

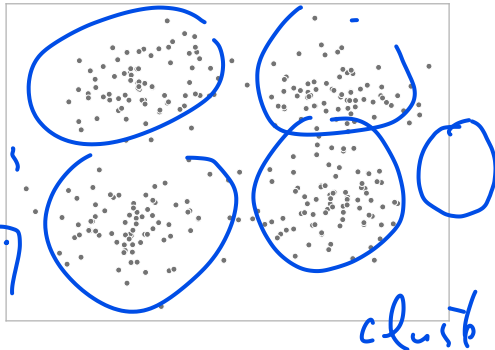
Unsupervised learning: no labels, just structure

Until now every task had a target y (rent, “bad cheese”, a digit). Now we get **only the features x** — no labels — and ask:

Are there natural groups in this data?

- ▶ customer segments (marketing)
- ▶ topics in a pile of documents
- ▶ colors in an image
- ▶ cell types in gene-expression data

Unlabeled data: how many groups? which?



Clustering: partition the data into groups so points in a group are similar, and points in different groups are not.

Where clustering sits

Supervised (had y)

- ▶ regression, classification
- ▶ learn a map $\mathbf{x} \rightarrow y$

Unsupervised (no y)

- ▶ **clustering** — group the *rows* (samples)
- ▶ dimensionality reduction — compress the *columns* (features) [next deck]

You can have labels and still cluster:

- ▶ find *sub-structure* inside a labeled class;
- ▶ QA the labels / spot mislabeled points;
- ▶ build cluster-id *features* for a supervised model;
- ▶ semi-supervised setups.

(That is also when the *external* metrics later apply.)

Outline

K-means ✓

From hard to soft: GMM & EM

Medoids: a real point at the center

Density-based clustering

Hierarchical clustering

Evaluation without labels ✓

Evaluation against labels ✓

In practice and beyond ✓

color RGB ✓
Hue Saturation



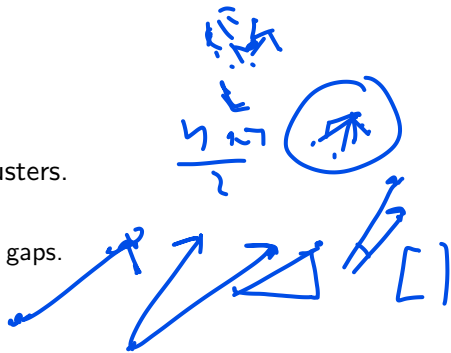
K-means

The workhorse: pick k , alternate “assign then update” until the centers stop moving.

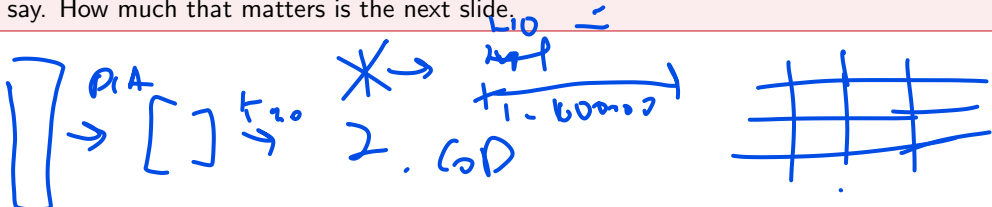
First: how do we measure “similar”?

Clustering needs a distance. The choice shapes the clusters.

- ▶ Euclidean $\|x - x'\|_2$ — the default (k-means uses it).
- ▶ Manhattan $\|x - x'\|_1$ — robust to large single-feature gaps.
- ▶ Cosine — angle, not magnitude (text / embeddings).



Trap — the units decide. Every distance above sums features in whatever units they arrive in. A monthly spend in drams (10,000s) and a visit count (single digits) do *not* get an equal say. How much that matters is the next slide.

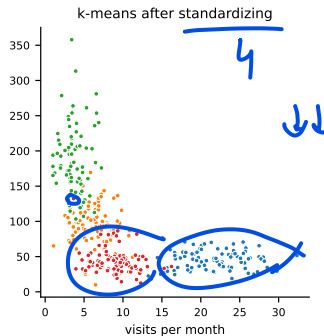
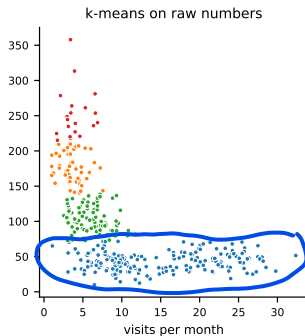
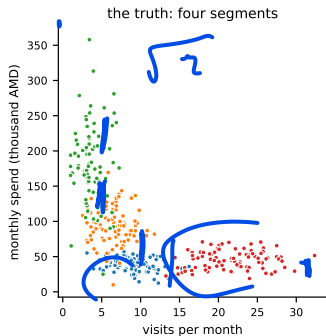


Predict-first: do the units change the answer?

Supermarket loyalty data — **age** (years), **monthly spend** (drams), **visits** per month, **share ordered online** (%) — with four real customer segments in it. Run k-means, $k = 4$: once on the raw columns, once after standardizing. **Same clusters?**

Predict-first: do the units change the answer?

Supermarket loyalty data — age (years), monthly spend (drams), visits per month, share ordered online (%) — with four real customer segments in it. Run k-means, $k = 4$: once on the raw columns, once after standardizing. **Same clusters?**



Spend has a standard deviation of **63,790**; visits, **7.3**. So raw k-means is clustering on spend and nothing else, and returns **spend bands** — one band swallows the students and the daily regulars together, and the big spenders get split in two. **Standardize first**, unless the features already share units.

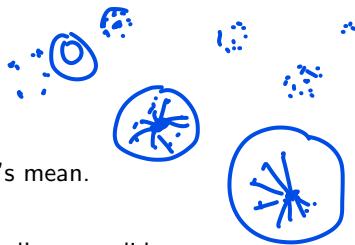
K-means: the objective

cent.

R_i, μ_j

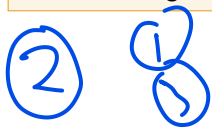
Choose k . Represent each cluster by a centroid μ_j . Minimize the within-cluster sum of squares (inertia):

$$\text{WCSS} = \sum_{j=1}^k \sum_{\mathbf{x} \in C_j} \|\mathbf{x} - \mu_j\|^2.$$



- ▶ Each point joins its nearest centroid; the centroid is the cluster's mean.
- ▶ You must pick k in advance (more on that soon).
- ▶ The name is MacQueen's (1967); the standard algorithm is Lloyd's, next slide.

Inertia is not normalized. It falls as k grows and scales with the data, so its absolute value means nothing — only useful for *relative* comparison (e.g. the elbow).



1000' → k → 1000' 5 5

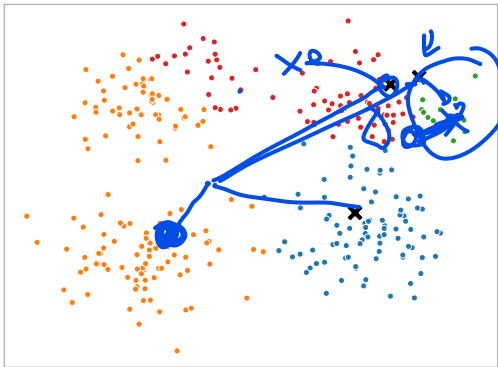


Lloyd's algorithm (1957, published 1982): assign, update, repeat

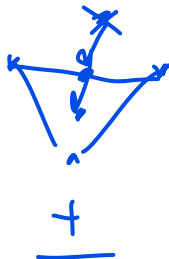
Two steps, each of which can only *lower* the WCSS (this is coordinate descent):
assign each point to its nearest centroid, then **update** each centroid to its cluster's mean. Watch it converge:

734

k-means --- initial centroids



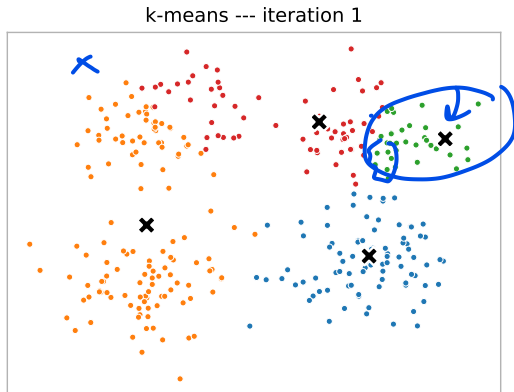
4



Converges — but only to a **local** optimum (Gaussian mixtures alternate the same way).

