

Application of *K-Means Clustering* to Group Islamic Boarding School Students' Kitab Kuning Reading Proficiency Based on *Nahwu* and *Sharraf*

Luluk Suhartini¹, Burhanis Sulthan², Milyanatus Saadah³
^{1,2,3} Teknologi Informasi, Fakultas Teknik, Universitas Annuqayah, Indonesia

Article Info

Article history:

Received : 05 10, 2026

Revised : 05 21, 2026

Accepted : 06 08, 2026

Keywords:

K-Means;
Clustering;
Data Mining;
Kitab Kuning;
Nahwu Sharraf;

Abstract

The ability to read classical Islamic texts (kitab kuning) is an important indicator for assessing students' mastery of Islamic sciences in Salafiyah Islamic boarding schools. However, the assessment of kitab kuning reading proficiency is often still conducted conventionally and tends to be subjective; therefore, a data-driven approach is needed to support a more objective student-grouping process. This study aims to apply the K-Means Clustering algorithm to group students based on their Nahwu and Sharraf scores as indicators of kitab kuning reading proficiency. The research methods included observation, interviews, literature review, data transformation, normalization using Min-Max Normalization, and clustering using Microsoft Excel and RapidMiner. The dataset consisted of 376 student records with the attributes of age, education level, Nahwu score, and Sharraf score. Manual grouping using Microsoft Excel produced three clusters: C1 with 57 students, C2 with 244 students, and C3 with 75 students. Meanwhile, the RapidMiner results showed C1 with 56 students, C2 with 244 students, and C3 with 76 students. The results indicate that K-Means can be used as a basis for grouping students' abilities in a more structured manner.

Corresponding Author:

Luluk Suhartini,
lulukkhafi@gmail.com

INTRODUCTION

Along with advances in information technology, the need to process information from collections of data has grown rapidly. Information technology facilitates the resolution of various complex problems in almost every area of life, including education (Ahmad Fajrul Falaah et al., 2022). In the educational context, Islamic boarding schools (pesantren) are a form of Islamic education that requires attention, particularly traditional Salaf Islamic boarding schools, which have a distinctive learning system and continue to preserve the classical Islamic scholarly tradition. According to KH. Imam Zarkasyi, pesantren are indigenous Indonesian Islamic educational institutions. Long before formal schools developed, pesantren education played an important role in disseminating Islam, providing religious education, transmitting knowledge, preserving Islamic traditions, and developing Islamic scholars and social activists (Larasati et al., 2022).

One important characteristic of education in traditional Salaf Islamic boarding schools is the study of classical Islamic texts (kitab kuning). Kitab kuning is one of the means through which earlier Islamic scholars transmitted knowledge to subsequent generations, while also constituting an enduring scholarly contribution (Sukron et al., 2022). The study of kitab kuning remains a central component of education in Salaf Islamic boarding schools because

these texts contain various Islamic sciences that form the foundation of students' understanding. In practice, however, the assessment of students' ability to read kitab kuning in Salaf Islamic boarding schools is still commonly conducted conventionally and subjectively. Such assessments generally depend on the examiners' observations and interpretations, which may lead to inconsistent evaluation results among examiners (Novianti et al., 2022).

Al-Is'af Salafiyah Islamic Boarding School is an Islamic educational institution that focuses on the development of Salaf-based Islamic sciences. Students' ability to read kitab kuning is an important indicator for measuring the extent to which they understand and apply Islamic knowledge during the learning process. Based on observations, the assessment of kitab kuning reading proficiency at Al-Is'af Salafiyah Islamic Boarding School is still largely conducted conventionally and tends to be subjective because it relies on the direct observations of teachers or ustaz. This condition may create disparities in the assessment and evaluation of students' abilities, particularly when the process is not supported by more objective data grouping. One approach that can assist in grouping students according to their abilities is data mining using a clustering method. The K-Means algorithm is a widely used clustering technique because it has a simple concept and can effectively group large amounts of data (Irfansyah & Zaehol Fatah, 2024). Several previous studies have used K-Means to address data-grouping problems, such as the study by Firza et al. on academic-major specialization among students at Pelita Raya Private School in Jambi City (Maori & Evanita, 2023). Another study implemented K-Means Clustering to classify PPTQ learning groups as a basis for group allocation so that the learning process could be conducted more efficiently (Sulistiyawati & Supriyanto, 2021).

