



ARTICLE

Cluster Overlap as Objective Function

Pasi Fränti^{1,*}, Claude Cariou² and Qinpei Zhao³

¹School of Computing, University of Eastern Finland, Joensuu, 80101, Finland

²Institut d'Electronique et des Technologies du numéRique, University of Rennes—ENSSAT, Lannion, 22305, France

³School of Computer Science and Technology, Tongji University, Shanghai, 200092, China

*Corresponding Author: Pasi Fränti. Email: franti@cs.uef.fi

Received: 10 April 2025; Accepted: 26 August 2025

ABSTRACT: K-means uses the sum-of-squared error as the objective function to minimize within-cluster distances. We show that, as a consequence, it also maximizes between-cluster variances. This means that the two measures do not provide complementary information and that using only one is enough. Based on this property, we propose a new objective function called *cluster overlap*, which is measured intuitively as the proportion of points shared between the clusters. We adopt the new function within k-means and present an algorithm called *overlap k-means*. It is an alternative way to design a k-means algorithm. A localized variant is also provided by limiting the overlap calculation to the neighboring points.

KEYWORDS: Clustering; k-means; overlap measure; within-cluster distance; between-cluster distance; arbitrary-shape clusters

1 Introduction

Clustering aims at grouping data points $X = \{x_1, x_2, \dots, x_n\}$ into c clusters by minimizing an objective function that estimates the goodness of the clusters. We assume that there is a distance function $dist(x_i, x_j)$ between the data points. Clustering is given as the partition index $p(i)$ of which cluster the object x_i is assigned to, and the clusters are represented by their centroids c_j . The most common objective function is to minimize the sum of squared error between the data points and their cluster centroid:

$$SSE = \sum_{i=1}^n dist(x_i, c_{p(i)})^2 \quad (1)$$

K-means [1–3] and many other algorithms [4–9] use this objective function. It requires that we can calculate the centroid (mean). It is trivial in Euclidean space but becomes more complex for distance functions designed for strings, sets, time series, and GPS trajectories [10]. SSE can also be converted to an analytical function that can be optimized by gradient descent [11], coordinate descent [12], or its split-and-merge variant [13].

Another popular class of algorithms is agglomerative clustering. The idea is to build the clustering hierarchically, starting with n clusters and then merging clusters until their number reduces to c . Agglomerative clustering applies a local optimization strategy by considering all pairs of clusters and merging the pair that increases the objective function the least. Its main advantages are better clustering accuracy than k-means and the fact that it is not necessary to calculate the centroids.



Regardless of the objective function, we can simply calculate the merge cost as the difference in objective function value before and after the merge:

$$\Delta f_{a,b} = f_{a+b} - f_a - f_b \quad (2)$$

where f_a and f_b are the objective function values of the clusters a and b before the merge, and f_{a+b} after their merge.

Many heuristics have been used to measure the *merge cost*. The most relevant to this paper is the classical *Ward's method* [14], which minimizes the same objective function as k-means (SSE). In the context of vector quantization, it is known as *Pairwise nearest neighbor* (PNN) [15]. Notably, the merge cost can be calculated only using the cluster sizes (n_a, n_b) and their centroids (c_a, c_b) in $O(1)$ time [15]:

$$\text{MergeCost}(a, b) = \frac{n_a + n_b}{(n_a \cdot n_b)} \cdot \text{dist}(c_a, c_b)^2 \quad (3)$$

A trivial implementation of the algorithm requires $O(n^3)$ time even if the distance matrix is stored. This is due to the time-consuming search for the best pair [16].

A more efficient variant maintains only the pointer (and merge cost) of the best pair for every cluster. It reduces the time complexity close to $O(n^2)$ with linear memory requirement [17]. A faster variant uses k-nearest neighbor graph, which reduces the time complexity to $O(n \log n)$ at a minor increase in the SSE-values: 0.0%, 0.2%, 2.6% (Birch datasets), and 1.0%, 2.6%, 4.3% (image datasets) [18].

Other objective functions have also been considered. Cut-based methods include *Normalized cut* [19] and *Cut ratio* [20] based on between-cluster distances. Graph-based methods first create a neighborhood graph and then use the links between the data points to define a cut. *Conductance* [21] calculates the ratio of between-cluster links and all links. *Modularity* [22] calculates the ratio of between-cluster and within-cluster distances. These objective functions have also been adopted into k-means when data is represented by a graph [23]. Between-cluster distances of each merge were measured during the agglomerative clustering process and then normalized by a generic function based on the cluster sizes [24].

Objective functions also exist in cases where the number of clusters is unknown. These are called *internal cluster validity indexes* and are used to compare clustering results with different values of c . Examples of indexes that work reasonably well according to our experiments include *Calinski-Harabasz* [25], *Silhouette coefficient* [26], *WB-index* [27], and *kCE-index* [28]. They all measure the ratio of within-cluster and between-cluster distances normalized by the number of clusters in a slightly different way.

The most striking observation is that almost all indices are based on the same two variables: *within-cluster* and *between-cluster* distances [27]. In this paper, we show the properties and relationship of these two variables. In particular, we show that the between-cluster distances can be calculated directly from the within-cluster distances and vice versa. Thus, the two variables are not mutually exclusive. Instead, they measure essentially the same property.

To demonstrate their relationship, we introduce a new overlap objective function based on the between-cluster distances. To keep the computation reasonable, we simplify the function using an overlap measure. We then apply the new objective function within k-means and compare the experimental results of the new Overlap k-means against the standard k-means based on within-cluster distances. We design a k-means based algorithm for the new objective function. A localized variant is also considered for clusters with arbitrary shapes.

Our results are as follows. Between-cluster distances provide an alternative way to construct k-means with similar results but higher computational complexity. The proposed overlap measure addresses this

problem by emphasizing the central points in the cluster. It has the advantage of providing more accurate centroid locations.

However, there is a drawback. The border points are important for two reasons. First, they allow centroids to travel from one cluster to another. Second, the far-away points play a significant role in pulling the centroids, enhancing their movement. The overlap k-means without this pulling force loses the dynamics of k-means and becomes too rigid.

The rest of the paper is organized as follows. We first discuss the properties of the SSE objective function in [Section 2](#) and show the connection between the within (inner) and between (outer) cluster distances. We then form a k-means variant using the between-cluster distances in [Section 3](#). It has mainly theoretical interest, but we also show how it can be used to create a new localized objective function that can potentially detect arbitrary shapes. The experimental setup is defined in [Section 4](#), and results are given in [Section 5](#). Conclusions are drawn in [Section 6](#).

