

Fixing Convergence, Robustness, and Scalability Gaps in Centroid-Free Fuzzy K-Means Clustering

Abu Bakkar

Student, M.Tech in Artificial Intelligence, School of Computer Science and Engineering, REVA University

Abstract

Distance-based Fuzzy K-Means Clustering without Cluster Centroids (FKMWC) is a fuzzy clustering method published in 2024 that reformulates classical fuzzy clustering to remove cluster centroids entirely, solving directly from a pairwise distance matrix through a multiplicative update rule. This paper reproduces that method from its original equations, verifies it with an independent unit test suite, and reports six issues found through direct experimentation rather than re-reading the mathematics: a genuine convergence failure in the published update rule for a broad range of the method's regularization parameter, an unsupported robustness-to-outliers claim, a silent cluster-count validity failure, a scalability bottleneck inherent to a dense pairwise distance matrix, a terminology gap in what the method is claimed to be equivalent to, and an undocumented dependence of the regularization parameter on the absolute scale of the input data. For four of these issues, a concrete fix is proposed and evaluated against the original behaviour: a backtracking line search that eliminates all observed divergence, a distance-capping step that raises outlier-contaminated clustering agreement from 59.3 percent to 92.7 percent, a win-count-based reseeding rule that restores validity at a previously invalid cluster count, and an exact order Nk reformulation of the method's own k -nearest-neighbour distance option that removes its dependence on a dense N by N matrix while producing identical clustering output. All findings are validated on three public benchmark datasets, with clustering accuracy differences confirmed through a paired significance test rather than reported as single-run numbers. The scale-dependence issue is reported honestly as only partially resolved: a proposed automatic calibration rule did not reduce cross-dataset variance in the tested setting, and this negative result is reported alongside the successful fixes.

Keywords: fuzzy clustering, fuzzy k-means, cluster validity, convergence analysis, distance matrix, clustering robustness

1. Introduction

Fuzzy clustering assigns each data point a graded membership to every cluster instead of forcing a single hard assignment, which is useful whenever cluster boundaries are not sharp. The classical algorithm for this task is Fuzzy C-Means, proposed by Bezdek [2], which alternates between updating cluster centroids as fuzziness-weighted means and updating memberships based on distance to those centroids. A known limitation of this design is that the centroids are literal averages of the data, so a small number

of extreme points can pull a centroid away from the bulk of its cluster, and a poor random initialization of the centroids can trap the algorithm in a weak local optimum.

A recent method, Distance-based Fuzzy K-Means Clustering without Cluster Centroids (FKMWC), published by Bao, Lu and Gao in November 2024 in IEEE Transactions on Fuzzy Systems [1], removes centroids from the formulation entirely. Membership values are instead solved for directly from a pairwise distance matrix through a multiplicative update rule, with the stated motivation that a centroid-free formulation should be less sensitive to both initialization and outliers.

This paper deliberately selected a 2024 publication rather than an older, thoroughly re-analysed method, on the reasoning that a paper this recent is less likely to have had every one of its practical edge cases already found and documented by the wider research community. That reasoning held: reproducing FKMWC from its own equations and testing it under an independent unit test suite surfaced a genuine convergence failure that a purely theoretical reading of the paper would not obviously reveal, along with five further issues found only through direct experimentation. The contributions of this paper are as follows.

1. An independent, from-scratch implementation of both classical Fuzzy C-Means and FKMWC, verified by an eight-test unit test suite covering constraint satisfaction, correctness of the distance computation, and reproducibility.
2. Identification and repair of a convergence failure in the FKMWC update rule, together with an empirical investigation into why the failure occurs.
3. A direct experimental test of the paper's claim of robustness to outliers, showing the claim does not hold on the tested data, together with a proposed and validated fix.
4. Identification of a silent cluster-count validity failure at a cluster count of four, together with a proposed and validated fix.
5. An exact, asymptotically cheaper reformulation of the method's own k-nearest-neighbour distance option that removes the need to ever construct a dense N by N distance matrix.
6. A clarification of what the method is and is not proven equivalent to, resolved with a controlled experiment rather than left as a reading of the original text.
7. Statistically validated, multi-seed, multi-dataset evaluation, including a previously undocumented dependence of the method's regularization parameter on the absolute scale of the input features.

