

Decision Support System for Determining the Eligibility of Economically Disadvantaged Students for Assistance Using the K-Means and MOORA Methods

Gilbert Johan Martin Sinaga
Accounting
Universitas Riau
Indonesia
Gilbertsinaga0303@gmail.com

Abstract— The determination of recipients for student financial aid often faces challenges related to subjectivity in the selection process, necessitating a system capable of conducting objective analysis. This study develops a Decision Support System using the K-Means method to cluster students based on similar socioeconomic characteristics and the MOORA method to rank aid recipients more accurately. The K-Means method is applied to classify students into three clusters based on parental income, number of dependents, and academic performance. The clustering results indicate that students in Cluster 1 belong to the lowest economic group, making them the top priority in the selection process. Subsequently, the MOORA method is used to rank students within Cluster 1 based on an optimal value calculated from the weighted benefit and cost criteria. This calculation produces a priority ranking that is more transparent and objective compared to conventional selection systems. The findings show that the combination of K-Means and MOORA methods enhances accuracy in selecting aid recipients while reducing subjectivity in the selection process. With this system, schools or relevant institutions can expedite decision-making and ensure that aid is distributed to the students most in need. This study is expected to serve as a solution for educational institutions in improving the effectiveness and efficiency of student welfare programs.

Keywords— Decision Support System, K-Means, MOORA, Eligibility for Aid for Economically Disadvantaged Students

Article Info : Date Submitted: 2025-03-25| Date Revised: 2025-03-26| Date Accepted:2025-04-19

I. INTRODUCTION

Poverty remains one of the major social issues faced by many countries, including Indonesia. Its impact is not limited to economic aspects but also affects the quality of human resources. The inability to meet basic needs, such as food, clothing, education, and healthcare, is one of the tangible consequences of poverty [1]. Education, as a key pillar in improving human resource quality, is often hindered by economic constraints [2]. However, a nation's progress heavily depends on the mastery of science and technology, which can only be achieved through equitable and high-quality education [3], [4].

To address economic barriers to education access, the Indonesian government has launched the Poor Student Assistance program. This program aims to provide financial aid to students from underprivileged families, enabling them to continue their education. In addition to assisting students facing financial difficulties, the program is also expected to encourage dropouts to return to school. However, the effectiveness of this program largely depends on the accuracy of recipient selection. Without a transparent and accurate system, the program risks being misallocated [5].

In reality, reports from the Ministry of Education, Culture, Research, and Technology through the Center for Education Financing Services indicate that many cases of misallocation still occur in the distribution of Poor Student Assistance funds. Some major factors contributing to this issue include individuals deliberately claiming aid that rightfully belongs to others in greater need, falsification of the Certificate of Economic Hardship by neighborhood officials, and a lack of transparency and accountability in the verification process at the school level [6]. Furthermore, the government has often been slow in addressing these problems, leading to suboptimal aid distribution. Therefore, a more objective, transparent, and efficient system is needed to determine aid recipients, ensuring that the Poor Student Assistance program operates more effectively [6], [7].

To address these challenges, this study aims to develop a Decision Support System that can assist in selecting Poor Student Assistance recipients with greater accuracy and fairness. A Decision Support System is a computer-based system designed to help decision-makers solve problems by generating the best possible alternatives based on available data. In this study, a combination of the K-Means and MOORA (Multi-Objective Optimization on the Basis of Ratio Analysis) methods is employed to improve the accuracy of recipient selection.

The K-Means method is used to cluster aid applicants based on two primary criteria: family dependents and parental income. This method is effective for clustering data with high accuracy, as it groups data into more homogeneous categories based on proximity values [8]. Once the eligible recipient group is formed, the next step involves ranking using the MOORA method. MOORA is chosen for its ability to handle multi-criteria decision-making problems [9]–[11]. In this study, an additional criterion, student academic performance, is incorporated into MOORA as a supporting factor in prioritizing aid recipients. This approach enables the system to provide more objective recommendations in selecting Poor Student Assistance recipients [7], [12].

Several previous studies have implemented the MOORA method in Decision Support Systems for scholarship selection. A study by E. Nahak [13] developed a web-based system to assist administrative staff in selecting recipients of the Indonesia Smart Program (PIP) scholarship, ensuring aid is allocated accurately. The findings indicated that a web-based system enhances efficiency and accuracy in recipient selection. Another study by Renny Puspita Sari [14] applied the MOORA method in a Decision Support System for Bidikmisi scholarship selection, demonstrating that MOORA provides a more systematic and accurate approach to scholarship recipient determination.

However, this study introduces significant differences compared to previous research. Here, the K-Means method is utilized in the initial selection stage to cluster applicants based on objective economic criteria before ranking them using the MOORA method. This approach offers the advantage of reducing bias in the preliminary selection and enhancing the efficiency of the selection process. Moreover, this study specifically focuses on the application of the combined K-Means and MOORA methods in determining Poor Student Assistance recipients—an area that has not been extensively explored in previous research. With this novel approach, the developed system is expected to provide a more accurate, fair, and transparent solution for the distribution of student financial aid, ensuring that assistance reaches those who need it most.

