



Implementation of K-Means Clustering for Student Achievement Classification Using Academic Performance Indicators

Dedi Candro Parulian Sinaga*, Endra Ary Prasasty Marpaung, Nera Mayana Br Tarigan, Vinsensius Fereri Purba, Dwicki Aditya

Informatics Engineering, STMIK Pelita Nusantara, Medan, Indonesia

Email: ^{1,*}dedisinaga27@gmail.com, ²marpaungendra83@gmail.com, ³neramayana658@gmail.com, ⁴vfpurba2121@gmail.com, ⁵dwikiadit20@gmail.com

(*: dedisinaga27@gmail.com)

Submitted: 29/06/2026; Accepted: 19/07/2026; Published: 29/07/2026

Abstract— Student achievement assessment in educational institutions is often conducted manually by relying on average scores or class rankings, making the results less objective and unable to represent students' overall academic conditions. This study aims to implement the K-Means Clustering algorithm to classify student achievement using academic performance indicators. The dataset consists of 10 student records with four variables: average score, examination score, attendance percentage, and assignment score. The research stages include data collection, data validation, Min-Max normalization, initial centroid selection, Euclidean Distance calculation, cluster formation, centroid updating, and interpretation of clustering results. The number of clusters was set to K=3, representing high, medium, and low achievement categories. The results show that Cluster C1 consists of 4 students with high achievement, Cluster C2 consists of 3 students with medium achievement, and Cluster C3 consists of 3 students with low achievement. The final centroid values indicate that Cluster C1 has the strongest academic performance, while Cluster C3 requires more intensive academic support. These findings demonstrate that K-Means Clustering can classify student achievement objectively and support data-driven educational decision-making, academic guidance, and targeted learning strategy development.

Keywords: K-Means Clustering; Student Achievement; Data Mining; Academic Performance; Educational Decision-Making.

1. INTRODUCTION

The rapid development of information technology has brought significant changes in various fields, including education. One positive impact of this development is the availability of large amounts of data that can be used to improve the quality of the learning process and decision-making[1]. Student academic data, such as grades, attendance, and activity, are invaluable sources of information when processed using appropriate methods[2]. However, in reality, many educational institutions have not utilized this data optimally, so the important information contained therein cannot be utilized to its full potential[3]. Student achievement assessment is an important aspect in the world of education. Student achievement not only reflects academic ability but also serves as the basis for determining various policies, such as awarding awards, determining development programs, and evaluating learning methods[4]. However, the process of assessing student achievement is often still carried out conventionally, namely by looking at average grades or class rankings. This approach tends to be simplistic and is less able to describe the actual conditions, because it does not consider various other factors that influence student achievement as a whole[5].

Furthermore, subjectivity in assessment is also a fairly common problem. Determining the categories of high-achieving students is often influenced by individual perceptions, resulting in less objective and consistent results[6]. On the other hand, the increasing amount of student data also presents challenges for schools in conducting manual analysis[7]. This makes the decision-making process less effective and potentially results in inappropriate policies[8]. To address these issues, a data-driven approach is needed that can systematically process and analyze data. One approach that can be used is data mining[9]. Data mining is the process of extracting patterns or important information from large data sets using specific techniques[10]. In the educational context, data mining can be used to identify student achievement patterns, group students based on specific characteristics, and support more accurate and objective decision-making[11].

One popular method in data mining is the K-Means Clustering algorithm. This method is used to group data into several clusters based on the level of similarity between the data[12]. K-Means works by determining a number of cluster centers (centroids), then grouping the data based on the closest distance to the centroid. This process is carried out iteratively until optimal clustering results are obtained[13]. The advantage of the K-Means method is its ability to process large amounts of data with a relatively fast and simple process[14]. In this study, the K-Means Clustering algorithm was used to analyze student achievement data with the aim of grouping students into several categories, such as high, medium, and low achievers[15]. This grouping is based on various relevant attributes, such as academic grades, attendance levels, and student activeness in learning activities. By using this approach, it is hoped that the clustering results can provide a clearer and more objective picture of the condition of student achievement[16].

