

ANALISIS KESADARAN MAHASISWA TERHADAP PRIVASI DATA DENGAN MENGGUNAKAN METODE NAÏVE BAYES

Kaman Septia¹, Loade Thoriq Fhadila², Muhammad Irvan Syahril Hidayat³, Chesario Sukarno⁴, Iman Kasih Nazara⁵, Fachri Amsury⁶

¹⁻⁶ Program Studi Teknologi Informasi, Universitas Bina Sarana Informatika
 Jl. Kramat Raya No. 98, Senen - DKI Jakarta, Indonesia
 Email: ¹17230312@bsi.ac.id, ²17230754@bsi.ac.id, ³17230642@bsi.ac.id, ⁴17230700@bsi.ac.id,
⁵17231042@bsi.ac.id, ⁶fachri.fcy@nusamandiri.ac.id

ABSTRAK

Privasi data merupakan aspek penting dalam aktivitas digital, terutama bagi mahasiswa yang aktif menggunakan berbagai *platform* daring. Penelitian ini bertujuan menganalisis tingkat kesadaran privasi data mahasiswa menggunakan algoritma Naïve Bayes. Data primer dikumpulkan melalui kuesioner Google Form yang berisi 13 indikator kesadaran privasi dan disebarluaskan melalui media sosial dengan teknik *voluntary response sampling*. Sebanyak 56 mahasiswa berpartisipasi sebagai sampel penelitian. Pengolahan data mengikuti tahapan *Knowledge Discovery in Database* (KDD), meliputi seleksi data, pembersihan, transformasi, pemodelan, serta evaluasi. Transformasi dilakukan dengan menghitung skor total per responden dan mengelompokkan tingkat kesadaran ke dalam kategori “Tinggi” dan “Standar” menggunakan *cut-off* empiris untuk menjaga keseimbangan kelas. Analisis klasifikasi dilakukan menggunakan algoritma Naïve Bayes melalui aplikasi Orange Data Mining, dengan evaluasi menggunakan *Test and Score* serta *Confusion Matrix*. Hasil penelitian menunjukkan bahwa model mampu mengklasifikasikan tingkat kesadaran privasi dengan akurasi 91.1%, *precision* 92.6%, *recall* 91.1%, *F1-score* 91.5%, *AUC* 0.976, dan *MCC* 0.738. Temuan ini menunjukkan bahwa Naïve Bayes efektif dalam mengenali pola kesadaran privasi mahasiswa dan layak digunakan sebagai dasar pengembangan program edukasi privasi data di lingkungan perguruan tinggi.

Kata kunci: privasi data, mahasiswa, Naïve Bayes, klasifikasi, KDD

ABSTRACT

Data privacy is a critical aspect of digital activity, particularly for university students who frequently engage with online platforms. This study aims to analyze students' awareness of data privacy using the Naïve Bayes classification algorithm. Primary data were collected through a Google Form questionnaire consisting of 13 indicators of privacy awareness and distributed via social media using a voluntary response sampling technique. A total of 56 students participated in this study. Data processing followed the Knowledge Discovery in Database (KDD) stages, including data selection, cleaning, transformation, modeling, and evaluation. The transformation process involved calculating the total awareness score for each respondent and categorizing awareness levels into "High" and "Standard" using an empirical cut-off to maintain class balance. The Naïve Bayes algorithm was applied using the Orange Data Mining application, with performance evaluated through the Test and Score and Confusion Matrix tools. The results indicate that the model performed effectively, achieving an accuracy of 91.1%, precision of 92.6%, recall of 91.1%, F1-score of 91.5%, AUC of 0.976, and MCC of 0.738. These findings demonstrate that Naïve Bayes is suitable for analyzing student privacy awareness patterns and can serve as a foundation for designing educational interventions to improve privacy literacy in academic environments.

Keywords: data privacy, students, Naïve Bayes, classification, KDD

1. PENDAHULUAN

Data pribadi adalah segala informasi yang dapat digunakan untuk mengidentifikasi seseorang, seperti nama, alamat, nomor telepon, alamat *email*, informasi keuangan, dan bahkan kebiasaan saat menggunakan internet [1]. Kebocoran data dapat menyebabkan berbagai konsekuensi, seperti penyalahgunaan informasi pribadi, pencurian identitas, hingga kerugian finansial [2]. Perlindungan data pribadi adalah isu penting di era digital yang memerlukan perhatian serius dari semua pihak. Dengan meningkatkan kesadaran dan edukasi masyarakat, diharapkan dapat mengurangi risiko kebocoran data dan melindungi privasi individu [3].



Mahasiswa merupakan kelompok demografis yang berada pada fase transisi dari remaja menuju dewasa, di mana mereka mengalami berbagai tantangan akademik, sosial, dan emosional [4]. Dalam aktivitas sehari-hari, mahasiswa sering merasa nyaman menggunakan berbagai *platform* digital tanpa sepenuhnya menyadari risiko yang ada, seperti pencurian identitas, peretasan akun, atau penipuan *online*. Oleh karena itu, sangat penting untuk meningkatkan kesadaran mereka mengenai pentingnya melindungi informasi pribadi.