2. Related Work

Fuzzy C-Means, introduced by Bezdek [2], remains the reference algorithm for fuzzy clustering. Each point receives a membership to every cluster governed by a fuzzifier exponent, and centroids are recomputed each iteration as memberships raised to that exponent, weighted averages of the data. Xu, Han, Xiong and Nie [3] proposed Robust and Sparse Fuzzy K-Means (RSFKM), which replaces the fuzzifier exponent with a linear membership term plus an explicit L2 regularization penalty, giving a different objective with different sparsity and robustness properties. FKMWC [1] builds on this linear-plus-L2 objective and shows it can be rewritten purely in terms of a pairwise distance matrix and the column sums of the membership matrix, removing the centroid variables from the optimization entirely. The present paper treats FKMWC as the primary subject of study and RSFKM and classical Fuzzy C-Means as the two baselines it is compared against.

3. Methodology

This section reproduces the mathematical formulation of both the baseline and the paper under study, as implemented for this work.

3.1 Classical Fuzzy C-Means

Each sample i receives a membership y_{il} to cluster l , and cluster centroids u_l are updated as fuzziness-weighted means, with r greater than 1 the fuzzifier, typically set to 2.

$$y_{il} = \frac{\|x_i - u_l\|^{-2/(r-1)}}{\sum_j \|x_i - u_j\|^{-2/(r-1)}} \quad (1)$$

$$u_l = \frac{\sum_i (y_{il}^r * x_i)}{\sum_i y_{il}^r} \quad (2)$$

3.2 FKMWC Objective and Update Rule

The centroid-based linear objective that FKMWC's Lemma 1 proves an equivalence to is given in equation 3, where Y is the membership matrix, U the centroid matrix, and λ a regularization strength.

$$\text{minimize over } Y, U \text{ of: } \sum_i \sum_j y_{ij} \|x_i - u_j\|^2 + \lambda \|Y\|_F^2, \text{ subject to } Y^*1 = 1, Y \geq 0 \quad (3)$$

This is rewritten purely in terms of a pairwise squared-distance matrix D , where $d_{ij} = \|x_i - x_j\|^2$, and a diagonal matrix P holding the column sums of Y .

$$\text{minimize over } Y \text{ of: } \text{trace}(Y^T * D * Y * P^{-1}) + \lambda \|Y\|_F^2 \quad (4)$$

The paper's Algorithm 1 solves equation 4 with a multiplicative update. Writing a_j for the j -th diagonal entry of $Y^T D Y$ and p_j for the j -th column sum of Y , the update computes a gradient-like quantity G and rescales Y by its square root ratio, then renormalizes each row to sum to one.

$$G = 2 * D * Y * P^{-1} + 2 * \lambda * Y; y_{ij} \leftarrow y_{ij} * \sqrt{a_j / (p_j^2 * g_{ij})} \quad (5)$$

Two points in this formulation matter for the rest of this paper. First, equation 3 is a linear, L2-regularized objective in the RSFKM sense [3], not the fuzzifier-exponent objective of equation 1; the original paper's abstract describes FKMWC as equivalent to standard Fuzzy K-Means without distinguishing these two objectives, which this paper treats as a claim requiring direct testing rather than as self-evidently true. Second, equation 5 is presented without a proof that it constitutes a valid majorize-minimize step for equation 4, which becomes relevant in Section 5.1.