The application of the K-Means Clustering method to student achievement data analysis has several benefits[17]. First, this method can help schools understand student characteristics more deeply. Second, the clustering results can be used as a basis for designing more effective and targeted learning strategies. Third, this method can also help determine policies related to awarding and coaching programs for students who require special attention. Furthermore, this research is expected to encourage the use of data mining technology in education, particularly in academic data management. With a data-based analysis system, decision-making can be carried out more quickly, precisely, and objectively. This will certainly have a positive impact on improving the overall quality of education.





Therefore, this research is crucial to optimize the use of student academic data through the application of the K-Means Clustering algorithm. The results of this study are expected to not only contribute to the development of science but also provide practical benefits for educational institutions in improving the quality of data management and decision-making. Therefore, analyzing student achievement data using the K-Means Clustering method is a relevant and effective solution to address the challenges of data management in today's digital era.

2. RESEARCH METHODOLOGY

2.1 Research Stages

The approach used in this study is a quantitative approach with data mining methods. This quantitative approach was chosen because this study focuses on processing numerical data in the form of student achievement scores, which are analyzed mathematically and statistically. Furthermore, this study also uses an unsupervised learning approach, as there are no initial category labels in the student data. The categories of high-achieving students will be formed automatically through a grouping process using the K-Means Clustering algorithm.

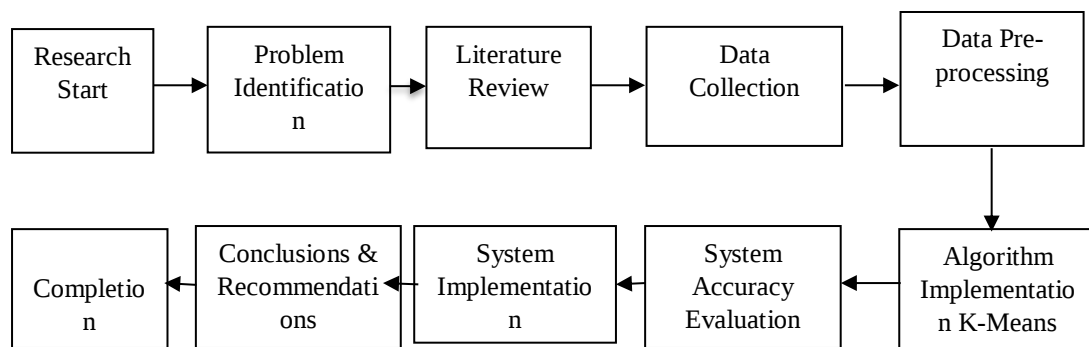


Fig 1. Research Framework

The research framework describes the flow of activities from the initial to the final stages of the research. In general, these stages include:

1. Starting the Research
The initial stage involves determining the topic, identifying research needs, and formulating the objectives to be achieved.
2. Problem Identification
Examining issues related to student achievement assessment, which is still manual, subjective, and not based on data analysis.
3. Literature Review
Collecting references related to data mining, clustering, and the K-Means algorithm, as well as relevant previous research as a theoretical basis.
4. Data Collection
Collecting student achievement data, such as academic grades, attendance, and learning activity from valid sources.
5. Data Preprocessing
Performing data cleaning, attribute selection, and normalization to prepare the data for processing.
6. Data Preprocessing
The K-Means algorithm is implemented through the following stages:
 - a. Initial Centroid Selection
Initial centroids are selected randomly from the normalized dataset.
 - b. Distance Calculation
The Euclidean Distance formula is used to calculate the distance between each data point and centroid.

$$d(x,y) = \sqrt{\sum_{i=1}^n (x_i - y_i)^2} \quad (1)$$
 where:
 $d(x,y)$ = distance between data point and centroid
 x_i = data value
 y_i = centroid value
7. Implementing the K-Means Clustering Algorithm
 - a. Determining the initial centroid
 - b. Calculating the distance between data (Euclidean Distance)
 - c. Grouping data into the closest cluster



