

ANALISIS KOMPARATIF ALGORITMA MACHINE LEARNING UNTUK KLASIFIKASI BINER STATUS KESEHATAN MENTAL PADA TEKS MEDIA SOSIAL

Ahmad Rijal Khairullah, Fuad Kanifudin, Habil Al-Farisy,
Windah Kurnia Sari, Ken Ditha Tania, Fathoni

Sistem Informasi, Universitas Sriwijaya
Jalan Srijaya Negara, Palembang, Indonesia
arijalk19@gmail.com

ABSTRAK

Kesehatan mental merupakan isu global yang krusial, di mana media sosial menjadi wadah utama bagi pengguna untuk mengekspresikan kondisi psikologisnya. Namun, tingginya volume data teks dan ketidakseimbangan kelas gangguan mental membuat deteksi secara manual menjadi tidak efisien dan rentan terhadap subjektivitas. Penelitian ini bertujuan untuk menganalisis dan membandingkan kinerja algoritma *Machine Learning*, yakni *Logistic Regression* dan *Random Forest*, dalam mengklasifikasikan status kesehatan mental secara biner. Metodologi yang digunakan meliputi pemrosesan 53.043 data teks berbahasa Inggris yang disederhanakan menjadi dua kelas, yaitu *normal* dan *not_normal*. Teks tersebut diekstraksi menggunakan *Term Frequency-Inverse Document Frequency* (TF-IDF) dengan batasan 8.000 fitur *unigram* dan *bigram*, kemudian dibagi menjadi 80% data latih dan 20% data uji secara stratifikasi. Hasil pengujian menunjukkan bahwa kedua model berkinerja stabil dan sangat baik. *Random Forest* sedikit lebih unggul dengan tingkat akurasi 93,68%, dibandingkan *Logistic Regression* yang mencapai 93,65%. Kedua model sama-sama mencatatkan *F1-Score* sebesar 94%. Pendekatan ini terbukti efisien dan sangat akurat sebagai landasan sistem deteksi dini kesehatan mental secara otomatis.

Kata kunci : Kesehatan Mental, Klasifikasi Teks, Machine Learning, Random Forest, Logistic Regression, TF-IDF.

1. PENDAHULUAN

Kesehatan mental merupakan salah satu isu kesehatan global yang krusial pada era digital saat ini. Seiring dengan peningkatan penggunaan internet, media sosial telah bertransformasi menjadi platform utama bagi banyak individu untuk mengekspresikan emosi, keluhan, tekanan psikologis, dan kondisi kehidupan sehari-hari [1]. Teks yang ditulis oleh pengguna di media sosial mengandung pola linguistik yang spesifik dan dapat menjadi sumber data sekunder yang berharga untuk mendeteksi berbagai jenis gangguan kesehatan mental, seperti depresi, kecemasan (*anxiety*), gangguan bipolar, stres, hingga kecenderungan bunuh diri (*suicidal*). Namun, lonjakan volume data teks (*big data*) yang sangat besar di media sosial membuat proses identifikasi dan pemantauan kondisi psikologis secara manual menjadi tidak efisien, memakan waktu, dan rentan terhadap subjektivitas.

Untuk mengatasi tantangan tersebut, penerapan *Natural Language Processing* (NLP) dan *Machine Learning* menjadi solusi otomatisasi yang sangat menjanjikan. Berbagai penelitian sebelumnya telah banyak mengeksplorasi penggunaan algoritma komputasi untuk mendeteksi masalah mental. Sejumlah studi terkini berfokus pada pendekatan *Deep Learning* seperti arsitektur berbasis *Transformer* yang memiliki performa tinggi, namun seringkali terkendala oleh kebutuhan komputasi yang sangat besar dan proses pelatihan model yang kompleks [2], [3]. Di sisi lain, beberapa penelitian komparatif membuktikan bahwa algoritma *Machine Learning* konvensional

seperti *Logistic Regression*, *Support Vector Machine* (SVM), dan *Random Forest* tetap tangguh (*robust*) dan mampu memberikan metrik akurasi yang kompetitif dan efisiensi yang tinggi pada dataset medis, di mana penelitian sebelumnya [9], [10] menunjukkan bahwa model ensemble seperti *Random Forest* seringkali mengungguli model tunggal dalam menangani data teks yang kompleks.

