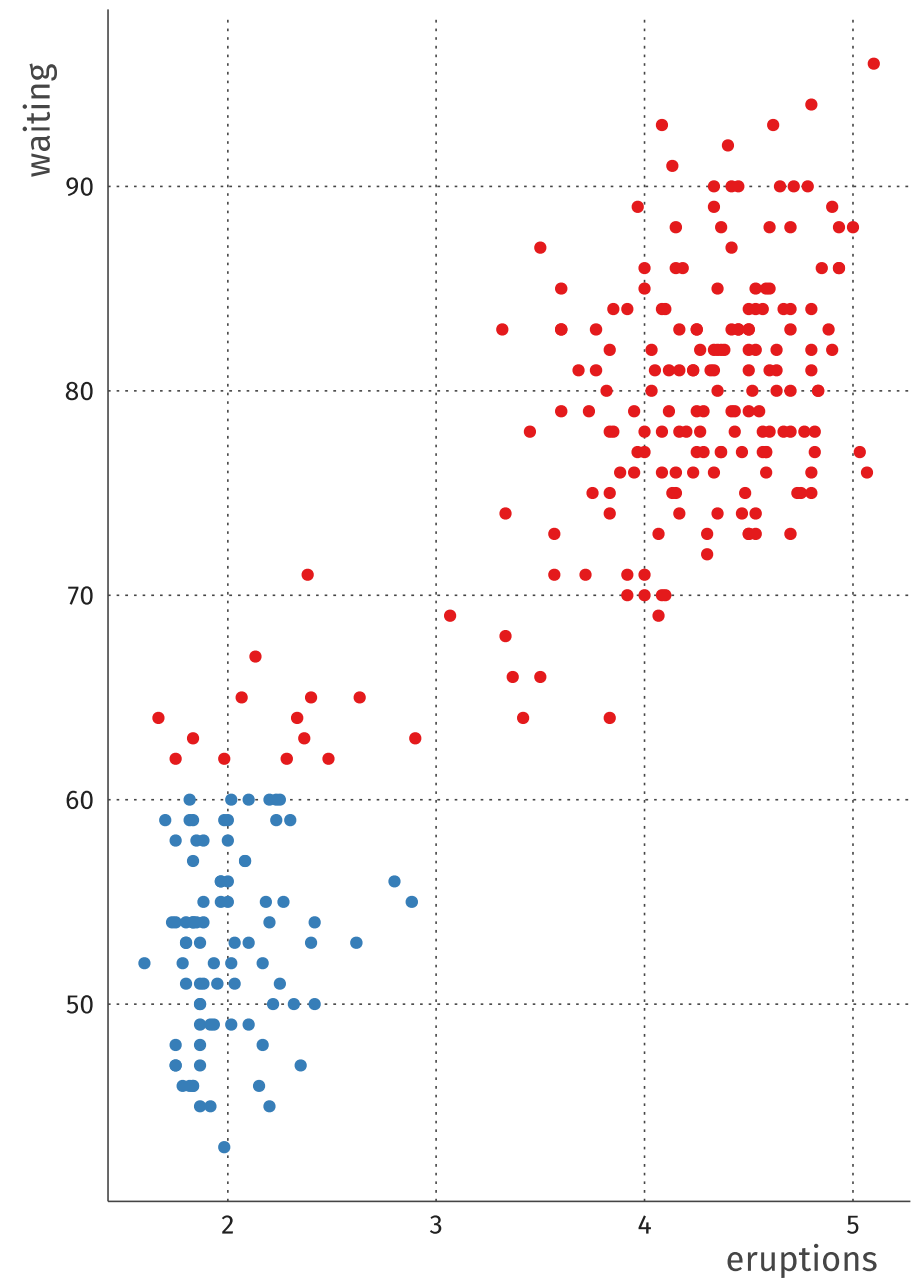


Silhouette plot of (x = clustering\_faithful, dist = distmat\_faithful)