Based on previous studies, the K-Means algorithm can be used as an effective and efficient data-grouping method. Therefore, this method is expected to be applicable to grouping students at Al-Is'af Salafiyah Islamic Boarding School according to their kitab kuning reading scores by using Nahwu and Sharraf scores. Nahwu and Sharraf play important roles in the ability to read kitab kuning because they are directly related to understanding sentence structure and changes in Arabic word forms. By applying the K-Means method, the assessment of students' kitab kuning reading proficiency can be conducted more quickly, systematically, and accurately (Yudistira & Andika, 2023). Based on these problems, this study examines the application of the K-Means algorithm to group students according to their kitab kuning reading proficiency using Nahwu and Sharraf scores. The study is expected to provide an alternative solution that helps Islamic boarding schools group students' abilities more objectively. The clustering results can serve as a consideration for teachers or ustaz when determining learning strategies, class placement, and guidance appropriate to each student's level of ability.

METHOD

Data were collected through observation, interviews, and a literature review of journals, books, and other references relevant to this study (Butsianto & Mayangwulan, 2020). Observation was conducted by examining the process used to assess kitab kuning reading proficiency at Al-Is'af Salafiyah Islamic Boarding School. The observations showed that the assessment process was still generally conducted conventionally, creating the possibility of imbalances when grouping students with high, moderate, or low levels of ability.

Interviews were conducted with individuals familiar with the teaching and assessment of kitab kuning reading, including ustaz, ustazah, and the leaders of the Islamic boarding school. An interview is a data-collection stage involving a question-and-answer process between researchers and informants who understand information related to the research object (Dacwanda & Nataliani, 2021). The interview results were used to strengthen the observational data and confirm that problems in assessing kitab kuning reading proficiency actually occurred in the field. In addition, a literature review was conducted to reinforce the theoretical foundation concerning data mining, clustering, the K-Means algorithm, and

relevant previous studies (Syukron et al., 2022).

The data used in this study consisted of student data from Al-Is'af Salafiyah Islamic Boarding School. The attributes used in the clustering process included age, education level, Nahwu score, and Sharraf score. Age was used as numerical data, while education level, Nahwu score, and Sharraf score were first transformed into numerical form so that they could be processed using the K-Means algorithm. These attributes were selected because they are related to students' ability to read kitab kuning. The education attribute consisted of several levels, namely Adna, Imrithi, and Alfiyah. The Nahwu score described a student's ability to understand the grammatical position of words and Arabic sentence structure, whereas the Sharraf score described the student's ability to understand changes in word forms. Both attributes are important indicators for determining students' kitab kuning reading proficiency.

Research Data and Attributes

Before the data were processed using the K-Means algorithm, a data-transformation stage was performed. Data transformation is the process of converting non-numerical data into numerical data so that they can be processed by a clustering algorithm (Saputra, 2021). In this study, the education level, Nahwu score, and Sharraf score attributes were transformed into numerical values. This process was necessary because the K-Means algorithm operates by calculating distances between data points; therefore, all attributes used must be numerical. Following the transformation process, the next stage was data normalization. Normalization was performed to equalize the scale of each attribute so that no particular attribute would dominate the distance-calculation process (Al-Fahmi et al., 2023). This study used the Min-Max Normalization method. This method transforms attribute values into a specified range so that each attribute has a comparable scale. Normalization is important because the K-Means algorithm relies heavily on distance calculations between data points (Najib & Rasiban, 2023).

Data Processing Using K-Means Clustering

The data that had undergone transformation and normalization were then processed using the K-Means Clustering algorithm. Data processing was conducted using Microsoft Excel 2010 and RapidMiner. Microsoft Excel was used to support manual calculations, while RapidMiner was used to test the clustering results automatically (Santoso, 2021). The manual results were compared with the RapidMiner results to determine the consistency of the grouping outcomes. The K-Means algorithm works by determining the number of clusters, selecting initial centroids, calculating the distance of each data point from each centroid, assigning each data point to the nearest cluster, and then updating the centroid values. This process is repeated until the centroid positions become stable or no longer undergo significant changes. Data with similar characteristics are placed in the same cluster, whereas data with different characteristics are placed in different clusters. In this study, the clustering process was designed to form three student-ability groups: adequate, good, and very good. The groups were formed based on similarities in Nahwu and Sharraf scores and were supported by the attributes of age and education level (Setiadi & Sikumbang, 2020). Thus, the clustering results can serve as a basis for helping the Islamic boarding school determine kitab kuning learning strategies that correspond to students' levels of ability.