- d. Updating the centroid until convergence
8. Clustering Results Analysis
Interpreting the student clustering results based on the characteristics of each cluster.
9. Clustering Results Evaluation
Assessing the quality of the clustering results to ensure that the formed clusters are representative.
10. Results Implementation
Using the clustering results as a basis for determining high-achieving student categories and supporting decision-making.
11. Conclusions and Recommendations
Drawing conclusions from the research results and providing recommendations for further research development.
12. Finished

2.2 Dataset Description

The dataset used in this study was obtained from academic records of students at a private secondary school. The data consist of 10 student records collected from the school academic administration database during one academic semester. Four variables were selected as clustering attributes, namely Average Score, Examination Score, Attendance Percentage, and Assignment Score. These variables were chosen because they represent both academic achievement and learning participation. The dataset was anonymized to protect student privacy by replacing student names with identification codes (S1–S10). Prior to clustering, the data were validated to ensure completeness and consistency. No missing values were found in the dataset. Subsequently, Min-Max Normalization was applied to transform all variables into the same scale range between 0 and 1 before performing the clustering process.

2.3 Data Normalization

Since the variables have different scales, Min-Max Normalization was applied using:

$$x' = \frac{x - x_{min}}{x_{max} - x_{min}} \quad (2)$$

Description:

x = original value

xmin = minimum value

xmax = maximum value

This transformation scales all attributes into the range of 0–1.

3. RESULT AND DISCUSSION

This section presents the results of data processing and discusses the implementation of the K-Means Clustering algorithm for classifying student achievement based on academic performance indicators. The analysis was conducted using student academic data consisting of four main variables: average score, examination score, attendance percentage, and assignment score. These variables were selected because they represent both academic achievement and student participation in the learning process. The clustering process was carried out through several stages, including data preparation, normalization, initial centroid selection, distance calculation, cluster formation, centroid updating, and interpretation of the final clustering results.

3.1 Research result

The dataset used in this study consists of 10 students identified using codes S1 to S10. Each student has four academic performance indicators: average score, examination score, attendance percentage, and assignment score. These indicators were used as input variables in the K-Means Clustering process. The initial dataset is presented in Table 1.

Table 1. Initial Dataset Used in This Study

No	Name	Average	Exam	Attendance (%)	Assignment
1	S1	85	88	95	90
2	S2	78	80	90	82
3	S3	60	65	80	68
4	S4	92	94	98	95
5	S5	70	72	85	75





No	Name	Average	Exam	Attendance (%)	Assignment
6	S6	88	90	96	92
7	S7	55	58	75	60
8	S8	76	78	88	80
9	S9	65	68	82	70
10	S10	90	92	97	94

Before the clustering process was performed, the minimum and maximum values of each variable were identified. These values were required for the normalization process because the variables used in the dataset have different value ranges. The minimum and maximum values for each variable are shown in Table 2.

Table 2. Minimum and Maximum Values

Variabel	Min	Max
Average	55	92
Exam	58	94
Attendance	75	98
Assignment	60	95

The data were then normalized using the Min-Max Normalization method to transform all variables into the same scale range between 0 and 1. This step was necessary because the K-Means algorithm uses distance-based calculation, and variables with larger numerical ranges may dominate the clustering result if normalization is not performed. The normalization results are presented in Table 3.

Table 3. Normalization Results

Name	Average	Exam	Attendance	Assignment
S1	0.81	0.83	0.87	0.86
S2	0.62	0.61	0.65	0.63
S3	0.14	0.19	0.22	0.23
S4	1.00	1.00	1.00	1.00
S5	0.41	0.39	0.43	0.43
S6	0.89	0.89	0.91	0.91
S7	0.00	0.00	0.00	0.00
S8	0.57	0.56	0.57	0.57
S9	0.27	0.28	0.30	0.29
S10	0.95	0.94	0.96	0.97

The normalized data show that students with values closer to 1 have stronger academic performance, while students with values closer to 0 have weaker academic performance. For example, S4 obtained a normalized value of 1.00 in all variables, indicating that this student achieved the highest score in average, exam, attendance, and assignment indicators. In contrast, S7 obtained a normalized value of 0.00 in all variables, indicating that this student had the lowest performance among the analyzed data. After the normalization process, the K-Means Clustering algorithm was implemented by determining the number of clusters. In this study, the number of clusters was set to $K = 3$, representing three categories of student achievement: high achievement, medium achievement, and low achievement. The initial centroids were selected from the normalized dataset to represent each performance category. The initial centroid selection is shown in Table 4.

