

Classification of Diabetes Mellitus Using Logistic Regression and MLP

I Ketut Putra Jaya^{a1}, I Nyoman Prayana Trisna^{a2}

^{a1}Information Technology Department, Faculty of Engineering,
Udayana University Bukit Jimbaran, Bali, Indonesia, telp. (0361) 701806
e-mail: putrajayaa002@gmail.com, prayana.trisna@unud.ac.id

Abstrak

Diabetes melitus atau penyakit kencing manis, adalah penyakit kronis yang berlangsung seumur hidup. Secara global, prevalensi diabetes pada orang dewasa berusia 20–79 tahun diperkirakan mencapai 10,5% atau sekitar 536,6 juta jiwa, dan diprediksi meningkat menjadi 12,2% atau sekitar 783,2 juta jiwa pada 2045. Peningkatan ini menjadikan diabetes sebagai salah satu penyebab utama kematian di dunia. Penelitian ini bertujuan mengembangkan model klasifikasi untuk skrining awal diabetes melitus menggunakan metode Regresi Logistik dan Multilayer Perceptron (MLP). Pendekatan yang digunakan meliputi perbandingan optimasi hyperparameter melalui Grid Search dan tanpa optimasi. Data yang digunakan berasal dari Rumah Sakit Bali Mandara dan Rumah Sakit Umum Tabanan Bali dengan 2.203 sampel dan 9 atribut. Hasil evaluasi menunjukkan bahwa hasil rasio terbaik yaitu 80:20 di MLP dengan optimasi Grid Search mencapai akurasi 92,97%, sementara Regresi Logistik memperoleh 88,21%.

Kata kunci: Diabetes melitus, regresi logistik, MultiLayer Perceptron, skrining dini, optimasi hyperparameter.

Abstract

Diabetes mellitus or diabetes, is a chronic disease that lasts a lifetime. Globally, the prevalence of diabetes in adults aged 20-79 years is estimated at 10.5% or around 536.6 million people, and is predicted to increase to 12.2% or around 783.2 million people by 2045. This increase makes diabetes one of the leading causes of death in the world. This study aims to develop a classification model for early screening of diabetes mellitus using Logistic Regression and Multilayer Perceptron (MLP) methods. The approach used includes comparison of hyperparameter optimization through Grid Search and without optimization. The data used came from Bali Mandara Hospital and Tabanan Bali General Hospital with 2,203 samples and 9 attributes. The evaluation results showed that the best ratio result of 80:20 in MLP with Grid Search optimization achieved 92.97% accuracy, while Logistic Regression obtained 88.21% accuracy.

Keywords: Diabetes mellitus, logistic regression, MultiLayer Perceptron, early screening, hyperparameter optimization.

1. Introduction

Diabetes mellitus, or better known as diabetes, is a chronic disease that can be suffered for life [1]. According to the IDF Diabetes Atlas report in 2021, the prevalence of diabetes in adults aged 20-79 years is estimated to reach 10.5% or around 536.6 million people, with a projected increase to 12.2% or around 783.2 million people in 2045. This increase shows a significant trend in the spread of diabetes in the world [2].

Diabetes mellitus is divided into several types, including type 1, type 2, gestational, and other specific diabetes. The most common diabetes case findings are type 1 and type 2 diabetes (Hardianto, 2021). According to the International Diabetes Federation (IDF), more than 90% of people with diabetes mellitus are type 2. Factors causing diabetes mellitus include economic conditions, demographics, environment, and genetics [2]. Common symptoms of diabetes include dehydration, weight loss, fatigue, decreased vision, cramps, candida infection, and constipation [3].

Advances in technology, especially in healthcare, have introduced Artificial Intelligence (AI) such as the application of machine learning and deep learning to improve the accuracy and efficiency of disease screening and diagnosis. Technology allows medical personnel to screen faster and more accurately, while supporting more informed decision-making. The application of machine learning algorithms to analyze medical data, such as laboratory results and symptoms, allows the screening process to be automated, reducing the workload of medical personnel and providing faster results. Implementing the use of appropriate algorithms, the screening system can be more easily customized to the needs and characteristics of patients. This paves the way for the development of more efficient screening systems by utilizing more diverse and relevant medical data [4].

Efforts to improve the effectiveness of diabetes mellitus disease classification have been made through the utilization of machine learning algorithms, such as logistic regression and MultiLayer Perceptron (MLP). Logistic regression, used in a dataset from the National Institute of Diabetes and Digestive and Kidney Diseases (NIDDK), showed 81% accuracy, and after optimization with Grid Search, the accuracy increased to 83.33% [5]. Meanwhile, MLP proved superior with an accuracy of up to 98.6%, beating other conventional methods that only reached around 85% [6]. Based on these findings, logistic regression and MLP both show great potential in diabetes mellitus classification. Therefore, this study aims to compare the two methods to determine which one is more effective in early screening of diabetes mellitus.

