

Mendeteksi Unsur Depresi pada Unggahan Media Sosial Menggunakan Metode *Machine Learning* dengan Optimasi Berbasis Inspirasi Alam

Zein Rizky Santoso^{1*}, Aji Hamim Wigena², Anang Kurnia³

¹²³Sekolah Sains Data, Matematika, dan Informatika - IPB University

*Corresponding author, email: rizkyzein@apps.ipb.ac.id

Abstract

Social media has now become an inseparable part of everyday life, including in expressing emotions and mental states. One popular platform is X (formerly Twitter), where many users indirectly share signs of depression. This study develops a classification model to detect indications of depression in social media posts, using machine learning algorithms and feature selection techniques based on nature-inspired algorithms. The classification algorithms used include Naïve Bayes, *k*-Nearest Neighbors (*k*-NN), Decision Tree, Random Forest, and XGBoost. Each algorithm is combined with feature selection techniques using Particle Swarm Optimization (PSO), Bat Algorithm (BA), and Flamingo Search Algorithm (FSA). The models are evaluated based on accuracy, precision, recall, *F1*-score, and the number of features used. The results show that the combination of the Random Forest method with FSA-based feature selection (RF-FSA) delivers the best performance, with an accuracy of 82.2%, balanced precision and recall, and efficient feature usage. Another strong alternative is XGBoost with FSA (XGB-FSA), although it requires more features and longer computational time. This study demonstrates that selecting the right feature selection algorithm, particularly FSA, can significantly improve both the accuracy and efficiency of depression text classification models. The resulting model is expected to serve as a useful tool for early detection of depression symptoms from social media posts, allowing for quicker and more targeted interventions.

Keywords: Depression, Social media, machine learning, feature selection, Flamingo Search Algorithm, text classification

Abstrak

Media sosial saat ini menjadi bagian yang tidak terpisahkan dari kehidupan masyarakat, termasuk dalam mengekspresikan perasaan dan kondisi mental. Salah satu platform populer adalah X (sebelumnya *Twitter*), di mana banyak pengguna secara tidak langsung membagikan tanda-tanda depresi. Model klasifikasi dibangun untuk mendeteksi indikasi depresi pada unggahan, menggunakan algoritma *machine learning* dan teknik seleksi fitur berbasis algoritma yang terinspirasi dari alam. Beberapa algoritma klasifikasi yang digunakan antara lain *Naïve Bayes*, *k*-Nearest Neighbors (*k*-NN), *Decision Tree*, *Random Forest*, dan *XGBoost*. Masing-masing algoritma akan dikombinasikan dengan teknik seleksi fitur menggunakan *Particle Swarm Optimization* (PSO), *Bat Algorithm* (BA), dan *Flamingo Search Algorithm* (FSA). Evaluasi model dilakukan berdasarkan akurasi, presisi, *recall*, *F1*-score, dan jumlah fitur yang digunakan. Hasilnya menunjukkan bahwa kombinasi metode *Random Forest* dengan seleksi fitur menggunakan FSA (RF-FSA) memberikan performa terbaik, dengan akurasi sebesar 82.2%, presisi dan *recall* yang seimbang, serta penggunaan fitur yang efisien. Alternatif lain yang juga cukup kuat adalah *XGBoost* dengan FSA (XGB-FSA), meskipun membutuhkan fitur yang lebih banyak dan waktu komputasi yang lebih besar. Penggunaan algoritma seleksi fitur yang tepat, khususnya FSA, mampu

meningkatkan akurasi dan efisiensi model dalam klasifikasi teks depresi. Model yang dihasilkan diharapkan dapat digunakan sebagai alat bantu deteksi dini gejala depresi dari unggahan media sosial, sehingga memungkinkan intervensi lebih cepat dan tepat sasaran.

Kata Kunci: Depression, Social Media, Machine Learning, Feature Selection, Flamingo Search Algorithm, Text Classification.

1. Pendahuluan

Perkembangan teknologi informasi di Indonesia telah memberikan dampak signifikan terhadap kehidupan masyarakat. Salah satu bentuk nyata dari kemajuan ini adalah kemunculan media sosial, yang memungkinkan penggunanya untuk berinteraksi, berbagi, dan menciptakan konten [1]. Platform seperti Facebook, Instagram, Line, WhatsApp, dan X (sebelumnya Twitter) kini digunakan oleh hampir semua kalangan. Media sosial menawarkan banyak manfaat, mulai dari memperlancar komunikasi hingga menjadi sarana hiburan dan pembelajaran. Namun, di balik manfaat tersebut, terdapat sisi negatif yang tidak bisa diabaikan, terutama terhadap kesehatan mental.

Penggunaan media sosial yang intens dapat berdampak buruk pada kesehatan mental, seperti stres, kecemasan, dan depresi [2]. Namun, isu kesehatan mental masih belum menjadi prioritas di negara berkembang seperti Indonesia [3]. Riset Kesehatan Dasar (Riskesdas) 2018 mencatat lebih dari 19 juta penduduk Indonesia mengalami gangguan mental emosional, dan lebih dari 12 juta di antaranya mengalami depresi [4]. Depresi sendiri merupakan gangguan psikologis serius yang dapat disebabkan oleh berbagai faktor, termasuk tekanan sosial, isolasi, serta penggunaan media sosial yang berlebihan [5].

Media sosial juga menjadi ruang ekspresi bagi individu yang mengalami tekanan emosional. Unggahan yang mencerminkan suasana hati negatif bisa menjadi indikator depresi [6]. Namun, mendeteksi hal tersebut secara manual sulit dilakukan, sehingga diperlukan pendekatan teknologi, seperti machine learning. Salah satu metode yang pernah digunakan adalah *Naïve Bayes*, namun akurasi yang dihasilkan masih rendah (70%) karena keterbatasan dalam pemilihan fitur yang relevan [6].

