
Comparison of K-Means and Hierarchical Clustering Using Silhouette Score on BPBD Data, Papua Barat

Najlah Putri Qudsiya^{1*}, Christian Dwi Suhendra², Josua Josen A. Limbong³

^{1,2,3} Universitas Papua, Fakultas Teknik, Teknik Informatika, Jl. Gunung Salju, Amban, Manokwari, West Papua 98314, Indonesia

Keywords

Data Mining; Disaster Clustering; Hierarchical Clustering; K-Means Clustering; Silhouette Score

***Corresponding Author:**

najlahputriqudsiya@gmail.com

Abstract

This study compares K-Means Clustering and Hierarchical Clustering to group disaster-prone areas in Papua Barat Province using official data from UPTD Pusdalops BPBD for the period 2020 to November 11, 2025. The dataset covers seven regencies with three aggregated variables: total floods and tidal floods, other disasters, and earthquakes. The research procedure includes data preprocessing, z-score normalization, determining the optimal number of clusters using the Elbow Method and Silhouette Score, and visualizing results with scatter plots and dendrograms. Both algorithms produced two clusters, with Hierarchical Clustering achieving a higher Silhouette Score of 0.4274 compared to 0.339 for K-Means. The scores indicate moderate clustering quality, reflecting the exploratory nature of the study due to the small dataset. Manokwari Regency formed a distinct cluster due to differing disaster characteristics. The novelty lies in employing official BPBD data and Silhouette Score evaluation to compare the two clustering algorithms, supporting data-driven prioritization for disaster mitigation.

1. Introduction

Disaster risk reduction requires the use of data-driven information so that preparedness planning, mitigation, and decision-making can be carried out more accurately, particularly in areas with interrelated multi-hazard events [1]. The availability of disaster data also highlights the need for analytical methods capable of summarizing events in a concise, structured, and relevant manner for both operational purposes and local government policy-making [2]. In the context of Papua Barat Province, the main challenge is not only the availability of disaster occurrence data but also the transformation of such data into interpretable information to distinguish areas with similar event patterns from those with different characteristics [1], [2]. The scope of this study is limited to seven regencies because the current administrative area of Papua Barat Province consists of only seven regencies following the establishment of Southwest Papua Province under Law Number 29 of 2022. This regional division caused several areas that were previously part of Papua Barat Province, namely Sorong Regency, South Sorong Regency, Raja Ampat Regency, Tambrau Regency, Maybrat Regency, and Sorong City, to become part of Southwest Papua Province. Therefore, the study area was adjusted to the administrative scope of Papua Barat Province after the regional division, namely Manokwari Regency,

Manokwari Selatan Regency, Pegunungan Arfak Regency, Teluk Wondama Regency, Fakfak Regency, Kaimana Regency, and Teluk Bintuni Regency [3].

One approach that can be used to summarize these patterns is clustering, an unsupervised learning technique that groups objects based on their similarity without requiring predefined class labels [4]. In disaster studies, clustering has been used to identify event patterns and group regions based on observed disaster characteristics [5]. In this study, clustering was used to explore patterns from aggregate disaster occurrence data across seven administrative regions in Papua Barat Province, with the main outputs consisting of regional grouping results and scatter plot visualizations to support the interpretation of proximity and separation among regions [4], [5].

This study compares K-Means Clustering and Hierarchical Clustering because they represent two different yet complementary clustering approaches [4]. K-Means Clustering is widely used for numerical data because it is efficient, simple, and tends to produce compact clusters [6]. In contrast, Hierarchical Clustering is useful for gradually illustrating the proximity structure among objects through a dendrogram, allowing relationships among regions to be observed in a more exploratory manner [7]. Comparing these two methods is relevant because clustering results can be influenced by algorithm characteristics, data structure, similarity measures, and the method used to determine the number of clusters [4], [7], [8].

Previous studies have shown that clustering has been applied in disaster-related contexts, including disaster statistical analysis, regional disaster resilience assessment, and the grouping of disaster-prone areas using Hierarchical Clustering, although its applications vary depending on the data characteristics and the objectives of each analysis [5], [9], [10]. Tin et al. [5] examined the application of several clustering algorithms to FEMA disaster statistical data, while Choi and Song [9] used a clustering approach to evaluate community disaster resilience based on GIS and census data. Studies on K-Means Clustering generally emphasize algorithm efficiency, initialization variations, and the characteristics of the resulting clusters, whereas Hierarchical Clustering tends to focus more on linkage strategies, dendrogram structures, and similarity measures [6], [7]. In addition, clustering results are influenced by feature normalization, the determination of the number of clusters, and, in Hierarchical Clustering, the dendrogram cutting decision; therefore, a clear and consistent internal evaluation procedure is required to ensure that the grouping results are not determined arbitrarily [4], [11], [12]. Based on these conditions, a research gap remains regarding the need for an empirical comparison between K-Means Clustering and Hierarchical Clustering using official regional disaster occurrence data, supported by a measurable and consistent procedure for determining the number of clusters [5], [6], [12].