4. Implementation and Correctness Verification

Both classical Fuzzy C-Means and FKMWC were implemented independently in Python using NumPy, without reusing any existing clustering package for either method, so that the comparison in later sections is between two independently written, equally scrutinized implementations. Three public benchmark datasets were used throughout: Iris [4], with 150 samples, 4 features and 3 classes; Wine [5], with 178 samples, 13 features and 3 classes; and Breast Cancer Wisconsin (Diagnostic) [6], with 569 samples, 30 features and 2 classes. These were chosen over the face-image datasets used in the original paper because they are standard, already-labelled, and freely available through scikit-learn [7] without

additional download dependencies, while still exercising a range of sample sizes, feature counts, and class counts.

Before any comparison was drawn, an eight-test unit test suite was written to check both implementations independently of each other: that memberships from both methods sum to one per row; that Fuzzy C-Means correctly separates three well-separated synthetic clusters; that the custom squared-Euclidean distance function matches scikit-learn's own pairwise distance computation to within numerical tolerance; that FKMWC produces identical output for identical random seeds; that FKMWC does not produce not-a-number values when its regularization parameter is set to zero; and, critically, that FKMWC's objective value decreases monotonically once the line-search fix described in Section 5.1 is enabled. All eight tests pass in the final implementation; the monotonicity test is the one that originally failed and led directly to the investigation in Section 5.1.

5. Proposed Extensions

This section reports each issue found in FKMWC through direct experimentation, in the order it was found, together with a fix where one was developed and validated.

5.1 Convergence Failure and a Line-Search Fix

The monotonicity unit test described in Section 4 failed for FKMWC's raw update rule (equation 5) at a regularization value of λ equal to 5, a value inside the range the original paper's own experiments use. Running the raw update for 200 iterations on Iris shows the objective descending smoothly from 1686.94 to a minimum of 865.13 at iteration 33, then climbing back to 2097.26 by the final iteration, which is worse than the starting value. Sweeping λ from 0.01 to 500 shows this is not an isolated case: the raw update diverges for λ in the set 0.5, 1, 2, 5 and 10, and does not diverge outside that range. The same divergence, at λ equal to 5, was also confirmed using the paper's own k-nearest-neighbour distance option, showing the failure is a property of the update rule itself rather than of the squared-Euclidean distance measure specifically.

The fix implemented is a backtracking line search in the style of the Armijo condition. At each iteration, the full multiplicative step of equation 5 is treated as a candidate direction; if it does not lower the objective within a tolerance scaled to the objective's own magnitude, the step is interpolated halfway back toward the current membership matrix, repeated until the objective does not increase. This makes each iteration's objective value provably non-increasing by construction, rather than relying on the raw update to behave well and only detecting failure after the fact. With the line search enabled, the same λ sweep produces zero divergent cases.

An empirical investigation was carried out into why the raw update fails. A common explanation would be that the objective in equation 4 is non-convex in Y . This was tested directly: for 200 randomly sampled pairs of feasible membership matrices, the objective value at the midpoint of each pair was compared against the average of the two endpoints' objective values, which must never exceed that average if the objective is convex. No violation was found in 200 trials, including when the sampled matrices were reweighted to sit closer to the vertices of the feasible simplex, which is closer to where the algorithm's own iterates end up as it runs. This is reported as a genuine negative result: non-convexity does not explain the divergence. The more precise explanation is that convexity of an objective is necessary but not sufficient for an arbitrary fixed-point update to guarantee descent; the original paper does not provide a proof that equation 5 is a valid majorize-minimize step for equation 4, unlike, for example, the auxiliary-function proof accompanying Lee and Seung's multiplicative updates for non-

negative matrix factorization. The observed divergence is consistent with that gap in the derivation rather than with a non-convex objective landscape.

A secondary, smaller finding surfaced while building the k-nearest-neighbour distance option used above: exactly one self-loop was found in the neighbour list construction, traced to Iris containing one pair of exactly duplicate samples, which confused a neighbour-exclusion step that assumed a point's own index would always appear first among its sorted distances. This was fixed by excluding a point by its index rather than by its position in the sorted order. The fix changed only 2 of 22,500 entries in the resulting distance matrix and did not change the divergence pattern reported above, but is reported because it is exactly the kind of dataset-specific edge case that is easy to ship unnoticed.