Untuk mengatasi hal ini, digunakan pendekatan seleksi fitur dengan algoritma metaheuristik seperti *Particle Swarm Optimization* (PSO), yang terbukti mampu meningkatkan akurasi hingga 91% [7]. Selain PSO, algoritma seperti Bat Algorithm (BA) [8] dan *Flamingo Search Algorithm* (FSA) [9] juga mulai dikembangkan karena potensinya dalam meningkatkan kinerja model. Metode klasifikasi lain seperti *K-Nearest Neighbors* (K-NN), *Decision Tree*, *Random Forest*, dan *XGBoost* juga menjadi alternatif karena keunggulan masing-masing dalam menangani data dan membuat prediksi.

Lima metode machine learning (*Naïve Bayes*, *k-NN*, *Decision Tree*, *Random Forest*, dan *XGBoost*) dibandingkan, dengan tiga algoritma metaheuristik (PSO, BA, dan FSA) untuk seleksi fitur dalam mendeteksi gejala depresi pada data dari platform X.

Evaluasi dilakukan menggunakan metrik *Confusion Matrix* seperti akurasi, presisi, dan recall, serta waktu komputasi. Melalui pendekatan ini, diharapkan dapat ditemukan kombinasi metode dan algoritma yang optimal untuk deteksi dini depresi secara otomatis di media sosial, yang dapat membantu intervensi lebih cepat dan tepat sasaran.

2. Material dan Metode

Depresi adalah gangguan suasana hati yang ditandai oleh kesedihan mendalam, putus asa, kehilangan minat, dan perasaan tidak berharga, serta berkaitan dengan ketidakseimbangan neuro- transmitter seperti serotonin dan norepinefrin [10]. Depresi diklasifikasikan menjadi ringan, sedang, dan berat, dengan dampak besar pada aspek sosial dan kesehatan pada tingkat. Trauma masa kecil sering menjadi pemicu, dengan gejala seperti suasana hati rendah, gangguan konsentrasi, isolasi sosial, dan pikiran bunuh diri [11]. Menurut DSM-V (2013), diagnosis depresi mencakup aspek afektif, kognitif, dan fisik, yang dirangkum dalam Tabel 1.

Analisis sentimen adalah metode untuk mengidentifikasi dan mengklasifikasikan sentimen dalam teks, apakah positif, negatif, atau netral. Teknik ini banyak digunakan dalam pengolahan bahasa alami, linguistik komputasional, dan penambangan teks. Tujuannya adalah memahami opini, evaluasi, emosi, atau sikap seseorang terhadap suatu topik, produk, layanan, atau individu. Sistem analisis sentimen telah menunjukkan tingkat presisi yang tinggi, antara 75–95% tergantung jenis data [12,13].

Tabel 1. Aspek Depresi menurut Diagnostic and Statistical Manual of Mental Disorders, Fifth Edition (DSM V)

No	Aspek	Indikator	Item
1	Kognitif	Menunjukkan adanya kesalahan berpikir terhadap diri sendiri, pengalaman, serta masa depan.	1. Perasaan gagal dalam hidup
			2. Perasaan bersalah
			3. Perasaan dihukum
			4. Kekecewaan terhadap diri
			5. Kritikan terhadap diri sendiri
			6. Ketidakmampuan dalam mengambil keputusan
			7. Persepsi mengenai penampilan
2	Afektif	Menunjukkan adanya emosi negative terhadap diri sendiri	1. Perasaan sedih
			2. Pesimisme terhadap masa depan
			3. Ketidakpuasan terhadap berbagai hal
			4. Keinginan untuk bunuh diri
			5. Keinginan untuk menangis
			6. Perasaan marah
			7. Kehilangan minat terhadap orang lain

3	Somatis	Munculnya gejala – gejala	1. Kesulitan berkonsentrasi
		fisik yang cenderung menunjukkan penurunan	2. Perubahan pola tidur 3. Kelelahan

Data akan diklasifikasikan menggunakan beberapa algoritma *machine learning*. Metode klasifikasi bertujuan untuk memetakan data *input* ke dalam kelas-kelas yang telah ditentukan berdasarkan pola yang ditemukan dalam data pelatihan. Lima algoritma klasifikasi yang umum digunakan dan efektif dalam berbagai studi teks, yaitu:

a. *Naïve Bayes*

Naïve Bayes adalah algoritma klasifikasi berbasis probabilistik yang mengandalkan teorema Bayes dengan asumsi independensi antar fitur. Metode ini menghitung probabilitas suatu data termasuk ke dalam kelas tertentu berdasarkan distribusi kata dalam data latih. Keunggulannya terletak pada kesederhanaan dan efisiensinya dalam memproses data teks berukuran besar.

b. *k-Nearest Neighbor* (k-NN)

k-NN adalah metode klasifikasi non-parametrik yang mengklasifikasikan data baru berdasarkan mayoritas kelas dari sejumlah tetangga terdekat (*k*) dalam data latih. Jarak antar data biasanya diukur menggunakan *Euclidean distance*. Kelebihan utamanya adalah kemudahan implementasi, meskipun performanya sensitif terhadap *data noise* dan fitur yang tidak relevan.

c. *Decision Tree*

Decision Tree membentuk struktur pohon berdasarkan pemilihan atribut yang paling informatif menggunakan metrik seperti *entropy* dan *information gain*. Proses pembentukan pohon dilakukan secara rekursif hingga tercapai klasifikasi yang optimal. Algoritma ini memberikan interpretasi yang intuitif terhadap proses pengambilan keputusan.