This study used aggregated data on total disaster events from 2020 to November 11, 2025, obtained from UPTD Pudaslops BPBD of Papua Barat Province. The analyzed variables consisted of the total number of floods and tidal floods, total other disasters, and total earthquakes. The initial data included several types of disaster events, namely floods, tidal floods, fires, abrasion, landslides, tornadoes, tidal waves, extreme weather, forest and land fires, and earthquakes. The research variables were then simplified into three numerical attributes, namely total floods and tidal floods, total other disasters, and total earthquakes. This simplification was carried out because the initial data contained many types of disasters, while some events had relatively low frequencies and were unevenly distributed across regions. If all disaster types were used as separate variables, the data structure would become too complex and could potentially affect the stability of the clustering results, especially because the number of regions analyzed was limited. Therefore, aggregated variables were used to maintain analytical stability without eliminating the general patterns of disaster vulnerability among regions [6]. The proposed approach compares K-Means and hierarchical clustering on the dataset by determining the number of clusters based on the internal evaluation of cluster structure quality. In K-Means, the Silhouette Score was used as the main basis for determining the number of clusters, while the elbow method was used as a supporting analysis [13]. In hierarchical clustering, the number of clusters was determined by evaluating the Silhouette Score for several possible dendrogram cut results, so that the cluster number decision was based on the quality of cluster cohesion and separation [12], [13]. Data processing was performed using Orange Data

Mining and Jupyter Notebook because both tools support clustering workflows, internal evaluation, and visualization of clustering results, while feature normalization was applied because distance-based methods are sensitive to differences in variable scales [11], [14], [15].

This study aims to group disaster-prone areas in Papua Barat Province by comparing K-Means Clustering and hierarchical clustering algorithms based on aggregated disaster event data [4], [7]. The novelty of this study lies in the empirical comparison of the two algorithms to determine the clustering method that is more suitable for grouping disaster-prone areas in Papua Barat Province, supported by internal evaluation using the Silhouette Score and scatter plot visualization as the basis for interpreting patterns of disaster vulnerability among regions [5], [14], [16].

2. Research Method

This study employs a quantitative-exploratory approach based on unsupervised learning data mining to compare the K-Means Clustering and Hierarchical Clustering algorithms on aggregated disaster event data in Papua Barat Province. The quantitative-exploratory approach is used because this study utilizes numerical data to explore regional grouping patterns based on similarities in disaster events, without aiming to test causal relationships among variables [4], [17], [18], [19].

Data processing was conducted using Orange Data Mining version 3.39.0 for the implementation of K-Means Clustering because it provides a visual workflow that facilitates the exploration and evaluation of clustering results [14], [15], [20]. Meanwhile, hierarchical clustering was implemented using Jupyter Notebook version 7.0.8 on the Windows 11 operating system to provide flexibility in parameter settings and dendrogram visualization [14], [15].

Based on the research workflow shown in Figure 1, the study began with the collection of disaster event data from UPTD Pusdalops BPBD of Papua Barat Province, followed by data preprocessing through dataset checking, variable selection, and z-score normalization. Next, the number of clusters was determined using the Elbow Method and Silhouette Score before implementing the K-Means Clustering and Hierarchical Clustering algorithms. The clustering results were then visualized using scatter plots and dendrograms and evaluated using the Silhouette Score to compare the quality of the resulting clusters.

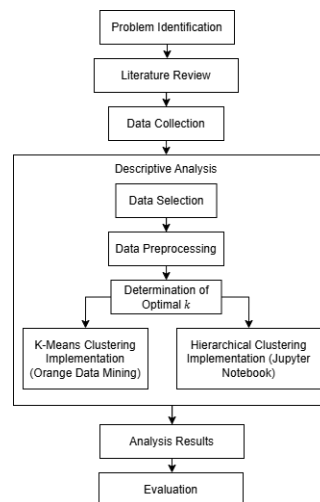


Figure 1. Research Method

2.1 Dataset

The research data were obtained from the Technical Implementation Unit (UPT) of the Operations Control Center (Pusdalops) at the Regional Disaster Management Agency (BPBD) of Papua Barat Province. The dataset consists of total disaster events from 2020 to November 11, 2025, across seven regions, namely Manokwari

Regency, Manokwari Selatan Regency, Pegunungan Arfak Regency, Teluk Wondama Regency, Fakfak Regency, Kaimana Regency, and Teluk Bintuni Regency. The number of study areas was limited to seven regencies because the administrative scope of Papua Barat Province after the establishment of Southwest Papua Province currently consists of only seven regencies [3].

Table 1. Dataset

Papua Barat Province	Total Flood and Coastal Flood Events	Total Other Disasters	Total Earthquake Events
Manokwari	28	65	311
Pegunungan Arfak	6	16	340
Manokwari Selatan	15	7	438
Teluk Wondama	7	4	312
Kaimana	0	3	569
Fakfak	4	24	70
Teluk Bintuni	17	4	479

The research variables consist of three numerical attributes formed from aggregated disaster event data, namely total floods and tidal floods (X1), total other disasters (X2), and total earthquakes (X3), as shown in Table 2. This grouping was carried out to simplify the initial data because some types of disasters had relatively low event frequencies and were unevenly distributed across regions, which could affect the stability of clustering results in a small dataset [6]. This study focuses on total disaster events as the basis for regional grouping; therefore, it does not yet include spatial factors, socio-economic conditions, population vulnerability levels, or the level of disaster damage, as the study is focused on an initial exploration of similarity patterns in disaster events across regions using aggregated data in Papua Barat Province [1], [2].

