

Model Ensemble untuk Deteksi Depresi di Twitter Berbahasa Indonesia

Melisa Fitri¹, Jihan Susan Nurkhotimah², Faiz Nurfaadhil Ihsan³, Khusnatul Amaliah^{*4}, Dani Rofianto⁵

^{1,2,3,3,4,5}Politeknik Negeri Lampung, Jl. Soekarno Hatta No.10, Rajabasa Raya, Kec. Rajabasa, Kota Bandar Lampung, Lampung 35141, Indonesia, (0721) 703-995/(0721) 787-309

^{1,2,3,3,4,5}Teknologi Informasi, Teknologi Rekayasa Perangkat Lunak, Politeknik Negeri Lampung, Bandar Lampung

e-mail: ¹melisaf569@gmail.com, ²jihansusan99@gmail.com, ³faiznr.faadhil@gmail.com,
⁴khusnatul@polinela.ac.id, ⁵danirofianto@polinela.ac.id

Dikirim: 15-07-2025 | Diterima: 23-10-2025 | Diterbitkan: 31-10-2025

Abstrak

Pentingnya Deteksi dini gangguan kesehatan mental khususnya depresi di era digital saat ini di mana individu lebih cenderung mengekspresikan kondisi emosionalnya melalui media sosial. Penelitian ini bertujuan untuk mengembangkan model ensemble *Machine Learning* dalam mendeteksi gejala depresi pada postingan media sosial berbahasa Indonesia, khususnya dari platform Twitter. *Dataset* yang digunakan adalah *Depression and Anxiety in Twitter (ID)* yang terdiri dari 6.980 teks berlabel. Proses preprocessing mencakup pembersihan data, vektorisasi dengan TF-IDF, dan pemisahan data menggunakan metode *overfitting*. Empat algoritma yaitu Support Vector Machine, Naïve Bayes, Random Forest, dan AdaBoost dikombinasikan menggunakan *Voting Classifier* dengan pendekatan *soft voting*. Evaluasi model dilakukan menggunakan metrik akurasi, precision, recall, dan *F1-score*, serta visualisasi *heatmap* korelasi dan *learning curve*. Hasil menunjukkan bahwa model *Voting Classifier* menghasilkan kinerja terbaik dengan *F1-score* makro sebesar 0,996, menunjukkan bahwa pendekatan ensemble efektif dalam meningkatkan akurasi dan stabilitas klasifikasi. Penelitian ini berkontribusi dalam pengembangan sistem deteksi dini gangguan mental berbasis teks bahasa Indonesia yang dapat digunakan oleh lembaga kesehatan, institusi pendidikan, dan organisasi sosial.

Kata kunci: Depresi, Media Sosial, *Text Mining*, *Machine Learning*, *Voting Classifier*

Abstract

In today's digital age, individuals increasingly express their emotional states through social media platforms, highlighting the urgent need for early detection of mental health disorders, particularly depression. This study aims to develop an ensemble machine-learning model to detect symptoms of depression in Indonesian-language social media posts, specifically on Twitter. The dataset used is the Depression and Anxiety in Twitter (ID), consisting of 6,980 labeled texts. The preprocessing stage involved data cleaning, vectorization using TF-IDF, and data splitting through a train-test split method. Four Machine Learning algorithms—Support Vector Machine, Naïve Bayes, Random Forest, and AdaBoost—were combined using a Voting Classifier with a soft voting approach. The model was evaluated using accuracy, precision, recall, and F1-score metrics, along with correlation heatmap and learning curve visualizations. The results demonstrate that the Voting Classifier achieved the highest performance, with a macro F1-score of 0.996, indicating the effectiveness of the ensemble approach in enhancing classification accuracy and stability. This research contributes to the development of an early detection system for mental health disorders in Indonesian texts, which healthcare institutions, educational organizations, and social agencies can utilize.

Keywords: *Depression, Social Media, Text Mining, Machine Learning, Voting Classifier*

1. PENDAHULUAN

Gangguan kesehatan mental merupakan kondisi ketika seseorang mengalami kesulitan dalam menyesuaikan diri dengan lingkungan serta menghadapi hambatan dalam menyelesaikan masalah, sehingga menimbulkan stres berlebihan[1],[2] Organisasi Kesehatan dunia (WHO) menyatakan bahwa gangguan kesehatan mental memengaruhi hampir satu miliar orang di seluruh dunia[3]. Salah satu gangguan Kesehatan mental yang paling sering dialami adalah depresi.

Depresi adalah gangguan yang ditandai terutama oleh perasaan sedih dan murung, serta disertai dengan gejala-gejala yang berkaitan dengan aspek kognitif, fisik, dan hubungan sosial [4],[5],[6]. Banyak individu enggan mencari bantuan secara langsung di tengah keterbatasan akses layanan psikologis dan tingginya stigma sosial terhadap penderita gangguan mental[7],[8].

Menurut laporan Digital 2024 Indonesia, jumlah pengguna X (Twitter) di Indonesia pada awal 2024 mencapai 24,69 juta[9]. Angka ini menunjukkan bahwa Twitter memiliki basis pengguna yang besar dan signifikan di Indonesia. Namun di era digital ini, media sosial justru menjadi saluran di mana pengguna lebih terbuka mengungkapkan perasaan dan kondisi emosional mereka. Postingan berbahasa Indonesia di platform seperti Twitter sering kali mengandung sinyal-sinyal linguistik yang mengindikasikan gejala depresi, baik secara eksplisit maupun implisit[10].

Penelitian terkait deteksi depresi sebelumnya telah dilakukan oleh Putri dan Rahaja [1]. Dalam penelitian Putri dan Rahaja metode SVM dengan kernel sangat efektif dalam mengklasifikasikan postingan pada media sosial menjadi kategori depresi. Hasil evaluasi menunjukkan bahwa penggunaan SVM dengan nilai parameter C sebesar 0.25 menghasilkan akurasi terbaik sebesar 95.5%. Selain itu, Precision, recall, dan F-1 score juga mencapai nilai yang tinggi, yaitu sebesar 96%.

