

Multi-Class Mental Health Classification Based on DASS-21 and Perceived Social Support Using Machine Learning Algorithms

Winny Dwita Sumbayak¹, Didi Supriyadi^{2,3*}

^{1,2}*Information Systems Program, Telkom University, Purwokerto, Indonesia*

³*Center of Excellence for Sustainability Cities, Village, and Food Security, Research Institute for Intelligent Business and Sustainable Economy, Telkom University, Purwokerto, Indonesia*

*corr-author: didisupriyadi@telkomuniversity.ac.id

Abstract - Mental health issues among students require data-driven approaches for early identification. This study aims to classify students' mental health levels using the Depression Anxiety Stress Scale (DASS-21) and perceived social support, measured by the Multidimensional Scale of Perceived Social Support (MSPSS), via machine learning algorithms. A supervised classification approach was employed using Random Forest, Support Vector Machine, and Logistic Regression on data collected from 450 respondents. The data were processed through scoring, labeling, encoding, balancing, and stratified 80:20 splitting. Model evaluation was conducted using hold-out testing and 5-fold cross-validation to ensure robust and reliable performance estimation. The results indicate that Random Forest achieved the best performance, with an accuracy of 0.97 on the test set, outperforming Support Vector Machine (0.81) and Logistic Regression (0.83). Improvements in recall and F1-score for minority classes demonstrate the effectiveness of the balancing process. These findings highlight the potential of machine learning for student mental health classification, although further validation on larger and more diverse datasets is required.

Keywords: mental health, classification, machine learning, DASS-21, perceived social support.

I. INTRODUCTION

Mental health is a fundamental component in supporting students' academic, social, and emotional development [1]. At the higher education level, increasing academic demands and social pressures make students a vulnerable group to stress, anxiety, and depression [2]. These conditions negatively affect academic achievement and learning motivation and may develop into more serious psychological disorders if not detected at an early stage [3].

In addition to academic factors, students' mental health is also significantly influenced by social support, including relationships with family and peers [4]. Low

levels of social support are associated with a higher risk of psychological distress, emphasizing the importance of incorporating psychosocial variables into mental health analysis [5]. Psychological conditions are commonly assessed using the DASS-21 instrument due to its strong validity and reliability in measuring depression, anxiety, and stress [6]. With advances in data analysis techniques, machine learning approaches have been increasingly applied to improve prediction accuracy and classification performance in mental health research [7].

From a technical perspective, machine learning algorithms such as Random Forest, Support Vector Machine, and Logistic Regression exhibit distinct characteristics in classification tasks. Random Forest, as an ensemble method, can reduce overfitting and capture complex nonlinear relationships in data [8]. Support Vector Machine is effective in high-dimensional spaces by constructing optimal hyperplanes to separate classes [9]. Logistic Regression is a probabilistic linear model with high interpretability and is widely used as a baseline method in various classification problems [10]. In general, the performance of machine learning models depends heavily on data quality, feature engineering, and preprocessing techniques [11].

Previous studies have shown that machine learning is highly effective across various classification problems. Machine learning techniques have been widely implemented to improve predictive performance through multi-algorithm approaches [12]. In addition, handling imbalanced data has been proven to enhance classification performance in medical prediction tasks [13]. Furthermore, Random Forest, combined with optimization techniques, has been shown to improve accuracy on complex datasets [14], while Support Vector Machine remains an effective method in data mining classification due to its ability to construct optimal decision boundaries [15]. These findings indicate that machine learning algorithms play an important role in

improving classification performance, particularly when supported by appropriate preprocessing techniques and optimal model configuration.

However, most existing studies focus primarily on psychological variables and do not integrate social support factors into classification models [16]. Moreover, comparative studies that combine psychological and social features and employ multiple machine learning algorithms remain limited, particularly for Indonesian university students [17]. To address this gap, the present study integrates psychological and social support variables within a unified classification framework. In addition, a class imbalance handling strategy is implemented, followed by a comparative evaluation of Random Forest, Support Vector Machine, and Logistic Regression to determine the most effective model. This approach is expected to provide a more comprehensive understanding of the relationship between psychological conditions and social support, while supporting early detection efforts and the development of more effective and sustainable psychosocial interventions in educational settings.

II. METHOD

Figure 1 illustrates a systematic process that begins with data collection, proceeds to methodological

computation, and culminates in evaluation. This depiction emphasizes the orderly progression from data acquisition to the final assessment phase, underscoring the methodology's structured nature.

A. Raw Questionnaire Data

This study employed a quantitative approach using a supervised classification scheme to predict students' mental health levels based on perceived social support, as measured by the Multidimensional Scale of Perceived Social Support (MSPSS) and the Depression Anxiety Stress Scale (DASS-21).