Research Stages

The research stages began with the collection of student scores, including Nahwu and Sharraf scores as the primary indicators of kitab kuning reading proficiency. The data were obtained through observations and interviews at Al-Is'af Salafiyah Islamic Boarding School. The next stage involved data preprocessing, namely selecting relevant data, converting categorical data into numerical values, and normalizing the data so that they were ready for the clustering process.

Once the data were ready, they were processed using the K-Means Clustering algorithm. At this stage, the number of clusters was determined, initial centroids were selected, and the distances between the data points and the centroids were calculated using Euclidean Distance (Suliman, 2021). Each data point was then assigned according to the shortest

distance. The centroids were subsequently updated by calculating the mean of the members in each cluster. Iterations were performed until the grouping results became stable.

The final stage was the analysis of the grouping results. The clustering results were analyzed to identify the characteristics of each student group. This analysis aimed to ensure that the grouping results corresponded to actual conditions in the field. The final grouping results were then used as a basis for providing recommendations for more effective kitab kuning instruction suited to the ability level of each student (Handayani, 2022).

Data Analysis Technique

The data were analyzed using a quantitative descriptive technique. The clustering results were analyzed based on the number of members in each cluster and the characteristics of students' abilities in each group. The manual calculations performed using Microsoft Excel were compared with the processing results from RapidMiner to assess the consistency of the grouping results. Each cluster was then interpreted according to students' levels of proficiency in Nahwu and Sharraf. The analysis results were used to describe groups of students requiring more intensive guidance, groups with moderate ability, and groups with high ability. Thus, the study not only demonstrates the application of the K-Means algorithm but also provides practical benefits for the Islamic boarding school in developing more appropriately targeted kitab kuning learning strategies.

RESULTS AND ANALYSIS

The data-processing procedure in this study consisted of several stages: data selection, data transformation, data normalization, manual application of the K-Means algorithm using Microsoft Excel 2010, and testing of the clustering results using RapidMiner. These stages were performed to ensure that the data met the requirements of the K-Means algorithm. This is important because K-Means operates based on distance calculations between data points; therefore, the data must be numerical and measured on comparable scales.

The data used in this study consisted of student data from Al-Is'af Salafiyah Islamic Boarding School. The initial dataset included several attributes: student name, age, education level, Nahwu score, and Sharraf score. The attributes used in the clustering process were age, education level, Nahwu score, and Sharraf score. The name attribute was not included in the calculations because it served only as the student's identity. During the data-selection stage, irrelevant data were excluded from the processing procedure so that the clustering results would focus more specifically on students' ability to read kitab kuning.

The next stage was data transformation. Transformation was required because several attributes were still categorical or textual, including education level, Nahwu score, and Sharraf score. To allow processing using the K-Means algorithm, these categorical data were converted into numerical values. The education attribute was transformed according to the Adna, Imrithi, and Alfiah levels. The Nahwu and Sharraf scores were also converted into numerical values corresponding to students' proficiency levels. This transformation enabled all attributes to be calculated using the Euclidean Distance formula. Following the transformation process, the data were normalized using the Min-Max Normalization method. Normalization was performed so that each attribute would have a balanced value scale. Without normalization, attributes with larger values could dominate the distance-calculation process. In this study, the age, education level, Nahwu score, and Sharraf score attributes were normalized so that all data fell within a more uniform range. Consequently, the clustering process could produce more proportional groupings.