Salah satu tantangan utama dalam klasifikasi teks kesehatan mental adalah kompleksitas dari kondisi psikologis itu sendiri. Seringkali, status kesehatan mental tidak hanya bersifat pengelompokan dua arah (biner), melainkan terdiri dari berbagai kategori gangguan yang saling beririsan secara linguistik. Oleh karena itu, terdiri dari berbagai kategori gangguan yang saling beririsan secara linguistik [6]. Namun, tingginya tumpang tindih leksikal pada kategori multikelas seringkali menurunkan akurasi model, sehingga pendekatan penyederhanaan menjadi klasifikasi biner sering digunakan untuk meningkatkan performa deteksi dini secara lebih stabil [11].

Penelitian ini bertujuan untuk melakukan analisis komparatif terhadap kinerja algoritma *Machine Learning*, khususnya *Logistic Regression* dan *Random Forest*, dalam melakukan analisis komparatif terhadap kinerja algoritma *Machine Learning*, khususnya *Logistic Regression* dan *Random Forest*, dalam melakukan klasifikasi biner pada data teks media sosial. Meskipun dataset awal mencakup tujuh kategori, penelitian ini memfokuskan pengujian pada pemetaan teks ke dalam status 'Normal' dan 'Not Normal' untuk mendapatkan model yang lebih presisi

dan efisien dalam implementasi nyata. Hasil dari penelitian ini diharapkan dapat memberikan wawasan empiris mengenai pemilihan algoritma klasifikasi teks yang paling optimal, stabil, dan efisien sebagai landasan pengembangan sistem deteksi dini kesehatan mental di masa depan.

2. TINJAUAN PUSTAKA

Penelitian terkait pemanfaatan data media sosial untuk mendeteksi masalah kesehatan mental telah banyak dilakukan seiring dengan pesatnya pertumbuhan jejak digital pengguna. Media sosial terbukti mengandung sinyal linguistik yang kuat terkait kondisi emosional dan psikologis seseorang. Beberapa studi sebelumnya berupaya melakukan klasifikasi berbagai jenis gangguan mental, seperti depresi, kecemasan, hingga kecenderungan bunuh diri dari unggahan pengguna [2]. Meskipun klasifikasi multi kelas pada tingkat granular memberikan informasi yang detail, pendekatan ini sering menghadapi kendala ketidakseimbangan data (*imbalanced data*) dan tumpang tindih fitur bahasa antara gangguan, sehingga penggabungan kelas menjadi biner (seperti normal dan tidak normal) sering digunakan untuk meningkatkan performa deteksi awal, meskipun beberapa studi terbaru [3], [4] tetap mengeksplorasi potensi klasifikasi multikelas pada tingkat granular.

Dalam ranah *Natural Language Processing* (NLP) untuk klasifikasi teks kesehatan mental, perbandingan antara algoritma *Machine Learning* tradisional dan *Deep Learning* menjadi fokus banyak peneliti. Model berbasis *Transformer* seperti BERT dan RoBERTa memang terbukti menghasilkan akurasi yang sangat tinggi berkat kemampuannya memahami konteks kalimat secara mendalam [5], [6], [7]. Namun, model *Deep Learning* tersebut memiliki kelemahan dalam hal tingginya beban komputasi dan kompleksitas waktu pelatihan (*training time*) [8], [9]. Sebagai alternatif yang efisien, algoritma *Machine Learning* konvensional masih sangat relevan. Studi komparatif pada berbagai domain, mulai dari prediksi churn pegawai [5] hingga diagnosis medis [9], [10], membuktikan bahwa algoritma seperti Logistic Regression dan Random Forest mampu mencapai kinerja komputasi yang cepat dengan metrik evaluasi yang sangat kompetitif.

Kinerja *Machine Learning* sangat bergantung pada teknik representasi teks. *Term Frequency-Inverse Document Frequency* (TF-IDF) merupakan metode ekstraksi fitur yang banyak diadopsi karena kemampuannya dalam memberikan bobot pada kata-kata yang paling diskriminatif dalam sebuah dokumen teks [11], [12]. Dengan memadukan TF-IDF dan algoritma seperti *Random Forest* atau *Logistic Regression*, model klasifikasi dapat secara efektif mengisolasi kata kunci spesifik yang mengindikasikan status kesehatan mental pengguna secara efisien tanpa memerlukan infrastruktur perangkat keras yang masif.