The research sample was selected using purposive sampling, with the criterion of active higher-education students, as this population was considered relevant to the study objectives due to their exposure to academic pressure, psychological stress, and varying levels of social support. However, this sampling approach may limit the generalizability of the findings, as the results are primarily applicable to similar student populations and may not fully represent other groups.

The minimum sample size was calculated using the Slovin formula based on a population of 5,283 students. Although this approach ensures data relevance, purposive sampling may limit generalizability. All responses were screened for completeness and validity prior to analysis.

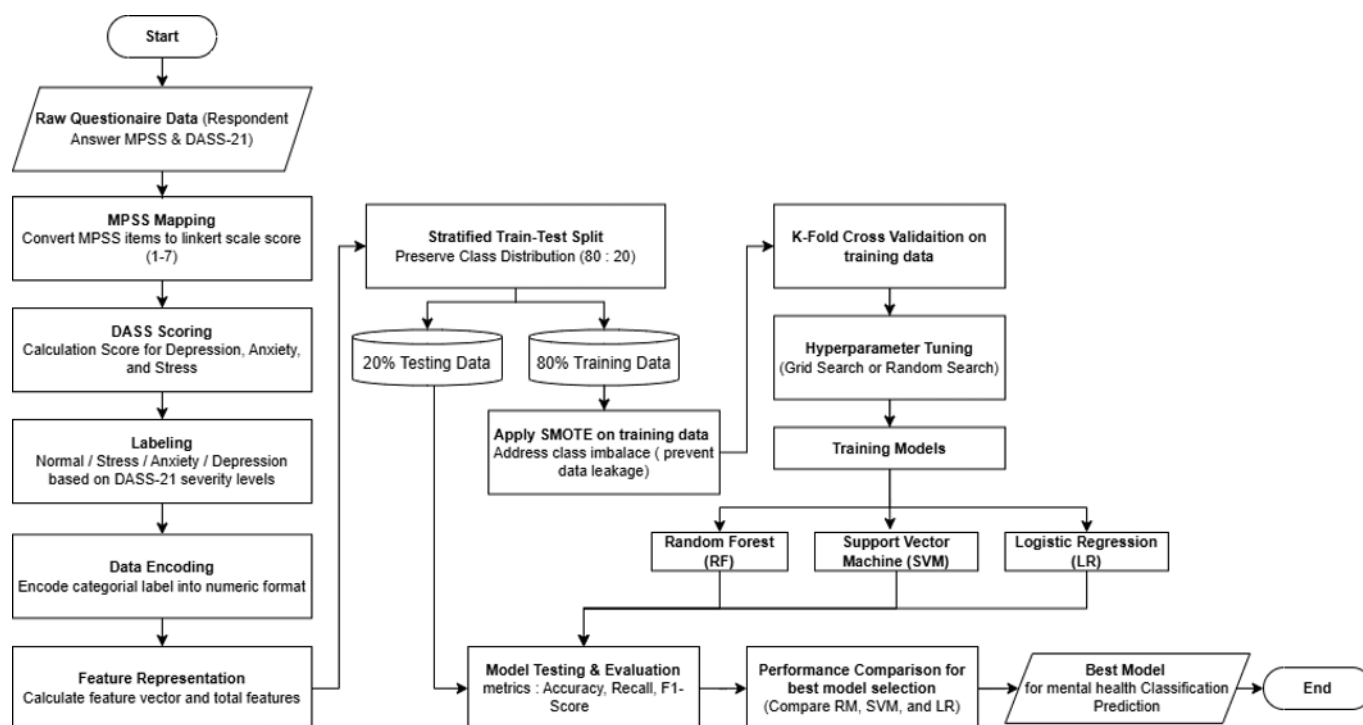


Fig. 1 Proposed workflow of multi-class mental health classification model

B. Data Preprocessing

1) *Likert and MSPSS Mapping*: Family support and social support variables were mapped into numerical form based on the Multidimensional Scale of Perceived Social Support (MSPSS) instrument with a Likert scale. Each response was converted to a numerical value for processing by the classification model.

2) *Grouping DASS-21 Items*: Items on the DASS-21 questionnaire were grouped into three dimensions, namely depression, anxiety, and stress. Each dimension consisted of seven items used to generate an initial score for each category.

3) *DASS-21 score Calculation*: The total score for each dimension is obtained by adding the values of the seven related items. The score is then multiplied by two to obtain the final score according to the DASS-21 assessment procedure using the formula in (1) [18]:

$$\text{Score}_{\text{category}} = \left(\sum_{Q=1}^n Q_i \right) \times 2 \quad (1)$$

Q denotes the score allocated to each item in the specified category, while n indicates the total number of items in the category.