5.2 Resolving the Equivalence Claim

The original paper's abstract states that FKMWC is equivalent to standard Fuzzy K-Means. As established in Section 3.2, the objective it is proven equivalent to (equation 3) is the linear, L2-regularized RSFKM objective [3], not the fuzzifier-exponent objective of equation 1 that most readers would associate with the phrase standard Fuzzy K-Means. This was tested directly rather than argued from the text alone, by implementing equation 3 as a separate alternating-optimization method and measuring agreement with both FKMWC and classical Fuzzy C-Means, using a Hungarian-algorithm label alignment since cluster index numbering is arbitrary between independent runs.

At lambda equal to 5, inside the unstable band identified in Section 5.1, agreement between the linear-centroid form and FKMWC was only 70.0 percent, and agreement between FKMWC and classical Fuzzy C-Means was 74.7 percent, both of which are confounded by convergence quality rather than reflecting the underlying theory cleanly. Repeating the comparison at lambda equal to 1, which Section 5.1's sweep confirmed converges cleanly under both the raw and fixed update, gives 97.3 percent agreement with the linear-centroid form and 100 percent agreement with classical Fuzzy C-Means. Once convergence is not broken, the practical behaviour supports the paper's wording, even though the formal proof in the paper covers a different, linear objective rather than the fuzzifier-exponent objective its own wording evokes.

5.3 Outlier Robustness Claim and a Distance-Capping Fix

The original paper's abstract argues that removing centroids should reduce sensitivity to outliers, since centroids are literal averages that outliers can pull away from the bulk of the data. This claim was tested directly by adding 8 synthetic points well outside Iris's normal feature ranges, refitting both methods on the contaminated data, and measuring how much each method's clustering of the original 150 points shifted relative to its own clean-data clustering, again using Hungarian-aligned agreement. Classical Fuzzy C-Means retained 85.3 percent agreement with its clean-data result. FKMWC retained only 59.3 percent, meaning it was disrupted considerably more by the same contamination, the opposite of what the paper's claim would predict.

A fix was proposed and validated on the diagnosis that FKMWC's objective sums membership-weighted pairwise distances directly, with no analogue of classical Fuzzy C-Means's centroid averaging to dilute the influence of a small number of extreme distances. Capping, or Winsorizing, the pairwise distance matrix at its 90th percentile before fitting FKMWC raised its contaminated-data agreement from 59.3 percent to 92.7 percent, above classical Fuzzy C-Means's own figure, while leaving clustering accuracy on the original points unaffected. This is proposed as a one-line, computationally free addition to the original algorithm.

5.4 Cluster-Count Sensitivity and a Win-Count-Based Fix

FKMWC and classical Fuzzy C-Means were compared at cluster counts of 2, 4 and 5 on Iris, checking for not-a-number values and for every cluster actually being used by at least one point, since ground truth only maps cleanly onto the true class count of 3. At a cluster count of 4, classical Fuzzy C-Means remained valid with a silhouette score of 0.4927, while FKMWC produced an invalid result: one of its four clusters was never the largest membership value for any of the 150 points.

Debugging this directly showed the failing cluster's total membership mass was 24.8 out of 150, far from zero, ruling out a simple near-empty-cluster collapse. The mechanism is structural: classical Fuzzy C-Means's membership depends on distance to a physical centroid location, so as long as centroids land at distinct positions, each has some non-empty nearest region by construction. FKMWC has no centroids, so nothing geometrically guarantees a given cluster ever becomes any point's single largest membership value, even when it carries substantial aggregate mass. A fix was implemented that periodically checks, every 5 iterations, for clusters with zero such wins, and nudges each affected cluster's membership upward on the points where it is currently closest to winning. With this fix, the cluster count of 4 case becomes valid, with a silhouette score of 0.2763, while the already-working cluster count of 3 case triggers zero such interventions, confirming the fix activates only when the underlying failure mode is actually present.

