

MODEL PREDIKSI GANGGUAN DEPRESI BERDASARKAN INTERAKSI SOSIAL DAN PENGGUNAAN MEDIA SOSIAL DENGAN ALGORITMA NAIVE BAYES

Bagus Riyadi, Aprilisa Arum Sari, Dwi Hartanti

Teknik Informatika, Universitas Duta Bangsa Surakarta

Jl. Bhayangkara No.55, Tipes, Kec. Serengan, Kota Surakarta, Jawa Tengah 57154.

Pejuangsejati22@gmail.com

ABSTRAK

Meningkatnya prevalensi aktivitas digital di kalangan remaja telah memicu transformasi pola interaksi sosial. Perubahan ini berdampak pada kerentanan terhadap gangguan psikologis seperti depresi, namun instrumen deteksi saat ini masih terbatas pada analisis variabel tunggal yang kurang komprehensif. Penelitian ini bertujuan mengusulkan model inovatif untuk mengklasifikasikan risiko depresi dengan mengintegrasikan dimensi perilaku digital dan dinamika interaksi sosial. Metodologi studi mengikuti siklus *Knowledge Discovery in Databases* (KDD) dengan menerapkan algoritma *Gaussian Naive Bayes* pada dataset berjumlah 8.000 observasi. Variabel inti yang dianalisis meliputi profil demografis, durasi penggunaan perangkat (*screen time*), pola istirahat, dan tingkat interaksi sosial. Temuan studi menunjukkan model berhasil mencapai tingkat akurasi sebesar 75,88%. Intensitas penggunaan media sosial dan durasi tidur teridentifikasi sebagai faktor determinan paling krusial dalam memicu risiko depresi. Sebagai kontribusi praktis, model ini diintegrasikan pada antarmuka berbasis web menggunakan *Streamlit* untuk memfasilitasi skrining kesehatan mental secara mandiri dan *real-time*.

Kata kunci : Depresi, *Gaussian Naive Bayes*, KDD, Media Sosial, Remaja

1. PENDAHULUAN

Perkembangan teknologi digital saat ini telah meningkatkan penggunaan media sosial di kalangan remaja yang secara signifikan mengubah pola interaksi dari komunikasi langsung ke ranah digital. Perubahan ini berpotensi menurunkan kualitas hubungan sosial serta memicu perasaan kesepian dan isolasi emosional yang berkontribusi terhadap meningkatnya risiko gangguan psikologis, seperti stres, kecemasan, dan depresi [1].

Meskipun aktivitas digital menghasilkan data yang kaya sebagai indikator kondisi mental seperti durasi penggunaan media sosial dan pola interaksi penelitian sebelumnya menunjukkan bahwa sebagian besar instrumen deteksi masih terbatas pada analisis variabel tunggal yang kurang komprehensif [2], [3]. Integrasi antara variabel aktivitas digital dan dinamika interaksi sosial dalam satu model prediksi masih jarang dilakukan, sehingga pendekatan yang ada saat ini dianggap belum mampu menggambarkan kondisi mental individu secara menyeluruh.

Berdasarkan permasalahan tersebut, penelitian ini mengusulkan pengembangan model klasifikasi risiko depresi pada remaja menggunakan algoritma *Gaussian Naive Bayes* yang mengintegrasikan variabel penggunaan media sosial dan interaksi sosial. Tujuan utama dari penelitian ini adalah untuk membangun model prediksi yang lebih inovatif serta menganalisis faktor-faktor dominan yang memengaruhi risiko depresi, khususnya intensitas penggunaan media sosial dan durasi tidur. Sebagai rencana penyelesaian masalah, model klasifikasi yang dihasilkan akan diimplementasikan ke dalam aplikasi berbasis web menggunakan *Streamlit*. Hal ini dilakukan guna menyediakan platform skrining kesehatan mental yang dapat diakses secara mandiri

dan *real-time*, sehingga mendukung upaya deteksi dini gangguan mental pada remaja secara lebih komprehensif dan efektif [5].

2. TINJAUAN PUSTAKA

Hasil sintesis literatur menunjukkan dominasi penelitian yang bersifat parsial dengan keterbatasan pada variabel tunggal. Riset ini hadir untuk menutup celah tersebut melalui integrasi variabel perilaku digital dan interaksi sosial ke dalam model *Gaussian Naive Bayes*. Pendekatan integratif ini diharapkan mampu mengoptimalkan performa klasifikasi dalam memprediksi risiko kesehatan mental remaja secara lebih akurat.

2.1. Penelitian Terdahulu

Studi ini dirangkai dengan mengacu berbagai studi sebelumnya yang relevan. Guna menjaga kebaruan dan validitas informasi, referensi yang digunakan dibatasi maksimal tiga tahun terakhir yaitu 2022 hingga 2025. Adapun kajian pustaka yang dipakai selaku landasan pada studi ini diuraikan sebagai berikut:

Efektivitas Deteksi Dini Gangguan Mental : Penelitian memperlihatkan algoritma *Naive Bayes* bisa mengklasifikasikan tingkat risiko kesehatan mental mahasiswa secara akurat. Model ini mengidentifikasi faktor-faktor seperti tekanan akademik, kondisi sosial, dan aktivitas harian sebagai variabel penting dalam deteksi dini gangguan mental [1].

Pengaruh Durasi *Screen time* : Penggunaan metode *Naive Bayes* juga diterapkan dalam analisis perilaku digital, termasuk durasi penggunaan perangkat dan dampaknya terhadap kondisi psikologis. Temuan studi memperlihatkan model ini mampu mengklasifikasikan tingkat risiko psikologis

secara selaras dengan tingginya paparan aktivitas digital [2].

