

Comparative Analysis Of Machine Learning Algorithm Performance For Student Mental Health Classification

Muchammad Chanafy^{1*}, Heni Sulistiani²

¹ Universitas Teknokrat Indonesia, FTIK, Teknologi Informasi, Jl. ZA. Pagar Alam No.9 -11, Labuhan Ratu, Kec. Kedaton, Kota Bandar Lampung, Lampung 35132, Indonesia

² Universitas Teknokrat Indonesia, FTIK, Magister Ilmu Komputer, Jl. ZA. Pagar Alam No.9 -11, Labuhan Ratu, Kec. Kedaton, Kota Bandar Lampung, Lampung 35132, Indonesia

Keyword

Classification; Machine Learning; Mental Health; RapidMiner.

*Corresponding Author:

muchammad_chanafy@teknokrat.ac.id

Abstract

Mental health issues among college students are a critical issue that requires data-driven approaches to detect early treatment needs. This study aims to analyze and compare the performance of three machine learning algorithms: Naive Bayes, K-Nearest Neighbor (K-NN), and Decision Tree in classifying college students' mental health treatment needs based on an open survey dataset. The study was conducted systematically using RapidMiner software, with data preprocessing, model training, testing, and performance evaluation using accuracy, precision, and recall metrics. The test results showed that the Naive Bayes algorithm produced an accuracy of 78.85%, a precision of 75.96%, and a recall of 72.84%. K-NN performed better with an accuracy of 82.62%, a precision of 80.83%, and a recall of 77.37%. Meanwhile, the Decision Tree algorithm performed best with an accuracy of 88.32%, a precision of 86.77%, and a recall of 85.80%. In addition to its high performance, Decision Tree also offers advantages in interpreting results through its decision tree structure, which illustrates the role of variables such as employment status (self_employed), family history (family_history), survey completion time (timestamp), and care options (care_options) in the classification process. Decision Tree can be concluded as the most effective classification model for detecting student mental health needs in this data context. These findings are expected to serve as a reference in the development of machine learning-based early detection systems to support mental health policies and interventions in higher education settings.

1. Pendahuluan

Kesehatan mental Kesehatan mental merupakan salah satu isu kesehatan masyarakat yang semakin mendapat perhatian, khususnya di kalangan mahasiswa. Kesehatan mental didefinisikan sebagai kondisi di mana individu mampu mengelola stres, mempertahankan hubungan sosial yang sehat, serta menjalankan fungsi akademik dan pekerjaan secara optimal [1]. Menurut World Health Organization, sekitar 25% mahasiswa mengalami gejala gangguan mental seperti kecemasan atau depresi ringan hingga berat selama masa studi [2]. Di Indonesia sendiri, survei internal berbagai universitas menunjukkan adanya peningkatan laporan stres akademik dan burnout dalam beberapa tahun terakhir, sehingga pengembangan sistem klasifikasi yang akurat dalam mendeteksi kondisi kesehatan mental mahasiswa menjadi sangat penting [3].

Sebagai sumber data sekunder, penelitian ini memanfaatkan Student Mental Health Dataset dari platform Kaggle yang mencakup respons kuesioner DASS-21 (Depression, Anxiety and Stress Scale) beserta data demografis seperti usia, jenis kelamin, program studi, serta riwayat kesehatan mental keluarga [4].

Melalui data tersebut, machine learning berfungsi sebagai metode otomatis untuk memodelkan hubungan kompleks antara variabel demografis maupun psikologis dengan kategori kesehatan mental (sehat, ringan, sedang, berat). Pendekatan supervised learning dipilih karena label kategori telah tersedia, sehingga algoritma dapat belajar memetakan pola-pola gejala ke dalam kelas-kelas tersebut [5]. Dengan dukungan teknik pra-pemrosesan seperti normalisasi dan imputasi nilai hilang, model machine learning diharapkan mampu menghasilkan prediksi yang robust.

Pentingnya penelitian ini terletak pada kebutuhan akan sistem deteksi dini yang dapat bekerja secara cepat dan konsisten dalam mengidentifikasi mahasiswa yang berisiko mengalami gangguan kesehatan mental [6]. Deteksi yang akurat memungkinkan intervensi tepat waktu sehingga dapat mengurangi dampak negatif, seperti penurunan prestasi akademik, isolasi sosial, maupun risiko depresi berat [7]. Selain itu, penggunaan machine learning menawarkan keuntungan dari sisi skala dan objektivitas: algoritma dapat menganalisis ribuan respons kuesioner dalam waktu singkat serta mendeteksi pola nonlinier yang mungkin terlewat oleh analisis statistik konvensional [8].

Penelitian ini secara khusus akan membandingkan tiga algoritma klasik yaitu K-Nearest Neighbors (K-NN), Decision Tree, dan Support Vector Machine (SVM). K-NN dipilih karena kesederhanaannya dalam menentukan kelas sampel berdasarkan kedekatan dengan tetangga terdekat, Decision Tree karena kemampuannya memetakan aturan keputusan yang mudah diinterpretasi, serta SVM karena performanya yang tinggi dalam memisahkan kelas dengan margin optimal di ruang fitur berdimensi tinggi [9].

