

Implementation of the Random Forest Algorithm for Classifying University Students' Stress Levels Based on DASS-21 Questionnaire Data

[Implementasi Random Forest untuk Klasifikasi Tingkat Stres Mahasiswa Berdasarkan Data Kuesioner DASS-21]

Syariful Rifky Bachtiar¹⁾, Mochamad Alfian Rosid^{*2)}, Suprianto³⁾, Nuril Lutvi Azizah¹⁾

^{1,2,3,4)}Program Studi Informatika, Universitas Muhammadiyah Sidoarjo, Indonesia

*Email Penulis Korespondensi: alfanrosid@umsida.ac.id

Abstract. Stress is a common psychological problem among students caused by academic pressure, achievement demands, and personal and social factors. Rapid and accurate identification of stress levels is essential to support timely interventions, yet manual assessment is often time-consuming and prone to errors. This study implements the Random Forest algorithm to classify student stress levels using DASS-21 questionnaire data. The dataset consists of 1,000 respondents and underwent preprocessing, including cleaning, scoring, labeling, and encoding. The data were divided into training and testing sets using an 80:20 ratio. Model training was performed using the Random Forest algorithm with hyperparameter optimization through GridSearchCV. The optimal parameters were $n_estimators = 200$, $max_depth = 10$, $min_samples_split = 2$, and $min_samples_leaf = 1$. The model achieved an accuracy of 73.50%, precision of 74.17%, recall of 73.50%, and an F1-score of 73.61%, demonstrating reliable performance in classifying student stress levels.

Keywords - Random Forest; classification; student stress; DASS-21; machine learning

Abstrak. Stres merupakan salah satu masalah psikologis yang banyak dialami mahasiswa akibat tekanan akademik, tuntutan prestasi, serta berbagai faktor pribadi dan sosial. Identifikasi tingkat stres secara cepat dan akurat diperlukan untuk mendukung penanganan yang lebih efektif, tetapi penilaian secara manual memerlukan waktu dan rentan terhadap kesalahan. Penelitian ini bertujuan mengimplementasikan algoritma Random Forest untuk mengklasifikasikan tingkat stres mahasiswa berdasarkan data kuesioner DASS-21. Dataset terdiri atas 1.000 data responden yang telah melalui tahapan preprocessing meliputi cleaning, scoring, labeling, dan encoding. Data dibagi menjadi data latih dan data uji dengan rasio 80:20. Pelatihan model dilakukan menggunakan algoritma Random Forest dengan optimasi parameter melalui metode GridSearchCV. Parameter terbaik adalah $n_estimators = 200$, $max_depth = 10$, $min_samples_split = 2$, dan $min_samples_leaf = 1$. Model menghasilkan akurasi sebesar 73,50%, precision sebesar 74,17%, recall sebesar 73,50%, dan F1-score sebesar 73,61%, menunjukkan kinerja yang baik dalam mengklasifikasikan tingkat stres mahasiswa.

Kata Kunci - Random Forest; classification; student stress; DASS-21; machine learning

I. PENDAHULUAN

Stres merupakan salah satu permasalahan psikologis yang banyak dialami mahasiswa selama menjalani perkuliahan. Tekanan akademik, beban tugas, tuntutan prestasi, serta berbagai faktor internal dan eksternal dapat memengaruhi kondisi psikologis mahasiswa sehingga menyebabkan tingkat stres yang berbeda pada setiap individu. Faktor internal meliputi efikasi diri, optimisme, motivasi belajar, dan prokrastinasi akademik, sedangkan faktor eksternal meliputi tekanan sosial dan kompleksitas perkuliahan [1]. Apabila tidak ditangani dengan baik, stres dapat memberikan dampak negatif terhadap kesehatan mental maupun performa akademik mahasiswa.

Fenomena stres pada mahasiswa juga ditemukan di lingkungan Universitas Muhammadiyah Sidoarjo. Penelitian Musikhah dan Nastiti menunjukkan bahwa 66% mahasiswa Universitas Muhammadiyah Sidoarjo yang menjalani perkuliahan sambil bekerja mengalami stres akademik pada tingkat sedang [2]. Temuan tersebut menunjukkan bahwa stres merupakan permasalahan yang perlu mendapat perhatian karena dapat memengaruhi proses pembelajaran dan kesejahteraan psikologis mahasiswa. Oleh karena itu, diperlukan metode yang mampu membantu identifikasi tingkat stres secara cepat, objektif, dan akurat.

Salah satu instrumen yang banyak digunakan untuk mengukur tingkat stres adalah Depression Anxiety Stress Scale-21 (DASS-21) yang dikembangkan oleh Lovibond dan Lovibond. Instrumen ini telah digunakan secara luas dan terbukti memiliki validitas serta reliabilitas yang baik dalam mengukur kondisi psikologis individu [3]. Selain itu, DASS-21 juga menunjukkan hasil yang baik dalam menilai tekanan psikologis mahasiswa pada berbagai konteks penelitian, termasuk selama masa pandemi COVID-19 [4]. Pada subskala stres, DASS-21 terdiri atas tujuh item

pernyataan yang digunakan untuk mengukur tingkat stres responden dan mengelompokkannya ke dalam kategori normal, ringan, sedang, berat, dan sangat berat.

Meskipun demikian, proses pengolahan data DASS-21 secara manual menjadi kurang efisien ketika jumlah responden yang dianalisis cukup besar. Proses perhitungan skor, pengelompokan kategori, serta analisis data membutuhkan waktu yang relatif lama dan berpotensi menimbulkan kesalahan. Oleh karena itu, diperlukan suatu pendekatan otomatis yang mampu melakukan klasifikasi tingkat stres secara lebih efektif dan konsisten melalui penerapan teknik machine learning yang telah banyak digunakan dalam penelitian prediksi dan klasifikasi tingkat stres mahasiswa [5], [6].

Perkembangan teknologi machine learning memberikan peluang untuk mengatasi permasalahan tersebut. Salah satu algoritma yang banyak digunakan pada tugas klasifikasi adalah Random Forest, yaitu metode ensemble learning yang menggabungkan hasil prediksi dari sejumlah pohon keputusan untuk menghasilkan klasifikasi yang lebih stabil dan akurat [6]. Berbagai penelitian menunjukkan bahwa Random Forest memiliki performa yang baik dalam berbagai kasus klasifikasi karena mampu mengurangi overfitting, menangani data dengan banyak atribut, serta menghasilkan tingkat akurasi yang tinggi [7], [8], [9], [10].

