

CESS

(Journal of Computer Engineering, System and Science)

Available online: <https://jurnal.unimed.ac.id/2012/index.php/cess>

ISSN: 2502-714x (Print) | ISSN: 2502-7131 (Online)



Random Forest dengan Hyperparameter Tuning untuk Sistem Prediksi Risiko Depresi Remaja

Random Forest with Hyperparameter Tuning for a Teen Depression Risk Prediction System

Suhendri^{1*}, Aep Saepuloh², Hegar Zalekania³, Ilma Ala Ulumillah⁴,
Raza Haan Fiddo Aryasturangga⁵, Regita Nuralvianti Pratiwi⁶

^{1,2,3,4,5,6}Informatika, Teknik, Universitas Majalengka, Indonesia

Jl. Raya K H Abdul Halim No.103, Majalengka Kulon, Kec. Majalengka, Kabupaten Majalengka, Jawa Barat 45418

Email: ¹theprof.suhendri@yahoo.co.id, ²aepsaepullah99@gmail.com,

³hegarzalekania111@gmail.com, ⁴ilmamillah21@gmail.com, ⁵razahaanfiddo@gmail.com,

⁶nrlvianti18@gmail.com

*Corresponding Author

ABSTRAK

Deteksi dini kerentanan depresi pada fase remaja memerlukan pendekatan komputasi yang presisi. Penelitian ini bertujuan merancang arsitektur sistem cerdas prediksi risiko depresi berbasis web dengan menganalisis parameter *lifestyle indicators*. Metodologi yang diusulkan mengimplementasikan algoritma klasifikasi *Random Forest* yang dipadukan dengan teknik penyeimbangan data *Synthetic Minority Over-sampling Technique (SMOTE)*. Untuk memaksimalkan performa prediktif, studi ini menghadirkan intervensi kebaruan berupa *hyperparameter tuning* menggunakan pendekatan optimasi Bayesian. Hasil pengujian komputasi membuktikan bahwa metode Bayesian sukses mendongkrak tingkat akurasi model secara solid, dari 77,20% pada konfigurasi *default* menjadi 79,20%. Nilai orisinalitas sekaligus dampak paling krusial dari pemodelan ini adalah kemampuannya menekan kesalahan deteksi fatal secara signifikan, yang tercermin dari capaian metrik sensitivitas (*Recall*) sebesar 0.9643 khusus pada kelas risiko tinggi. Integrasi penuh dari algoritma pasca optimasi ini ke dalam antarmuka perangkat lunak interaktif mengonfirmasi bahwa penelitian ini tidak sekadar menyajikan luaran analitis di atas kertas, melainkan berhasil memproduksi instrumen skrining preventif yang fungsional.

Kata Kunci: *prediksi risiko depresi; random forest; optimasi bayesian, machine learning, SMOTE.*



ABSTRACT

Early detection of vulnerability to depression during adolescence requires a precise computational approach. This study aims to design a web-based intelligent system architecture for predicting depression risk by analyzing lifestyle indicators and genetic parameters. The proposed methodology implements the Random Forest classification algorithm combined with the Synthetic Minority Over-sampling Technique (SMOTE) for data balancing. To maximize predictive performance, this study introduces an innovative intervention in the form of hyperparameter tuning using a Bayesian optimization approach. Computational testing results demonstrate that the Bayesian method successfully and significantly boosts the model's accuracy, from 77.20% in the default configuration to 79.20%. The most crucial aspect of the originality and impact of this modeling is its ability to significantly reduce false-positive errors, as reflected by a Recall metric of 0.9643 specifically for the high-risk class. The full integration of this post-optimization algorithm into an interactive software interface confirms that this research does not merely present analytical outputs on paper but has successfully produced a functional, reliable, and ready-to-implement preventive screening tool for direct use in the community.

Keywords: *depression risk prediction; random forest; bayesian optimization; machine learning; SMOTE.*

1. PENDAHULUAN

Depresi merupakan salah satu tantangan kesehatan mental paling serius yang mengancam generasi muda di seluruh dunia. Organisasi Kesehatan Dunia (WHO) mendefinisikan depresi sebagai gangguan mental yang ditandai dengan perubahan suasana hati yang tidak terduga, kesedihan yang mendalam, serta penurunan kualitas fungsi kognitif yang memengaruhi interaksi sosial. Berdasarkan rilis data global terbaru dari WHO, prevalensi depresi secara global diperkirakan terjadi pada 3,4% populasi remaja berusia 15–19 tahun. Kondisi ini tidak bisa diabaikan karena merupakan pemicu rentetan masalah perilaku, mulai dari penurunan akademik hingga risiko menyakiti diri sendiri [1].

Di tingkat nasional, krisis kesehatan mental remaja juga menunjukkan angka yang sangat mengkhawatirkan. Laporan resmi dari Indonesia National Adolescent Mental Health Survey (I-NAMHS) tahun 2022 mengungkap bahwa 34,9% (sekitar 1 dari 3) remaja di Indonesia memiliki setidaknya satu masalah kesehatan mental [2]. Angka ini diperparah dengan temuan bahwa 5,5% remaja di Indonesia terdiagnosis mengalami gangguan mental dalam 12 bulan terakhir [3]. Fakta ini selaras dengan laporan Kementerian Kesehatan Republik Indonesia pada tahun 2023 yang menyoroti bahwa prevalensi depresi di Indonesia mencapai puncaknya pada kelompok anak muda (15–24 tahun) di angka 2% [4]. Ironisnya, hanya sebagian kecil dari mereka yang mendapatkan akses pengobatan yang layak, sehingga mendesak adanya sistem pendeteksian dini yang terukur agar intervensi dapat dilakukan sebelum gejala memburuk.