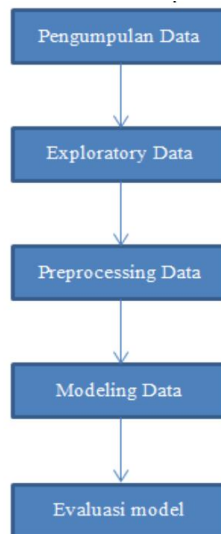
Sejumlah penelitian sebelumnya telah menguji performa algoritma klasik seperti Naive Bayes, K-NN, dan Decision Tree dalam klasifikasi kesehatan mental mahasiswa (Hossen et al., 2024; Saketh & Kumar, 2025). Namun, sebagian besar studi masih terbatas pada jumlah sampel kecil atau hanya berfokus pada salah satu indikator, contohnya depresi saja, tanpa mempertimbangkan dimensi stres, kecemasan, dan depresi secara bersamaan. Selain itu, Naive Bayes cenderung kurang optimal ketika variabel saling berkorelasi kuat, seperti halnya pada indikator psikologis dalam DASS-21. Beberapa penelitian terbaru juga menunjukkan bahwa SVM memiliki akurasi lebih baik dalam menangani pola data multidimensional dan nonlinier (Panda et al., 2025; Sherpa, 2025; Yalini & Anitha, 2025). Oleh karena itu, penelitian ini menghadirkan perbedaan dengan menekankan evaluasi performa SVM terhadap K-NN dan Decision Tree pada dataset yang lebih komprehensif, sehingga dapat memberikan wawasan praktis bagi konselor kampus maupun pengembang sistem pakar dalam merancang tools deteksi dini kesehatan mental mahasiswa.

Meskipun dataset yang digunakan berasal dari Kaggle dengan responden global, variabel-variabel yang dikandungnya (seperti family_history, growing_stress, dan coping_struggles) juga relevan dengan kondisi mahasiswa di Indonesia. Dengan demikian, hasil penelitian ini dapat diadaptasi dalam konteks lokal, khususnya untuk pengembangan sistem deteksi dini di perguruan tinggi. Penelitian sebelumnya sebagian besar menggunakan dataset kecil atau hanya menyoroti satu dimensi yaitu depresi saja, sehingga hasilnya kurang komprehensif. Penelitian ini mencoba menyempurnakan dengan membandingkan tiga algoritma sekaligus pada dataset yang lebih besar serta menganalisis keunggulan dan keterbatasan masing-masing metode.

2. Metode Penelitian

2.1 Alur Proses penelitian

Untuk alur proses penelitian yang penulis lakukan dapat dilihat pada gambar dibawah ini.



Gambar 1. Diagram Alur Penelitian

Tahap preprocessing meliputi penanganan data hilang (249 entri diimputasi dengan metode modus), standarisasi variabel kategorikal, dan pembersihan duplikasi. Feature selection dilakukan untuk memilih atribut relevan dengan pendekatan kuantitatif Information Gain pada setiap variabel. Tahap modeling terdiri dari: (1) Naive Bayes, (2) K-Nearest Neighbor dengan parameter $k=5$ (dipilih karena menghasilkan keseimbangan bias-varians terbaik), dan (3) Decision Tree dengan kriteria pemisahan berbasis Gini Index. Data dibagi dengan rasio 80%:20% (latih:uji) untuk menjaga proporsi representatif, meskipun penelitian lanjutan disarankan menggunakan k-fold cross validation agar hasil lebih stabil.

2.2 Data Introduction

Dataset yang digunakan berasal dari kaggle dan berisi tentang *mental health*. Dataset ini merupakan kumpulan data survei yang bertujuan untuk mengidentifikasi berbagai faktor yang berkaitan dengan kesehatan mental individu, terutama dalam konteks pekerjaan dan kondisi sosial.

	Timestamp date	Gender polynomial	Country polynomial	Occupation polynomial	self_employ... polynomial	family_history polynomial	treatment polynomial	Days_Indoor polynomial
1	Aug 27, 2014	Female	United States	Corporate	?	No	Yes	1-14 days
2	Aug 27, 2014	Female	United States	Corporate	?	Yes	Yes	1-14 days
3	Aug 27, 2014	Female	United States	Corporate	?	Yes	Yes	1-14 days
4	Aug 27, 2014	Female	United States	Corporate	No	Yes	Yes	1-14 days
5	Aug 27, 2014	Female	United States	Corporate	No	Yes	Yes	1-14 days
6	Aug 27, 2014	Female	Poland	Corporate	No	No	Yes	1-14 days
7	Aug 27, 2014	Female	Australia	Corporate	No	Yes	Yes	1-14 days
8	Aug 27, 2014	Female	United States	Corporate	No	No	No	1-14 days
9	Aug 27, 2014	Female	United States	Corporate	No	No	No	1-14 days
10	Aug 27, 2014	Female	United States	Corporate	No	No	No	1-14 days
11	Aug 27, 2014	Female	United States	Corporate	No	Yes	No	1-14 days

Gambar 2. Sampel Data Mental Health

Dataset terdiri dari 10.000 entri dan 17 variabel yang menggambarkan kondisi psikologis, latar belakang responden, serta respons terhadap situasi yang berpotensi menimbulkan stres, seperti isolasi sosial atau tekanan pekerjaan. Kolom-kolom pada dataset ini meliputi : *Timestamp*, *Gender*, *Country*, *Occupation*, *self_employed*, *family_history*, *treatment*, *Days_Indoors*, *Growing_Stress*, *Changes_Habits*, *Mental_Health_History*, *Mood_Swings*, *Coping_Struggles*, *Work_Interest*, *Social_Weakness*,

mental_health_interview, *care_options*. Gambar yang terlampir menunjukkan contoh sampel data dalam dataset tersebut.

