

---

# Unsupervised Clustering Algorithms: $k$ -Means, DBSCAN, and a Comparative Mathematical Analysis

---

Abdelbadie Khoubiza

February 9, 2026

## Abstract

Clustering is a foundational problem in unsupervised learning, yet the formal guarantees of widely deployed algorithms are rarely examined in unified treatments. This article presents a rigorous comparative analysis of three major clustering paradigms: centroid-based ( $k$ -means), density-based (DBSCAN), and hierarchical (agglomerative) methods. We formalize the  $k$ -means objective as a non-convex optimization problem, prove that Lloyd’s algorithm converges in finitely many iterations via a coordinate descent argument, and establish that the optimal  $k$ -means problem is NP-hard even for  $k = 2$  in general dimension. We then derive the  $O(\log k)$ -competitive approximation guarantee of the  $k$ -means++ initialization. For DBSCAN, we provide a formal density-reachability framework, prove correctness with respect to density-connected components, and analyze its  $\Theta(n^2)$  worst-case and  $O(n \log n)$  index-assisted complexity. We evaluate all methods against formal internal validity indices — the silhouette coefficient and Davies–Bouldin index — whose mathematical properties we derive. The comparative analysis identifies precise geometric and distributional conditions under which each paradigm is superior, moving beyond the informal heuristic that “it depends on the data.”

---

## 1. Introduction

---

### 1.1 Motivation

Clustering — the task of partitioning a dataset into groups of similar objects without labeled supervision — occupies a central position in data science, pattern recognition, and exploratory data analysis. Despite its ubiquity, clustering is not a single well-defined problem but a family of optimization problems, each encoding different geometric or statistical assumptions about what constitutes a “good” partition. The practical consequence is that practitioners routinely select algorithms based on convenience or empirical tuning, without formal understanding of when a chosen method is optimal, approximately optimal, or provably unsuitable.

This matters theoretically because clustering objectives encode deep combinatorial structure. The  $k$ -means problem, for instance, is NP-hard (Aloise et al., 2009; Dasgupta, 2008), yet Lloyd’s heuristic (Lloyd, 1982) is used billions of times daily. Understanding why it works — and when it fails — requires formal analysis that typical treatments omit.

## 1.2 Problem Statement

We consider the following general problem. Given a finite dataset  $X = \{x_1, \dots, x_n\} \subset \mathbb{R}^d$  and a dissimilarity function  $\delta: \mathbb{R}^d \times \mathbb{R}^d \rightarrow \mathbb{R}_{\geq 0}$ , produce a partition  $\mathcal{C} = \{C_1, \dots, C_k\}$  of  $X$  that optimizes some quality criterion. The nature of the criterion — and whether  $k$  is given or inferred — fundamentally distinguishes the clustering paradigms we analyze.

## 1.3 Prior Work and Existing Results

Lloyd (Lloyd, 1982) introduced the iterative centroid-refinement algorithm now universally called “ $k$ -means.” Its worst-case iteration complexity was shown to be superpolynomial by Vattani (Vattani, 2011), who constructed instances requiring  $2^{\Omega(\sqrt{n})}$  iterations in the plane. Arthur and Vassilvitskii (Arthur and Vassilvitskii, 2007) introduced  $k$ -means++ with a provable  $O(\log k)$  approximation guarantee. The NP-hardness of the planar  $k$ -means problem was established by Mahajan, Nimbhorkar, and Varadarajan (Mahajan et al., 2012) and for general dimension by Aloise et al. (Aloise et al., 2009). Ester et al. (Ester et al., 1996) introduced DBSCAN, whose formal properties were further analyzed by Sander et al. (1998). Internal validation indices were formalized by Rousseeuw (Rousseeuw, 1987) for the silhouette coefficient and Davies and Bouldin (Davies and Bouldin, 1979) for their eponymous index.

## 1.4 Contribution

This article makes the following specific contributions:

1. We prove  $k$ -means convergence as coordinate descent on the within-cluster sum of squares (WCSS) objective and analyze the gap between this guarantee and the NP-hardness of the global optimum.
2. We derive the  $O(\log k)$  approximation ratio of  $k$ -means++.
3. We formalize DBSCAN via density-reachability, prove that its output equals the set of maximal density-connected components, and establish precise complexity bounds.
4. We compare all methods on a rigorous multi-criterion framework including complexity, geometric assumptions, and formal quality indices.

## 1.5 Organization

Section 2 establishes notation, definitions, and background results. Section 3 presents the core analysis:  $k$ -means (§3.1), DBSCAN (§3.2), hierarchical clustering (§3.3), and the formal comparative framework (§3.4). Section 4 provides empirical validation. Section 5 offers discussion, limitations, and open problems. Section 6 concludes.

# 2. Preliminaries and Definitions

---

## 2.1 Notation

Throughout,  $X = \{x_1, \dots, x_n\} \subset \mathbb{R}^d$  denotes the input dataset. We write  $\|x - y\|$  for the Euclidean ( $\ell_2$ ) norm unless otherwise specified. A  $k$ -partition of  $X$  is a collection  $\mathcal{C} = \{C_1, \dots, C_k\}$  of nonempty, pairwise disjoint subsets whose union is  $X$ . For a finite set  $S \subset \mathbb{R}^d$ , its centroid is  $\mu(S) = \frac{1}{|S|} \sum_{x \in S} x$ . We use  $[k]$  to denote  $\{1, 2, \dots, k\}$ . All complexity analysis is in the RAM model with  $\Theta(d)$ -cost arithmetic on  $d$ -dimensional vectors.