Langkah mitigasi yang efektif sangat bergantung pada identifikasi variabel-variabel risiko secara akurat. Klasifikasi risiko depresi pada remaja kini tidak lagi bertumpu semata-mata pada asesmen psikologis pasca-trauma [5]. Kondisi ini dapat diprediksi sejak awal melalui pengawasan pada indikator gaya hidup, seperti durasi tidur, tingkat stres akademik, dan pola

makan [6]. Pengolahan data yang melibatkan kombinasi interaksi antara indikator gaya hidup ini sangat kompleks serta *nonlinear*. Oleh karena itu, analisis statistik tradisional seringkali tidak mampu mengenali pola laten yang ada, menuntut adanya pemanfaatan komputasi kecerdasan buatan.

Sebuah tinjauan empiris menyimpulkan bahwa pengimplementasian *machine learning* berhasil mereduksi bias observasi dan mengekstraksi pola dari set data berdimensi tinggi secara sangat efisien, menjadikannya instrumen prediksi risiko depresi yang jauh lebih tajam dibandingkan metode skrining konvensional [7]. Kemampuan komputasi ini sangat diandalkan karena terbukti mampu memproses set data yang besar dan multidimensi secara efisien [8]. Dari berbagai algoritma yang tersedia, *Random Forest* (RF) sering dijadikan model dasar yang andal dalam klasifikasi kesehatan mental. Keunggulan *Random Forest* terletak pada kemampuannya menyusun ansambel pohon keputusan yang membuatnya resistan terhadap data pencilan (*outliers*) [9]. Dalam implementasinya pada data *multivariabel* kesehatan mental, algoritma *Random Forest* terbukti unggul dalam menjaga stabilitas prediksi dan meminimalisir kesalahan varians ketika memproses interaksi dari puluhan indikator psikososial secara bersamaan [10]. Berdasarkan evaluasi performa komparatif terhadap beberapa arsitektur kecerdasan buatan, model *Random Forest* secara meyakinkan mencatatkan metrik akurasi dan presisi tertinggi berkat kemampuannya meminimalkan varians saat memproses indikator-indikator kesehatan mental yang kompleks [11]. Selain pada kasus depresi, pendekatan berbasis ansambel pohon keputusan ini juga terbukti secara signifikan efektif dalam mengekstraksi dan memetakan pola tingkatan stres akademik pada subjek usia remaja [12]. Meski begitu, model ini sangat sensitif terhadap pengaturan parameter bawaannya tanpa kalibrasi yang tepat, arsitektur RF sangat rentan terjebak pada komputasi yang berat dan *overfitting*.

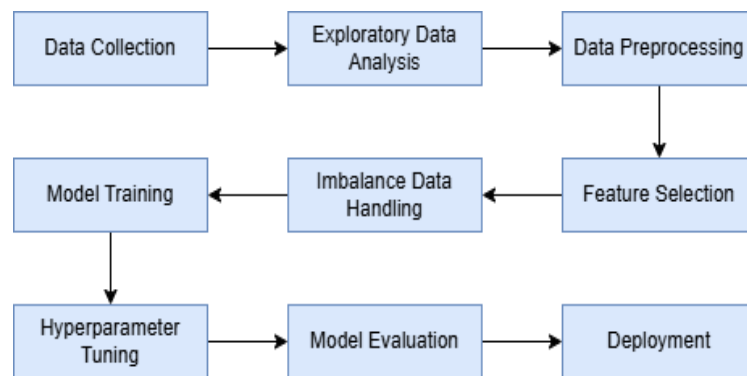
Dalam mengoptimalkan performa arsitektur tersebut, *hyperparameter tuning* memegang peran sentral. Penggunaan metode konvensional seperti *Grid Search* seringkali terhambat oleh inefisiensi komputasi, di mana iterasi pencarian parameter dilakukan secara ekstensif namun buta, tanpa memanfaatkan informasi historis dari hasil evaluasi parameter sebelumnya [13]. Oleh karena itu, pendekatan *Bayesian Optimization* menawarkan efisiensi iterasi yang jauh lebih superior. Memanfaatkan model probabilistik, metode berbasis Bayesian secara cerdas merekam hasil dari hyperparameter sebelumnya untuk menebak arah evaluasi berikutnya [14]. Berdasarkan analisis hasil pengujian pada set data klinis, penerapan *Bayesian Optimization* pada *Random Forest* terbukti mampu mengalahkan kinerja parameter bawaan (*default*) dengan meningkatkan skor akurasi secara signifikan sekaligus menekan biaya dan waktu komputasi secara drastis [15].

Berdasarkan kajian literatur, meskipun algoritma *machine learning* telah diaplikasikan dalam ranah psikologi klinis, masih terdapat celah penelitian dalam mengintegrasikan *Hyperparameter Bayesian* pada *Random Forest* spesifik untuk kasus kesehatan mental remaja. Pemilihan *Bayesian Optimization* didasarkan pada keunggulan surrogate model probabilistiknya yang jauh lebih efisien dan terarah dibandingkan *Grid* atau *Random Search* dalam menavigasi kompleksitas dimensi data indikator gaya hidup. Kebaruan dari penelitian ini berada pada optimasi *hyperparameter* sistem prediksi yang diukur berdasarkan korelasi indikator gaya hidup, serta implementasinya ke dalam lingkungan pengembangan sistem

aplikasi yang siap pakai. Transformasi model teroptimasi menjadi mesin inferensi ke dalam sistem aplikasi terapan ini membedakannya secara signifikan dari studi terdahulu yang umumnya terhenti pada sebatas evaluasi performa metrik. Solusi yang diusulkan adalah merancang sebuah sistem cerdas yang mampu mengalkulasi tingkat risiko depresi dengan bobot parameter yang sudah ditala dengan fungsi Bayesian. Oleh karena itu, penelitian ini bertujuan untuk membangun model sistem prediksi risiko depresi, mengevaluasi matriks peningkatan performanya dibandingkan algoritma standar, serta menghasilkan arsitektur sistem peringatan dini yang berkontribusi pada penanganan kesehatan mental preventif.