3. METODE PENELITIAN

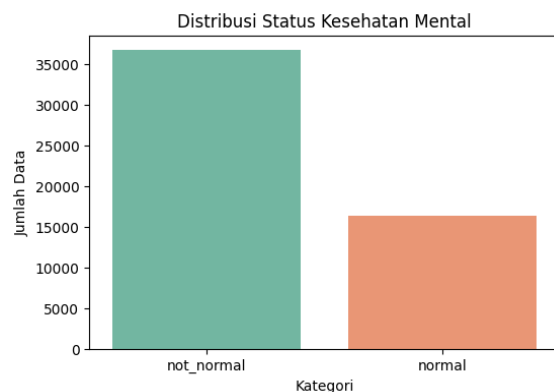
Penelitian ini dilakukan melalui beberapa tahapan sistematis untuk mengklasifikasikan status kesehatan mental dari teks media sosial. Tahapan tersebut meliputi pengumpulan dataset, pra pemrosesan teks, ekstraksi fitur menggunakan *Term Frequency-Inverse Document Frequency* (TF-IDF), pembagian data latih dan data uji, pelatihan model *Machine Learning*, hingga evaluasi kinerja model.



Gambar 1. Alur metodologi penelitian

3.1. Pengumpulan Data

Dataset yang digunakan dalam penelitian ini merupakan data sekunder berupa teks pernyataan (*statement*) berbahasa Inggris yang berkaitan dengan kondisi kesehatan mental pengguna media sosial. Dataset diperoleh dari “Dataset Sentiment Analysis for Mental Health” yang tersedia secara publik pada platform Kaggle.. Total keseluruhan data berjumlah 53.043 baris teks. Pada data aslinya, terdapat tujuh kategori label status psikologis, yaitu *Normal*, *Depression*, *Suicidal*, *Anxiety*, *Bipolar*, *Stress*, dan *Personality disorder*. Untuk mengoptimalkan kinerja model, keenam label yang menunjukkan indikasi gangguan mental digabungkan ke dalam satu kelas terpadu, yaitu *not_normal* (tidak normal), sedangkan kelas *Normal* tetap dipertahankan. Pemetaan ini menghasilkan klasifikasi biner dengan jumlah 36.692 data untuk kelas *not_normal* dan 16.351 data untuk kelas *normal*.



Gambar 2. Distribusi data status kesehatan mental kelas biner

Berdasarkan Gambar 1, dapat dilihat bahwa terdapat ketidakseimbangan jumlah data antara kelas normal dan tidak normal pada dataset sebelum dilakukan pra pemrosesan. Proporsi kelas *not normal* terlihat jauh lebih mendominasi dibandingkan kelas *normal*. Kondisi ketidakseimbangan (*imbalanced data*) ini menjadi pertimbangan penting dalam proses pembagian data selanjutnya untuk memastikan model tidak bias terhadap kelas mayoritas.

3.2. Data Preprocessing

Data teks yang diperoleh dari media sosial umumnya masih berupa teks mentah yang mengandung berbagai karakter yang tidak relevan dengan proses analisis. Oleh karena itu, sebelum dilakukan proses pemodelan machine learning, data terlebih dahulu melalui tahap data preprocessing untuk meningkatkan kualitas teks. Tahapan preprocessing yang dilakukan dalam penelitian ini meliputi case folding, yaitu mengubah seluruh huruf menjadi huruf kecil, penghapusan URL, mention, hashtag, angka, dan tanda baca, serta stopword removal untuk menghilangkan kata-kata umum yang tidak memiliki kontribusi signifikan terhadap proses klasifikasi. Setelah proses pembersihan teks selesai dilakukan, label target kemudian dikonversi ke dalam bentuk numerik menggunakan Label Encoder agar dapat diproses oleh algoritma machine learning.

3.3. Implementasi TF-IDF

Metode Term Frequency-Inverse Document Frequency (TF-IDF) merupakan teknik ekstraksi fitur yang digunakan untuk mengubah data teks menjadi representasi numerik sehingga dapat diproses oleh algoritma machine learning. Metode ini memberikan bobot pada setiap kata berdasarkan tingkat kemunculannya dalam suatu dokumen serta frekuensinya dalam keseluruhan kumpulan dokumen. Term Frequency (TF) digunakan untuk mengukur seberapa sering suatu kata muncul dalam dokumen, sedangkan Inverse Document Frequency (IDF) digunakan untuk menentukan tingkat kepentingan kata terhadap seluruh dokumen. Dengan menggabungkan kedua komponen tersebut, TF-IDF mampu menghasilkan representasi fitur yang lebih informatif sehingga dapat membantu meningkatkan performa model klasifikasi.