Table 2. Research Variables

No	Variable	Symbol	Description
1	Total Flood and Coastal Flood Events	X ₁	Total number of flood and coastal flood events in each region of Papua Barat Province
2	Total Other Disaster Events	X ₂	Total number of other disaster events, including fire, abrasion, landslide, tornado, tidal wave, extreme weather, and forest and land fires, in each region of Papua Barat Province
3	Total Earthquake Events	X ₃	Total number of earthquake events in each region of Papua Barat Province

2.2 Preprocessing

Preprocessing was carried out through data checking, dataset tabulation, variable selection, and feature normalization. At this stage, the initial data consisting of several types of disaster events were organized into three aggregated variables, namely total floods and tidal floods, total other disasters, and total earthquakes. The formation of aggregated variables was conducted because some types of disasters had relatively low event frequencies and were unevenly distributed across regions; therefore, using all disaster types as separate variables could affect the stability of clustering results when the amount of data is limited [6]. Data completeness checks, regional duplication checks, and exploratory examinations of possible outliers were also performed before the clustering process. Data with different characteristics were retained if they still represented the actual conditions of the administrative areas analyzed.

In this study, data normalization for K-Means Clustering and Hierarchical Clustering was carried out using the z-score method so that each variable had a comparable scale and the distance calculation was not dominated by variables with larger value ranges [4].

The z-score standardization formula is shown in (1).

$$z = \frac{x - \mu}{\sigma} \quad (1)$$

where z is the standardized value, x is the original value, μ is the mean, and σ is the standard deviation.

The distance between data points was calculated using Euclidean Distance, as shown in (2) [4].

$$d(x, y) = \sqrt{\sum_{m=1}^p (x_m - y_m)^2} \quad (2)$$

where $d(x, y)$ denotes the distance between objects x and y , while p denotes the number of variables used in the distance calculation. This formula is used to measure the level of proximity among data points based on differences in each variable [4], [6].

2.3 Determination of the Optimal Number of Clusters

The number of clusters was determined using the Elbow Method and the Silhouette Score. The Elbow Method was used to observe the decrease in the Within-Cluster Sum of Squares (WCSS) value in K-Means Clustering [6].

The WCSS formula is shown in (3).

$$WCSS(k) = \sum_{k=1}^k \sum_{x_i \in C_k} \|x_i - \mu_k\|^2 \quad (3)$$

The Silhouette Score was used to assess cluster quality based on cohesion and separation in data without class labels [13], [18], [20], [21].

The Silhouette Score formula is shown in (4).

$$s(i) = \frac{b(i) - a(i)}{\max\{a(i), b(i)\}} \quad (4)$$

where $a(i)$ denotes the average distance between object i and other members within the same cluster, while $b(i)$ denotes the minimum average distance between object i and members of the nearest neighboring cluster [21].

2.4 K-Means Clustering Algorithm

The implementation of K-Means Clustering was carried out using the default configuration in Orange Data Mining, including automatic centroid initialization, iterative centroid updating, and process termination when the cluster composition reached convergence. The default configuration was used because this study focuses on comparing clustering results rather than optimizing algorithm parameters [17], [22], [23]. The number of clusters in K-Means was determined using the Silhouette Score as the main basis, while the Elbow Method was used as a supporting analysis [6], [20], [21], [24].

K-Means groups data by minimizing the sum of squared distances between data points and cluster centroids [6], [17], [22], [23].

The objective function of K-Means Clustering is shown in (5).

$$\arg \min_c \sum_{j=1}^k \sum_{x_i \in C_j} \|x_i - \mu_j\|^2 \quad (5)$$

where C denotes the set of clusters, x_i denotes the i -th data point, C_j denotes the j -th cluster, and μ_j denotes the centroid of cluster C_j . This objective function is used to minimize the total squared distance between data points and the centroid of each cluster [6].

2.5 Hierarchical Clustering Algorithm

Hierarchical Clustering was implemented using Jupyter Notebook with an agglomerative hierarchical clustering approach. This study used complete linkage, in which the distance between clusters is determined based on the farthest pair of members [7], [12], [25]. Complete linkage was selected because it produces clearer separation between clusters and tends to form more compact clusters than other linkage methods on small datasets, making the resulting dendrogram easier to interpret for grouping disaster-prone areas [7], [25].

The complete linkage formula is shown in (6).

$$D(C_i, C_j) = \max_{x \in C_i, y \in C_j} d(x, y) \quad (6)$$

where $D(C_i, C_j)$ denotes the distance between clusters C_i and C_j , while $d(x, y)$ denotes the distance between objects x and y . The number of clusters in Hierarchical Clustering was determined by evaluating the Silhouette Score for several possible dendrogram cuts [13], [21].