2. METODE PENELITIAN

Penelitian ini menerapkan kerangka kerja komprehensif yang mencakup tahapan akuisisi data terkait indikator gaya hidup, pra-pemrosesan data, hingga perancangan model prediktif berbasis algoritma Random Forest. Guna mencapai kinerja model yang optimal, tahapan klasifikasi ini diintegrasikan secara sistematis dengan teknik *hyperparameter tuning*. Secara lebih terperinci, keseluruhan alur metodologi yang diusulkan dalam sistem prediksi ini diilustrasikan pada Gambar 1.



Gambar 1. Alur Penelitian

2.1. Data Collection

Penelitian ini memanfaatkan kumpulan data sekunder yang diperoleh secara komprehensif dari repositori publik Kaggle [16]. Dataset tersebut mencakup total 2.500 rekam data yang direpresentasikan melalui 11 variabel pengamatan. Kesebelas variabel ini diformulasikan untuk mengukur berbagai aspek kondisi remaja, yang meliputi dimensi demografi, interaksi sosial, indikator gaya hidup, rekam akademik, hingga asesmen kesehatan mental. Rincian deskripsi dari masing-masing variabel tersebut diuraikan secara lengkap pada Tabel 1.

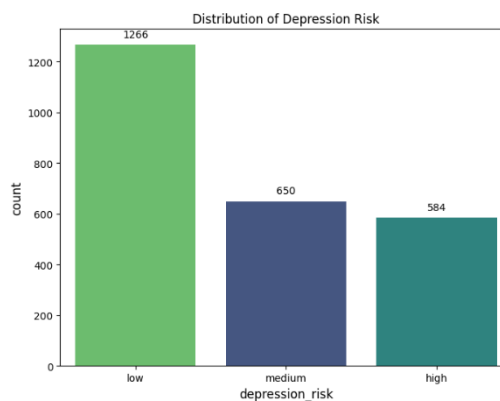
Tabel 1. Informasi Dataset

Fitur	Deskripsi
<i>age</i>	<i>Age of the teenager (13–19)</i>
<i>gender</i>	<i>Male or Female</i>
<i>daily_social_media_hours</i>	<i>Time spent on social media per day</i>
<i>platform_usage</i>	<i>Instagram, TikTok, Both, or Other</i>

<i>sleep_hours</i>	<i>Average daily sleep duration</i>
<i>screen_time_before_sleep</i>	<i>Phone usage before bedtime</i>
<i>academic_performance</i>	<i>GPA score (2.0–4.0 scale)</i>
<i>physical_activity</i>	<i>Daily exercise (hours)</i>
<i>social_interaction_level</i>	<i>Low, Medium, High</i>
<i>stress_level</i>	<i>Stress score (1–10)</i>
<i>anxiety_level</i>	<i>Anxiety score (1–10)</i>
<i>depression_risk</i>	<i>Low, Medium, High</i>

2.2. Exploratory Data Analysis

Exploratory Data Analysis merupakan tahapan peninjauan awal yang esensial untuk mengarakterisasi variabel-variabel utama penelitian dengan mengevaluasi distribusi frekuensi, persentase, serta tendensi pemusatan data sebelum tahap pemodelan komputasi dilakukan [17].



Gambar 2. Distrbusi Risiko Depresi

Berdasarkan visualisasi grafik batang pada Gambar 2, dapat diamati distribusi frekuensi secara spesifik dari variabel target tingkat risiko depresi. Dari total 2.500 observasi data yang secara keseluruhan diekstraksi dari profil demografi, sosial, akademik, hingga indikator gaya hidup, terdapat dominasi signifikan pada kategori risiko rendah dengan jumlah 1.266 rekam data. Sementara itu, kelompok risiko sedang dan risiko tinggi mencakup porsi yang lebih kecil, yakni berturut-turut sebanyak 650 dan 584 rekam data. Profil distribusi ini secara eksplisit menunjukkan adanya ketidakseimbangan kelas yang moderat dalam himpunan data.

2.3. Data Preprocessing

Secara konseptual, pra-pemrosesan data merupakan tahapan analitis yang krusial untuk mentransformasi serta membersihkan data mentah dari berbagai anomali maupun gangguan (*noise*), guna memastikan tingkat akurasi dan keandalan performa model komputasi saat memprediksi pola kesehatan mental [18]. Pada tahapan implementasi dalam penelitian ini, serangkaian teknik pra-pemrosesan diterapkan secara sistematis terhadap keseluruhan 2.500 rekam observasi. Langkah investigasi awal difokuskan pada pembersihan himpunan data melalui identifikasi nilai kosong (*missing values*) serta eliminasi entri ganda (*duplicate records*) demi mengamankan integritas informasi. Lebih lanjut, mengingat variabel prediktor yang dianalisis mencakup ragam dimensi, tahapan konversi data atau *encoding* diimplementasikan.

2.4. Feature Selection

Secara analitis, *feature selection* atau seleksi fitur merupakan mekanisme komputasi yang dirancang untuk mengevaluasi dan mengisolasi atribut prediktor yang paling relevan, sekaligus mengeliminasi dimensi data yang *redundan* [19]. Proses reduksi dimensi ini merupakan tahapan esensial dalam rancang bangun model kecerdasan buatan untuk prediksi depresi, mengingat kemampuannya dalam menekan kompleksitas komputasi serta secara persisten mengamplifikasi tingkat akurasi dan interpretabilitas pada algoritma klasifikasi yang diterapkan.