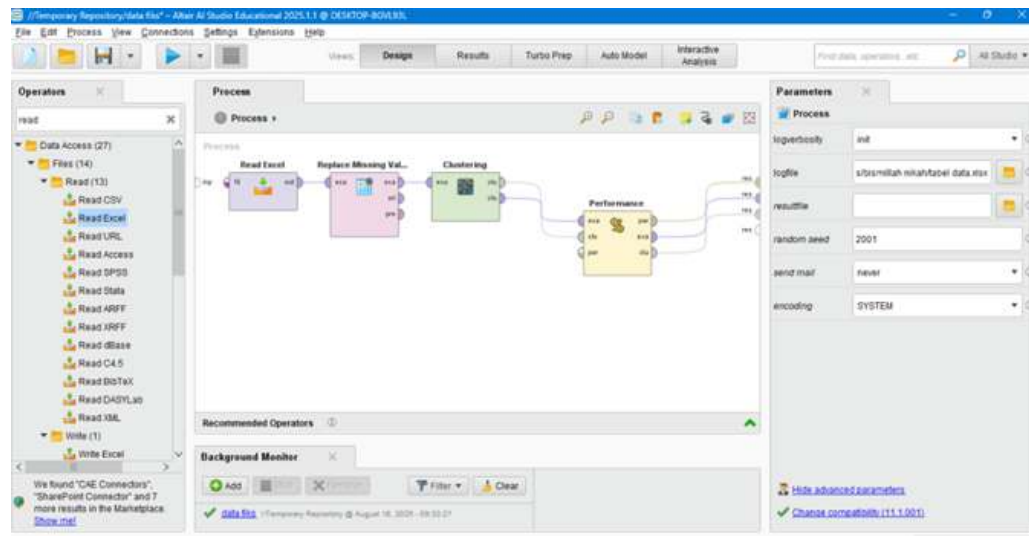
The application of the K-Means algorithm began by determining the number of clusters. This study used three clusters as the basis for grouping students according to their abilities. The clusters were used to represent groups of kitab kuning reading proficiency based on similarities in Nahwu and Sharraf scores. After the number of clusters had been determined, the initial centroids were selected. The initial centroids were chosen from the normalized data, after which the distance of each data point from each centroid was calculated

using Euclidean Distance. The calculation process was repeated through several iterations. In each iteration, the data were assigned to the cluster whose centroid was nearest. After all data had been assigned to their respective clusters, new centroids were calculated based on the mean of the members in each cluster. This process continued until there were no significant changes in cluster membership. In this study, the iteration process reached a stable condition in the sixth iteration. This indicates that the data grouping had converged and no further iterations were required. The final results of the manual calculations using Microsoft Excel showed that the student data were divided into three clusters. Cluster C0 contained 57 students, cluster C1 contained 244 students, and cluster C2 contained 75 students. These results indicate that most students belonged to cluster C1. The final manual grouping results are presented in Table 1.

Table 1. Manual Clustering Results Using Microsoft Excel

Cluster	Number of Students	Percentage
C0	57	15,16%
C1	244	64,89%
C2	75	19,95%
Total	376	100%

Based on Table 1, cluster C1 had the largest number of members, comprising 244 students or 64.89% of the entire dataset. Cluster C0 contained 57 students or 15.16%, while cluster C2 contained 75 students or 19.95%. The differences in the number of members in each cluster indicate that students' abilities were not evenly distributed. Most students had score characteristics that tended to place them in cluster C1, whereas the remaining students were assigned to clusters C0 and C2. After the manual calculations were completed, testing was also conducted using RapidMiner. RapidMiner was used as a tool to compare the manual clustering results with the automated clustering results. In RapidMiner, the prepared data were imported through the Read Excel operator and then processed using the K-Means and Performance operators. After the process was executed, RapidMiner produced three clusters



with membership totals that differed slightly from the manual calculations.

Figure. 1. K-Means Operator Arrangement in RapidMiner

The processing results obtained using RapidMiner showed that cluster C0 contained 56 students, cluster C1 contained 244 students, and cluster C2 contained 76 students. These results indicate a slight difference between the manual calculations using Microsoft Excel and the automated calculations using RapidMiner. The differences occurred primarily in clusters

C0 and C2, whereas the number of members in cluster C1 remained the same at 244 students.