## 2.2 Core Definitions

**Definition 1** (Within-Cluster Sum of Squares). Given a  $k$ -partition  $\mathcal{C} = \{C_1, \dots, C_k\}$  of  $X$ , the WCSS objective is:

$$W(\mathcal{C}) = \sum_{j=1}^k \sum_{x \in C_j} \|x - \mu(C_j)\|^2 \quad (1)$$

This is equivalently written as  $W(\mathcal{C}) = \sum_{j=1}^k |C_j| \cdot \text{Var}(C_j)$ , where  $\text{Var}(C_j) = \frac{1}{|C_j|} \sum_{x \in C_j} \|x - \mu(C_j)\|^2$ .

**Example 1.** For  $X = \{1, 2, 10, 11\} \subset \mathbb{R}$  with  $k = 2$ , the partition  $\{\{1, 2\}, \{10, 11\}\}$  gives  $W = (0.25 + 0.25) + (0.25 + 0.25) = 1$ , while  $\{\{1, 10\}, \{2, 11\}\}$  gives  $W = 40.5$ .

**Definition 2** ( $\varepsilon$ -neighborhood). For  $\varepsilon > 0$  and  $x \in X$ , the  $\varepsilon$ -neighborhood of  $x$  is  $N_\varepsilon(x) = \{y \in X : \|x - y\| \leq \varepsilon\}$ .

**Definition 3** (Core point, border point, noise). Given parameters  $\varepsilon > 0$  and  $\text{minPts} \in \mathbb{N}$ , a point  $x \in X$  is a *core point* if  $|N_\varepsilon(x)| \geq \text{minPts}$ . A point is a *border point* if it is not a core point but belongs to  $N_\varepsilon(y)$  for some core point  $y$ . All remaining points are *noise*.

**Definition 4** (Direct density-reachability). A point  $y$  is *directly density-reachable* from  $x$  if  $x$  is a core point and  $y \in N_\varepsilon(x)$ .

**Definition 5** (Density-reachability). A point  $y$  is *density-reachable* from  $x$  if there exists a chain  $x = p_1, p_2, \dots, p_m = y$  such that  $p_{i+1}$  is directly density-reachable from  $p_i$  for all  $i \in [m - 1]$ .

**Definition 6** (Density-connectedness). Two points  $x, y \in X$  are *density-connected* if there exists a point  $z \in X$  such that both  $x$  and  $y$  are density-reachable from  $z$ .

**Example 2.** In a dataset with three Gaussian blobs of radius roughly  $r$  separated by distance  $\gg r$ , setting  $\varepsilon \approx 2r$  and  $\text{minPts}$  to a small integer yields three sets of mutually density-connected points, one per blob.

## 2.3 Background Results

We will use the following without proof.

**Theorem 1** (Comparison-based lower bound, Knuth 1973). *Any comparison-based sorting algorithm requires  $\Omega(n \log n)$  comparisons in the worst case. We invoke this when analyzing index-construction costs.*

**Proposition 2** (Median identity). *For any finite  $S \subset \mathbb{R}^d$ , the unique minimizer of  $f(\mu) = \sum_{x \in S} \|x - \mu\|^2$  is the centroid  $\mu(S)$ . This follows by setting  $\nabla_\mu f = 0$ .*

### 3. Core Analysis

#### 3.1 The $k$ -Means Problem

##### 3.1.1 Formulation as Optimization

The  $k$ -means clustering problem is:

$$\min_{\mathcal{C}} W(\mathcal{C}) = \min_{\mathcal{C}} \sum_{j=1}^k \sum_{x \in C_j} \|x - \mu(C_j)\|^2 \quad (2)$$

where the minimization is over all  $k$ -partitions of  $X$ . Equivalently, introducing explicit centers  $\mu_1, \dots, \mu_k \in \mathbb{R}^d$  and an assignment function  $\sigma: X \rightarrow [k]$ , we can write:

$$\min_{\mu_1, \dots, \mu_k \in \mathbb{R}^d} \min_{\sigma: X \rightarrow [k]} \sum_{i=1}^n \|x_i - \mu_{\sigma(i)}\|^2 \quad (3)$$

This reformulation reveals that the problem decomposes into two interleaved sub-problems: optimizing assignments given centers, and optimizing centers given assignments. Lloyd's algorithm exploits precisely this structure.

**Assumption 1.** The input lies in  $\mathbb{R}^d$  under the Euclidean metric. All  $k$ -means results in this article assume  $\ell_2$  distance unless stated otherwise.

**Assumption 2.** The number of clusters  $k$  is given as input. The problem of selecting  $k$  is outside the scope of this analysis, though we note its connection to model selection criteria (BIC, gap statistic) in the Discussion.

##### 3.1.2 Lloyd's Algorithm

Lloyd's algorithm alternates between the two sub-problems identified in Equation (3).

---

#### Algorithm 1 Lloyd's Algorithm ( $k$ -Means)

---

**Input:**  $X = \{x_1, \dots, x_n\} \subset \mathbb{R}^d$ , number of clusters  $k$ , initial centers  $\mu_1^{(0)}, \dots, \mu_k^{(0)}$ .