The research developed aims to develop a classification model for early screening of diabetes mellitus using logistic regression and MLP methods, and compare the performance of the two methods. Previous research has used these methods with adequate results, but the developed study provides a more in-depth look at their effectiveness in the context of early screening for diabetes mellitus. The data used comes from Bali Mandara Hospital and Bali Tabanan General Hospital, which provide data related to diabetes mellitus indicators. Based on the evaluation, this study will determine which method is superior in terms of accuracy for use in early screening of diabetes mellitus.

2. Research Method / Proposed Method

This study aims to analyze the classification of diabetes mellitus using logistic regression and Multi Layer Perceptron (MLP) methods. The data used comes from Bali Mandara Hospital and Tabanan General Hospital, with a total sample size of 2,203 and includes 8 medical attributes. The research process is depicted in the attached flowchart, as seen in Figure 1, which explains each step starting from data collection to the result visualization stage.

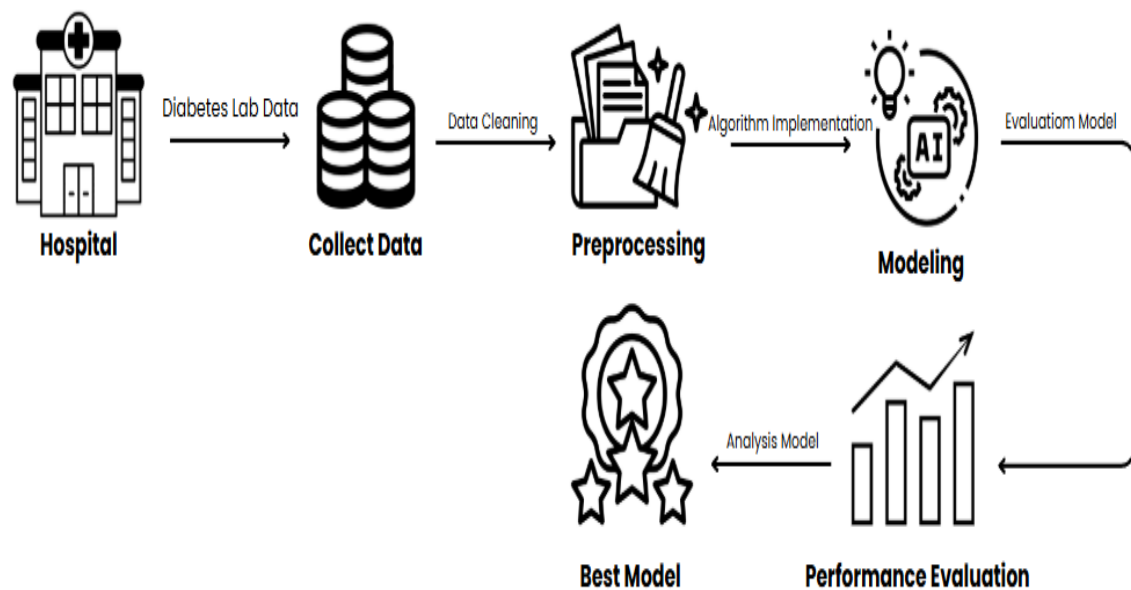


Figure 1. Research Overview

Figure 1 is the beginning of the study, which started with the collection of medical data from Bali Mandara Hospital and Tabanan General Hospital, with a total sample size of 2,203 data. The data collected included 9 key medical attributes used for diabetes mellitus classification. Table 1 below shows the indicators obtained from the laboratory data:

Table 1. Research Datasets

Indicator	Description
Gender	Patient's gender (male/female)
Age	Patient's age at the time of data collection
Random Blood Glucose	Blood glucose level at the time of examination without considering the time
2-Hour Postprandial Blood Glucose	Blood glucose level 2 hours after oral glucose intake
Fasting Blood Glucose	Blood glucose level after 8 hours of fasting
HbA1C	Hemoglobin A1C level, an indicator of long-term blood glucose control
HDL Cholesterol	High-density lipoprotein cholesterol level in the blood
LDL Cholesterol	Low-density lipoprotein cholesterol level in the blood
Total Cholesterol	Total Cholesterol in the Blood

Table 1 is the data of diabetes mellitus-related indicators that have been collected from the laboratory system of Bali Mandara Hospital and Tabanan General Hospital. Data preprocessing is the next stage, which includes several important steps, starting with data filtering, which aims to filter out relevant information and reduce distractions that can affect the accuracy of the analysis. Record selection is performed to select data that matches certain criteria, facilitating analysis and identification of patterns in the dataset. Next, feature engineering transforms raw data into more informative features to improve the quality of predictive models. Pivoting changes the orientation of the data to make it easier to understand and analyze. Handling missing values addresses missing or empty data to maintain consistency and accuracy of analysis. Finally, label transformation changes the labels in the dataset to make it easier to categorize, making it easier for the model to identify patterns.

Once the data is processed, classification models are built using two main methods: Logistic Regression and Multi Layer Perceptron (MLP). Both methods were tested using two optimization techniques: Grid Search and without optimization. The data sharing scenarios used were 60:40, 70:30, and 80:20 ratios, aiming to evaluate the effect of training set size on model performance. Model evaluation is performed using evaluation metrics such as Precision, Recall, and F1-Score. Using these metrics, each model will be analyzed to determine which one provides the best results in terms of accuracy and reliability. Based on the analysis, the best model will be selected for use in early screening of diabetes mellitus.