Lloyd's algorithm (1957, published 1982): assign, update, repeat

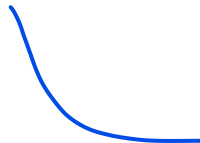
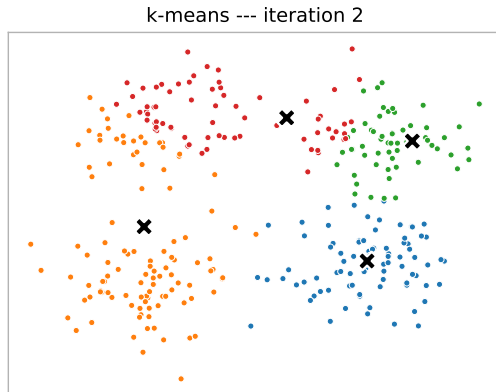
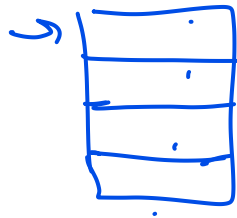
Two steps, each of which can only *lower* the WCSS (this is coordinate descent):
assign each point to its nearest centroid, then **update** each centroid to its cluster's mean. Watch it converge:



Converges — but only to a **local** optimum (Gaussian mixtures alternate the same way).

Lloyd's algorithm (1957, published 1982): assign, update, repeat

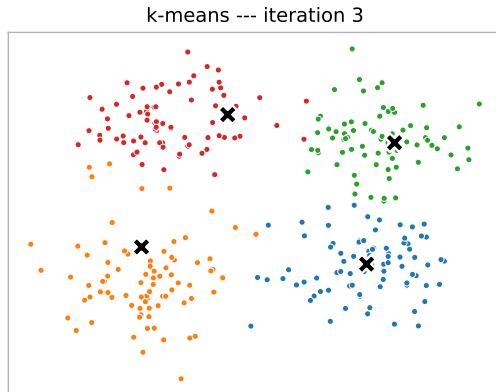
Two steps, each of which can only *lower* the WCSS (this is coordinate descent):
assign each point to its nearest centroid, then **update** each centroid to its cluster's mean. Watch it converge:



Converges — but only to a **local** optimum (Gaussian mixtures alternate the same way).

Lloyd's algorithm (1957, published 1982): assign, update, repeat

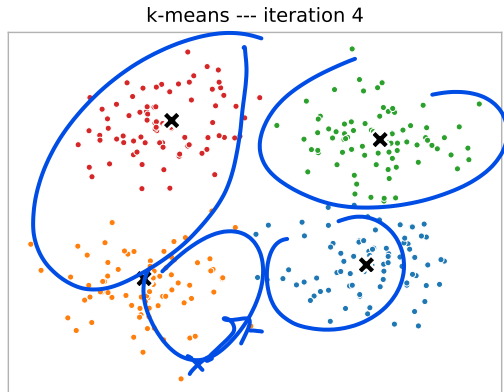
Two steps, each of which can only *lower* the WCSS (this is coordinate descent):
assign each point to its nearest centroid, then **update** each centroid to its cluster's mean. Watch it converge:



Converges — but only to a **local** optimum (Gaussian mixtures alternate the same way).

Lloyd's algorithm (1957, published 1982): assign, update, repeat

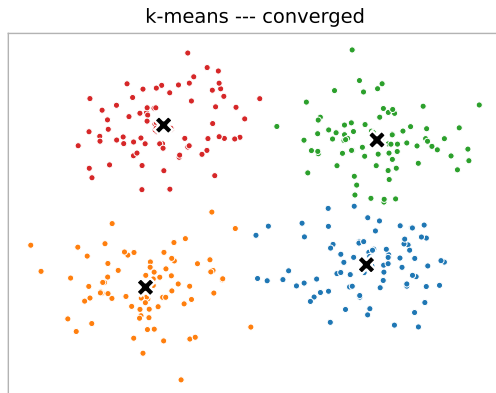
Two steps, each of which can only *lower* the WCSS (this is coordinate descent):
assign each point to its nearest centroid, then **update** each centroid to its cluster's mean. Watch it converge:



Converges — but only to a **local** optimum (Gaussian mixtures alternate the same way).

Lloyd's algorithm (1957, published 1982): assign, update, repeat

Two steps, each of which can only *lower* the WCSS (this is coordinate descent):
assign each point to its nearest centroid, then **update** each centroid to its cluster's mean. Watch it converge:



Converges — but only to a **local** optimum (Gaussian mixtures alternate the same way).

Worked iteration (by hand), $k = 2$

Points

$P_1(1, 1)$, $P_2(2, 1)$, $P_3(4, 3)$, $P_4(5, 4)$, $P_5(2, 2)$; init centroids (placed freely, *not* on data points)

$\mu_1 = (1.5, 1)$, $\mu_2 = (4, 4)$.

Assign (nearest centroid): $P_1, P_2, P_5 \rightarrow \mu_1$;

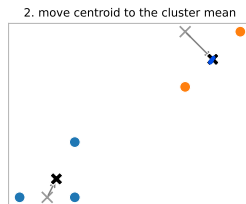
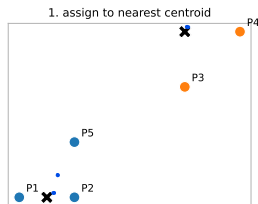
$P_3, P_4 \rightarrow \mu_2$.

Update (cluster means):

$$\mu_1 = \frac{(1,1)+(2,1)+(2,2)}{3} = (1.67, 1.33)$$

$$\mu_2 = \frac{(4,3)+(5,4)}{2} = (4.5, 3.5)$$

Then repeat until nothing moves.



Predict-first: does initialization matter?

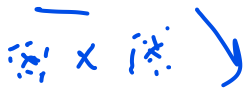
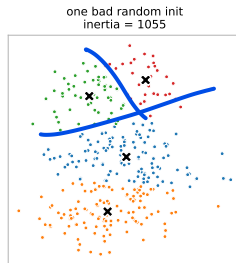
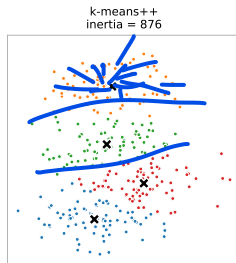
Run k-means twice from two *different* random starts. Same answer?

Predict-first: does initialization matter?

Run k-means twice from two *different* random starts. Same answer?

No. It only finds a *local* optimum — a bad start gives a worse clustering (higher inertia).

Fix: k-means++ (Arthur & Vassilvitskii, 2007) picks each new centroid with probability proportional to its **squared distance** from the nearest centroid already chosen — so the seeds spread out. Then run it several times (`n_init`) and keep the lowest-inertia result.



WCS
MSF



Choosing k : the elbow

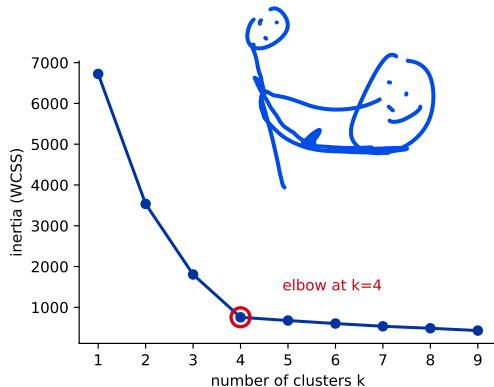
Plot inertia vs k . It always decreases; look for the **elbow** where extra clusters stop paying off.

- ▶ Below the elbow: real structure being captured.
- ▶ Past it: diminishing returns (splitting real clusters).

The elbow is a heuristic — often fuzzy. The silhouette (evaluation section) gives a more principled choice.



Diminishing ↙

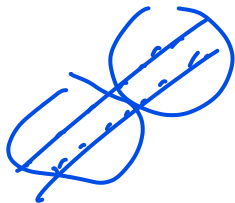


K

known

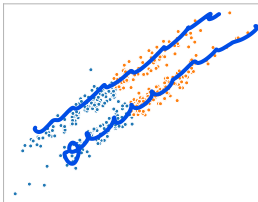


Assumptions = failure modes

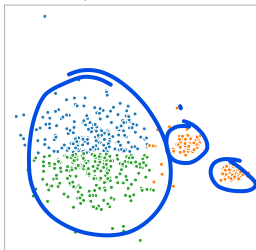


k-means forces round, balanced clusters

elongated clusters



unequal sizes / variances



K-means assumes **round, similarly-sized, convex** clusters — it breaks on elongated shapes and on clusters of very different size/variance.

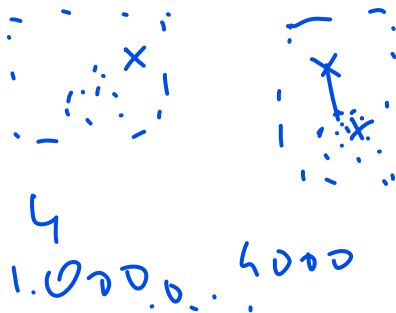
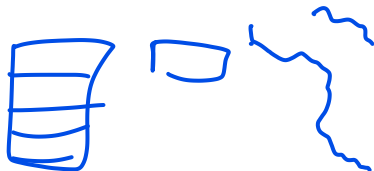
Reach for k-means when: large n , round/equal clusters, you can pick k , fast baseline.

Mini-batch k-means: scaling up

For very large n , update centroids from small random mini-batches instead of the full data each step (Sculley, 2010).