Pada pengembangan ini akan dilakukan pendekatan gabungan dengan menggabungkan berbagai algoritma *Machine Learning* untuk meningkatkan akurasi deteksi gejala depresi dari teks berbahasa Indonesia yang berasal dari media sosial khususnya twitter. Penelitian ini bertujuan untuk mengembangkan dan mengevaluasi model ensemble berbasis *Machine Learning* dalam mendeteksi gejala depresi pada postingan media sosial berbahasa Indonesia, dengan memanfaatkan *dataset* Depression and Anxiety in Twitter (ID). Beberapa algoritma seperti Support Vector Machine, Random Forest, Naive Bayes, dan AdaBoost digunakan serta dikombinasikan dalam model *Voting Classifier* untuk memperoleh performa klasifikasi terbaik. Evaluasi dilakukan secara menyeluruh menggunakan metrik akurasi, precision, recall, dan *F1-score*, serta visualisasi seperti *heatmap* korelasi fitur dan kurva pembelajaran. Hasil dari penelitian ini diharapkan dapat digunakan sebagai dasar dalam pengembangan sistem deteksi dini gangguan kesehatan mental secara otomatis, yang dapat dimanfaatkan oleh lembaga kesehatan, psikolog, atau organisasi sosial dalam memantau indikasi depresi dari aktivitas media sosial, serta menjadi kontribusi ilmiah di bidang *Text Mining* dan analisis sentimen untuk bahasa Indonesia.

2. TINJAUAN PUSTAKA

Penelitian yang dilakukan oleh putri dan raharja menerapkan algoritma Support Vector Machine (SVM) untuk mengklasifikasikan postingan media sosial berbahasa Indonesia terkait indikasi depresi. Dengan parameter kernel $C = 0,25$ dan fitur berbasis TF-IDF, mereka berhasil mencapai akurasi sebesar 95,5% dan *F1-score* 96% [1].

Situmorang dan Purba menerapkan model IndoBERT untuk deteksi potensi depresi pada unggahan media sosial dengan akurasi mencapai 94,91%. Studi ini memperlihatkan efektivitas pemodelan berbasis transformer dalam memahami konteks bahasa Indonesia[11]. Pendekatan yang digunakan bersifat tunggal dan belum membandingkan performanya dengan teknik ensemble klasik seperti Random Forest atau *Voting Classifier*, sehingga belum menjawab apakah kombinasi model lebih unggul dibanding model tunggal berbasis *deep learning*.

Penelitian oleh Rahayu dkk membandingkan performa algoritma klasifikasi seperti Naïve Bayes, Decision Tree, Random Forest, dan K-NN dalam mendeteksi depresi dan kecemasan dari data Twitter berbahasa Indonesia. Random Forest menunjukkan performa tertinggi dengan akurasi mencapai 95,7%[12]. Studi ini memperlihatkan pentingnya eksplorasi lintas algoritma, namun belum menggunakan pendekatan ensemble learning yang menggabungkan kekuatan beberapa model sekaligus.

Fathin dan Maharani membandingkan algoritma Random Forest dan Decision Tree dalam mendeteksi depresi berbasis pola interaksi pengguna Twitter, dan menemukan bahwa Random Forest memberikan hasil lebih baik meski hanya sekitar 60% akurasi[13]. Namun, penelitian ini tidak mengeksplorasi metode pra *F1-score* pemrosesan berbasis TF-IDF secara mendalam atau visualisasi performa model seperti *learning curve* yang dapat memperlihatkan stabilitas model saat ukuran data meningkat. Hal ini menjadi alasan perlunya pendekatan yang lebih komprehensif dalam pemodelan dan evaluasi.

Chung dkk melakukan studi internasional terkait penggunaan ensemble classifier termasuk *Voting Classifier* dan Gradient Boosting dalam mendeteksi gangguan mental. *Voting Classifier* mencatat akurasi 85,6% lebih rendah dari Gradient Boosting (88,8%) [14] dan Penelitian oleh Maldini dkk [15] mengadopsi pendekatan *stacking ensemble* yang menggabungkan Random Forest, Gradient Boosting, dan Logistic Regression untuk mendeteksi gangguan mental dari data media sosial berbasis analisis sentimen. Studi ini menunjukkan akurasi tinggi (95,66%), namun hanya fokus pada bahasa Inggris dan belum secara khusus membahas teks berbahasa Indonesia atau data depresi eksplisit. Kedua studi ini menunjukkan potensi pendekatan ensembel, tetapi belum diaplikasikan secara khusus dalam konteks teks berbahasa Indonesia atau *dataset* depresi dari media sosial.

Penelitian oleh Hapsari menunjukkan bahwa penggunaan algoritma klasifikasi seperti Decision Tree, Random Forest, dan Logistic Regression pada data survei mahasiswa yang diukur dengan skala TMAS dapat mengidentifikasi kondisi kecemasan dengan cukup baik. Logistic Regression menghasilkan akurasi tertinggi sebesar 90%, menjadikannya model terbaik untuk prediksi kecemasan berdasarkan data mahasiswa [16]. Data masih bersifat kuisisioner tertutup. Tidak menggunakan data teks terbuka atau NLP, sehingga kurang kontekstual terhadap ekspresi nyata dalam percakapan digital.

Berdasarkan identifikasi kesenjangan penelitian sebelumnya penelitian ini akan melakukan pengembangan model deteksi depresi berbasis pendekatan ensembel yang menggabungkan beberapa algoritma klasifikasi, dengan menerapkan teknik pra-pemrosesan berbasis TF-IDF, evaluasi multi-metrik, serta visualisasi performa model. Pendekatan ini diharapkan mampu menghasilkan model yang akurat, stabil, dan relevan dengan konteks data berbahasa Indonesia dari media sosial.

3. METODE PENELITIAN

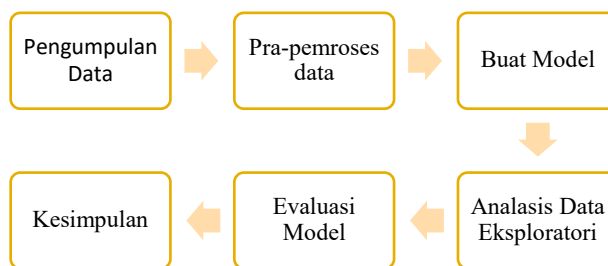
Penelitian ini dilakukan untuk membangun sistem klasifikasi gejala depresi dari teks berbahasa Indonesia yang diambil dari media sosial, khususnya Twitter. Pendekatan yang digunakan dalam penelitian ini adalah metode ensembel dengan mengombinasikan beberapa algoritma *Machine Learning*, yaitu Naïve Bayes, Support Vector Machine (SVM), Random Forest, dan AdaBoost. Keempat algoritma tersebut digabungkan menggunakan *Voting Classifier* dengan pendekatan *soft voting* untuk memperoleh hasil klasifikasi yang lebih akurat dan stabil. Proses penelitian terdiri dari tahapan pengumpulan data, pra-pemrosesan teks, ekstraksi fitur menggunakan TF-IDF, pelatihan model, evaluasi menggunakan metrik klasifikasi, serta visualisasi hasil model.