4) *Labeling Mental Health Classes*: The DASS-21 labeling process is carried out by classifying the mental health condition of respondents based on three measurement dimensions, namely depression, anxiety, and stress. Respondents are categorized as Normal if all scores are within the normal range. If any score exceeds the normal limit, the system determines the severity level for each dimension and assigns a label based on the highest severity level. If there is more than one dimension with the same severity level, the label is assigned based on the higher dimension score [19].

5) *Data Encoding*: All non-numeric attributes are converted into numeric format to ensure compatibility with the machine learning algorithms used in the research.

C. Feature Representation

The feature vector X consists of six numerical variables: three social support dimensions from the MSPSS and three psychological scores from the DASS-21. The target y is the DASS Label with four classes: Normal, Stress, Anxiety, and Depression.

D. Stratified Train-Test Split (80:20)

The data is divided into training and test sets using an 80:20 ratio via a stratified split. This technique is used to maintain a consistent class proportion across both subsets and reduce potential model bias. The training data is used to train the classification model, while the test data is used to measure the model's performance on data it has never seen before.

E. Data Balancing

The classification model was trained using three algorithms, namely Random Forest, Support Vector Machine (SVM), and Logistic Regression, by using family support and social support variables as input features, while the mental health categories generated from DASS-21 labeling were used as output labels. To address class imbalance, oversampling was applied only to the training data after dataset splitting, while the test data retained its original distribution to prevent data leakage and reduce the risk of overfitting, particularly on a relatively small dataset. Furthermore, model performance was evaluated using cross-validation to assess consistency across multiple data subsets, while model complexity was controlled through parameter settings to maintain better generalization ability.

To prevent data leakage, the dataset was first divided into training and testing subsets using a stratified 80:20 split. The SMOTE balancing was then applied only to the training data, while the testing data retained its original class distribution and was not exposed to any synthetic samples. This approach helps prevent information leakage during the training process and reduces the risk of overfitting, thereby providing a more objective and reliable evaluation of the model's generalization performance.

TABLE I
DEPRESSION ANXIETY AND STRESS SCALE (DASS-21)

Label	Score				
Anxiety	0-7	8-9	10-14	15-19	20+
Depression	0-9	10-13	14-20	21-27	28+
Stress	0-14	15-18	19-25	26-33	34+

F. Classification Model Training

The classification model was trained using three algorithms, namely Random Forest, Support Vector Machine (SVM), and Logistic Regression. All three models accepted family support and social support variables as input features, while the mental health category resulting from DASS-21 labeling was used as the output label. The training process was performed on the training data after the class-balancing stage. Random Forest was selected as the primary ensemble classifier for its ability to model non-linear relationships and to reduce overfitting by aggregating multiple decision trees. A Support Vector Machine was implemented as a comparative model because of its demonstrated effectiveness at maximizing class separation in high-dimensional feature spaces. Logistic Regression was employed as the baseline model for its simplicity, computational efficiency, and interpretability in multi-class classification tasks.

To ensure a fair comparison among models, hyperparameter tuning was performed prior to training. Grid Search with cross-validation was employed to identify the optimal parameter combinations for each algorithm. For Random Forest, parameters such as the number of trees and maximum depth were optimized; for Support Vector Machine, the regularization parameter and kernel settings were adjusted; while for Logistic Regression, the regularization strength and solver configuration were tuned. This process allowed each model to be evaluated under more representative and competitive settings.

G. Cross Validation

In the data validation stage, 5-fold cross-validation was applied to all classification models to evaluate performance consistency. This process was used to compute the mean accuracy, standard deviation, and 95% confidence interval as indicators of model reliability.

H. Evaluation Model

Model performance was evaluated using four metrics, namely accuracy, recall, F1-score, and support, which were calculated for each class using a one-vs-all approach to assess the model's ability to distinguish DASS-21 label categories. The improvement in recall and F1-score for the Stress and Normal classes was considered the primary indicator of balancing success, as these classes had previously been underrepresented relative to the others. Higher recall values indicate that the model became more effective in identifying minority-class samples, while improved F1-scores

reflect a better balance between precision and recall. These findings are consistent with the primary objective of balancing: enhancing sensitivity to minority classes and reducing prediction bias toward majority classes.