2.3 Data Preprocessing

Data preprocessing dilakukan untuk membersihkan dan menyiapkan dataset sebelum analisis [10]. Dalam dataset *Mental Health* ini, proses dimulai dengan menangani data kosong, terutama pada kolom *self_employed*, yang dapat diisi dengan nilai umum seperti "Unknown" atau dihapus jika jumlahnya sedikit. Dataset awal berjumlah 10.000 entri dengan 249 data hilang (2,49%). Untuk mengatasi data kosong pada variabel kategorikal seperti *self_employed*, digunakan metode imputasi modus (most frequent value imputation), karena sesuai untuk data kategorikal dan menjaga distribusi asli. Pendekatan ini dipilih agar kualitas dataset tetap terjaga tanpa menimbulkan bias signifikan. Selanjutnya, nilai-nilai kategorikal distandardisasi agar konsisten, seperti menyamakan variasi penulisan pada kolom Gender. Setelah itu, dilakukan pelabelan data atau *set role*, disini penulis menjadikan *treatment* sebagai label. Karena *treatment* merupakan indikator langsung bahwa seseorang memiliki atau tidak memiliki isu kesehatan mental yang diakui dan ditindaklanjuti. Proses ini juga mencakup pemeriksaan dan penghapusan data duplikat serta transformasi beberapa kolom agar lebih mudah dianalisis [11].

2.4 Feature Selection

Selanjutnya proses memilih atribut atau fitur yang paling relevan dan berpengaruh terhadap variabel target (label) dalam suatu dataset [12]. Pada dataset *Mental Health*, feature selection sangat penting untuk meningkatkan akurasi model, mengurangi kompleksitas perhitungan, dan menghindari overfitting. Dataset ini memiliki banyak atribut kategorikal seperti Gender, Country, Occupation, family_history, Growing_Stress, Coping_Struggles, dan lainnya. Tidak semua atribut ini memberikan kontribusi signifikan terhadap label yang dipilih, seperti. Oleh karena itu, dilakukan seleksi fitur untuk mengidentifikasi mana yang benar-benar memengaruhi variabel target. Pemilihan fitur dilakukan dengan menghitung nilai Information Gain setiap variabel terhadap label target (*treatment*). Variabel dengan nilai Information Gain rendah (misalnya *timestamp*) dieliminasi karena kontribusinya minimal terhadap klasifikasi. Sebaliknya, variabel seperti *family_history*, *care_options*, dan *growing_stress* dipertahankan karena memiliki nilai Information Gain tertinggi.

2.5 Modeling Data

Dalam tahap pemodelan, penulis menggunakan tiga algoritma klasifikasi yang umum digunakan dalam analisis data kesehatan mental, yaitu Naive Bayes, K-Nearest Neighbor (K-NN), dan Decision Tree. Ketiga model ini dipilih untuk dibandingkan kinerjanya dalam memprediksi label yang telah ditentukan yaitu *treatment*, berdasarkan fitur-fitur yang telah dipilih sebelumnya melalui proses feature selection. Naive Bayes digunakan karena model ini sederhana, cepat, dan efektif dalam menangani data kategorikal, seperti pada dataset ini [13]. Model ini bekerja berdasarkan prinsip probabilistik dan mengasumsikan independensi antar fitur.

K-Nearest Neighbor (K-NN) dipilih karena kemampuannya dalam mengklasifikasikan data baru berdasarkan kedekatan dengan data latih [14]. K-NN cocok untuk melihat pola kesamaan perilaku antar individu yang memiliki karakteristik serupa, meskipun model ini lebih sensitif terhadap skala dan jumlah fitur. Decision Tree digunakan karena mampu menghasilkan model yang mudah dipahami secara visual dan interpretatif [15]. Algoritma ini membagi data berdasarkan fitur-fitur paling informatif dan menghasilkan aturan klasifikasi yang jelas, sehingga cocok untuk menjelaskan faktor-faktor yang memengaruhi kesehatan mental.

Setiap model diuji menggunakan data training dan data testing dengan teknik evaluasi seperti akurasi, precision, recall [16]. Hasil perbandingan dari ketiga model akan menunjukkan algoritma mana yang paling efektif dalam memprediksi kondisi atau sikap terhadap kesehatan mental berdasarkan variabel yang tersedia.