Random Forest merupakan salah satu algoritma *machine learning* yang termasuk dalam metode *ensemble learning* berbasis kumpulan *decision tree*. Algoritma ini bekerja dengan membangun sejumlah pohon keputusan dari subset data yang berbeda, kemudian menghasilkan prediksi akhir melalui mekanisme *majority voting*. *Random Forest* memiliki keunggulan dalam menangani data berdimensi besar, mengurangi risiko *overfitting*, tahan terhadap *noise*, serta mampu menghasilkan performa klasifikasi yang baik pada berbagai permasalahan. Selain itu, algoritma ini juga dapat memberikan informasi mengenai tingkat kepentingan setiap fitur (*feature importance*) sehingga hasil klasifikasi lebih mudah diinterpretasikan [11]. Namun demikian, *Random Forest* memiliki kelemahan berupa kebutuhan komputasi yang lebih tinggi ketika jumlah pohon semakin banyak serta performanya dipengaruhi oleh pemilihan *hyperparameter* [12]. Oleh karena itu, algoritma ini banyak digunakan dalam berbagai bidang penelitian, termasuk kesehatan, pendidikan, dan prediksi perilaku manusia [13].

Penelitian terkait penerapan *Random Forest* telah banyak dilakukan pada berbagai bidang, termasuk kesehatan mental. Penelitian yang dilakukan oleh Ririn et al. [14] menunjukkan bahwa algoritma *Random Forest* mampu memprediksi kondisi kesehatan mental dengan akurasi sebesar 93,8%. Hasil penelitian tersebut menunjukkan bahwa *Random Forest* efektif dalam menangani data yang kompleks dan beragam untuk menghasilkan prediksi yang akurat.

Sementara itu, Astari et al. [15] menerapkan algoritma *Random Forest* untuk mengklasifikasikan tingkat stres saat tidur dan memperoleh akurasi sebesar 93,65%. Hasil penelitian menunjukkan bahwa model mampu mengklasifikasikan tingkat stres ke dalam lima kategori dengan performa yang baik, sehingga *Random Forest* dinilai efektif dalam menangani permasalahan klasifikasi tingkat stres.

Pada bidang pendidikan, Oon Wira Yuda et al. [16] menerapkan *Random Forest* untuk klasifikasi kelulusan mahasiswa dan memperoleh akurasi sebesar 98%. Sementara itu, Adriansyah et al. [8] membandingkan *Random Forest* dan *Support Vector Machine* (SVM) dalam klasifikasi kemampuan adaptasi siswa pada pembelajaran jarak jauh, di mana *Random Forest* menghasilkan akurasi yang lebih tinggi yaitu sebesar 91,5%.

Beberapa penelitian telah menerapkan *Random Forest* untuk klasifikasi tingkat stres maupun kesehatan mental. Rois et al. [5] memperoleh akurasi 89,72% pada prediksi tingkat stres mahasiswa menggunakan faktor fisiologis dan gaya hidup, sedangkan Astari et al. [15] memperoleh akurasi 93,65% pada klasifikasi tingkat stres saat tidur. Ringkasan penelitian terdahulu ditunjukkan pada Tabel 1.

Tabel 1. Perbandingan Penelitian Terdahulu

Penelitian	Metode	Hasil	Pembeda dengan Penelitian Lain
Rois et al. [5]	Decision Tree, Random Forest, SVM, Logistic Regression	Random Forest memberikan performa terbaik dengan akurasi 89,72%	Menggunakan faktor fisiologis dan gaya hidup, bukan instrumen DASS-21.
Astari et al. [15]	Random Forest	Random Forest memperoleh akurasi klasifikasi tingkat stres sebesar 93,65%.	Berfokus pada klasifikasi tingkat stres saat tidur, bukan mahasiswa.
Penelitian ini	Random Forest	Mengevaluasi model menggunakan accuracy, precision, recall, dan F1-score.	Menggunakan subskala stres DASS-21 pada mahasiswa Universitas Muhammadiyah Sidoarjo.

Berdasarkan berbagai penelitian terdahulu, *Random Forest* menunjukkan kemampuan yang baik dalam menangani permasalahan klasifikasi, termasuk pada bidang kesehatan mental dan tingkat stres, dengan tingkat akurasi yang tinggi [17]. Penelitian klasifikasi kesehatan mental dan tingkat stres mahasiswa telah dilakukan menggunakan berbagai algoritma *machine learning*, seperti *Naïve Bayes*, *K-Nearest Neighbor* (K-NN), *Decision Tree*, dan *Support Vector Machine* (SVM) [18], [19]. Meskipun demikian, penerapan *Random Forest* untuk klasifikasi tingkat stres mahasiswa berdasarkan data kuesioner DASS-21 pada konteks mahasiswa di Indonesia masih relatif terbatas. Pada penelitian ini, *GridSearchCV* digunakan untuk memperoleh kombinasi *hyperparameter* yang optimal melalui proses pencarian parameter secara sistematis menggunakan *cross-validation*, sehingga parameter model diperoleh secara objektif dan mampu meningkatkan performa klasifikasi [20]. Meskipun algoritma seperti *Support Vector Machine* (SVM) dan *Extreme Gradient Boosting* (XGBoost) juga banyak digunakan dalam tugas klasifikasi, penelitian ini difokuskan pada implementasi dan evaluasi performa *Random Forest* sesuai dengan tujuan penelitian, sehingga algoritma lain tidak menjadi bagian dari evaluasi karena berada di luar ruang lingkup penelitian.

Temuan dari penelitian terdahulu menunjukkan masih terdapat beberapa celah penelitian. Pertama, penerapan *Random Forest* untuk klasifikasi tingkat stres menggunakan subskala stres DASS-21 pada mahasiswa Indonesia masih terbatas. Kedua, belum banyak penelitian yang menggunakan dataset mahasiswa Universitas Muhammadiyah Sidoarjo. Ketiga, optimasi *hyperparameter Random Forest* menggunakan *GridSearchCV* pada kasus ini masih jarang dikaji. Oleh karena itu, penelitian ini dilakukan untuk mengisi kesenjangan tersebut.