1) *Accuracy*: Accuracy measures the proportion of correct predictions against all test data. In multi-class classification, accuracy is calculated by summing the number of correct predictions for each class and dividing by the total number of samples. This metric describes the model's ability to produce accurate predictions overall. Accuracy is calculated using the formula shown in (2) [20]:

$$\text{Accuracy}_i = \frac{\text{TPA} + \text{TPB} + \text{TPC}}{N} \quad (2)$$

2) *Recall (Sensitivity)*: Recall measures the model's ability to recognize samples that truly belong to a class. In this study, recall was used to evaluate the prediction performance in each mental health condition category based on the DASS-21 label. A high recall value indicates that the model can detect most samples in that class, making this metric important for analyzing performance on minority classes that are often difficult to recognize. Recall is calculated using the formula shown in (3) [20]:

$$\text{Recall}_i = \frac{\text{TP}_i}{\text{TP}_i + \text{FN}_i} \quad (3)$$

3) *F1-Score*: F1-Score is used to evaluate the balance of model performance between prediction accuracy and sensitivity in each mental health category. This metric is prioritized in the model selection process because it is more responsive to prediction errors between classes than accuracy, especially for labels with relatively fewer samples. The use of F1-Score allows for a more comprehensive assessment when class distributions are uneven. The F1-Score calculation uses the formula shown in (4) [20]:

$$\text{F1-Score}_i = \frac{2 \times \text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}} \quad (4)$$

4) *Support*: Support indicates the number of actual samples in each class in the test data. This value is used to provide context for calculating other evaluation metrics. Support is important in multi-class cases because it shows the class distribution, which can affect the sensitivity of other metrics, especially when there is class imbalance in the DASS-21 labels. The Support calculation uses the formula shown in (5) [20]:

$$\text{Support}_i = \text{TP}_i + \text{FN}_i \quad (5)$$

I. Model Comparative Analysis

The results were presented using confusion matrices and evaluation metrics for each algorithm to facilitate the interpretation of prediction outcomes. Model comparison was conducted using accuracy, recall, F1-score, and support across the three DASS-21 classes. The F1-score was used as the primary indicator because it balances precision and recall in the presence of imbalanced class distributions. In addition, the results were analyzed to identify patterns of classification error and performance consistency across classes, thereby reducing potential model bias. These findings were then used to determine the most appropriate algorithm for predicting students' mental health conditions based on family and social support.

III. RESULT AND DISCUSSION

A. Raw Questionnaire Data

Data collection yielded 450 valid respondents who met the criteria as active students, exceeding the minimum sample requirement and strengthening the dataset for classification. Although the majority of respondents were from Telkom University Purwokerto, participation from other universities ensured a more diverse dataset.

The DASS-21 results indicate that Anxiety and Depression categories dominate compared to Stress and Normal, suggesting a higher tendency of students toward these conditions. This imbalance may affect classification performance, as minority classes (Stress and Normal) have fewer samples and are more prone to misclassification.

Meanwhile, MSPSS results show variability in perceived social support across family, friends, and significant others. Higher levels of social support are generally associated with the Normal category, while lower support tends to correspond with Anxiety and Depression categories.

B. Data Preprocessing Result

The data preparation step involved converting questionnaire responses into a format the classification model could use. First, scores from the Multidimensional Scale of Perceived Social Support (MSPSS) were converted to numerical scores to measure support from family, friends, and other important people. This showed differences among the three types of social support, reflecting how people felt supported, and suggested these features might help distinguish mental health conditions.

Also, questions from the Depression Anxiety Stress Scales-21 (DASS-21) were grouped into three parts: depression, anxiety, and stress. This created three sets of

scores, one for each part. The results showed that anxiety and depression scores varied more than stress scores, suggesting these conditions were more common among the respondents.

The final DASS-21 scores were used to label mental health categories by choosing the highest severity level from the three parts. This created four categories: Depression, Anxiety, Stress, and Normal. The labels showed that Anxiety and Depression were the most common. The smaller number of Stress and Normal cases, as shown in Fig. 2, was an important feature of the dataset and made the classification task more difficult in subsequent analysis.

The final stage is the encoding process, which converts all non-numeric attributes into a numeric format so they can be processed by the classification algorithm. The encoding results ensure that all features produced are standardized and compatible as model inputs. Overall, the data pre-processing sequence produces a structured dataset that is ready for use in modeling and evaluation.

C. Feature Representation

The input data consisted of six numerical features, namely Significant Others Support Score, Family Support Score, Friends Support Score, Stress Score, Anxiety Score, and Depression Score. These variables were arranged into a feature vector $X = [SO_Score, Family_Score, Friends_Score, Stress_Score, Anxiety_Score, Depression_Score]$, where the first three features represent perceived social support dimensions from MSPSS, while the remaining three features represent psychological condition scores derived from DASS-21. The target variable y was the DASS_Label, which classifies respondents into four categories: Normal, Stress, Anxiety, and Depression.