2 Properties of the Objective Function

The intuitive goal of clustering is to minimize the distances within the clusters and maximize those between the clusters. These can be measured by the following three quantities:

- Sum of squared within-cluster distances (*SSW*)
- The sum of squared between-cluster distances (*SSB*)
- All pairwise (squared) distances (*APD*)

They can be calculated as follows:

$$SSW = \sum_{i,j} dist(x_i, x_j)^2 \forall i, j | p(i) = p(j) \quad (4)$$

$$SSB = \sum_{i,j} dist(x_i, x_j)^2 \forall i, j | p(i) \neq p(j) \quad (5)$$

APD is the sum of the other two:

$$APD = SSW + SSB = \sum_{i,j} dist(x_i, x_j)^2 \quad (6)$$

SSW and *SSB* depend on the clustering result, whereas *APD* is constant and depends only on the data. *SSW* can also be calculated for each cluster separately:

$$SSW = \sum_{a=1}^c SSW_a \quad (7)$$

where SSW_a denotes the sum of squared distances within cluster a :

$$SSW_a = \sum_{i,j} dist(x_i, x_j)^2 \forall i, j | p(i) = p(j) = a \quad (8)$$

2.1 Use of *SSW* and *SSB*

SSW is widely used in agglomerative clustering by calculating the change (Δ) in *SSW* if any pair of clusters a and b are merged:

$$\Delta SSW_{a,b} = SSW_{a+b} - SSW_a - SSW_b = SSB_{a+b} \quad (9)$$

Here, it is important to notice that the change equals the sum of squared distances between the clusters, SSB_{a+b} , defined as follows:

$$SSB_{a+b} = \sum_{i,j} dist(x_i, x_j)^2 \forall i, j | p(i) = a \wedge p(j) = b \quad (10)$$

In other words, while the algorithm aims to minimize SSW , it proceeds merely by evaluating the between-cluster values (SSB). SSW can also be normalized by calculating the sum of distances divided by the cluster size:

$$\widehat{SSW} = \sum_{a=1}^c \frac{SSW_a}{n_a} \quad (11)$$

This version is used in Ward's method and is shown to equal SSE [15]. The unnormalized version in (10) has been shown to provide more balanced cluster sizes compared to SSE [29]. Average linkage clustering eliminates the cluster sizes completely from the equation by dividing by $n_a \cdot n_b$. Its behavior is opposite to (10) by creating more biased cluster sizes and favoring outlier clusters.

Normalized cut [19], cut ratio [20], and conductance [21] are all based on SSB but normalized in different ways. The ratio of SSW and SSB is used in [30], and their difference in [22]. Calinski and Harabasz [25] and WB-index [27] also use the ratio of SSW and SSB , but they also include c in the equation to allow comparison with different numbers of clusters. This is a useful variation when the number of clusters also needs to be solved.

2.2 Relationship between SSE and SSW

In the case of Euclidean distance, SSE is also directly related to SSW according to Huygens's theorem [29,31]:

$$SSW_a = \sum_{x_i, x_j \in a} dist(x_i, x_j)^2 = n_a SSE_a \quad (12)$$

where SSW_a and SSE_a refer to the SSW and SSE values in cluster a , and n_a its size. The main consequence of this relationship is that minimizing SSE and SSW within a cluster differs only by a scaling factor n_a . In other words, SSE_a sums up n_a squared distances, whereas SSW_a sums up n_a^2 distances. It is possible to normalize them into the same scale simply as $SSE_a = SSW_a / n_a$.

Using (12), we can reformulate Eq. (6) as:

$$SSB = APD - SSW = APD - \sum_{a=1}^c n_a SSE_a \quad (13)$$

Since APD is constant, regardless of the clustering, minimizing SSE (and SSW) is effectively the same as maximizing SSB . This implies that the objective function does not need to contain SSE or SSW , but SSB alone is sufficient.

If clusters are of equal sizes ($n_a = n_b$, for all a, b), then it simplifies even further to:

$$SSB = APD - n_a \cdot \sum_{a=1}^c SSE_a = APD - \frac{n}{c} \cdot SSE \quad (14)$$

Fig. 1 shows an example of two clusters with their corresponding SSW and SSB values.

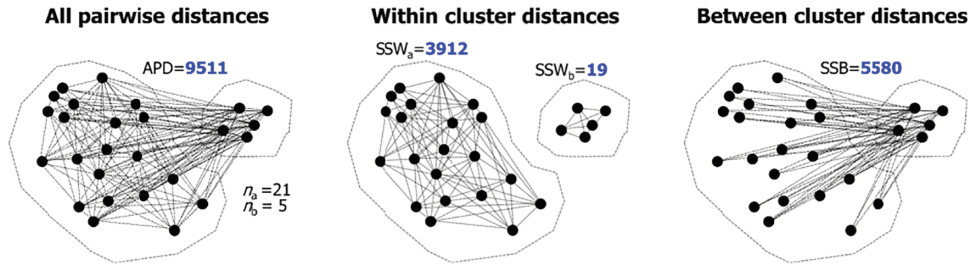


Figure 1: Example of *APD* (left), *SSW* (middle), and *SSB* (right). There are $O(n^2)$ links, but not all of them are drawn. The links have the following relationship: $\{APD\} = \{SSW\} \cup \{SSB\}$

2.3 Number of SSB Calculations

Motivated by the observations, we reformulate the k-means algorithm using between-cluster distances (*SSB*). However, this can be time-consuming because there are $O(n^2)$ pairwise distances to be calculated at every iteration, which is too slow. An alternative solution is to store the $O(n^2)$ size distance matrix, but this would reduce the practicality of the method due to quadratic memory consumption. Neither option looks very promising for scalability.

Fortunately, most *SSB* distances are between points located far away from each other, so they will almost never be assigned to the same cluster. In other words, they have no contribution to clustering. The distances that matter are points near each other so that they can sometimes belong to the same cluster depending on the operation of the algorithm. They contribute to the objective function most.

One possibility to make the method faster is to limit the distance calculations only between points in neighboring clusters. However, if the number of clusters is small ($c \ll n$), almost all clusters are neighbors. This can also happen when the dimension increases high because the distances tend to become equidistant when using Euclidean (L_2) distance [32]. Minkowski distances (L_p) with other values of p have also been considered and L-norms with fractional values ($0 < p < 1$) have been shown to work better for high-dimensional data [33].