Table 4. Initial Centroid Selection

Cluster	Student	Average	Exam	Attendance	Assignment
C1 (High)	S4	1.00	1.00	1.00	1.00
C2 (Medium)	S2	0.62	0.61	0.65	0.63
C3 (Low)	S7	0.00	0.00	0.00	0.00

The next stage was calculating the distance between each student and each centroid using Euclidean Distance. Each student was assigned to the cluster with the closest centroid value. After all students were grouped, the centroid values were updated based on the average value of the members in each cluster. This process was repeated until there was no change in cluster membership or until the centroid values became stable. The final grouping results are presented in Table 5.

Table 5. Grouping Results



Name	Cluster	Category
S1	C1	High
S2	C2	Medium
S3	C3	Low
S4	C1	High
S5	C2	Medium
S6	C1	High
S7	C3	Low
S8	C2	Medium
S9	C3	Low
S10	C1	High

Based on Table 5, the 10 students were successfully classified into three achievement categories. Students S1, S4, S6, and S10 were grouped into Cluster C1 as high-achievement students. Students S2, S5, and S8 were grouped into Cluster C2 as medium-achievement students. Meanwhile, students S3, S7, and S9 were grouped into Cluster C3 as low-achievement students. The distribution of students in each cluster is shown in Table 6.

Table 6. Cluster Distribution

Cluster	Category	Amount
C1	High	4
C2	Medium	3
C3	Low	3

Table 6 shows that the high-achievement cluster consists of four students, while the medium-achievement and low-achievement clusters consist of three students each. This result indicates that the student performance data are not homogeneous. There are clear differences in academic performance among students based on the four indicators used in the analysis. The final clustering result is summarized in Table 7.

Table 7. Final Clustering Result

Cluster	Category	Students	Total
C1	High	S1, S4, S6, S10	4
C2	Medium	S2, S5, S8	3
C3	Low	S3, S7, S9	3

The final centroid values obtained after the clustering process reached convergence are presented in Table 8. These centroid values represent the average characteristics of each cluster based on the four academic performance indicators.

Table 8. Final Centroid Values

Cluster	Average	Exam	Attendance	Assignment
C1	88.75	91.00	96.50	92.75
C2	74.67	76.67	87.67	79.00
C3	60.00	63.67	79.00	66.00

The centroid values in Table 8 show clear differences among the three clusters. Cluster C1 has the highest centroid values in all indicators, with an average score of 88.75, examination score of 91.00, attendance rate of 96.50, and assignment score of 92.75. These values indicate that students in this cluster have strong academic performance, high discipline in attendance, and good consistency in completing assignments. Therefore, Cluster C1 can be interpreted as the high-achievement group.

Cluster C2 has moderate centroid values, with an average score of 74.67, examination score of 76.67, attendance rate of 87.67, and assignment score of 79.00. Students in this cluster show acceptable academic performance, but their scores are still lower than those in Cluster C1. This cluster can be interpreted as the medium-achievement group. Students in this category have the potential to improve if they receive proper academic guidance, learning motivation, and monitoring from teachers.

Cluster C3 has the lowest centroid values, with an average score of 60.00, examination score of 63.67, attendance rate of 79.00, and assignment score of 66.00. These values indicate that students in this cluster experience academic



difficulties compared with students in the other clusters. Low attendance and assignment scores may contribute to weaker academic performance. Therefore, Cluster C3 can be interpreted as the low-achievement group that requires special attention, such as additional tutoring, academic mentoring, and closer evaluation by teachers.