2.5. Imbalance Data Handling

Secara komputasional, SMOTE beroperasi dengan memperluas ruang sampel melalui sintesis data buatan pada instans minoritas, yang secara algoritmik diekstraksi dari mekanisme interpolasi linier antar titik data yang berdekatan [20]. Implementasi strategi augmentasi ini terbukti secara empiris mampu mengamplifikasi tingkat sensitivitas dan performa deteksi pada model pembelajaran mesin, menjadikannya pendekatan yang krusial untuk mengklasifikasikan himpunan data yang asimetris pada kasus prediksi dini risiko depresi.

2.6. Model Training

Secara konseptual, Random Forest merupakan algoritma komputasi berbasis *ensemble learning* yang beroperasi dengan menyusun agregasi dari ratusan pohon keputusan (*decision tree*) independen selama fase pelatihan data, guna mengkalkulasi prediksi akhir secara mayoritas sekaligus menekan risiko *overfitting* [21]. Dari total 2500 observasi, data dipartisi dengan rasio 80:20 menjadi 500 data uji independen dan 2000 data latih yang terdiri dari 1010 risiko rendah, 518 risiko sedang, 472 risiko tinggi, yang selanjutnya diseimbangkan SMOTE menjadi 3030 baris dengan distribusi merata yakni 1010 data per kelas risiko.

2.7. Hyperparameter Tuning

Hyperparameter tuning merupakan mekanisme evaluasi iteratif komputasional yang ditujukan untuk mencari dan mengkalibrasi konfigurasi parameter eksternal pada algoritma klasifikasi, di mana nilai-nilai tersebut secara inheren tidak dapat diestimasi secara otonom oleh model selama proses pelatihan data. Penerapan tahapan optimasi ini berimplikasi esensial dalam memaksimalkan kapabilitas prediktif model sekaligus meminimalisasi probabilitas terjadinya distorsi *overfitting*, sehingga sistem analitik memiliki kemampuan generalisasi yang keandalannya tinggi saat dihadapkan pada observasi data baru [22]. Secara teknis, ruang pencarian pada optimasi ini dikonfigurasi ke dalam lima parameter utama *n_estimators* (50, 100, 150, 200, 300), *max_depth* (5, 10, 15, 20, 25, None), *min_samples_split* (2, 5, 10), *min_samples_leaf* (1, 2, 4), serta *criterion* (*gini* dan *entropy*). Guna menjamin keandalan model yang dihasilkan, setiap fungsi objektif divalidasi secara ketat menggunakan skema 5-fold *cross-validation* pada himpunan data latih yang telah melalui tahapan penyeimbangan kelas, dengan total pencarian ruang parameter dieksekusi sebanyak 30 iterasi.

2.8. Model Evaluation

Evaluasi model merupakan tahapan analitis kuantitatif yang berfokus pada pengukuran tingkat validitas, keandalan, serta kemampuan generalisasi sebuah algoritma pembelajaran mesin dalam mengklasifikasikan kelas target secara presisi pada himpunan data uji [23].

Evaluasi komputasional terhadap performa klasifikasi algoritma Random Forest pasca optimasi diekstraksi dari tabulasi *Confusion Matrix*, yang mencakup empat parameter dasar, *True Positive* (TP), *True Negative* (TN), *False Positive* (FP), dan *False Negative* (FN). Berpijak pada komponen struktural tersebut, penelitian ini menggunakan empat metrik pengujian utama yang direpresentasikan melalui perumusan matematis berikut [24].

$$Accuracy = \frac{\sum_{i=1}^K TP}{N} \quad (1)$$

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

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

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

2.9. Deployment

Deployment atau penyebaran model didefinisikan sebagai fase transisi teknis di mana algoritma *machine learning* yang telah tervalidasi pada tahap eksperimental diintegrasikan secara penuh ke dalam infrastruktur atau lingkungan produksi. Proses ini memungkinkan model komputasi untuk beroperasi secara otonom, menerima masukan data, dan mengeksekusi prediksi secara langsung untuk mendukung pengambilan keputusan klinis atau personal [25].

3. HASIL DAN PEMBAHASAN

Bab ini memaparkan secara komprehensif hasil pengujian empiris serta analisis performa dari arsitektur sistem prediktif yang telah dirancang. Pembahasan difokuskan secara spesifik pada evaluasi metrik klasifikasi dari algoritma Random Forest pasca penerapan optimasi *hyperparameter*, analisis signifikansi variabel prediktor khususnya yang mencakup indikator gaya hidup serta efektivitas implementasi komputasi tersebut dalam memetakan spektrum risiko depresi remaja secara sistematis.

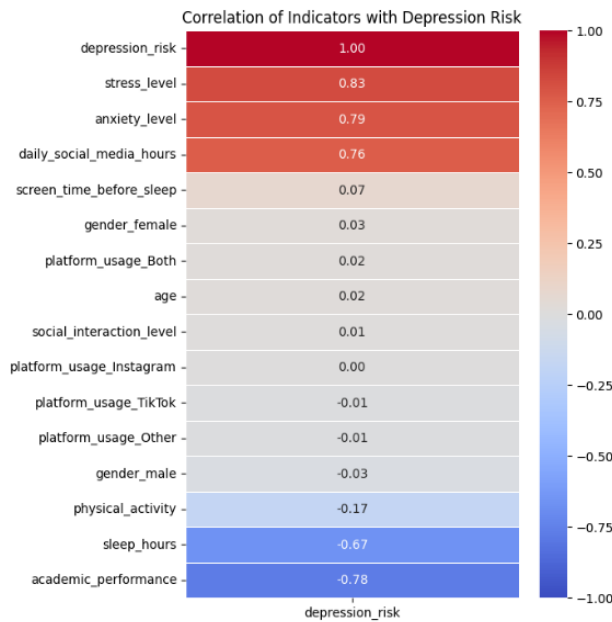
3.1. Data Preprocessing

Berdasarkan implementasi komputasi yang telah dieksekusi, tahapan pra-pemrosesan difokuskan pada purifikasi himpunan observasi, di mana sistem secara otonom mendeteksi dan mengeliminasi instans data yang memuat nilai kosong (*missing values*) sekaligus mereduksi entri ganda (*duplicate records*). Setelah integritas data dipastikan, proses standardisasi format variabel direalisasikan melalui dua pendekatan encoding yang disesuaikan dengan hierarki tipe data kategorikal. Pertama, teknik Ordinal Encoding diterapkan secara spesifik pada variabel ordinal seperti *social interaction level* dan *depression risk*, dengan mentranslasikannya ke dalam representasi numerik berjenjang komputasi, yakni low=0, medium=1, dan high=2. Kedua, mekanisme *One-Hot Encoding* diaplikasikan pada

atribut nominal, mencakup gender serta *platform usage*, yang kemudian dikonversi secara matematis menjadi representasi biner utuh angka 0 dan 1.