To demonstrate the problem, assume that the points are equally divided so that all clusters have size n/c points. In this case, there are $(n/c)^2$ distances between every pair of clusters. If we compute distances only between the nearest cluster, there would still be $c \cdot (n/c)^2 = n^2/c$ distance calculations. If $c = \sqrt{n}$, the number of distance calculations becomes $n^{1.5}$, but when c is a small constant, $c = O(1)$, the time complexity remains $O(n^2)$. See Fig. 2 for an example. In other words, we cannot afford to calculate all between-cluster distances if we want to keep the algorithm fast.

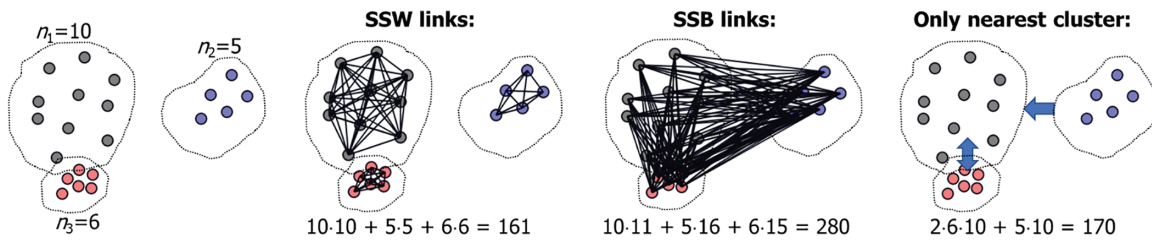


Figure 2: Total number of *SSB* calculations (280) and the number of calculations if only considering the nearest cluster (170)

2.4 Euclidean Distances

When minimizing sum-of-squared *Euclidean* distances, the calculations can be done more efficiently. From (12), we know that the sum of squared distances (SSE), multiplied by n , equals all pairwise distances within the clusters (SSW):

$$SSW_a = n_a \cdot SSE_a \quad (15)$$

This is already utilized in k-means by using the centroids to calculate the within-cluster distances (SSW). However, the result also generalizes between-cluster distances as shown in [27,34].

$$SSB = APD - \sum_{a=1}^c n_a \cdot SSE_a = \sum_{a=1}^c n_a \cdot \left(\frac{APD}{n} - SSE_a \right) = \sum_{a=1}^c \left(\frac{n_a}{n} \cdot APD - SSW_a \right) = APD - \sum_{a=1}^c SSW_a \quad (16)$$

In other words, both SSW and SSB can be done in $O(n)$ time via SSE. In the following, we do not utilize this property as we are interested in overlap. This requires point-level information that would be lost if we used centroids to represent the clusters. However, this property can be beneficial for others, and thus, it is presented here.

3 Overlap k-Means (OK)

Among all the point pairs, the most meaningful ones are those that are near to each other but in different clusters. In other words, border pixels. In the proposed method, we, therefore, calculate only one distance per point to its nearest neighbor in another cluster. This distance is put into the context of its relative location in its own cluster.

3.1 Overlap Value

We formulate a so-called *overlap value* for each data point based on the overlap measure proposed in [35]. The overlap score for a point x_i is calculated as follows:

$$O_1(x_i) = \frac{\text{dist}(x_i, c_{p(i)})}{\text{dist}(x_i, \overline{NN}(x_i))} \quad (17)$$

where $\overline{NN}(x_i)$ is the nearest neighbor of x_i in another cluster:

$$\overline{NN}(x_i) = \text{argmin} \text{dist}(x_i, x_j), \forall x_j | p(i) \neq p(j) \quad (18)$$

The lower the overlap value, the better the cluster assignment of the point. The new objective function then minimizes the sum of all overlap values in the data:

$$\text{Overlap} = \sum_{i=1}^n O_1(x_i) \quad (19)$$

The overlap value includes two terms. The first term (numerator) represents a sample of the SSW function, and the second term (divisor) is a sample of the SSB function. The overlap is, therefore, a strongly sampled version of the Calinski and Harabasz measure [25] and closely related to other sum-of-squared based indexes [27]. The difference is that the overlap measure is not normalized by the cluster size.

Example calculations are shown in Fig. 3, where the large cluster (gray points) has points near the small dense cluster (red points). While all border points have high distances to the centroids, only those near the red cluster have a high overlap value.

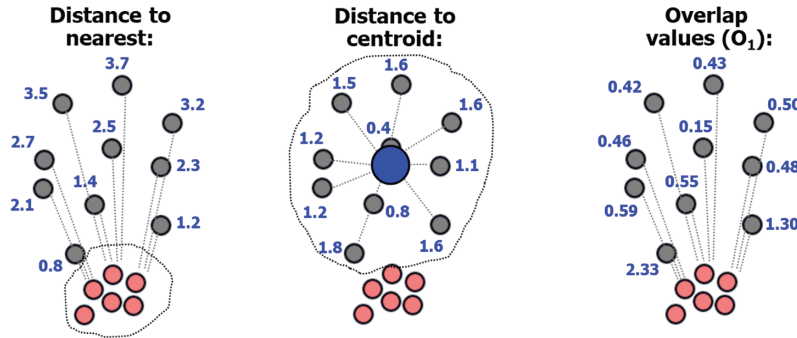


Figure 3: The proposed objective function calculates the distance from a point only to its nearest neighbor in another cluster and to its own centroid

3.2 K-Means Using Overlap

K-means has two steps that should be updated according to the new objective function. Theoretically, the nearest cluster should be calculated by minimizing (16). However, we face computational problems. Finding the nearest neighbor point takes $O(n)$ time for one point, and we have n points to process. The time complexity of the partition would, therefore, grow to $O(n^2)$. To avoid this, we keep the original k-means assignment and find the nearest centroid.

The second step is to calculate the centroid for the clusters. We calculate the weighted average so that every point is inversely weighted by its overlap value. In this way, the border points with high overlap values will have less impact on the centroid location, which will move closer to the non-border points. The exact update rule is as follows:

$$c_j = \frac{\sum_{p(i)=j} w_i \cdot x_i}{\sum_{p(i)=j} w_i} \quad (20)$$

$$w_i = \exp \left[- (O_1(x_i) / \gamma)^2 \right] \quad (21)$$

where γ is a (constant) bandwidth parameter. In other words, the calculation of the cluster centroid is an overlap-weighted average of the points belonging to the same cluster. The only difference to traditional k-means is the addition of the Step 2 (calculation of overlap). Pseudo code of the Overlap k-means is given in Algorithm 1.