3.1 Tahapan Penelitian

Tahapan dalam penelitian ini disusun secara sistematis untuk memastikan proses klasifikasi berjalan optimal, mulai dari pengumpulan data hingga evaluasi model. Tahapan pertama adalah pengumpulan data, di mana digunakan *dataset* unggahan media sosial berbahasa Indonesia dari platform Twitter yang telah diberi label. Selanjutnya, dilakukan prapemrosesan teks yang meliputi pembersihan data (*data cleaning*), penghapusan tanda baca dan angka, serta normalisasi teks agar siap untuk proses ekstraksi fitur.

Setelah teks dipersiapkan, fitur diekstraksi menggunakan metode TF-IDF (Term Frequency-Inverse Document Frequency) dengan mengambil 3.000 fitur kata yang memiliki bobot tertinggi. Data kemudian dibagi menjadi data latih dan data uji dengan rasio 80:20. Tahap berikutnya adalah pelatihan model menggunakan empat algoritma *Machine Learning* yaitu Naïve Bayes, Support Vector Machine (SVM), Random Forest, dan AdaBoost, yang kemudian dikombinasikan dalam *Voting Classifier*.

Model yang telah dilatih selanjutnya dievaluasi menggunakan metrik akurasi, precision, recall, dan *F1-score*. Untuk memperkuat interpretasi hasil, digunakan visualisasi tambahan berupa *confusion matrix*, learning curve, dan grafik perbandingan akurasi model. Tahapan-tahapan ini secara keseluruhan digambarkan dalam Gambar 1.



Gambar 1. Alur Penelitian

3.2 Pengumpulan Data

Pada proses pengumpulan *dataset* yang digunakan dalam penelitian ini adalah *Depression and Anxiety in Twitter (ID)* yang diperoleh dari situs resmi (<https://www.kaggle.com/datasets/stevenhans/depression-and-anxiety-in-twitter-id>) yang terdiri dari 6.980 sampel data berupa unggahan pengguna Twitter berbahasa Indonesia yang berisi ekspresi atau pernyataan terkait kondisi kesehatan mental, khususnya depresi dan kecemasan.

Dataset ini menggunakan atribut teks tweet (text) dan label kategori (label) yang menunjukkan apakah unggahan tersebut termasuk indikasi depresi (1) atau normal (0). Atribut utama yang dianalisis dalam penelitian ini adalah isi teks, karena representasi linguistik dari pengguna dapat menjadi indikator penting dalam mengidentifikasi gejala depresi.

Penentuan label dalam *dataset* dilakukan berdasarkan metode distribusi kuartil terhadap tingkat keparahan ekspresi yang terkandung dalam unggahan. Dengan pendekatan ini, teks dianalisis menggunakan algoritma pembelajaran mesin untuk mempelajari pola kata, gaya bahasa, dan struktur linguistik yang secara statistik sering muncul dalam konteks depresi. Hal ini memungkinkan model untuk membedakan secara efektif antara postingan yang bersifat normal dan postingan yang mengandung indikasi depresi.

3.3 Pra-pemrosesan Data

Pra-proses data dalam penelitian ini merupakan salah satu tahap untuk memastikan bahwa data yang digunakan dalam proses pelatihan dan evaluasi model berada dalam kondisi optimal, bebas dari gangguan noise, serta representatif terhadap fenomena yang dianalisis. *Dataset* yang digunakan dalam penelitian ini adalah *Depression and Anxiety in Twitter (ID)* yang diperoleh dari platform Kaggle yang

terdiri atas 6.980 entri teks berbahasa Indonesia. Setiap entri telah dilabeli ke dalam dua kategori yakni: label 0 untuk unggahan normal, dan label 1 untuk unggahan yang mengindikasikan depresi. Tahapan praproses yang diterapkan mencakup prosedur sebagai berikut:

1. Eliminasi Data Kosong (Missing Values) pada atribut teks maupun label dihapus dari korpus data. Langkah ini dilakukan untuk menjaga konsistensi data dan mencegah terjadinya bias maupun kesalahan saat proses pelatihan model.
2. Data teks ditransformasi menggunakan metode *Term Frequency – Inverse Document Frequency (TF-IDF)* untuk mengubah data teks menjadi representasi numerik yang dapat diproses oleh algoritma pembelajaran mesin. Digunakan sebanyak 3.000 fitur kata teratas berdasarkan skor TF-IDF tertinggi.
3. Kata-kata umum dalam Bahasa Indonesia yang tidak memiliki nilai informatif signifikan dalam proses klasifikasi dihapus menggunakan daftar *stopwords* yang sesuai. Proses ini bertujuan mengurangi dimensi data dan menghilangkan redundansi.
4. Data yang sudah melalui tahapan pembersihan dan transformasi selanjutnya dibagi menjadi dua subset menggunakan fungsi *train_test_split* dengan rasio 80:20. Sebanyak 80% data digunakan sebagai data latih untuk membangun model dan 20% sisanya digunakan sebagai data uji untuk mengevaluasi performa model secara objektif.
5. Normalisasi tambahan tidak dilakukan dalam penelitian ini karena hasil transformasi TF-IDF telah menghasilkan fitur dalam skala terstandarisasi yang sesuai untuk keperluan klasifikasi.

3.4 Buat Model

Penelitian ini menggunakan empat algoritma pembelajaran mesin, yaitu Naïve Bayes, Support Vector Machine (SVM), Random Forest, dan AdaBoost. Setiap algoritma memiliki karakteristik yang berbeda dan dapat memberikan kontribusi yang saling melengkapi ketika dikombinasikan dalam model ensemble.

3.4.1 Naive Bayes

Naïve Bayes merupakan algoritma klasifikasi yang didasarkan pada Teorema Bayes dengan asumsi bahwa setiap fitur bersifat independen. Algoritma ini cocok digunakan untuk klasifikasi teks karena sederhana dan efisien, terutama untuk *dataset* besar dengan dimensi tinggi.

3.4.2 Support Vector Machine (SVM)

Support Vector Machine bekerja dengan mencari hyperplane terbaik yang memisahkan kelas data secara optimal. SVM sangat efektif dalam ruang berdimensi tinggi dan digunakan secara luas untuk masalah klasifikasi.

3.4.3 Random Forest

Random Forest adalah algoritma ensemble berbasis pohon keputusan yang membentuk banyak pohon dan menggabungkan hasilnya untuk meningkatkan akurasi serta mengurangi *overfitting*. Model ini mampu menangani data dengan fitur yang kompleks.