II. RELATED WORKS

A. K-Means Clustering

The K-Means clustering method is employed in this study to group student data based on economic and academic characteristics, allowing for a more objective selection of Poor Student Assistance recipients. K-Means is a non-hierarchical algorithm that partitions data into a predetermined number of clusters (k) based on similarity in

characteristics [15], [16]. This algorithm is chosen because it is faster in clustering data compared to hierarchical methods, which require a structured hierarchy [17], [18].

B. Multi-Objective Optimization on the Basis of Ratio Analysis (MOORA)

The Multi-Objective Optimization on the Basis of Ratio Analysis (MOORA) method is applied in this study to optimize multiple conflicting attributes within a defined constraint. This method involves assigning weights to each attribute and performing a ranking process to select the best alternative for decision-making [19]. MOORA is advantageous due to its simplicity, flexibility, and high selectivity, making it widely used in various fields, including the selection of poor student aid recipients [20].

III. METHOD

This study aims to develop a Decision Support System to objectively determine the eligibility of Poor Student Assistance recipients at SMP Negeri 6 Pematangsiantar. The methodology used in this research combines the K-Means method for the clustering process and the MOORA method for ranking the aid recipients.

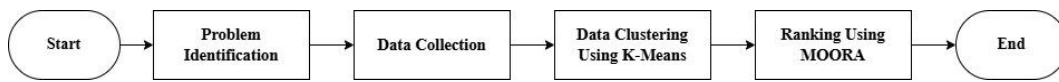


Figure 1. Research Framework

A. Problem Identification

At this stage, an analysis is conducted to identify issues in determining recipients of Poor Student Assistance. The primary problem is the misallocation of aid due to a lack of transparency, subjectivity in the selection process, and potential data manipulation. Additionally, the selection process for assistance recipients at SMP Negeri 6 Pematangsiantar is still carried out manually, making it prone to errors. The problem identification is based on interviews with school representatives, literature reviews, and data related to the assistance program.

B. Data Collection

The data collection phase involves gathering information from various relevant sources, primarily student data from SMP Negeri 6 Pematangsiantar who applied for Poor Student Assistance. The collected data includes economic and academic factors, which play a crucial role in the selection process. The economic factors analyzed include parental income and family dependents, which serve as primary indicators in determining students' level of need. Meanwhile, academic factors are assessed based on students' report card grades, which are considered in the ranking process for aid recipients.

C. Data Clustering Using K-Means

The K-Means clustering process in this study involves the following steps:

- Determining the number of clusters (k).
The number of clusters is set to three ($k=3$), consisting of groups categorized as highly in need of assistance, in need of assistance, and less in need of assistance.
- Initializing initial centroids randomly.
The initial centroids (V) are selected randomly from the available dataset, using the following equation:

$$V = \sum_{i=1}^n x_i \quad (1)$$

- Calculating the Distance Between Each Data Point and the Centroid Using Euclidean Distance.

This algorithm measures the proximity of data points to the cluster center using the following equation:

$$d(x_i, c_j) = \sqrt{\sum_{j=1}^n (x_j, c_j)^2} \quad (2)$$

- d. Assigning Each Data Point to the Nearest Cluster.
Once the distances are calculated, each data point is assigned to the cluster with the closest centroid.
- e. Recalculating the New Centroid Positions (k).
After all data points have been assigned to clusters, the centroid positions are recalculated for each cluster.
- f. Iterating Until No Further Centroid Changes Occur
If the new centroid positions continue to change, the iteration is repeated from step (c) until the centroids remain stable.

D. Ranking Using MOORA

The steps in implementing the MOORA method are as follows:

- a. Determining the Decision Matrix Values.
The first step involves constructing a decision matrix based on the clustering results obtained from the K-Means method. This matrix represents the values of each alternative (student) based on predefined criteria, such as economic and academic factors.

$$X = \begin{bmatrix} x_{11} & x_{12} & \cdots & x_{1n} \\ x_{21} & x_{22} & \cdots & x_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ x_{m1} & x_{m2} & \cdots & x_{m3} \end{bmatrix} \quad (3)$$

- b. Normalization of the Matrix.
After constructing the decision matrix, the next step is normalization to ensure that all values are on a comparable scale. Normalization is performed using the following formula:

$$X_{ij}^* = \frac{X_{ij}}{\sqrt{\sum_{i=1}^m x_{ij}^2}} \quad (4)$$

- c. Optimizing Attributes for Multi-Objective Analysis.
In this stage, each normalized value is multiplied by its predefined weight. The calculation is done by summing the values of benefit criteria and subtracting the values of cost criteria.

$$Y_i = \sum_{j=1}^g W_j \times x_{ij}^* - \sum_{j=g+1}^n W_j \times x_{ij}^* \quad (5)$$

- d. Ranking Alternatives.
The final result from the previous calculations is used to rank the alternatives. Students with the highest Y^* score are given the highest priority for receiving aid, as they best meet the economic and academic selection criteria.

IV. RESULT AND DISCUSSION

A. Data Collection

The data utilized in this study were obtained from multiple sources, primarily from students of SMP Negeri 6 Pematangsiantar who applied for the Poor Student Assistance program. For analysis purposes, the study utilized a sample of 30 students. To determine the aid recipients, three main criteria were used:

- a. Parent's Income, serving as an indicator of the family's economic condition.

- b. Number of Dependents, reflecting the financial burden within the household.
- c. Report Card Grades, an academic factor supporting the selection process.