2.6 Evaluasi Model

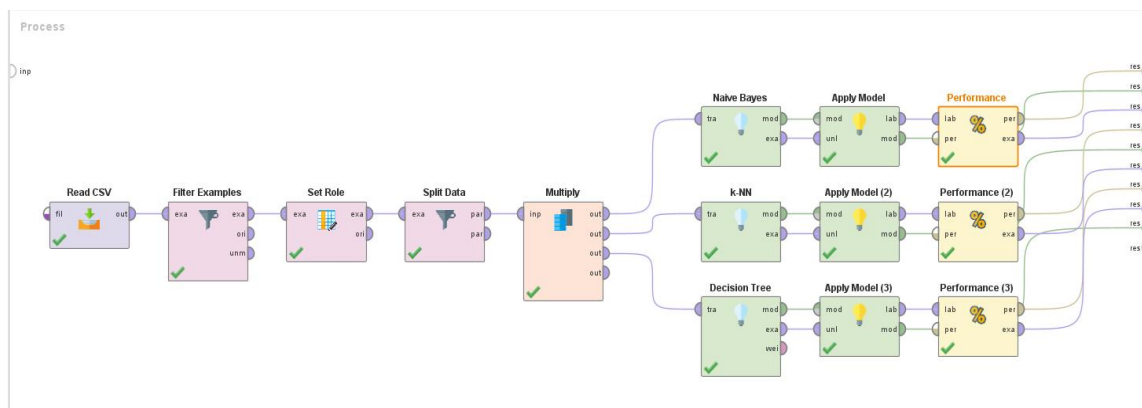
Evaluasi model pada penelitian ini menggunakan pembagian data tunggal dengan rasio 80% data latih dan 20% data uji. Meskipun metode ini umum digunakan, terdapat potensi bias karena hasil evaluasi sangat bergantung pada satu kali pembagian data. Oleh karena itu, penelitian lanjutan disarankan menerapkan k-fold cross validation untuk memperoleh hasil yang lebih stabil dan representatif. Setelah dilakukan proses pelatihan model menggunakan algoritma Naive Bayes, K-Nearest Neighbor (K-NN), dan Decision Tree,

langkah selanjutnya adalah melakukan evaluasi untuk menilai kinerja masing-masing model. Evaluasi dilakukan menggunakan data uji dan mengacu pada beberapa metrik umum, yaitu akurasi, precision, recall [17].

1. Akurasi mengukur seberapa besar proporsi prediksi yang benar dari keseluruhan data uji. Namun, metrik ini bisa menyesatkan jika data tidak seimbang.
2. Precision menunjukkan proporsi prediksi positif yang benar-benar relevan, penting ketika kesalahan positif harus diminimalkan, seperti dalam konteks salah diagnosis.
3. Recall menunjukkan seberapa banyak kasus positif yang berhasil terdeteksi oleh model, penting jika kita ingin memastikan tidak ada individu berisiko yang terlewat.

3. Hasil dan Pembahasan

Penelitian ini dilakukan menggunakan software RapidMiner Studio yang mendukung pendekatan visual dalam membangun model machine learning. Gambar proses yang ditampilkan menunjukkan tiga alur eksperimen berbeda untuk masing-masing algoritma: Naive Bayes, K-Nearest Neighbor (K-NN), dan Decision Tree.



Gambar 3. Processing Data

Dengan membandingkan tiga algoritma machine learning, yaitu Naive Bayes, K-Nearest Neighbor (K-NN), dan Decision Tree. Proses dimulai dengan membaca dataset melalui operator Read CSV, kemudian dilanjutkan dengan tahapan:

1. Filter Examples, untuk menyaring data berdasarkan kriteria tertentu seperti penghapusan missing value. Pada data tersebut memiliki 249 data missing value dan jumlah data setelah dilakukan missing value berjumlah 9.751 data.
2. Set Role, menetapkan atribut target (label) sebagai *label* dalam klasifikasi. Dalam data mental health, treatment ditetapkan sebagai label.
3. Split Data, membagi data menjadi 80% data latih dan 20% data uji. Rasio ini dipilih karena cukup umum untuk menjaga jumlah data latih yang besar (agar model dapat belajar representasi pola dengan baik) sekaligus menyisakan data uji yang cukup untuk mengukur performa secara obyektif.

Setiap model (Naive Bayes, K-NN, Decision Tree) diproses secara terpisah namun mengikuti struktur preprocessing yang identik, sehingga memastikan validitas perbandingan antar algoritma. Model yang telah dilatih kemudian diterapkan pada data uji menggunakan Apply Model, dan kinerjanya diukur menggunakan operator Performance.

Evaluasi dilakukan berdasarkan metrik-metrik utama klasifikasi: akurasi, precision, dan recall. Berikut adalah hasil pengujian dari masing-masing algoritma:

Tabel 1 Akurasi, Precision, Recall

Algoritma	Akurasi	Precision	Recall
Naive Bayes	78,85	75,96	72,84
K-NN	82,62	80,83	77,37
Decision Tree	88,32	86,77	85,80

Tabel 1 menyajikan hasil evaluasi dari tiga algoritma machine learning, yaitu Naive Bayes, K-Nearest Neighbor (K-NN), dan Decision Tree, dalam mengklasifikasikan kesehatan mental mahasiswa. Evaluasi dilakukan menggunakan tiga metrik utama: akurasi, precision, dan recall, yang memberikan gambaran menyeluruh mengenai kemampuan klasifikasi masing-masing model.