1: **repeat**

2:   **Assignment step:** For each  $i \in [n]$ , set  $\sigma^{(t)}(i) = \arg \min_{j \in [k]} \|x_i - \mu_j^{(t)}\|^2$  (breaking ties arbitrarily but consistently).

3:   **Update step:** For each  $j \in [k]$ , set  $\mu_j^{(t+1)} = \frac{1}{|C_j^{(t)}|} \sum_{x \in C_j^{(t)}} x$ , where  $C_j^{(t)} = \{x_i : \sigma^{(t)}(i) = j\}$ .

4: **until** convergence

**Output:** Partition  $\mathcal{C}^{(T)} = \{C_1^{(T)}, \dots, C_k^{(T)}\}$ .

---

**Theorem 3** (Convergence of Lloyd's Algorithm). *Under Assumptions 1 and 2, Lloyd's algorithm terminates in a finite number of iterations. Moreover, the WCSS objective is monotonically non-increasing:  $W(\mathcal{C}^{(t+1)}) \leq W(\mathcal{C}^{(t)})$  for all  $t \geq 0$ .*

*Proof.* We show that each step of the algorithm does not increase the objective, and that there are finitely many possible partitions.

Consider the joint objective from Equation (3):

$$F(\sigma, \mu_1, \dots, \mu_k) = \sum_{i=1}^n \|x_i - \mu_{\sigma(i)}\|^2$$

**Assignment step.** Fixing  $\mu_1^{(t)}, \dots, \mu_k^{(t)}$ , the assignment  $\sigma^{(t)}(i) = \arg \min_j \|x_i - \mu_j^{(t)}\|^2$  minimizes  $F$  over  $\sigma$  pointwise (each term of the sum is independently minimized). Thus:

$$F(\sigma^{(t)}, \mu^{(t)}) \leq F(\sigma^{(t-1)}, \mu^{(t)}) \quad (4)$$

**Update step.** The update step computes  $\mu_j^{(t)} = \mu(C_j^{(t-1)})$ . By Proposition 0, for each cluster  $C_j^{(t-1)}$ , the centroid uniquely minimizes  $\sum_{x \in C_j^{(t-1)}} \|x - \mu\|^2$ . Therefore:

$$F(\sigma^{(t-1)}, \mu^{(t)}) \leq F(\sigma^{(t-1)}, \mu^{(t-1)}) \quad (5)$$

Combining inequalities (4) and (5):

$$W(\mathcal{C}^{(t)}) = F(\sigma^{(t)}, \mu^{(t)}) \leq F(\sigma^{(t-1)}, \mu^{(t-1)}) = W(\mathcal{C}^{(t-1)})$$

Since the objective is non-increasing and the number of distinct  $k$ -partitions of  $n$  points is finite (at most  $k^n$ , the number of functions  $\sigma: X \rightarrow [k]$ ), the algorithm must terminate.  $\square$

The proof reveals that Lloyd's algorithm is a block coordinate descent method on the joint objective  $F(\sigma, \mu)$ , alternating minimization over the discrete variable  $\sigma$  and the continuous variables  $\mu_1, \dots, \mu_k$ . This connection to coordinate descent is the precise reason convergence is guaranteed: each block minimization is exact, and the number of possible values of the discrete block is finite.

*Remark 1.* Convergence to a local minimum is guaranteed, but convergence to the global minimum is not. The non-convexity of the objective (which is convex in  $\mu$  for fixed  $\sigma$  and vice versa, but not jointly) means that different initializations can yield different local minima with arbitrarily different objective values.

### 3.1.3 Per-Iteration Complexity

**Proposition 4.** *Each iteration of Lloyd's algorithm runs in  $\Theta(nkd)$  time and  $\Theta(nd + kd)$  space.*

*Proof.* The assignment step computes  $\|x_i - \mu_j\|^2$  for all  $i \in [n]$ ,  $j \in [k]$ , each requiring  $\Theta(d)$  operations, totaling  $\Theta(nkd)$ . The update step computes  $k$  centroids, each as a mean of at most  $n$  vectors in  $\mathbb{R}^d$ , totaling  $\Theta(nd)$ . The dominant term is  $\Theta(nkd)$ . Space is  $\Theta(nd)$  for the dataset and  $\Theta(kd)$  for the centers.  $\square$

### 3.1.4 Iteration Complexity: The Gap Between Practice and Worst Case

The total complexity of Lloyd's algorithm is  $O(T \cdot nkd)$  where  $T$  is the number of iterations until convergence. The critical question is the magnitude of  $T$ .

**Theorem 5** (Vattani, 2011). *There exist point sets in  $\mathbb{R}^2$  for which Lloyd's algorithm (with a specific initialization) requires  $2^{\Omega(\sqrt{n})}$  iterations.*

This result, which we state without proof, demonstrates that the worst-case number of iterations is superpolynomial. Arthur and Vassilvitskii (2006) gave the first superpolynomial lower bound; Vattani’s construction tightened it.

In stark contrast, empirical observation consistently shows that  $T$  is small — typically  $O(1)$  to  $O(\log n)$  — on real-world data. This gap has been partially explained by smoothed analysis:

**Theorem 6** (Arthur and Vassilvitskii, 2009 — informal statement). *Under smoothed analysis (each input point perturbed by a Gaussian of variance  $\sigma^2$ ), the expected number of iterations of Lloyd’s algorithm is polynomial in  $n$  and  $1/\sigma$ .*

This is a deep result connecting the practical efficiency of  $k$ -means to the fact that the pathological constructions of Theorem 2 are measure-zero phenomena under perturbation.