Table 1. Overall Data

ID Alternative (Student)	Dependents	Income (per month) (Rp)	Report Card Score
A-01	3	1.000.000	75
A-02	4	1.500.000	80
A-03	4	1.000.000	90
A-04	5	1.500.000	78
A-05	3	500.000	88
A-06	5	1.000.000	65
A-07	6	1.000.000	79
A-08	3	700.000	80
A-09	5	1.000.000	78
A-10	4	1.200.000	85
A-11	3	1.000.000	75
A-12	2	500.000	90
A-13	4	1.000.000	70
A-14	5	700.000	80
A-15	3	1.300.000	78
A-16	5	1.000.000	76
A-17	4	900.000	88
A-18	3	1.500.000	90
A-19	3	1.200.000	90
A-20	4	1.000.000	88
A-21	2	1.500.000	9
A-22	5	1.500.000	79
A-23	4	1.000.000	70
A-24	3	1.000.000	65
A-25	3	1.400.000	70
A-26	4	700.000	70
A-27	5	750.000	69
A-28	6	1.000.000	78
A-29	5	1.000.000	70
A-30	4	1.500.000	76

B. Processing Data Using K-Means

In the K-Means calculation, two main attributes are used in the clustering process. First, Dependents, categorized as X (the number of family members financially supported by the parents). Second, Parental Income, categorized as Y (the total monthly income of the parents).

Table 2. K-Means Calculation Data

ID Alternative (Student)	Dependents	Income (per month) (Rp)
A-01	3	1.000.000
A-02	4	1.500.000
A-03	4	1.000.000
A-04	5	1.500.000
A-05	3	500.000
A-06	5	1.000.000
A-07	6	1.000.000
A-08	3	700.000
A-09	5	1.000.000
A-10	4	1.200.000
A-11	3	1.000.000
A-12	2	500.000
A-13	4	1.000.000
A-14	5	700.000
A-15	3	1.300.000
A-16	5	1.000.000
A-17	4	900.000

A-18	3	1.500.000
A-19	3	1.200.000
A-20	4	1.000.000
A-21	2	1.500.000
A-22	5	1.500.000
A-23	4	1.000.000
A-24	3	1.000.000
A-25	3	1.400.000
A-26	4	700.000
A-27	5	750.000
A-28	6	1.000.000
A-29	5	1.000.000
A-30	4	1.500.000

C. Determining Initial Cluster Centers

Based on the data in Table 1 and Table 2, the clustering process is conducted by dividing the data into three clusters ($k=3$). The initial centroid points, selected randomly, are determined as follows:

- Data point 5 as the center of Cluster 1
- Data point 20 as the center of Cluster 2
- Data point 30 as the center of Cluster 3

Table 3. Cluster Center Points

Cluster Center	Dependents	Income (per month) (Rp)
C1	3	500.000
C2	4	1.000.000
C3	4	1.500.000

D. Iteration 1 Calculation Process

The calculation process begins with Iteration 1, where the distance between each data point and all centroids is computed using the Euclidean Distance formula.

$$d(x_i, c_j) = \sqrt{(X - c_x)^2 + (Y - c_y)^2} \quad (6)$$

This iteration aims to form the initial clusters by assigning each data point to the cluster with the nearest centroid. To determine the nearest distance between each data point and the centroid, the Euclidean Distance method is used as follows:

- If $C1 < C2$ and $C1 < C3$, then the data point is assigned to Cluster 1.
- If $C2 < C1$ and $C2 < C3$, then the data point is assigned to Cluster 2.
- If $C3 < C1$ and $C3 < C2$, then the data point is assigned to Cluster 3.

The calculation results for selected data points are presented in Table 4. As shown in Table 4, data point 1 has the shortest distance to centroid 2, so it is grouped into Cluster 2. Using the same process, the nearest distance for all data points can be determined. This distance serves as the basis for classifying each data point into Cluster 1, Cluster 2, or Cluster 3.

Table 4. Cluster Determination for All Data in Iteration 1

ID Alternative	Dependents	Income (per month) (Rp)	$D(x_1, c_1)$	$D(x_1, c_2)$	$D(x_1, c_3)$	Cluster Grouping
A-01	3	1.000.000	500.000	1	500.000	2
A-02	4	1.500.000	1.000.000	500.000	0	3
A-03	4	1.000.000	500.000	0	500.000	2
A-04	5	1.500.000	1.000.000	500.000	1	3
A-05	3	500.000	0	500.000	1.000.000	1
A-06	5	1.000.000	500.000	1	500.000	2
A-07	6	1.000.000	500.000	2	500.000	2
A-08	3	700.000	200.000	300.000	800.000	1
A-09	5	1.000.000	500.000	1	500.000	2

A-10	4	1.200.000	700.000	200.000	300.000	2
A-11	3	1.000.000	500.000	1	500.000	2
A-12	2	500.000	1	500.000	100.000	1
A-13	4	1.000.000	500.000	0	500.000	2
A-14	5	700.000	200.000	300.000	800.000	1
A-15	3	1.300.000	800.000	300.000	200.000	3
A-16	5	1.000.000	500.000	1	500.000	2
A-17	4	900.000	400.000	100.000	600.000	2
A-18	3	1.500.000	1.000.000	500.000	1	3
A-19	3	1.200.000	700.000	200.000	300.000	2
A-20	4	1.000.000	500.000	0	500.000	2
A-21	2	1.500.000	1.000.000	500.000	2	3
A-22	5	1.500.000	1.000.000	500.000	1	3
A-23	4	1.000.000	500.000	0	500.000	2
A-24	3	1.000.000	500.000	1	500.000	2
A-25	3	1.400.000	900.000	400.000	100.000	3
A-26	4	700.000	200.000	300.000	800.000	1
A-27	5	750.000	250.000	250.000	750.000	2
A-28	6	1.000.000	500.000	2	500.000	2
A-29	5	1.000.000	500.000	1	500.000	2
A-30	4	1.500.000	1.000.000	500.000	0	3