2.6 Metric Evaluation

The research procedure was systematically designed so that it can be replicated on similar datasets. The clustering results were evaluated using the Silhouette Score as a comparative metric between K-Means Clustering and Hierarchical Clustering. The Silhouette Score was selected because it can be used for data without labels, can be applied to more than one clustering algorithm, and can simultaneously assess cohesion and separation [4], [13], [18], [21]. The Davies-Bouldin Index and Calinski-Harabasz Index were not used in this study because the evaluation focused on the Silhouette Score, which was considered more suitable for comparing cluster cohesion and separation quality in unlabeled data and small datasets [13], [26], [27]. The evaluation results were then explained together with the scatter plots to examine the patterns of proximity, separation, and cluster composition among regions.

3. Result and Discussion

3.1 Data Preprocessing Results

The normalization results for K-Means Clustering are shown in Table 3. Normalization using the z-score method produced a more balanced data scale across variables, thereby reducing the dominance of variables with larger value ranges in the process of calculating distances between data points [4].

Table 3. Normalization Results for K-Means Clustering

Papua Barat Province	Total Flood and Coastal Flood Events	Total Other Disaster Events	Total Earthquake Events
Manokwari	1.914	2.296	-0.3303
Teluk Bintuni	0.676	-0.657	0.8055
Manokwari Selatan	0.45	-0.512	0.5283
Teluk Wondama	-0.45	-0.657	-0.3236
Pegunungan Arfak	-0.563	-0.076	-0.1343
Fakfak	-0.788	0.311	-1.9598
Kaimana	-1.239	-0.706	1.414

Figure 2 shows the comparison of data distribution before and after normalization in the Hierarchical Clustering process. Normalization using the z-score method produces a more balanced data scale among variables, thereby reducing the dominance of variables with larger value ranges and supporting a more optimal distance-based clustering process.

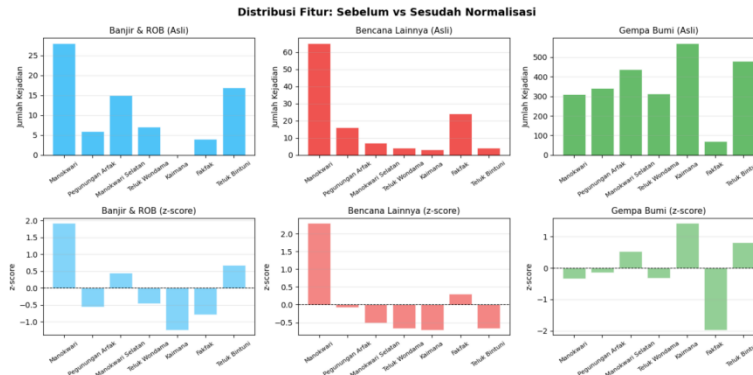


Figure 2. Data Normalization Visualization in Hierarchical Clustering

3.2 Determination of the Optimal Number of Clusters

The number of clusters using the Elbow Method was determined by evaluating the Within-Cluster Sum of Squares (WCSS) value for each number of clusters (k). The WCSS value was used to observe the level of variation within clusters, where a smaller WCSS value indicates that the objects within a cluster are more homogeneous. The WCSS calculation results for each value of k are shown in Table 4.

Table 4. WCSS Value

Number of Clusters (k)	WCSS
1	2.095783203
2	0.971818187
3	0.45314111
4	0.21892919
5	0.028039699
6	0.007097164
7	0

Figure 3 shows a sharp decrease in the WCSS value from $k=1$ to $k=2$, followed by a tendency to level off for the subsequent numbers of clusters. This pattern indicates that two clusters are a reasonable number of clusters to consider in K-Means Clustering because, after this point, the decrease in WCSS is no longer significant. In clustering analysis, the elbow pattern in the Elbow Method graph is commonly used as an initial indication for determining the optimal number of clusters [6], [20].

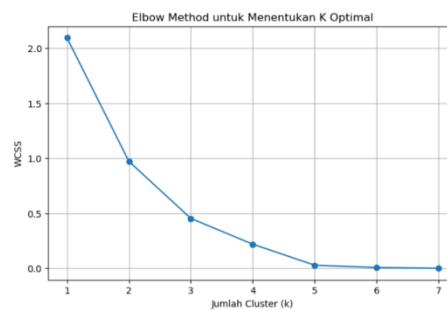


Figure 3. Elbow Method Graph

The comparison of Silhouette Score values is shown in Table 5. Based on the evaluation results, the highest Silhouette Score was obtained at $k = 2$, with a value of 0.4274 for Hierarchical Clustering and 0.339 for K-Means Clustering. These results indicate that $k = 2$ produced better cluster cohesion and separation than the other numbers of clusters [13], [18], [20].

Table 5. Comparison of Silhouette Score Values for K-Means Clustering and Hierarchical Clustering

Number of Clusters k	Hierarchical Clustering (Jupyter Notebook)	K-Means Clustering (Orange Data Mining)
2	0.4274	0.339
3	0.3149	0.22
4	0.2319	0.126
5	0.3714	0.146
6	0.3714	0.076

Based on both approaches, the results consistently show that the optimal number of clusters is k=2. The agreement between the Elbow Method and the Silhouette Score indicates that the selection of the number of clusters in this study has a good level of reliability.