3.4. Training Model

Pada tahap ini dilakukan pelatihan model machine learning untuk melakukan klasifikasi status kesehatan mental berdasarkan fitur teks yang telah diekstraksi. Penelitian ini membandingkan dua algoritma klasifikasi yaitu Logistic Regression dan Random Forest Classifier.

a. Logistic Regression

Logistic Regression merupakan algoritma klasifikasi berbasis pendekatan statistik yang digunakan untuk memodelkan hubungan antara variabel independen dengan probabilitas suatu kejadian pada data biner. Metode ini bekerja dengan menghitung kemungkinan suatu data termasuk ke dalam salah satu kelas berdasarkan nilai fitur yang dimiliki, kemudian menentukan kelas dengan probabilitas tertinggi sebagai hasil prediksi [11]. Logistic Regression dikembangkan dari metode regresi linier dengan memanfaatkan fungsi logistik untuk mengubah hasil perhitungan menjadi nilai probabilitas antara 0 dan 1 sehingga dapat digunakan dalam proses klasifikasi [10].

Dalam penelitian ini, algoritma Logistic Regression digunakan sebagai model pembandingan untuk melakukan prediksi klasifikasi serta mengevaluasi kinerjanya menggunakan metrik accuracy, precision, recall, dan F1-score guna mengetahui efektivitas model dalam mengklasifikasikan data.

b. Random Forest

Random Forest merupakan metode ensemble learning yang dikembangkan dari algoritma pohon keputusan (decision tree) dengan membangun banyak pohon keputusan pada subset data yang berbeda, kemudian menggabungkan hasil prediksinya untuk menghasilkan keputusan akhir yang lebih akurat. Pendekatan ini mampu meningkatkan kinerja model serta mengurangi risiko overfitting karena setiap pohon dibangun dari sampel data yang berbeda. Dalam penelitian ini, Random Forest digunakan sebagai salah satu algoritma klasifikasi untuk memprediksi status kesehatan mental pada teks media sosial yang telah direpresentasikan dalam bentuk fitur numerik. Kinerja model kemudian dievaluasi menggunakan metrik accuracy, precision, recall, dan F1-score untuk mengetahui efektivitas Random Forest dalam mengklasifikasikan data dibandingkan dengan algoritma lain seperti Logistic Regression [10], [11].

3.5. Evaluasi Model

Kinerja dari kedua algoritma klasifikasi dievaluasi menggunakan Confusion Matrix serta beberapa metrik evaluasi standar, yaitu Accuracy, Precision, Recall, dan F1-Score. Confusion Matrix digunakan untuk menggambarkan hasil prediksi model terhadap data uji berdasarkan empat komponen utama, yaitu True Positive (TP), True Negative (TN), False Positive (FP), dan False Negative (FN). Evaluasi ini bertujuan untuk mengukur tingkat keberhasilan serta kemampuan model dalam mengklasifikasikan teks yang mengindikasikan status kesehatan mental secara tepat. Tabel 1 menunjukkan Confusion Matrix sebagai berikut.

Tabel 1. Confusion Matrix

Prediksi	Prediksi positif	Prediksi negatif
Positive	True Positive(TP)	False Positive(FP)
Negative	False Negative(FN)	True Negative(TN)

Berdasarkan nilai pada Confusion Matrix, metrik evaluasi model dapat dihitung menggunakan persamaan berikut.