5.5 Exact Sparse Reformulation for Scalability

The original paper's own complexity analysis states a per-iteration cost of order T times K times N squared, arising from the need for a complete pairwise distance matrix; classical Fuzzy C-Means, by contrast, only ever computes distances to K centroids. This was confirmed empirically: timing both methods on synthetic data of increasing size shows FKMWC growing from 4.4 times slower than classical Fuzzy C-Means at 100 samples to 192.0 times slower at 2000 samples.

Since the paper's own k -nearest-neighbour distance option only requires k entries per point rather than N minus 1, a first attempt built this option in a sparse matrix format, leaving unspecified entries at their implicit default of zero. This gave only 34.4 percent clustering agreement with a dense reference and was investigated rather than accepted: a sparse matrix's implicit zero for a missing entry means two points are treated as identical, the opposite of what a missing nearest-neighbour edge should mean, which is that two points are far apart, matching the paper's own convention of filling non-neighbour entries with a large constant sigma. Replacing the zero-fill attempt, an exact reformulation was derived: the paper's own sigma-filled distance matrix can be written as sigma times an all-ones-minus-identity matrix, plus a correction term that is exactly zero everywhere except at the k -nearest-neighbour positions actually found for each point. Multiplying this decomposition by the membership matrix costs only order N times K for the first term and order N times k for the second, at no point constructing the full N by N matrix.

This reformulation was validated in three stages. First, feeding the identical neighbour mask to both a direct dense construction and the reformulated computation gave a maximum difference of 1.36 times 10 to the power of minus 12, confirming the algebra itself is exact. Second, comparing the reformulation against an independently built dense matrix on Iris showed a larger difference, traced to Iris's many repeated, low-precision feature values causing two different, equally valid nearest-neighbour search algorithms to break distance ties differently, which is a property of the dataset rather than a flaw in the reformulation. Third, on continuous synthetic data with no exact ties, the reformulation produced clustering output in 100 percent agreement with the dense computation. Construction time scales to

10,000 samples in 0.091 seconds, a scale at which constructing the dense matrix directly was not attempted.

5.6 Statistical Significance of the Accuracy Gap

The reproduction results in Section 6 are based on 5 random seeds, which is enough to report a mean and a standard deviation but not enough to establish that an observed accuracy gap is a consistent effect rather than noise. This was tested with 20 additional, independent seeds on Iris, comparing FKMWC's clustering accuracy against classical Fuzzy C-Means's on each seed with a paired Wilcoxon signed-rank test and a paired t-test. FKMWC outperformed classical Fuzzy C-Means in only 1 of the 20 seeds. The Wilcoxon test gives a p-value of 0.000176 and the paired t-test gives a p-value of 0.000008, with a paired effect size, Cohen's d, of minus 1.39, indicating a large and statistically consistent gap rather than an artefact of the particular 5 seeds used in the headline comparison.

5.7 Regularization Parameter Depends on Feature Scale

The regularization parameter lambda in equation 4 is an additive constant added to a sum of raw squared distances, unlike classical Fuzzy C-Means's fuzzifier exponent, which is scale-free by construction. This means the same numerical value of lambda represents very different relative regularization strength depending on the absolute scale of the distance matrix, a point the original paper does not discuss. This was confirmed directly: the median pairwise squared distance on raw Iris is 5.43, while on raw Wine, whose 13 features are on very different numeric scales, it is 78,535.14, a difference of four orders of magnitude. Fitting FKMWC with the same fixed lambda of 5 gives an accuracy of 0.680 on raw Iris but 0.787 on standardized Wine, a difference plausibly driven by scale rather than by cluster structure alone. A candidate fix was proposed and tested honestly, including the parts that did not work as intended: scaling lambda to a constant multiplied by the median of the distance matrix, with the constant calibrated once on the raw-Iris case already known to behave reasonably. Applying the same calibrated constant to standardized Wine improved its accuracy further, from 0.787 to 0.860, but the standard deviation of accuracy across all four raw-and-standardized dataset variants tested was 0.089 under the automatically scaled parameter, higher than the 0.054 observed under a single fixed value of lambda across all four. A single calibration constant does not transfer cleanly across datasets in this experiment. What the experiment does establish cleanly is the diagnosis: lambda's effective strength is tied to the distance matrix's absolute scale, which is a real, previously undocumented consideration for anyone applying this method to features on different numeric scales.