d. *Random Forest*

Random Forest adalah metode *ensemble* yang menggabungkan banyak pohon keputusan (*decision tree*) menggunakan teknik *bagging*. Setiap pohon dibentuk dari sampel acak dan subset fitur. Prediksi akhir ditentukan oleh *voting* mayoritas dari seluruh pohon. Metode ini terkenal karena stabilitas dan akurasi yang tinggi, terutama pada dataset besar dan kompleks.

e. XGBoost

XGBoost (*Extreme Gradient Boosting*) merupakan metode *boosting* yang membangun pohon keputusan secara bertahap untuk memperbaiki kesalahan dari model sebelumnya. Dengan teknik regularisasi dan pengoptimalan gradien, XGBoost

menawarkan performa yang sangat baik dan efisiensi tinggi dalam tugas klasifikasi dan regresi.

Seleksi fitur (*feature selection*) adalah proses memilih fitur-fitur yang paling relevan dari sekumpulan data untuk digunakan dalam proses pemodelan. Tujuan utamanya adalah meningkatkan kinerja model dengan mengurangi dimensi data, mempercepat waktu komputasi, serta menghindari *overfitting* yang disebabkan oleh fitur yang tidak penting.

Dalam konteks klasifikasi teks, seleksi fitur sangat dibutuhkan karena data teks biasanya memiliki jumlah fitur yang sangat banyak. Tanpa proses ini, model cenderung bekerja kurang optimal dan menghasilkan akurasi yang rendah. Oleh karena itu, seleksi fitur menjadi salah satu langkah penting dalam membangun model klasifikasi yang efektif.

Beberapa metode seleksi fitur memanfaatkan algoritma optimasi, seperti *Particle Swarm Optimization* (PSO), *Bat Algorithm* (BA), dan *Flamingo Search Algorithm* (FSA). Metode-metode ini digunakan karena kemampuannya dalam mencari kombinasi fitur terbaik secara efisien, sehingga dapat meningkatkan performa klasifikasi secara keseluruhan.

Particle Swarm Optimization (PSO) adalah algoritma yang terinspirasi dari perilaku sosial sekawanan burung atau ikan. PSO, yang diperkenalkan oleh Eberhart dan Kennedy pada tahun 1995, menggunakan sekelompok partikel yang masing-masing memiliki posisi dan kecepatan dalam ruang pencarian. Setiap partikel bergerak dengan mengingat posisi terbaik yang pernah dicapai, serta mengikuti informasi dari partikel terbaik lainnya, untuk mendekati solusi optimal [14].

Bat Algorithm (BA) merupakan algoritma *swarm intelligence* yang dikembangkan oleh [8], terinspirasi dari kemampuan ekolokasi kelelawar. Ekolokasi memungkinkan kelelawar mengeluarkan gelombang ultrasonik dan menerima gema pantulannya untuk mengenali lingkungan sekitar secara akurat. Dalam BA, kelelawar buatan merepresentasikan solusi potensial yang memperbarui posisinya berdasarkan frekuensi, tingkat kenyaringan, dan laju emisi suara. Mekanisme ini memungkinkan BA untuk menjelajahi ruang solusi secara efisien, menjaga keseimbangan antara eksplorasi dan eksploitasi, serta menghindari solusi lokal. BA telah digunakan secara luas dalam optimasi fungsi, penjadwalan, dan berbagai aplikasi industri [15].

Flamingo Search Algorithm (FSA) adalah algoritma *swarm intelligence* yang diperkenalkan oleh [9], terinspirasi dari perilaku makan dan migrasi burung flamingo. Flamingo mencari makan dengan membungkukkan leher dan menyapu dasar air menggunakan paruhnya, serta berkomunikasi menggunakan suara untuk berbagi informasi lokasi dan ketersediaan makanan. FSA mengadopsi dua perilaku utama flamingo: mencari makan dan migrasi. Perilaku mencari makan mencerminkan eksploitasi lokal, sedangkan migrasi digunakan untuk eksplorasi global saat sumber

makanan menipis. Kedua perilaku ini dimodelkan secara matematis dalam FSA untuk mencari solusi optimal dalam ruang pencarian.

Kinerja lima model machine learning dengan seleksi fitur berbasis *nature-inspired optimization* dievaluasi menggunakan *Confusion Matrix* dengan parameter utama:

a. *Accuracy*

Mengukur seberapa tepat model mengklasifikasikan data secara keseluruhan.

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

b. *Precision*

Menunjukkan proporsi prediksi positif yang benar-benar relevan (unsur depresi).

$$Precision = \frac{TN}{FP + FN} \times 100\% \quad (2)$$

c. *Recall*

Mengukur seberapa baik model menemukan data positif yang sebenarnya.

$$Recall = \frac{TN}{FN + TN} \times 100\% \quad (3)$$

d. *F1-Score*

Rata-rata dari *precision* dan *recall*, berguna saat distribusi kelas tidak seimbang.

$$F_1Score = \frac{2 \times Precision \times Recall}{Precision + Recall} \quad (4)$$

Data dikumpulkan dari *platform X* menggunakan alat “tweet-harvest” berdasarkan *term* “depresi”, “bunuh diri”, dan “sedih” pada Agustus 2023. Data yang diambil berbahasa Indonesia mencakup *reply*, *reply*, dan *post* asli. Total 1105 data disimpan dalam format .xlsx setelah duplikat dan data tanpa sentimen dihapus. Data masih mentah dan mengandung elemen seperti *hashtag*, *mention*, *link*, dan *username*, sebagaimana ditunjukkan pada Tabel 2.