3. Literature Study

Contains all libraries used as references in this study. All references in the text are written w Literature study is a reference source that contains material relevant to the topic of the research being conducted. The section discusses basic concepts and previous studies related to diabetes classification using logistic regression and Multi Layer Perceptron (MLP) algorithms. The review is organized to provide an understanding of the techniques applied, evaluation metrics used, and algorithmic considerations applied in diabetes classification in the healthcare sector. The main focus of the review is to assess the effectiveness of the algorithms in detecting diabetes at an early stage and how data processing can affect model performance.

3.1. Classification

Classification is a method of grouping samples based on similar characteristics using target variables as categories. In data science, classification is used to build models from unclassified data to classify new data [7]. This technique is important in various fields, including health, such as in diagnosis determination to accelerate patient treatment through medical record data [8]. Classification helps group patient data, making it easier to identify groups in new data. It is widely used to detect diseases such as heart disease, diabetes, tumors, and other health conditions, improving early diagnosis, treatment, patient care, and

overall health outcomes.

3.2. Preprocessing

Preprocessing is the initial stage in managing raw data that will be used in further analysis. At this stage, the data will be processed through several steps, such as filtering, record selection, feature engineering, pivoting, handling missing values, and label transformation, to ensure the data used is relevant, consistent, and ready for further analysis. This process helps to improve the accuracy and performance of models by eliminating noise and irrelevant information [9].

3.3. Logistic Regression

Logistic Regression in Machine Learning is an algorithm used to model the relationship between predictor variables (features) and response variables (targets). It is often used in classification problems, where the main goal is to predict the category or class of the data based on the features. Logistic Regression works by calculating the probability that a data belongs to a certain category, using a logistic function [10].

3.4. Multilayer Perceptron

Multilayer Perceptron (MLP) is a type of artificial neural network (ANN) consisting of multiple layers of neurons connected by weights. It includes an input layer, at least one hidden layer, and an output layer. The data flows forward from the input to the output layer without feedback, known as feedforward propagation [11]. MLP is trained through backpropagation, allowing it to learn complex patterns in data.

3.5. Grid Search CV

Grid Search is a hyperparameter tuning method that allows users to perform a systematic search for combinations of various hyperparameters that have been selected for the model [12]. By using Grid Search, it can test various combinations of hyperparameter values to find the combination that produces the best model performance, thereby improving the accuracy and effectiveness of the model.

3.6. Evaluation Methods

Performance evaluation is the stage of testing and measuring the performance of the applied classification model. In this research, the performance of the model will be measured and compared based on several evaluation metrics, namely accuracy, precision, recall, and f1-score. These values are obtained by calculating the results of the confusion matrix used to assess the extent to which the model can classify data correctly [13].

4. Result and Discussion

The Logistic Regression and Multilayer Perceptron (MLP) algorithms used in this study will be implemented on a diabetes dataset consisting of 2203 patient samples. The primary objective of this research is to effectively classify diabetes by utilizing two distinct machine learning models: Logistic Regression and Multilayer Perceptron (MLP). In this study, a series of experiments will be conducted to evaluate the performance of both models under various conditions. These conditions include different data split scenarios as well as the implementation of Grid Search optimization to determine the most suitable hyperparameters for each model. The effectiveness of each model will be assessed based on the classification accuracy achieved under these varying scenarios.

Table 2. Analysis Testing Scenarios

Testing	Scenario Split
Not using grid search	60:40, 70:30, 80:20
Using Grid Search	60:40, 70:30, 80:20

Table 2 is the analysis testing scenario, where the scenario without the use of grid search optimization will be applied to the Logistic Regression and Multilayer Perceptron (MLP) algorithms with splits of 60:40, 70:30, and 80:20. Meanwhile, Grid Search optimization will be

used on both algorithms with the same split variations. The testing scenarios will be categorized and outlined in a table 2, providing a comprehensive overview of the different configurations used throughout the experiments. The following is a table of parameters used in the grid search trials of the logistic regression and multilayer perceptron methods presented in tables 3 and 4.

Table 3. Logistic Regression Search Parameters and Grid Values

Hyperparameters	Usefulness of Hyperparameters	Grid Search Value
C	Represents a regulatory parameter that controls the strength of regularization	[0.0001, 0.001, 0.01, 0.1, 1, 10, 100, 1000]
Solver	This solver helps the process of finding coefficients that minimize the loss function.	['liblinear', 'saga']
Penalty	Determine the type of regularization applied to prevent overfitting	['l1', 'l2']
Max_iter	Determines the maximum number of iterations that the solver performs during the optimization process to achieve convergence.	[500, 1000, 2000]

Table 3 presents the search parameters and grid values for logistic regression. It outlines four key hyperparameters: C, which regulates the strength of regularization, Solver, responsible for finding the optimal coefficients, Penalty, which determines the type of regularization to avoid overfitting, and Max_iter, specifying the maximum number of iterations for convergence. The grid search values provided allow testing various configurations to optimize the model's performance. These values are essential for fine-tuning logistic regression to achieve better accuracy. The grid search process helps in selecting the optimal combination of hyperparameters to improve the model's generalization ability. By adjusting the C parameter, the model can balance between underfitting and overfitting. The choice of Solver affects the efficiency of finding the optimal coefficients, while the Penalty determines the type of regularization used, such as L1 or L2, to reduce model complexity. The Max_iter hyperparameter ensures that the solver has enough iterations to converge to the best solution. The grid search values for logistic regression guide the process of exploring different possibilities [14]. These values and their effects on performance are vital for creating a robust model. The parameters and values from the grid search for multilayer perceptron can be seen in Table 4.