Accuracy

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN}$$

Precision

$$Precision = \frac{TP}{TP + FP}$$

Recall

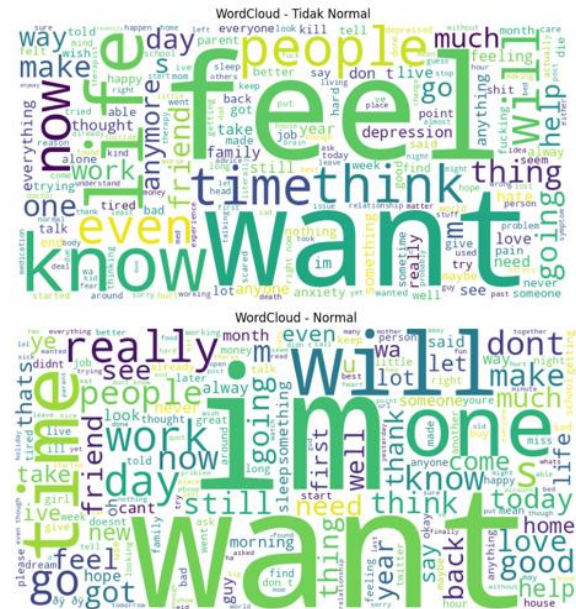
$$Recall = \frac{TP}{TP + FN}$$

F1-Score

$$F1Score = 2 \times \frac{Precision \times Recall}{Precision + Recall}$$

4. HASIL DAN PEMBAHASAN

Sebelum melakukan pemodelan, dilakukan analisis visual terhadap kata-kata yang paling sering muncul pada dataset menggunakan metode *Word Cloud*. Visualisasi ini bertujuan untuk mengamati perbedaan pola linguistik antara pengguna dengan indikasi kondisi mental "normal" dan "tidak normal".



Gambar 3. Visualisasi word cloud pada kelas normal dan tidak normal

Dari grafik pada Gambar 2, dapat dilihat kemunculan kata-kata yang paling dominan pada masing-masing kelas. Pada kelas *not_normal* (tidak normal), kata-kata yang sering muncul umumnya merepresentasikan fokus pada emosi internal, beban pikiran, atau luapan perasaan pengguna secara mendalam (seperti munculnya kata '*feel*', '*like*', '*want*', dan '*think*' secara dominan). Sebaliknya, pada kelas *normal*, kata-kata yang mendominasi cenderung lebih berfokus pada aktivitas keseharian, orientasi waktu saat ini, serta ekspresi yang lebih positif (seperti dominasi kata '*day*', '*today*', '*love*', dan '*good*'). Perbedaan pola kata yang signifikan ini mengonfirmasi bahwa representasi fitur berbasis TF-IDF mampu menangkap konteks kalimat yang membedakan kedua kelas tersebut secara efektif sebelum dilakukan klasifikasi oleh algoritma.

4.1. Hasil Pengujian Model

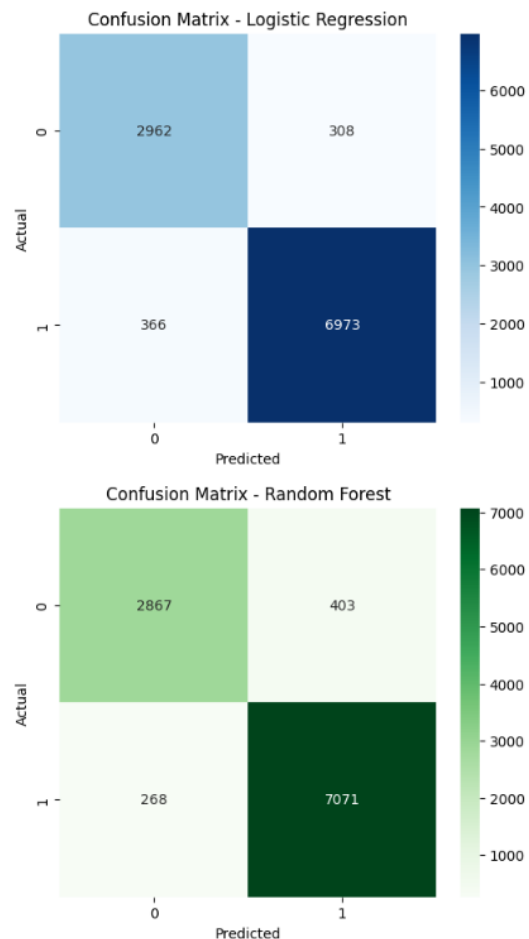
Pengujian dilakukan menggunakan dua algoritma *Machine Learning*, yaitu *Logistic Regression* dan *Random Forest*. Model dilatih menggunakan 80% data latih dan diuji pada 20% data uji (sebanyak 10.609 data teks). Kinerja model dievaluasi secara komparatif menggunakan metrik akurasi, presisi, *recall*, dan *f1-score*.