After ensuring that all data points (from data 1 to data 30) have been assigned to a cluster, the next step is to determine the new centroid based on the data grouped within each cluster. As shown in Table 5, this process begins with Cluster 1, which consists of 5 data points. The new centroid for this cluster is calculated by taking the average value of the grouped data, resulting in a new centroid value of 3.4 for the number of dependents and Rp620,000 for parental income.

Table 5. Determination of New Centroid for Cluster 1

ID Alternative	Dependents	Income (per month) (Rp)
A-05	3	500.000
A-08	3	700.000
A-12	2	500.000
A-14	5	700.000
A-26	4	700.000
Total	17	3.100.000
Average	3,4	620.000

In Cluster 2, there are 17 data points grouped based on the initial calculation results. After recalculating, the new centroid is obtained by averaging the number of dependents and parental income within the cluster. The calculation results show that the new centroid value for the number of dependents is 4.29, while the new centroid value for parental income is Rp1,002,778.

Table 6. Determination of New Centroid for Cluster 2

ID Alternative	Dependents	Income (per month) (Rp)
A-01	3	1.000.000
A-03	4	1.000.000
A-06	5	1.000.000
A-07	6	1.000.000
A-09	5	1.000.000
A-10	4	1.200.000
A-11	3	1.000.000
A-13	4	1.000.000
A-16	5	1.000.000
A-17	4	900.000
A-19	3	1.200.000
A-20	4	1.000.000
A-23	4	1.000.000
A-24	3	1.000.000
A-27	5	750.000
A-28	6	1.000.000

A-29	5	1.000.000
Total	73	17.050.000
Average	4,294117647	1.002.777,778

In Cluster 3, there are 8 data points grouped after the initial clustering process. After recalculating, the new centroid values are obtained by averaging the number of dependents and parental income within this cluster. The calculation results show that the new centroid for the number of dependents is 3.63 (rounded from 3.625), while the new centroid for parental income is Rp1,462,500.

Table 7. Determination of New Centroid for Cluster 3

ID Alternative	Dependents	Income (per month) (Rp)
A-02	4	1.500.000
A-04	5	1.500.000
A-15	3	1.300.000
A-18	3	1.500.000
A-21	2	1.500.000
A-22	5	1.500.000
A-25	3	1.400.000
A-30	4	1.500.000
Total	26	11.700.000
Average	3,625	1.462.500

The average obtained from each cluster after the clustering process represents the new centroid. This new centroid value will be used in Iteration 2 to recalculate the distance of each data point from the updated centroid. These centroid values will be utilized in the next iterations to evaluate whether there are still changes in cluster assignments or if the process has reached a stable condition (convergence).

Table 8. New Centroid in Iteration 1

Cluster 1	3,4	620.000
Cluster 2	4,294117647	1.002.777,778
Cluster 3	3,625	1.462.500

E. Iteration 2 Calculation Process

In Iteration 2, the calculation process is repeated by measuring the distance between each data point and the new centroids obtained from the previous iteration. This calculation aims to determine whether any changes occur in the cluster assignments. Each data point is compared with the new centroid of each cluster, and the data is reassigned to the cluster with the closest distance.

Table 9. Cluster Determination for All Data in Iteration 2

ID Alternative	$D(P, C_1)$	$D(P, C_2)$	$D(P, C_3)$	Cluster Grouping
A-01	380.000	2.777,778079	462.500	2
A-02	880.000	497.222,2222	37.500	3
A-03	380.000	2.777,777793	462.500	2
A-04	880.000	497.222,2222	37.500	3
A-05	120.000	502.777,7778	962.500	1
A-06	380.000	2.777,777867	462.500	2
A-07	380.000	2.777,778302	462.500	2
A-08	80.000	302.777,7778	762.500	1
A-09	380.000	2.777,777867	462.500	2
A-10	580.000	197.222,2222	262.500	2
A-11	380.000	2.777,778079	462.500	2
A-12	120.000	502.777,7778	962.500	1
A-13	380.000	2.777,777793	462.500	2
A-14	80.000	302.777,7778	762.500	1
A-15	680.000	297.222,2222	162.500	3
A-16	380.000	2.777,777867	462.500	2
A-17	280.000	102.777,7778	562.500	2
A-18	880.000	497.222,2222	37.500	3
A-19	580.000	197.222,2222	262.500	2
A-20	380.000	2.777,777793	462.500	2
A-21	880.000	497.222,2222	37.500	3

A-22	880.000	497.222,2222	37.500	3
A-23	380.000	2.777,777793	462.500	2
A-24	380.000	2.777,778079	462.500	2
A-25	780.000	397.222,2222	62.500	3
A-26	80.000	302.777,7778	762.500	1
A-27	130.000	252.777,7778	712.500	1
A-28	380.000	2.777,778302	462.500	2
A-29	380.000	2.777,777867	462.500	2
A-30	880.000	497.222,2222	37.500	3

After ensuring that all data points from data 1 to data 30 have been assigned to a cluster, the next step is to determine the new centroids based on the data grouped within each cluster.