Korelasi Penggunaan Media Sosial : Analisis terhadap data media sosial menunjukkan adanya hubungan antara intensitas penggunaan platform digital dengan kondisi kesehatan mental. Metode *Naive Bayes* terbukti efektif dalam mengidentifikasi kecenderungan gangguan mental berbasis data teks dan aktivitas pengguna [3].

Signifikansi Pola Interaksi Sosial : Faktor interaksi sosial menjadi salah satu variabel penting dalam klasifikasi kesehatan mental. Penelitian menunjukkan bahwa kombinasi variabel sosial dengan pendekatan machine learning bisa menaikkan akurasi terhadap mendeteksi tingkat stres dan depresi individu [6].

Integrasi Variabel Sosial : Integrasi variabel interaksi sosial ke dalam model klasifikasi berbasis *Naive Bayes* terbukti meningkatkan performa sistem, khususnya dalam hal akurasi dan kemampuan identifikasi individu berisiko [7].

Stabilitas pada Data Perilaku Digital: Dalam pengolahan data kontinu seperti aktivitas digital pengguna, algoritma *Naive Bayes* menunjukkan performa stabil dan konsisten. Hal ini menjadikannya metode yang cocok untuk analisis data perilaku berbasis waktu dan intensitas penggunaan [8].

Implementasi di Lingkungan Pendidikan : Penerapan *Naive Bayes* dalam lingkungan pendidikan menunjukkan hasil yang efektif sebagai alat skrining dini gangguan mental pada siswa. Model ini mampu membantu mengidentifikasi siswa yang memerlukan perhatian khusus secara lebih cepat dan akurat [4].

2.2. Landasan Teori

a. Data Mining dan Machine Learning

Data mining ialah komponen siklus KDD atau “*Knowledge Discovery in Databases*” yang berfungsi mengekstraksi pola tersembunyi dari dataset besar. Dalam konteks kesehatan mental, teknik klasifikasi telah terbukti efektif untuk mendeteksi kondisi psikologis dalam ranah pendidikan [1], [4]. Penelitian ini memfokuskan pendekatan data mining pada fungsi prediktif dengan mengintegrasikan data aktivitas digital dan interaksi sosial untuk mengklasifikasikan risiko depresi remaja secara lebih menyeluruh.

b. Deteksi Dini Kesehatan Mental (*Mental Health Screening*)

Skrining dini bertujuan mengidentifikasi potensi gangguan psikologis sebelum mencapai tahap klinis yang serius. Meskipun penelitian terdahulu telah memanfaatkan indikator durasi perangkat dan aktivitas media sosial [2], [3], serta variabel interaksi sosial [6], [7], pendekatannya masih bersifat parsial. Riset ini mengusulkan integrasi variabel-variabel tersebut secara simultan untuk menghasilkan analisis deteksi dini yang lebih komprehensif.

c. Algoritma *Naive Bayes* dan *Gaussian Naive Bayes*

Naive Bayes ialah metode klasifikasi *probabilistik* populer pada riset kesehatan mental karena kesederhanaan dan stabilitas performanya [1], [8]. Guna mengatasi keterbatasan variabel tunggal atau data teks saja, penelitian ini menerapkan varian *Gaussian Naive Bayes*. Penggunaan varian ini bertujuan untuk mengakomodasi data numerik kontinu, seperti durasi penggunaan media sosial, sehingga model dapat mengolah berbagai jenis variabel secara fleksibel dan meningkatkan akurasi klasifikasi akhir.

d. Implementasi Sistem dan Teknologi yang Digunakan

Implementasi sistem pada studi ini mengacu pada tahapan KDD atau “*Knowledge Discovery in Databases*”, tersusun atas :

➤ Data Selection (Seleksi Data)

Tahap awal studi ini ialah melakukan seleksi data yang relevan tujuan penelitian. Dataset yang digunakan adalah *social_media_mental_health.csv*. Dataset yang diperoleh dari Platform Kaggle [9]. Selanjutnya, dilakukan pemilihan atribut yang berkaitan dengan kesehatan mental remaja, yaitu durasi penggunaan media sosial (*daily social media hours*), jenis platform yang digunakan (*platform usage*), dan tingkat interaksi sosial (*social interaction level*). Pemilihan variabel tersebut didasarkan pada hasil tinjauan pustaka yang menunjukkan bahwa aktivitas digital dan interaksi sosial memiliki keterkaitan dengan kondisi kesehatan mental remaja.

➤ Pre-processing (Pembersihan Data)

Untuk menjamin kualitas dan konsistensi data, pada tahap ini dilakukan proses pembersihan data. Proses tersebut mencakup penanganan missing values, penghapusan data duplikat, serta identifikasi dan perbaikan inkonsistensi data. Dengan demikian, tahap ini diarahkan untuk meminimalkan noise yang berpotensi memengaruhi kinerja model, sehingga data yang digunakan pada proses berikutnya memiliki tingkat validitas yang lebih baik.

➤ Transformation (Transformasi Data)

Data yang sudah melewati tahapan pembersihan kemudian ditransformasikan agar sesuai dengan kebutuhan algoritma pembelajaran mesin. Proses transformasi meliputi konversi data kategorikal menjadi data numerik memakai teknik *label encoding*, serta pembagian dataset menjadi data latih atau “*training set*” serta data uji atau “*testing set*”. Tahap ini bertujuan untuk memastikan data dapat diolah secara optimal oleh algoritma klasifikasi.