Salah satu pendekatan yang dapat digunakan untuk menganalisis kesadaran mahasiswa terhadap privasi data adalah metode klasifikasi. Klasifikasi adalah proses untuk menemukan model atau fungsi yang dapat menggambarkan dan membedakan kelas data atau konsep, dengan tujuan agar model tersebut dapat digunakan untuk memprediksi kelas yang belum diketahui dari suatu objek pengamatan [5]. Pada konteks *machine learning*, Naïve Bayes merupakan sebuah pengklasifikasian probabilistik sederhana yang menghitung sekumpulan probabilitas dengan menjumlahkan frekuensi dan kombinasi nilai dari dataset yang diberikan. Algoritma menggunakan teorema Bayes dan mengasumsikan semua atribut independen atau tidak saling ketergantungan yang diberikan oleh nilai pada variabel kelas [6]. Algoritma ini memiliki keuntungan karena hanya membutuhkan jumlah data pelatihan yang relatif kecil untuk mengestimasi parameter yang dibutuhkan dalam proses klasifikasi. Dalam banyak situasi dunia nyata yang kompleks, Naïve Bayes sering memberikan hasil yang lebih baik daripada yang diharapkan [7]. Naïve Bayes termasuk metode *supervised learning* yang diketahui mempunyai tingkat akurasi yang baik dengan perhitungan sederhana [8].

Terdapat beberapa penelitian sebelumnya yang menggunakan algoritma Naïve Bayes untuk menganalisis isu-isu terkait keamanan dan privasi data. Salah satu penelitian melaporkan bahwa model Naïve Bayes mampu mencapai akurasi sebesar 88,80% dengan *precision* 29,27% dan *recall* 30,77% dalam mengklasifikasikan sentimen masyarakat terhadap isu kebocoran data di media sosial [9]. Selain itu, penelitian lain juga membandingkan performa algoritma Support Vector Machine (SVM) dan Naïve Bayes dalam analisis sentimen. Pada data Twitter, SVM mampu menghasilkan performa klasifikasi yang cukup baik dengan nilai *precision* sebesar 80% dan *recall* 93%. Namun, Naïve Bayes menunjukkan kinerja yang lebih unggul dengan *precision* mencapai 97% dan *recall* 97%, sehingga memberikan hasil yang jauh lebih stabil dan akurat dalam proses klasifikasi teks [10]. Studi lain juga menemukan performa yang sangat tinggi dari Naïve Bayes dalam klasifikasi sentimen terkait kasus pembobolan data, dengan capaian akurasi 98.33%, *precision* 100.00%, dan *recall* 97.13% [11]. Temuan-temuan tersebut mengonfirmasi bahwa Naïve Bayes merupakan algoritma yang andal untuk analisis klasifikasi data berbasis teks maupun survei.

Meskipun demikian, sebagian besar penelitian tersebut lebih berfokus pada analisis sentimen media sosial atau kelompok masyarakat umum. Belum banyak penelitian yang secara khusus menganalisis tingkat kesadaran privasi data pada mahasiswa menggunakan data primer dari kuesioner terstruktur. Selain itu, riset sebelumnya lebih menekankan isu keamanan siber secara umum, bukan secara spesifik pada aspek kesadaran privasi data. Hal ini menunjukkan adanya celah penelitian (*research gap*) yang perlu diisi.

Berdasarkan latar belakang tersebut, penelitian ini dilakukan untuk menganalisis tingkat kesadaran mahasiswa terhadap privasi data dengan menggunakan metode Naïve Bayes. Berbeda dari penelitian sebelumnya, penelitian ini menggunakan data kuesioner yang disebarakan langsung kepada mahasiswa sebagai data primer. Data kemudian diolah menggunakan aplikasi Orange, yaitu *software data mining* merupakan aplikasi *open source* yang mampu membantu penelitian dalam menganalisa suatu data [12]. Pendekatan ini diharapkan dapat memberikan gambaran yang lebih spesifik, terukur, dan relevan mengenai sejauh mana mahasiswa memahami dan menyadari pentingnya perlindungan data pribadi dalam lingkungan akademik.

2. METODE PENELITIAN

Jenis Penelitian

Penelitian ini menggunakan pendekatan kuantitatif dengan metode *data mining*. Secara khusus, penelitian menerapkan algoritma Naïve Bayes untuk melakukan klasifikasi tingkat kesadaran mahasiswa terhadap privasi data. Pengolahan data mengikuti tahapan *Knowledge Discovery in Database* (KDD) yang mencakup proses seleksi data, pembersihan, transformasi, analisis dengan Naïve Bayes, serta evaluasi hasil.