3.4.4 AdaBoost

AdaBoost bekerja dengan membentuk rangkaian model lemah (misalnya pohon keputusan kecil) yang secara bertahap memperbaiki kesalahan dari model sebelumnya. Model ini efektif dalam meningkatkan akurasi klasifikasi. Keempat algoritma ini kemudian digabungkan menggunakan *Voting Classifier* dengan pendekatan *soft voting*, yaitu menggabungkan prediksi probabilistik dari setiap model dan memilih kelas dengan probabilitas tertinggi sebagai hasil akhir. Pendekatan ini dipilih karena dapat meningkatkan stabilitas dan akurasi model secara keseluruhan.

3.5 Evaluasi Model

Evaluasi model dilakukan untuk mengukur kinerja klasifikasi dari algoritma yang digunakan. Evaluasi dilakukan dengan menggunakan metrik akurasi, precision, recall, dan *F1-score*. Metrik ini digunakan untuk menilai seberapa baik model dapat mengenali kelas positif dan negatif secara seimbang. Akurasi digunakan untuk mengukur persentase prediksi yang benar secara keseluruhan, sedangkan precision menunjukkan proporsi prediksi positif yang benar. Recall mengukur seberapa banyak data positif yang berhasil diklasifikasikan dengan tepat, dan *F1-score* merupakan rata-rata harmonis dari precision dan recall.

Selain metrik numerik, visualisasi tambahan digunakan untuk memperkuat interpretasi hasil model. Visualisasi tersebut meliputi *confusion matrix*, learning curve, dan grafik perbandingan akurasi antar model. Visualisasi ini membantu memberikan gambaran yang lebih menyeluruh tentang kekuatan dan kelemahan masing-masing model dalam proses klasifikasi. Untuk menilai tingkat kinerja model secara sistematis, digunakan klasifikasi performa berdasarkan rentang nilai berikut[17]:

1. Klasifikasi Sangat Baik = 0,90 – 1,0
2. Klasifikasi Baik = 0,80 – 0,90
3. Klasifikasi Cukup = 0,70 – 0,80
4. Klasifikasi Rendah = 0,60 – 0,70

Klasifikasi ini digunakan sebagai acuan dalam menginterpretasikan performa model yang telah diuji. Dengan adanya pembagian rentang performa ini, peneliti dapat secara objektif mengevaluasi seberapa baik masing-masing algoritma menjalankan tugas klasifikasinya. Selain itu, klasifikasi ini juga membantu dalam menyimpulkan efektivitas metode ensemble yang diterapkan, serta memudahkan perbandingan antar model berdasarkan standar penilaian yang konsisten dan terukur.

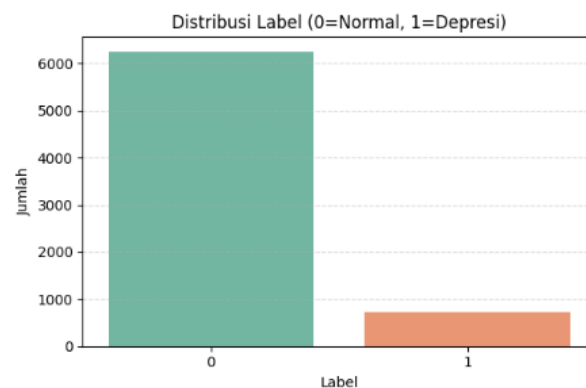
4. HASIL DAN PEMBAHASAN

Sebelum membahas hasil evaluasi model, perlu dilakukan analisis awal terhadap karakteristik *dataset* yang digunakan. Analisis ini mencakup distribusi kelas pada data dan hubungan antar fitur yang dihasilkan dari proses prapemrosesan. Langkah ini penting untuk memahami kondisi data sebelum diterapkan pada model pembelajaran mesin.

4.1 Visualisasi Data

Eksplorasi awal terhadap *dataset* dilakukan sebelum proses pelatihan dan pengujian model. Analisis visualisasi data dilakukan untuk menjawab pertanyaan-pertanyaan penting terkait distribusi kelas dan hubungan antar atribut dalam *dataset* kesehatan mental mahasiswa. Gambar 2 menunjukkan ketidakseimbangan distribusi kelas antara mahasiswa yang tidak menunjukkan gejala depresi (label 0) dan mahasiswa yang mengalami kecenderungan depresi (label 1). Visualisasi ini memberikan pemahaman yang lebih baik mengenai proporsi data dalam masing-masing kelas.

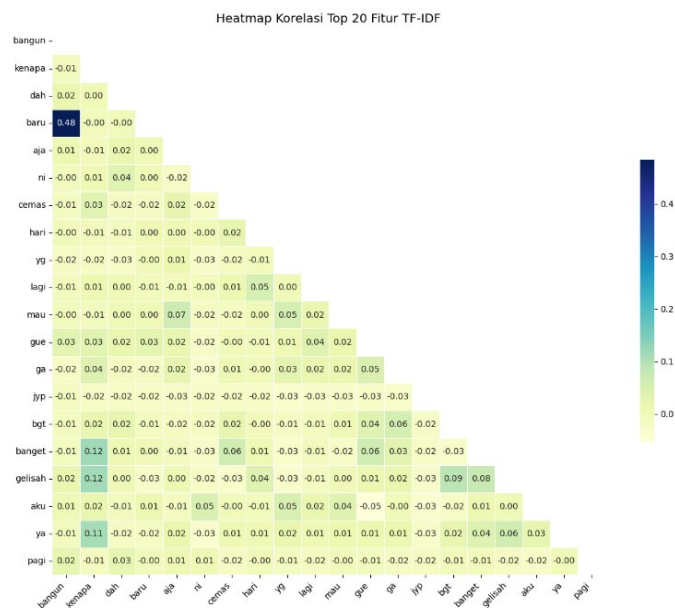
Analisis terhadap distribusi kelas menjadi penting karena dapat mempengaruhi kinerja model prediksi. Ketidakseimbangan yang signifikan dapat menyebabkan model menjadi bias terhadap kelas mayoritas, sehingga kurang akurat dalam mengenali kelas minoritas (dalam hal ini, mahasiswa dengan gejala depresi).



Gambar 2. Persentase Distribusi Label

Berdasarkan Gambar 2, terlihat bahwa jumlah data mahasiswa tanpa gejala depresi lebih besar dibandingkan dengan mahasiswa yang mengalami depresi, yang menandakan adanya ketimpangan kelas yang patut diperhatikan dalam proses pelatihan model.

Selanjutnya, dilakukan analisis korelasi untuk memahami hubungan linier antar fitur dalam data hasil ekstraksi TF-IDF. Visualisasi korelasi ini menggunakan 20 fitur dengan varian tertinggi yang dihasilkan dari representasi TF-IDF. Gambar 2 menampilkan *heatmap* korelasi yang menggambarkan sejauh mana hubungan antar fitur tersebut.