3.3 K-Means Clustering Results

The clustering results using the K-Means algorithm are shown in Table 6. Based on these results, the data were divided into two main clusters, where most regions were grouped into one cluster, while one region formed a separate cluster with significantly different disaster event characteristics.

Table 6. K-Means Clustering Results

Papua Barat Province	Cluster	Silhouette Score	Total Flood and Coastal Flood Events	Total Other Disaster Events	Total Earthquake Events
Manokwari	C2	0.5	1.914	2.296	-0.3303
Pegunungan Arfak	C1	0.633696	-0.563	-0.076	-0.1343
Manokwari Selatan	C1	0.61877	0.45	-0.512	0.5283
Teluk Wondama	C1	0.64463	-0.45	-0.657	-0.3236
Kaimana	C1	0.636126	-1.239	-0.706	1.414
Fakfak	C1	0.571805	-0.788	0.311	-1.9598
Teluk Bintuni	C1	0.609626	0.676	-0.657	0.8055

3.4 Hierarchical Clustering Results

The clustering results using Hierarchical Clustering are shown in Figure 4 in the form of a dendrogram. Based on the dendrogram cut at k=2, two main clusters were obtained, representing the grouping of regions based on data similarity. This approach enables grouping to be carried out gradually based on the proximity among objects [7], [8].

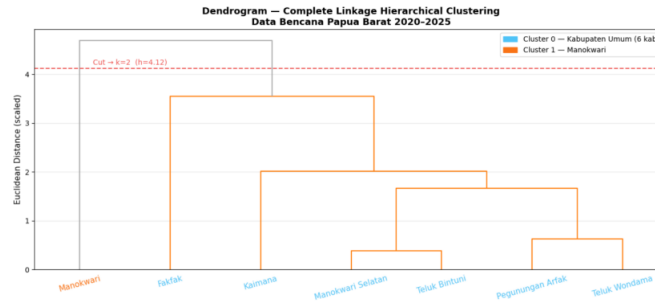


Figure 4. Hierarchical Clustering Dendrogram

The dendrogram structure shows that several regions have a high level of proximity and merge at relatively small distances, while other regions merge at larger distances. This pattern indicates variations in disaster event characteristics among regions, which can be observed through the hierarchical cluster merging structure [12].

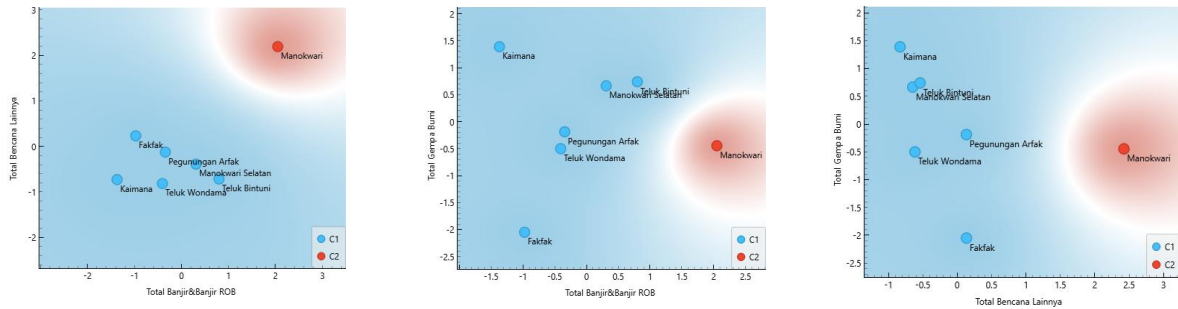
The Hierarchical Clustering results are presented in Table 7. The data in the table are presented in normalized form (z-score), so that each variable is on a comparable scale. Based on these results, most regions were grouped into cluster 0, while Manokwari Regency formed a separate cluster in cluster 1 [7].

Table 7. Hierarchical Clustering Results

Papua Barat Province	Total Flood and Coastal Flood Events	Total Other Disaster Events	Total Earthquake Events	Cluster	Silhouette Score
Pegunungan Arfak	-0.563053	-0.076089	-0.134256	0	0.577946
Manokwari Selatan	0.450443	-0.511869	0.528332	0	0.528448
Teluk Wondama	-0.450443	-0.657129	-0.323567	0	0.613913
Kaimana	-1.238717	-0.705549	1.414037	0	0.522816
Fakfak	-0.788275	0.311272	-1.959754	0	0.273823
Teluk Bintuni	0.675664	-0.657129	0.805537	0	0.474749
Manokwari	1.914381	2.296492	-0.330328	1	0

3.5 Cluster Comparison Scatter Plot

The visualization of clustering results using the K-Means algorithm is shown in Figure 5. The scatter plot shows the data distribution based on the feature pairs used in this study, namely Floods & Tidal Floods, Other Disasters, and Earthquakes. Based on this visualization, K-Means Clustering produces relatively compact and clear cluster separation in several feature pairs, allowing the data grouping patterns to be clearly observed [22], [28], [29].