Tabel 2. Perbandingan hasil evaluasi model klasifikasi

Algoritma	Akurasi	Presisi	Recall	F1-Score
Logistic Regression	93.65%	94.00%	94.00%	94.00%
Random Forest	93.68%	94.00%	94.00%	94.00%

Dari Tabel 2 dapat dilihat bahwa kedua model klasifikasi memberikan kinerja yang sangat baik dan stabil dalam membedakan teks, dengan nilai akurasi secara keseluruhan di atas 93%. *Random Forest* unggul secara tipis pada metrik akurasi dengan nilai 93.68%, sedangkan *Logistic Regression* memperoleh akurasi sebesar 93.65%. Kedua model juga mencatatkan nilai rata-rata terbobot (*weighted average*) untuk presisi, *recall*, dan *f1-score* yang identik, yakni sebesar 94%.

4.2. Pembahasan



Gambar 3. *Confusion Matrix* model *Logistic Regression* dan *Random Forest*.

Meskipun kedua algoritma menunjukkan performa yang hampir setara, *Random Forest* sedikit lebih tangguh dalam memprediksi kelas teks berkat mekanisme *ensemble learning* yang membangun ratusan pohon keputusan (*decision trees*) untuk menekan varians prediksi. Di sisi lain, *Logistic Regression* membuktikan efisiensinya. Walaupun

merupakan algoritma yang lebih sederhana dan ringan secara komputasi, algoritma ini tetap mampu memisahkan kelas teks secara optimal berkat kombinasi ekstraksi fitur unigram dan bigram dari TF-IDF yang menghasilkan batas keputusan (*decision boundary*) yang jelas.

Berdasarkan Gambar 3, terlihat detail distribusi klasifikasi pada 10.609 data uji. Pada model *Logistic Regression*, terdapat 2.969 data normal dan 6.966 data *not_normal* yang berhasil diprediksi dengan benar (*True Positives* dan *True Negatives*). Namun, model ini melakukan kesalahan prediksi dengan mengklasifikasikan 301 data *normal* sebagai *not_normal* (*False Positive*) dan 373 data *not_normal* sebagai *normal* (*False Negative*).

Di sisi lain, model *Random Forest* menunjukkan pola prediksi yang sedikit berbeda. Meskipun *Random Forest* lebih banyak salah menebak data *normal* menjadi *not_normal* (386 *False Positive*), model ini terbukti jauh lebih peka dalam mendeteksi kelas *not_normal*. Hal ini dibuktikan dengan jumlah prediksi benar pada kelas *not_normal* yang lebih tinggi (7.054 data) dan tingkat kesalahan *False Negative* yang lebih rendah (285 data) dibandingkan *Logistic Regression*. Dalam konteks deteksi dini kesehatan mental, kemampuan *Random Forest* untuk menekan angka *False Negative* ini sangat bernilai, karena risiko membiarkan individu yang berpotensi memiliki gangguan mental terdeteksi sebagai "normal" dapat diminimalisir dengan lebih baik.

5. KESIMPULAN DAN SARAN

Berdasarkan hasil penelitian dan pengujian yang telah dilakukan, dapat disimpulkan bahwa penerapan algoritma *Machine Learning* terbukti sangat efektif dalam mengklasifikasikan status kesehatan mental pengguna media sosial berdasarkan teks. Pendekatan penggabungan kelas menjadi biner (*normal* dan *tidak normal*) yang dipadukan dengan teknik ekstraksi fitur *Term Frequency-Inverse Document Frequency* (TF-IDF) mampu menghasilkan performa deteksi yang sangat baik dan stabil. Dari kedua algoritma yang dibandingkan, *Random Forest Classifier* menunjukkan kinerja yang sedikit lebih unggul dengan tingkat akurasi sebesar 93,68%, sedangkan *Logistic Regression* memperoleh akurasi sebesar 93,65%, di mana keduanya sama-sama mencatatkan *F1-Score* sebesar 94%. Sebagai saran untuk penelitian selanjutnya, peneliti direkomendasikan untuk mengeksplorasi penggunaan teknik penyeimbangan data lanjutan (seperti SMOTE) atau arsitektur *Deep Learning* agar model komputasi dapat melakukan klasifikasi secara spesifik pada level multikelas (seperti membedakan depresi, kecemasan, stres, dan bipolar secara langsung) tanpa perlu menyederhanakan kategori, serta menguji metode representasi teks lain seperti *Word Embeddings* guna menangkap relasi semantik antarkata yang lebih kompleks.