Table 4. Multilayer Perceptron Search Parameters and Grid Values

Hyperparameters	Usefulness of Hyperparameters	Grid Search Value
mlp__hidden_layer_sizes	Size and number of hidden layers in the model	(32,), (64,), (64, 32), (100,), (128,)
mlp__activation	Determines the activation function used in each neuron	['relu', 'tanh']
mlp__solver	Determine the optimization algorithm used to train the model	['adam']
mlp__alpha	Is a regularization parameter that helps prevent overfitting	[0.01, 0.1, 1.0]
mlp__learning_rate	Determines how fast or slow the model learns by controlling the rate at which the weights are updated during training.	['constant', 'adaptive']
mlp__learning_rate_init	Determine the initial value for the learning rate	[0.001, 0.01]
mlp__max_iter	Determines the maximum number of iterations performed during training to achieve convergence	[300]

Table 4 presents the search parameters and grid values for multilayer perceptron

(MLP). It outlines several key hyperparameters: `mlp_hidden_layer_sizes`, which defines the size and number of hidden layers in the model, `mlp_activation`, which determines the activation function used in each neuron, and `mlp_solver`, which specifies the optimization algorithm used to train the model. Additionally, `mlp_alpha` is a regularization parameter that helps prevent overfitting, while `mlp_learning_rate` controls the speed at which the model learns during training by adjusting the rate at which the weights are updated. `mlp_learning_rate_init` determines the initial value for the learning rate, and `mlp_max_iter` specifies the maximum number of iterations performed during training to achieve convergence [15]. The grid search values provided allow testing various configurations to fine-tune the multilayer perceptron model for optimal performance. These parameters are essential for customizing the learning process and improving the generalization ability of the model.

4.1. Performance Comparison Without Grid Search

The performance comparison is conducted by examining the accuracy levels produced by the Logistic Regression and Multilayer Perceptron (MLP) models across various data split scenarios. This analysis aims to evaluate the effectiveness of both algorithms in classifying diabetes without the use of grid search optimization. The results provide insights into how each model performs under different train-test ratios, as detailed in Table 4, offering a comprehensive overview of their classification capabilities in each scenario.

Table 5. Comparison Accuracy Without Grid Search

Testing	Scenario Split	Accuracy
Logistic Regression Without Grid Search	60:40	0.6281
Logistic Regression Without Grid Search	70:30	0.6293
Logistic Regression Without Grid Search	80:20	0.6576
Multilayer Perceptron Without Grid Search	60:40	0.8844
Multilayer Perceptron Without Grid Search	70:30	0.8971
Multilayer Perceptron Without Grid Search	80:20	0.9002

Based on the results in Table 5, the highest accuracy achieved by the Logistic Regression model without grid search optimization is 0.6576 with an 80:20 train-test split. In comparison, the Multilayer Perceptron (MLP) model demonstrates superior performance, reaching a maximum accuracy of 0.9002 with the same 80:20 split. This indicates that, under the tested scenarios, the Multilayer Perceptron (MLP) outperforms Logistic Regression in classifying diabetes data. A more specific evaluation is carried out using additional performance metrics—precision, recall, and f1-score—to provide a deeper understanding of the classification results. These metrics highlight how well each model performs not just overall, but also in identifying each class correctly. The detailed evaluation results for both models are in Table 6.

Table 6. Performance Metrics Without Grid Search

Testing	Split	Class	Precision	Recall	F1-score
Logistic Regression Without Grid Search	60:40	0	0.25	1.00	0.41
		1	0.82	0.54	0.65
		2	0.62	0.71	0.67
Logistic Regression Without Grid Search	70:30	0	0.27	1.00	0.42
		1	0.83	0.52	0.64
		2	0.62	0.74	0.67
Logistic Regression Without Grid Search	80:20	0	0.30	1.00	0.46
		1	0.84	0.57	0.68
		2	0.63	0.74	0.68
Multilayer Perceptron Without Grid Search	60:40	0	0.75	0.71	0.73
		1	0.89	0.91	0.90
		2	0.89	0.87	0.88
Multilayer Perceptron Without Grid Search	70:30	0	0.77	0.76	0.75
		1	0.90	0.91	0.91
		2	0.90	0.89	0.90
		0	0.94	0.76	0.84

Multilayer Perceptron Without Grid Search	80:20	1	0.89	0.94	0.91
		2	0.91	0.86	0.89

Based on the results presented in Table 6, the Logistic Regression model without grid search optimization achieved the best performance at a train-test split of 80:20, especially in the prediabetes class with a precision value of 0.84, recall of 0.57, and f1-score of 0.68. On the other hand, the Multilayer Perceptron (MLP) model showed the highest performance in the diabetes class with precision 0.91, recall 0.86, and f1-score 0.89 at the same ratio. The results show that although both models perform optimally at a ratio of 80:20, the Multilayer Perceptron (MLP) excels overall, especially in detecting diabetes cases.