Table 10. Determination of New Centroid for Cluster 1

ID Alternative	Dependents	Income (per month) (Rp)
A-05	3	500.000
A-08	3	700.000
A-12	2	500.000
A-14	5	700.000
A-26	4	700.000
A-27	5	750.000
Total	22	3.850.000
Average	3,666	641.666,666

Table 11. Determination of New Centroid for Cluster 2

ID Alternative	Dependents	Income (per month) (Rp)
A-01	3	1.000.000
A-03	4	1.000.000
A-06	5	1.000.000
A-07	6	1.000.000
A-09	5	1.000.000
A-10	4	1.200.000
A-11	3	1.000.000
A-13	4	1.000.000
A-16	5	1.000.000
A-17	4	900.000
A-19	3	1.200.000
A-20	4	1.000.000
A-23	4	1.000.000
A-24	3	1.000.000
A-28	6	1.000.000
A-29	5	1.000.000
Total	68	16.300.000
Average	4,25	1.018.750

Table 12. Determination of New Centroid for Cluster 3

ID Alternative	Dependents	Income (per month) (Rp)
A-02	4	1.500.000
A-04	5	1.500.000
A-15	3	1.300.000
A-18	3	1.500.000
A-21	2	1.500.000
A-22	5	1.500.000
A-25	3	1.400.000
A-30	4	1.500.000
Total	26	11.700.000
Average	3,625	1462500

The averages obtained from the three clusters serve as the new centroids, which are then used in Iteration 3 to recalculate the distance of each data point to the updated centroids.

Table 13. New Centroid in Iteration 2

Cluster 1	3,666666667	641.666,6667
Cluster 2	4,25	1.018.750
Cluster 3	3,625	1.462.500

F. Iteration 3 Calculation Process

The calculation process is repeated in Iteration 3 because the centroids still changed in the previous iteration. Therefore, a recalculation is necessary to ensure that each data point is assigned to the most appropriate cluster.

Table 14. Cluster Determination for All Data in Iteration 3

ID Alternative	$D(P, C_1)$	$D(P, C_2)$	$D(P, C_3)$	Cluster Grouping
A-01	358,333	18.750	462.500	2
A-02	858,333	481.250	37.500	3
A-03	358,333	18.750	462.500	2
A-04	858,333	481.250	37.500	3
A-05	141,667	518.750	962.500	1
A-06	358,333	18.750	462.500	2
A-07	358,333	18.750	462.500	2
A-08	58,333	318.750	762.500	1
A-09	358,333	18.750	462.500	2
A-10	558,333	181.250	262.500	2
A-11	358,333	18.750	462.500	2
A-12	141,667	518.750	962.500	1
A-13	358,333	18.750	462.500	2
A-14	58,333	318.750	762.500	1
A-15	658,333	281.250	162.500	3
A-16	358,333	18.750	462.500	2
A-17	258,333	118.750	562.500	2
A-18	858,333	481.250	37.500	3
A-19	558,333	181.250	262.500	2
A-20	358,333	18.750	462.500	2
A-21	858,333	481.250	37.500	3
A-22	858,333	481.250	37.500	3
A-23	358,333	18.750	462.500	2
A-24	358,333	18.750	462.500	2
A-25	758,333	381.250	62.500	3
A-26	58,333	318.750	762.500	1
A-27	108,333	268.750	712.500	1
A-28	358,333	18.750	462.500	2
A-29	358,333	18.750	462.500	2
A-30	858,333	481.250	37.500	3

In Iteration 3, the nearest distance calculation is performed again to ensure that each data point is assigned to the most appropriate cluster based on the updated centroids from Iteration 2.

Table 15. Determination of New Centroid for Cluster 1

ID Alternative	Dependents	Income (per month) (Rp)
A-05	3	500.000
A-08	3	700.000
A-12	2	500.000
A-14	5	700.000
A-26	4	700.000
A-27	5	750.000
Total	22	3.850.000

Average	3,666	641.666,666
---------	-------	-------------

Table 16. Determination of New Centroid for Cluster 2

ID Alternative	Dependents	Income (per month) (Rp)
A-01	3	1.000.000
A-03	4	1.000.000
A-06	5	1.000.000
A-07	6	1.000.000
A-09	5	1.000.000
A-10	4	1.200.000
A-11	3	1.000.000
A-13	4	1.000.000
A-16	5	1.000.000
A-17	4	900.000
A-19	3	1.200.000
A-20	4	1.000.000
A-23	4	1.000.000
A-24	3	1.000.000
A-28	6	1.000.000
A-29	5	1.000.000
Total	68	16.300.000
Average	4,25	1.018.750

Table 17. Determination of New Centroid for Cluster 3

ID Alternative	Dependents	Income (per month) (Rp)
A-02	4	1.500.000
A-04	5	1.500.000
A-15	3	1.300.000
A-18	3	1.500.000
A-21	2	1.500.000
A-22	5	1.500.000
A-25	3	1.400.000
A-30	4	1.500.000
Total	26	11.700.000
Average	3,625	1.462.500

After calculating the average to determine the new centroids, the results show that the new centroids in Iteration 3 are the same as those in Iteration 2. Thus, the iteration process is stopped because there are no further changes in cluster assignments, indicating that the algorithm has reached convergence.