Tabel 2. Contoh data *post*

<i>Post</i>
@lilaccountz Aku contohnya, aku kirain nikah bakal ngilangin traumaku, gak ada yang bilang kalau nikah itu setidak nyangka gini, mama juga udah meninggal, bapak sibuk nyari pasangan baru. Akhirnya ketika nikah duar.. Aku depresi sedang menuju ambang batas (?) Gitu kata dokternya -c-
Kamu bisa melakukan relaksasi dengan yoga, pijat, atau meditasi. Relaksasi juga bisa membantumu untuk meredakan depresi, meningkatkan konsentrasi, mengurangi stres,

dan membuat otot jadi lebih rileks

@tanyakanrl lu mending kuliah, gw ga kuliah kuliah ggra harus kerja jdi tulang punggung keluarga dri umur 14, harus rela bagi makanan padhal laper banget, harus tpp baik baik aja pdhal mau teriak, pdhal bru 18 thun, harus ngerasain depresi tp keluarga taunya gue ketempelan

Data teks dilabeli secara manual oleh dua orang, dengan profil singkatnya ditampilkan pada Tabel 3. Proses pelabelan dilakukan secara objektif mengacu pada aspek depresi dalam DSM V (2013). Setiap post dikategorikan ke dalam dua label, yaitu “Depresi” dan “Bukan Depresi”. Contoh hasil pelabelan dapat dilihat pada Tabel 4.

Tahapan *preprocessing* dilakukan pada 1105 data *post* berlabel “Depresi” dan “Bukan Depresi”. Langkah-langkahnya meliputi: pembersihan teks (menghapus *hashtag*, *mention*, *link*, *username*, tanda baca, angka, dan spasi berlebih), tokenisasi, normalisasi dengan kamus slang “kamus-alay” dari GitHub, penghapusan *stopwords*, dan *stemming* menggunakan pustaka Sastrawi. Setelah itu, token digabung kembali menjadi string untuk memudahkan analisis.

Tabel 3. Profil pemberi label pada *post*

No	Nama	Jenis Kelamin	Pekerjaan
1	Dwi Ratno Priyambodo, S.Psi., M.Psi.	Laki-laki	Dosen Psikologi Wisnuwardhana Malang
2	Zein Rizky Santoso, S.Mat	Laki-laki	Mahasiswa

Tabel 4. Contoh pelabelan data secara manual

<i>Post</i>	Kategori
@lilaccountz Aku contohnya, aku kirain nikah bakal ngilangin traumaku, gak ada yang bilang kalau nikah itu setidak nyangka gini, mama juga udah meninggal, bapak sibuk nyari pasangan baru. Akhirnya ketika nikah duar.. Aku depresi sedang menuju ambang batas (?) Gitu kata dokternya -c-	Depresi
Kamu bisa melakukan relaksasi dengan yoga, pijat, atau meditasi. Relaksasi juga bisa membantumu untuk meredakan depresi, meningkatkan konsentrasi, mengurangi stres, dan membuat otot jadi lebih rileks	Bukan Depresi
@tanyakanrl lu mending kuliah, gw ga kuliah kuliah ggra harus kerja jdi tulang punggung keluarga dri umur 14, harus rela bagi makanan padhal laper banget, harus tpp baik baik aja pdhal mau teriak, pdhal bru 18 thun, harus ngerasain depresi tp keluarga taunya gue ketempelan	Depresi

Selanjutnya, dilakukan pembobotan menggunakan TF-IDF untuk memberi nilai penting pada kata-kata yang relevan dalam mendeteksi depresi. Hasilnya berupa matriks TF-IDF berukuran $m \times n$ seperti pada Tabel 5. Data kemudian dibagi menjadi data latih dan uji (70:30). Jika terdapat ketidakseimbangan kelas, digunakan teknik SMOTE untuk

menyeimbangkan jumlah data, sehingga model dapat belajar secara lebih adil dan akurat.

Tabel 5. Ilustrasi pembobotan kata

	Kata₁	Kata₂	Kata₃	...	Kata_n
Dok₁	TFIDF ₁₁	TFIDF ₁₂	TFIDF ₁₃	...	TFIDF _{1n}
Dok₂	TFIDF ₂₁	TFIDF ₂₂	TFIDF ₂₃	...	TFIDF _{2n}
Dok₃	TFIDF ₃₁	TFIDF ₃₂	TFIDF ₃₃	...	TFIDF _{3n}
...	...	...	...	...	...
Dok_m	TFIDF _{m1}	TFIDF _{m2}	TFIDF _{m3}	...	TFIDF _{mn}

Penelitian ini menggunakan tiga metode seleksi fitur berbasis algoritma alam, yaitu *Particle Swarm Optimization* (PSO), *Bat Algorithm* (BA), dan *Flamingo Search Algorithm* (FSA). Proses seleksi fitur dilakukan dengan input matriks TF-IDF untuk memilih kata-kata paling relevan dalam mendeteksi indikasi depresi. TF-IDF membantu model mengenali pentingnya suatu kata dalam keseluruhan data.