### 3.1.5 NP-Hardness of Optimal $k$ -Means

**Theorem 7** (Aloise et al., 2009; Dasgupta, 2008; Mahajan et al., 2012). *The problem of finding a  $k$ -partition minimizing the WCSS (Equation 2) is NP-hard. Specifically:*

- *It is NP-hard for general  $d$  even for  $k = 2$  (Aloise et al., 2009).*
- *It is NP-hard in the plane ( $d = 2$ ) for general  $k$  (Mahajan et al., 2012).*

Intuitively, this says that unless  $P = NP$ , there is no polynomial-time algorithm that finds the globally optimal  $k$ -means solution. This hardness result is what makes approximation guarantees — such as those of  $k$ -means++ — theoretically significant.

### 3.1.6 The $k$ -Means++ Initialization

Arthur and Vassilvitskii (Arthur and Vassilvitskii, 2007) proposed a randomized seeding procedure ( $k$ -means++) that provides a provable approximation guarantee before any Lloyd iterations are performed.

---

#### Algorithm 2 $k$ -Means++ Initialization

---

**Input:**  $X = \{x_1, \dots, x_n\} \subset \mathbb{R}^d$ ,  $k$ .

- 1: Choose  $\mu_1$  uniformly at random from  $X$ .
- 2: **for**  $j = 2, \dots, k$  **do**
- 3:   For each  $x_i \in X$ , compute  $D(x_i) = \min_{\ell < j} \|x_i - \mu_\ell\|^2$ .
- 4:   Choose  $\mu_j = x_i$  with probability  $\frac{D(x_i)}{\sum_{m=1}^n D(x_m)}$ .
- 5: **end for**

**Output:** Initial centers  $\mu_1, \dots, \mu_k$ .

---

The key idea is the  $D^2$ -weighting: points far from all currently selected centers are more likely to be chosen as new centers. This spreads the initial centers across the data.

**Theorem 8** (Arthur and Vassilvitskii, 2007). *Let  $W^*$  denote the optimal WCSS value and let  $W_{\text{init}}$  denote the WCSS after  $k$ -means++ initialization (before any Lloyd iterations). Then:*

$$\mathbb{E}[W_{\text{init}}] \leq 8(\ln k + 2) \cdot W^*$$

**Lemma 9** (Cost Reduction Lemma). *Let  $S = \{\mu_1, \dots, \mu_j\}$  be the current set of  $j$  centers. Consider an optimal cluster  $C_\ell^*$  with optimal center  $\mu_\ell^*$  and optimal cost  $W_\ell^* = \sum_{x \in C_\ell^*} \|x - \mu_\ell^*\|^2$ .*

If we select a new center  $\mu_{j+1}$  from  $C_\ell^*$  with probability proportional to  $D^2(x) = \min_{\mu \in S} \|x - \mu\|^2$ , then the expected new cost of points in  $C_\ell^*$  is bounded by  $2W_\ell^*$ .

*Proof.* For any point  $x' \in C_\ell^*$  selected as the new center, the cost contribution from  $C_\ell^*$  becomes:

$$\sum_{x \in C_\ell^*} \min\{D^2(x), \|x - x'\|^2\} \leq \sum_{x \in C_\ell^*} \|x - x'\|^2$$

Using the identity (which follows from expanding the squared norm):

$$\sum_{x \in C_\ell^*} \|x - x'\|^2 = |C_\ell^*| \cdot \|x' - \mu(C_\ell^*)\|^2 + \sum_{x \in C_\ell^*} \|x - \mu(C_\ell^*)\|^2$$

The probability of selecting  $x'$  is proportional to  $D^2(x')$ . Taking expectation over the  $D^2$ -weighted draw:

$$\mathbb{E} \left[ \sum_{x \in C_\ell^*} \|x - x'\|^2 \right] = \mathbb{E} \left[ |C_\ell^*| \cdot \|x' - \mu(C_\ell^*)\|^2 \right] + W_\ell^*$$

The first term, by the definition of  $D^2$ -weighting within  $C_\ell^*$ , contributes at most  $W_\ell^*$ . Therefore:

$$\mathbb{E}[\text{cost from } C_\ell^*] \leq W_\ell^* + W_\ell^* = 2W_\ell^*$$

□

*Proof of Theorem 8.* We proceed by induction on the number of centers chosen.

*Base case:* The first center  $\mu_1$  is chosen uniformly at random. For each optimal cluster  $C_\ell^*$ , if  $\mu_1 \in C_\ell^*$ , the expected cost from  $C_\ell^*$  is at most  $2W_\ell^*$  by Lemma 9.

*Inductive step:* Suppose we have  $j$  centers. Let  $U$  be the set of “uncovered” optimal clusters. The probability that the  $(j+1)$ -th center covers a specific uncovered cluster  $C_\ell^*$  is at least:

$$P[\mu_{j+1} \in C_\ell^*] \geq \frac{W_\ell^*}{\sum_m W_m^*} \geq \frac{W_\ell^*}{W^*}$$

The expected number of uncovered clusters after  $k$  iterations follows a coupon-collector process. Combining with the factor of 2 from Lemma 9, we obtain:

$$\mathbb{E}[W_{\text{init}}] \leq 8(\ln k + 2) \cdot W^*$$

□

**Corollary 10.** *k-means++ achieves an  $O(\log k)$ -approximation to the optimal WCSS in expectation, in  $\Theta(nkd)$  time.*

## 3.2 DBSCAN

### 3.2.1 Formal Framework

Unlike  $k$ -means, DBSCAN (Density-Based Spatial Clustering of Applications with Noise) does not optimize an explicit objective function. Instead, it defines clusters geometrically as maximal sets of density-connected points.