DAFTAR PUSTAKA

- [1] A. Hussain and A. Aslam, "Hate speech against women and immigrants: A comparative analysis of machine learning and text embedding techniques", *JART*, vol. 22, no. 4, pp. 548–559, Aug 2024, doi: 10.22201/icat.24486736e.2024.22.4.2466
- [2] M. Kumar, L. Khan, and A. Choi, "RAMHA: A Hybrid Social Text-Based Transformer with Adapter for Mental Health Emotion Classification," *Mathematics*, vol. 13, no. 18, Sep. 2025, doi: 10.3390/math13182918.
- [3] M. Sao and H. J. Lim, "MIROBERTa: Mental Illness Text Classification With Transfer Learning on Subreddits," *IEEE Access*, vol. 12, pp. 197454–197466, 2024, doi: 10.1109/ACCESS.2024.3522465.
- [4] L. Bendebane, Z. Laboudi, A. Saighi, H. Al-Tarawneh, A. Ouannas, and G. Grassi, "A Multi-Class Deep Learning Approach for Early Detection of Depressive and Anxiety Disorders Using Twitter Data," *Algorithms*, vol. 16, no. 12, Dec. 2023, doi: 10.3390/a16120543.
- [5] I. N. A. Palupi, M. F. F. Mardianto, I. Yuadi, and B. Mariyadi, "STUDI KOMPARATIF MODEL MACHINE LEARNING DALAM MEMPREDIKSI KETERLAMBATAN PEGAWAI: LOGISTIC REGRESSION, SVM, DAN RANDOM FOREST," *J@ti Undip: Jurnal Teknik Industri*, vol. 21, no. 1, pp. 76–87, Jan 2026 doi: 10.14710/jati.21.1.76-87
- [6] R. Dash, S. Udgata, R. K. Mohapatra, V. Dash, and A. Das, "A Deep Learning Approach to Unveil Types of Mental Illness by Analyzing Social Media Posts," *Mathematical and Computational Applications*, vol. 30, no. 3, Jun. 2025, doi: 10.3390/mca30030049.
- [7] M. Nouman, S. Y. Khoo, M. A. P. Mahmud, and A. Z. Kouzani, "A hybrid BERT-BiRNN framework for mental health prediction using textual data," *Natural Language Processing Journal*, vol. 12, Sep. 2025, doi: 10.1016/j.nlp.2025.100165.
- [8] Z. Ding *et al.*, "Trade-offs between machine learning and deep learning for mental illness detection on social media," *Sci. Rep.*, vol. 15, no. 1, Dec. 2025, doi: 10.1038/s41598-025-99167-6.
- [9] W. Zhu, "Comparison of Prediction Models for Diabetes: Accuracy Analysis of Logistic Regression, Random Forest, and Back Propagation Neural Network," *Theoretical and Natural Science*, vol. 119, no. 1, pp. 69–76, Jun. 2025, doi: 10.54254/2753-8818/2025.au24412.
- [10] Z. Z. Hulaifah Al Abrori and E. R. Subhiyakti, "Analisis Komparatif Akurasi Prediksi Kanker Payudara Menggunakan Algoritma Random Forest dan Logistic Regression," *Jurnal Algoritma*, vol. 22, no. 1, pp. 300–311, May 2025, doi: 10.33364/algoritma/v.22-1.2164.

- [11] A. M. Majid, K. Imelda, and I. Nawangsih, "Komparasi Algoritma Klasifikasi Machine Learning dan Penerapan Metode Ensemble Stacking untuk Menganalisis Sentimen Kesehatan Mental," *SKANIKA: Sistem Komputer dan Teknik Informatika*, vol. 8, no. 2, pp. 293–304, 2025, [Online]. Available: <https://www.kaggle.com/datasets/suchintikasarkar/sentiment-analysis-for-mental-health>
- [12] A. Karamat, M. Imran, M. U. Yaseen, R. Bukhsh, S. Aslam, and N. Ashraf, "A Hybrid Transformer Architecture for Multiclass Mental Illness Prediction Using Social Media Text," *IEEE Access*, vol. 13, pp. 12148–12167, 2025, doi: 10.1109/ACCESS.2024.3519308