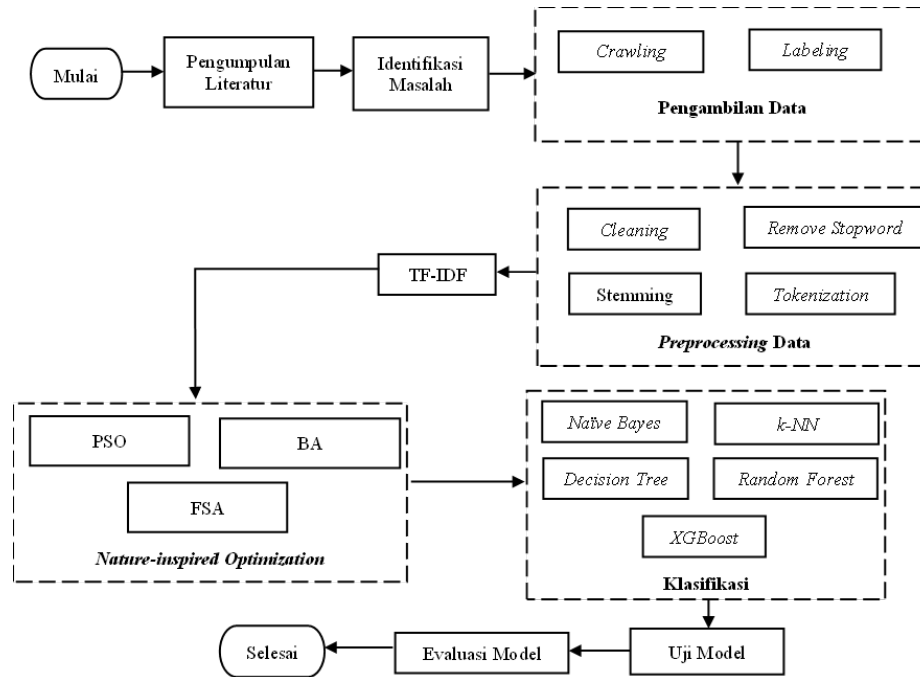
Eksperimen dilakukan dengan variasi jumlah partikel (20, 30, 40, 50) dan iterasi sebanyak 150, dengan penghentian saat hasil konvergen. Tujuannya adalah mengevaluasi performa model dalam berbagai kondisi dan menemukan konfigurasi terbaik.

Setelah seleksi fitur, data diklasifikasikan menggunakan lima metode machine learning: *Naïve Bayes*, k-NN, *Random Forest*, *Decision Tree*, dan *XGBoost*. Data dibagi menjadi data latih dan uji, dengan penerapan *Grid Search* untuk *hyperparameter tuning* guna mengoptimalkan performa. Model dibangun untuk klasifikasi dua kelas: “depresi” dan “bukan depresi”. Hasil tiap metode dibandingkan berdasarkan *confusion matrix*. Rincian parameter ada pada Tabel 6, dan alur penelitian ditunjukkan pada Gambar 1.

Tabel 6. *Hyperparameter Tuning* Kelima Metode Klasifikasi

No	Model	Hyperparameter	Nilai Kandidat
1	<i>Naïve Bayes</i>	alpha (α)	[0.01, 0.1, 1, 10, 100]
		fit_prior	[True, False]
2	<i>K-Nearest Neighbours</i>	n_neighbors	[3,5,7,9,11]
		weights	[uniform, distance]
3	<i>Decision Tree</i>	max_depths	[None, 10,20,30,40,50]
		min_samples_split	[2,5,10]
		min_samples_leaf	[1,2,4]
4	<i>Random Forest</i>	n_estimators	[10,50,100,200]
		max_depth	[None, 10,20,30,40,50]
		min_samples_split	[2,5,10]
		min_samples_leaf	[1,2,4]
		bootstrap	[True, False]

5	XGBoost	n_estimators	[100,200,300]
		max_depth	[3,4,5,6]
		learning_rate	[0.01,0.1,0.3]
		subsample	[0.8,0.9,1]
		colsample_bytree	0.8,0.9,1]
		min_child_weight	[1,3,5]



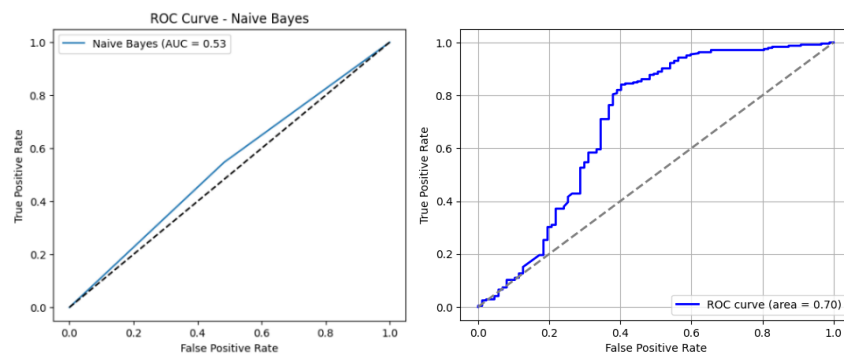
Gambar 1. Diagram Alir Metode Penelitian

3. Hasil dan Diskusi

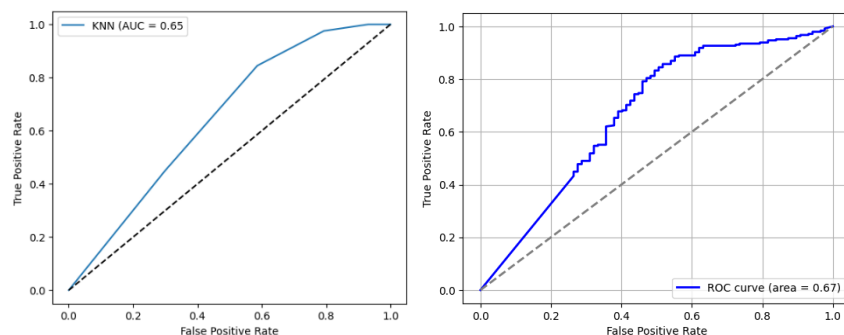
Pada tahap eksplorasi data, dilakukan analisis sentimen terhadap 1105 *post* asli, terdiri dari 819 *post* dengan sentimen depresi (74,1%) dan 286 *post* dengan sentimen bukan depresi (25,9%), sebagaimana ditunjukkan pada Gambar 2. Hasil ini menunjukkan adanya ketidakseimbangan kelas, di mana jumlah *post* kelas depresi jauh lebih banyak. Ketidakseimbangan ini berisiko membuat model lebih sering memprediksi kelas mayoritas. Untuk mengatasinya, diterapkan *Synthetic Minority Over-sampling Technique* (SMOTE) agar jumlah data di kedua kelas menjadi seimbang, sehingga model dapat mengenali kedua sentimen dengan lebih baik.