3.2 Discussion

The results of this study show that the K-Means Clustering algorithm can classify students into achievement groups based on similarities in academic performance indicators. The use of four variables, namely average score, examination score, attendance percentage, and assignment score, provides a broader perspective than assessment based only on final grades. This is important because student achievement is not determined solely by examination results, but also by attendance discipline and consistency in completing assignments.

The clustering results demonstrate that Cluster C1 contains students with the strongest academic characteristics. Students in this group achieved high scores in all variables, especially in attendance and assignment completion. This indicates that students with consistent attendance and assignment performance tend to have better academic results. Therefore, students in Cluster C1 can be considered suitable candidates for academic appreciation, enrichment programs, or participation in advanced learning activities. Cluster C2 represents students with moderate performance. Although students in this group did not achieve the highest scores, their academic indicators are still relatively stable. This cluster is important because students in this group have the potential to move into the high-achievement category if they receive appropriate support. Schools can design improvement programs for this group, such as additional exercises, learning motivation programs, and periodic academic monitoring. Cluster C3 consists of students with the lowest performance across the four indicators. The centroid values show that this cluster has the lowest average score, examination score, attendance percentage, and assignment score. This condition indicates that students in Cluster C3 need more intensive academic support. The school can provide special guidance, remedial learning, counseling, or communication with parents to identify the factors affecting students' academic performance. The findings also show that the K-Means Clustering method is useful for reducing subjectivity in student achievement classification. In conventional assessment, students are often classified only by ranking or average score. Such an approach may not fully represent the overall academic condition of students. Through clustering, each student is grouped based on the similarity of several performance indicators, making the classification process more objective and data-driven.

The results of this study are also consistent with previous research that applied the K-Means Clustering algorithm to classify student achievement based on academic data. Previous studies showed that K-Means can group students into several achievement categories, such as high, medium, and low achievement, by using indicators such as academic scores, attendance, and learning performance [18], [19]. The similarity between this study and previous research lies in the use of clustering to identify student performance patterns objectively and to support educational decision-making.

However, this study has a different focus from previous studies. In this research, student classification was carried out using four academic performance indicators, namely average score, examination score, attendance percentage, and assignment score. These variables provide a more comprehensive view of student achievement because the assessment does not only depend on examination results, but also considers student discipline and learning consistency. In addition, this study presents final centroid values for each cluster, making the characteristics of high, medium, and low achievement groups clearer and easier to interpret. Compared with previous research, the clustering results in this study show that attendance and assignment completion have an important role in forming student achievement groups. Students in the high-achievement cluster showed strong performance across all indicators, while students in the low-achievement cluster had lower values in academic scores, attendance, and assignments. Therefore, this study strengthens previous findings that K-Means Clustering is effective for student achievement classification and can be used as a practical basis for designing academic guidance, enrichment programs, and targeted learning interventions. In addition, the cluster distribution provides practical information for educational decision-making. The high-achievement cluster can be used as a reference for reward and enrichment programs, the medium-achievement cluster can be targeted for improvement programs, and the low-achievement cluster can be prioritized for academic intervention. Thus, the results of this study can help schools design more targeted learning strategies and allocate educational support more effectively. Overall, the implementation of K-Means Clustering in this study successfully produced meaningful student achievement categories. The method was able to identify patterns in academic data and group students according to their performance characteristics. The results support the use of data mining methods in education, especially for academic evaluation, student guidance, and data-based decision-making. However, the interpretation of clustering results should still consider the context of each student, because academic performance may also be influenced by non-academic factors such as motivation, family background, learning environment, and psychological conditions.

4. CONCLUSION

This study concludes that the implementation of the K-Means Clustering algorithm can effectively classify student achievement based on academic performance indicators. The main problem addressed in this research is the limitation of conventional student assessment, which often relies only on average scores or manual ranking and may not fully represent students' overall academic conditions. By using four variables, namely average score, examination score, attendance



percentage, and assignment score, the clustering process provides a more objective and data-driven approach to identifying student achievement patterns. The results show that students were successfully grouped into three clusters: Cluster C1 as the high-achievement group, Cluster C2 as the medium-achievement group, and Cluster C3 as the low-achievement group. Cluster C1 consisted of four students with the highest academic performance, Cluster C2 consisted of three students with moderate performance, and Cluster C3 consisted of three students who required additional academic support.