3.2. Feature Selection

Pada tahapan seleksi fitur, evaluasi teknis dilakukan dengan mengukur seberapa kuat hubungan antara setiap variabel prediktor dengan target klasifikasi utama, yakni *depression risk*. Pemetaan skor korelasi dari berbagai indikator gaya hidup serta demografi terhadap variabel target tersebut dapat diamati secara detail pada matriks visual (*heatmap*) di Gambar 2.



Gambar 3. Korelasi Indikator dengan Risiko Depresi

Berdasarkan hasil kalkulasi pada Gambar 3, terlihat adanya pemisahan yang jelas antara fitur yang dominan dan fitur yang tidak signifikan. Variabel seperti *stress level* 0.83, *anxiety level* 0.79, dan *daily social media hours* 0.76 menunjukkan korelasi positif yang sangat kuat. Hal ini secara teknis mengindikasikan bahwa lonjakan pada metrik-metrik tersebut berkontribusi langsung terhadap peningkatan risiko depresi. Sebaliknya, fitur seperti *academic performance* -0.78 dan *sleep hours* -0.67 memperlihatkan korelasi negatif yang tajam, yang berarti kondisi yang lebih baik pada indikator tersebut cenderung berbanding terbalik dengan tingkat risiko.

3.3. Imbalance Data Handling

Untuk mengatasi kendala dominasi kelas mayoritas yang teridentifikasi pada tahap eksplorasi awal, tahapan augmentasi data dieksekusi secara spesifik pada himpunan data latih menggunakan teknik SMOTE. Hasil dari proses penyeimbangan ini dapat diamati secara komparatif pada Gambar 4.



Gambar 4. Distribusi Kelas Sebelum dan Sesudah SMOTE

Pada grafik sebelah kiri, terlihat kondisi distribusi kelas sebelum SMOTE diimplementasikan. Kelas *Low* sangat mendominasi dengan jumlah 1.010 rekam observasi, sedangkan kelas *Medium* dan *High* memiliki proporsi yang jauh lebih sedikit, yakni 518 dan 472 data. Melalui intervensi SMOTE, sistem secara algoritmik mensintesis titik-titik data buatan untuk memperbanyak sampel pada kelompok yang kekurangan data. di mana ketiga kategori kelas risiko target kini telah mencapai titik keseimbangan yang sempurna, yaitu masing-masing sebanyak 1.010 rekam data.

3.4. Model Training

Berdasarkan Tabel 2 hasil pengujian komputasi, model dasar ini berhasil mencatatkan *Accuracy* sebesar 0.7720 atau 77,20%. Nilai tersebut diikuti oleh metrik *Precision* yang berada pada level 0.7887, serta *Recall* yang ekuivalen dengan akurasinya, yakni di angka 0.7720. Sementara itu, F1-Score, yang berfungsi sebagai indikator keseimbangan performa klasifikasi, tercatat pada nilai 0.7738. Secara teknis, perolehan metrik awal ini mengindikasikan bahwa model secara bawaan sudah memiliki kemampuan yang cukup solid dalam mendeteksi risiko depresi, namun masih menyisakan ruang yang sangat potensial untuk ditingkatkan lebih tajam melalui proses *hyperparameter tuning*.

Tabel 2. Model *Random Forest Default*

Model	<i>Accuracy</i>	<i>Precision</i>	<i>Recall</i>	<i>F1-Score</i>
Random Forest	0.7720	0.7887	0.7720	0.7738

3.5. Hyperparameter Tuning

Untuk meningkatkan performa dari model dasar, dilakukan penalaan parameter menggunakan metode optimasi Bayesian. Pendekatan ini dipilih karena kemampuannya mencari konfigurasi paling ideal secara lebih efisien dibandingkan pencarian acak. Fokus penalaan mencakup lima komponen utama, yaitu *n_estimators*, *max_depth*, *min_samples_split*, *min_samples_leaf*, dan *criterion*. Kalibrasi kelima parameter komputasi ini sangat penting agar algoritma memiliki fleksibilitas tinggi dalam membaca pola data prediktor, khususnya saat mengevaluasi variabel indikator gaya hidup yang memicu risiko depresi.

Tabel 3. Model Random Forest Hyperparameter Tuning

Metode	Hyperparameter	Accuracy
Bayesian	n_estimator: 100	0.7920
	max_depth: 25	
	min_samples_split: 2	
	min_samples_leaf: 1	
	criterion: entropy	

Rincian hasil identifikasi konfigurasi parameter terbaik dari iterasi Bayesian tersebut dapat diamati pada Tabel 3. Kinerja puncak sistem tercapai pada pengaturan *n_estimator* = 100, *max_depth* = 25, *min_samples_split* = 2, *min_samples_leaf* = 1, dan *criterion entropy*. Penerapan konfigurasi ini terbukti sukses mendongkrak performa klasifikasi sistem secara keseluruhan. Tingkat akurasi komputasi model Random Forest pasca optimasi kini meningkat secara signifikan menjadi 0.7920 (79,20%) jika dibandingkan dengan capaian pengujian model bawaan.