mencapai 0.9, tertinggi di antara metode lain. Menariknya, performa optimal tidak selalu butuh populasi besar, seperti pada NB-BA yang terbaik dengan populasi 20. Secara keseluruhan, NB-FSA terbukti efektif untuk analisis sentimen, mampu mengurangi noise dan tetap memberikan prediksi akurat, dengan peningkatan AUC-ROC dari 0.53 ke 0.7.

Eksperimen k-NN menunjukkan bahwa seleksi fitur berdampak bervariasi pada performa model. KNN-FSA dengan populasi 50 mencatat akurasi tertinggi (0.81), mengungguli KNN-PSO (0.77) dan KNN-BA (0.753), seperti pada Tabel 7. Meski sedikit di bawah NB-FSA (0.82), KNN-FSA menunjukkan F1-Score yang lebih seimbang: 0.88 untuk “Yes” dan 0.6 untuk “No” (vs 0.9 dan 0.37 pada NB-FSA). KNN-FSA juga sangat efisien, hanya membutuhkan 1 iterasi dibandingkan 98 pada NB-FSA. Meski menggunakan lebih banyak fitur (967 vs 511), hal ini mencerminkan karakter algoritma: NB lebih sensitif terhadap seleksi fitur, sedangkan k-NN mengandalkan struktur data. KNN-BA terlihat kurang stabil dibanding NB-BA.



Gambar 4. (a) Grafik ROC klasifikasi *Naïve Bayes* (b) Grafik ROC klasifikasi *Naïve Bayes* dengan *Flamingo Search Algorithm*

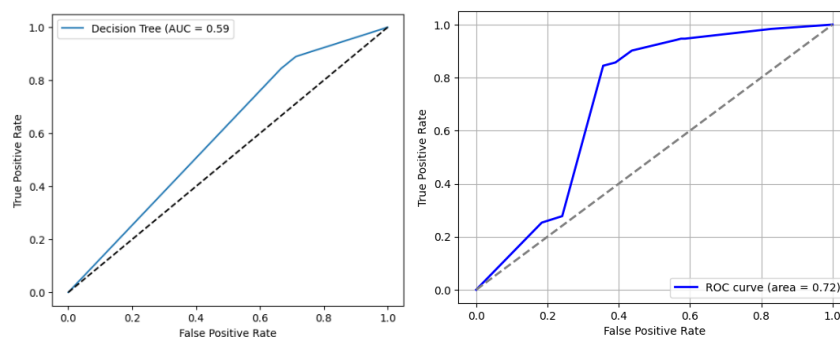


Gambar 5. (a) Grafik ROC klasifikasi K-NN (b) Grafik ROC klasifikasi K-NN dengan *Flamingo Search Algorithm*

Seleksi fitur meningkatkan akurasi k-NN, namun analisis AUC-ROC (Gambar 5) menunjukkan keterbatasan dalam membedakan kelas, berbeda dengan peningkatan signifikan pada *Naïve Bayes* (dari 0.53 ke 0.7). Hal ini menunjukkan bahwa meski terbantu seleksi fitur, k-NN masih kesulitan saat batas antar kelas tidak jelas. Secara

keseluruhan, KNN-FSA cocok untuk implementasi praktis karena efisien, cukup stabil, dan seimbang. Namun, jika prioritas adalah akurasi tertinggi dan reduksi fitur maksimal, NB-FSA tetap menjadi pilihan utama.

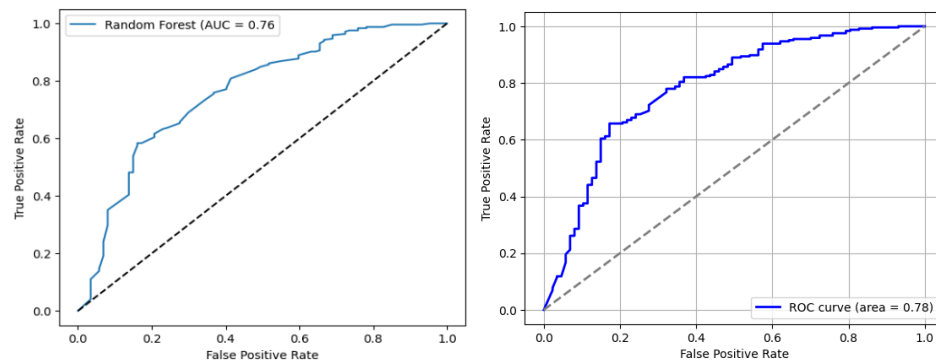
Hasil pada Tabel 7 menunjukkan bahwa DT-FSA dengan populasi 50 adalah kombinasi optimal, mencatat akurasi tertinggi (0.813) dibanding DT-PSO (0.78) dan DT-BA (0.795). Keunggulannya terletak pada efisiensi, hanya membutuhkan 1 iterasi dibanding 9–14 pada metode lain. Meski menggunakan lebih banyak fitur (1769 vs 1153 pada DT-PSO), DT-FSA tetap mempertahankan performa tinggi dengan waktu komputasi optimal. AUC-ROC meningkat dari 0.59 ke 0.72 (Gambar 6), menunjukkan peningkatan stabilitas model. F1-Score juga seimbang: 0.88 untuk “Yes” dan 0.6 untuk “No”. Meski sedikit di bawah NB-FSA dalam akurasi (0.813 vs 0.82), DT-FSA lebih unggul dalam keseimbangan prediksi. Secara keseluruhan, DT-FSA merupakan pilihan ideal untuk implementasi Decision Tree karena menggabungkan akurasi, efisiensi, dan keseimbangan prediksi. Konsistensi performa FSA di berbagai algoritma semakin menegaskan keandalannya sebagai metode seleksi fitur.