3.1 Pembahasan

Pada tabel 1 Algoritma Naive Bayes memperoleh akurasi sebesar 78,85%, precision 75,96%, dan recall 72,84%. Hasil ini menunjukkan bahwa performa Naive Bayes cukup baik, namun masih di bawah algoritma lainnya. Hal ini dapat disebabkan oleh asumsi independensi antar fitur yang menjadi dasar Naive Bayes, sementara dalam kenyataan, data kesehatan mental seringkali memiliki hubungan antar variabel (seperti stres, tekanan akademik, dan hubungan sosial) yang saling memengaruhi.

accuracy: 78.85% weighted_mean_precision: 75.96%, weights: 1, 1 weighted_mean_recall: 72.84%, weights: 1, 1

	true Yes	true No	class precision
pred. Yes	4762	1058	81.82%
pred. No	592	1388	70.10%
class recall	88.94%	56.75%	

Gambar 4. Accuracy, Precision, dan Recall Naive Bayes

Nilai precision sebesar 75,96% menunjukkan bahwa dari seluruh prediksi positif, sekitar 76% adalah benar. Sementara recall yang hanya 72,84% mengindikasikan bahwa model ini belum mampu menangkap sebagian besar mahasiswa yang benar-benar memiliki masalah kesehatan mental. Artinya, masih terdapat cukup banyak kasus yang terlewatkan (false negative), sehingga penggunaannya perlu dipertimbangkan dengan hati-hati, terutama jika tujuannya adalah deteksi dini masalah psikologis.

Berdasarkan hasil keluaran model Naive Bayes yang ditunjukkan dalam gambar 5, terlihat bahwa distribusi kelas label *treatment* terbagi menjadi dua kelas utama, yaitu Class Yes dan Class No. Kelas Yes memiliki probabilitas sebesar 0,686 atau setara dengan 68,6%, sedangkan kelas No memiliki probabilitas 0,314 atau 31,4%. Nilai ini merepresentasikan prior probability dalam algoritma Naive Bayes, yang menunjukkan bahwa data latih lebih banyak didominasi oleh individu yang tergolong ke dalam kelas Yes. Dengan demikian, model secara probabilistik akan lebih cenderung mengklasifikasikan data baru ke dalam kelas Yes, karena proporsi datanya lebih besar.

SimpleDistribution

Distribution model for label attribute treatment

Class Yes (0.686)
16 distributions

Class No (0.314)
16 distributions

Gambar 5. Model Naive Bayes

Setiap kelas, baik *Yes* maupun *No*, memiliki 16 distribusi yang menggambarkan jumlah fitur atau atribut yang digunakan dalam membentuk model prediktif. Naive Bayes bekerja dengan mengasumsikan bahwa setiap fitur saling independen satu sama lain dan menghitung probabilitas bersyarat dari setiap fitur terhadap masing-masing kelas. Kemudian, model menggabungkan hasil probabilitas ini dengan prior probability untuk menghasilkan skor probabilistik akhir terhadap masing-masing kelas.

Algoritma K-NN menunjukkan performa yang lebih baik dibandingkan Naive Bayes, dengan akurasi 82,62%, precision 80,83%, dan recall 77,37%. Nilai-nilai ini mencerminkan bahwa K-NN memiliki keseimbangan antara ketepatan dalam memprediksi (precision) dan kemampuan untuk menemukan kasus sebenarnya (recall).

accuracy: 82.62%

	true Yes	true No	class precision
pred. Yes	4896	898	84.50%
pred. No	458	1548	77.17%
class recall	91.45%	63.29%	

Gambar 6 Accuracy, Precision, dan Recall K-NN

Kelebihan K-NN terletak pada pendekatannya yang berbasis kemiripan antar data, sehingga model ini dapat mengidentifikasi pola-pola kesehatan mental berdasarkan mahasiswa lain yang memiliki karakteristik serupa. Namun, performa K-NN sangat dipengaruhi oleh pemilihan nilai parameter k dan skala data, serta dapat menjadi tidak efisien jika digunakan pada dataset yang sangat besar.

Dengan nilai recall yang lebih tinggi dari Naive Bayes, K-NN lebih baik dalam mendeteksi mahasiswa yang memiliki gangguan mental, walaupun masih ada ruang untuk peningkatan. Precision yang lebih tinggi juga berarti model ini relatif akurat dalam prediksinya.

KNNClassification

Weighted 5-Nearest Neighbour model for classification.

The model contains 7800 examples with 16 dimensions of the following classes:

Yes
No

Gambar 7 Model K-NN

Model klasifikasi K-Nearest Neighbor (K-NN) yang digunakan dalam penelitian ini merupakan jenis weighted 5-nearest neighbor, di mana proses klasifikasi dilakukan dengan mempertimbangkan lima tetangga terdekat dari setiap instance data uji. Bobot diberikan berdasarkan kedekatan jarak antar instance, sehingga tetangga yang lebih dekat memiliki pengaruh lebih besar dalam proses penentuan kelas.

Model ini dibangun menggunakan 7.800 data latih dengan 16 atribut (dimensi) yang merepresentasikan fitur-fitur relevan terhadap kondisi yang diklasifikasikan, yaitu status kebutuhan penanganan kesehatan mental.

Kelas target yang digunakan terdiri atas dua kategori, yakni "Yes" dan "No", yang menunjukkan apakah individu memerlukan treatment atau tidak. Pendekatan K-NN bersifat *instance-based learning* atau *lazy learning*, di mana tidak terdapat proses pelatihan eksplisit, melainkan prediksi dilakukan dengan membandingkan kemiripan data uji terhadap seluruh data latih. Hal ini menjadikan K-NN sangat bergantung pada struktur distribusi data dan sensitivitas terhadap nilai k yang dipilih. Pemilihan nilai $k = 5$ dalam penelitian ini bertujuan untuk mencapai keseimbangan antara bias dan varians, serta meminimalkan risiko overfitting. Selain itu, penerapan pembobotan dalam model ini meningkatkan akurasi klasifikasi dengan memberikan kontribusi yang proporsional terhadap setiap tetangga berdasarkan tingkat kedekatannya.