4.2. Performance Comparison With Grid Search

The performance evaluation is carried out by analyzing the accuracy results obtained from the Logistic Regression and Multilayer Perceptron (MLP) models under several data split configurations. This evaluation highlights the effectiveness of both algorithms in classifying diabetes cases, utilizing grid search optimization to enhance model performance. The detailed results across different train-test ratios are presented in Table 7, providing a thorough comparison of their classification accuracy in each scenario.

Table 7. Comparison Accuracy With Grid Search

Testing	Scenario Split	Accuracy
Logistic Regression With Grid Search	60:40	0.8367
Logistic Regression With Grid Search	70:30	0.8427
Logistic Regression With Grid Search	80:20	0.8821
Multilayer Perceptron With Grid Search	60:40	0.8844
Multilayer Perceptron With Grid Search	70:30	0.9107
Multilayer Perceptron With Grid Search	80:20	0.9297

Based on the results in Table 7, the Logistic Regression model with grid search optimization achieved an accuracy of 0.8821 with an 80:20 train-test split, while the Multilayer Perceptron (MLP) model outperformed it with an accuracy of 0.9297. The best parameters for the MLP model were: activation function "tanh", alpha 0.1, hidden layer sizes (64, 32), adaptive learning rate, learning rate 0.001, 300 iterations, and solver "adam". For Logistic Regression, the best parameters were: C=1000, 500 iterations, L1 penalty, and solver "liblinear". Further evaluation using precision, recall, and F1-score is available in Table 8.

Table 8. Performance Metrics With Grid Search

Testing	Split	Class	Precision	Recall	F1-score
Logistic Regression With Grid Search	60:40	0	0.85	0.83	0.84
		1	0.87	0.84	0.85
		2	0.79	0.84	0.81
		0	0.90	0.84	0.87
Logistic Regression With Grid Search	70:30	1	0.87	0.84	0.86
		2	0.80	0.84	0.82
		0	0.90	0.86	0.88
		1	0.90	0.90	0.90
Logistic Regression With Grid Search	80:20	2	0.86	0.87	0.86
		0	0.76	0.60	0.67
Multilayer Perceptron With Grid Search	60:40	1	0.87	0.93	0.90
		2	0.92	0.85	0.88
		0	0.81	0.81	0.81
		1	0.92	0.92	0.92
Multilayer Perceptron With Grid Search	70:30	2	0.91	0.90	0.91
		0	0.94	0.76	0.84
	80:20	1	0.92	0.96	0.94
		2	0.95	0.90	0.93

Based on Table 8, Logistic Regression achieves the best performance at a ratio of 80:20, especially in the prediabetes class with precision, recall, and f1-score values of 0.90 each. On the other hand, Multilayer Perceptron (MLP) showed the highest performance in the diabetes class with precision 0.95, recall 0.90, and f1-score 0.93 at the same ratio. The results show that although both models perform optimally at a ratio of 80:20, the Multilayer Perceptron (MLP) excels overall, especially in detecting diabetes cases.

4.3. Vizualitation

The best model in diabetes classification analysis is Multilayer Perceptron (MLP) without grid search optimization with 80:20 data division, which achieved an accuracy of 0.9297. This result shows a very good performance in classifying diabetes data. Visualization of the results is presented in the form of bar graphs depicting the amount of training and validation data, as well as confusion matrix in the form of heatmap and learning curve. These graphs can be seen in Figure 2.