(a) Total Flood and Coastal Flood Events vs Total Other Disasters

(b) Total Flood and Coastal Flood Events vs Total Earthquake Events

(c) Total Other Disasters vs Total Earthquake Events

Figure 5. K-Means Clustering Scatter Plot

The visualization of the clustering results using Hierarchical Clustering is shown in Figure 6. Based on the scatter plot, the data distribution shows a more gradual grouping pattern according to the level of proximity among regions. Several regions appear to be positioned close to one another, whereas other regions tend to be separated from the main group [7].

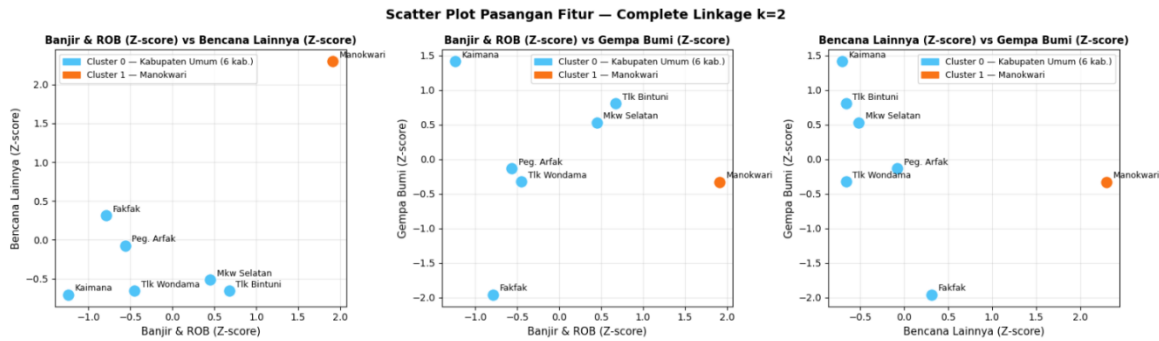


Figure 6. Hierarchical Clustering Scatter Plot

Based on the figure, the resulting cluster separation appears compact and relatively clear, indicating that K-Means Clustering was able to group the data effectively using the predetermined number of clusters [5], [6]. Meanwhile, Hierarchical Clustering showed a more gradual data distribution, in which relationships among data points were formed based on different levels of proximity.

3.6 Results Analysis

Based on the research results, both algorithms produced the same optimal number of clusters, namely $k = 2$. However, the Silhouette Score values indicate that Hierarchical Clustering had better cluster separation quality than K-Means Clustering, with a Silhouette Score of 0.4274, whereas K-Means Clustering obtained a score of 0.339. A higher Silhouette Score indicates that the level of cohesion and separation among clusters in Hierarchical Clustering better represented the structure of the research data [13].

The difference in results is influenced by the different cluster formation mechanisms. Centroid-based K-Means Clustering tends to produce compact clusters, whereas Hierarchical Clustering considers relationships among data points gradually based on the level of proximity among objects [6], [7], [17], [22]. This result is consistent with previous studies showing that Hierarchical Clustering can better represent hierarchical relationships among regions compared to centroid-based methods such as K-Means Clustering [8].

The scatter plot and dendrogram visualizations show that Manokwari Regency tends to form a separate cluster, indicating disaster event characteristics that differ significantly from other regions. This difference is mainly observed in the Floods and Tidal Floods variable and the Other Disasters variable, which have relatively higher

values than those of other regions. This finding is consistent with previous studies showing that clustering methods can identify regional vulnerability patterns based on data characteristics [5], [8], [10].

Although the scope of seven regencies is consistent with the administrative boundaries of Papua Barat Province after the regional division, the limited number of objects remains a methodological limitation. The small dataset makes the cluster structure sensitive to the characteristics of certain regions and limits the generalizability of the results. In addition, the use of three aggregated numerical variables limits the representation of more specific disaster vulnerability. Therefore, the clustering results should be understood as an initial exploratory analysis and not as the sole basis for determining disaster vulnerability levels. Future studies are recommended to include spatial/GIS variables, socio-economic variables, or other clustering methods to strengthen the generalizability and applicability of the analysis results [8], [12]. The use of spatial data, such as regional coordinates, distance to rivers or coasts, slope, and population density, can help improve clustering quality and provide a deeper understanding of regional vulnerability.

The practical implication of this study is that Manokwari Regency can be prioritized for disaster mitigation through resource allocation, improved preparedness, strengthened early warning systems, and disaster education programs. Overall, both algorithms are able to group disaster-prone areas in Papua Barat Province effectively. However, based on the evaluation using the Silhouette Score, Hierarchical Clustering shows better performance on the dataset used in this study. This finding indicates that the selection of a clustering method should be adjusted to the characteristics of the data and the objectives of the analysis in order to obtain more optimal grouping results [1].

4. Conclusions and Future Works

Based on the research results, the K-Means Clustering and Hierarchical Clustering algorithms were able to group disaster-prone areas in Papua Barat Province into two main clusters, with Hierarchical Clustering showing better performance due to its gradual grouping mechanism that preserves proximity among objects, making it more effective than centroid-based K-Means, particularly on a small dataset with significant variations in regional characteristics.

The scatter plot and dendrogram results show that Manokwari Regency formed a separate cluster as an outlier, particularly in the Floods, Tidal Floods, and Other Disasters variables. This indicates that clustering methods can help identify data-driven regional vulnerability patterns to support disaster mitigation and decision-making.