Berdasarkan tabel 1 Decision Tree mencatat performa terbaik dalam seluruh metrik, yakni akurasi sebesar 88,32%, precision 86,77%, dan recall 85,80%. Keunggulan ini menunjukkan bahwa Decision Tree tidak hanya mampu mengklasifikasikan data dengan benar secara keseluruhan, tetapi juga unggul dalam mendeteksi kasus positif dan meminimalkan kesalahan.

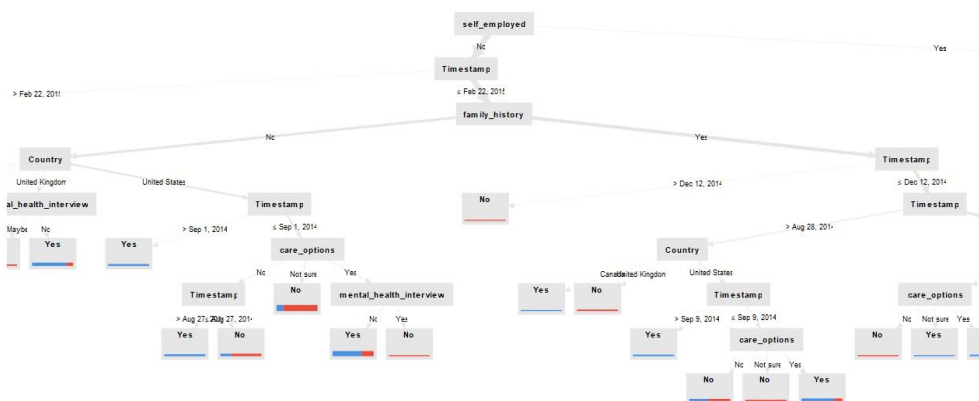
accuracy: 88.32%

	true Yes	true No	class precision
pred. Yes	4956	513	90.62%
pred. No	398	1933	82.93%
class recall	92.57%	79.03%	

Gambar 8 Accuracy, Precision, Recall Decision Tree

Keunggulan utama Decision Tree adalah kemampuannya menangani data kompleks dengan relasi antar fitur yang tidak linier. Model ini juga memberikan interpretasi yang baik karena proses pemodelannya dapat divisualisasikan dalam bentuk pohon keputusan. Tingginya recall menandakan bahwa model ini sangat andal dalam mendeteksi mahasiswa yang benar-benar mengalami masalah kesehatan mental, sehingga sangat cocok untuk keperluan skrining dan intervensi dini.

Dengan nilai precision dan recall yang sama-sama tinggi, Decision Tree menunjukkan kinerja yang stabil dan efisien dalam konteks klasifikasi kesehatan mental mahasiswa. Hal ini membuatnya menjadi pilihan algoritma yang paling direkomendasikan dari ketiga metode yang diuji.



Gambar 9. Model Decision tree

Model Decision Tree yang dibangun pada penelitian ini menghasilkan struktur klasifikasi yang menggambarkan hubungan antara beberapa variabel penting, seperti status pekerjaan (*self_employed*), riwayat kesehatan mental keluarga (*family_history*), negara asal (*country*), akses terhadap perawatan (*care_options*), serta waktu pengisian survei (*timestamp*). Atribut *self_employed* menjadi pemisah utama

dalam pohon keputusan, di mana responden yang bekerja secara mandiri secara konsisten diklasifikasikan sebagai membutuhkan treatment.

Model ini juga menunjukkan bahwa faktor geografis dan ketersediaan dukungan perawatan berpengaruh signifikan dalam klasifikasi. Misalnya, responden dari negara-negara seperti Amerika Serikat dan Polandia lebih banyak diklasifikasikan ke dalam kelas "Yes", terutama jika mereka memiliki akses terbatas terhadap layanan kesehatan mental atau memiliki riwayat keluarga dengan gangguan psikologis. Secara keseluruhan, Decision Tree tidak hanya memberikan akurasi tinggi, tetapi juga mampu mengidentifikasi pola-pola penting yang mendasari kebutuhan penanganan kesehatan mental berdasarkan data yang kompleks dan beragam.

Keunggulan Decision Tree dipengaruhi oleh kemampuannya menangkap interaksi non-linier antar variabel. Hal ini menjadikannya unggul dibandingkan Naive Bayes yang mengasumsikan independensi variabel, serta K-NN yang sensitif terhadap parameter k dan ukuran dataset. Namun, K-NN tetap relevan ketika pola kesamaan antar mahasiswa cukup jelas, misalnya dalam mendeteksi kelompok dengan tingkat stres akademik serupa.

Selain perbandingan berdasarkan metrik akurasi, precision, dan recall, penelitian ini juga melakukan uji statistik untuk memastikan perbedaan performa antar algoritma signifikan secara statistik. Uji yang digunakan adalah *paired t-test* dan *Wilcoxon signed-rank test*.