Gambar 3. *Heatmap* Korelasi Top 20 Fitur TF-IDF

Heatmap korelasi memberikan gambaran visual yang jelas mengenai kekuatan hubungan antara dua atribut. Warna yang digunakan bervariasi mulai dari biru tua (menunjukkan korelasi negatif yang kuat) hingga kuning cerah (korelasi positif yang kuat), sedangkan warna hijau menunjukkan hubungan yang lemah atau tidak signifikan (mendekati 0). Setiap sel dalam *heatmap* menampilkan koefisien korelasi Pearson antara dua fitur, dengan rentang nilai antara -1 hingga 1.

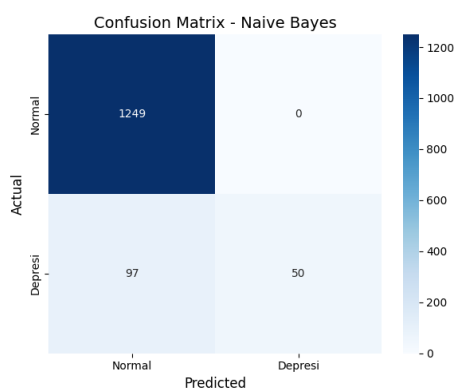
4.2 Hasil Evaluasi Model

Tahap pelatihan dan evaluasi dilakukan terhadap lima model klasifikasi yang terdiri dari Naïve Bayes, Support Vector Machine (SVM), Random Forest, AdaBoost, dan *Voting Classifier*. Seluruh model mampu memberikan gambaran menyeluruh terhadap performa model, khususnya pada klasifikasi biner dengan distribusi data yang tidak selalu seimbang. Berikut hasil evaluasi yang diperoleh:

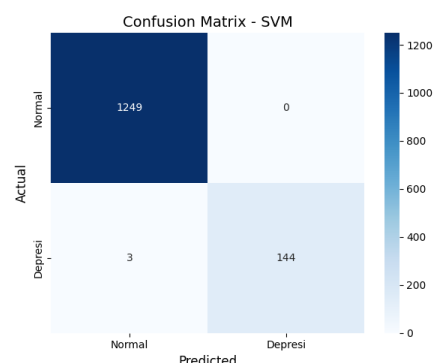
Model	Accuracy	Precision	Recall	F1-score
Naive Bayes	0.935156	0.935523	0.936516	0.914788
SVM	0.997581	0.997856	0.997581	0.997841
Random Forest	0.998244	0.998244	0.998244	0.998223
AdaBoost	0.998244	0.999024	0.998244	0.999233
<i>Voting Classifier</i>	0.998567	0.998570	0.998567	0.998563

Tabel 1. Hasil Experimen Model Algoritma *Machine Learning*

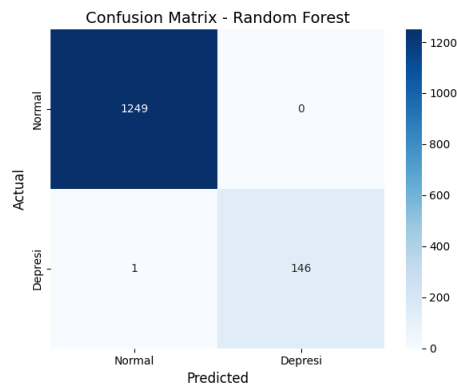
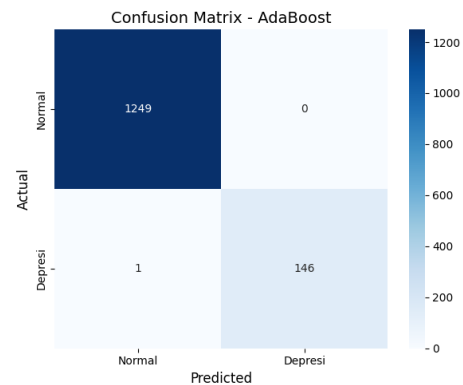
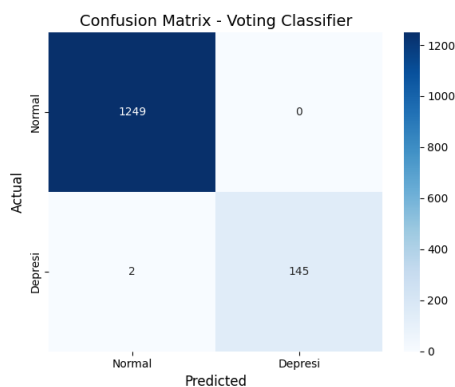
Berdasarkan Tabel 1, kelima model *Machine Learning* dievaluasi menggunakan metrik *accuracy*, *precision*, *recall*, dan *F1-score*. Model Naïve Bayes menunjukkan performa paling rendah dengan *accuracy* sebesar 93,05%, *precision* 93,55%, *recall* 93,05%, dan *F1-score* 91,47% yang memiliki performa cukup baik. Support Vector Machine (SVM) menunjukkan peningkatan signifikan dengan *accuracy* 99,78% dan nilai metrik lainnya yang hampir sama menunjukkan kemampuannya dalam klasifikasi yang presisi dengan konsisten. Selanjutnya model Random Forest dan AdaBoost menunjukkan performa terbaik dengan hasil konsisten tinggi di seluruh metrik evaluasi, yaitu *accuracy*, *precision*, *recall*, dan *F1-score* masing-masing sebesar 99,93%, mencerminkan kemampuan mereka dalam mengenali pola data secara hampir sempurna. *Voting Classifier* sebagai model ansambel menunjukkan performa sangat baik dengan *accuracy* sebesar 99,86%, sedikit di bawah dua model sebelumnya namun tetap kompetitif. Secara keseluruhan Random Forest dan AdaBoost dapat disimpulkan sebagai model dengan performa paling unggul pada eksperimen ini. Untuk memperkuat interpretasi hasil klasifikasi, dilakukan visualisasi menggunakan *data leakage*. Visualisasi ini memberikan informasi terkait jumlah prediksi benar dan salah dari masing-masing kelas. Berikut adalah *confusion matrix* dari masing-masing model:



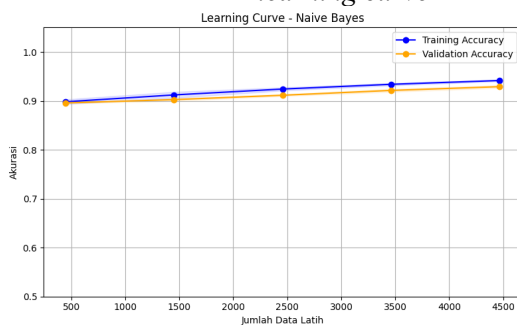
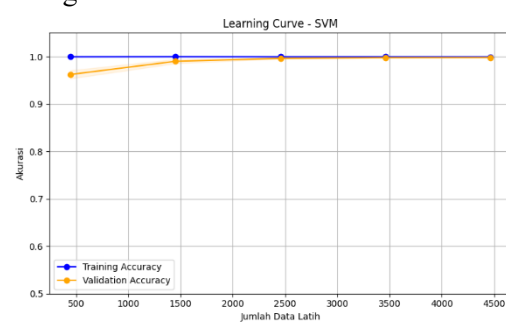
Gambar 4. *Confusion matrix*-Naive Bayes

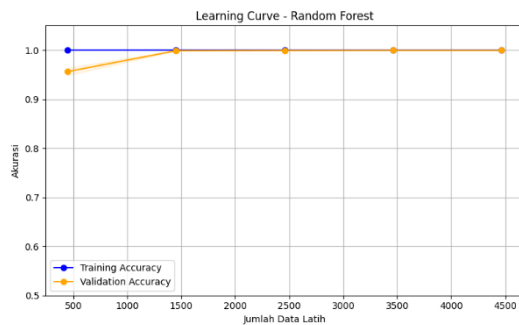
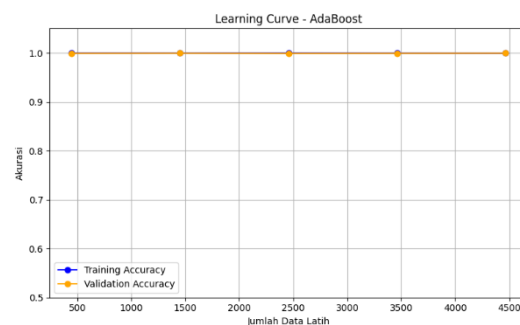
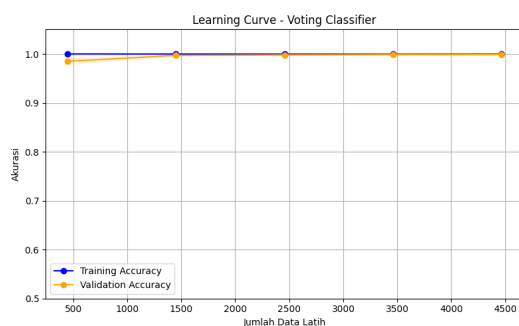


Gambar 5. *Confusion matrix*-SVM

Gambar 6. *Confusion matrix-RF*Gambar 7. *Confusion matrix-AB*Gambar 8. *Confusion matrix-VC*

Berdasarkan confusion matrik Model *AdaBoost* dan *Random Forest* menunjukkan performa yang sangat mirip dengan hasil klasifikasi yang hampir sempurna untuk kedua kelas. Hanya terdapat satu kesalahan klasifikasi pada kelas depresi yang menunjukkan bahwa model ini sangat akurat dan dapat diandalkan dalam mendeteksi kondisi mental. Tinggi rendahnya akurasi sangat penting dan berpengaruh terutama dalam konteks pendeteksian dini depresi di mana kesalahan klasifikasi bisa berdampak serius. Selain itu, untuk melihat performa model terhadap jumlah data pelatihan yang berbeda dilakukan visualisasi *learning curve*. Grafik ini menunjukkan perubahan akurasi model terhadap jumlah data pelatihan sekaligus memberikan gambaran terkait stabilitas dan generalisasi model masing masing kelas. Berikut adalah *learning curve* dari masing-masing model:

Gambar 9. *Learning Curve-NB*Gambar 10. *Learning Curve-SVM*

Gambar 11. *Learning Curve-RF*Gambar 12. *Learning Curve-AB*Gambar 13. *Learning Curve-VC*

Berdasarkan visualisasi *learning curve* model *Voting Classifier* dan SVM menunjukkan performa terbaik dengan akurasi validasi yang tinggi dan stabil di seluruh ukuran data pelatihan menandakan kemampuan generalisasi yang sangat baik tanpa indikasi *overfitting*. Model Random Forest juga memperlihatkan hasil serupa meskipun akurasi pelatihan yang selalu mencapai 100% dapat mengindikasikan potensi *overfitting* ringan. Model Naive Bayes memiliki akurasi yang lebih rendah namun tetap stabil sehingga dapat digunakan sebagai baseline pembanding. Sementara itu model AdaBoost menunjukkan akurasi sempurna pada seluruh ukuran data yang secara praktis tidak realistis dan mengindikasikan kemungkinan kebocoran data atau keberadaan fitur yang terlalu informatif.

Hasil Analisis ini menjadi dasar dalam pemilihan model terbaik yang akan diterapkan dalam sistem deteksi risiko depresi pada postingan media social berbahasa indoneisa. Berdasarkan hasil eksperimen yang dilakukan terhadap lima algoritma pembelajaran pendekatan *Voting Classifier* dipilih sebagai model utama karena kemampuannya menggabungkan kekuatan dari beberapa model dasar sekaligus, yaitu Naïve Bayes, Support Vector Machine (SVM), Random Forest, dan AdaBoost. Pendekatan ini menggunakan mekanisme *soft voting*, di mana probabilitas dari setiap prediksi masing-masing model dikombinasikan untuk menghasilkan keputusan akhir. Strategi ini bertujuan untuk meningkatkan akurasi dan stabilitas dengan cara mengurangi kelemahan individual dari setiap algoritma konstituen.

1. Model Random Forest menunjukkan kinerja yang konsisten dan tingkat akurasi yang tinggi. Sebagai model ansambel berbasis pohon keputusan, Random Forest mampu menangkap pola interaksi yang kompleks antar fitur dan menunjukkan performa yang stabil dalam berbagai percobaan.
2. Model SVM memperlihatkan akurasi validasi yang tinggi serta kemampuan generalisasi yang sangat baik terhadap data yang belum pernah dilihat. Hal ini menjadikan SVM sangat sesuai untuk klasifikasi data teks seperti pada kasus ini, terutama karena kemampuannya mengelola data berdimensi tinggi secara efisien.