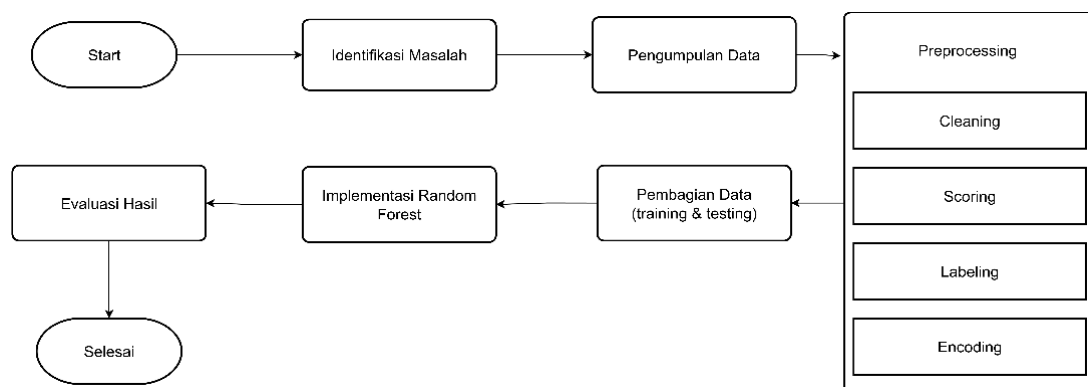
Sebagai upaya menjawab permasalahan tersebut, penelitian ini bertujuan mengimplementasikan algoritma *Random Forest* dengan optimasi *hyperparameter* menggunakan *GridSearchCV* untuk mengklasifikasikan tingkat stres mahasiswa berdasarkan data kuesioner DASS-21. Model yang dibangun diharapkan mampu mengidentifikasi tingkat stres mahasiswa secara otomatis ke dalam kategori normal, ringan, sedang, berat, dan sangat berat. Selanjutnya, performa model dievaluasi menggunakan metrik *accuracy*, *precision*, *recall*, dan *F1-score* untuk mengetahui kemampuan *Random Forest* dalam mengklasifikasikan tingkat stres mahasiswa secara akurat [21].

II. METODE

Pendekatan Penelitian

Penelitian ini menggunakan pendekatan kuantitatif dengan metode klasifikasi berbasis machine learning. Tujuan penelitian adalah membangun model klasifikasi tingkat stres mahasiswa menggunakan algoritma *Random Forest* berdasarkan data hasil pengisian kuesioner *Depression Anxiety Stress Scale-21* (DASS-21).

Tahapan Penelitian



Gambar 1. Tahapan Penelitian

Alur penelitian pada Gambar 1 diawali dengan tahap identifikasi masalah, yaitu mengidentifikasi permasalahan terkait proses klasifikasi tingkat stres mahasiswa yang masih dilakukan secara manual sehingga kurang efisien ketika jumlah responden cukup besar. Selanjutnya dilakukan pengumpulan data menggunakan kuesioner *DASS-21* yang disebarkan kepada mahasiswa aktif Universitas Muhammadiyah Sidoarjo. Data yang diperoleh kemudian melalui tahap *preprocessing* yang meliputi *cleaning*, *scoring*, *labeling*, dan *encoding*. Dataset selanjutnya dibagi menjadi data latih (*training set*) dan data uji (*testing set*) menggunakan metode *train-test split*. Data latih digunakan untuk membangun model klasifikasi menggunakan algoritma *Random Forest*, sedangkan data uji digunakan untuk mengevaluasi performa model. Tahap evaluasi dilakukan menggunakan metrik *accuracy*, *precision*, *recall*, dan *F1-score*, serta *confusion matrix* untuk menganalisis hasil klasifikasi pada setiap kategori tingkat stres.

Pengumpulan Data

Data penelitian diperoleh melalui penyebaran kuesioner DASS-21 secara daring menggunakan Google Form kepada mahasiswa aktif Universitas Muhammadiyah Sidoarjo dari berbagai program studi. Sebanyak 1.000 respons yang memenuhi kriteria penelitian digunakan sebagai dataset dalam penelitian ini. Teknik pengambilan sampel menggunakan convenience sampling, yaitu pemilihan responden berdasarkan kemudahan akses dan kesediaan responden untuk berpartisipasi dalam penelitian [22]. Teknik ini dipilih karena kuesioner disebarkan secara daring hingga diperoleh 1.000 respons yang memenuhi kriteria penelitian. Penelitian ini hanya menggunakan subskala stres yang terdiri atas tujuh item pertanyaan. Setiap item dinilai menggunakan skala Likert empat poin sebagaimana ditunjukkan pada Tabel 2.

Tabel 2. Skala Respon DASS-21

No	Keterangan
0	Tidak Pernah
1	Kadang-Kadang
2	Cukup Sering
3	Sangat Sering

Subskala stres pada instrumen DASS-21 terdiri dari sejumlah item pernyataan yang digunakan untuk mengukur tingkat stres responden. Tujuh item pertanyaan yang digunakan sebagai fitur penelitian ditunjukkan pada Tabel 3.

Tabel 3. Item subskala Stres DASS-21

No	Keterangan
1	Saya merasa sulit beristirahat
2	Saya cenderung bereaksi berlebihan terhadap suatu situasi
3	Saya merasa menggunakan banyak energi karena gugup
4	Saya merasa mudah gelisah
5	Saya merasa sulit bersantai
6	Saya menemukan diri saya menjadi tidak sabar ketika mengalami penundaan
7	Saya merasa bahwa saya mudah tersinggung

Preprocessing Data

Tahap preprocessing dilakukan untuk memastikan kualitas data sebelum digunakan dalam proses klasifikasi. Proses ini meliputi pembersihan data (cleaning), perhitungan skor (scoring), pelabelan kategori (labeling), dan pengubahan label ke bentuk numerik (encoding). Pada tahap cleaning, data diperiksa untuk mengidentifikasi data kosong (missing value), data duplikat, maupun nilai yang berada di luar rentang skala DASS-21. Selanjutnya dilakukan scoring dengan menjumlahkan skor dari tujuh item pertanyaan dan mengalikannya dengan dua sesuai pedoman DASS-21. Tahap labeling dilakukan dengan mengelompokkan skor total ke dalam kategori tingkat stres sesuai pedoman interpretasi DASS-21 [23]. Kategori tingkat stres yang digunakan dalam penelitian ini ditunjukkan pada Tabel 4.

Tabel 4. Kategori Tingkat Stres DASS-21

Kategori	Rentang Skor
Normal	0-14

Ringan	15-18
Sedang	19-25
Berat	26-33
Sangat Berat	≥ 34

Berdasarkan Tabel 4, skor total responden dikelompokkan ke dalam lima kategori tingkat stres yang selanjutnya digunakan sebagai label kelas (class label) pada proses klasifikasi menggunakan algoritma Random Forest. Selanjutnya kategori tingkat stres diubah menjadi bentuk numerik menggunakan teknik label encoding seperti pada Tabel 5.