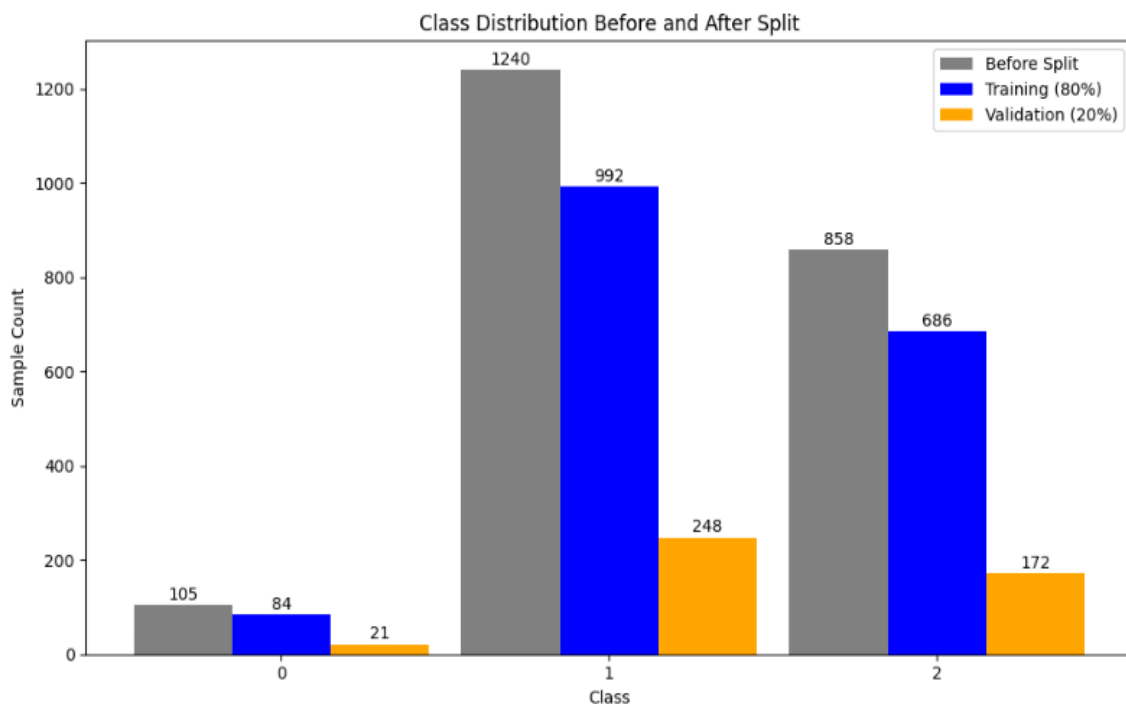


Figure 2. Class Distribution Before and After Data Split

Figure 2 is the class distribution before and after splitting the data with a ratio of 80:20 for training and validation. Before the split, the class distribution was fairly balanced, and after the split, the samples were proportionally split between training (80%) and validation (20%) data. The gray bars represent the distribution before the split, while the blue and orange bars represent the training and validation data. This 80:20 split is the best split for the Multilayer Perceptron (MLP) model with optimization, resulting in the highest accuracy of 92.97%. With this ratio, MLP can maximize the training data and still have enough validation data for accurate evaluation. This split ratio has proven to provide a more robust model compared to other split variations. The results highlight the effectiveness of using a larger proportion of data for training without compromising the validation process. The model's performance is more reliable with this configuration, which further supports the choice of an 80:20 split. The next visualization is the confusion metric heatmap, which is shown in Figure 3.

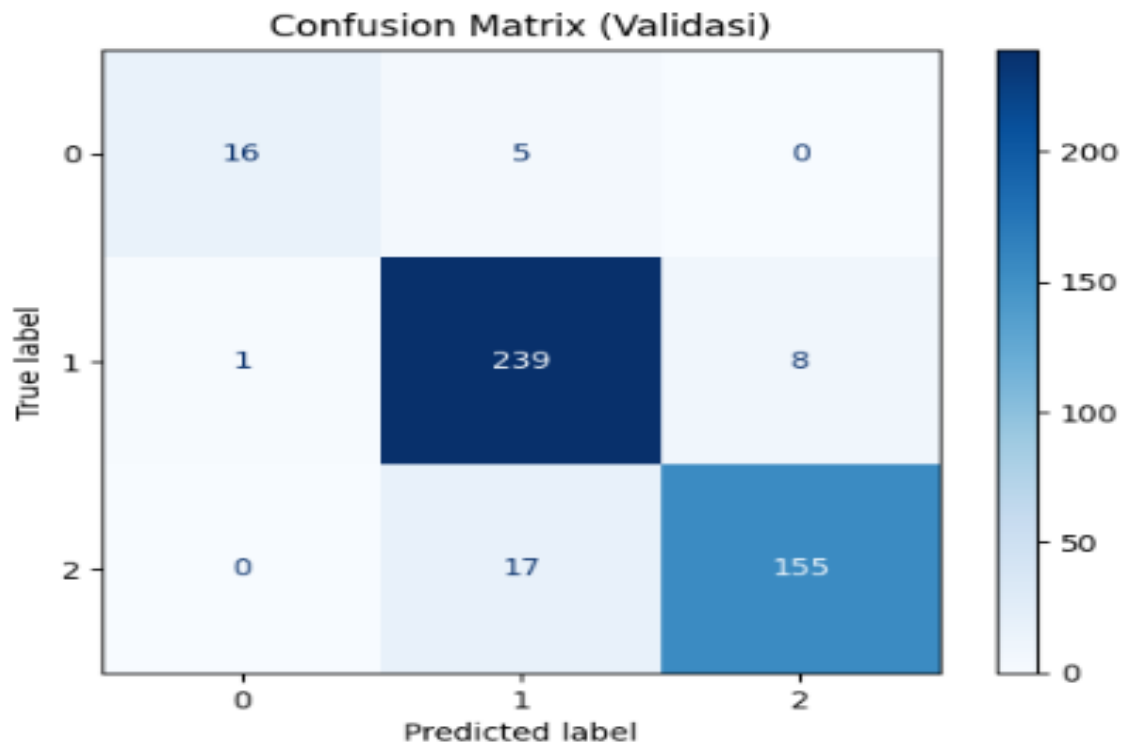


Figure 3. Confusion Matrix

Figure 3 shows the Confusion Matrix for the model validation results. For class 0, the model correctly predicted 16 samples and made 5 incorrect predictions. For class 1, the model correctly predicted 239 samples and was wrong 9 times. As for class 2, the model correctly predicted 155 samples and made 17 incorrect predictions. Overall, the model performed well, with the highest number of errors occurring in class 2. The next visualization is the learning curve shown in Figure 4.