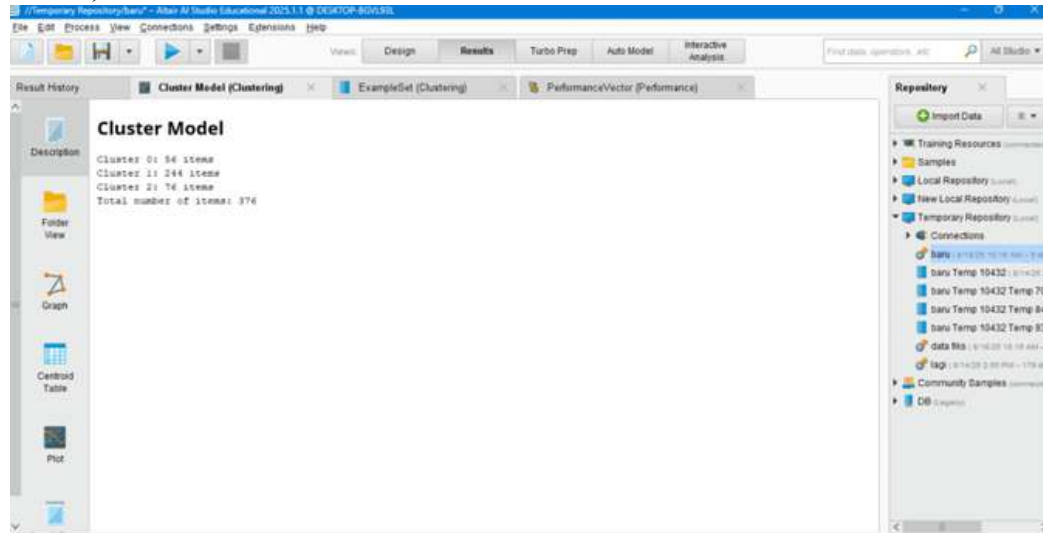


Figure 2. Cluster Model Results in RapidMiner

A comparison of the manual calculation results and the RapidMiner results is presented in Table 2.

Table 2. Comparison of Manual and RapidMiner Clustering Results

Cluster	Microsoft Excel	RapidMiner	Difference
C0	57	56	1
C1	244	244	0
C2	75	76	1
Total	376	376	-

Based on Table 2, the manual clustering and RapidMiner results differed only slightly. The differences occurred only in C0 and C2, with one data record in each case. Meanwhile, C1 had the same number of members in both methods, namely 244 students. This indicates that the manual grouping results obtained using Microsoft Excel were relatively consistent with the testing results obtained using RapidMiner. Therefore, the application of the K-Means algorithm in this study can be considered sufficiently stable for forming groups of students according to their abilities.

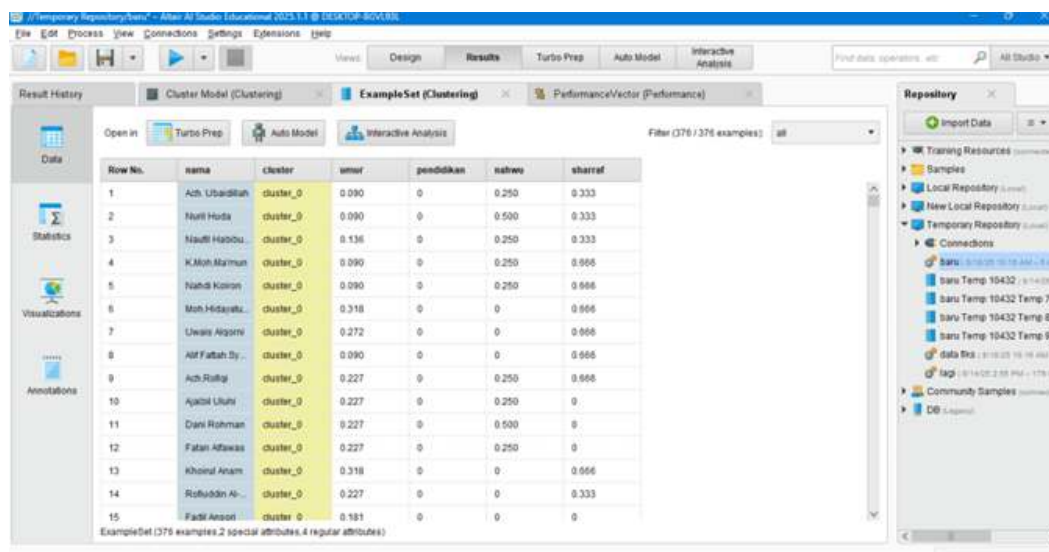


Figure 3. Detailed Data Grouping Results in RapidMiner

RapidMiner also displayed the Average Centroid Distance result. This information can be used to observe the average distance of the data points from the formed centroids. The better the cluster formation, the more closely the data within a cluster will resemble the characteristics of their respective centroid. This result further supports the use of the K-Means algorithm to group student data based on similarities among score attributes.