3.6. Model Evaluation

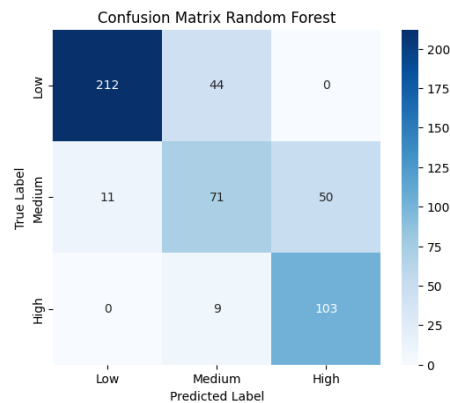
Berdasarkan komparasi data pada Tabel 4, terlihat dengan jelas adanya margin peningkatan performa yang positif berkat penerapan optimasi. Model Random Forest dengan pengaturan *default* awalnya hanya mampu menghasilkan skor akurasi dan *recall* masing-masing di level 0.7720, diiringi oleh nilai presisi 0.7887 dan F1-Score sebesar 0.7738. Namun, setelah algoritma diintervensi menggunakan metode *hyperparameter tuning Bayesian*, kinerjanya mengalami lonjakan yang cukup solid. Tingkat akurasi dan *recall* model hasil penalaan berhasil naik menjadi 0.7920, sementara kemampuan presisinya meningkat menyentuh angka 0.8059, serta metrik keseimbangan F1-Score terdongkrak ke 0.7914.

Tabel 4. Perbandingan Model

Model	Accuracy	Precision	Recall	F1-Score
Random Forest Default	0.7720	0.7887	0.7720	0.7738
Random Forest Hyperparameter Tuning Bayesian	0.7920	0.8059	0.7920	0.7914

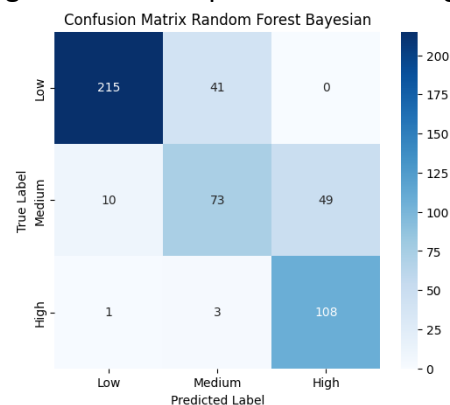
Rincian performa klasifikasi untuk masing-masing tingkatan risiko dapat diamati secara lebih spesifik melalui visualisasi *Confusion Matrix* pada Gambar 5. Pada kelompok responden yang secara aktual memiliki risiko rendah, model komputasi berhasil menebak dengan tepat sebanyak 212 observasi. Meskipun terdapat 44 data yang meleset dan diprediksi sebagai risiko sedang, algoritma sama sekali tidak melakukan kesalahan dengan memprediksi kelompok ini ke tingkat risiko tinggi. Selanjutnya, untuk observasi yang aslinya berada pada kategori risiko sedang, sistem mampu mengenali 71 data secara akurat, sementara sisanya mengalami pergeseran tebakan yakni 11 data diklasifikasikan sebagai risiko rendah dan 50 data diprediksi masuk ke kelompok risiko tinggi. Terakhir, pada kelompok target dengan kerentanan risiko tinggi, kapabilitas deteksi model menunjukkan hasil yang tangguh dengan 103 rekam data berhasil ditebak secara presisi, 9 data yang keliru diklasifikasikan sebagai risiko sedang dan

tidak ada satupun responden berisiko tinggi yang secara fatal ditebak memiliki risiko rendah atau bernilai 0.



Gambar 5. *Confusion Matrix Random Forest*

Rincian performa klasifikasi model pasca optimasi Bayesian pada Gambar 6 memperlihatkan sebaran prediksi yang sangat terstruktur untuk setiap tingkatan risiko. Pada kelompok aktual risiko rendah, algoritma berhasil menebak jitu 215 observasi, dengan 41 data bergeser ke tebakan risiko sedang dan sama sekali tidak ada atau 0 yang keliru diprediksi sebagai risiko tinggi. Selanjutnya, untuk kelas risiko sedang, sistem mampu mengenali 73 data secara tepat, sementara sisanya meleset menjadi 10 tebakan risiko rendah dan 49 tebakan risiko tinggi. Pada kelompok paling krusial yakni risiko tinggi, kapabilitas model terbukti sangat memuaskan dengan 108 tebakan akurat, menyisakan margin kesalahan minor berupa 3 data yang diprediksi sedang dan hanya 1 data yang keliru ditebak sebagai risiko rendah. Secara keseluruhan, konfigurasi angka pada matriks ini secara gamblang menegaskan bahwa optimasi Bayesian tidak hanya menaikkan akurasi, tetapi juga efektif menekan tingkat kesalahan ekstrem seperti gagal mendeteksi pasien berisiko tinggi hanya meleset 1 observasi.



Gambar 6. *Confusion Matrix Random Forest Bayesian*

Analisis performa komprehensif dari arsitektur model terbaik ini juga direkapitulasi secara terperinci melalui metrik *Classification Report* pada Tabel 5. Berdasarkan rincian tersebut, algoritma menunjukkan keunggulan presisi komputasi yang sangat tajam pada kelompok risiko rendah dengan skor mencapai 0.9513. Namun, temuan paling substansial justru terletak pada kelas risiko tinggi, di mana *Recall* melonjak sangat optimal hingga angka 0.9643. Secara teknis, perolehan ini bermakna bahwa sistem berhasil mengidentifikasi hampir

seluruh responden yang secara aktual berada di ambang risiko tinggi dari total 112 sampel uji, tanpa banyak yang terlewat. Walaupun performa pemetaan pada kelas sedang cenderung moderat dengan F1-Score 0.5863, kalkulasi *weighted average* secara keseluruhan yang stabil di kisaran 0.79 membuktikan bahwa model ini memiliki keseimbangan prediktif yang sangat mumpuni sebagai instrumen deteksi dini.