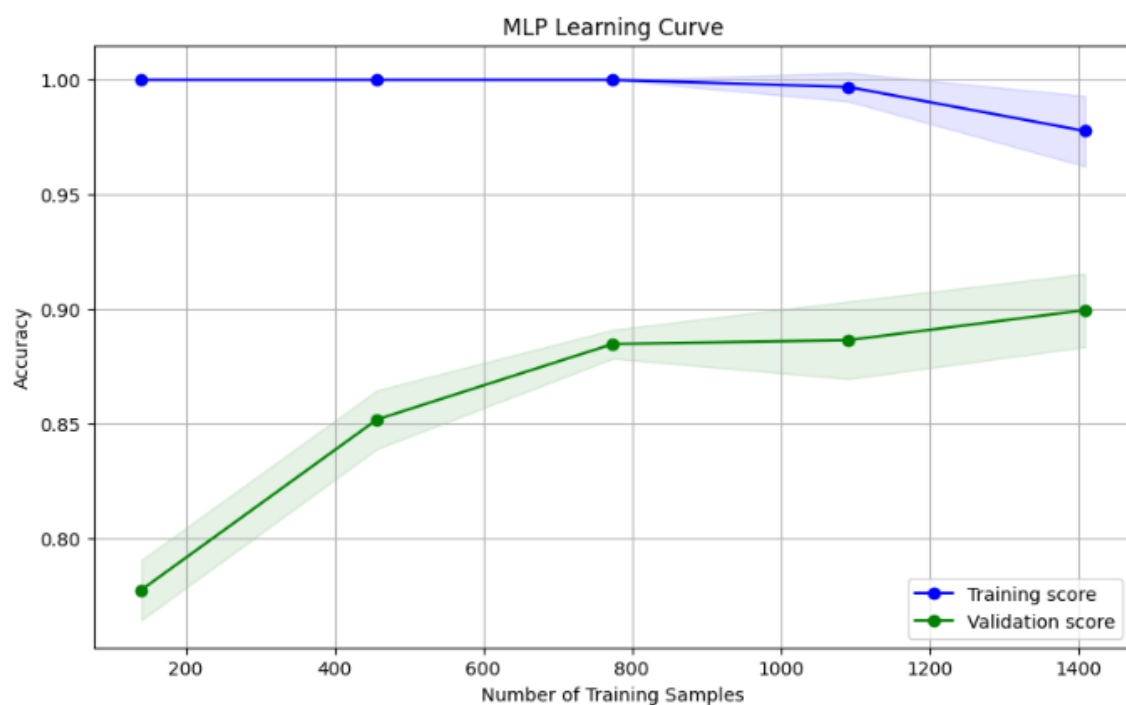


Figure 4. Confusion Matrix

Figure 4 is the learning curve of the best Multilayer Perceptron (MLP) model. Based on the graph, it can be seen that the training accuracy (blue line) is very high when the amount of training data is still small, even reaching 100%, but tends to decrease as the amount of training data increases. Meanwhile, the validation accuracy (green line) increases as the data increases until it stabilizes in the range of 0.88-0.89.

5. Conclusion

Based on diabetes mellitus classification research using Logistic Regression and Multilayer Perceptron (MLP) methods, the data obtained comes from Bali Mandara Hospital and Tabanan Hospital with a total of 2203 diabetes patient samples collected. The data has 9 attributes, and the pre-processing process includes filtering, record selection, feature engineering, pivoting, handling missing values, and label transformation. In the study, the grid search optimization technique and without grid search optimization were used in both models to evaluate the model performance. Data sharing is done with various split scenarios of 60:40, 70:30, and 80:20. The best results of each model showed that Multilayer Perceptron (MLP) produced the highest validation accuracy of 92.97% at a data split ratio of 80:20 with optimization, while Logistic Regression obtained a validation accuracy of 88.21% at a ratio of 80:20 with grid search optimization. The evaluation results show that Multilayer Perceptron (MLP) with optimization provides a relatively small improvement, but still effective in diabetes classification when compared to Logistic Regression even with grid search optimization. This study revealed that the MLP technique yields better performance in the classification of diabetes mellitus data. On the other hand, the application of grid search optimization to Logistic Regression resulted in a significant increase in accuracy, providing better performance than the model without optimization. Nonetheless, although Logistic Regression showed clear improvements after optimization, MLP remained superior in terms of overall performance, especially in classifying diabetes data with higher accuracy.