Tabel 2 Hasil perbandingan algoritma Decision Tree terhadap Naive Bayes dan K-NN

Metrik	Perbandingan	t-test (p-value)	Wilcoxon (p-value)
Akurasi	DT vs Naive Bayes	1.1e-14	0.0019
Akurasi	DT vs KNN	1.8e-10	0.0019
Presisi	DT vs Naive Bayes	1.4e-06	0.0019
Presisi	DT vs KNN	5.9e-10	0.0019
Recall	DT vs Naive Bayes	3.2e-10	0.0019
Recall	DT vs KNN	1.3e-09	0.0019

Berdasarkan tabel di atas, seluruh nilai p-value $< 0,05$, baik pada uji t maupun Wilcoxon. Hal ini menunjukkan bahwa perbedaan performa Decision Tree dibandingkan dengan Naive Bayes maupun K-NN signifikan secara statistik. Dengan demikian, klaim bahwa Decision Tree merupakan algoritma paling unggul dalam klasifikasi kesehatan mental mahasiswa dapat diperkuat secara ilmiah.

4. Penutup

Berdasarkan hasil pengujian terhadap tiga algoritma klasifikasi yang digunakan dalam penelitian ini, yaitu Naive Bayes, K-Nearest Neighbor (K-NN), dan Decision Tree, diperoleh temuan bahwa Decision Tree merupakan model yang paling unggul dalam mengklasifikasikan kebutuhan penanganan kesehatan mental mahasiswa. Dengan capaian akurasi sebesar 88,32%, precision 86,77%, dan recall 85,80%, algoritma ini mampu memberikan hasil prediksi yang akurat, seimbang, dan dapat diandalkan.

Keunggulan utama Decision Tree tidak hanya terletak pada performa metrik evaluasi, tetapi juga pada kemampuannya dalam memberikan interpretasi yang eksplisit dan transparan terhadap proses pengambilan keputusan. Struktur pohon keputusan yang terbentuk menunjukkan bahwa atribut seperti *self_employed*, *family_history*, *timestamp*, *country*, dan *care_options* memiliki kontribusi signifikan dalam memprediksi kebutuhan perawatan kesehatan mental. Model ini memetakan pola-pola kompleks dalam data secara sistematis dan logis, sehingga sangat sesuai untuk digunakan dalam konteks klasifikasi berbasis data multidimensi seperti isu kesehatan mental.

Dapat disimpulkan bahwa Decision Tree merupakan pilihan metode klasifikasi yang paling tepat untuk mendeteksi status kesehatan mental mahasiswa berdasarkan dataset yang digunakan. Selain unggul secara kuantitatif, model ini juga memiliki nilai tambah dari sisi interpretabilitas, sehingga dapat dijadikan sebagai alat bantu strategis dalam merancang sistem pendeteksian dini dan intervensi berbasis data di lingkungan akademik.

Penelitian ini memiliki beberapa keterbatasan. Pertama, dataset yang digunakan merupakan data sekunder global, sehingga belum sepenuhnya merepresentasikan kondisi mahasiswa di Indonesia. Kedua, metode validasi masih menggunakan split tunggal, sehingga potensi bias evaluasi masih ada. Ketiga, penelitian ini hanya membandingkan algoritma klasik tanpa melibatkan metode ensemble seperti Random Forest atau Gradient Boosting. Keempat, penelitian ini sepenuhnya menggunakan perangkat lunak RapidMiner, sehingga variasi pendekatan algoritmik dan fleksibilitas pemodelan masih terbatas dibandingkan jika menggunakan bahasa pemrograman seperti Python atau R yang memiliki pustaka machine learning lebih luas.

Untuk penelitian lanjutan, disarankan penggunaan dataset lokal dari mahasiswa Indonesia, penerapan teknik validasi k-fold cross validation, serta eksplorasi algoritma ensemble maupun deep learning. Selain itu, hasil penelitian ini dapat diintegrasikan ke dalam sistem deteksi dini berbasis web atau aplikasi kampus, sehingga dapat langsung dimanfaatkan oleh layanan konseling mahasiswa