Tabel 5. Label Encoding Tingkat Stres

Kategori	Rentang Skor
Normal	0
Ringan	1
Sedang	2
Berat	3
Sangat Berat	4

Pembagian Dataset

Dataset dibagi menggunakan metode train-test split dengan rasio 80:20, yaitu 80% sebagai data latih dan 20% sebagai data uji. Rasio ini dipilih karena memberikan keseimbangan antara jumlah data yang digunakan untuk pelatihan model dan jumlah data untuk pengujian, sehingga model memiliki data yang cukup untuk mempelajari pola sekaligus dapat dievaluasi kemampuan generalisasinya secara memadai [24]. Meskipun terdapat ketidakseimbangan data, penelitian ini memilih untuk menggunakan dataset asli untuk melihat performa model pada kondisi nyata.

Pembangunan Model Random Forest

Pada tahap ini, algoritma Random Forest dilatih menggunakan data latih yang telah disiapkan sebelumnya. Proses pelatihan melibatkan beberapa parameter penting, yaitu `n_estimators` untuk menentukan jumlah pohon keputusan dalam ensemble, `max_depth` untuk membatasi kedalaman pohon guna mengurangi risiko overfitting, serta `min_samples_split` dan `min_samples_leaf` yang mengatur jumlah minimum sampel pada proses pembentukan node. Untuk memperoleh kombinasi parameter yang optimal, dilakukan hyperparameter tuning menggunakan GridSearchCV. Metode ini bekerja dengan menguji berbagai kombinasi parameter melalui proses cross-validation untuk menemukan konfigurasi yang menghasilkan performa terbaik. Model yang telah dilatih menggunakan parameter optimal kemudian digunakan pada tahap evaluasi untuk mengukur performa algoritma Random Forest dalam mengklasifikasikan tingkat stres mahasiswa berdasarkan data kuesioner DASS-21.

Evaluasi Model

Performa model dievaluasi menggunakan data uji dengan metrik evaluasi berupa accuracy, precision, recall, dan F1-score. Selain itu, digunakan confusion matrix untuk menganalisis hasil klasifikasi pada setiap kategori tingkat stres.

1) Akurasi (Accuracy)

Menghitung rasio antara jumlah prediksi yang tepat dengan keseluruhan jumlah data yang diuji. Akurasi memberikan gambaran umum kinerja model dalam keseluruhan klasifikasi.

$$\text{Accuracy} = \frac{TP+TN}{TP+TN+FP+FN} \quad (1)$$

2) Precision

Menggambarkan ketepatan model dalam mengklasifikasikan suatu kelas tertentu, yaitu seberapa banyak prediksi positif yang benar dari seluruh prediksi positif yang dilakukan.

$$\text{Precision} = \frac{TP}{TP+FP} \quad (2)$$

3) Recall (Sensitivity)

Mengukur sejauh mana model dapat menangkap semua data aktual dari setiap kelas. Metrik ini penting untuk mengidentifikasi apakah model mampu menemukan semua kasus yang benar.

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

4) F1-Score

Rata-rata gabungan dari precision dan recall yang menunjukkan keseimbangan antara keduanya. F1-Score sangat berguna ketika terjadi ketidakseimbangan antar kelas dalam data.

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

Penelitian ini menggunakan Confusion Matrix sebagai alat bantu analisis untuk mengevaluasi hasil klasifikasi model. Confusion Matrix digunakan untuk menggambarkan jumlah prediksi yang benar maupun salah pada setiap kategori tingkat stres, yaitu Normal, Ringan, Sedang, Berat, dan Sangat Berat. Matriks ini menyajikan informasi mengenai:

1. True Positive (TP), yaitu prediksi benar untuk kelas tertentu.
2. False Positive (FP), yaitu data diprediksi termasuk dalam suatu kelas, padahal sebenarnya tidak.
3. False Negative (FN), yaitu data sebenarnya termasuk dalam suatu kelas, tetapi diprediksi berbeda.
4. True Negative (TN), yaitu prediksi benar untuk kelas selain yang diamati.

Melalui Confusion Matrix, peneliti dapat melakukan evaluasi lebih mendalam terhadap kesalahan klasifikasi yang terjadi serta mengidentifikasi kelas yang paling sering keliru diklasifikasikan. Selain menggunakan metrik evaluasi seperti accuracy, precision, recall, F1-score, dan confusion matrix, penelitian ini juga melakukan analisis feature importance untuk mengetahui kontribusi masing-masing variabel terhadap hasil klasifikasi. Feature importance pada algoritma Random Forest digunakan untuk mengukur tingkat pengaruh setiap fitur dalam proses pengambilan keputusan model. Nilai importance dihitung berdasarkan besarnya kontribusi suatu fitur dalam menurunkan impurity pada setiap pohon keputusan.

III. HASIL DAN PEMBAHASAN

A. Preprocessing Data

Data penelitian diperoleh melalui penyebaran kuesioner Depression Anxiety Stress Scale-21 (DASS-21) kepada mahasiswa aktif Universitas Muhammadiyah Sidoarjo. Sebanyak 1.000 data responden digunakan sebagai dataset penelitian. Sebelum proses klasifikasi, dilakukan tahapan preprocessing yang meliputi cleaning, scoring, labeling, dan encoding. Hasil data cleaning menunjukkan bahwa seluruh atribut S1–S7 yang berasal dari subskala stres DASS-21 tidak memiliki missing values (0%), sehingga seluruh data dapat digunakan pada proses pelatihan model tanpa penghapusan maupun imputasi. Kondisi ini menunjukkan bahwa kualitas data yang digunakan telah memadai dan meminimalkan potensi bias akibat data yang tidak lengkap. Selanjutnya, dilakukan scoring dengan menjumlahkan skor tujuh item subskala stres DASS-21 dan mengalikannya dengan faktor dua sesuai pedoman DASS-21. Skor akhir kemudian digunakan pada tahap labeling untuk mengelompokkan responden ke dalam lima kategori tingkat stres, yaitu Normal, Ringan, Sedang, Berat, dan Sangat Berat. Distribusi data pada setiap kategori ditunjukkan pada Tabel 6.

Tabel 6. Distribusi Kategori Tingkat Stres

Kategori	Rentang Skor
Normal	261
Ringan	231
Sedang	219
Berat	191
Sangat Berat	98

Berdasarkan Tabel 6, kategori Normal memiliki jumlah data terbanyak yaitu 261 responden (26,1%), diikuti kategori Ringan sebanyak 231 responden (23,1%), kategori Sedang sebanyak 219 responden (21,9%), dan kategori Berat sebanyak 191 responden (19,1%). Sementara itu, kategori Sangat Berat memiliki jumlah data paling sedikit yaitu 98 responden (9,8%).