Algorithm 1: Overlap K-means (OK)

OK(X, c) $\rightarrow \{p_i\}$

Input: X, c, γ

Output: $p_i, 1 \leq i \leq n, p_i \in \{1, \dots, c\}$

– Select c points randomly from X as centroids $c_j, 1 \leq j \leq c$

– While $\exists i: p_i \neq p_i^{\text{old}}$

 Step 1: For each x_i , assign label p_i as the index of its nearest centroid;

 Step 2: For each x_i , calculate the overlap value $O_1(x_i)$ according to (17);

 Step 3: For each cluster, calculate new centroid according to (20).

3.3 Localized Variant

Despite the alternative formulation, Overlap k-means is merely another variant of k-means that optimizes essentially the same variables: minimizing SSW and maximizing SSB. The overlap score itself is localized regarding the SSB calculations by considering only the nearest neighbor in another cluster. The motivation was merely to reduce the calculations.

Next, we extend the localization to the SSW part by removing the dependency on the centroid. We want the overlap to be based only on local properties. To achieve this, we calculate a local neighborhood of a point defined by its k -nearest neighbors (KNN). *Localized* overlap value is then defined as:

$$O_2(x_i) = \frac{\text{Meanshift}(x_i)}{\text{dist}(x_i, \overline{NN}(x_i))} \quad (22)$$

where $\text{Meanshift}(x_i)$ denotes the distance of x_i to the mean of its KNN neighbors in the *same cluster*. It is calculated as follows:

$$\text{Meanshift}(x_i) = \frac{\sum_{y \in \text{KNN}(x_i)} \text{dist}(x_i, y)}{k} \quad (23)$$

where $\text{KNN}(x_i)$ is the set of k nearest neighbors of x_i in the same cluster.

The divider is the same in both O_1 and O_2 , i.e., the distance to the nearest neighbor in another cluster. However, the numerator is different. In O_2 , it depends on the local neighborhood alone no longer requires centroid. The overlap score O_1 minimizes SSW, which implies spherical clusters, whereas O_2 minimizes *Meanshift*, which allows clusters with arbitrary non-convex shapes.

It is both an advantage and a disadvantage. The idea is that the higher the mean-shift distance, the more points differ from their neighborhoods. In specific, border points have high mean-shift values. This idea was originally developed for outlier detection in [36]. An example is shown in Fig. 4. Pseudo code of the local variant is given in Algorithm 2.

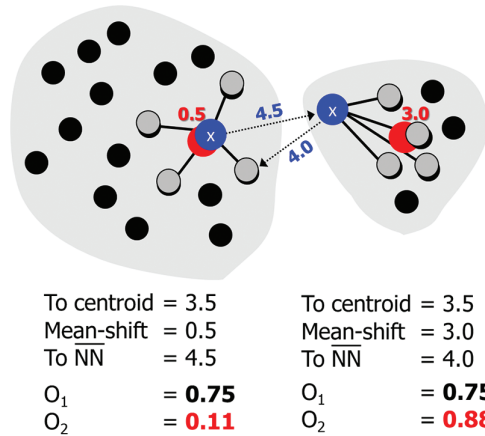


Figure 4: Two points (blue) having the same distance to their cluster centroid (3.5), almost the same distance to the nearest point in another cluster (4.5 and 4.0), but with very different mean shift values (0.5 and 3.0)

Algorithm 2: Localized variant (OK-local)OK-local(X, k) $\rightarrow \{p_i\}$ Input: X, c, γ Output: $p_i, 1 \leq i \leq n, p_i \in \{1, \dots, c\}$

- Select c points randomly from X as centroids $c_j, 1 \leq j \leq c$
- While $\exists i: p_i \neq p_i^{\text{old}}$
 - Step 1: For each x_i , assign label p_i as the index of its nearest centroid;
 - Step 2: For each x_i , calculate the overlap value $O_2(x_i)$ according to (17);
 - Step 3: For each cluster, calculate new centroid according to (20).

Note: A more detailed pseudo code is in [Appendix A](#).**4 Experimental Setup**

The datasets used in the experiments are listed in [Table 1](#) and demonstrated in [Fig. 5](#). We selected datasets from the clustering basic benchmark [35] and from [37]. The datasets have varying cluster overlaps, shapes, and densities.

We cluster the datasets with the following algorithms:

- K-means (KM)
- K-means++ (KM++)
- Random swap (RS)
- Density peaks (DensP)
- DBSCAN
- Spectral clustering (Spectral)
- Overlap k-means (OK)
- Localized variant (OK-local)

Table 1: Datasets used in the experiments

Dataset	Reference	n	d	C
Flame	[38]	240	2	2
Aggregation	[39]	788	2	7
S1	[40]	5000	2	15
S2	[40]	5000	2	15
S3	[40]	5000	2	15
S4	[40]	5000	2	15
A1	[34]	3000	2	20
A2	[34]	5250	2	35
A3	[34]	7500	2	50
Dim32	[18]	1024	32	16
Unbalance	[41]	6500	2	8
Birch1	[42]	100,000	2	100
Birch2	[42]	100,000	2	100
Birch3	[42]	100,000	2	100

(Continued)

Table 1 (continued)

Dataset	Reference	n	d	C
Worms-2d	[37]	105,600	2	35
Worms-64d	[37]	105,600	64	25

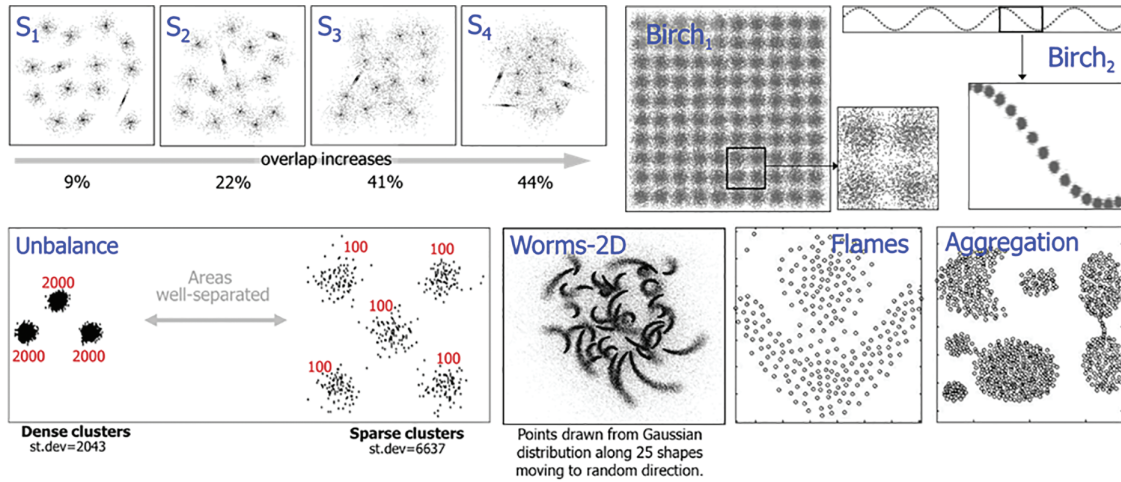


Figure 5: Visualization of the selected datasets. Some datasets are seemingly simple with regular structure ($Birch_{1-3}$), whereas others have challenges like varying cluster overlap (S_1 - S_4), density (Unbalance), or shapes (Worms-2D, Worms-64D)

K-means is the main point of comparison to see whether Overlap k-means provides similar clustering results to the standard k-means with the same seeding method. K-means itself can be improved by simple tricks like repeats and better seeding [43]. We, therefore, also tested the most popular seeding called K-means++ [44], which is based on a randomized further point heuristic.