Table 18. New Centroid in Iteration 3

Cluster 1	3,666666667	641.666,6667
Cluster 2	4,25	1.018.750
Cluster 3	3,625	1.462.500

G. MOORA Method Calculation

In the MOORA method, the data used is derived from the K-Means clustering results, where Cluster 1 consists of students with the lowest parental income, thus having the highest priority compared to Cluster 2 and Cluster 3.

1. Defining Objectives and Identifying Evaluation Attributes

The weight of each criterion is determined based on its importance in the selection process. MOORA considers two types of criteria:

- Benefit (Max): A criterion where a higher value is more beneficial for the recipient.
- Cost (Min): A criterion where a lower value is more beneficial for the recipient.

Table 19. K-Means Calculation Results in Cluster 1

ID Alternative (Student)	Income (per month) (Rp)	Dependents	Report Card Score
A-05	500.000	3	88
A-08	700.000	3	80

A-12	500.000	2	90
A-14	700.000	5	80
A-26	700.000	4	70
A-27	750.000	5	69

Tabel 20. Alternative

ID Alternative	C1	C2	C3
A-05	500.000	3	88
A-08	700.000	3	80
A-12	500.000	2	90
A-14	700.000	5	80
A-26	700.000	4	70
A-27	750.000	5	69

Tabel 21. Criteria Weights

Criteria	Description	Weight	Type	Explanation
C1	Parental Income	4,5	Cost (Min)	The lower the value, the higher the eligibility.
C2	Dependents	3,5	Benefit (Max)	The greater the number, the higher the eligibility.
C3	Report Card Grades	2,0	Benefit (Max)	The higher the value, the higher the eligibility.

2. Creating a Decision Matrix

The decision matrix contains the values of each alternative (student) based on the predetermined criteria.

$$X_{ij} = \begin{bmatrix} 500.000 & 3 & 88 \\ 700.000 & 3 & 80 \\ 500.000 & 2 & 90 \\ 700.000 & 5 & 80 \\ 700.000 & 4 & 70 \\ 750.000 & 5 & 69 \end{bmatrix}$$

3. Calculating the Normalized Matrix

a. Criterion C1 (Parental Income) – Cost

$$\begin{aligned} X_{1,1}^* &= \frac{500.000}{\sqrt{500.000^2 + 700.000^2 + 500.000^2 + 700.000^2 + 700.000^2 + 750.000^2}} \\ &= \frac{500.000}{1.591.383,046} = 0,3141 \\ X_{2,1}^* &= \frac{700.000}{\sqrt{500.000^2 + 700.000^2 + 500.000^2 + 700.000^2 + 700.000^2 + 750.000^2}} \\ &= \frac{700.000}{1.591.383,046} = 0,4398 \\ X_{3,1}^* &= \frac{500.000}{\sqrt{500.000^2 + 700.000^2 + 500.000^2 + 700.000^2 + 700.000^2 + 750.000^2}} \\ &= \frac{500.000}{1.591.383,046} = 0,3141 \\ X_{4,1}^* &= \frac{700.000}{\sqrt{500.000^2 + 700.000^2 + 500.000^2 + 700.000^2 + 700.000^2 + 750.000^2}} \\ &= \frac{700.000}{1.591.383,046} = 0,4398 \\ X_{5,1}^* &= \frac{700.000}{\sqrt{500.000^2 + 700.000^2 + 500.000^2 + 700.000^2 + 700.000^2 + 750.000^2}} \\ &= \frac{700.000}{1.591.383,046} = 0,4398 \\ X_{6,1}^* &= \frac{750.000}{\sqrt{500.000^2 + 700.000^2 + 500.000^2 + 700.000^2 + 700.000^2 + 750.000^2}} \\ &= \frac{750.000}{1.591.383,046} = 0,4712 \end{aligned}$$

b. Criterion C2 (Number of Dependents) – Benefit

$$X_{1,2}^* = \frac{3}{\sqrt{3^2 + 3^2 + 2^2 + 5^2 + 4^2 + 5^2}} = \frac{3}{9,38083152} = 0,3198$$

$$X_{2,2}^* = \frac{3}{\sqrt{3^2 + 3^2 + 2^2 + 5^2 + 4^2 + 5^2}} = \frac{3}{9,38083152} = 0,3198$$

$$X_{3,2}^* = \frac{2}{\sqrt{3^2 + 3^2 + 2^2 + 5^2 + 4^2 + 5^2}} = \frac{2}{9,38083152} = 0,2132$$

$$X_{4,2}^* = \frac{5}{\sqrt{3^2 + 3^2 + 2^2 + 5^2 + 4^2 + 5^2}} = \frac{5}{9,38083152} = 0,5330$$

$$X_{5,2}^* = \frac{4}{\sqrt{3^2 + 3^2 + 2^2 + 5^2 + 4^2 + 5^2}} = \frac{4}{9,38083152} = 0,4264$$

$$X_{6,2}^* = \frac{5}{\sqrt{3^2 + 3^2 + 2^2 + 5^2 + 4^2 + 5^2}} = \frac{5}{9,38083152} = 0,5330$$