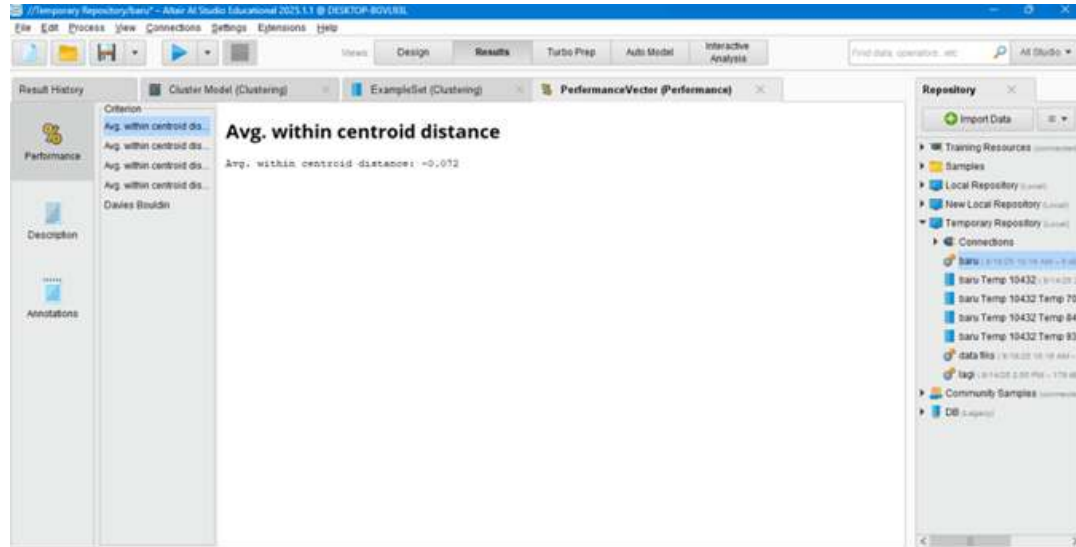


Figure 4. Average Centroid Distance Results in RapidMiner

These grouping results can be interpreted as a basis for understanding students' levels of proficiency in reading kitab kuning. The resulting clusters show differences in students' ability characteristics based on their Nahwu and Sharraf scores. Students in the lower-ability group can be provided with additional guidance or instruction, whereas students with better abilities can be directed toward advanced material. Thus, the clustering results not only provide numerical outputs but can also be used as a consideration in developing kitab kuning learning strategies at the Islamic boarding school. In general, the application of the K-Means algorithm in this study demonstrates that student-score data can be grouped more systematically. Whereas the assessment of kitab kuning reading proficiency previously depended largely on the subjective observations of examiners, the clustering approach can serve as a supporting tool for presenting a more objective overview of students' abilities. The clustering results can assist ustaz or Islamic boarding school administrators in dividing learning groups, developing guidance programs, and evaluating students' abilities based on the available data.

Based on the manual and RapidMiner results, it can be concluded that the K-Means algorithm was able to group the students into three clusters. The difference between the manual calculations and the RapidMiner results was very small, involving only one data record in C0 and one data record in C2. Therefore, these grouping results can be used as a reference to assist Al-Is'af Salafiyah Islamic Boarding School in mapping students' kitab kuning reading proficiency based on Nahwu and Sharraf scores.

CONCLUSIONS

Based on the research results, the K-Means Clustering algorithm can be applied to group students according to their kitab kuning reading proficiency using the Nahwu and Sharraf score attributes. The grouping process was conducted through several stages: data selection, data transformation, data normalization, distance calculation from the centroids, centroid updating, and interpretation of the cluster results. The manual calculations using Microsoft Excel showed that the student data were divided into three clusters: C1 with 57

students, C2 with 244 students, and C3 with 75 students. Meanwhile, processing using RapidMiner produced C1 with 56 students, C2 with 244 students, and C3 with 76 students. The difference between the Microsoft Excel and RapidMiner results was relatively small, involving only one data record in C1 and C3, while C2 had the same number of members. Thus, the K-Means algorithm can be used as an alternative method to help Islamic boarding schools group students' abilities more objectively, systematically, and on the basis of data.