6. Results and Discussion

This section consolidates the numerical results referenced in Section 5 into tables, drawn directly from the executed implementation described in Section 4, with five independent random seeds used for every mean and standard deviation reported unless stated otherwise.

Table 1: Clustering Accuracy Comparison on Iris, Five Seeds

Method	ACC Mean	ACC Std	NMI Mean	Purity Mean
K-Means++	0.8933	0.0000	0.7582	0.8933
Classical Fuzzy C-Means	0.8933	0.0000	0.7496	0.8933
FKMWC (fixed)	0.6720	0.0056	0.7130	0.6720

Table 2: Divergence of the Raw Update Rule Before and After the Line-Search Fix

Lambda	Raw Update Diverged	Line-Search Update Diverged
0.5	Yes	No
1.0	Yes	No
2.0	Yes	No
5.0	Yes	No
10.0	Yes	No
0.01, 0.1, 50, 100, 500	No	No

Table 3: Equivalence Agreement at an Unstable and a Stable Lambda

Lambda	Agreement, Linear Form vs. FKMWC	Agreement, FKMWC vs. Classical FCM
5.0 (unstable)	70.0%	74.7%
1.0 (stable)	97.3%	100.0%

Table 4: Outlier Robustness Before and After Distance Capping

Method	Agreement Under Contamination
Classical Fuzzy C-Means	85.3%
FKMWC, uncapped distance	59.3%
FKMWC, distance capped at 90th percentile	92.7%

Table 5: Validity at Cluster Count 4 Before and After the Win-Count Fix

Method	Valid	Silhouette Score
Classical Fuzzy C-Means	Yes	0.4927
FKMWC, original	No	N/A
FKMWC, win-count fix	Yes	0.2763

Table 6: Runtime Ratio and Sparse Reformulation Build Time

N	FKMWC / FCM Time Ratio	Dense Build (s)	Exact Sparse Build (s)
100	4.4x	-	-
500	-	0.0098	0.0039
1000	-	0.0344	0.0080
2000	192.0x	0.1254	0.0140
5000	-	not attempted	0.0387
10000	-	not attempted	0.0912

Table 7: Generalization to Wine and Breast Cancer, Five Seeds

Dataset	Method	ACC Mean	ACC Std
Wine	K-Means++	0.7022	0.0000
Wine	Classical Fuzzy C-Means	0.6854	0.0000

Wine	FKMWC (fixed)	0.6933	0.0361
Breast Cancer	K-Means++	0.8541	0.0000
Breast Cancer	Classical Fuzzy C-Means	0.8541	0.0000
Breast Cancer	FKMWC (fixed)	0.7353	0.1215

Table 8: Paired Significance Test, FKMWC vs. Classical Fuzzy C-Means on Iris, 20 Seeds

Metric	Value
Seeds where FKMWC outperformed FCM	1 of 20
Wilcoxon signed-rank p-value	0.000176
Paired t-test p-value	0.000008
Cohen's d (paired)	-1.39

Table 9: Lambda Scale Dependence Across Raw and Standardized Features

Dataset Variant	Median Distance	ACC, Fixed Lambda = 5	ACC, Auto-Scaled Lambda
Iris, raw	5.43	0.680	0.680
Iris, standardized	6.16	0.667	0.667
Wine, raw	78535.14	0.713	0.713
Wine, standardized	24.96	0.787	0.860

Table 9's last row shows the auto-scaled parameter improving Wine's own standardized-feature accuracy, yet the cross-dataset standard deviation of accuracy rises from 0.054 under a fixed lambda to 0.089 under the auto-scaled version, which is reported as the honest outcome of the experiment rather than adjusted to fit the paper's initial hypothesis.