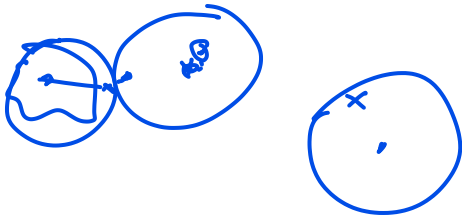
- ▶ Much faster, far less memory; streaming-friendly.
- ▶ Slightly worse clusters than full k-means — usually a fine trade.

When: same problem as k-means, but n is huge.



From hard to soft

K-means says “you ARE in cluster A”. A Gaussian mixture says “70% A, 30% B” — and allows elliptical clusters.



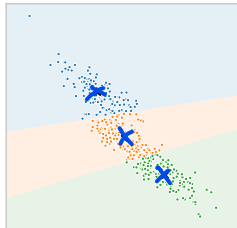
Gaussian mixtures: soft, elliptical clusters

A Gaussian Mixture Model (GMM) treats each cluster as a **Gaussian** with its own mean *and* covariance Σ_k , so clusters can be **elliptical** and different sizes — those ellipses are the Mahalanobis contours you met in LDA/QDA. Each point gets a **soft** membership (responsibility) in every cluster.

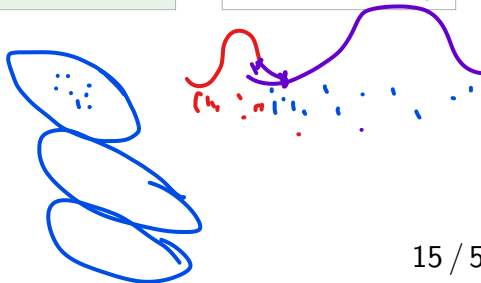
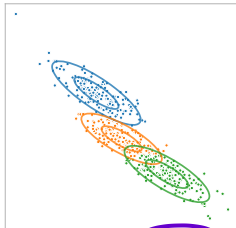
K-means is the special case of a GMM with *spherical, equal-variance* covariances and *hard* (winner-take-all) assignments.

When: overlapping / elliptical clusters, soft memberships, or density estimation.

k-means: round, hard (Voronoi)



GMM: elliptical, soft



EM, step 1: the E-step (responsibilities)

Chicken-and-egg: we need assignments to fit the Gaussians, but the Gaussians to assign points. **Expectation–Maximization (EM; Dempster, Laird & Rubin, 1977)** breaks it by **alternating**. Start from random Gaussians, then:

E-step — **soft-assign every point**. For point i and cluster k , the **responsibility** is the posterior probability it belongs to k :

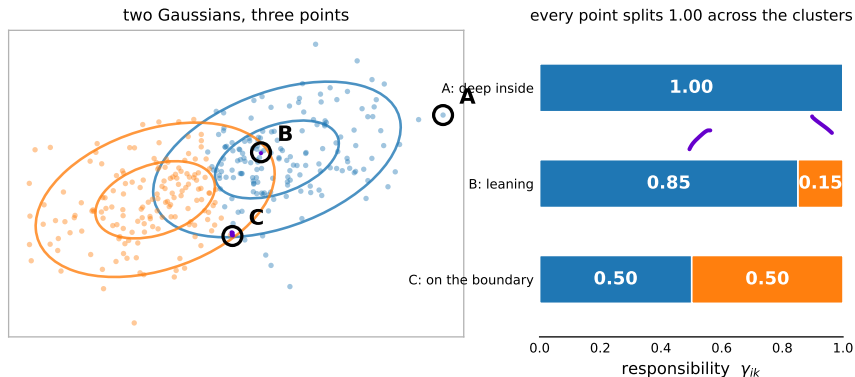
$$\gamma_{ik} = \frac{\pi_k \mathcal{N}(\mathbf{x}_i \mid \boldsymbol{\mu}_k, \boldsymbol{\Sigma}_k)}{\sum_{j=1}^K \pi_j \mathcal{N}(\mathbf{x}_i \mid \boldsymbol{\mu}_j, \boldsymbol{\Sigma}_j)}.$$

- ▶ π_k is the cluster's weight (its prior size); $\mathcal{N}$ is the Gaussian density.
- ▶ $\gamma_{ik} \in [0, 1]$ and $\sum_k \gamma_{ik} = 1$ — a soft membership, not a hard label.

Compare k-means: its “assign” step is the *hard* version — $\gamma_{ik} \in \{0, 1\}$.

What a responsibility actually looks like

Three points, the same two Gaussians. Every point's responsibilities sum to 1 — but *how* that 1 splits is the whole difference from k-means:



K-means answers all three the same way: a hard 1 for the nearer centroid. C is a coin flip and A is not — hard assignment throws that difference away.

EM, step 2: the M-step, and convergence

M-step — **refit each Gaussian**, weighting every point by its responsibility. With $N_k = \sum_i \gamma_{ik}$ (the soft count in cluster k):

$$\mu_k = \frac{1}{N_k} \sum_i \gamma_{ik} \mathbf{x}_i, \quad \Sigma_k = \frac{1}{N_k} \sum_i \gamma_{ik} (\mathbf{x}_i - \mu_k)(\mathbf{x}_i - \mu_k)^\top, \quad \pi_k = \frac{N_k}{n}.$$

Each is the ordinary mean / covariance / proportion, just *responsibility-weighted*.

- **Repeat** E $\rightarrow$ M until the log-likelihood stops rising — this is **maximum likelihood estimation**, the same objective you maximised for logistic regression. Each round increases it $\Rightarrow$ convergence to a **local** optimum (so, like k-means, run several inits).

Pick the number of components with the **Bayesian** or **Akaike Information Criterion (BIC / AIC)** — likelihood criteria (they need a probabilistic model, so they work for GMM, *not* for plain k-means). EM recurs across ML wherever there are latent variables.





1.5 σ

Beyond the mean

Keep the k-means loop, but let the center be an *actual data point* — and then any distance you like becomes legal.



...



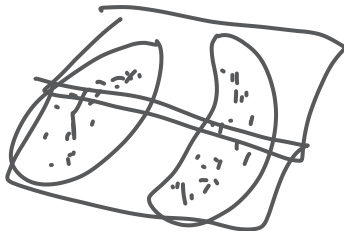
K-medoids — Partitioning Around Medoids (PAM)

PAM (Kaufman & Rousseeuw, 1987) works like k-means, but the center is an **actual data point** — a **medoid** — chosen to minimize the total distance to its cluster (it never *averages* points).

- ▶ **Any distance / dissimilarity** works (cosine, Gower for mixed types, edit distance. . .) — because it never has to compute a mean.
- ▶ **Robust to outliers** — a real central point, not a mean dragged by extremes.
- ▶ **Interpretable center** — “this customer is typical of the segment”.

The real difference from k-means is the *medoid* and the *free choice of metric* — not the exact objective. *When*: non-Euclidean data, outliers, interpretable centers. *Not in core* sklearn — sklearn_extra.cluster.KMedoids.

0. 0-1000
1.
2.



Density-based clustering

Stop asking “which center is nearest?” and start asking “is it crowded here?”
— any shape, no k , and outliers for free.



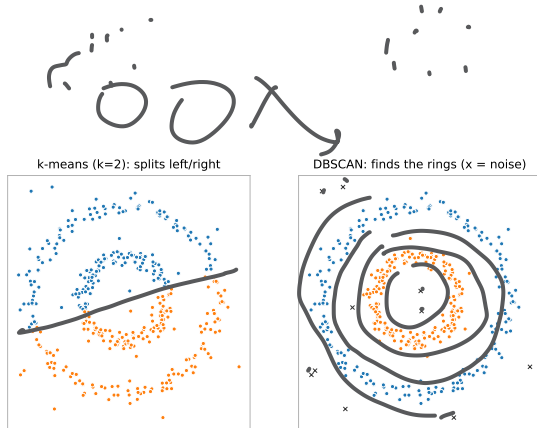
DBSCAN: the three kinds of point

Density-Based Spatial Clustering of Applications with Noise (DBSCAN;
Ester et al., 1996): clusters are **dense regions** separated by sparse ones. Two knobs: radius ϵ (ϵ) and `min_samples`.

- ▶ core: $\geq \text{min_samples}$ points within ϵ ;
- ▶ border: within ϵ of a core, but not dense itself;
- ▶ noise: neither — an outlier.

Finds arbitrary shapes, needs **no** k , and **flags outliers**.

Predict-first: what can DBSCAN do that k-means/GMM cannot?



DBSCAN: the three kinds of point

Density-Based Spatial Clustering of Applications with Noise (DBSCAN; Ester et al., 1996): clusters are **dense regions** separated by sparse ones. Two knobs: radius ϵ (ϵ) and min_samples .