➤ Data Mining (Pemodelan)

Tahapan pemodelan ialah inti tahapan KDD atau “*Knowledge Discovery in Databases*”,

yaitu penerapan algoritma guna mengekstraksi pola dari data. Pada studi ini digunakan algoritma *Gaussian Naive Bayes* yang diimplementasikan memakai bahasa pemrograman Python dengan bantuan pustaka *Scikit-Learn*. Algoritma ini dipakai guna menghitung probabilitas kecenderungan depresi berdasarkan variabel aktivitas digital dan interaksi sosial, dengan asumsi bahwa data numerik mengikuti distribusi normal.

➤ Interpretasi dan Evaluasi (*Interpretation and Evaluation*)

Tahap final dalam siklus KDD ini melibatkan pengujian performa model menggunakan data uji melalui metrik *Accuracy*, *Precision*, *Recall*, dan *Confusion Matrix*. Hasil evaluasi diterapkan pada dashboard interaktif berbasis Streamlit yang beroperasi dengan *real-time*. Antarmuka ini merepresentasikan luaran prediksi dalam dua kategori risiko, yaitu Risiko Tinggi atau “*High Risk*” serta Risiko Rendah atau “*Low Risk*”, yang disertai dengan visualisasi tingkat probabilitas hasil perhitungan Teorema Bayes.

Selain menampilkan hasil klasifikasi individu berdasarkan variabel durasi media sosial dan interaksi sosial, dashboard ini juga menyajikan grafik distribusi data serta metrik performa algoritma secara transparan. Integrasi ini bertujuan untuk menghasilkan sebuah *early warning system* yang informatif, sehingga memudahkan pengguna dalam memahami profil risiko kesehatan mental secara cepat dan akurat.

e. Teorema Bayes

Merupakan dasar dari algoritma prediksi probabilitas. Secara matematis, formula ini didefinisikan sebagai berikut:

$$P(C|X) = \frac{P(X|C) \cdot P(C)}{P(X)} \quad (1)$$

Keterangan:

- $P(C|X)$: Probabilitas kelas C (misal : Depresi) jika diketahui sampel uji memiliki kondisi X (*Posterior probability*).
- $P(X|C)$: Probabilitas atribut X jika diketahui kelasnya adalah C (*Likelihood*).
- $P(C)$: Probabilitas kemunculan kelas C secara keseluruhan pada data historis (*Prior probability*).
- $P(X)$: Probabilitas kemunculan kondisi pengujian X (*Evidence*).

f. Fungsi Kepadatan Peluang *Gaussian* (*Gaussian PDF*)

Metode ini diaplikasikan khusus untuk menghitung atribut yang memiliki nilai kontinu (seperti durasi *screen time*) ke dalam *likelihood* probabilitas dengan distribusi normal. Rumusnya:

$$f(x; \mu, \sigma) = \frac{1}{\sqrt{2\pi\sigma^2}} e^{-\frac{(x-\mu)^2}{2\sigma^2}} \quad (2)$$

Formula ini menghitung nilai probabilitas untuk fitur x (nilai durasi media sosial) dengan

memanfaatkan nilai μ (*mean* atau rata-rata durasi pada kelas tertentu) dan nilai σ^2 (*varians* dari data pada kelas tersebut). Fokus utama dari rumus ini adalah membaca fluktuasi data numerik secara halus tanpa harus mengelompokkannya secara paksa, sehingga probabilitas yang dihasilkan sangat presisi.

g. Metrik Evaluasi Model (*Confusion Matrix*)

Sebagai alat fundamental dalam evaluasi kinerja klasifikasi, *Confusion Matrix* digunakan untuk membandingkan hasil prediksi sistem dengan kondisi aktual. Berdasarkan matriks tersebut, kemudian diturunkan berbagai metrik untuk mengukur sejauh mana performa model klasifikasi probabilistik.

➤ Akurasi (*Accuracy*)

Mengukur proporsi tebakan yang benar secara keseluruhan dari total seluruh sampel yang diuji.

$$Accuracy = \frac{(TP + TN)}{TP + TN + FP + FN} \quad (3)$$

Keterangan :

- TP (*True Positive*): Jumlah remaja yang secara aktual berada pada kategori depresi serta berhasil diklasifikasikan tepat oleh model sebagai *High Risk*.
- TN (*True Negative*): Jumlah remaja yang secara aktual berada dalam kondisi normal/stabil serta berhasil diklasifikasikan tepat oleh model sebagai *Low Risk*.
- FP (*False Positive*): Kesalahan klasifikasi di mana model memprediksi remaja sebagai *High Risk*, padahal kondisi aktualnya adalah *Low Risk* (*False Alarm*).
- FN (*False Negative*): Kesalahan klasifikasi di mana model memprediksi remaja sebagai *Low Risk*, padahal secara aktual remaja tersebut berada pada kategori *High Risk*.

➤ Presisi (*Precision*)

Mengukur rasio ketepatan prediksi positif di antara keseluruhan prediksi positif yang dilakukan model.

$$Precision = \frac{TP}{TP + FP} \quad (4)$$

Tujuan utama presisi adalah meminimalkan false alarm. Nilai presisi yang tinggi sangat penting dalam aplikasi kesehatan agar konselor tidak membuang waktu mengintervensi remaja yang pada kenyataannya tidak memiliki indikasi depresi (FP).