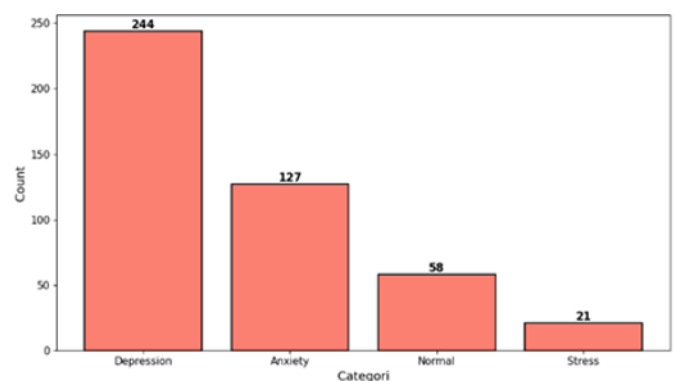


Fig. 2 DASS-21 label distribution

D. Splitting Data Result

The preprocessed dataset is then divided into two subsets using a stratified split, with 80% for training and 20% for testing. The stratification technique is used to maintain consistent label proportions across subsets, ensuring the class distribution does not change significantly relative to the initial dataset.

The results of the division show that both subsets maintain a similar class distribution, with the Anxiety and Depression classes remaining dominant, while the Stress and Normal classes are in the minority. This consistency in distribution is important to ensure that the evaluation process is not biased towards certain classes due to worsening imbalances after data division.

Numerically, the training data consists of 360 samples, and the test data consists of 90 samples from a total of 450 respondents. This pattern allows the model to learn data variations during training while providing sufficient test data to measure its generalization ability more objectively.

The data distribution at this stage reinforces the dataset's characteristic: multi-class cases with low minority counts can be challenging when the model makes predictions for Stress and Normal. Therefore, the next stage requires data handling to ensure the model does not exhibit biased performance on minority classes.

E. Balancing Data

The imbalance in label distribution in the training data caused the Stress and Normal classes to be in a

minority position compared to the Anxiety and Depression classes. This imbalance had the potential to reduce the model's sensitivity to minority classes and increase the model's tendency to predict dominant classes.

To overcome this problem, class balancing was performed using the oversampling technique. The balancing results showed that the label distribution in the training data became more proportional between classes. This change in distribution provided a more balanced opportunity for the model to learn patterns in each class, including minority classes that were previously underrepresented.

The impact of data balancing was evident in the increases in recall and F1-score for the Stress and Normal classes during evaluation. This indicates that the model can recognize minority classes more accurately after the balancing process. However, balancing does not eliminate the dominance of the majority class altogether, but it is effective enough to reduce the prediction bias that was previously heavier towards Anxiety and Depression.

F. Evaluation Model

Evaluation was conducted using four metrics: Accuracy, Recall, F1-Score, and Precision. The test results for Random Forest, SVM, and Logistic Regression are shown in Table II.

TABLE II
MODELS' PERFORMANCE

Model	Label	Precision	Recall	F1-Score
Random Forest	Anxiety	0.96	0.96	0.94
	Depression	0.98	0.98	0.98
	Stress	0.80	1.00	0.98
	Normal	1.00	1.00	1.00
	Accuracy			0.97
Logistic Regression	Anxiety	0.70	0.84	0.76
	Depression	0.90	0.78	0.84
	Stress	0.80	1.00	0.89
	Normal	0.92	1.00	0.96
	Accuracy			0.83
Support Vector Machine	Anxiety	0.67	0.80	0.73
	Depression	0.88	0.78	0.83
	Stress	1.00	0.85	0.86
	Normal	0.86	1.00	0.92
	Accuracy			0.81

The results show that Random Forest achieved the highest performance with an accuracy of 0.97 compared to the other models. Logistic Regression showed consistent performance in the Anxiety class but lower results in the Normal and Depression classes, while Support Vector Machine produced lower recall values in the Normal and Stress classes, as shown in Fig. 3. To further validate robustness, 5-fold cross-validation was conducted, where Random Forest again achieved the highest mean accuracy of 0.92 with a low standard deviation of 0.01, followed by Logistic Regression (88.22%) and Support Vector Machine (82.00%). The consistent performance between hold-out testing and cross-validation indicates that the model did not exhibit significant overfitting. These findings confirm that Random Forest remained the most reliable model.

G. Cross Validation

To evaluate model robustness, 5-fold cross-validation was conducted, where the mean accuracy ($\bar{x}$), standard deviation (s), and 95% confidence intervals were calculated based on five folds ($n = 5$) as in (6) and Table III.

$$CI = \bar{x} \pm 1.96 \left(\frac{s}{\sqrt{n}} \right) \quad (6)$$

Table III compares accuracy across previous studies and this study for three classification algorithms: Support Vector Machine, Random Forest, and Logistic Regression. In general, the results of this study indicate increased accuracy across all models compared with previous studies. Support Vector Machine and Logistic Regression each increased in accuracy from 0.69 and 0.79 to 0.81 and 0.83, respectively, while Random Forest showed the most substantial improvement, from 0.80 to 0.97. These results indicate that the proposed approach and combination of variables enhanced the classification performance of all models.