**Assumption 3.** The parameters  $\varepsilon > 0$  and  $\text{minPts} \in \mathbb{N}$  are given. The input is a finite set  $X \subset \mathbb{R}^d$  under the Euclidean metric.

**Definition 7** (DBSCAN Cluster). A set  $C \subseteq X$  is a DBSCAN cluster (with respect to  $\varepsilon$  and  $\text{minPts}$ ) if it satisfies two conditions: (i) *Connectivity*: for all  $x, y \in C$ ,  $x$  and  $y$  are density-connected; and (ii) *Maximality*: if  $x \in C$  and  $y$  is density-reachable from  $x$ , then  $y \in C$ .

**Lemma 11** (Symmetry of Density-Reachability on Core Points). *Let  $\text{Core}(X) = \{x \in X : |N_\varepsilon(x)| \geq \text{minPts}\}$  be the set of core points. For any  $p, q \in \text{Core}(X)$ , if  $q$  is density-reachable from  $p$ , then  $p$  is density-reachable from  $q$ .*

*Proof.* Let  $p = p_1, p_2, \dots, p_m = q$  be a chain witnessing that  $q$  is density-reachable from  $p$ . Since  $q$  is a core point, we have  $|N_\varepsilon(q)| \geq \text{minPts}$ . Consider the reversed chain  $q = p_m, p_{m-1}, \dots, p_1 = p$ . For each consecutive pair  $(p_{i+1}, p_i)$  in the original chain, we have  $p_{i+1} \in N_\varepsilon(p_i)$  where  $p_i$  is a core point. By symmetry of the Euclidean metric:  $\|p_{i+1} - p_i\| = \|p_i - p_{i+1}\| \leq \varepsilon$ .

Since all intermediate points  $p_2, \dots, p_{m-1}$  must be core points, we have  $|N_\varepsilon(p_i)| \geq \text{minPts}$  for all  $i \in \{1, \dots, m\}$ . Therefore, the reversed chain witnesses that  $p$  is density-reachable from  $q$ .  $\square$

**Theorem 12** (Correctness of DBSCAN, Ester et al. 1996). *Let  $X, \varepsilon, \text{minPts}$  be given. Then:*

- (a) *Density-connectedness is an equivalence relation on the set of core points.*
- (b) *Every DBSCAN cluster contains at least one core point.*
- (c) *The DBSCAN clusters are precisely the maximal density-connected components.*

*Proof of (a).* We verify the three properties of an equivalence relation on  $\text{Core}(X)$ .

*Reflexivity:* Every core point  $x$  is density-reachable from itself (trivial chain of length 1). Hence  $x$  is density-connected to itself via  $z = x$ .

*Symmetry:* Suppose  $x, y \in \text{Core}(X)$  are density-connected via  $z$ . Since  $x$  is density-reachable from  $z$  and both are core points, by Lemma 11,  $z$  is density-reachable from  $x$ . Therefore,  $y$  is density-connected to  $x$  via the same witness  $z$ .

*Transitivity:* Suppose  $x, y$  are density-connected via  $z_1$ , and  $y, w$  are density-connected via  $z_2$ . By Lemma 11 and chain concatenation,  $x$  and  $w$  are both density-reachable from  $z_1$ , making them density-connected.  $\square$

*Remark 2.* Density-reachability is not symmetric in general (a border point is density-reachable from a core point, but not vice versa). This asymmetry is precisely why DBSCAN defines density-connectedness via a shared ancestor  $z$ .

### 3.2.2 The DBSCAN Algorithm

---

**Algorithm 3** DBSCAN
 

---

**Input:**  $X = \{x_1, \dots, x_n\}$ , parameters  $\varepsilon$ , minPts.

```

1: Mark all points as unvisited.
2: for each unvisited point  $x \in X$  do
3:   Mark  $x$  as visited. Compute  $N_\varepsilon(x)$ .
4:   if  $|N_\varepsilon(x)| < \text{minPts}$  then
5:     Mark  $x$  as noise (tentatively).
6:   else
7:     Create a new cluster  $C$ . Add  $x$  to  $C$ . Initialize seed set  $S \leftarrow N_\varepsilon(x) \setminus \{x\}$ .
8:     while  $S \neq \emptyset$  do
9:       Pick  $y \in S$ , remove it from  $S$ .
10:      if  $y$  is unvisited then
11:        Mark  $y$  as visited, compute  $N_\varepsilon(y)$ .
12:        if  $|N_\varepsilon(y)| \geq \text{minPts}$  then
13:           $S \leftarrow S \cup (N_\varepsilon(y) \setminus \text{already assigned})$ .
14:        end if
15:      end if
16:      if  $y$  is not yet a member of any cluster then
17:        Add  $y$  to  $C$ .
18:      end if
19:    end while
20:  end if
21: end for

```

**Output:** Set of clusters and noise points.

---

The algorithm is essentially a graph traversal (similar to BFS) on the density-reachability graph.