We also include random swap [5], which is almost as simple as k-means and provides a state-of-the-art reference to how good clustering accuracy one can expect by minimizing SSE. It is one of the few clustering algorithms that find the correct cluster allocation for all the benchmark data. Other algorithms included are density peaks [45], DBSCAN [46,47], and spectral clustering [48].

We run all algorithms 100 times and report the average results, with the exceptions of Birch and Worms datasets, which are repeated only 10 times due to their higher processing times. Clustering results are measured by clustering accuracy (ACC) [49] and centroid index (CI) [50].

We have two main hypotheses to test: (1) Does OK provide similar results to k-means, and (2) can OK-local find clusters having non-uniform shapes and density?

5 Results

The main results are summarized in Tables 2 and 3. Good algorithms like Random swap and Density peaks can find correct clustering (CI = 0) for most datasets. They only fail with the non-spherical datasets $Birch_3$, Worms, and Aggregation, which cannot be clustered correctly by the SSE objective function. K-means++ manages to solve Dim32 and Unbalance, whereas K-means does not solve any of the datasets (average CI = 5.8).

Table 2: Clustering results measured by centroid index (CI)[†]

	S1	S2	S3	S4	A1	A2	A3	Dim32	Unb	B1	B2	B3	W2	W64	Agg	Av
KM	2.0	1.1	1.2	1.3	2.0	4.2	6.5	3.5	3.5	8.8	16.1	24	9.0	10.2	0.7	5.8
KM++	0.2	0.5	0.6	0.4	0.6	0.7	1.7	0.0	0.0	3.3	1.5	21	8.2	9.5	0.7	3.1
RS	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	16	7.0	3.0	1.0	1.7
DensP	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	21	8.0	3.0	1.0	2.1
DBSCAN	11	9	18	22	12	20	28	0.0	1.0	11	5	38	9	11	2	12.4
Spectral	0.0	0.0	0.4	1.1	0.0	0.0	0.3	1.2	2.1	1.0	0.0	18	6.0	0.0	2.0	2.1
OK	2.3	2.3	2.5	2.2	2.9	6.0	7.5	4.0	4.0	12	20	31	11	11.3	1.8	7.6
OK-local	2.1	1.2	1.3	0.9	2.5	4.7	5.6	3.7	4.0	8.3	16	24	8.0	9.9	1.0	5.8

Note: Perfect clustering result (CI = 0) is emphasized by boldface. Average results are shaded.

Table 3: Clustering results measured by clustering accuracy (ACC)[†] in %

	S1	S2	S3	S4	A1	A2	A3	Dim32	Unb	B1	B2	B3*	W2	W64	Agg	Av
KM	87	90	80	74	90	87	86	82	60	87	84	–	51	52	82	74
KM++	98	94	83	78	97	98	96	100	100	95	98	–	52	54	83	82
RS	99	97	86	80	100	100	100	100	100	97	100	–	52	76	83	85
DensP	99	97	85	79	98	99	99	95	100	100	100	–	55	74	86	85
DBSCAN	95	77	56	58	88	84	85	91	98	78	68	–	34	25	87	62
Spectral	99	96	86	77	98	99	99	94	58	95	100	–	52	85	74	81
OK	83	82	74	71	84	82	84	75	58	83	78	–	50	48	73	68
OK-local	86	89	79	76	87	86	88	77	60	88	84	–	52	53	83	73

Note: *Ground truth partitions are not available. Perfect clustering result (ACC = 100%) is emphasized by boldface. Average results are shaded.

Beyond the clustering accuracy, the interest here is more on the comparison between k-means (KM) and its overlap variant (OK). The results of these two are indeed similar, but there is a noticeable difference favoring k-means. The localized variant provides results closer to k-means. The corresponding average CI values are 5.8 (K-means), 7.6 (OK), and 5.8 (Local-OK).

Our first hypothesis was that Overlap k-means would provide similar results to k-means, but the results do not support it. We analyze the reasons in Fig. 6 with the S1 dataset, which is easy for good algorithms but troublesome for k-means due to small cluster overlap [35]. The results show that k-means can move all centroids to their correct location except two. Overlap k-means, however, performs even worse with four clusters unsolved.

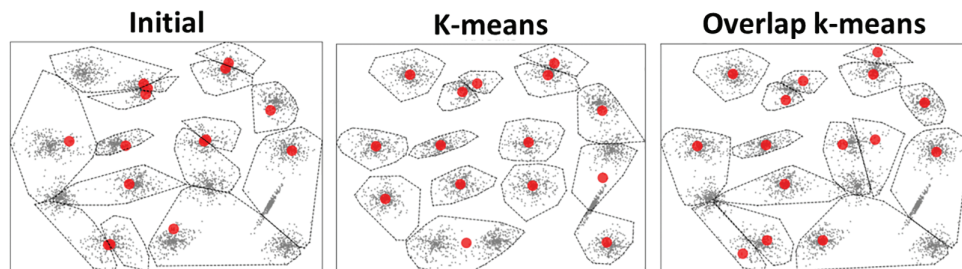


Figure 6: K-means and Overlap k-means result with S1 for the same initial solution (left). Overlap k-means positions the centroids better within the clusters but has problems locating them correctly globally. The CI-values of the solutions are CI = 2 (k-means) and CI = 4 (Overlap k-means) accordingly

The reason is studied in Fig. 7, where the pulling effect of the faraway points attracts a centroid in just three iterations with k-means. However, the same does not happen with Overlap k-means. The reason is due

to the low weight of the faraway points in the centroid calculation. This leads to nicely balanced locations of the centroids *within* the same cluster but leaving the one above empty.