Gambar 6. (a) Grafik ROC klasifikasi *Decision Tree* (b) Grafik ROC klasifikasi *Decision Tree* dengan *Flamingo Search Algorithm*

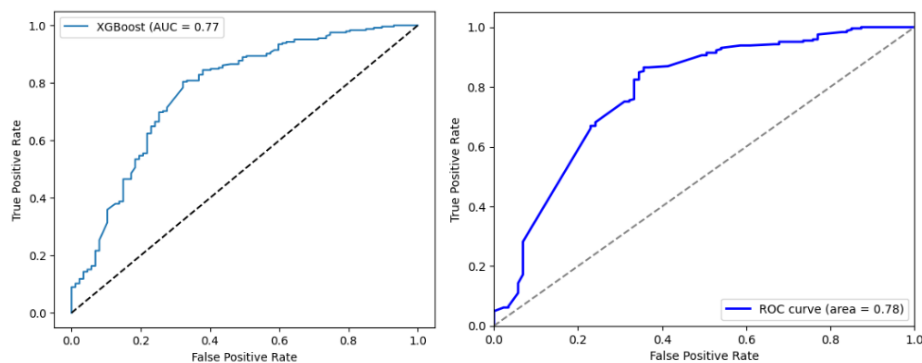
Eksperimen Random Forest menunjukkan bahwa seleksi fitur berdampak besar pada keseimbangan performa model. Sebelum seleksi, akurasi cukup tinggi (0.77), namun F1-Score sangat timpang (0.87 untuk “Yes”, 0.29 untuk “No”), menunjukkan bias terhadap satu kelas. AUC sebesar 0.76 pun belum mencerminkan performa sebenarnya.

Setelah seleksi fitur dengan FSA, akurasi meningkat menjadi 0.822, F1-Score menjadi lebih seimbang (0.9 dan 0.6), dan AUC naik ke 0.78 (Gambar 7). FSA juga lebih efisien, hanya menggunakan 1280 fitur dan 20 iterasi, lebih hemat dibanding DT-FSA (1769 fitur). Dibanding PSO dan BA, FSA tetap unggul dalam hasil dan efisiensi. Secara keseluruhan, kombinasi RF-FSA dengan populasi 50 menjadi yang terbaik, memberikan performa yang seimbang, akurat, dan efisien—sangat ideal untuk deteksi dini gangguan mental.



Gambar 7. (a) Grafik ROC klasifikasi *Random Forest* (b) Grafik ROC klasifikasi *Random Forest* dengan *Flamingo Search Algorithm*

Berdasarkan Tabel 7, XGB-FSA dengan populasi 50 mencatat akurasi tertinggi (0.8102), mengungguli XGB-PSO (0.786) dan XGB-BA (0.78). F1-Score kelas “Yes” sebesar 0.88 menunjukkan prediksi yang seimbang. Meski menggunakan lebih banyak fitur (1712) dibanding RF-FSA (1280), XGB-FSA jauh lebih efisien dengan hanya 1 iterasi untuk konvergensi. Dari evaluasi model (Gambar 8), FSA berhasil meningkatkan recall dan menjaga presisi, sehingga AUC naik dari 0.76 ke 0.78. Meski peningkatannya kecil, perubahan ini penting karena mengurangi false negative dan menjadikan model lebih adil dalam mengenali kedua kelas-krusial dalam deteksi dini gangguan mental. Secara keseluruhan, XGB-FSA menawarkan performa stabil dan cocok untuk data kompleks atau saat kecepatan penting. Namun, RF-FSA tetap unggul secara keseluruhan berkat akurasi lebih tinggi (0.822), fitur lebih sedikit, dan keseimbangan prediksi yang lebih baik, menjadikannya pilihan utama dalam eksperimen ini.



Gambar 8. (a) Grafik ROC klasifikasi XGBoost (b) Grafik ROC klasifikasi XGBoost dengan *Flamingo Search Algorithm*

Tabel 7. Hasil eksperimen klasifikasi XGBoost

Metode	Jumlah Populasi Terbaik	Akurasi	F1-Score		Fitur	Iterasi
			Yes	No		
NB	-	0.54	0.64	0.37	-	-
NB-PSO	30	0.76	0.8	0.5	1779	150
NB-BA	20	0.75	0.84	0.5	1715	101
NB-FSA	40	0.82	0.9	0.6	511	98
KNN	-	0.77	0.86	0.32	-	-
KNN-PSO	50	0.77	0.86	0.45	1148	6
KNN-BA	50	0.753	0.84	0.4	1802	25
KNN-FSA	50	0.81	0.88	0.6	967	1
DT	-	0.73	0.83	0.36	-	-
DT-PSO	20	0.78	0.86	0.48	1153	14
DT-BA	50	0.795	0.87	0.54	1698	9
DT-FSA	50	0.813	0.88	0.6	1769	1
RF	-	0.77	0.87	0.29	-	-
RF-PSO	50	0.77	0.8	0.3	1795	52
RF-BA	50	0.8	0.8	0.4	1678	26
RF-FSA	50	0.822	0.9	0.6	1280	20
XGB	-	0.79	0.87	0.49	-	-
XGB-PSO	50	0.786	0.86	0.56	1740	1
XGB-BA	50	0.78	0.85	0.56	1801	1
XGB-FSA	50	0.8102	0.88	0.6	1712	1