Distribusi tersebut menunjukkan bahwa kondisi stres mahasiswa dalam dataset penelitian cukup beragam dan mencakup seluruh kategori tingkat stres yang terdapat pada DASS-21. Meskipun jumlah data pada setiap kategori tidak sepenuhnya seimbang, seluruh kelas memiliki representasi yang cukup untuk digunakan dalam proses pelatihan model klasifikasi. Keberadaan data pada setiap kategori memungkinkan algoritma Random Forest mempelajari karakteristik masing-masing tingkat stres sehingga dapat melakukan klasifikasi dengan lebih baik.

Tahap terakhir pada proses preprocessing adalah encoding. Pada tahap ini, kategori tingkat stres diubah ke dalam bentuk numerik menggunakan teknik label encoding agar dapat diproses oleh algoritma Random Forest. Kategori Normal dikodekan sebagai 0, Ringan sebagai 1, Sedang sebagai 2, Berat sebagai 3, dan Sangat Berat sebagai 4. Hasil encoding menghasilkan variabel target dalam bentuk numerik yang selanjutnya digunakan pada proses pelatihan dan pengujian model.

B. Pembagian Data

Setelah seluruh tahapan preprocessing selesai dilakukan, dataset dibagi menjadi data latih (training set) dan data uji (testing set) menggunakan metode train-test split dengan rasio 80:20. Dari total 1.000 data responden, sebanyak 800 data digunakan sebagai data latih dan 200 data sebagai data uji. Data latih digunakan untuk membangun model Random Forest dan mempelajari pola hubungan antara atribut masukan dengan kategori tingkat stres, sedangkan data uji digunakan untuk mengevaluasi performa model terhadap data yang tidak dilibatkan selama proses pelatihan. Rasio 80:20 dipilih karena memberikan jumlah data latih yang cukup untuk membentuk model, sekaligus menyediakan data uji yang memadai untuk mengevaluasi kemampuan generalisasi model secara objektif terhadap data baru.

C. Pelatihan Model Random Forest

Pelatihan model dilakukan menggunakan algoritma Random Forest dengan memanfaatkan 800 data latih yang diperoleh dari tahap pembagian dataset. Untuk memperoleh performa model yang optimal, dilakukan proses hyperparameter tuning menggunakan metode GridSearchCV dengan validasi silang (cross-validation) sebanyak 5 lipatan (5-fold cross validation). Parameter yang diuji dalam proses hyperparameter tuning meliputi jumlah pohon keputusan (*n_estimators*), kedalaman maksimum pohon (*max_depth*), jumlah minimum sampel untuk melakukan pemisahan node (*min_samples_split*), dan jumlah minimum sampel pada node daun (*min_samples_leaf*). Kombinasi parameter yang diuji ditunjukkan pada Tabel 7.

Tabel 7. Parameter Hyperparameter Tuning

Parameter	Nilai yang Diuji
<i>n_estimators</i>	100, 200, 300
<i>max_depth</i>	None, 10, 20, 30
<i>min_samples_split</i>	2, 5, 10

Parameter	Nilai yang Diuji
<i>min_samples_leaf</i>	1, 2, 4

Kombinasi nilai dari keempat parameter tersebut menghasilkan sebanyak 108 kombinasi hyperparameter yang dievaluasi menggunakan metode GridSearchCV. Berdasarkan proses evaluasi tersebut, diperoleh beberapa kombinasi parameter dengan nilai mean test score yang berbeda. Ringkasan hasil hyperparameter tuning ditunjukkan pada Tabel 8.

Tabel 8. Hasil Hyperparameter Tuning Menggunakan GridSearchCV

No	Rank	Mean Test Score	Parameter			
			<i>n estimators</i>	<i>max depth</i>	<i>min samples split</i>	<i>min samples leaf</i>
1	1	0,78000	200	10	2	1
2	2	0,77750	200	10	2	2
3	3	0,77625	200	10	5	1
4	4	0,77500	200	<i>None</i>	2	2
5	4	0,77500	200	30	2	2
...						
108	105	0,74375	300	30	10	4

Berdasarkan Tabel 8, kombinasi hyperparameter terbaik yang diperoleh melalui proses GridSearchCV terdiri atas *n_estimators* sebesar 200, *max_depth* sebesar 10, *min_samples_split* sebesar 2, dan *min_samples_leaf* sebesar 1 dengan mean test score sebesar 0,78000. Nilai tersebut menunjukkan bahwa kombinasi parameter tersebut memberikan performa validasi tertinggi dibandingkan kombinasi lainnya, sehingga dipilih sebagai parameter optimal untuk membangun model Random Forest. Selanjutnya, untuk memastikan bahwa performa model tidak hanya bergantung pada satu pembagian data latih, dilakukan validasi menggunakan metode 5-fold cross-validation. Hasil validasi ditunjukkan pada Tabel 9.

Tabel 9. Hasil Validasi Menggunakan 5-Fold Cross Validation

Fold	Accuracy (%)
1	76,88
2	75,00
3	80,00
4	73,75
5	75,62
Rata-rata ± Standar Deviasi	76,25 ± 2,13

Berdasarkan Tabel 9, hasil validasi menggunakan metode 5-fold cross-validation menunjukkan bahwa model Random Forest memperoleh rata-rata akurasi sebesar 76,25% dengan standar deviasi sebesar 2,13%. Nilai standar deviasi yang relatif kecil menunjukkan bahwa performa model konsisten pada setiap fold, sehingga proses pelatihan tidak dipengaruhi secara signifikan oleh variasi pembagian data. Hasil ini mengindikasikan bahwa kombinasi

hyperparameter yang diperoleh melalui GridSearchCV mampu menghasilkan model yang stabil serta memiliki kemampuan generalisasi yang baik terhadap data yang belum pernah dilihat sebelumnya.

D. Evaluasi Model

Evaluasi model dilakukan menggunakan 200 data uji yang telah dipisahkan pada tahap pembagian dataset. Pengujian bertujuan untuk mengukur kemampuan model Random Forest dalam mengklasifikasikan tingkat stres mahasiswa berdasarkan data kuesioner DASS-21. Metrik evaluasi yang digunakan meliputi *accuracy*, *precision*, *recall*, dan *F1-score*. Hasil evaluasi model ditunjukkan pada Tabel 10.