➤ Recall (*Sensitivity*)

Mengukur rasio seberapa baik model dalam menemukan (mendeteksi) seluruh kasus positif dari keseluruhan populasi positif yang sebenarnya.

$$Recall = \frac{TP}{TP + FN} \quad (5)$$

Metrik ini merupakan metrik paling krusial dalam pendeteksian dini kesehatan mental. Tujuannya adalah memastikan model *se-sensitif* mungkin dalam mendeteksi gejala. Dalam skenario nyata, nilai *Recall* harus

dioptimalkan untuk memastikan tidak ada remaja yang terlewat dari deteksi (menekan angka FN) sebelum depresi tersebut berkembang menjadi tahap klinis yang berbahaya.

3. METODOLOGI PENELITIAN

Penelitian ini mengadopsi struktur metodologi sistematis yang mengintegrasikan teknik data mining dalam kerangka kerja KDD atau “*Knowledge Discovery in Databases*”. Pemanfaatan instrumen ini difokuskan pada analisis mendalam terhadap variabel interaksi digital guna menghasilkan model klasifikasi dengan akurasi dan validitas tinggi dalam mendeteksi potensi depresi di kalangan remaja.

3.1. Jenis dan Sumber Data

Objek observasi pada studi ini adalah pola perilaku remaja dengan fokus pada variabel aktivitas digital dan profil demografis. Data yang digunakan merupakan data sekunder berjumlah 8.000 observasi yang bersumber dari repositori *Kaggle*.

Tabel 1. Definisi Operasional Variabel

No	Variabel	Jenis Data	Keterangan
1	Umur	Numerik	Usia responden dalam satuan tahun.
2	Jenis Kelamin	Kategorikal	Identitas gender responden (<i>Male/Female</i>).
3	Screen time	Numerik	Durasi harian penggunaan media sosial (jam).
4	Durasi Tidur	Numerik	Rata-rata waktu istirahat harian (jam).
5	Interaksi Sosial	Numerik	Tingkat keterlibatan sosial atau interaksi pengguna (misal: skala 1-5).
6	Label Risiko	Kategorikal	Target output: <i>High Risk</i> atau <i>Low Risk</i> .

3.2. Metode Pengumpulan Data

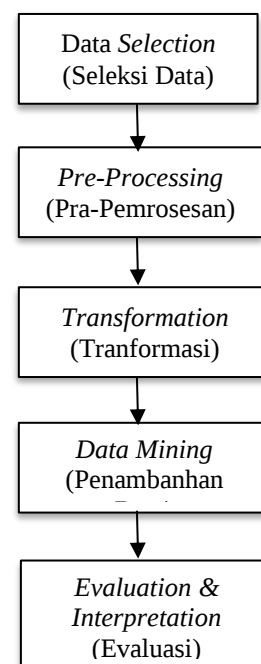
- Strategi Perolehan Data : Mengingat studi ini berbasis pada pemodelan pola historis digital, maka pengumpulan data tidak dilakukan melalui instrumen primer seperti kuesioner fisik maupun observasi lapangan. Fokus utama diarahkan pada ekstraksi data digital yang telah tervalidasi secara publik.
- Sumber dan Ekstraksi : Data sampel diekstraksi dari repositori publik *Kaggle* yang menyediakan dataset riset machine learning secara terbuka (*open source*). Dataset yang dipilih memiliki relevansi variabel yang kuat terhadap kesehatan mental dan perilaku penggunaan media sosial remaja.
- Lingkungan Kerja dan Platform: Berkas mentah berformat *Comma Separated Values* (CSV) yang telah diperoleh diintegrasikan ke dalam lingkungan kerja *Visual Studio Code* (VS Code)

menggunakan bahasa pemrograman *Python*. Platform ini dipilih karena stabilitasnya dalam mengelola pustaka data dan efisiensi dalam penulisan skrip komputasi.

- Strukturisasi Data melalui *DataFrame* : Guna mengoptimalkan proses analisis, dataset tersebut ditransformasikan ke dalam struktur tabel memori (*DataFrame*) dengan bantuan pustaka *Pandas*. Tahapan ini memfasilitasi peneliti dalam melakukan *pra-pemrosesan*, sterilisasi data (*data cleaning*), serta seleksi kolom-kolom prediktor yang krusial bagi algoritma *Gaussian Naive Bayes*.

3.3. Metode Pengembangan Sistem

Arsitektur pengembangan sistem klasifikasi risiko depresi ini dirancang secara terstruktur memakai pendekatan KDD atau “*Knowledge Discovery in Databases*” yang terintegrasi dengan teknik data mining. Kerangka kerja ini dipilih karena kemampuannya dalam mentransformasi data mentah menjadi pola pengetahuan prediktif yang valid melalui lima *fase integral*. Secara visual, alur kerja sistem diilustrasikan pada Gambar 3.1.



Gambar 1. Flowchart Tahapan Knowledge Discovery in Databases (KDD)

Secara mendetail, implementasi teknis pada setiap tahapan pengembangan sistem dijelaskan sebagai berikut :

3.4. Seleksi Data (*Data Selection*)

Fase ini bertujuan guna mengidentifikasi dan mengekstraksi atribut-atribut yang memiliki relevansi teoretis tinggi terhadap indikator depresi klinis dari total 8.000 observasi. Peneliti melakukan inspeksi terhadap dataset sekunder *social_media_mental_health.csv*

untuk memilih kolom-kolom strategis yang akan dijadikan fitur prediktor dalam model. Atribut yang dipilih meliputi variabel demografis Umur, Jenis Kelamin serta variabel perilaku digital dan istirahat (*Screen time*), Durasi Tidur dan Interaksi Sosial), dengan target klasifikasi berupa Label Risiko.