The clustering results can be used by Al-Is'af Salafiyah Islamic Boarding School as a consideration when determining kitab kuning learning strategies. Cluster C1 can be interpreted as a group of students with adequate ability who therefore require more intensive guidance, cluster C2 as a group with good ability, and cluster C3 as a group with very good ability. This grouping can assist ustaz or Islamic boarding school administrators in designing instruction that is more appropriate to students' levels of ability. Future studies are advised to include additional variables, such as memorization scores, learning participation, length of study at the Islamic boarding school, or kitab examination results, so that the grouping results become more accurate. In addition, future research could compare K-Means with other clustering methods and develop a web-based or mobile application so that student grouping can be performed automatically and more practically.

REFERENCES

- G. W. Saraswati, M. S. Rohman, N. A. S. Winarsih and M. A. Larasati, "Penerapan Metode K-Means Clustering Dalam Menganalisis Sentimen Masyarakat Terhadap K-Poppers pada Twitter," *Jurnal Ilmiah Komputer*, vol. XVIII, pp. 201-210, 2022. <https://doi.org/10.35889/progresif.v18i2.877>
- F. A. Falaah and dkk, "Pendampingan Metode Pembelajaran Baca Kitab Salaf Untuk Santri Pondok Pesantren Sunan Drajat Kedungsantren Campurejo Bojonegoro," *Student engagement*, vol. 1, p. 41, 2022. <https://doi.org/10.55352/santri.v1i1.193>
- K. Irfansyah and dkk, "Implementasi Algoritma Clustering K-Means Pada Pengguna Wartel Di Pondok Pesantren Salafiyah Syafi'iyah Sukorejo," *ilmiah multi disiplin ilmu*, pp. 81-86, 2024. <https://doi.org/10.69714/55xet429>
- Sukron, M., Supriadi, A., & Sulton, R. (2021). Optimasi Metode Naïve Bayes Menggunakan Algoritma Particle Swarm Optimization (Pso) Untuk Prediksi Penyakit Diabetes Mellitus. *COREAI: Jurnal Kecerdasan Buatan, Komputasi dan Teknologi Informasi*, 2(2), 18-24. <https://doi.org/10.33650/coreai.v2i2.3304>
- Y. Nataliani and D. O. Dacwanda, "Implementasi K-Means Clustering untuk Analisis Nilai Akademik Siswa Berdasarkan Nilai Pengetahuan dan Keterampilan," *Jurnal Teknologi Informasi*, vol. XXVIII, pp. 125-138, 2021. <https://doi.org/10.24246/aiti.v18i2.125-138>
- H. Syukron and dkk, "Perbandingan K-Means K-Medoids dan Fuzzy C-Means untuk Pengelompokan Data Pelanggan dengan <https://journal.irpi.or.id/index.php/malcom>, vol. II, pp. 76-83, 2022. Model LRFM," <https://doi.org/10.57152/malcom.v2i2.442>
- Suliman, "Implementasi Data Mining Terhadap Prestasi Belajar Mahasiswa Berdasarkan Pergaulan Dan Sosial Ekonomi Dengan Algoritma K-Means Clustering," *Sistem Informasi dan Sistem Komputer*, vol. VI, p. 1, 2021. <https://doi.org/10.51717/simkom.v6i1.48>
- N. Kurniawati, M. A. Fahtoni and D. T. Yuliana, "Penentuan Penerima Kartu Indonesia Pintar KIP Kuliah dengan Menggunakan Metode K-Means Clustering," *Focus ACTion Of Research Mathematic*, vol. V, p. 129, 2022.