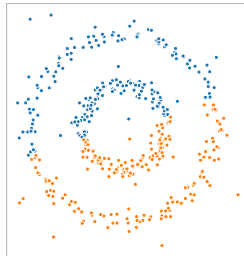
- **core**: $\geq \text{min_samples}$ points within ϵ ;
- **border**: within ϵ of a core, but not dense itself;
- **noise**: neither — an outlier.

Finds **arbitrary shapes**, needs **no** k , and **flags outliers**.

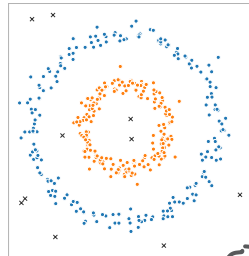
Predict-first: what can DBSCAN do that k-means/GMM cannot? *Non-convex shapes, and automatic noise/outlier flagging.*



k-means ($k=2$): splits left/right



DBSCAN: finds the rings (x = noise)



min_samples

3

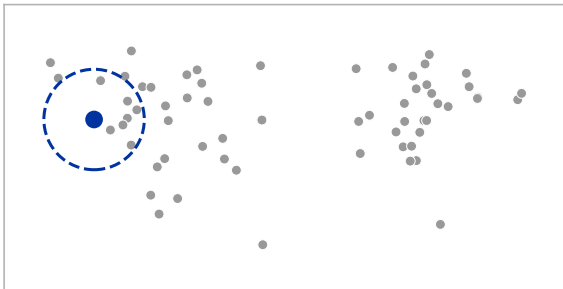
ϵ



DBSCAN, step by step

A core point “grows” a cluster by absorbing everything density-reachable from it:

1. pick a point: count neighbors within ϵ

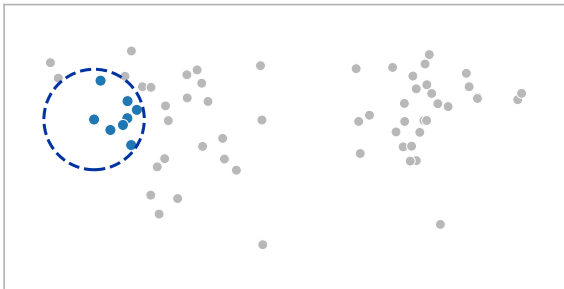


Core points chain together into a cluster; border points attach; whatever is left in sparse regions is **noise** (-1).

DBSCAN, step by step

A core point “grows” a cluster by absorbing everything density-reachable from it:

2. $\geq \text{min_samples}$ $\rightarrow$ core; neighbors join

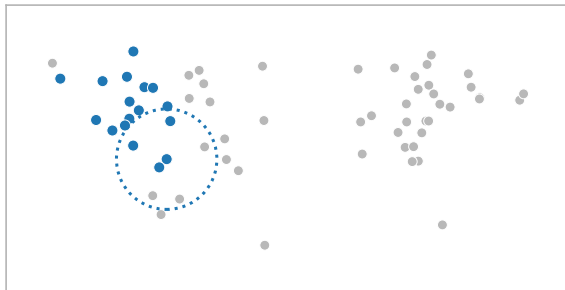


Core points chain together into a cluster; border points attach; whatever is left in sparse regions is **noise** (-1).

DBSCAN, step by step

A core point “grows” a cluster by absorbing everything density-reachable from it:

3. expand: each core pulls in its neighbors

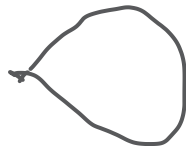
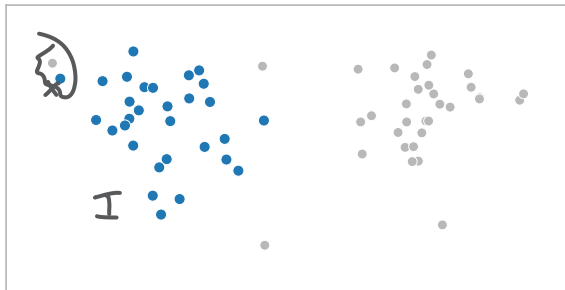


Core points chain together into a cluster; border points attach; whatever is left in sparse regions is **noise** (-1).

DBSCAN, step by step

A core point “grows” a cluster by absorbing everything density-reachable from it:

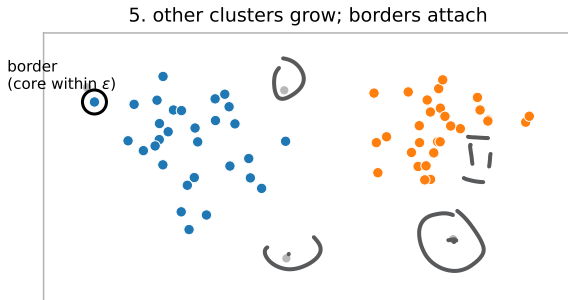
4. cluster 0 complete (density-connected)



Core points chain together into a cluster; border points attach; whatever is left in sparse regions is **noise** (-1).

DBSCAN, step by step

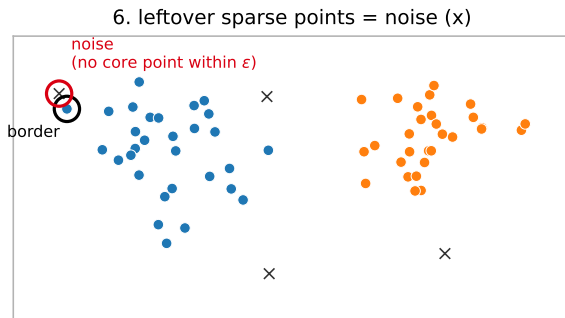
A core point “grows” a cluster by absorbing everything density-reachable from it:



Core points chain together into a cluster; border points attach; whatever is left in sparse regions is **noise** (-1).

DBSCAN, step by step

A core point “grows” a cluster by absorbing everything density-reachable from it:

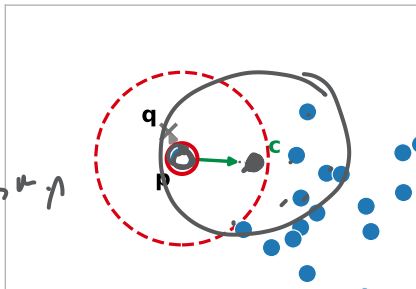


Core points chain together into a cluster; border points attach; whatever is left in sparse regions is **noise** (-1).

Border or noise? The rule that gets missed

Those two top-left points sit 0.22 apart and **neither** has enough neighbours to be core. One still joins the cluster; the other is an outlier:

BORDER: a core point is in reach
3 points in its ϵ -ball (needs 4)



$d(p, q) = 0.55 \leq \epsilon$: a core is inside
 $\Rightarrow p$ joins the cluster

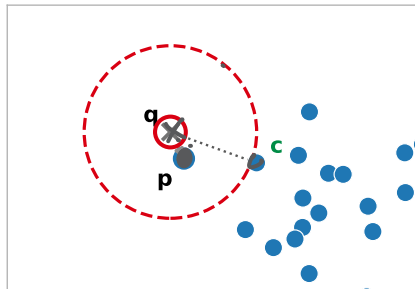
● core

○ border

× noise

--- the point's ϵ -ball

NOISE: only a border point is in reach
2 points in its ϵ -ball (needs 4)



$d(q, p) = 0.69 > \epsilon$: no core is inside
 $\Rightarrow q$ stays noise

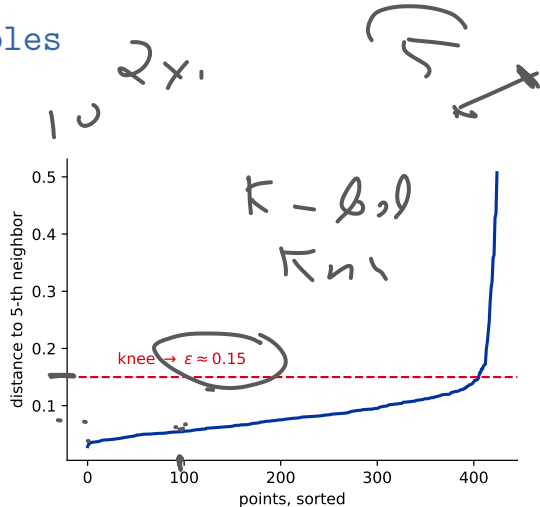
A border point needs a **core** point inside its ϵ -ball. Being next to p is not enough — **reachability flows only through core points.**

DBSCAN: choosing eps and min_samples

min_samples: rule of thumb $\approx 2 \times$ the number of features ($\geq D+1$, and at least 3); raise it for large or noisy data.

eps: make a k-distance plot — sort every point's distance to its k -th neighbor ($k = \text{min_samples}$) and look for the knee. The distance there is a good eps. *That is a k-NN query, and reading the knee is the same trick as the elbow.*

The two are coupled: change one and you usually re-tune the other.



Limits: one global eps can't fit clusters of very different density; the naive algorithm is $O(n^2)$. Sander et al.; see the routine in stats.stackexchange.com/q/88872.