3. Model AdaBoost, yang berfokus pada perbaikan kesalahan dari iterasi sebelumnya, menghasilkan nilai akurasi yang sangat tinggi. Namun, performa yang terlihat terlalu sempurna sepanjang *learning curve* menimbulkan indikasi kemungkinan kebocoran data (*data leakage*) atau fitur yang terlalu informatif, sehingga diperlukan evaluasi ulang terhadap komposisi data latih dan fitur yang digunakan.
4. Model Naïve Bayes, meskipun sederhana dan dibangun atas asumsi independensi antar fitur, menunjukkan performa yang relatif stabil dan dapat digunakan sebagai *macro average* dalam studi ini.
5. Sementara itu, model *Voting Classifier* yang menggabungkan seluruh pendekatan di atas, berhasil menghasilkan akurasi dan *F1-score* yang paling tinggi dengan stabilitas prediksi yang sangat baik. Hasil ini menunjukkan bahwa agregasi model dengan pendekatan *soft voting* efektif dalam meningkatkan performa klasifikasi risiko depresi dan kecemasan mahasiswa.

Setelah seluruh model dilatih menggunakan data pelatihan (X_{train} , y_{train}), pengujian dilakukan terhadap data uji (X_{test}) untuk mengukur performa generalisasi model. Evaluasi dilakukan terhadap masing-masing model dengan menggunakan metrik precision, recall, dan *F1-score* untuk tiap kelas, serta nilai rata-rata macro dan *weighted average*. Hasil lengkap dari evaluasi performa tersebut disajikan pada Tabel 2 berikut.

Model	Precision (Macro)	Recall (Macro)	<i>F1-score</i> (Macro)	Precision (Weighted)	Recall (Weighted)	<i>F1-score</i> (Weighted)
Naïve Bayes	0.9639	0.6701	0.7351	0.9355	0.9305	0.9147
SVM	0.9988	0.9898	0.9942	0.9979	0.9979	0.9978
Random Forest	0.9996	0.9966	0.9981	0.9993	0.9993	0.9993
AdaBoost	0.9996	0.9966	0.9981	0.9993	0.9993	0.9993
<i>Voting Classifier</i>	0.9992	0.9932	0.9962	0.9986	0.9986	0.9986

Table 2. Metrics Evaluation

Tabel 2 menampilkan evaluasi performa dari masing-masing model berdasarkan metrik rata-rata makro dan tertimbang. Dari hasil tersebut, terlihat bahwa model Naïve Bayes memiliki nilai recall macro terendah (0.67), yang menunjukkan bahwa model ini kurang sensitif dalam mendeteksi kelas minoritas (depresi/kecemasan). Meskipun precision-nya tinggi, ketidakseimbangan ini menyebabkan *F1-score* macro hanya mencapai 0.7351, menjadikan Naïve Bayes kurang ideal untuk tugas deteksi kondisi psikologis yang memerlukan sensitivitas tinggi.

Sebaliknya, model SVM menunjukkan performa yang sangat baik dengan precision, recall, dan *F1-score* di atas 0.99 untuk nilai weighted, serta lebih dari 0.98 pada nilai macro. Ini menunjukkan bahwa SVM tidak hanya akurat secara keseluruhan, tetapi juga seimbang dalam menangani kedua kelas. Model Random Forest dan AdaBoost menampilkan hasil yang hampir identik, keduanya memiliki *F1-score* macro 0.9981 dan weighted 0.9993. Kinerja hampir sempurna ini menunjukkan bahwa kedua model tersebut sangat andal dalam klasifikasi teks untuk deteksi risiko depresi dan kecemasan. *Voting Classifier*, sebagai agregasi dari keempat model sebelumnya, mencapai nilai tertinggi kedua setelah Random Forest dan AdaBoost, dengan *F1-score* macro 0.9962 dan weighted 0.9986. Meskipun sedikit di bawah dua model sebelumnya, *Voting Classifier* memberikan keseimbangan yang sangat baik antara akurasi, stabilitas, dan generalisasi, menjadikannya pilihan yang aman dan kuat untuk digunakan dalam sistem prediksi.

Secara keseluruhan, hasil evaluasi ini memperkuat temuan bahwa metode ensemble seperti Random Forest, AdaBoost, dan *Voting Classifier* unggul dalam klasifikasi risiko psikologis berbasis

teks. Sementara SVM dapat dijadikan alternatif kompetitif, dan Naïve Bayes lebih cocok digunakan sebagai baseline atau untuk skenario dengan kebutuhan komputasi yang rendah.

5. KESIMPULAN

Penelitian ini berhasil mengembangkan dan mengevaluasi model ensemble *Machine Learning* untuk deteksi gejala depresi pada postingan media sosial berbahasa Indonesia dengan menggunakan dataset “Depression and Anxiety in Twitter (ID)” yang berisi 6.980 data teks. Proses klasifikasi dilakukan dengan membandingkan lima algoritma, yaitu Naïve Bayes, SVM, Random Forest, AdaBoost, dan *Voting Classifier*. Model-model tersebut dilatih menggunakan teknik representasi teks TF-IDF dengan 3.000 fitur, serta dievaluasi menggunakan metrik akurasi, precision, recall, dan *F1-score*.

Hasil evaluasi menunjukkan bahwa Random Forest dan AdaBoost memberikan performa klasifikasi yang sangat tinggi dan hampir sempurna dengan nilai *F1-score* di atas 0,99 yang tergolong dalam kategori “Sangat Baik”. Model SVM juga menunjukkan performa yang kompetitif dengan *F1-score* mendekati 0,99. Naïve Bayes memiliki *F1-score* yang lebih rendah (sekitar 0,73 pada rata-rata makro) tetapi tetap menunjukkan stabilitas yang baik dan dapat digunakan sebagai model baseline. Sementara itu, *Voting Classifier* yang merupakan gabungan dari keempat model tersebut, memberikan kinerja terbaik secara menyeluruh dengan *F1-score* rata-rata makro sebesar 0,996 dan weighted sebesar 0,998 menunjukkan bahwa pendekatan ensemble mampu meningkatkan akurasi sekaligus menjaga kestabilan model. Keunggulan ini diperkuat oleh hasil visualisasi menggunakan *Confusion matrix* yang menunjukkan distribusi prediksi yang seimbang dan kesalahan klasifikasi yang sangat minimal. Matrix tersebut memperlihatkan bahwa *Voting Classifier* mampu meminimalkan kesalahan pada kedua kelas, yang sangat penting dalam konteks deteksi gangguan mental, di mana kesalahan klasifikasi dapat berimplikasi serius terhadap keputusan intervensi.

Selanjutnya Visualisasi menggunakan *learning curve* memperlihatkan bahwa *Voting Classifier*, SVM, dan Random Forest memiliki kemampuan generalisasi yang tinggi tanpa *overfitting* sedangkan performa AdaBoost yang terlalu sempurna menimbulkan indikasi kemungkinan kebocoran data atau fitur yang terlalu dominan.