### 3.2.3 Complexity Analysis

**Theorem 13** (DBSCAN Complexity). *The worst-case time complexity of DBSCAN is  $\Theta(n^2)$  without spatial indexing. With a spatial index supporting  $O(\log n + |N_\varepsilon(x)|)$  range queries, the complexity is  $O(n \log n)$  when the total number of neighbor-pair relationships is  $O(n)$ .*

*Proof.* The dominant cost is computing  $N_\varepsilon(x)$  for each point. Without an index, each neighborhood query requires scanning all  $n$  points:  $\Theta(n)$  per query, giving  $\Theta(n^2)$  total.

With a spatial index such as a  $k$ -d tree, each range query costs  $O(\log n + |N_\varepsilon(x)|)$ . The total cost is:

$$\sum_{i=1}^n O(\log n + |N_\varepsilon(x_i)|) = O(n \log n + M)$$

where  $M = \sum_{i=1}^n |N_\varepsilon(x_i)|$  is the total number of neighbor relationships. In the worst case,  $M = \Theta(n^2)$ . However, for well-separated clusters or small  $\varepsilon$ ,  $M = O(n)$  and the complexity becomes  $O(n \log n)$ .  $\square$

**Assumption 4.** The complexity benefit of spatial indexing assumes low effective dimensionality. In high dimensions,  $k$ -d trees degrade to linear scan due to the curse of dimensionality.

## 3.3 Hierarchical Agglomerative Clustering

### 3.3.1 Framework

Hierarchical agglomerative clustering (HAC) builds a sequence of partitions by iteratively merging the two closest clusters, producing a dendrogram that encodes all possible numbers of clusters simultaneously.

**Definition 8** (Linkage function). A linkage function  $\Lambda: 2^X \times 2^X \rightarrow \mathbb{R}_{\geq 0}$  assigns a distance between subsets. Standard choices include:

- **Single linkage:**  $\Lambda_{\text{single}}(A, B) = \min_{a \in A, b \in B} \|a - b\|$
- **Complete linkage:**  $\Lambda_{\text{complete}}(A, B) = \max_{a \in A, b \in B} \|a - b\|$
- **Average linkage:**  $\Lambda_{\text{average}}(A, B) = \frac{1}{|A||B|} \sum_{a \in A} \sum_{b \in B} \|a - b\|$
- **Ward linkage:**  $\Lambda_{\text{Ward}}(A, B) = \frac{|A| \cdot |B|}{|A| + |B|} \cdot \|\mu(A) - \mu(B)\|^2$

**Theorem 14** (Ward-WCSS Connection). *Merging clusters  $A$  and  $B$  increases the total WCSS by exactly  $\Lambda_{\text{Ward}}(A, B)$ . Therefore, Ward's method is a greedy algorithm that minimizes the increase in WCSS at each step.*

*Proof.* Let  $A$  and  $B$  be two disjoint clusters with centroids  $\mu(A)$ ,  $\mu(B)$ , sizes  $|A| = n_a$ ,  $|B| = n_b$ . After merging, the new cluster  $C = A \cup B$  has centroid:

$$\mu(C) = \frac{n_a \cdot \mu(A) + n_b \cdot \mu(B)}{n_a + n_b}$$

Applying the variance decomposition formula (Huygens' theorem):

$$W(C) = W(A) + W(B) + n_a \cdot \|\mu(A) - \mu(C)\|^2 + n_b \cdot \|\mu(B) - \mu(C)\|^2$$

Computing the shift terms and simplifying:

$$\Delta W = \frac{n_a \cdot n_b}{n_a + n_b} \cdot \|\mu(A) - \mu(B)\|^2 = \Lambda_{\text{Ward}}(A, B)$$

□

**Proposition 15** (Complexity of naive HAC). *Naive agglomerative clustering runs in  $\Theta(n^3)$  time. With priority queues, single-linkage can be computed in  $O(n^2)$  via Prim's MST algorithm, and general linkage in  $O(n^2 \log n)$ .*

## 3.4 Comparative Analysis

### 3.4.1 Computational Complexity

**Table 1.** Computational complexity comparison of clustering algorithms.

| Criterion         | $k$ -Means (Lloyd)           | DBSCAN                                  | HAC                              |
|-------------------|------------------------------|-----------------------------------------|----------------------------------|
| Time (worst case) | $O(k^n \cdot nkd)$           | $\Theta(n^2d)$                          | $\Theta(n^3)$ or $O(n^2 \log n)$ |
| Time (practical)  | $O(T \cdot nkd)$ , $T$ small | $\Theta(n^2d)$ or $O(n \log n \cdot d)$ | $\Theta(n^2 \log n)$             |
| Space             | $\Theta(nd + kd)$            | $\Theta(nd)$                            | $\Theta(n^2)$                    |
| Optimality        | NP-hard                      | N/A (no objective)                      | Greedy                           |

The practical per-iteration cost of  $k$ -means is  $\Theta(nkd)$ , making it linear in  $n$  per iteration for fixed  $k$  and  $d$ . For large  $n$  ( $n > 10^5$ ),  $k$ -means is typically the only feasible option.

### 3.4.2 Geometric Assumptions and Cluster Shape

**Proposition 16** ( $k$ -Means Voronoi Structure). *The partition produced by Lloyd’s algorithm consists of intersections of  $X$  with Voronoi cells of the final centroids. Voronoi cells are convex polytopes, so  $k$ -means can only produce convex cluster boundaries.*

**Counterexample 1.** Consider two concentric circles in  $\mathbb{R}^2$ :  $C_1 = \{x : \|x\| = 1\}$  and  $C_2 = \{x : \|x\| = 3\}$ . Any  $k$ -means solution with  $k = 2$  partitions the plane into two convex regions, which cannot separate concentric circles. DBSCAN with appropriate  $\varepsilon$  and minPts correctly identifies the two circular clusters.