4. Kesimpulan

Seleksi fitur secara signifikan meningkatkan akurasi serta keseimbangan presisi dan recall dibanding model tanpa seleksi. Di antara metode yang diuji, Flamingo Search Algorithm (FSA) terbukti paling unggul dalam menghasilkan model yang akurat, stabil, dan efisien. FSA meningkatkan AUC-ROC di semua algoritma, dengan peningkatan terbesar pada Naïve Bayes. Kombinasi terbaik ditemukan pada RF-FSA karena menawarkan akurasi tinggi dengan jumlah fitur yang efisien. Sementara itu, XGB-FSA menjadi alternatif andal untuk data kompleks, meski memerlukan daya komputasi lebih besar. Secara keseluruhan, pemilihan metode seleksi fitur sangat krusial dalam klasifikasi teks. RF-FSA direkomendasikan sebagai model optimal untuk deteksi depresi melalui teks karena hasilnya yang akurat, efisien, dan stabil. Penelitian selanjutnya disarankan mengeksplorasi kombinasi deep learning dan seleksi fitur, serta menguji pada dataset yang lebih besar dan beragam. Metode seleksi atau optimasi

terbaru dan model klasifikasi yang lebih canggih juga layak dipertimbangkan untuk meningkatkan akurasi dan efisiensi.

Daftar Pustaka

- [1] Cahyono, A. S. Pengaruh Media Sosial Terhadap Perubahan Sosial Masyarakat Di Indonesia. *Publiciana*, 9, 140–157, 2016.
- [2] Zsila, Á., Reyes, M. Pros & cons: impacts of social media on mental health. *BMC Psychol*, 11(1), 201, 2023. doi:10.1186/s40359-023-01243-x.
- [3] Ridlo, I. A. Pandemi COVID-19 dan Tantangan Kebijakan Kesehatan Mental di Indonesia. *INSAN Jurnal Psikologi dan Kesehatan Mental*, 5(2), 162–171, 2020. doi : 10.20473/jpkm.V5I22020.162-171.
- [4] Muslimahayati, M., & Rahmy, H. A. Depresi dan Kecemasan Remaja Ditinjau dari Perspektif Kesehatan dan Islam. *DEMOS: Journal of Demography, Ethnography and Social Transformation*, 2021.
- [5] Brody, H. A journey into the causes and effects of depression. *Nature*, 608(7924), S35–S35, 2022.
- [6] Nurrochman, M. M., Prasasti A. I., & Nugrahaeni, R. A. Implementasi Machine Learning Untuk Mendeteksi Unsur Depresi Pada Tweet Menggunakan Metode Naïve Bayes. *eProceedings of Engineering*, 8, 6250–6257, 2021.
- [7] Dainamang, S. A., Hayatin, N., & Chandranegara, D. R. Analisis Sentimen Media Sosial Twitter terhadap RUU Omnibus Law dengan Metode Naive Bayes dan Particle Swarm Optimization. *Komputika : Jurnal Sistem Komputer*, 11(2), 211–218, 2022. doi:10.34010/komputika.v11i2.6037.
- [8] Yang, X. S. A New Metaheuristic Bat-Inspired Algorithm. Di dalam: González JR, Pelta DA, Cruz C, Terrazas G, Krasnogor N, editor. *Nature Inspired Cooperative Strategies for Optimization (NICSO 2010)*. Berlin, Heidelberg: Springer Berlin Heidelberg, 65–74, 2010.
- [9] Zhiheng, W., & Jianhua, L. Flamingo Search Algorithm: A New Swarm Intelligence Optimization Algorithm. *IEEE Access*, 9, 88564–88582, 2021. doi:10.1109/ACCESS.2021.3090512.
- [10] Dirgayunita, A. Depresi: Ciri, Penyebab dan Penangannya. *Journal An-Nafs: Kajian Penelitian Psikologi*, 1(1), 1–14, 2016. doi:https://doi.org/10.33367/psi.v1i1.235.

- [11] Hadi, I., Wijayanti, F., Devianti, R., & Rosyanti, L. GANGGUAN DEPRESI MAYOR (MAYOR DEPRESSIVE DISORDER) MINI REVIEW. *Health Information : Jurnal Penelitian*, 9, 25–40, 2017. doi:10.36990/hijp.v9i1.102.
- [12] Nasukawa, T., & Yi, J. Sentiment Analysis: Capturing Favorability Using Natural Language Processing. *Proceedings of the 2nd International Conference on Knowledge Capture*, 2003.
- [13] Liu, X., dkk. SocialCube: A text cube framework for analyzing social media data. Di dalam: *Proceedings of the 2012 ASE International Conference on Social Informatics, SocialInformatics 2012. IEEE Computer Society*, 252–259, 2012.
- [14] Freitas, D., Lopes, L. G., & Morgado-Dias F. Particle Swarm Optimisation: A Historical Review Up to the Current Developments. *Entropy*, 22(3), 2020. doi:10.3390/e22030362
- [15] Yang, X. Nature-Inspired Optimization Algorithms. Di dalam: Yang X-S, editor. *Nature-Inspired Optimization Algorithms. Oxford: Elsevier*, 2014.