Heuristics:



HDBSCAN: density handled automatically

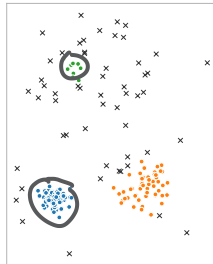
A *single* ϵ ps cannot fit clusters of **different density** — left, the sparse top group is lost to noise.

Hierarchical DBSCAN (HDBSCAN; Campello, Moulavi & Sander, 2013) fixes it. The **H** is the whole idea: rather than one ϵ , build a **hierarchy** over all of them.

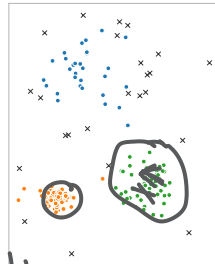
- ▶ it *scans every density level itself* and keeps the clusters that **persist** (are most stable);
- ▶ **no** ϵ ps — you set only `min_cluster_size`;
- ▶ so each cluster is dense at **its own scale**.

Core sklearn since 1.3
(`sklearn.cluster.HDBSCAN`).

DBSCAN, one ϵ ps=0.4 -> 3 clusters
(sparse group lost to noise)



HDBSCAN -> 3 clusters
(adapts to each density)



HDBSCAN: how it works (sketch)

1. **Core distance** of a point = distance to its `min_samples`-th neighbor (how dense it is locally).
2. **Mutual reachability distance** between two points = the max of their two core distances and their actual distance — this pushes sparse points apart.
3. Build the **minimum spanning tree** of that distance, then sweep a threshold from low to high to get a **hierarchy** (a tree) of clusters.
4. **Condense** the tree and keep the clusters that **persist** over the widest range of densities (most “stable”) — those are the output; the rest is noise.

The payoff: clusters chosen by *stability*, not by one hand-picked radius.

DBSCAN

Hierarchical clustering

Do not pick k at all. Merge the closest pair over and over, and cut the resulting tree wherever you like afterwards.

A_2



Hierarchical (agglomerative) clustering



Bottom-up — build clusters by repeated merging:

1. Start: every point is its own cluster (n of them).
 2. Find the **two closest** clusters (“closest” = the linkage, defined shortly).
 3. **Merge** them into one.
 4. Repeat until a single cluster remains, recording each merge and its distance.
- ▶ No need to fix k up front — the merge history is a tree you *cut* afterwards.
 - ▶ Gives **nested** structure (taxonomies).

De 4,1
↓

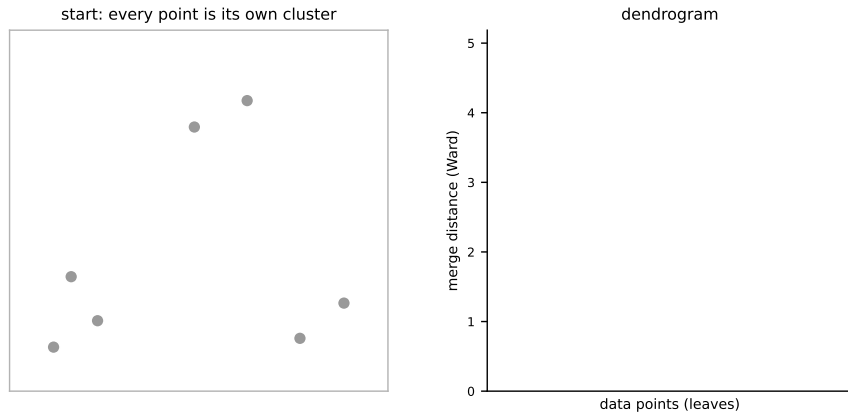
When: nested structure / a dendrogram, exploring #clusters, connectivity constraints; smaller n .



Agglomerative, step by step

Merge the two clusters that cost the least **extra spread** — *Ward* linkage, the common default. *Left*: points and cluster midpoints ($\times$), dashed line = the current merge.

Right: the dendrogram builds up, one bar per merge (latest in red):

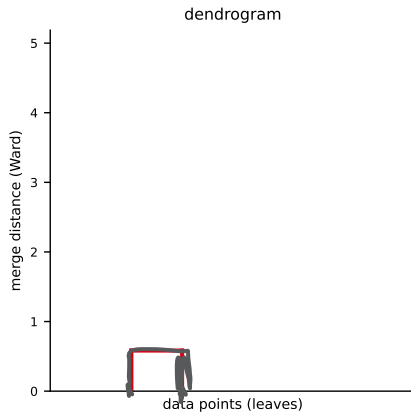
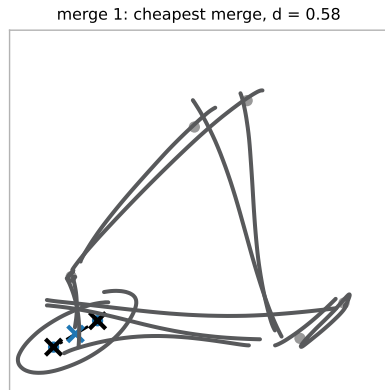


Each merge's cost d is that bar's height. Ward never dips: the tree only grows up.

Agglomerative, step by step

Merge the two clusters that cost the least **extra spread** — *Ward* linkage, the common default. *Left*: points and cluster midpoints ($\times$), dashed line = the current merge.

Right: the dendrogram builds up, one bar per merge (latest in red):

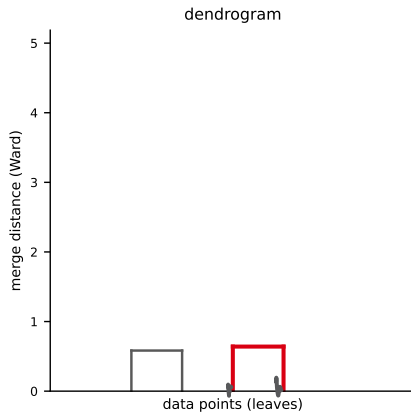
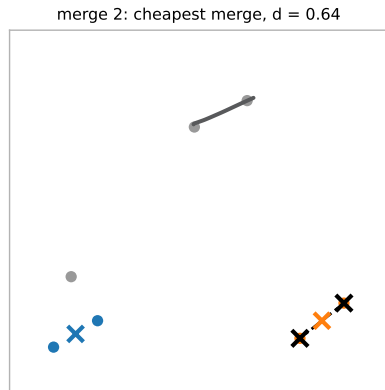


Each merge's cost d is that bar's height. Ward never dips: the tree only grows up.

Agglomerative, step by step

Merge the two clusters that cost the least **extra spread** — *Ward* linkage, the common default. *Left*: points and cluster midpoints ($\times$), dashed line = the current merge.

Right: the dendrogram builds up, one bar per merge (latest in red):

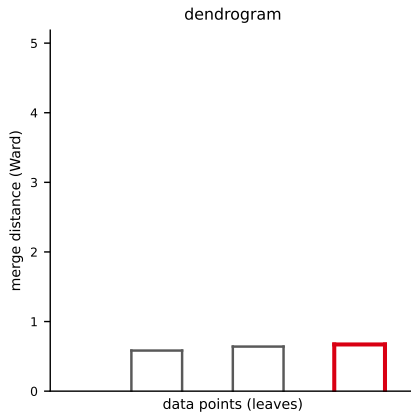
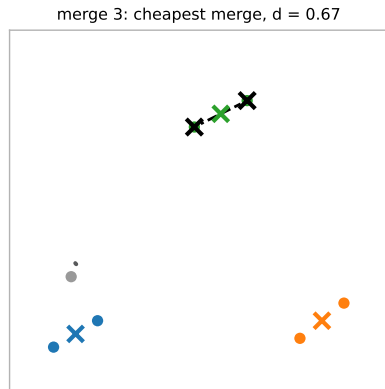


Each merge's cost d is that bar's height. Ward never dips: the tree only grows up.

Agglomerative, step by step

Merge the two clusters that cost the least **extra spread** — *Ward* linkage, the common default. *Left*: points and cluster midpoints ($\times$), dashed line = the current merge.

Right: the dendrogram builds up, one bar per merge (latest in red):

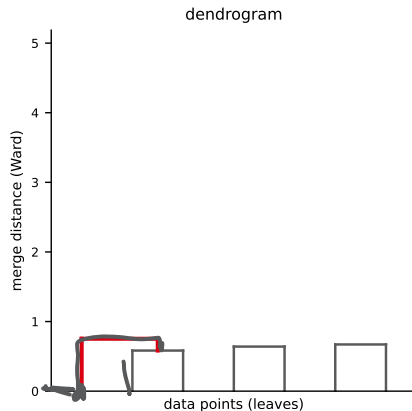
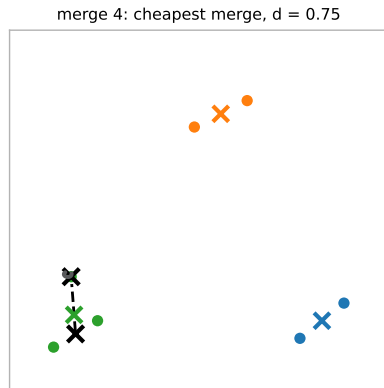


Each merge's cost d is that bar's height. Ward never dips: the tree only grows up.