Although Random Forest achieved a hold-out accuracy of 0.97, this result should be interpreted with caution, given the relatively small dataset size. To obtain a more robust performance estimate, 5-fold cross-validation was conducted, in which Random Forest again achieved the highest mean accuracy of 0.92 with a low standard deviation of 0.01, indicating stable performance across different data partitions. In addition, oversampling was applied only to the training data to prevent data leakage. These findings indicate that Random Forest remained the most reliable model, although further validation on a larger dataset is still recommended.

Table III compares accuracy across previous studies and this study for three classification algorithms: Support Vector Machine, Random Forest, and Logistic Regression. In general, the results of this study indicate increased accuracy across all models compared with previous studies. Support Vector Machine and Logistic Regression each increased in accuracy from 0.69 and 0.79 to 0.81 and 0.83, respectively, while Random Forest showed the most substantial improvement, from 0.80 to 0.97. These results indicate that the proposed approach and combination of variables enhanced the classification performance of all models.

Although Random Forest achieved a hold-out accuracy of 0.97, this result should be interpreted with caution, given the relatively small dataset size. To obtain a more robust performance estimate, 5-fold cross-validation was conducted, in which Random Forest again achieved the highest mean accuracy of 0.92 with a low standard deviation of 0.01, indicating stable performance across different data partitions. In addition, oversampling was applied only to the training data to prevent data leakage. These findings indicate that Random Forest remained the most reliable model, although further validation on a larger dataset is still recommended.

TABLE III
CROSS-VALIDATION PERFORMANCE WITH 95% CONFIDENCE

Model	Mean Accuracy	Std Dev	95% Confidence Interval
Random Forest	0.92	0.01	0.91 – 0.93
Support Vector Machine	0.82	0.02	0.80 – 0.90
Logistic Regression	0.88	0.02	0.86 – 0.90

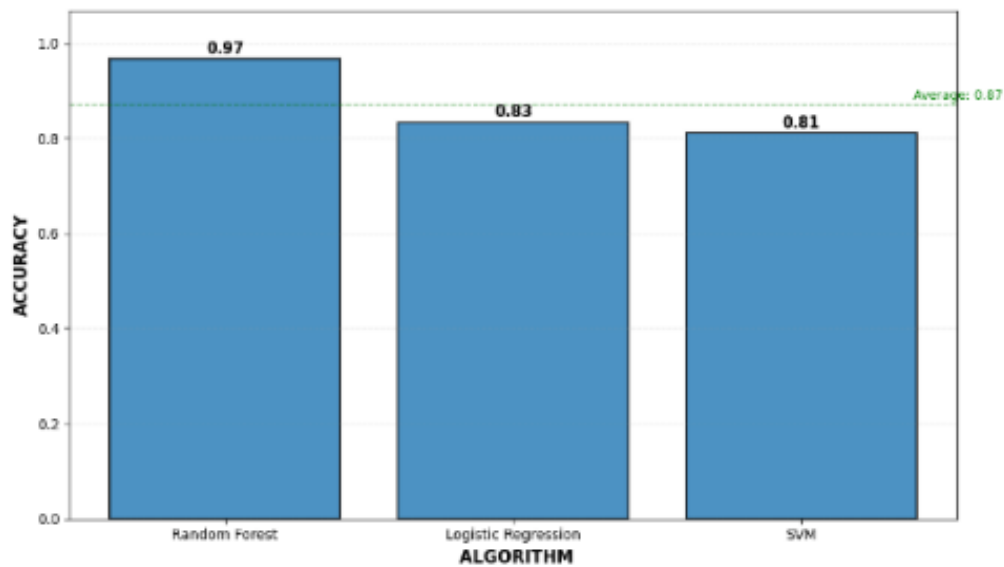


Fig. 3 Accuracy comparison of three classification algorithms

Random Forest outperformed the other models by capturing non-linear relationships and interactions among features through its ensemble mechanism, resulting in more robust predictions. Furthermore, feature importance analysis suggests that psychological variables (Stress, Anxiety, and Depression scores) play a dominant role in the classification process, as they are directly associated with the DASS-21 target labels, making them more discriminative predictors. Meanwhile, social support features contribute as complementary factors that further enhance overall model performance.