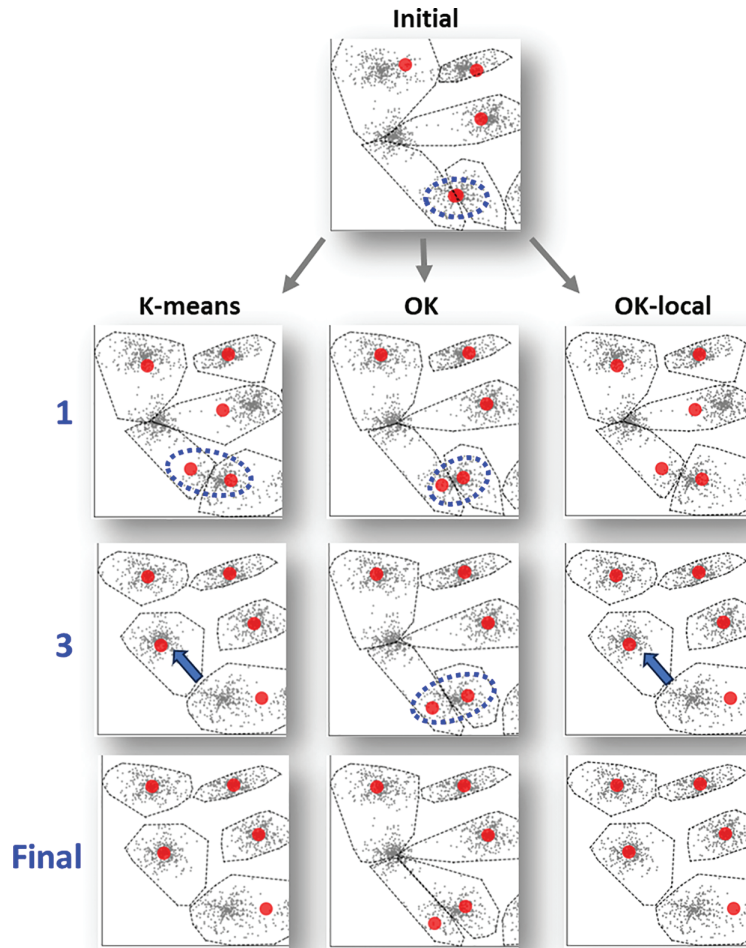


Figure 7: Snapshots of the bottom left corner of S1 after the 1st iteration, 3rd iteration, and the final solution. The pulling effect of the faraway points in k-means manages to move one of the centroids to the cluster missing a centroid. Overlap k-means gives less weight to the border points, which stabilizes the centroids close to their original location and makes the algorithm less dynamic

We conclude that it is possible to design k-means via the between-cluster distances, but the chosen design using the overlap is not good. It makes the algorithm practical by avoiding the overwhelming computations of the between-cluster distances, but it significantly reduces the pulling effect of the remote points, which is vital for k-means. In other words, the algorithm becomes too rigid and tends to get stuck in a local minimum. An effective algorithm should do the opposite: enhance the vital faraway boundary points instead of restricting their use.

The second hypothesis was that the localized variant could also work for non-spherical datasets. This does not happen either. The results of OK-local are closer to k-means than those of OK, but they do not perform any better with the non-spherical datasets (Birch3, Worms, and Aggregation). We cannot conclude whether the idea is not good enough or the potential is not fulfilled by the inferior local optimization.

We did not carry out further tests other than this proof-of-concept. The main contribution is theoretical, showing the connection between the within-cluster and between cluster distances. If one wants to achieve a better algorithm based on the overlap concept, we outline two potential ways to do it here.

1. Design a stochastic variant starting with the opposite: give high weight for the high overlap points. This would further enhance the dynamics of k-means, hopefully pushing the centroids even further than the standard SSE function would do. Then, decrease the weights and gradually converge to the Overlap k-means variant, which would focus on the central points.
2. Replace k-means with a better local optimizer: random swap. This would completely avoid the problem of local minima. If there are any benefits of the OK or OK-local, their random swap counterparts have a higher potential to exploit them.

6 Conclusions

We have the following conclusions:

First, within-cluster (SSW) and between-cluster distances (SSB) measure essentially the same quantity and do not offer any complementary information. K-means uses the sum-of-squared (SSE) objective function, which is essentially a scaled version of SSW. We have shown that the same objective can be constructed using between-cluster distances.

Second, if we calculate the centroids for the data, both approaches can be calculated in $O(n)$ time [27,34]. Otherwise, we need to calculate all pairwise distances. Within-cluster distances (SSW) take $k \cdot (n/k)^2 = O(n^2/k)$ time in case of balanced cluster sizes. Between-cluster distances, on the other hand, require the calculation of all pairwise distances in the entire data, which leads to $O(n^2)$ time complexity.

Third, Calinski–Harabasz, silhouette score, WB-index, and DBI and most other internal indexes assume that both the within-cluster and between-cluster variances are needed. This paper demonstrates that this is not the case. The consequence is that all these indexes are heuristics based on wrong assumption, which implies that a better approach should be designed based on a different principle.

Fourth, we presented a new algorithm called overlap k-means (OK), which utilizes between-cluster distances. Instead of minimizing all distances, it focuses on the distances near the cluster borders. However, this construction lost the dynamics of the k-means algorithm, leading to a more rigid algorithm behavior. The results of the localized variant were more consistent with that of k-means.

Fifth, the overlap idea has further potential. We outlined two possibilities. The first one would be a stochastic variant, which would emphasize the border points more in the beginning and gradually reduce their effect, leading to the Overlap k-means. The second possibility would be simply to replace the inferior k-means with a more robust random swap algorithm. These two are points for future studies.

Acknowledgement: We thank Fei Wang and Le Li for their help.

Funding Statement: The authors received no specific funding for this study.

Author Contributions: Study conception and design: Pasi Fränti; theoretical part: Pasi Fränti and Qinpei Zhao; algorithm design and experiments: Claude Cariou; analysis and interpretation of results: Pasi Fränti and Claude Cariou; draft manuscript preparation: Pasi Fränti and Claude Cariou. All authors reviewed the results and approved the final version of the manuscript.

Availability of Data and Materials: The source code and the datasets generated during and/or analyzed during the current study are available in the <https://cs.uef.fi/ml/software/>, <https://cs.uef.fi/sipu/datasets/> (accessed on 05 September 2025).

Ethics Approval: Not applicable.

Conflicts of Interest: The authors declare no conflicts of interest to report regarding the present study.