Agglomerative, step by step

Merge the two clusters that cost the least **extra spread** — *Ward* linkage, the common default. *Left*: points and cluster midpoints ($\times$), dashed line = the current merge.

Right: the dendrogram builds up, one bar per merge (latest in red):

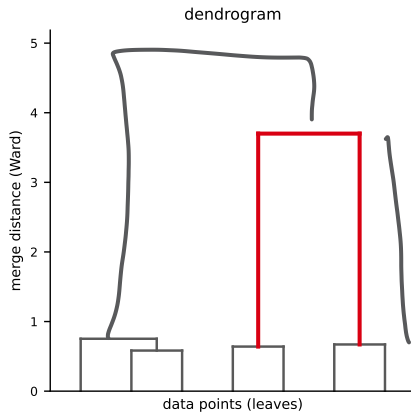
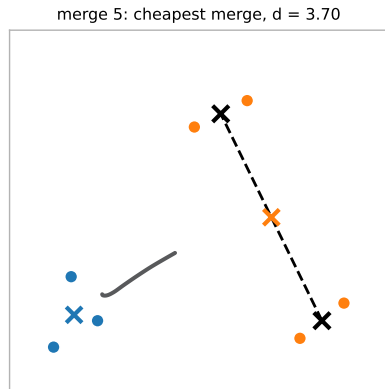


Each merge's cost d is that bar's height. Ward never dips: the tree only grows up.

Agglomerative, step by step

Merge the two clusters that cost the least **extra spread** — *Ward* linkage, the common default. *Left*: points and cluster midpoints ($\times$), dashed line = the current merge.

Right: the dendrogram builds up, one bar per merge (latest in red):

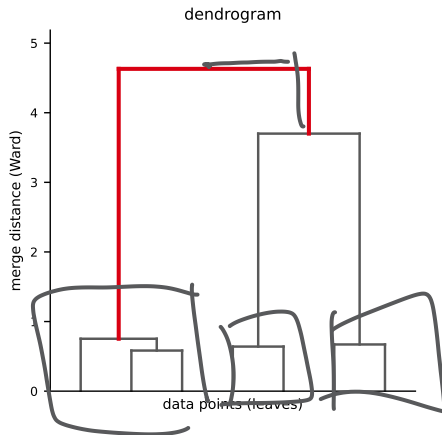
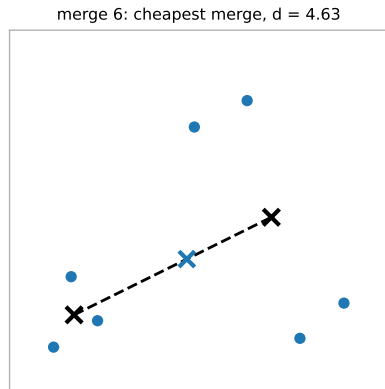


Each merge's cost d is that bar's height. Ward never dips: the tree only grows up.

Agglomerative, step by step

Merge the two clusters that cost the least **extra spread** — *Ward* linkage, the common default. *Left*: points and cluster midpoints ($\times$), dashed line = the current merge.

Right: the dendrogram builds up, one bar per merge (latest in red):

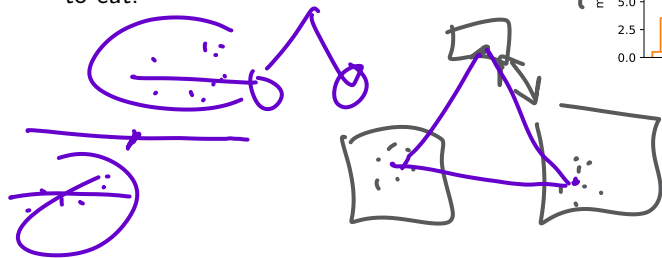


Each merge's cost d is that bar's height. Ward never dips: the tree only grows up.

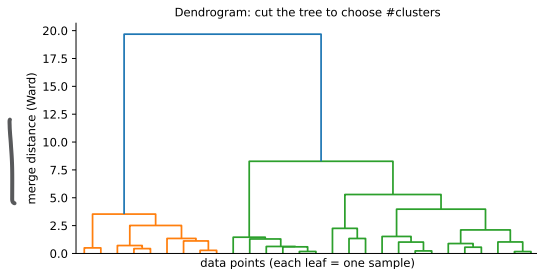
The dendrogram: cut the tree

Every merge is a bar at its **merge distance**. **Cut** the tree at a height to read off clusters:

- ▶ low cut = many small clusters; high cut = few big ones;
- ▶ the largest vertical gap is a natural place to cut.



$$E[X - E[X]] = 0$$



Linkage: what is the distance between two clusters?

“Merge the two closest clusters” — but *closest* can mean different things. Linkage is that definition:

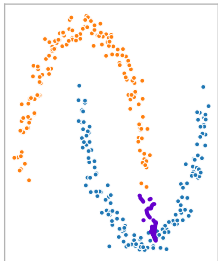
- ▶ Single — distance between the *nearest* pair of points (one in each). Chains along shapes; can “snake”.
- ▶ Complete — distance between the *farthest* pair. Makes compact, equal blobs.
- ▶ **Average** — mean over all cross-cluster pairs. A compromise.
- ▶ Ward (Ward, 1963) — merge the pair that increases total within-cluster variance the least. The common default; behaves like k-means.



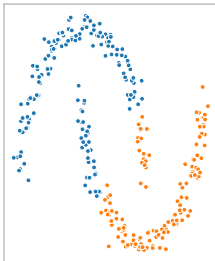
Linkage: same data, different answers

Linkage changes the clusters (same data, $k=2$)

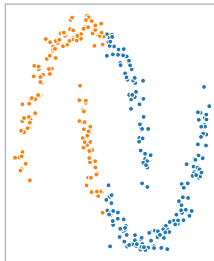
single



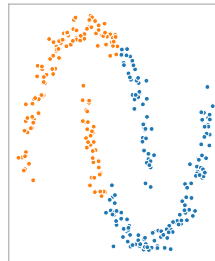
complete



average



ward

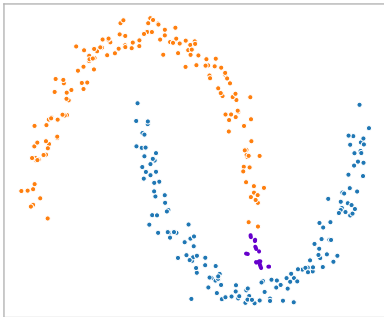


Two moons, $k = 2$: **single** linkage follows the crescents; **complete** / **average** / **Ward** cut them into compact halves. The linkage choice changes everything.

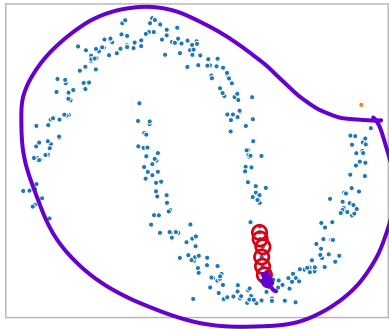
Single linkage's weakness: chaining

Single linkage just won that slide — but it merges on the *nearest pair*, so a thin trail of points is enough to fuse two clusters that are obviously separate:

clean data: single linkage follows the crescents



6 points in the gap: both moons become one cluster (sizes 305 and 1)



Six added points — 2% of the data — collapse both moons into one cluster of **305**, leaving a single leftover point as “cluster 2”. Single linkage won the last slide because that data was **clean**. For odd shapes *with* noise, reach for DBSCAN instead.



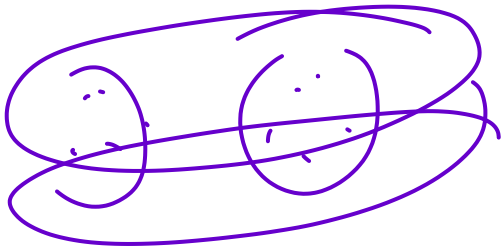
How good is a clustering?

The hard part: there is usually *no ground truth*. Internal metrics judge the shape; external ones need labels.

internal.



WCSS



Two regimes: internal vs external

- ▶ Internal (no labels — the usual case): judge cohesion + separation from the data alone. Silhouette, Davies–Bouldin, Calinski–Harabasz, inertia/elbow.
- ▶ External (labels exist — the “labels but still cluster” cases): compare the clustering to known labels. ARI, AMI, V-measure.

y

There is no single “accuracy” for clustering — the cluster numbers are arbitrary, and usually there is nothing to be right *about*.