3.5. Pra-pemrosesan Data (Data Pre-processing)

Tahapan pra-pemrosesan esensial dilakukan guna menjamin mutu input data sebelum eksekusi algoritma. Proses sterilisasi data mencakup dua langkah utama :

- Penanganan Nilai Kosong (*Missing Values*) : Dilakukan pemeriksaan menyeluruh untuk mendeteksi adanya redundansi atau data kosong. Peneliti menerapkan teknik imputasi median atau menghapus baris data yang tidak lengkap guna menghindari bias sistematis pada model komputasi.
- Kodifikasi Data Kategorikal: Karena algoritma *Gaussian Naïve Bayes* sekedar dapat memproses input numerik, atribut kategorikal “Jenis Kelamin” (*Female/Male*) ditransformasikan menjadi bentuk kuantitatif (misalnya, 0 untuk Female dan 1 untuk Male) menggunakan pustaka *Scikit-Learn LabelEncoder*.

3.6. Transformasi Data (Data Transformation)

Pada fase ini, data ditransformasikan ke dalam format yang optimal untuk fase pelatihan. Peneliti melakukan pemisahan dataset (*Data Splitting*) secara acak (*stratified shuffle*) guna membangun model yang stabil sekaligus mengukur kemampuan generalisasinya terhadap data baru. Data didistribusikan dengan proporsi 80% selaku data latih atau “*training set*” serta 20% sebagai data uji atau “*testing set*”.

Tabel 2. Distribusi Kuantitatif Dataset

Kategori Data	Persentase	Jumlah Record (Baris)	Kegunaan
Data Latih	80%	6.400	Membangun dan melatih model komputasi Bayes.
Data Uji	20%	1.600	Mengevaluasi performa dan akurasi model.
Total	100%	8.000	Populasi data yang valid

3.7. Data Mining (Gaussian Naïve Bayes)

Sebagai inti dari kerangka kerja KDD, pada tahap ini model klasifikasi dibangun dengan menggunakan algoritma *Gaussian Naïve Bayes*. Pemilihan model ini didasarkan pada efektivitasnya dalam memproses fitur-fitur kontinu (numerik) seperti durasi penggunaan perangkat dan waktu tidur tanpa harus melakukan diskritisasi data secara paksa.

Secara teknis, model mengasumsikan bahwa distribusi probabilitas setiap fitur kontinu mengikuti

pola distribusi normal (*Gaussian*). Model menghitung parameter rata-rata (μ) serta standar deviasi (σ) untuk semua fitur pada setiap kelas target (*High Risk/Low Risk*), lalu menggunakan rumus fungsi kepadatan peluang *Gaussian* untuk menentukan *probabilitas prior* dan *likelihood*. Klasifikasi akhir ditentukan berdasarkan *probabilitas posterior* tertinggi.

3.8. Evaluasi dan Implementasi Antarmuka

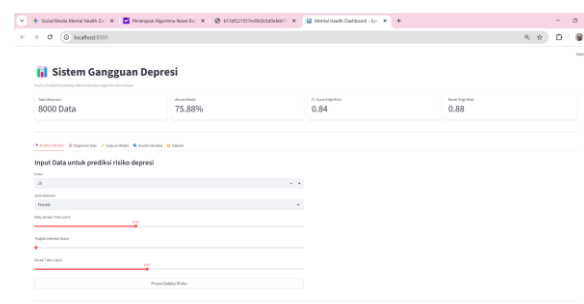
Fase akhir dari KDD adalah evaluasi model dan *deployment* sistem:

- Evaluasi Model : Kinerja model diuji menggunakan data testing (1.600 record) dan diukur melalui metrik *Confusion Matrix*. Berdasarkan pengujian kuantitatif, model yang sudah diciptakan memakai *classification report* dari *library* *sklearn*. Didapat temuan Akurasi sebesar 75,88% dan tingkat Sensitivitas (*Recall*) sebesar 88%. Hal ini menunjukkan bahwa sistem memiliki reliabilitas tinggi, khususnya dalam meminimalisir risiko penderita yang tidak terdeteksi (*false negative*).

	precision	recall	F1 score	support	
high Risk	0.89	0.88	0.84	1128.00	
low Risk	0.62	0.46	0.53	472.00	
accuracy	0.76	0.76	0.76	0.76	
macro avg	0.71	0.67	0.68	3000.00	
weighted avg	0.75	0.76	0.75	3000.00	

Gambar 2. Hasil *Classification report*

- Implementasi Antarmuka Interaktif : Sebagai bentuk hilirisasi riset, model yang sudah tervalidasi diintegrasikan ke dalam *dashboard* web interaktif memakai *framework* *Streamlit* dan platform *VS Code*. Antarmuka ini dirancang secara detail guna memudahkan pemakainya dalam melaksanakan prediksi risiko dengan *real-time*, sebagaimana diilustrasikan pada Gambar 2.



Gambar 3. Implementasi Antarmuka Sistem Berbasis *Streamlit*,

4. HASIL DAN PEMBAHASAN

Bab ini memaparkan temuan implementasi model *Gaussian Naïve Bayes* bagi klasifikasi risiko depresi pada remaja, serta pembahasan mendalam mengenai kinerja model tersebut. Analisis dimulai dari prapemrosesan data, perhitungan probabilitas manual, hingga evaluasi model dan analisis kontribusi variabel terhadap hasil prediksi.

4.1. Pra-pemrosesan Data Latih

Sebelum dilakukan pelatihan model, dataset sekunder yang bersumber dari *Kaggle* melewati tahap sterilisasi. Penelitian ini menggunakan total 8.000 observasi dengan sebaran kelas target yang relatif seimbang. Atribut kategorikal, misalnya Jenis Kelamin, ditransformasikan menjadi wujud numerik (*Female* : 0, *Male* : 1) guna memfasilitasi komputasi matematis.