Tabel 10. Hasil Evaluasi Model Random Forest

Metrik	Nilai (%)
<i>Accuracy</i>	73,50
<i>Precision</i>	74,17
<i>Recall</i>	73,50
<i>F1-Score</i>	73,61

Berdasarkan Tabel 10, model Random Forest memperoleh nilai *accuracy* sebesar 73,50%, *precision* sebesar 74,17%, *recall* sebesar 73,50%, dan *F1-score* sebesar 73,61%. Hasil tersebut menunjukkan bahwa model mampu mengklasifikasikan tingkat stres mahasiswa dengan performa yang cukup baik serta memiliki keseimbangan antara ketepatan prediksi (*precision*) dan kemampuan mengidentifikasi setiap kategori (*recall*). Performa model ini sejalan dengan penelitian sebelumnya yang menunjukkan bahwa Random Forest efektif digunakan pada permasalahan klasifikasi di bidang kesehatan mental dan tingkat stres [11], [17]. Perbedaan nilai akurasi dengan penelitian terdahulu kemungkinan dipengaruhi oleh karakteristik dataset, jumlah responden, distribusi kelas, serta proses preprocessing dan optimasi hyperparameter yang digunakan. Untuk memperoleh gambaran yang lebih rinci mengenai performa model pada setiap kategori tingkat stres, dilakukan analisis menggunakan classification report yang menyajikan nilai *precision*, *recall*, dan *F1-score* pada masing-masing kelas. Hasil klasifikasi ditunjukkan pada Gambar 2.

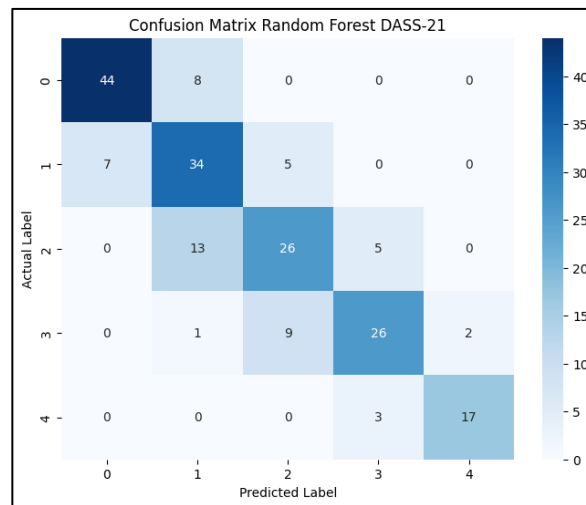
	precision	recall	f1-score	support
Normal	0.86	0.85	0.85	52
Ringan	0.61	0.74	0.67	46
Sedang	0.65	0.59	0.62	44
Berat	0.76	0.68	0.72	38
Sangat Berat	0.89	0.85	0.87	20
accuracy			0.73	200
macro avg	0.76	0.74	0.75	200
weighted avg	0.74	0.73	0.74	200

Gambar 2. Hasil Klasifikasi Menggunakan Algoritma Random Forest

Berdasarkan Gambar 2, kategori Sangat Berat memiliki performa terbaik dengan nilai *precision* sebesar 0,89, *recall* sebesar 0,85, dan *F1-score* sebesar 0,87. Kategori Normal juga menunjukkan performa yang tinggi dengan nilai *F1-score* sebesar 0,85, yang mengindikasikan bahwa model mampu mengenali karakteristik responden pada kedua kategori tersebut dengan baik. Sebaliknya, kategori Sedang memperoleh nilai *F1-score* terendah sebesar 0,62, menunjukkan bahwa model masih mengalami kesulitan dalam membedakannya dari kategori tingkat stres lain yang memiliki karakteristik serupa. Selain itu, kategori Ringan juga menunjukkan performa yang relatif lebih rendah dengan nilai *F1-score* sebesar 0,67 dibandingkan kategori lainnya.

Secara keseluruhan, model Random Forest memperoleh nilai weighted average *precision* sebesar 0,74, weighted average *recall* sebesar 0,73, dan weighted average *F1-score* sebesar 0,74. Hasil tersebut menunjukkan bahwa model mampu mengklasifikasikan tingkat stres mahasiswa dengan performa yang cukup baik pada seluruh kategori yang

diuji. Untuk memperoleh analisis yang lebih rinci mengenai hasil klasifikasi pada setiap kategori, digunakan confusion matrix yang disajikan pada Gambar 3.

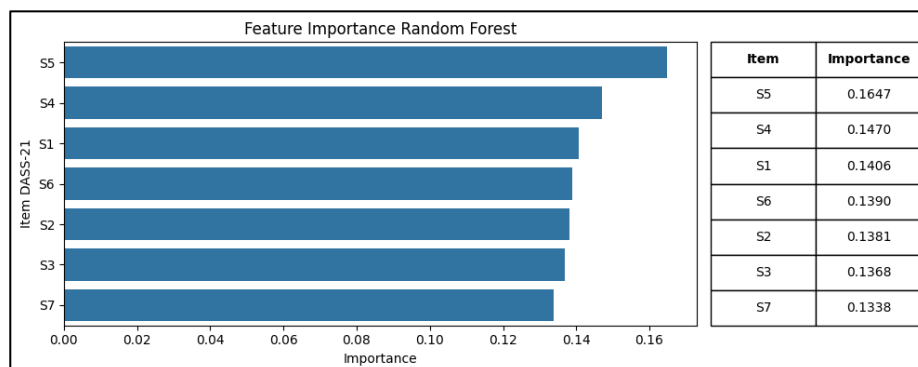


Gambar 3. Confusion Matrix Random Forest

Berdasarkan Gambar 3, model mengklasifikasikan kategori Normal dengan baik, dengan 44 dari 52 data diprediksi benar dan 8 data lainnya diprediksi sebagai Ringan. Pada kategori Ringan, 34 dari 46 data berhasil diklasifikasikan dengan benar, sementara sisanya diprediksi sebagai Normal dan Sedang. Pada kategori Sedang, 26 dari 44 data diprediksi dengan tepat, sedangkan 13 data diprediksi sebagai Ringan dan 5 data sebagai Berat. Hal ini menunjukkan bahwa kategori Sedang merupakan salah satu kategori yang paling sulit dibedakan karena berada di antara kategori Ringan dan Berat. Akibatnya, beberapa data pada kategori Ringan diprediksi sebagai Sedang, dan sebaliknya. Untuk kategori Berat, model berhasil mengklasifikasikan 26 dari 38 data dengan benar, sementara sisanya diprediksi sebagai Sedang, Ringan, dan Sangat Berat. Adapun kategori Sangat Berat menunjukkan performa yang baik dengan 17 dari 20 data berhasil diklasifikasikan secara tepat dan hanya 3 data diprediksi sebagai Berat.