This study is still limited by the small amount of data, three aggregated numerical variables, and the use of a single internal evaluation metric. Therefore, the results should be understood as an initial exploratory analysis and not as the sole basis for determining disaster vulnerability levels.

Future studies are recommended to use larger datasets, add spatial/GIS variables such as regional coordinates, distance to rivers or coasts, slope, and population density, include socio-economic variables, and compare additional clustering methods or evaluation metrics. A spatio-temporal clustering approach can also be applied to make disaster-prone area analysis more comprehensive and applicable.

Overall, this study provides theoretical and practical contributions by showing that the selection of an appropriate clustering algorithm, supported by scatter plot and dendrogram visualizations, can support the systematic identification of disaster-prone areas and more effective data-driven disaster mitigation strategies.

5. References

- [1] S. Z. S. Abdalla, "Data-driven innovations in disaster risk management: Advancing resilience and sustainability through big data analytics," *Progress in Disaster Science*, vol. 27, p. 100451, Oct. 2025, doi: 10.1016/J.PDISAS.2025.100451.

- [2] A. Kondraganti, G. Narayanamurthy, and H. Sharifi, "A systematic literature review on the use of big data analytics in humanitarian and disaster operations," *Annals of Operations Research* 2023 335:3, vol. 335, no. 3, pp. 1015–1052, Nov. 2022, doi: 10.1007/S10479-022-04904-Z.
- [3] BPK Perwakilan Provinsi Papua Barat, "Provinsi Papua Barat." Accessed: May 15, 2026. [Online]. Available: <https://papuabarabpk.go.id/wilayah-pemeriksaan/provinsi-papua-barat/>
- [4] A. Jaeger and D. Banks, "Cluster analysis: A modern statistical review," *Wiley Interdisciplinary Reviews: Computational Statistics*, vol. 15, no. 3, p. e1597, May 2023, doi: 10.1002/WICS.1597.
- [5] T. T. Tin *et al.*, "Natural Disaster Clustering Using K-Means, DBSCAN, SOM, GMM, and Mean Shift: An Analysis of Fema Disaster Statistics," *International Journal of Advanced Computer Science and Applications*, vol. 15, no. 9, pp. 667–681, Sep. 2024, doi: 10.14569/IJACSA.2024.0150968.
- [6] A. M. Ikotun, A. E. Ezugwu, L. Abualigah, B. Abuhaija, and J. Heming, "K-means clustering algorithms: A comprehensive review, variants analysis, and advances in the era of big data," *Information Sciences*, vol. 622, pp. 178–210, Apr. 2023, doi: 10.1016/J.INS.2022.11.139.
- [7] T. Li, A. Rezaeipannah, and E. S. M. Tag El Din, "An ensemble agglomerative hierarchical clustering algorithm based on clusters clustering technique and the novel similarity measurement," *Journal of King Saud University - Computer and Information Sciences*, vol. 34, no. 6, pp. 3828–3842, Jun. 2022, doi: 10.1016/J.JKSUCI.2022.04.010.
- [8] S. Brina and W. Rizky, "Perbandingan Algoritma K-Means dan Hierarchical Clustering dalam Segmentasi Kabupaten/Kota di Jawa Timur Berdasarkan Data Perjalanan dan Pergerakan Wisatawan," *JATI (Jurnal Mahasiswa Teknik Informatika)*, vol. 9, no. 5, pp. 8499–8506, Jul. 2025, doi: 10.36040/JATI.V9I5.15091.
- [9] E. Choi and J. Song, "Clustering-based disaster resilience assessment of South Korea communities building portfolios using open GIS and census data," *International Journal of Disaster Risk Reduction*, vol. 71, p. 102817, Mar. 2022, doi: 10.1016/J.IJDRR.2022.102817.
- [10] G. Setapati, D. Hafidh Zulfikar, R. Fatah Palembang, P. Sistem Informasi, F. Sains dan Teknologi, and U. Raden Intan Lampung, "Algoritma K-Means untuk Mengelompokkan Daerah Rawan Bencana di Sumatera Selatan," *Journal of software engineering and computational intelligence*, vol. 2, no. 2, p. 11, 2024.
- [11] C. Wongoutong, "The impact of neglecting feature scaling in k-means clustering," *PLOS ONE*, vol. 19, no. 12, p. e0310839, Dec. 2024, doi: 10.1371/JOURNAL.PONE.0310839.
- [12] J. Ge and R. Tibshirani, "Optimal pruning of hierarchical clustering dendrograms," *Communications in Statistics - Theory and Methods*, 2025, doi: 10.1080/03610926.2025.2543191.
- [13] M. Shutaywi, "Silhouette Analysis for Performance Evaluation in Machine Learning with Applications to Clustering," pp. 1–17, 2021, doi: 10.3390/e23060759.
- [14] J. Demšar and B. Zupan, "Hands-on training about data clustering with orange data mining toolbox," *PLOS Computational Biology*, vol. 20, no. 12, p. e1012574, Dec. 2024, doi: 10.1371/JOURNAL.PCBI.1012574.
- [15] Z. Dobesova, "Evaluation of Orange data mining software and examples for lecturing machine learning tasks in geoinformatics," *Computer Applications in Engineering Education*, vol. 32, no. 4, p. e22735, Jul. 2024, doi: 10.1002/CAE.22735.
- [16] J. C. N. Bittencourt, D. G. Costa, P. Portugal, and F. Vasques, "A data-driven clustering approach for assessing spatiotemporal vulnerability to urban emergencies," *Sustainable Cities and Society*, vol. 108, p. 105477, Aug. 2024, doi: 10.1016/J.SCS.2024.105477.