Tabel 3. Sampel Data untuk Simulasi Perhitungan Manual

Variabel	Data Input	Rata-rata (μ) Kelas High	Standar Deviasi (σ) Kelas High	Rata-rata (μ) Kelas Low	Standar Deviasi (σ) Kelas Low
Umur	20	22.5	3.5	19.2	2.8
Screen time	9 jam	8.5	1.5	4.2	1.2
Durasi Tidur	5 jam	5.4	0.8	7.6	0.7
Interaksi Sosial	1	1.8	0.6	4.1	0.8

- Menghitung *Probabilitas Prior* berdasarkan distribusi kelas pada data latih 6.400 record :

$$a. P(High Risk) = \frac{4.512}{6.400} = 0,705$$

$$b. P(Low Risk) = \frac{1.888}{6.400} = 0,295$$

- Menghitung *likelihood* untuk setiap fitur numerik menggunakan fungsi kepadatan peluang *Gaussian*. Sebagai contoh, *likelihood* untuk fitur *Screen time*:

a. High Risk

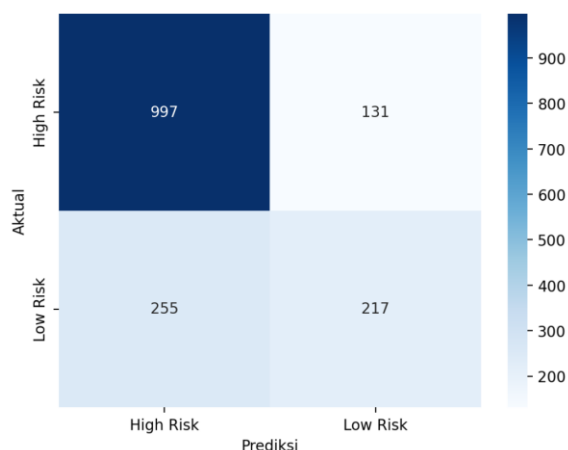
$$f(9; 8,5,1,5) = \frac{1}{\sqrt{2\pi(1,5)^2}} e^{\left(-\frac{(9-8,5)^2}{2(1,5)^2}\right)} = 0,251$$

b. Low Risk

$$f(9; 4,2,1,2) = \frac{1}{\sqrt{2\pi(1,2)^2}} e^{\left(-\frac{(9-4,2)^2}{2(1,2)^2}\right)} = 0,001$$

Hasil simulasi manual ini menunjukkan bahwa penggunaan fungsi *Gaussian PDF* efektif dalam memberikan bobot probabilitas yang kontras pada fitur-fitur yang memiliki perbedaan rata-rata signifikan antara kedua kelas. Hal ini berbanding lurus dengan hasil evaluasi model yang menunjukkan nilai F1-Score sebesar 0,84 untuk kategori *High Risk*.

4.3. Hasil Pelatihan dan Evaluasi Model



Gambar 4. Confusion Matrix Hasil Pengujian Model

4.2. Perhitungan Probabilitas Manual (Simulasi Teorema Bayes)

Guna memvalidasi logika kerja algoritma *Gaussian Naïve Bayes*, dilakukan simulasi perhitungan probabilitas manual menggunakan rumus *Teorema Bayes*. Simulasi ini penting untuk membuktikan asumsi distribusi normal pada fitur kontinu.

Pelatihan model dilaksanakan menggunakan bahasa pemrograman *Python* pada platform VS Code. Model dilatih memakai 80% data latih (6.400 observasi) serta diuji memakai 20% data uji (1.600 observasi). Performa model dievaluasi menggunakan metrik akurasi dari *Confusion Matrix*.

Berlandaskan analisis *Confusion Matrix* dalam Gambar 4, didapatkan metrik performa sebagai berikut :

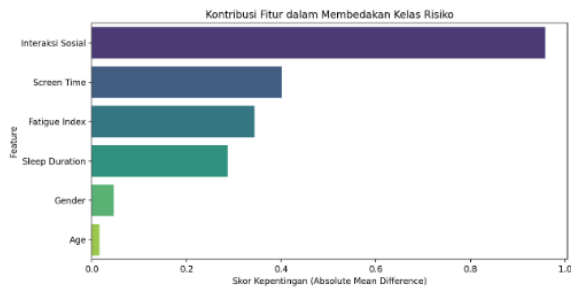
Tabel 4. Metrik Evaluasi Performa Model

Metrik	Nilai	Deskripsi
Akurasi	75,88%	Kemampuan sistem mengklasifikasikan data uji (<i>High Risk & Low Risk</i>) secara benar.
Presisi (<i>Precision</i>)	79,63%	Keakuratan model dalam mengidentifikasi kelas <i>High Risk</i> secara tepat.
<i>Recall</i>	88,38%	Sensitivitas model dalam mendeteksi sampel <i>High Risk</i> yang sesungguhnya.
<i>F1-Score</i>	83,78%	Rata-rata harmonis antara presisi dan <i>recall</i> .

Akurasi sebesar 75,88% menunjukkan tingkat reliabilitas yang memadai dalam klasifikasi risiko kesehatan mental remaja, mengingat kompleksitas variabel psikologis yang bersifat *multifaktorial*. Poin keunggulan utama dari model ini terletak pada nilai *Recall* sebesar 88,38%, yang mengindikasikan sensitivitas tinggi sistem dalam mendeteksi individu yang benar-benar berada pada kategori *High Risk*.