References

- [1] M. Soleh, N. Ammar, and I. Sukmadi, "Website-Based Application for Classification of Diabetes Using Logistic Regression Method," *J. Ilm. Merpati (Menara Penelit. Akad. Teknol. Informasi)*, vol. 9, no. 1, p. 23, 2021, doi: 10.24843/jim.2021.v09.i01.p03.
- [2] I. D. Federation, "Facts & Figures: Diabetes Around the World," *International Diabetes Federation*, 2021. <https://idf.org/about-diabetes/diabetes-facts-figures/> (accessed Oct. 20, 2024).
- [3] D. Hardianto, "Telaah Komprehensif Diabetes Melitus: Klasifikasi, Gejala, Diagnosis, Pencegahan, Dan Pengobatan," *J. Bioteknol. Biosains Indones.*, vol. 7, no. 2, pp. 304–317, 2021, doi: 10.29122/jbbi.v7i2.4209.
- [4] L. W. Astuti, I. Saluza, E. Yulianti, and D. Dhamayanti, "Feature Selection Menggunakan Binary Wheal Optimizatoin Algorithm (BWOA) pada Klasifikasi Penyakit Diabetes," *J. Ilm. Inform. Glob.*, vol. 13, no. 1, pp. 1–6, 2022, doi: 10.36982/jiig.v13i1.2057.
- [5] Z. Mutaqin, C. Rozikin, and Y. A. Tomo, "Jurnal Sistem dan Teknologi Informasi (JSTI) KLASIFIKASI PENYAKIT DIABETES MENGGUNAKAN ALGORITMA LOGISTIC REGRESSION," vol. 06, no. 3, pp. 320–329, 2024, [Online]. Available: <https://journalpedia.com/1/index.php/jsti>
- [6] N. C. Majid, F. R. Mamun, H. Geovani, and I. Nurhamidah, "Penggunaan Multilayer Perceptron Untuk Klasifikasi Diabetes Mellitus," vol. 2, no. 7, pp. 1101–1107, 2024.
- [7] D. Dalbergio, M. N. Hayati, and Y. N. Nasution, "Klasifikasi Lama Studi Mahasiswa Menggunakan Metode C5.0 pada Studi Kasus Data Kelulusan Mahasiswa Fakultas Matematika Dan Ilmu Pengetahuan Alam Universitas Mulawarman Tahun 2017," *Pros. Semin. Nas. Mat. Stat. dan Apl. 2019*, vol. 1, no. 1, pp. 36–42, 2019.
- [8] A. Anggraini, L. Widjaja, L. Indawati, and D. Rosmala Dewi, "Analisis Ketepatan Kode Diagnosis Kasus Persalinan Secara Sectio Caesarea Di Rumah Sakit Pelabuhan Jakarta," *Cerdika J. Ilm. Indones.*, vol. 3, no. 1, pp. 6–11, 2023, doi: 10.59141/cerdika.v3i1.505.
- [9] F. Alghifari and D. Juardi, "Penerapan Data Mining Pada Penjualan Makanan Dan

- Minuman Menggunakan Metode Algoritma Naïve Bayes,” *J. Ilm. Inform.*, vol. 9, no. 02, pp. 75–81, 2021, doi: 10.33884/jif.v9i02.3755.
- [10] A. Pratama, A. C. Nurcahyo, and L. Firgia, “Penerapan Machine Learning dengan Algoritma Logistik Regresi untuk Memprediksi Diabetes,” *Pros. CORISINDO 2023*, pp. 116–121, 2023, [Online]. Available: <https://stmikpontianak.org/ojs/index.php/corisindo/article/view/30%0Ahttps://stmikpontianak.org/ojs/index.php/corisindo/article/download/30/22>
- [11] P. Githa Pratiwi, I. Ketut Gede Darma Putra, and D. Purnami Singgih Putri, “Peramalan Jumlah Tersangka Penyalahgunaan Narkoba Menggunakan Metode Multilayer Perceptron,” *J. Ilm. Merpati (Menara Penelit. Akad. Teknol. Informasi)*, vol. 7, no. 2, p. 143, 2019, doi: 10.24843/jim.2019.v07.i02.p06.
- [12] P. R. Togatorop, M. Sianturi, D. Simamora, and D. Silaen, “Optimizing Random Forest using Genetic Algorithm for Heart Disease Classification,” *Lontar Komput. J. Ilm. Teknol. Inf.*, vol. 13, no. 1, p. 60, 2022, doi: 10.24843/lkjiti.2022.v13.i01.p06.
- [13] R. Nurhidayat and K. E. Dewi, “Penerapan Algoritma K-Nearest Neighbor Dan Fitur Ekstraksi N-Gram Dalam Analisis Sentimen Berbasis Aspek,” *Komputa J. Ilm. Komput. dan Inform.*, vol. 12, no. 1, pp. 91–100, 2023, doi: 10.34010/komputa.v12i1.9458.
- [14] R. M. Daffaa, D. Santika, and F. Mahardika, “Perbandingan Xgboost dan Logistic Regression dalam Memprediksi Credit Card Customer Churn,” vol. 3, 2025.
- [15] L. A. Zahir, “METODE MULTI LAYER PERCEPTRON (OPTIMIZATION OF CONCRETE COMPRESSIVE STRENGTH WITH ARTIFICIAL NEURAL NETWORKS USING MULTI LAYER PERCEPTRON),” pp. 45–55, 2024.