Appendix A: Pseudo Code of the Overlap k-Means

```

X = dataset
K = Size of KNN
NN = Nearest neighbors
NC = number of clusters
PerformOK(X, NC, K)
{
/* initial solution */
NN:= FindKNN(X,K);
C:= SelectRandomRepresentatives(X,NC);
P:= OptimalPartition(C,X);
P_old:= zeros(size(P));
O:= Overlap-v1(P,X,NN,C);
gamma = Mean(O);
WHILE P <> P_old
{
P_old = P;
(P,C,O):= OK-means(P,C,X,O,NN,gamma);
}
RETURN (P,C,O);
}
OK-means(P,C,X,O,NN,gamma)
{
/* performs one K-means iteration with KNN-overlap weighting */
C:= OptimalRepresentatives_w_Overlap(P,X,O,gamma);
P:= OptimalPartition(C,X);
O:= Overlap-v1(P,X,NN,C);
RETURN (P,C,O);
}
OptimalPartition(C,X)
{
FOR i:= 1 TO N DO
{
P[i]:= FindNearestRepresentative(C,X[i]);
}
}

```

```

RETURN P;
}
FindNearestRepresentative(C,x)
{
j: = 1;
FOR i: = 2 TO k DO
{
IF Dist(x,C[i]) < Dist(x,C[j]) THEN
{
j: = i;
}
}
RETURN j;
}
Overlap-v1(P,X,NN,C)
{
FOR i: = 1 TO N DO
{
B: = NN[i];
Q: = Points in X with label P[i];
T: = B  $\cap$  (X \ Q) /*The nearest neighbors having a different label
D2: = SmallestDistance(T,X[i]); /*Smallest distance between X[i] and T*/
D1: = Dist(X[i],C[P[i]]); /*Dist. between X[i] and its centroid*/
O[i]: = D1/D2; /*Overlap cost*/
}
RETURN O;
}
OptimalRepresentatives_w_Overlap(P,X,O,gamma)
{
/* initialize Sum[1..k] and Count[1..k] by zero values! */
/* sum vector and count for each partition */
FOR i: = 1 TO N DO
{
j: = P[i];
Sum[j]: = Sum[j] + X[i]*exp(-(O[i]/gamma)**2);
Count[j]: = Count[j] + exp(-(O[i]/gamma)**2);
}
/* optimal representatives are average vectors */

```



```

FOR i: = 1 TO k DO
{
IF Count[i] <> 0 THEN
{
C[i]: = Sum[i]/Count[i];
}
}
RETURN C;
}
SmallestDistance(Data,Query)
{
Dmin: = some large positive value;
FOR i: = 1 TO N DO
{
D: = Dist(Data[i],Query);
IF D < Dmin THEN
{
Dmin: = D;
}
}
RETURN Dmin;
}
AvgDistance(D)
{
Sum = 0;
FOR i: = 1 TO K DO
{
Sum = Sum + D[i];
}
Davg: = Sum/K;
RETURN Davg;
}

```

References

1. Forgy E. Cluster analysis of multivariate data: efficiency vs. interpretability of classification. *Biometrics*. 1965;21:768–80.
2. MacQueen J. Some methods for classification and analysis of multivariate observations. In: *Berkeley Symposium on Mathematical Statistics and Probability*; 1965 Jun 21–Jul 18; Berkeley, CA, USA. p. 281–97.

3. Lloyd S. Least squares quantization in PCM. *IEEE Trans Inf Theory*. 1982;28(2):129–37. doi:10.1109/TIT.1982.1056489.
4. Fränti P. Genetic algorithms for large-scale clustering problems. *Comput J*. 1997;40(9):547–54. doi:10.1093/comjnl/40.9.547.
5. Fränti P. Efficiency of random swap clustering. *J Big Data*. 2018;5(1):13. doi:10.1186/s40537-018-0122-y.
6. Malinen MI, Mariescu-Istodor R, Fränti P. K-means: clustering by gradual data transformation. *Pattern Recognit*. 2014;47(10):3376–86. doi:10.1016/j.patcog.2014.03.034.
7. Fritzke B. Breathing k-means. *arXiv:2006.15666*. 2020.
8. Lee J, Perkins D. A simulated annealing algorithm with a dual perturbation method for clustering. *Pattern Recognit*. 2021;112(4):107713. doi:10.1016/j.patcog.2020.107713.
9. Baldassi C. Recombinator-k-means: an evolutionary algorithm that exploits k-means for recombination. *IEEE Trans Evol Comput*. 2022;26(5):991–1003. doi:10.1109/TEVC.2022.3144134.
10. Fränti P, Mariescu-Istodor R. Averaging GPS segments competition 2019. *Pattern Recognit*. 2021;112(2):107730. doi:10.1016/j.patcog.2020.107730.
11. Malinen MI, Fränti P. Clustering by analytic functions. *Inf Sci*. 2012;217:31–8. doi:10.1016/j.ins.2012.06.018.
12. Nie F, Xue J, Wu D, Wang R, Li H, Li X. Coordinate descent method for kk -means. *IEEE Trans Pattern Anal Mach Intell*. 2022;44(5):2371–85. doi:10.1109/TPAMI.2021.3085739.
13. Qu F, Shi Y, Yang Y, Liu Y, Hu Y. Coordinate descent K-means algorithm based on split-merge. *Comput Mater Contin*. 2024;81(3):4875–93. doi:10.32604/cmc.2024.060090.
14. Ward JH Jr. Hierarchical grouping to optimize an objective function. *J Am Stat Assoc*. 1963;58(301):236–44. doi:10.1080/01621459.1963.10500845.
15. Equitz WH. A new vector quantization clustering algorithm. *IEEE Trans Acoust Speech Signal Process*. 1989;37(10):1568–75. doi:10.1109/29.35395.
16. Shanbehzadeh J, Ogunbona PO. On the computational complexity of the LBG and PNN algorithms. *IEEE Trans Image Process*. 1997;6(4):614–6. doi:10.1109/83.563327.
17. Fränti P, Kaukoranta T, Shen DF, Chang KS. Fast and memory efficient implementation of the exact PNN. *IEEE Trans Image Process*. 2000;9(5):773–7. doi:10.1109/83.841516.
18. Fränti P, Virtamäki O, Hautamäki V. Fast agglomerative clustering using a k-nearest neighbor graph. *IEEE Trans Pattern Anal Mach Intell*. 2006;28(11):1875–81. doi:10.1109/TPAMI.2006.227.
19. Shi J, Malik J. Normalized cuts and image segmentation. *IEEE Trans Pattern Anal Mach Intell*. 2000;22(8):888–905. doi:10.1109/34.868688.
20. Hagen L, Kahng AB. New spectral methods for ratio cut partitioning and clustering. *IEEE Trans Comput Aided Des Integr Circuits Syst*. 1992;11(9):1074–85. doi:10.1109/43.159993.
21. Leskovec J, Lang KJ, Mahoney M. Empirical comparison of algorithms for network community detection. In: *Proceedings of the 19th International Conference on World Wide Web*; 2010 Apr 26–30. Raleigh, NC, USA. p. 631–40. doi:10.1145/1772690.1772755.
22. Newman MEJ, Girvan M. Finding and evaluating community structure in networks. *Phys Rev E*. 2004;69(2):026113. doi:10.1103/physreve.69.026113.
23. Sieranoja S, Fränti P. Adapting k-means for graph clustering. *Knowl Inf Syst*. 2022;64(1):115–42. doi:10.1007/s10115-021-01623-y.
24. Cohen-Addad V, Kanade V, Mallmann-Trenn F, Mathieu C. Hierarchical clustering: objective functions and algorithms. *J ACM*. 2019;66(4):26. doi:10.1137/1.9781611975031.26.
25. Calinski T, Harabasz J. A dendrite method for cluster analysis. *Commun Stat Simul Comput*. 1974;3(1):791519860. doi:10.1080/03610917408548446.
26. Rousseeuw PJ. Silhouettes: a graphical aid to the interpretation and validation of cluster analysis. *J Comput Appl Math*. 1987;20:53–65. doi:10.1016/0377-0427(87)90125-7.
27. Zhao Q, Fränti P. WB-index: a sum-of-squares based index for cluster validity. *Data Knowl Eng*. 2014;92:77–89. doi:10.1016/j.datak.2014.07.008.