- [17] D. P. Dewi, I. H. Santi, and W. D. Puspitasari, "Perhitungan Penilaian Tingkat Kepuasan Pelanggan Dengan Menerapkan Algoritma K-Means," *J-INTECH*, vol. 11, no. 2, pp. 257–265, Dec. 2023, doi: 10.32664/J-INTECH.V11I2.981.
- [18] V. P. Ramadhan and G. M. Namung, "Klasterisasi Komentar Cyberbullying Masyarakat di Instagram berdasarkan K-Means Clustering," *J-INTECH*, vol. 11, no. 1, pp. 32–39, Jul. 2023, doi: 10.32664/J-INTECH.V11I1.846.
- [19] P. Rinekso Andriyanto, J. Santoso, and E. Setyati, "Identifikasi Viseme Untuk Fonem Bahasa Madura Berbasis Clustering Berdasarkan Facial Landmark Point," *J-INTECH*, vol. 11, no. 1, pp. 73–82, Jul. 2023, doi: 10.32664/J-INTECH.V11I1.835.
- [20] E. Yuniar, S. Ramadhania, P. S. Universitasari, M. Hermansyah, and A. B. Setiawan, "Segmentation and Prediction of Store Performance on the Shopee Marketplace Using a Hybrid Clustering Approach, Spatial Analysis, and Feature Importance," *J-INTECH*, vol. 14, no. 01, pp. 121–129, Mar. 2026, doi: 10.32664/J-INTECH.V14I01.2256.
- [21] I. T. Utami, F. Suryaningrum, D. Ispriyanti, I. T. Utami, F. Suryaningrum, and D. Ispriyanti, "K-Means Cluster Count Optimization with Silhouette Index Validation and Davies Bouldin Index (Case Study: Coverage of Pregnant Women, Childbirth, and postpartum health services in Indonesia in 2020)," *BAREKENG: Jurnal Ilmu Matematika dan Terapan*, vol. 17, no. 2, pp. 0707–0716, Jun. 2023, doi: 10.30598/BAREKENG.VOL17ISS2PP0707-0716.
- [22] H. Herayati, A. Siregar, and H. Hariyanto, "K-Means Clustering Algorithm Measuring the Satisfaction Level of MNC TV Muslim I'murojaah Program Viewers," *J-INTECH*, vol. 13, no. 01, pp. 179–191, Jul. 2025, doi: 10.32664/J-INTECH.V13I01.1926.
- [23] D. Wintana, H. Sulaiman, R. S. Rohman, G. Gunawan, and M. A. Ghani, "Analyzing Students' Interest in Mathematics Through the Implementation of the K-Means Clustering Algorithm," *J-INTECH*, vol. 13, no. 01, pp. 71–77, Jun. 2025, doi: 10.32664/J-INTECH.V13I01.1861.
- [24] J. J. A. Limbong and F. M. Likumahwa, "Implementation of the K-Means Algorithm for Clustering Students' Web Programming Course Grades Using Silhouette Score," vol. 6, no. 3, pp. 271–280, 2025, doi: <https://doi.org/10.61628/jsce.v6i3.2100>.
- [25] S. Wijitkosum, "Integrated spatial analysis of drought risk factors using agglomerative hierarchical clustering and correlation," *Environmental Advances*, vol. 21, p. 100646, Oct. 2025, doi: 10.1016/J.ENVADV.2025.100646.
- [26] D. Chicco, A. Campagner, A. Spagnolo, D. Ciucci, and G. Jurman, "The Silhouette coefficient and the Davies-Bouldin index are more informative than Dunn index, Calinski-Harabasz index, Shannon entropy, and Gap statistic for unsupervised clustering internal evaluation of two convex clusters," *PeerJ Computer Science*, vol. 11, p. e3309, Nov. 2025, doi: 10.7717/PEERJ-CS.3309/TABLE-18.
- [27] M. Gagolewski, M. Bartoszek, and A. Cena, "Are cluster validity measures (in) valid?," *Information Sciences*, vol. 581, pp. 620–636, Dec. 2021, doi: 10.1016/J.INS.2021.10.004.
- [28] A. M. Ikotun, F. Habyarimana, and A. E. Ezugwu, "Cluster validity indices for automatic clustering: A comprehensive review," *Heliyon*, vol. 11, no. 2, p. e41953, Jan. 2025, doi: 10.1016/J.HELIVON.2025.E41953.
- [29] Q. Xin, "Self-Supervised Customer Representation Learning for Segmentation and Next-Purchase Prediction on UCI Online Retail," *J-INTECH*, vol. 14, no. 01, pp. 20–37, Apr. 2026, doi: 10.32664/J-INTECH.V14I01.2229.