4.4. Analisis Variabel Prediktor (*Feature Importance*)

Penelitian ini menganalisis kontribusi relatif setiap variabel prediktor (fitur) terhadap hasil klasifikasi risiko depresi menggunakan *Teorema Bayes*.



Gambar 5. Feature Importance

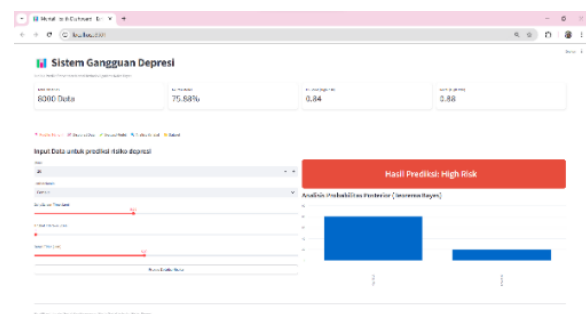
Dari Gambar 5 menjelaskan mengenai faktor risiko utama pada remaja :

- Durasi *Screen time* (Faktor Dominan) Variabel ini menempati posisi teratas dalam hierarki kontribusi fitur. Hal ini mengindikasikan bahwa dalam algoritma *Gaussian Naïve Bayes*, distribusi probabilitas pada kelas *High Risk* sangat dipengaruhi oleh intensitas penggunaan media sosial. Secara teknis, rata-rata (?) penggunaan layar yang tinggi pada data latih menjadi pembeda utama yang signifikan dibandingkan dengan kelompok *Low Risk*.
- Durasi Tidur (Pola Istirahat) merupakan variabel yang paling berpengaruh. Dalam model ini, pola tidur yang tidak teratur ataupun kurang dari standar kesehatan memiliki likelihood yang besar untuk diklasifikasikan ke dalam kategori risiko tinggi. Sinergi antara *Screen time* yang tinggi dan Durasi Tidur yang rendah menjadi pola yang paling sering dikenali oleh model dalam mendeteksi potensi depresi.
- Interaksi Sosial (Dinamika Relasi Nyata) Variabel ini memberikan kontribusi yang signifikan dalam menyeimbangkan data perilaku digital. Dalam algoritma *Gaussian Naïve Bayes*, tingkat interaksi sosial yang rendah (skor rendah) memiliki nilai likelihood yang tinggi untuk berasosiasi dengan kelas *High Risk*. Hal ini membuktikan asumsi penelitian bahwa penurunan kualitas interaksi sosial secara langsung di dunia nyata merupakan indikator kuat dalam mendeteksi risiko depresi pada remaja. Integrasi variabel ini memperkuat model agar tidak hanya terfokus pada durasi penggunaan gawai, tetapi juga mempertimbangkan aspek psikososial responden.
- Jenis Kelamin (Variabel Demografis) Variabel kategorikal ini memberikan kontribusi moderate. Melalui proses *Label Encoding*, model mampu memetakan perbedaan probabilitas prior antara laki-laki dan perempuan. Meskipun kontribusinya tidak sebesar faktor perilaku digital, variabel ini tetap esensial dalam mempertajam akurasi prediksi berdasarkan tren gender pada dataset 8.000 records tersebut.
- Umur (Variabel Pelengkap) Umur memiliki bobot kontribusi yang paling rendah disandingkan variabel lainnya. Memperlihatkan dalam cakupan data remaja ini, faktor usia tidak memiliki variasi yang cukup ekstrem untuk mengubah hasil prediksi

secara drastis. Dengan kata lain, risiko depresi pada sistem ini lebih ditentukan oleh gaya hidup digital (*screen time* dan tidur) daripada sekadar faktor usia biologis.

4.5. Implementasi Antarmuka Sistem

Sebagai manifestasi hilir pragmatis dari kerangka kerja komputasi *Bayesian* yang diverifikasi, antarmuka interaktif (*dashboard*) dibangun menggunakan arsitektur *Streamlit* [10]. Antarmuka dibuat dengan cermat untuk memfasilitasi pengguna akhir (remaja, orang tua, atau konselor) dalam melakukan simulasi risiko kesehatan mental secara mandiri.

Gambar 6. Dashboard *Streamlit*

Pemeriksaan komprehensif penyebaran antarmuka yang digambarkan pada Gambar 6 menjelaskan atribut operasional berikut :

- Panel Dashboard Utama : Bagian ini menyajikan metrik performa model untuk menjamin validitas riset, mencakup penggunaan dataset sebanyak 8.000 sampel, capaian Akurasi 75,88%, serta Recall 88%. Visualisasi ini menegaskan keandalan algoritma *Gaussian Naïve Bayes* dalam mengidentifikasi risiko secara akurat.
- Modul Input Data : Formulir dinamis di sisi kiri memfasilitasi pengguna untuk menginput data personal yang terdiri dari Usia, Jenis Kelamin, Durasi Waktu Layar, Durasi Tidur, dan Interaksi Sosial. Struktur formulir ini dirancang secara sistematis guna memastikan alur pengisian data berlangsung sekuensial dan minim ambiguitas interpretatif dari pengguna. Penggunaan kontrol slider memberikan pengalaman pengguna (*User Experience*) yang lebih modern dan praktis dalam memasukkan data.
- Output Analisis: Setelah aktivasi tombol “Proses Deteksi Risiko”, sistem melakukan kalkulasi *probabilitas posterior* secara otomatis berdasarkan *Teorema Bayes*. Proses komputasional ini berlangsung secara deterministik dengan memanfaatkan parameter yang telah dipelajari pada fase pelatihan model sebelumnya. Hasil akhir ditampilkan melalui dua format: label status risiko (Tinggi/Rendah) dan grafik batang probabilitas, guna memudahkan interpretasi data secara cepat dan objektif.