Berdasarkan hasil tersebut, kategori Sedang memiliki nilai recall paling rendah, yaitu 59%. Hal ini menunjukkan bahwa model masih mengalami kesulitan dalam mengidentifikasi data pada kategori tersebut. Secara psikologis, kategori Sedang merupakan kondisi transisi antara tingkat stres Ringan dan Berat, sehingga respons pada beberapa item DASS-21 cenderung saling tumpang tindih. Akibatnya, model menghasilkan prediksi yang berpindah ke kategori yang berdekatan dibandingkan ke kategori yang memiliki tingkat stres jauh berbeda. Pola ini menunjukkan bahwa kesalahan klasifikasi lebih dipengaruhi oleh kemiripan karakteristik data daripada ketidakmampuan model dalam mengenali pola secara umum.

Berdasarkan hasil analisis confusion matrix, masih terdapat beberapa kesalahan klasifikasi pada setiap kategori tingkat stres mahasiswa. Kesalahan tersebut menunjukkan bahwa beberapa kategori memiliki karakteristik yang saling berdekatan sehingga sulit dibedakan oleh model Random Forest. Oleh karena itu, dilakukan analisis feature importance untuk mengidentifikasi kontribusi masing-masing variabel, yaitu S1 hingga S7, dalam proses pengambilan keputusan model Random Forest. Hasil analisis feature importance ditunjukkan pada Gambar 4.



Gambar 4. Feature Importance Random Forest

Gambar 4 menunjukkan bahwa seluruh item DASS-21 memiliki nilai feature importance yang relatif berdekatan, yaitu berkisar antara 0,133 hingga 0,165. Item S5 (Saya merasa sulit bersantai) memiliki nilai feature importance tertinggi sebesar 0,1647, diikuti oleh S4 (Saya merasa mudah gelisah) sebesar 0,1470 dan S1 (Saya merasa sulit beristirahat) sebesar 0,1406. Hal ini menunjukkan bahwa ketiga item tersebut memberikan kontribusi paling besar dalam membedakan tingkat stres mahasiswa pada dataset penelitian. Sementara itu, item S7 (Saya merasa bahwa saya mudah tersinggung) memiliki nilai feature importance terendah sebesar 0,1338, meskipun perbedaannya relatif kecil dibandingkan item lainnya. Temuan bahwa item sulit bersantai menjadi indikator dengan kontribusi terbesar sejalan dengan teori Lovibond dan Lovibond yang menyatakan bahwa ketegangan serta kesulitan untuk rileks merupakan karakteristik utama dari dimensi stres pada DASS-21 [3]. Secara keseluruhan, hasil ini menunjukkan bahwa seluruh item pada subskala stres DASS-21 berkontribusi terhadap proses klasifikasi tanpa adanya satu variabel yang mendominasi secara signifikan.

IV. SIMPULAN

Hasil penelitian yang telah dilakukan, algoritma Random Forest berhasil diimplementasikan untuk mengklasifikasikan tingkat stres mahasiswa berdasarkan data kuesioner DASS-21 ke dalam lima kategori, yaitu Normal, Ringan, Sedang, Berat, dan Sangat Berat. Tahapan preprocessing yang meliputi cleaning, scoring, labeling, dan encoding berhasil menghasilkan dataset yang siap digunakan, sedangkan hyperparameter tuning menggunakan GridSearchCV menghasilkan parameter terbaik yang mendukung kinerja model. Hasil pengujian menunjukkan bahwa model memperoleh nilai accuracy sebesar 73,50%, precision sebesar 74,17%, recall sebesar 73,50%, dan F1-score sebesar 73,61%, sehingga tujuan penelitian telah tercapai. Model menunjukkan performa yang baik dalam mengidentifikasi kategori Normal dan Sangat Berat, meskipun masih terdapat kesalahan klasifikasi pada kategori Ringan, Sedang, dan Berat. Oleh karena itu, penelitian selanjutnya disarankan menggunakan dataset yang lebih besar dan seimbang, menerapkan teknik penanganan ketidakseimbangan data seperti Synthetic Minority Over-sampling Technique (SMOTE) atau class weighting, serta membandingkan performa Random Forest dengan algoritma klasifikasi lainnya untuk memperoleh hasil yang lebih optimal.

REFERENSI

- [1] Thasya Nur Oktaviona, Herlina, and T. H. Sari, "Faktor Yang Mempengaruhi Tingkat Stres Pada Mahasiswa Akhir," *J. Keperawatan Trop. Papua*, vol. 6, no. 1, pp. 26–32, 2023, doi: 10.47539/jktp.v6i1.344.
- [2] W. Musikhah and D. Nastiti, "Academic Stress of University of Muhammadiyah Sidoarjo Students Who Study While Working Class of 2021," *J. Islam. Muhammadiyah Stud.*, vol. 2, pp. 1–6, 2022, doi: <https://doi.org/10.21070/jims.v2i0.1543>.
- [3] Q. Nada, I. Herdiana, and F. Andriani, "Testing the validity and reliability of the Depression Anxiety Stress Scale (DASS)-21 instrument for individuals with Psychodermatology," *Psikohumaniora*, vol. 7, no. 2, pp. 153–168, 2022, doi: 10.21580/pjpp.v7i2.11802.
- [4] A. S. N. Isha *et al.*, "Validation of 'Depression, Anxiety, and Stress Scales' and 'Changes in Psychological Distress during COVID-19' among University Students in Malaysia," *Sustain.*, vol. 15, no. 5, 2023, doi: 10.3390/su15054492.
- [5] R. Rois, M. Ray, A. Rahman, and S. K. Roy, "Prevalence and predicting factors of perceived stress among Bangladeshi university students using machine learning algorithms," *J. Heal. Popul. Nutr.*, vol. 40, no. 1, pp. 1–12, 2021, doi: 10.1186/s41043-021-00276-5.
- [6] S. Arya, Anju, and N. A. Ramli, "Predicting the stress level of students using Supervised Machine Learning and Artificial Neural Network (ANN)," *Indian J. Eng.*, vol. 21, no. 56, 2024, doi: 10.54905/disssi.v21i55.e9ije1684.
- [7] K. Rahayu, V. Fitria, D. Septhya, and ..., "Klasifikasi Teks untuk Mendeteksi Depresi dan Kecemasan pada Pengguna Twitter Berbasis Machine Learning: Text Classification for Detecting Depression and ...," *MALCOM Indones. J. Mach. Learn. Comput. Sci.*, vol. 3, no. 2, pp. 108–114, 2023, doi: <https://doi.org/10.57152/malcom.v3i2.780>.
- [8] I. Adriansyah, M. D. Mahendra, E. Rasywir, and Y. Pratama, "Perbandingan Metode Random Forest Classifier dan SVM Pada Klasifikasi Kemampuan Level Beradaptasi