The heatmap shows that social support (MSPSS Total) is negatively correlated with the DASS-21 subscales, with the strongest correlation in the Stress variable ($r = -0.22$), followed by Anxiety ($r = -0.09$) and Depression ($r = -0.06$). This indicates that the higher the social support, the lower the tendency for students to experience depression, anxiety, and stress. In addition, the correlations between the DASS-21 subscales are fairly strong, particularly in the Anxiety–Depression ($r = 0.73$) and Anxiety–Stress ($r = 0.60$) pairs, indicating that these psychological symptoms are interrelated, as shown in Fig. 4.

Fig. 5 shows the confusion matrix of the Random Forest model. Most samples were correctly classified across all classes, with the highest accuracy in the Depression class (48 samples). Minor misclassifications occurred in the Anxiety class, where a few samples were predicted as Depression and Stress. Overall, the model demonstrates strong and consistent performance.

The results indicate that the proposed model achieves good classification performance; however, the findings should be interpreted with caution given the limited

dataset size and experimental setting. The model is intended as a decision-support tool for preliminary mental health classification rather than a definitive diagnostic system. From a technical perspective, the superior performance of Random Forest is attributed to its ensemble structure, which effectively captures feature interactions and reduces overfitting compared to single-model approaches.

TABLE IV
COMPARISON OF CLASSIFICATION MODEL
ACCURACY BETWEEN PRIOR STUDIES AND THE
PRESENT STUDY

Model	Previous Study	This Study
Random Forest	0.80	0.97
Logistic Regression	0.79	0.83
Support Vector Machine	0.69	0.81

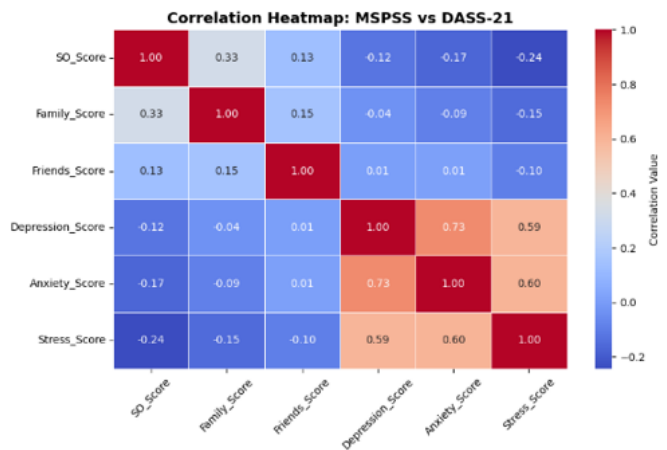


Fig. 4 Correlation heatmap between MSPSS and DASS-21

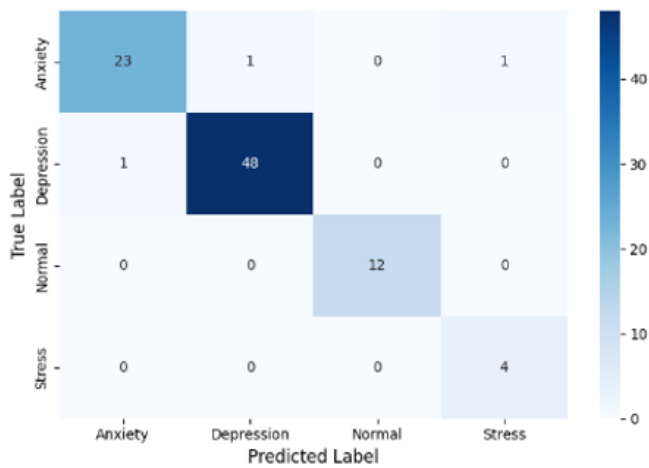


Fig. 4 Random forest confusion matrix

IV. CONCLUSION

This study shows that integrating DASS-21 and MSPSS can support multi-class classification of student mental health using machine learning. The results indicate that Random Forest achieved the best performance, with an accuracy of 0.97, outperforming Support Vector Machine and Logistic Regression, and that it maintained consistent evaluation metrics across all classes. These findings suggest that, within the current dataset, ensemble models may be more effective at capturing the relationship between social support and psychological conditions. However, the proposed model remains a preliminary approach and requires further validation on larger and more diverse datasets before practical implementation.