Tabel 5. *Classification Report Random Forest Bayesian*

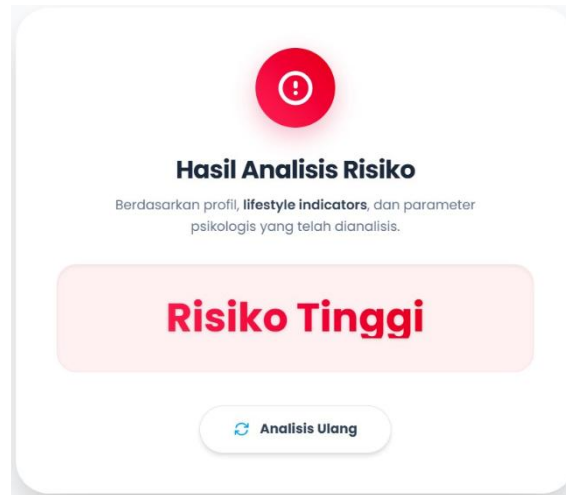
<i>Class</i>	<i>Precision</i>	<i>Recall</i>	<i>F1-Score</i>	<i>Support</i>
<i>Low</i>	0.9513	0.8398	0.8921	256
<i>Medium</i>	0.6239	0.5530	0.5863	132
<i>High</i>	0.6879	0.9643	0.8030	112
<i>Accuracy</i>			0.7920	500
<i>Macro Avg</i>	0.7544	0.7857	0.7605	500
<i>Weighted Avg</i>	0.8059	0.7920	0.7914	500

3.7. Deployment

Tahapan final dari penelitian ini adalah mentransformasikan model prediktif ke dalam lingkungan operasional. Berdasarkan visualisasi pada Gambar 7, arsitektur antarmuka dirancang secara terstruktur untuk mengumpulkan data variabel prediktor secara mandiri dari pengguna. Form isian tersebut dikategorikan ke dalam tiga segmen observasi utama, data profil dan rekam akademik seperti usia, jenis kelamin, nilai IPK, observasi *lifestyle indicators* seperti durasi penggunaan media sosial, jam tidur, aktivitas fisik, serta asesmen interaksi dan psikologis seperti tingkat stres dan kecemasan. Seluruh parameter masukan ini dirancang agar selaras dengan format data latih pada tahap pra-pemrosesan sebelumnya. Saat pengguna mengeksekusi tombol Analisis Risiko Depresi, sistem akan melakukan pengiriman data secara *real-time* menuju model Random Forest di *backend*.

Gambar 7. Antarmuka Halaman Prediksi Depresi

Visualisasi respons keluaran dari proses inferensi algoritma tersebut direpresentasikan melalui antarmuka halaman hasil analisis, sebagaimana yang diperlihatkan pada Gambar 8. Pada skenario simulasi ini, setelah *backend* memproses seluruh rangkaian masukan data yang mencakup profil demografi, lifestyle indicators, hingga asesmen psikologis, model Random Forest Bayesian secara komputasional mengklasifikasikan pengguna ke dalam kategori Risiko Tinggi.



Gambar 8. Antarmuka Halaman Hasil Prediksi Depresi

4. KESIMPULAN

Penelitian ini telah berhasil mengembangkan sistem cerdas untuk deteksi dini risiko depresi remaja dengan mengintegrasikan analisis *lifestyle indicators*. Melalui pra-pemrosesan data menggunakan teknik SMOTE, pemodelan klasifikasi dibangun menggunakan algoritma Random Forest yang dioptimasi melalui pendekatan Bayesian. Intervensi metode penalaan tersebut secara empiris terbukti sukses mendongkrak akurasi sistem dari 77,20% pada konfigurasi bawaan menjadi 79,20%. Keandalan model ini semakin diperkuat oleh kemampuannya memitigasi kesalahan deteksi fatal, yang ditandai dengan pencapaian *Recall* sebesar 0.9643 pada kategori risiko tinggi tanpa adanya satupun responden berisiko tinggi yang keliru ditebak sebagai risiko rendah. Pada akhirnya, keberhasilan *deployment* model prediktif ini ke dalam bentuk antarmuka aplikasi interaktif menegaskan bahwa luaran penelitian tidak sekadar berupa evaluasi komputasi, melainkan telah berwujud instrumen skrining preventif yang fungsional.

DAFTAR PUSTAKA

- [1] World Health Organization, "Mental health of adolescents," World Health Organization. Accessed: May 25, 2026. [Online]. Available: <https://www.who.int/news:room/fact-sheets/detail/adolescent-mental-health>
- [2] Gadjah Mada University & The University of Queensland, "Indonesia - National Adolescent Mental Health Survey (I-NAMHS) Report," Yogyakarta, 2022. Accessed: May 25, 2026. [Online]. Available: <https://acmhr.org/outputs/reports/12-i-namhs-report-bahasa-indonesia>