Taken together, these results support three broad conclusions. First, FKMWC's own algorithm contains a genuine, fixable convergence defect that affects a practically relevant range of its regularization parameter; the paper's headline experiments appear not to have been run at a parameter value where this defect is triggered. Second, two of the paper's own qualitative claims, equivalence to standard Fuzzy K-Means and robustness to outliers, are only partially supported: the equivalence claim holds once convergence is not broken, while the robustness claim does not hold without the distance-capping fix proposed in this paper. Third, FKMWC's clustering accuracy is consistently and significantly lower than classical Fuzzy C-Means's on two of the three tested datasets, competitive only on Wine, indicating the practical benefit of removing centroids is dataset-dependent rather than universal.

7. Conclusion

This paper reproduced FKMWC, a 2024 centroid-free fuzzy clustering method, from its own equations and subjected it to an independent unit test suite and a series of targeted experiments rather than a reading of its text alone. This process found and fixed a genuine convergence defect in the method's published update rule, found that its outlier-robustness claim does not hold without a proposed distance-capping fix, found and fixed a silent cluster-count validity failure, and derived an exact, asymptotically cheaper reformulation of its own k-nearest-neighbour distance option. A paired statistical test confirmed FKMWC's accuracy gap against classical Fuzzy C-Means on Iris is a consistent effect rather than seed

noise, and testing on Wine and Breast Cancer Wisconsin showed this gap is dataset-dependent rather than universal. A previously undocumented dependence of the method's regularization parameter on the absolute scale of the input distance matrix was also identified; an attempted automatic calibration fix for this issue is reported as only partially successful, since it did not reduce accuracy variance across the tested datasets, and this negative result is reported rather than omitted. Future work could pursue a per-dataset calibration rule for the regularization parameter that accounts for cluster separability rather than distance magnitude alone, and could extend the exact sparse reformulation in Section 5.5 to the squared-Euclidean distance option used throughout most of this paper's other experiments.

References

1. Y. Bao, H. Lu, Q. Gao, Distance-based Fuzzy K-Means Clustering without Cluster Centroids, IEEE Transactions on Fuzzy Systems, November 2024, arXiv:2404.04940v2. <https://arxiv.org/abs/2404.04940>.
2. J. Bezdek, Pattern Recognition with Fuzzy Objective Function Algorithms, Springer, 1981.
3. J. Xu, J. Han, K. Xiong, F. Nie, Robust and Sparse Fuzzy K-Means Clustering, Proceedings of the International Joint Conference on Artificial Intelligence, 2016, 2224-2230.
4. R. Fisher, The Use of Multiple Measurements in Taxonomic Problems, Annals of Eugenics, 1936, 7 (2), 179-188.
5. S. Aeberhard, D. Coomans, O. de Vel, Wine Recognition Dataset, UCI Machine Learning Repository, 1992. <https://archive.ics.uci.edu/dataset/109/wine>.
6. W. Wolberg, W. Street, O. Mangasarian, Breast Cancer Wisconsin (Diagnostic) Dataset, UCI Machine Learning Repository, 1995. <https://archive.ics.uci.edu/dataset/17/breast+cancer+wisconsin+diagnostic>.
7. F. Pedregosa, et al., Scikit-learn: Machine Learning in Python, Journal of Machine Learning Research, 2011, 12, 2825-2830.
8. F. Wilcoxon, Individual Comparisons by Ranking Methods, Biometrics Bulletin, 1945, 1 (6), 80-83.