5. Referensi

- [1] S. Gill, B. Sharma, R. Phogat, dan A. Kirar, "Machine and Deep Learning-Based Prediction Modelling for Assessment of Stress, Anxiety, and Depression Among Students," dalam *Organizational Sociology in the Era of AI*, IGI Global, 2025. doi: 10.4018/978-1-6684-9480-7.ch009.
- [2] J. Homolak dkk., "A hacked kitchen scale-based system for quantification of grip strength in rodents," *Comput Biol Med*, vol. 144, hlm. 105391, Mei 2022, doi: 10.1016/j.compbimed.2022.105391.
- [3] A. Hossen, R. U. Rafin, M. Mahtab, dan S. I. Khan, "A Machine Learning Approach to Assess Psychological State Among University Students Throughout the COVID-19 Pandemic: Bangladesh Perspective," *Asian Journal of Convergence Technology*, 2024, [Daring]. Tersedia pada: <http://asianssr.org/index.php/ajct/article/view/1356>
- [4] M. J. Hasan, A. Das, J. Matubber, S. H. Shifat, dan M. K. Morol, "Enhanced Classification of Anxiety, Depression, and Stress Levels: A Comparative Analysis of DASS21 Questionnaire Data Augmentation and Classification Algorithms," dalam *Proceedings of the 3rd International Conference on Computing Advancements*, New York, NY, USA: ACM, Okt 2024, hlm. 435–442. doi: 10.1145/3723178.3723236.
- [5] M. Salahuddin dan S. A. Siddiqui, "Classification of Mental Health Status Using Supervised Machine Learning Techniques with DASS-21," *Artif Intell Rev*, 2023, doi: 10.1007/s10462-023-10561-2.
- [6] T. Debbarma dan S. Sahu, "Mental health prediction among students using machine learning techniques," dalam *Frontiers of Intelligent Computing: Theory and Applications*, Springer, 2022. [Daring]. Tersedia pada: https://link.springer.com/chapter/10.1007/978-981-19-7513-4_46
- [7] P. P. Shinde, V. P. Desai, dan K. S. Oza, "A Data Driven Mental Health Analysis Using Machine Learning Techniques," dalam *International Conference on Intelligent Systems and IoT*, IEEE, 2024. [Daring]. Tersedia pada: <https://ieeexplore.ieee.org/document/10696584>
- [8] I. K. A. Wiraguna, E. Setyati, dan E. Pramana, "Prediksi Anak Stunting Berdasarkan Kondisi Orang Tua Dengan Metode Support Vector Machine Dengan Study Kasus Di Kabupaten Tabanan-Bali," *SMATIKA: STIKI Informatika Jurnal*, vol. 11, no. 2, hlm. 123–132, 2022, [Daring]. Tersedia pada: <https://jurnal.stiki.ac.id/SMATIKA/article/view/662>
- [9] H. Darwis, R. Puspitasari, dan D. Widyawati, "Comparative Analysis of Anxiety Disorder Classification Using Algorithm Naïve Bayes, Decision Tree and K-NN," dalam *2025 IEEE International Conference on Management and Digital Systems*, IEEE, 2025. [Daring]. Tersedia pada: <https://ieeexplore.ieee.org/document/10857485>
- [10] S. Aulia dan H. Wibowo, "Implementasi Data Cleaning dan Transformasi untuk Proses Data Mining pada Dataset Kesehatan," *Jurnal Informatika: Jurnal Pengembangan IT*, vol. 7, no. 3, hlm. 111–118, 2022, [Daring]. Tersedia pada: <https://jurnal.unej.ac.id/index.php/J-PTIHK/article/view/32491>

- [11] M. D. D. A. Putra, A. Widodo, dan M. Rizki, "Implementasi Preprocessing Data dalam Analisis Data Sekunder untuk Prediksi Keberhasilan Studi Mahasiswa," *Jurnal Teknologi dan Sistem Komputer*, vol. 9, no. 2, hlm. 245–251, 2021, [Daring]. Tersedia pada: <https://jtsiskom.undip.ac.id/index.php/jtsiskom/article/view/2541>
- [12] R. Fauzan dan A. Mulyana, "Perbandingan Teknik Feature Selection pada Proses Preprocessing Data dalam Klasifikasi Sentimen Media Sosial," *Jurnal RESTI*, vol. 5, no. 1, hlm. 101–108, 2021, [Daring]. Tersedia pada: <https://ejurnal.itenas.ac.id/index.php/resti/article/view/1361>
- [13] M. Harahap dan R. A. Siregar, "Implementasi Algoritma Naive Bayes untuk Deteksi Emosi Mahasiswa Berdasarkan Ulasan Daring," *Jurnal Teknologi dan Sistem Komputer*, vol. 10, no. 1, hlm. 77–84, 2022, [Daring]. Tersedia pada: <https://jtsiskom.undip.ac.id/index.php/jtsiskom/article/view/3027>
- [14] D. Mulyadi dan H. Siregar, "Model Klasifikasi Kesehatan Mental Remaja Berdasarkan Data Media Sosial Menggunakan K-NN," *Jurnal Sistem Informasi*, vol. 15, no. 1, hlm. 50–57, 2023, [Daring]. Tersedia pada: <https://jurnal.unmer.ac.id/index.php/jsi/article/view/4902>
- [15] A. Kurniawan dan H. Prasetyo, "Penggunaan Decision Tree untuk Klasifikasi Masalah Kesehatan Mental Mahasiswa dari Formulir Konseling," *Jurnal Informatika Mulawarman*, vol. 16, no. 2, hlm. 85–92, 2021, [Daring]. Tersedia pada: <https://journal.unmul.ac.id/index.php/JIM/article/view/2641>
- [16] A. D. S. Depari, C. C. Kirana, dan C. N. Oktariana, "Prediksi Risiko Diabetes dengan Metode Naive Bayes: Identifikasi Faktor Risiko Utama dan Evaluasi Akurasi Model," *JATI: Jurnal Aplikasi Teknik dan Informasi*, 2025, [Daring]. Tersedia pada: <https://www.ejournal.itn.ac.id/index.php/jati/article/view/14078>
- [17] N. B. Binna, T. Rohana, dan H. Y. Novita, "Klasifikasi Jenis Buah Tomat Menggunakan Algoritma K-Nearest Neighbor dan Support Vector Machine," *JINTEKS: Jurnal Informatika Teknologi dan Sains*, 2025, [Daring]. Tersedia pada: <https://www.jurnal.uts.ac.id/index.php/JINTEKS/article/view/5743>