28. Hämmäläinen J, Jauhainen S, Kärkkäinen T. Comparison of internal clustering validation indices for prototype-based clustering. *Algorithms*. 2017;10(3):105. doi:10.3390/al0030105.
29. Malinen MI, Fränti P. All-pairwise squared distances lead to more balanced clustering. *Appl Comput Intell*. 2023;3(1):93–115. doi:10.3934/aci.2023006.
30. Bubeck S, von Luxburg U. Nearest neighbor clustering: a baseline method for consistent clustering with arbitrary objective functions. *J Mach Learn Res*. 2009;10:657–98.
31. Aloise D, Deshpande A, Hansen P, Popat P. NP-hardness of Euclidean sum-of-squares clustering. *Mach Learn*. 2009;75(2):245–8. doi:10.1007/s10994-009-5103-0.
32. Beyer K, Goldstein J, Ramakrishnan R, Shaft U. When is “nearest neighbor” meaningful? In: *International Conference on Database Theory*; 1999 Jan 10–12; Jerusalem, Israel. Berlin/Heidelberg, Germany: Springer; 1999. p. 217–35. doi:10.1007/3-540-49257-7_15.
33. Aggarwal CC, Hinneburg A, Keim DA. On the surprising behavior of distance metrics in high dimensional space. In: *International Conference on Database Theory (ICDT’01)*; 2001 Jan 4–6; London, UK. Berlin/Heidelberg, Germany: Springer-Verlag; 2001. p. 420–34. doi:10.1007/3-540-44503-x_27.
34. Kärkkäinen I, Fränti P. Dynamic local search algorithm for the clustering problem. Joensuu, Finland: University of Joensuu; 2002. Report No.: A-2002-6.
35. Fränti P, Sieranoja S. K-means properties on six clustering benchmark datasets. *Appl Intell*. 2018;48(12):4743–59. doi:10.1007/s10489-018-1238-7.
36. Yang J, Rahardja S, Fränti P. Mean-shift outlier detection and filtering. *Pattern Recognit*. 2021;115(406–21):107874. doi:10.1016/j.patcog.2021.107874.
37. Sieranoja S, Fränti P. Fast and general density peaks clustering. *Pattern Recognit Lett*. 2019;128(3–5):551–8. doi:10.1016/j.patrec.2019.10.019.
38. Fu L, Medico E. FLAME, a novel fuzzy clustering method for the analysis of DNA microarray data. *BMC Bioinform*. 2007;8(1):3. doi:10.1186/1471-2105-8-3.
39. Gionis A, Mannila H, Tsaparas P. Clustering aggregation. *ACM Trans Knowl Discov Data*. 2007;1(1):4. doi:10.1145/1217299.1217303.
40. Fränti P, Virmajoki O. Iterative shrinking method for clustering problems. *Pattern Recognit*. 2006;39(5):761–75. doi:10.1016/j.patcog.2005.09.012.
41. Rezaei M, Fränti P. Set matching measures for external cluster validity. *IEEE Trans Knowl Data Eng*. 2016;28(8):2173–86. doi:10.1109/TKDE.2016.2551240.
42. Zhang T, Ramakrishnan R, Livny M. BIRCH: a new data clustering algorithm and its applications. *Data Min Knowl Discov*. 1997;1(2):141–82. doi:10.1023/A:1009783824328.
43. Fränti P, Sieranoja S. How much can k-means be improved by using better initialization and repeats? *Pattern Recognit*. 2019;93(2):95–112. doi:10.1016/j.patcog.2019.04.014.
44. Arthur D, Vassilvitskii S. K-means++: the advantages of careful seeding. In: *ACM-SIAM Symposium on Discrete Algorithms (SODA’07)*; 2007 Jan 7–9; New Orleans, LA, USA.
45. Rodriguez A, Laio A. Machine learning. Clustering by fast search and find of density peaks. *Science*. 2014;344(6191):1492–6. doi:10.1126/science.1242072.
46. Ester M, Kriegel HP, Sander J, Xu X. A density-based algorithm for discovering clusters in large spatial databases with noise. *KDD*. 1996;96(34):226–31.
47. Schubert E, Sander J, Ester M, Kriegel HP, Xu X. DBSCAN revisited, revisited: why and how you should (still) use DBSCAN. *ACM Trans Database Syst*. 2017;42(3):1–21. doi:10.1145/3068335.
48. von Luxburg U. A tutorial on spectral clustering. *Stat Comput*. 2007;17(4):395–416. doi:10.1007/s11222-007-9033-z.
49. Fränti P, Sieranoja S. Clustering accuracy. *Appl Comput Intell*. 2024;4(1):24–44. doi:10.3934/aci.2024003.
50. Fränti P, Rezaei M, Zhao Q. Centroid index: cluster level similarity measure. *Pattern Recognit*. 2014;47(9):3034–45. doi:10.1016/j.patcog.2014.03.017.