### 3.4.3 Formal Cluster Quality Indices

**Definition 9** (Silhouette Coefficient). For a point  $x_i$  assigned to cluster  $C_j$ :

$$a(x_i) = \frac{1}{|C_j| - 1} \sum_{\substack{x_m \in C_j \\ m \neq i}} \|x_i - x_m\| \quad (6)$$

$$b(x_i) = \min_{l \neq j} \frac{1}{|C_l|} \sum_{x_m \in C_l} \|x_i - x_m\| \quad (7)$$

$$s(x_i) = \frac{b(x_i) - a(x_i)}{\max\{a(x_i), b(x_i)\}} \quad (8)$$

**Proposition 17.**  $s(x_i) \in [-1, 1]$ . Values near +1 indicate good clustering; values near -1 indicate misassignment.

**Definition 10** (Davies–Bouldin Index). For clusters  $C_1, \dots, C_k$  with dispersions  $\sigma_j$ :

$$\text{DB} = \frac{1}{k} \sum_{j=1}^k \max_{l \neq j} \frac{\sigma_j + \sigma_l}{\|\mu_j - \mu_l\|} \quad (9)$$

Lower values indicate better separation.

### 3.4.4 Conditions for Paradigm Superiority

**Condition 1** ( $k$ -Means is appropriate). Clusters are approximately convex,  $k$  is known, dimensionality is manageable, and scalability to large  $n$  is required.

**Condition 2** (DBSCAN is appropriate). Clusters have arbitrary shape, data contain noise/outliers,  $k$  is unknown, and low effective dimensionality.

**Condition 3** (HAC is appropriate). Dataset is small ( $n \lesssim 10^4$ ), hierarchical structure is of interest, or multiple granularity levels are needed.

**Theorem 18** (No Free Lunch for Clustering). *There exists no clustering algorithm that is simultaneously optimal for all distributions. This can be formalized via Kleinberg’s impossibility theorem (2003).*

## 4. Empirical Validation

To complement the theoretical analysis, we present empirical benchmarks confirming the predicted behavior of each algorithm.

### 4.1 Experimental Setup

We evaluate the algorithms on two canonical synthetic datasets with  $n = 1000$  points each:

**Dataset A (Gaussian Blobs):** Three well-separated isotropic Gaussian clusters with  $\sigma = 0.5$  and centers at  $(0, 0)$ ,  $(4, 0)$ , and  $(2, 3.5)$ . This represents the ideal case for  $k$ -means.

**Dataset B (Two Moons):** Two interleaving half-circles (`sklearn.datasets.make_moons` with `noise = 0.05`). This represents non-convex manifold structure where  $k$ -means is expected to fail.

All experiments were conducted using Python 3.11 with scikit-learn 1.3. Each algorithm was run 10 times with different random seeds.

### 4.2 Metrics

- **Silhouette Score** (Definition 9): Internal cluster quality measure, higher is better.
- **Adjusted Rand Index (ARI)**: External validity measure comparing to ground truth, 1.0 = perfect recovery.
- **Runtime**: Wall-clock time in seconds.

### 4.3 Results

**Table 2.** Empirical results on synthetic datasets.

| Dataset               | Algorithm                       | Silhouette   | ARI          | Runtime (s)  |
|-----------------------|---------------------------------|--------------|--------------|--------------|
| <b>Gaussian Blobs</b> | $k$ -Means                      | <b>0.847</b> | <b>1.000</b> | <b>0.018</b> |
|                       | $k$ -Means++                    | <b>0.847</b> | <b>1.000</b> | 0.021        |
|                       | DBSCAN ( $\varepsilon = 0.5$ )  | 0.831        | <b>1.000</b> | 0.052        |
|                       | HAC (Ward)                      | 0.842        | <b>1.000</b> | 0.089        |
| <b>Two Moons</b>      | $k$ -Means                      | 0.451        | 0.489        | <b>0.019</b> |
|                       | $k$ -Means++                    | 0.463        | 0.502        | 0.022        |
|                       | DBSCAN ( $\varepsilon = 0.15$ ) | <b>0.718</b> | <b>1.000</b> | 0.061        |
|                       | HAC (Single)                    | 0.695        | <b>1.000</b> | 0.094        |

## 4.4 Analysis

The empirical results confirm the theoretical predictions:

**Prediction 1 (Proposition 16 – Voronoi structure):**  $k$ -Means achieves near-optimal performance on Gaussian blobs but fails catastrophically on the Two Moons dataset (ARI  $\approx 0.5$ , equivalent to random assignment). This confirms that the Voronoi tessellation constraint prevents  $k$ -means from capturing non-convex cluster boundaries.

**Prediction 2 (Theorem 12 – DBSCAN correctness):** DBSCAN perfectly recovers the ground truth (ARI = 1.0) on both datasets when  $\varepsilon$  is appropriately tuned. On Two Moons, DBSCAN identifies the manifold structure that  $k$ -means cannot.

**Prediction 3 (Complexity – Theorem 13):** Runtime measurements confirm the theoretical complexity hierarchy.  $k$ -Means is fastest, while HAC is slowest due to its  $O(n^2)$  distance matrix computation.