- N. N. Fransiska R, D. S. Anggraeni and U. Enri, "Pengelompokan Data Kemiskinan Provinsi Jawa Barat Menggunakan Algoritma K-Means dengan Silhouett Coefficiente," <https://jurnal.plb.ac.id/index.php/tematik/index>, vol. V, pp. 29-33, 2023. 79
- A. Sulistiyawati and E. Supriyanto, "Implementasi Algoritma K-means Clustering dalam Penentuan Siswa kelas Unggulan," Jurnal TEKNO KOMPAK, vol. 15, pp. 25-36, 2021. <https://doi.org/10.33365/jtk.v15i2.1162>
- S. Romelah, P. H. Nopita and D. Syaputri, "Implemnetation of K-Means Algorithm for Economic Distribution Clustering Base on Demographics of Population," Indonesian Journal of Machine Learning and Computer Science, vol. I, pp. 2-6, 2021.
- F. Handayani, "Aplikasi Aplikasi Data Mining Menggunakan Algoritma K-Means Clustering untuk Mengelompokan Mahasiswa Berdasarkan Gaya Belajar," Jurnal Teknologi dan Informasi (JATI) , vol. XII, p. 48, 2022. <https://doi.org/10.34010/jati.v12i1.6733>
- A. Setiadi and E. D. Sikumbang, "K-Means Clustering Dalam Penerimaan Karyawan Baru," Informatics For Educators And Professionals, vol. IV, p. 104, 2020. <https://doi.org/10.51211/itbi.v4i2.1304>
- R. Andika and A. Yudhistira, "Pengelompokan Data Nilai Siswa Menggukan Metode KMeans Clustering," Journal of Artificial Intelligence and Technology Information (JAITI), vol. I, pp. 20-28, 2023. <https://doi.org/10.58602/jaiti.v1i1.22>
- N. A. Maori, "Metode Elbow Dalam Optimasi Jumlah Cluster Pada Kmeans Clustering," Jurnal SIMETRIS,, vol. XIV, p. 28, 2023. <https://doi.org/10.24176/simet.v14i2.9630>
- Z. Abidin , Z. Nabila, A. R. Isnain and Permata, "Analisis Data Mining Untuk Clustering Kasus Covid-19 Di Provinsi Lampung Dengan Algoritma K-Means," Jurnal Teknologi dan Sistem Informasi (JTSI), vol. II, pp. 100-108, 2021.
- H. M. Santoso, "Application of Association Rule Method Using Apriori Algorithm to Find Sales Patterns Case Study of Indomaret Tanjung Anom," Brilliance, vol. I, p. 58, 2021. <https://doi.org/10.47709/brilliance.v1i2.1228>
- B. Sufajar and N. T. Mayangwulan, "Penerapan Data Mining Untuk Prediksi Penjualan Mobil Menggunakan Metode K-Means Clustering," Nasional Komputasi dan Teknologi Informasi, vol. 3, p. 187, 2020. <https://doi.org/10.32672/jnkti.v3i3.2428>
- M. B. Al-Fahmi and dkk, "Penerapan K-Means Clustering pada pariwisata Kabupaten Bojonegoro untuk Mrndukung Strategi Pemasaran," Jurnal Nasional Teknologi dan Sistem Informasi, vol. 09, p. 1, 2023. <https://doi.org/10.25077/TEKNOSI.v9i2.2023.141-149>
- M. A. Najib and Rasiban, "Penerapan Algoritma K-Means untuk Analisis Nilai Tahfidz Santri Berdasarkan Nilai Fashohah dan Tajwid (studi kasus: Pesantren Modern At-Taqwa Bogor)," Jurnal Indonesia: Manajemen Informatika dan KOMunikasi, vol. 4, pp. 1670-1672, 2023. 80 <https://doi.org/10.35870/jimik.v4i3.392>
- P. S. Saputra and dkk, "Perbandingan Algoritma Fuzzy C-Means dan Algoritma Naive Bayes dalam Menentukan Keluarga Penerima Manfaat (KPM) Berdasarkan status Sosial Ekonomi (SSE) Terendah," Jurnal Sain dan Teknologi, vol. X, pp. 1-7, 2021. <https://doi.org/10.23887/jstundiksha.v10i1.23340>
- F. Novianti, Y. R. Yasmin and D. C. Novitasari, "Penerepan algoritma Fuzzy C-Maen (FCM) dalam Pengelompokan Provinsi di Indonesia berdasarkan Indikator Penyakit Menular manusia," Jumanji, vol. VI, pp. 23-33, 2022. <https://doi.org/10.26874/jumanji.v6i1.103>