The final centroid values confirm that each cluster has different academic characteristics. Students in Cluster C1 showed strong performance in all indicators, especially in examination scores, attendance, and assignment completion. Students in Cluster C2 showed stable but not optimal performance, while students in Cluster C3 had the lowest values across all indicators. These findings indicate that K-Means Clustering can support schools in making more accurate decisions related to student guidance, enrichment programs, remedial learning, and academic intervention. However, this study still has limitations, particularly the small dataset and the use of only four academic variables. Future research is recommended to use larger datasets, include non-academic factors such as motivation and learning environment, and apply cluster evaluation methods such as Silhouette Score or Davies–Bouldin Index to improve the validity and reliability of the clustering results.

REFERENCES

- [1] R. Awalia, “Penerapan Metode K-Means Clustering untuk Pengelompokan Prestasi Siswa menggunakan Orange Data Mining: Studi Kasus di MTs Muhammadiyah Tallo Makassar,” *MAPLE*, vol. 6, no. 2, pp. 36–41, Dec. 2024.
- [2] M. F. Firdaus, G. Abdillah, and E. Ramadhan, “KLAUSTERISASI AKADEMIK PRESTASI SISWA MENGGUNAKAN ALGORITMA K-PROTOTYPE,” *JURNAL LOCUS: Penelitian & Pengabdian*, vol. 4, no. 8, pp. 7563–7582, Aug. 2025.
- [3] S. Haviyola and M. Jajuli, “PENGELOMPOKAN PRESTASI SISWA GUNA KUALIFIKASI BEASISWA BERDASARKAN DATA NILAI MENGGUNAKAN ALGORITMA K-MEANS,” Aug. 2023.
- [4] A. Wira Permana, R. Saragih, and G. Prahmana, “IMPLEMENTASI CLUSTERING DATA NILAI SISWA MENGGUNAKAN ALGORITMA K-MEANS: SEBUAH STUDI KASUS DI SMK NASIONAL NAMOTERASI,” 2025. [Online]. Available: <https://journaledutech.com/index.php/great>
- [5] M. Syukron Ramadani and Z. Fatah, “ANALISIS PENGELOMPOKAN DATA NILAI SISWA UNTUK MENENTUKAN SISWA BERPRESTASI MENGGUNAKAN METODE CLUSTERING K-MEANS,” *JURNAL RISET SISTEM INFORMASI*, vol. 1, no. 4, pp. 103–110, Oct. 2024, doi: 10.69714/hq2bsy84.
- [6] H. Sampurno, M. N. Wafi, and N. Nurdin, “Perbandingan Algoritma K-Means dan Hierarchical Clustering dalam Pengelompokan Prestasi Akademik Siswa,” *Format : Jurnal Ilmiah Teknik Informatika*, vol. 15, no. 1, p. 72, Jan. 2026, doi: 10.22441/format.2026.v15.i1.008.
- [7] D. Oktario Dacwanda and Y. Nataliani, “Implementasi k-Means Clustering untuk Analisis Nilai Akademik Siswa Berdasarkan Nilai Pengetahuan dan Keterampilan,” *AITI: Jurnal Teknologi Informasi*, vol. 18, no. Agustus, pp. 125–138, 2021.
- [8] A. F. Habibie and R. K. R., “Implementasi K-Means Clustering dalam Mengklasifikasi Pengaruh Les Terhadap Prestasi Siswa dengan Metode Elbow,” *Journal of Computer System and Informatics (JoSYC)*, vol. 6, no. 1, pp. 133–147, Nov. 2024, doi: 10.47065/josyc.v6i1.6201.