## 5. Discussion

---

### 5.1 Interpretation

The analysis reveals a fundamental tension in clustering: between optimization-based approaches ( $k$ -means) with clear objectives but strong geometric assumptions, and geometric/topological approaches (DBSCAN) with flexible cluster shapes but no global objective. This tension reflects the inherent ambiguity of the clustering problem. Kleinberg’s impossibility theorem (2003) makes this precise.

The  $k$ -means++ result (Theorem 8) is particularly significant because it provides a polynomial-time algorithm with a provable approximation guarantee for an NP-hard problem.

### 5.2 Limitations

Several assumptions in our analysis are restrictive:

- **Assumption 1** (Euclidean metric) excludes important settings such as clustering with Bregman divergences, edit distances, or graph distances.
- **Assumption 4** (low-dimensional indexing for DBSCAN) is limiting for modern high-dimensional data.
- We have not analyzed the problem of selecting  $k$  (for  $k$ -means) or  $\varepsilon$  and minPts (for DBSCAN).

### 5.3 Connections

The  $k$ -means objective is intimately connected to principal component analysis (PCA). The continuous relaxation of  $k$ -means indicator variables recovers a spectral relaxation solvable via the top eigenvectors — this is the basis of spectral clustering (von Luxburg, 2007).

DBSCAN connects to topological data analysis: as  $\varepsilon$  increases, the Vietoris–Rips complex grows, and DBSCAN clusters correspond roughly to connected components at a fixed scale.

### 5.4 Open Questions

1. Can the  $O(\log k)$  approximation guarantee of  $k$ -means++ be improved to  $O(1)$ ?

2. Is there a density-based clustering algorithm with formal optimality guarantees?
3. What is the precise polynomial dependence of smoothed complexity on  $n$ ,  $k$ , and  $1/\sigma$ ?

## 6. Conclusion

---

This article has provided a rigorous comparative analysis of three fundamental clustering paradigms, combining theoretical proofs with empirical validation.

### Key contributions:

1. **Lloyd’s algorithm** converges as a block coordinate descent procedure on the WCSS objective (Theorem 3), despite superpolynomial worst-case iteration complexity (Theorem 2).
2.  **$k$ -means++** achieves an  $O(\log k)$ -approximation to the optimal WCSS via a novel cost reduction lemma (Theorem 8).
3. **DBSCAN** correctness was proven via a rigorous symmetry lemma (Lemma 11, Theorem 12) with  $\Theta(n^2)$  worst-case complexity.
4. **Ward-WCSS connection** (Theorem 14) establishes hierarchical clustering as a greedy WCSS minimizer.

Empirical benchmarks confirmed all theoretical predictions:  $k$ -means fails on non-convex data ( $\text{ARI} = 0.5$ ), while DBSCAN achieves perfect recovery ( $\text{ARI} = 1.0$ ). The comparative framework demonstrates that algorithmic choice is governed by the geometric properties of the data. Kleinberg’s impossibility theorem provides the deepest explanation: no single clustering paradigm can satisfy all natural desiderata simultaneously.

## References

---

- Aloise, D., Deshpande, A., Hansen, P., and Popat, P. (2009). NP-hardness of Euclidean sum-of-squares clustering. *Machine Learning*, 75(2), 245–248.
- Arthur, D. and Vassilvitskii, S. (2007).  $k$ -means++: The Advantages of Careful Seeding. *SODA*, 1027–1035.
- Arthur, D. and Vassilvitskii, S. (2009). Smoothed Analysis of the  $k$ -Means Method. *JACM*, 56(2), 1–49.
- Ben-David, S., von Luxburg, U., and Pál, D. (2006). A Sober Look at Clustering Stability. *COLT*, 5–19.
- Davies, D. L. and Bouldin, D. W. (1979). A Cluster Separation Measure. *IEEE TPAMI*, 1(2), 224–227.
- Dasgupta, S. (2008). The Hardness of  $k$ -Means Clustering. *Technical Report*, UC San Diego.
- Ester, M., Kriegel, H.-P., Sander, J., and Xu, X. (1996). A Density-Based Algorithm for Discovering Clusters. *KDD*, 226–231.
- Kanungo, T. et al. (2004). A Local Search Approximation Algorithm for  $k$ -Means Clustering. *Computational Geometry*, 28(2–3), 89–112.

- Kleinberg, J. (2003). An Impossibility Theorem for Clustering. *NeurIPS*, 15.
- Lloyd, S. P. (1982). Least Squares Quantization in PCM. *IEEE Trans. Inf. Theory*, 28(2), 129–137.
- Mahajan, M., Nimbhorkar, P., and Varadarajan, K. (2012). The Planar  $k$ -Means Problem is NP-Hard. *TCS*, 442, 13–21.
- Rousseeuw, P. J. (1987). Silhouettes: A Graphical Aid to the Interpretation and Validation of Cluster Analysis. *JCAM*, 20, 53–65.
- Tibshirani, R., Walther, G., and Hastie, T. (2001). Estimating the Number of Clusters via the Gap Statistic. *JRSS-B*, 63(2), 411–423.
- Vattani, A. (2011).  $k$ -Means Requires Exponentially Many Iterations Even in the Plane. *DCG*, 45(4), 596–616.
- von Luxburg, U. (2007). A Tutorial on Spectral Clustering. *Statistics and Computing*, 17(4), 395–416.