n = 272



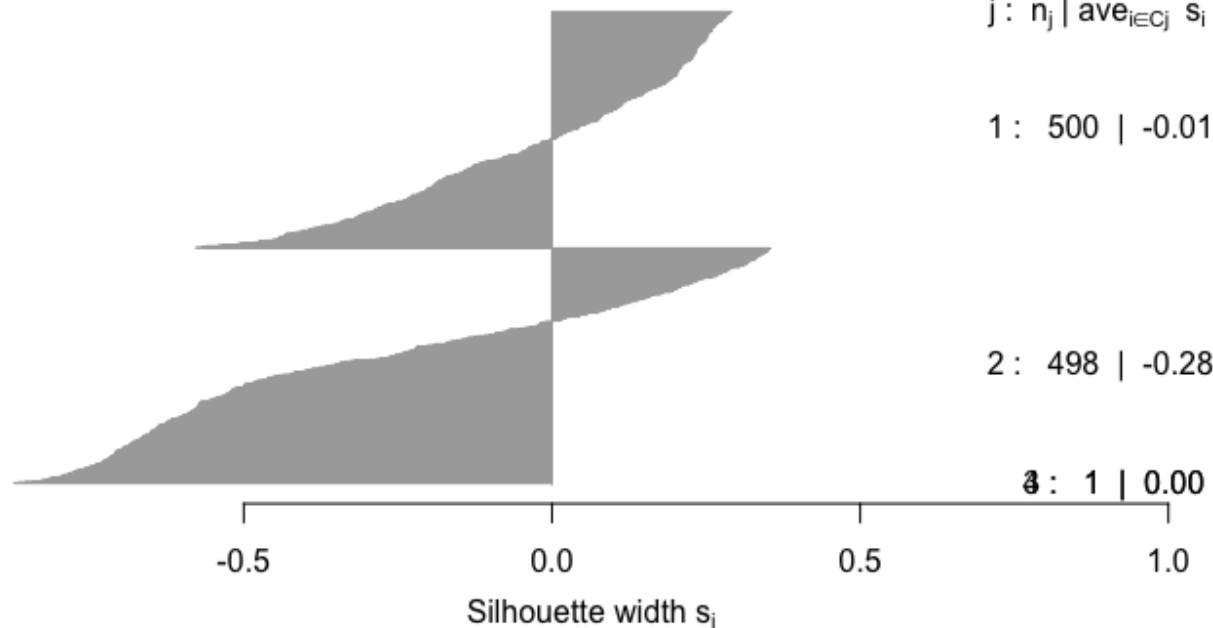
Average silhouette width : 0.68



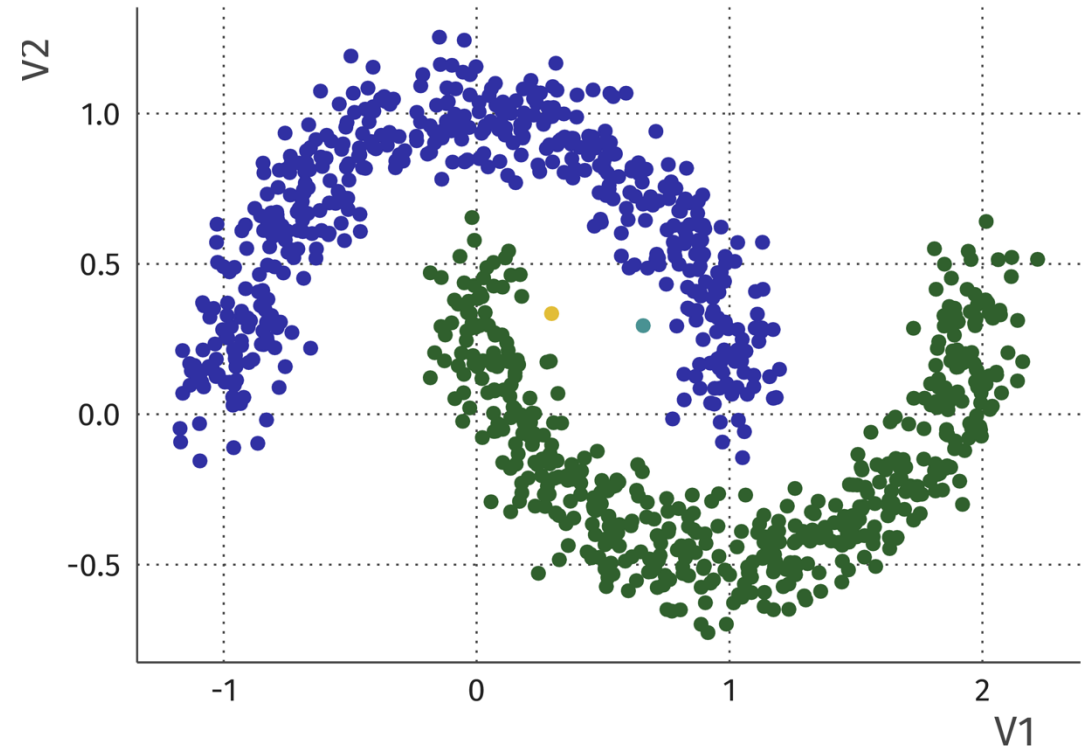
What do you think the average silhouette width (ASW) of this solution will approximately be?

Silhouette plot of (x = clus\_hclust, dist = dist(x))

n = 1000



hclust [single linkage, 4 cluster cutoff]



# Conclusion: clustering

- Clustering looks for “similar” groups of observations
- Two basic clustering methods:
  1. **Hierarchical** clustering (e.g. bottom-up agglomerative, top-down divisive, ...)
  2. **Partitional** clustering (e.g. k-means, DBSCAN, ...).
- Cluster evaluation is an important and subtle topic;
- No way to get rid of critical thought.