c. Criterion C3 (Report Card Score) – Benefit

$$X_{1,3}^* = \frac{88}{\sqrt{88^2 + 80^2 + 90^2 + 80^2 + 70^2 + 69^2}} = \frac{88}{195,7166319} = 0,4496$$

$$X_{2,3}^* = \frac{80}{\sqrt{88^2 + 80^2 + 90^2 + 80^2 + 70^2 + 69^2}} = \frac{80}{195,7166319} = 0,4087$$

$$X_{3,3}^* = \frac{90}{\sqrt{88^2 + 80^2 + 90^2 + 80^2 + 70^2 + 69^2}} = \frac{90}{195,7166319} = 0,4598$$

$$X_{4,3}^* = \frac{80}{\sqrt{88^2 + 80^2 + 90^2 + 80^2 + 70^2 + 69^2}} = \frac{80}{195,7166319} = 0,4087$$

$$X_{5,3}^* = \frac{70}{\sqrt{88^2 + 80^2 + 90^2 + 80^2 + 70^2 + 69^2}} = \frac{70}{195,7166319} = 0,3576$$

$$X_{6,3}^* = \frac{69}{\sqrt{88^2 + 80^2 + 90^2 + 80^2 + 70^2 + 69^2}} = \frac{69}{195,7166319} = 0,3525$$

Thus, the normalized matrix can be presented as follows.

$$X_{ij}^* = \begin{bmatrix} 0,3141 & 0,3198 & 0,4496 \\ 0,4398 & 0,3198 & 0,4087 \\ 0,3134 & 0,2132 & 0,4598 \\ 0,4398 & 0,5330 & 0,4087 \\ 0,4398 & 0,4264 & 0,3576 \\ 0,4712 & 0,5330 & 0,3525 \end{bmatrix}$$

4. Calculating the Optimum Value

After obtaining the Normalized Matrix, the next step in the MOORA method is to calculate the optimization value. This is done by summing the benefit factors and subtracting the cost factor to determine the ranking of the most eligible students for financial aid.

$$Y_1^* = (3,5 \times 0,3198) + (2,0 \times 0,4496) - (4,5 \times 0,3141) = 0,60505$$

$$Y_2^* = (3,5 \times 0,3198) + (2,0 \times 0,4087) - (4,5 \times 0,4398) = -0,0424$$

$$Y_3^* = (3,5 \times 0,2132) + (2,0 \times 0,4598) - (4,5 \times 0,3134) = 0,2555$$

$$Y_4^* = (3,5 \times 0,5330) + (2,0 \times 0,4087) - (4,5 \times 0,4398) = 0,7038$$

$$Y_5^* = (3,5 \times 0,4264) + (2,0 \times 0,3576) - (4,5 \times 0,4398) = 0,2285$$

$$Y_6^* = (3,5 \times 0,5330) + (2,0 \times 0,3525) - (4,5 \times 0,4712) = 0,4501$$

5. Ranking

After calculating the optimization value Y^* for each alternative (student), the final step is to rank them based on their Y^* values. Students with the highest Y^* values are given top priority for financial aid, as they best meet the selection criteria based on economic and academic factors.

Table 22. Ranking Results

Student Name	Nilai Y^* Value	Ranking
A-14	0,7038	1
A-05	0,60505	2
A-27	0,4501	3
A-12	0,2555	4
A-26	0,2285	5
A-08	-0,0424	6

V. CONCLUSION

This study successfully developed a Decision Support System for selecting recipients of financial aid for underprivileged students by integrating the K-Means and MOORA methods. The K-Means method was employed to cluster students based on parental income, number of dependents, and academic performance, where students in Cluster 1—representing those with the lowest economic conditions—were prioritized for selection. Subsequently, the MOORA method was applied to rank eligible recipients more objectively, ensuring a more transparent and accurate selection process. The findings indicate that the developed system effectively reduces subjectivity, accelerates the analysis process, and enhances the efficiency of aid distribution. However, certain limitations exist, particularly in the number of variables used and the system's reliance on the accuracy of initial data. Therefore, future improvements may involve incorporating additional selection criteria, such as housing conditions and prior social assistance received, as well as integrating web-based or mobile applications to enhance the efficiency of the selection process. Furthermore, the implementation of artificial intelligence techniques could serve as a potential solution to improve the system's accuracy in assessing recipient eligibility. With further advancements, this system is expected to enhance the transparency, effectiveness, and efficiency of aid distribution for underprivileged students, thereby assisting educational institutions in optimizing student welfare programs.