Berdasarkan keseluruhan hasil *Voting Classifier* terbukti sebagai pendekatan paling efektif dalam mendeteksi risiko depresi dan kecemasan dari teks media sosial berbahasa Indonesia dengan kemampuannya dalam menggabungkan keunggulan dari masing-masing model dasar. Temuan ini memberikan kontribusi signifikan dalam pengembangan sistem pemantauan psikologis berbasis *Machine Learning* yang adaptif dan dapat digunakan oleh instansi pendidikan, lembaga kesehatan, maupun komunitas sosial untuk deteksi dini gangguan kesehatan mental.

6. SARAN

Distribusi label pada data menunjukkan ketidakseimbangan yang signifikan antara kelas normal dan depresi, dengan dominasi jumlah data pada kelas normal. Kondisi ini berpotensi menyebabkan bias model terhadap kelas mayoritas dan menurunkan akurasi dalam mendeteksi kasus depresi. Oleh karena itu, disarankan pada penelitian selanjutnya untuk menerapkan teknik penanganan data tidak seimbang, seperti *oversampling* (misalnya menggunakan SMOTE), *undersampling*, atau pemanfaatan algoritma yang sensitif terhadap ketidakseimbangan kelas guna meningkatkan performa model secara menyeluruh.

7. UCAPAN TERIMA KASIH

Penulis menyampaikan terima kasih kepada Politeknik Negeri Lampung khususnya Jurusan Teknologi Informasi Program Studi Teknologi Rekayasa Perangkat Lunak atas dukungan bimbingan akademik dan lingkungan penelitian yang kondusif selama proses penyusunan dan pelaksanaan penelitian ini. Penulis juga mengucapkan terima kasih kepada platform Kaggle yang telah menyediakan *dataset* Depression and Anxiety in Twitter (ID) sebagai dasar analisis dalam penelitian ini.

Ucapan terima kasih juga ditujukan kepada seluruh pihak yang telah memberikan dukungan, masukan, serta motivasi baik secara langsung maupun tidak langsung. Semua bantuan dan kontribusi yang diberikan sangat berarti dalam mendukung keberhasilan penelitian ini

DAFTAR PUSTAKA

- [1] K. Putri and M. A. Raharjaa *et al.*, “Implementasi Algoritma Support Vector Machine dalam Klasifikasi Deteksi Depresi dari Postingan pada Media Sosial,” *Jurnal Nasional Teknologi Informasi dan Aplikasinya*, vol. 2, no. 1, pp. 193–202, 2023.
- [2] B. Gavin, J. Lyne, and F. McNicholas, “Mental health and the COVID-19 pandemic,” 2020. doi: 10.1017/ipm.2020.72.
- [3] World Health Organization, “Mental health,” World Health Organization. Accessed: May 09, 2025. [Online]. Available: <https://www.who.int/health-topics/mental-health>
- [4] I. R. Ramadani, T. Fauziyah, and B.K. Rozzaq., “Depresi, Penyebab Dan Gejala Depresi Tryana Fauziyah,” *Bersatu : Jurnal Pendidikan Bhinneka Tunggal Ika*, vol. 2, no. 2, pp. 89–99, 2024.
- [5] A. Acuña Alvarado, E. Ramírez Zumbado, and M. F. Azofeifa Zumbado, “Depresión postparto,” *Revista Medica Sinergia*, vol. 6, no. 9, 2021, doi: 10.31434/rms.v6i9.712.
- [6] S. R. Vreijling *et al.*, “Sociodemographic, lifestyle and clinical characteristics of energy-related depression symptoms: A pooled analysis of 13,965 depressed cases in 8 Dutch cohorts,” *J Affect Disord*, vol. 323, 2023, doi: 10.1016/j.jad.2022.11.005.
- [7] L. Munira, P. Liamputtong, and P. Viwattanakulvanid, “Barriers and facilitators to access mental health services among people with mental disorders in Indonesia: A qualitative study,” *Belitung Nurs J*, vol. 9, no. 2, 2023, doi: 10.33546/bnj.2521.
- [8] M. A. Subu *et al.*, “Types of stigma experienced by patients with mental illness and mental health nurses in Indonesia: a qualitative content analysis,” *Int J Ment Health Syst*, vol. 15, no. 1, 2021, doi: 10.1186/s13033-021-00502-x.
- [9] We Are Social and DataReportal, “DIGITAL 2024 INDONESIA.”
- [10] S. T. Rabani, Q. R. Khan, and A. M. U. D. Khanday, “Quantifying Suicidal Ideation on Social Media using *Machine Learning*: A Critical Review,” *Iraqi Journal of Science*, vol. 62, no. 11, 2021, doi: 10.24996/ij.s.2021.62.11.29.
- [11] G. F. Situmorang and R. Purba, “Deteksi Potensi Depresi dari Unggahan Media Sosial X Menggunakan IndoBERT,” *Building of Informatics, Technology and Science (BITS)*, vol. 6, no. 2, pp. 649–661, Sep. 2024, doi: 10.47065/bits.v6i2.5496.
- [12] K. S. Nugroho, I. Akbar, A. N. Suksmawati, and Istiadi, “Deteksi Depresi dan Kecemasan Pengguna Twitter Menggunakan Bidirectional LSTM,” no. Ciastech, pp. 287–296, 2023.
- [13] F. T. Fathin and W. Maharani, “Comparison of Random Forest and Decision Tree for Depression Detection Using Interaction Patterns,” *Technology and Science (BITS)*, vol. 6, no. 4, 2025, doi: 10.47065/bits.v6i4.6660.
- [14] S. Method, “Deteksi Dini Gangguan Kesehatan Mental berdasarkan Sentimen menggunakan Metode Stacking Early Detection of Mental Health Disorders based on Sentiment using,” vol. 14, pp. 271–280, 2025.
- [15] N. Maldini *et al.*, “Sistemasi: Jurnal Sistem Informasi Deteksi Dini Gangguan Kesehatan Mental berdasarkan Sentimen menggunakan Metode Stacking Early Detection of Mental Health

- Disorders based on Sentiment using Stacking Method,” 2025. [Online]. Available: <http://sistemasi.ftik.unisi.ac.id>
- [16] A. Ananda Hapsari, A. Syafei Nursuwanda, H. Zuhriyah, and D. Junesco Vresdian, “Klasifikasi Kesehatan Mental Mahasiswa Model TMAS dengan Algoritma Decision Tree, Logistic Regression, dan Random Forest,” *Jurnal INTEK*, vol. 7, 2024.
- [17] D. Rofianto, E. Safitri, K. Amaliah, J. Fitra, and A. Hijriani, “Cyber Threat Detection Using an Ensemble Model Approach for Phishing Website Identification ARTICLE INFORMATION ABSTRACT,” 2024. [Online]. Available: <http://innovatics.unsil.ac.id>