- Pembelajaran Jarak Jauh Siswa,” *Bull. Informatics Data Sci.*, vol. 1, no. 2, pp. 98–103, 2022, doi: <http://dx.doi.org/10.61944/bids.v1i2.49>.
- [9] D. Haganta Depari, Y. Widiastiwi, and M. Mega Santoni, “Perbandingan Model Decision Tree, Naive Bayes dan Random Forest untuk Prediksi Klasifikasi Penyakit Jantung,” *J. Inform.*, vol. 18, no. 3, pp. 239–248, Dec. 2022, doi: <https://doi.org/10.52958/iftk.v18i3.4694>.
- [10] A. Yusuf Permana, H. Noer Fazri, M. Fakhrizal Nur Athoilah, M. Robi, and R. Firmansyah, “Penerapan Data Mining Dalam Analisis Prediksi Kanker Paru Menggunakan Algoritma Random Forest,” *J. Ilm. Tek. Inform. dan Komun.*, vol. 3, no. 2, pp. 27–41, Jun. 2023, doi: <https://doi.org/10.35970/infotekmesin.v14i1.1751>.
- [11] M. L. Wallace *et al.*, “Use and misuse of random forest variable importance metrics in medicine: demonstrations through incident stroke prediction,” *BMC Med. Res. Methodol.*, vol. 23, no. 1, pp. 1–12, 2023, doi: 10.1186/s12874-023-01965-x.
- [12] L. Langsetmo *et al.*, “Advantages and Disadvantages of Random Forest Models for Prediction of Hip Fracture Risk Versus Mortality Risk in the Oldest Old,” *JBMR Plus*, vol. 7, no. 8, pp. 1–11, 2023, doi: 10.1002/jbm4.10757.
- [13] Y. Rothacher and C. Strobl, *Identifying Informative Predictor Variables With Random Forests*, vol. 49, no. 4. 2024. doi: 10.3102/10769986231193327.
- [14] R. Ririn, A. Rizki Rinaldi, and F. M Basysyar, “Prediksi Kesehatan Mental Menggunakan Random Forest Berdasarkan Faktor Demografi Dan Lingkungan Kerja,” *JATI (Jurnal Mhs. Tek. Inform.)*, vol. 9, no. 3, pp. 4515–4522, 2025, doi: 10.36040/jati.v9i3.13720.
- [15] D. Fuji Astari, Y. Herry Chrisnanto, and M. Melina, “KLASIFIKASI TINGKAT STRES SAAT TIDUR MENGGUNAKAN ALGORITMA RANDOM FOREST,” *JATI (Jurnal Mhs. Tek. Inform.)*, vol. 7, no. 5, pp. 3676–3684, 2024, doi: 10.36040/jati.v7i5.7750.
- [16] Oon Wira Yuda, Darmawan Tuti, Lim Sheih Yee, and Susanti, “Penerapan Data Mining Untuk Klasifikasi Kelulusan Mahasiswa Tepat Waktu Menggunakan Metode Random Forest,” *SATIN - Sains dan Teknol. Inf.*, vol. 8, no. 2, pp. 122–131, Dec. 2022, doi: 10.33372/stn.v8i2.885.
- [17] A. Hapsari, A. S. Nursuwanda, H. Zuhriyah, and D. J. Vresdian, “Klasifikasi Kesehatan Mental Mahasiswa Model TMAS dengan Algoritma Decision Tree, Logistic Regression, dan Random Forest,” *INTEK J. Inform. dan Teknol. Inf.*, vol. 7, no. 2, pp. 55–64, 2024, doi: 10.37729/intek.v7i2.5690.
- [18] L. Pefrianti, I. R. Munthe, I. Irmayanti, and M. Masrizal, “Klasifikasi Tingkat Stres Mahasiswa Dalam Penyelesaian Tugas Akhir Menggunakan Naïve Bayes Dan K-Nearest Neighbor,” *J. Comput. Sci. Inf. Syst.*, vol. 7, no. 1, pp. 129–137, 2026, doi: <https://doi.org/10.36987/jcoins.v7i1.9060>.
- [19] M. Chanafy and H. Sulistiani, “Comparative Analysis Of Machine Learning Algorithm Performance For Student Mental Health Classification,” *SMATIKA STIKI Inform. J.*, vol. 16, no. 1, pp. 19–29, 2026, doi: <https://doi.org/10.32664/smatika.v16i01.2080>.
- [20] I. M. M. Matin, “Hyperparameter Tuning Menggunakan GridsearchCV pada Random Forest untuk Deteksi Malware,” *Multinetics*, vol. 9, no. 1, pp. 43–50, 2023, doi: 10.32722/multinetics.v9i1.5578.
- [21] M. A. Rosid, F. Muharram, and G. R. Affandi, “Perbandingan Kinerja Algoritma Machine Learning Untuk Mendeteksi Kalimat Sarkasme Dalam Bahasa Indonesia,” *Procedia Soc. Sci. Humanit.*, vol. 3, no. c, pp. 1192–1195, 2022, doi: <https://doi.org/10.21070/pssh.v3i.253>.
- [22] J. W. Creswell and J. D. Creswell, *Research Design: Qualitative, Quantitative, and Mixed Methods Approaches*. SAGE Publications, 2022. [Online]. Available:

- <https://books.google.co.id/books?id=Rkh4EAAAQBAJ>
- [23] M. U. Nadeem, S. J. Kulich, and I. H. Bokhari, “The assessment and validation of the depression, anxiety, and stress scale (DASS-21) among frontline doctors in Pakistan during fifth wave of COVID-19,” *Front. Public Heal.*, vol. 11, 2023, doi: 10.3389/fpubh.2023.1192733.
- [24] Y. A. Prasetyo, E. Utami, and A. Yaqin, “Pengaruh Komposisi Split Data Terhadap Performa Akurasi Analisis Sentimen Algoritma Naïve Bayes dan SVM,” *J. Electr. Eng. Comput.*, vol. 6, no. 2, pp. 382–390, 2024, doi: 10.33650/jeecom.v6i2.9188.

Conflict of Interest Statement:

The author declares that the research was conducted in the absence of any commercial or financial relationships that could be construed as a potential conflict of interest.