REFERENCES

- [1] World Health Organization, *World Mental Health Today*, First Edit. Geneva: World Health Organization, 2025. [Online]. Available: <https://iris.who.int/>
- [2] A. Suryanto and S. Nada, "Evaluation of Mental Health Analysis of College Students at the early Covid-19 Outbreak in Indonesia," *Jurnal Citizenship Virtues*, no. 2, pp. 83–97, 2021.
- [3] N. A. Kloping, T. Citraningtyas, R. Lili, S. M. Farrell, and A. Molodynski, "Mental health and wellbeing of Indonesian medical students: A regional comparison study," *International Journal of Social Psychiatry*, vol. 68, no. 6, pp. 1295–1299, 2022, doi: 10.1177/002076402111057732.
- [4] S. Howlader, S. Abedin, and M. M. Rahman, "Social support, distress, stress, anxiety, and depression as predictors of suicidal thoughts among selected university students in Bangladesh," *PLOS Global Public Health*, vol. 4, no. 4, pp. 1–13, 2024, doi: 10.1371/journal.pgph.0002924.
- [5] F. M. Ekawati, "The health and wellbeing of undergraduate students in Indonesia: descriptive results of a survey in three public universities," *Sci. Rep.*, vol. 15, no. 1, pp. 1–21, 2025, doi: 10.1038/s41598-025-90527-w.
- [6] B. Lee and Y. E. Kim, "Validity of the depression, anxiety, and stress scale (DASS-21) in a sample of Korean university students," *Current Psychology*, vol. 41, no. 6, pp. 3937–3946, 2022, doi: 10.1007/s12144-020-00914-x.
- [7] T. Wahyono and Y. Heryadi, *Machine Learning Konsep dan Implementasi*. Yogyakarta: Gava Media, 2020.
- [8] A. El Attaoui, Y. Koulou, and N. El Hami, "At the Core of Artificial Intelligence: Leveraging Machine Learning through Random Forest," in *Methods and Applications of Artificial Intelligence: Dynamic Response, Learning, Random Forest, Linear Regression, Interoperability, Additive Manufacturing and Mechatronics: Volume 2*, vol. 2, ENSA-Kénitra, Université Ibn Tofail, Morocco: wiley, 2025, pp. 77–119. doi: 10.1002/9781394351817.ch5.
- [9] V. Bruni, F. Pelosi, and D. Vitulano, "Integrating subdivision schemes into SVM for improved signal classification," *J. Comput. Appl. Math.*, vol. 477, p. 117142, 2026, doi: 10.1016/j.cam.2025.117142.
- [10] R. Nodir and K. Dilmurod, "The Mathematical Essence of Logistic Regression for Machine Learning," no. 102, pp. 102–105, 2022.
- [11] U. Braga-Neto, *Fundamentals of Pattern Recognition and Machine Learning*. Cham: Springer, 2020. doi: 10.1007/978-3-030-27656-0.
- [12] M. Kunta Biddinika, A. Masitha, and V. Arfiana Nurul Fatimah, "Machine Learning Techniques for Heart Disease Prediction Using a Multi-Algorithm Approach," *JUITA*, vol. 12, no. 2, pp. 149–158, Nov. 2024.
- [13] W. P. Indahwati and F. M. Afendi, "Improving Stroke Detection with Hybrid Sampling and Cascade Generalization," *JUITA*, vol. 12, no. 1, pp. 9–18, 2024.
- [14] H. Herlinda, M. Itqan Mazdadi, D. Kartini, and I. Budiman, "Implementation of Particle Swarm Optimization on Sentiment Analysis of Cyberbullying using Random Forest," *JUITA*, vol. 11, no. 2, pp. 301–309, 2023.
- [15] A. Handayanto, K. Latifa, N. D. Saputro, and R. R. Waliyansyah, "Analysis and Application of Algorithm Support Vector Machine (SVM) in Data Mining to Support Promotional Strategies," vol. 7, no. 2, pp. 71–79, Nov. 2019.
- [16] X. Jia and W. Chu, "Analysis and Prevention Strategies of the Causes of College Students' Mental Health Problems," *Academic Journal of Management and*

- Social Sciences*, vol. 5, no. 3, pp. 96–98, 2023, doi: 10.54097/1or428w3.
- [17] A. H. Putra and M. Mudjiran, “*Mental Health Disorders in Students University as a Challenge in Higher Education Era Society 5.0*,” *International Journal of Applied Counseling and Social Sciences*, vol. 4, no. 2, pp. 73–80, 2023, doi: 10.24036/005966ijaccs.
- [18] World Rugby, “*Mindset-A Mental Health Resource for Team Doctors*,” 2009.
- [19] G. D. Zimet, N. W. Dahlem, S. G. Zimet, and G. K. Farley, “*The Multidimensional Scale of Perceived Social Support*,” *J. Pers. Assess.*, vol. 52, no. 1, pp. 30–41, 1988, doi: 10.1207/s15327752jpa5201_2.
- [20] M. Fahmy Amin, “*Confusion Matrix in Three-class Classification Problems: A Step-by-Step Tutorial*,” *Journal of Engineering Research*, vol. 7, no. 1, pp. 3–5, 2023, doi: 10.21608/erjeng.2023.296718.