- [9] A. Anilshi, R. T. Abineno, and A. A. Pekuwal, “PENERAPAN ALGORITMA K-MEANS CLUSTERING NILAI KEMAMPUAN SISWA DALAM PELAJARAN MATEMATIKA,” Aug. 2024. [Online]. Available: <https://ojs.unkriswina.ac.id/index.php/semnas-FST>
- [10] N. D. Rahayu, A. H. Anshor, I. Afriantoro, and A. Halim Anshor, “Penerapan Data Mining untuk Pemetaan Siswa Berprestasi menggunakan Metode Clustering K-Means,” *JUKI : Jurnal Komputer dan Informatika*, vol. 6, no. 1, pp. 71–83, May 2024.
- [11] E. A. Saputra and Y. Nataliani, “Analisis Pengelompokan Data Nilai Siswa untuk Menentukan Siswa Berprestasi Menggunakan Metode Clustering K-Means,” *Journal of Information Systems and Informatics*, vol. 3, no. 3, pp. 424–439, Sep. 2021, [Online]. Available: <http://journal-isi.org/index.php/isi>
- [12] M. Azzam Al Fauzie and J. Akhir Putra, “Clustering Data Menggunakan Metode K-Means untuk Rekomendasikan Pembelajaran Akademik bagi Siswa Aktif dalam Ekstrakurikuler,” *KLIK: Kajian Ilmiah Informatika dan Komputer*, vol. 4, no. 1, pp. 642–648, 2023, doi: 10.30865/klik.v4i1.1116.
- [13] D. Apriandi, R. M. Sari, and M. I. Sarif, “Analisis Clustering Untuk Menentukan Siswa Berprestasi di SMK Swasta TI Panca Dharma Stungki Menggunakan Metode K-Means,” *Jurnal Minfo Polgan*, vol. 13, no. 1, pp. 1117–1129, Aug. 2024, doi: 10.33395/jmp.v13i1.13959.
- [14] T. Widyanti, S. S. Hilabi, A. Hananto, Tukino, and E. Novalia, “Implementasi K-Means dan K-Nearest Neighbors pada Kategori Siswa Berprestasi,” *Jurnal Informasi dan Teknologi*, vol. 5, no. 1, pp. 75–82, 2023.
- [15] N. Suarna, N. Rahaningsih, and A. A. Suarna, “OPTIMALISASI PRESTASI AKADEMIK SISWA MELALUI PENGELOMPOKAN INDEKS PRESTASI DENGAN K-MEANS CLUSTERING,” *Jurnal Kecerdasan Buatan dan Teknologi Informasi*, vol. 4, no. 2, pp. 198–207, May 2025, doi: 10.69916/jkbt.v4i2.321.
- [16] N. Jannah and T. Yulianto, “Mengelompokkan Siswa Berprestasi Akademik dengan Menggunakan Metode K Means Kelas VII MT,” *Zeta Math Journal*, vol. 2, no. 2, pp. 41–45, Nov. 2016.





- [17] M. Sholeh and D. Andayati, “Penerapan Metode Clustering dengan Algoritma K-Means Pada Pengelompokan Indeks Prestasi Akademik Mahasiswa,” *SKANIKA: Sistem Komputer dan Teknik Informatika*, vol. 6, no. 1, pp. 51–60, Jan. 2023.
- [18] N. D. Rahayu, A. H. Anshor, I. Afriantoro, and A. Halim Anshor, “Penerapan Data Mining untuk Pemetaan Siswa Berprestasi menggunakan Metode Clustering K-Means Oleh : Penerapan Data Mining untuk Pemetaan Siswa Berprestasi menggunakan Metode Clustering K-Means,” *JUKI : Jurnal Komputer dan Informatika*, vol. 6, 2024.
- [19] E. A. Saputra and Y. Nataliani, “Analisis Pengelompokan Data Nilai Siswa untuk Menentukan Siswa Berprestasi Menggunakan Metode Clustering K-Means,” *Journal of Information Systems and Informatics*, vol. 3, no. 3, 2021, [Online]. Available: <http://journal-isi.org/index.php/isi>