5. KESIMPULAN DAN SARAN

Studi ini membuktikan efektivitas algoritma *Gaussian Naive Bayes* sebagai instrumen skrining risiko depresi remaja dengan capaian akurasi 75,88% dan *recall* 88,38%, yang secara signifikan mampu meminimalisir risiko penderita tidak terdeteksi (*false negative*). Temuan ini sekaligus mengindikasikan bahwa pendekatan komputasional berbasis probabilitas memiliki kapabilitas yang memadai dalam memetakan kerentanan psikologis secara dini, bahkan dalam lanskap data yang bersifat heterogen. Berdasarkan analisis pada 8.000 observasi, ditemukan bahwa rendahnya interaksi sosial serta tingginya kelelahan digital akibat ketimpangan durasi layar dan waktu tidur merupakan prediktor risiko utama yang kini telah diimplementasikan secara praktis melalui *dashboard* Streamlit. Korelasi tersebut memperlihatkan adanya dinamika relasional antara perilaku digital dan kondisi afektif remaja yang kian kompleks dalam ekosistem kehidupan modern. Untuk pengembangan ke depan, disarankan pengintegrasian data secara *real-time* dan penambahan variabel psikososial yang lebih komprehensif guna meningkatkan presisi model, sehingga model tidak hanya bersifat deskriptif, melainkan juga mampu berfungsi sebagai instrumen prediktif yang adaptif terhadap perubahan temporal, sementara pihak terkait diimbau memanfaatkan sistem ini sebagai langkah mitigasi preventif kesehatan mental di era digital. Dengan demikian, kontribusi penelitian ini tidak hanya terletak pada aspek teknis, tetapi juga pada implikasi praktisnya dalam mendukung intervensi kebijakan berbasis data.

DAFTAR PUSTAKA

- [1] F. A. Sumantri, Y. H. Chrisnanto, and M. Melina, "Prediksi Risiko Kesehatan Mental Mahasiswa Menggunakan Klasifikasi Naive Bayes," *JURIKOM (Jurnal Riset Komputer)*, vol. 12, no. 3, pp. 383–393, Jun. 2025, doi: 10.30865/jurikom.v12i3.8648.
- [2] P. Elisa and A. R. Isnain, "Comparison Of Random Forest, Support Vector Machine And Naive Bayes Algorithms To Analyze Sentiment Towards Mental Health Stigma," *Jurnal Teknik Informatika (Jutif)*, vol. 5, no. 1, pp. 2024, doi: 10.52436/1.jutif.2024.5.1.1817.
- [3] S. Shevira, I. M. A. D. Suarjaya, and P. W. Buana, "Lexicon and Naive Bayes Algorithms to Detect Mental Health Situations from Twitter Data," *Journal of Information Systems Engineering and Business Intelligence*, vol. 8, no. 2, pp. 142–148, Oct. 2022, doi: 10.20473/jisebi.8.2.142-148.
- [4] Febria Sera Darnefi and M. Riko Anshori Prasetya, "Naïve Bayes Approach to Understanding Students' Psychological Well-Being," *Install: Information System and Technology Journal*, vol. 1, no. 1, pp. 27–34, Jun. 2024, doi: 10.33859/install.v1i1.547.
- [5] A. Sazwati, D. Pratama, K. Anam, E. Wahyudin, and A. Rifa'i, "Penerapan Data Mining Untuk Mengestimasi Persentase Penduduk Miskin Di Jawa Barat Menggunakan Regresi Linier Berganda," *JATI (Jurnal Mahasiswa Teknik Informatika)*, vol. 8, no. 1, pp. 1103–1108, Mar. 2024, doi: 10.36040/jati.v8i1.8214.
- [6] R. Alfarezy, E. Ermatita, and R. M. B. Wadu, "Implementasi Algoritma Naïve Bayes Untuk Analisis Klasifikasi Survei Kesehatan Mental (Studi Kasus: Open Sourcing Mental Illness)," *Informatik : Jurnal Ilmu Komputer*, vol. 19, no. 1, pp. 1–10, May 2023, doi: 10.52958/iftk.v19i1.4696.
- [7] A. I. Maulana and R. Rismayani, "Mental Health Analysis at the University of Dipa Makassar using Naïve Bayes Classifier," *IT Journal Research and Development*, vol. 8, no. 1, pp. 72–80, Dec. 2023, doi: 10.25299/itjrd.2023.11444.
- [8] D. Rizqa Khalaida, S. Shofia, T. Tukino, and E. Novalia, "Klasifikasi Karakteristik Pengguna Mediasosial Menggunakan Metode Naive Bayes," *JATI (Jurnal Mahasiswa Teknik Informatika)*, vol. 9, no. 5, pp. 8906–8912, Jul. 2025, doi: 10.36040/jati.v9i5.15225.
- [9] B. M. Kono, "social_media_mental_health.csv," *Kaggle*. Accessed: Apr. 21, 2026. [Online]. Available: <https://www.kaggle.com/datasets/bertnardo/mariousskono/social-media-and-mental-health>
- [10] K. N. Oktaviarini, E. D. Wahyuni, and R. P. Sari, "Transformasi Data Statistik Menjadi Visual Interaktif Menggunakan Streamlit: Studi Kasus BPS Kota Mojokerto," *JURNAL KRIDATAMA SAINS DAN TEKNOLOGI*, vol. 6, no. 02, pp. 804–822, Dec. 2024, doi:10.53863/kst.v6i02.1426.