- [3] P. Y. Sari *et al.*, "KONDISI KESEHATAN MENTAL REMAJA." [Online]. Available: <http://jurnal.globalhealthsciencegroup.com/index.php/JPPP>
- [4] Badan Kebijakan Pembangunan Kesehatan, "Fact Sheet Kesehatan Jiwa: Kesehatan Jiwa di Indonesia," Jakarta, 2023. Accessed: May 25, 2026. [Online]. Available: <https://repository.badankebijakan.kemkes.go.id/id/eprint/5532/>
- [5] J. Zhang, D. Liu, L. Ding, and G. Du, "Prevalence of depression in junior and senior adolescents," *Front. Psychiatry*, vol. 14, 2023, doi: 10.3389/fpsyt.2023.1182024.
- [6] A. Wahyudi Oktavia Gama, A. Degni Melsy Grren, I. Gusti Ngurah Darma Paramartha, G. Humaswara Prathama, N. Made Widnyani, and M. Wira Putra Dananjaya, "PENERAPAN ALGORITMA NAÏVE BAYES UNTUK PREDIKSI PENYAKIT DEPRESI PADA MAHASISWA." [Online]. Available: <http://jurnal.globalhealthsciencegroup.com/index.php/JLH>
- [7] Y. Cao *et al.*, "Machine Learning Approaches for Mental Illness Detection on Social Media: A Systematic Review of Biases and Methodological Challenges," 2025, *International Society for Data Science and Analytics*. doi: 10.35566/jbds/caoyc.
- [8] E. Murakami, T. Shionoya, S. Komenoi, Y. Suzuki, and F. Sakane, "Cloning and characterization of novel testis-Specific diacylglycerol kinase η splice variants 3 and 4," *PLoS One*, vol. 11, no. 9, Sep. 2016, doi: 10.1371/journal.pone.
- [9] Fikri Muhamad Fahmi, Budiman Budiman, and Nur Alamsyah, "Optimizing IT Remote Workers Mental Health Prediction using Feature Engineering," *Int. J. Sci. Math. Educ.*, vol. 2, no. 2, pp. 01–09, Apr. 2025, doi: 10.62951/ijsme.v2i2.193.
- [10] A. Alfahrizi Putra Arifiansyah, M. Afandi, D. Dwi Riskianto, and K. Kunci, "Prediksi Depresi Mahasiswa Menggunakan Algoritma Random Forest Berbasis Data Psikososial Depression Prediction Among University Students Using a Random Forest Algorithm Based on Psychosocial Data," *Jurnal Riset Sistem dan Teknologi Informasi (RESTIA)*, vol. 4, no. 1, pp. 12–21, 2026.
- [11] D. Nickson, C. Meyer, L. Walasek, and C. Toro, "Prediction and diagnosis of depression using machine learning with electronic health records data: a systematic review," *BMC Med. Inform. Decis. Mak.*, vol. 23, no. 1, Dec. 2023, doi: 10.1186/s12911-023-02341-x.
- [12] D. S. Sonnadara, S. D. Viswakula, and M. D. T. Attygalle, "A machine learning approach for predicting the depression status among undergraduates: A case study from Sri Lanka," Nov. 16, 2023. doi: 10.21203/rs.3.rs-3555106/v1.
- [13] M. Arifin and S. Adiyono, "Hyperparameter Tuning in Machine Learning to Predicting Student Academic Achievement," 2024. [Online]. Available: <http://ijair.id>
- [14] B. Ramadhan and S. F. Pane, "Pengaruh Hyperparameter Tuning untuk Efektivitas pada Pendekatan Hybrid dalam Mendiagnosis Stres dan Depresi : Tinjauan Studi Literatur," *Jurnal Tekno Insentif*, vol. 18, no. 2, pp. 104–118, Dec. 2024, doi: 10.36787/jti.v18i2.1516.
- [15] Z. Chen, D. Liu, W. J. Marrero, N. C. Jacobson, and T. Thesen, "A hierarchical Bayesian approach for predictive analytics of depression severity in medical students," *Healthcare Analytics*, Jun. 2026, doi: 10.1016/j.health.2026.100453.

- [16] S. Beekhiani, "Teen social media usage and mental health," Kaggle. Accessed: May 25, 2026. [Online]. Available: <https://www.kaggle.com/datasets/sureshbeekhiani/teen-social-media-usage-and-mental->
- [17] G. M. Kim, S. Cha, M. Jung, and S. Kim, "Machine learning prediction of depression in culturally diverse families: Findings from the Korea Community Health Survey," *Front. Public Health*, vol. 13, 2025, doi: 10.3389/fpubh.2025.1666084.
- [18] H. U. Khan, A. Naz, F. K. Alarfaj, and N. Almusallam, "Analyzing student mental health with RoBERTa-Large: a sentiment analysis and data analytics approach," *Front. Big Data*, vol. 8, 2025, doi: 10.3389/fdata.2025.1615788.
- [19] D. Enkhbayar *et al.*, "Explainable Artificial Intelligence Models for Predicting Depression Based on Polysomnographic Phenotypes," *Bioengineering*, vol. 12, 2025, doi: 10.3390/bioengineering12020186.
- [20] B. Zhang *et al.*, "A refined SMOTE-ENN optimization method based on machine learning for heart rate variability data classification," *Front. Digit. Health*, vol. 8, 2026, doi: 10.3389/fdgth.2026.1749570.
- [21] C. Yu, X. Kong, W. Yu, X. Ni, J. Chen, and X. Liao, "Machine learning models for predicting the risk of depressive symptoms in Chinese college students," 2025.
- [22] K. O. Asare, Y. Terhorst, J. Vega, E. Peltonen, E. Lagerspetz, and D. Ferreira, "Predicting Depression From Smartphone Behavioral Markers Using Machine Learning Methods, Hyperparameter Optimization, and Feature Importance Analysis: Exploratory Study," *JMIR Form. Res.*, 2025.
- [23] P. W. Simarmata and P. T. Prasetyaningrum, "Development of a Student Depression Prediction Model Based on Machine Learning with Algorithm Performance Evaluation," *Journal of Information Systems and Informatics*, vol. 7, 2025, doi: 10.51519/journalisi.v7i2.1087.
- [24] G. Zeng, "Invariance Properties and Evaluation Metrics Derived from the Confusion Matrix in Multiclass Classification," *Mathematics*, vol. 13, 2025, doi: 10.3390/math13162609.
- [25] A. P. Yan *et al.*, "A roadmap to implementing machine learning in healthcare: from concept to practice," *Front. Digit. Health*, vol. 7, 2025, doi: 10.3389/fdgth.2025.1462751.