REFERENCES

- [1] M. Reza, "Sistem Pendukung Keputusan Penerima Program Indonesia Pintar (PIP) Menggunakan Metode MOORA," 2022.
- [2] M. Suginan, E. Suryani, and U. L. Sapria, "Sistem Pendukung Keputusan Penerima Bantuan Siswa Miskin Menerapkan Metode WASPAS dan MOORA," *Nas. Sains Teknol. Inf.*, pp. 719â, vol. 727, 2018.
- [3] K. Kurniawan, J. Prayudha, and Z. Panjaitan, "Sistem Pendukung Keputusan Penentuan Penerimaan Beasiswa Menggunakan Metode Moora," *J. Cyber Tech*, vol. 1, no. 2, pp. 232–245, 2018.
- [4] L. Cahyani, M. Arif, and F. Ningsih, "Sistem Pendukung Keputusan Pemilihan Mahasiswa Berprestasi Menggunakan Metode Moora (Studi Kasus Fakultas Ilmu Pendidikan Universitas Trunojoyo Madura)," *J. Ilm. Edutic Pendidik. Dan Inform.*, vol. 5, no. 2, pp. 108–114, 2019.
- [5] H. Haryanto, "Pembuatan Aplikasi Sistem Penunjang Keputusan Untuk Pemilihan Penerima Beasiswa Siswa KMS dengan Metode MOORA," *J. Inf. J. Penelit. Dan Pengabd. Masy.*, vol. 4, no. 1, pp. 15–19, 2018.
- [6] Y. Rohmatin, W. Kusriani, A. Noor, and F. Fathurrahmani, "Sistem Pendukung Keputusan Penentuan Calon Penerima Beasiswa Menggunakan Metode Simple Additive Weighting (SAW) Berbasis Web," *J. Sains dan Inform.*, vol. 6, no. 1, pp. 102–111, 2020.
- [7] A. Romadhona, U. Ulfiah, S. Sukardi, and N. I. Sari, "Implementasi Seleksi Penerimaan Bantuan Siswa Miskin Dengan Metode Moora," *JTIK (Jurnal Tek. Inform. Kaputama)*, vol. 7, no. 1, pp. 128–135, 2023.
- [8] D. Darlinda and J. N. Utamajaya, "Sistem Pendukung Keputusan Penerima Beasiswa Program Indonesia Pintar Menggunakan Metode Algoritma K-Means Clustering," *JURIKOM (Jurnal Ris. Komputer)*, vol. 9, no. 2, pp. 167–175, 2022.
- [9] Y. Amaliah, "Sistem Pendukung Keputusan Penerima Beasiswa Tidak Mampu

- Menggunakan Metode Moora,” *J. Teknol. Inf.*, vol. 5, no. 1, pp. 12–18, 2021.
- [10] N. W. A. Ulandari, “Implementasi metode MOORA pada proses seleksi beasiswa bidikmisi di Institut Teknologi dan Bisnis STIKOM Bali,” *J. Eksplora Inform.*, vol. 10, no. 1, pp. 53–58, 2020.
 - [11] Y. S. Siregar, “Analisis Penerima Bantuan Beasiswa Program Studi Teknik Informatika Menggunakan Metode Moora Dan Topsis,” *JiTEKH*, vol. 9, no. 1, pp. 58–64, 2021.
 - [12] R. F. Sinaga, S. R. Andani, and S. Suhada, “Penentuan Penerima Kip Dengan Menggunakan Metode Moora Pada Sd Negeri 124395 Pematang Siantar,” *KOMIK (Konferensi Nas. Teknol. Inf. dan Komputer)*, vol. 2, no. 1, pp. 278–285, 2018.
 - [13] E. Nahak, F. Tedy, Y. C. H. Siki, E. Ngaga, E. Jando, and S. D. B. Mau, “Implementasi Metode MOORA dalam Sistem Pendukung Keputusan bagi Calon Penerima Beasiswa Program Indonesia Pintar di SMPN Satu Atap Nununamat,” *KONSTELASI Konvergensi Teknol. dan Sist. Inf.*, vol. 4, no. 1, pp. 83–98, 2024, doi: 10.24002/konstelasi.v4i1.8972.
 - [14] R. P. Sari and A. M. Alliandaw, “Penerapan Metode MOORA Pada Sistem Penentuan Penerimaan Bidikmisi UNTAN,” *J. Sisfokom (Sistem Inf. dan Komputer)*, vol. 11, no. 2, pp. 242–250, 2022, [Online]. Available: <https://jurnal.atmaluhur.ac.id/index.php/sisfokom/article/view/1420>
 - [15] E. Patel and D. S. Kushwaha, “Clustering cloud workloads: K-means vs gaussian mixture model,” *Procedia Comput. Sci.*, vol. 171, pp. 158–167, 2020.
 - [16] A. Žiberna, “K-means-based algorithm for blockmodeling linked networks,” *Soc. Networks*, vol. 61, pp. 153–169, 2020.
 - [17] C. L. Clayman, S. M. Srinivasan, and R. S. Sangwan, “K-means clustering and principal components analysis of microarray data of L1000 landmark genes,” *Procedia Comput. Sci.*, vol. 168, pp. 97–104, 2020.
 - [18] J. Hutagalung, N. L. W. S. R. Ginantra, G. W. Bhawika, W. G. S. Parwita, A. Wanto, and P. D. Panjaitan, “COVID-19 cases and deaths in Southeast Asia clustering using K-Means algorithm,” in *Journal of Physics: Conference Series*, 2021, vol. 1783, no. 1, p. 12027.
 - [19] A. P. U. Siahaan, R. Rahim, and M. Mesran, “Student Admission Assesment using Multi-Objective Optimization on the Basis of Ratio Analysis,” 2017.
 - [20] M. Moradian, V. Modanloo, and S. Aghaiee, “Comparative analysis of multi criteria decision making techniques for material selection of brake booster valve body,” *J. traffic Transp. Eng. (English Ed.)*, vol. 6, no. 5, pp. 526–534, 2019.