$x_1 x_2 x_3 | y$
 $\sqrt{1} \sqrt{7} \sqrt{5}$
36 / 56

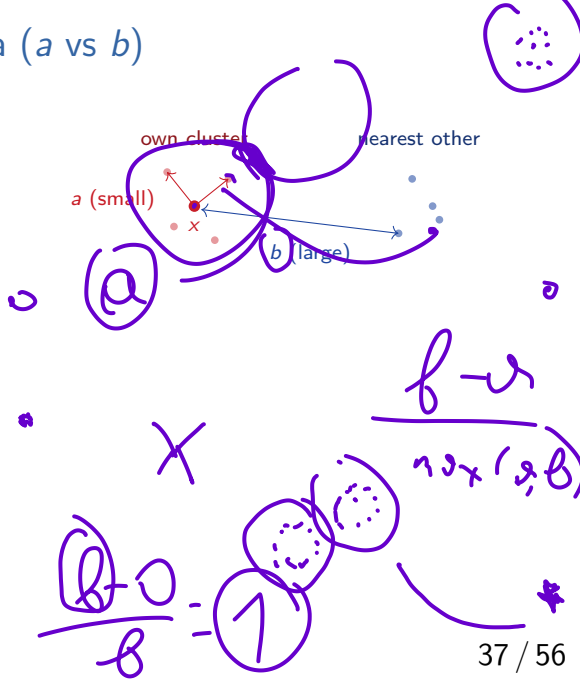
Silhouette (Rousseeuw, 1987): the idea (a vs b)

For one point: a = mean distance to its *own* cluster; b = mean distance to the *nearest other* cluster.

$$s = \frac{b - a}{\max(a, b)} \in [-1, 1]$$

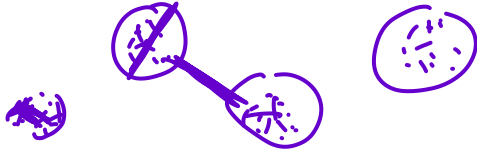
- ▶ $s \approx 1$: snug in its own cluster, far from the rest (ideal).
- ▶ $s \approx 0$: sitting on the boundary between two clusters.
- ▶ $s < 0$: closer to a neighbouring cluster than to its own — probably misassigned.

What would $s = -1$ mean? Only $b = 0$ gives it: the point sits *exactly on top of* the whole nearest cluster while its own is spread out. A theoretical floor — a genuinely misassigned point in real data scores about -0.4 , not -1 .



Silhouette: reading the plot

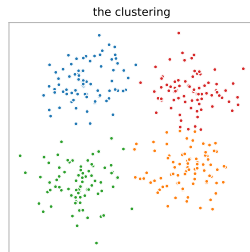
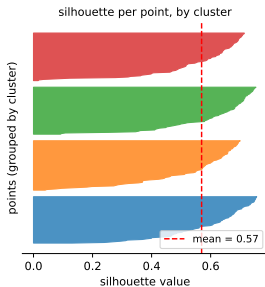
b-u b-a



Each bar is one point's s , **grouped by cluster** and sorted; the dashed line is the **mean** silhouette.

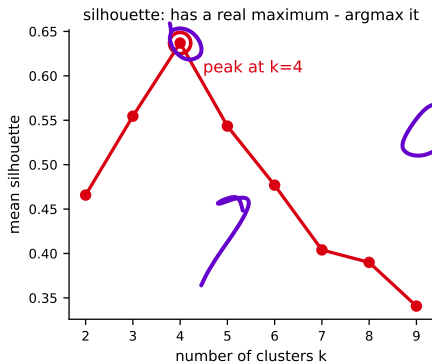
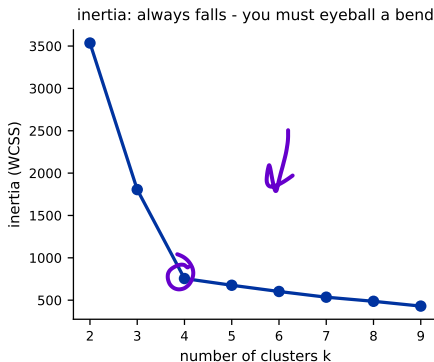
- Wide, tall, mostly-right bars = a good cluster.
- A cluster dipping below 0 = likely merged/over-split.

Pick the k with the highest mean silhouette.
(Also: Davies–Bouldin, lower better;
Calinski–Harabasz, higher better.)



Choosing k : elbow vs silhouette

The same four-blob data as the elbow slide, now scored both ways:



- ▶ Inertia has **no maximum** — it falls forever, so “the elbow” stays a judgement call.
- ▶ The silhouette **does** have one: you can take the argmax and be done.
- ▶ Here they agree on $k = 4$. **When they disagree, trust neither** — that is the data telling you the groups are not clean.

g

v -
| -

When you do have labels

Now there is something to be right about — but plain “accuracy” still will not work.



External metrics: when you do have labels

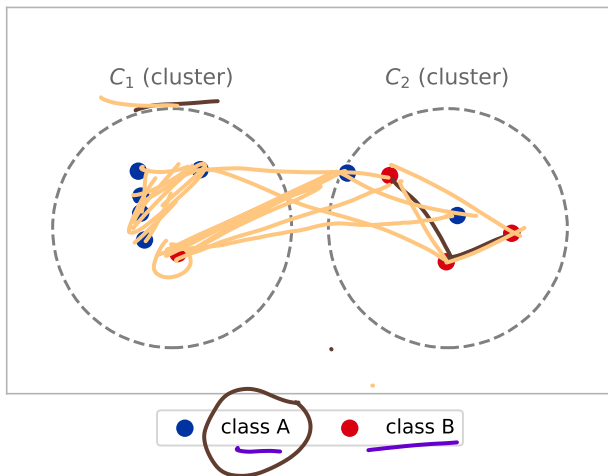
Compare the clustering to a known labeling (which cluster = which label does *not* matter):

- ▶ **ARI** (adjusted Rand index; Hubert & Arabie, 1985) and **AMI** (adjusted mutual information): **chance-adjusted** — 0 means “no better than random”, 1 perfect.
- ▶ **V-measure** (homogeneity + completeness); also Fowlkes–Mallows.

Prefer the **adjusted** versions (ARI, AMI). Raw or **normalized mutual information (NMI)** is not corrected for chance, so it rewards simply using more clusters. Plain accuracy fails — the labels are a *permutation* of the clusters.

The data behind the Rand index

Eleven points, each with a true **class** (color) and a **cluster** (the dashed groups). The next slide just *counts* these into a table:



r 1
c 1

Cluster C_1 : 5 class-A + 1 class-B; cluster C_2 : 2 class-A + 3 class-B.

Step 1: stop comparing labels, compare *pairs*

Cluster names are arbitrary, so “ $C_1 = \text{class A}$ ” is meaningless. Instead take every **pair** of points and ask one question both labelings can answer:

Do these two points belong together?

The two labelings **agree** on a pair if they *both* say “together”, or *both* say “apart”.



	C_1	C_2	Σ
A	5	2	7
B	1	3	4
Σ	6	5	11

The **5**: five points are *both* class A and in cluster C_1 .



With 11 points there are $\binom{11}{2} = 55$ pairs. Any group of m points contributes $\binom{m}{2}$ “together” pairs, so read the counts straight off the table:

- ▶ together in **both** (inside one cell):

$$\binom{5}{2} + \binom{2}{2} + \binom{1}{2} + \binom{3}{2} = 10 + 1 + 0 + 3 = \mathbf{14}$$

- ▶ together by **class** (row totals): $\binom{7}{2} + \binom{4}{2} = \mathbf{27}$

- ▶ together by **cluster** (column totals): $\binom{6}{2} + \binom{5}{2} = 25$
- 

Step 2: Rand looks fine, ARI does not

Apart in both is what is left over: $55 - 27 - 25 + 14 = 17$ (subtract each "together" count, then add back the 14 removed twice).

0.61

$$\text{Rand} = \frac{\text{pairs they agree on}}{\text{all pairs}} = \frac{14 + 17}{55} = 0.56$$

That sounds respectable. But **even random labels get a lot of pairs right by luck**, so subtract what chance alone would have scored:

$$\text{ARI} = \frac{14 - \overbrace{27 \cdot 25 / 55}^{= 12.3 \text{ expected}}}{\underbrace{(27 + 25) / 2}_{= 26 \text{ best possible}} - 12.3} = \frac{1.7}{13.7} = 0.13$$

Of the 14 "together in both" pairs, **12.3 were free** — random labels would have managed almost as well. Rand says 0.56; ARI says **barely above chance**. Always report the **adjusted** one.

Evaluation in practice

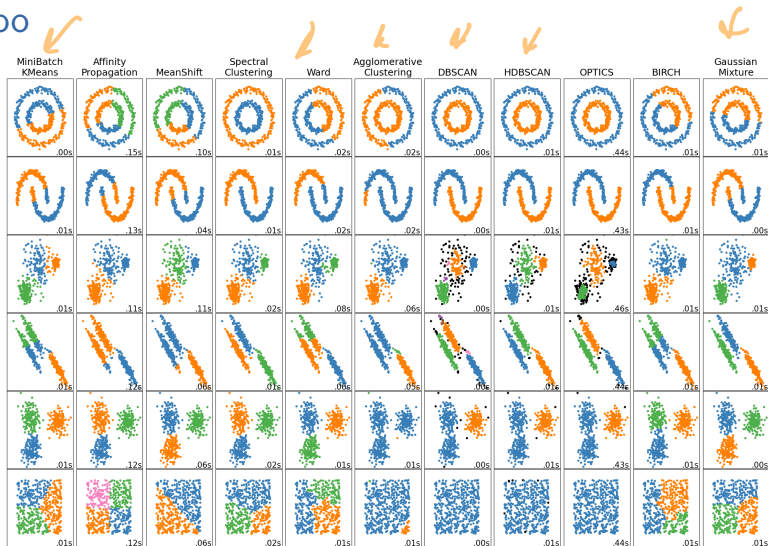


- ▶ Track more than one metric. No single number captures a clustering — report, say, silhouette *and* Davies–Bouldin (and ARI/AMI if labels exist); trust the result only when they agree.
- ▶ Pick k by the **silhouette** (or the elbow), not by eye alone.
- ▶ Check **stability**: do clusters survive a re-run / a resample / a different seed?
- ▶ **Scale first** — every internal metric is distance-based too.
- ▶ Sanity-check by *looking* at the clusters.

Choosing, and a wider view

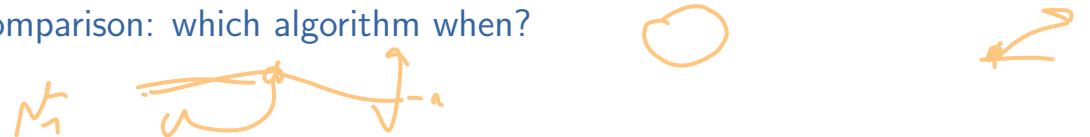
Many more algorithms exist — here is how they all behave, and the traps to avoid.

The wider zoo



Source: scikit-learn (BSD-3), clustering-algorithm comparison.

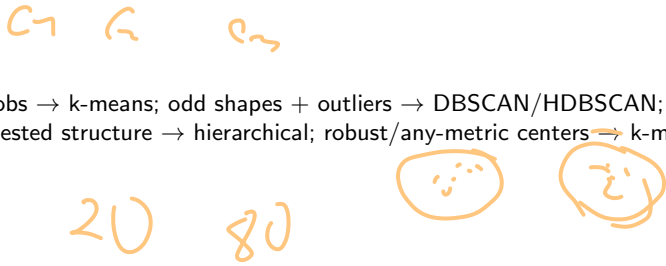
Comparison: which algorithm when?



Algorithm	Needs k ?	Cluster shape	Noise?	Scales?
k-means / MiniBatch	yes	round, equal	no	very well
k-medoids	yes	any metric	no (robust)	poorly
GMM (EM)	yes (+BIC)	elliptical	no	medium
DBSCAN	no	arbitrary	yes	medium
HDBSCAN	no	arbitrary, varying density	yes	medium
Hierarchical	no (cut tree)	flexible (linkage)	no	poorly



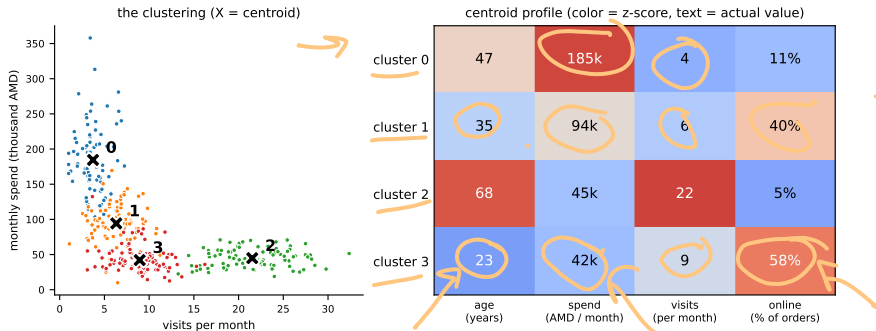
Use when: fast round blobs $\rightarrow$ k-means; odd shapes + outliers $\rightarrow$ DBSCAN/HDBSCAN;
soft/elliptical $\rightarrow$ GMM; nested structure $\rightarrow$ hierarchical; robust/any-metric centers $\rightarrow$ k-medoids.



The step after the algorithm: read the clusters

`fit_predict` hands you a column of 0,1,2,3. That is not an answer, it is homework.

Profile the centroids — one row per cluster — and see what actually separates them.



Only now can they be *named*: **0** the monthly stock-up (185k drams, 4 visits), **1** young families, **2** daily regulars (68, 22 visits), **3** students (23, 58% online).

The algorithm found the groups; you supplied the meaning. So the names are **hypotheses** — check them on customers the clustering never saw.

Bonus use: anomaly detection

isolated.
-1
#5

Clustering doubles as outlier detection:

- ▶ **DBSCAN / HDBSCAN** label low-density points as **noise** (-1) — those are your anomalies, for free.
- ▶ **GMM** gives each point a **likelihood**; flag the lowest-probability points.

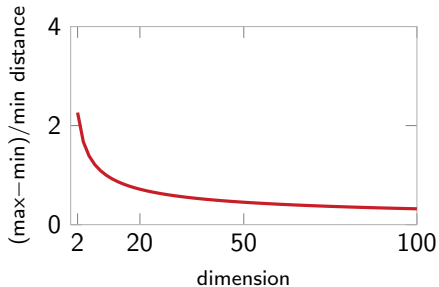
Useful when “normal” has structure but “weird” does not — fraud, faults, intrusions.

The curse of dimensionality

As dimensions grow, points spread out and **pairwise distances concentrate** — the nearest and farthest neighbors become almost equally far.

Distance-*based* clustering (all of today's methods) degrades: “near” stops meaning much.

Fix: **reduce dimensions first** (PCA), then cluster in the low-dimensional space.



Contrast between nearest and farthest collapses toward 0.

Cluster after PCA

A standard pipeline for high-dimensional data:

standardize → PCA (e.g. to 95% variance) → cluster.

- ▶ Distances behave better; clustering is faster and often cleaner.
- ▶ You can finally *plot* the result in 2D.

Dimensionality reduction (PCA, t-SNE, UMAP) is the **next deck** — it pairs naturally with clustering.

All of it in sklearn

```
from sklearn.preprocessing import StandardScaler
from sklearn.cluster import KMeans, DBSCAN, AgglomerativeClustering

X = StandardScaler().fit_transform(X_raw)           # scale first!

labels = KMeans(n_clusters=4, n_init=10, random_state=509).
    fit_predict(X)
labels = DBSCAN(eps=0.3, min_samples=5).fit_predict(X)      # -1 =
    noise
labels = AgglomerativeClustering(n_clusters=4).fit_predict(X)
```

Same `fit_predict` pattern for every algorithm. Set `random_state` for reproducibility; evaluate with `silhouette_score` / `adjusted_rand_score`.

Easy to get wrong — watch out

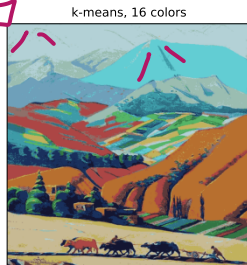
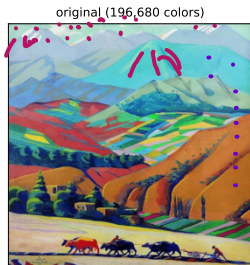
- ▶ Results depend on **scaling, k , the metric, and the algorithm** — change any one and the clusters change.
- ▶ There is **no ground truth** to check against.
- ▶ It is easy to **see groups that are not real**. Run k-means with $k = 4$ on 300 points drawn *uniformly at random* and it returns four tidy “clusters” with a silhouette of 0.41 — from data with no structure whatsoever.

Always validate (silhouette, stability, a hard look) before trusting a clustering.

Next time: shrinking image files with k-means

Color quantization. An image has up to ~ 16 million colors. Run k-means on the *pixels* (in RGB space); repaint each pixel with its centroid color.

Result: only k **colors** — a much smaller palette / file, for a tiny quality cost. We build this end-to-end in the **practical**.



Martiros Saryan, *Mountains* (1923).

Recap

- ▶ **Clustering** groups unlabeled data; you choose the distance and (often) k .
- ▶ **k-means** — fast, round, equal clusters; Lloyd's alternation, k-means++, pick k by elbow/silhouette. **GMM** generalizes it (soft, elliptical) via EM.
- ▶ **DBSCAN / HDBSCAN** — density-based: arbitrary shapes, no k , free outliers. **Hierarchical** — a dendrogram you cut. **k-medoids** — robust, any metric.
- ▶ **Evaluation** has no ground truth: silhouette (internal); ARI/AMI (external); track more than one metric, always scale, check stability.
- ▶ Then **read the clusters** — profile the centroids and name the groups. Labels are the algorithm's output; meaning is yours.

Next: dimensionality reduction (PCA, t-SNE, UMAP) — compress the features, visualize the data, and cluster better in high dimensions.