Sumber dan Teknik Pengumpulan Data

Data penelitian merupakan data primer yang diperoleh melalui penyebaran kuesioner Google Form. Tautan kuesioner dibagikan melalui berbagai media sosial seperti WhatsApp, Instagram, dan Telegram dengan sasaran utama yaitu mahasiswa. Kuesioner berisi pernyataan mengenai pemahaman, sikap, dan perilaku mahasiswa terkait privasi data. Respons diberikan menggunakan skala Likert (Sangat Tidak Setuju – Sangat Setuju).



Gambar 1. Ilustrasi proses pengolahan data

Gambar 1 menunjukkan alur pengolahan data pada penelitian ini, dimulai dari pengisian kuesioner melalui Google Form, kemudian diekspor ke Excel untuk *pre-processing*, dan selanjutnya data digunakan dalam aplikasi Orange Data Mining untuk proses klasifikasi.

Populasi dan Sampel

Populasi penelitian adalah mahasiswa yang aktif menggunakan layanan digital. Sampel diperoleh melalui *voluntary response sampling*, yaitu mahasiswa yang bersedia mengisi kuesioner yang disebarakan melalui media sosial. Jumlah responden yang berpartisipasi adalah sebanyak 56 mahasiswa.

Metode *voluntary response sampling* memiliki keterbatasan karena responden yang berpartisipasi cenderung berasal dari kelompok yang lebih aktif secara digital, sehingga tidak sepenuhnya merepresentasikan seluruh populasi mahasiswa. Oleh sebab itu, generalisasi temuan penelitian perlu dilakukan secara hati-hati.

Variabel Penelitian

Variabel penelitian terdiri atas variabel input dan variabel *output*. Variabel input berupa 13 indikator kesadaran privasi yang diukur pada 56 responden melalui skala Likert. Sementara itu, variabel output merupakan kategori tingkat kesadaran yang terdiri dari dua kelas, yaitu Kesadaran Tinggi dan Kesadaran Standar. Penentuan kelas dilakukan berdasarkan total skor Likert setiap responden yang kemudian dikelompokkan secara kategorikal sesuai batas *cut-off* yang telah ditetapkan.

Uji Reliabilitas Instrumen

Sebelum digunakan dalam proses transformasi data, instrumen kuesioner diuji reliabilitasnya untuk memastikan konsistensi internal antar butir pernyataan. Uji reliabilitas dilakukan menggunakan metode Cronbach's Alpha terhadap 13 item indikator kesadaran privasi. Pada Tabel 1, hasil pengujian menunjukkan nilai alpha sebesar 0,856, yang berada di atas batas minimal 0,70, sehingga instrumen dapat dinyatakan reliabel dan layak digunakan dalam analisis.

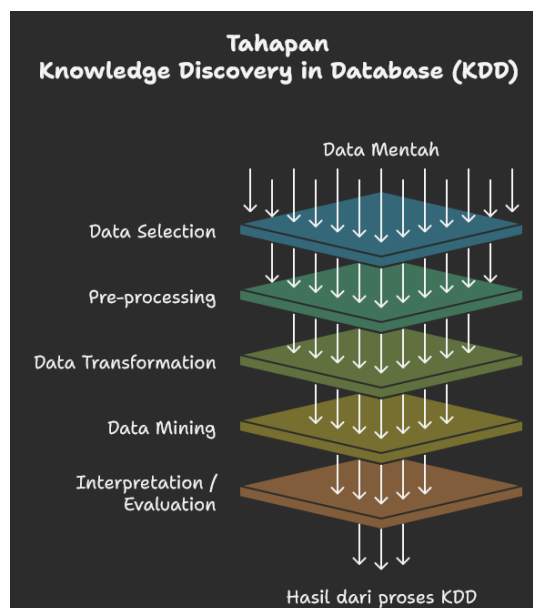
Tabel 1. uji reliabilitas kuesioner

Komponen	Nilai
Jumlah Item	13
Cronbach's Alpha	0,856

Pada Tabel 1, nilai tersebut mengindikasikan bahwa seluruh item memiliki tingkat konsistensi yang baik dan mampu mengukur konstruk kesadaran privasi secara stabil. Dengan demikian, data hasil kuesioner valid untuk dilanjutkan pada tahap transformasi dan klasifikasi menggunakan algoritma Naïve Bayes.

Tahapan Penelitian KDD (*Knowledge Discovery in Database*)

KDD merupakan sebuah metode yang digunakan untuk mencari pengetahuan dalam database. *Knowledge discovery in Database* (KDD) didefinisikan sebagai ekstraksi informasi potensial, implisit dan tidak dikenal dari sekumpulan data [13]. Proses KDD dapat dilihat pada Gambar. 2.



Gambar 2. Tahapan KDD

Data Selection

Data diekspor dari Google Form menjadi *file spreadsheet*. Dari seluruh data yang masuk, 56 *instance* dinyatakan lengkap dan relevan sehingga tidak ada data yang dieliminasi.

Preprocessing

Tahapan ini meliputi pemeriksaan adanya duplikasi data (dan tidak ditemukan duplikasi), verifikasi bahwa tidak terdapat *missing value*, penyeragaman penamaan kolom agar konsisten, serta konversi skala Likert menjadi nilai numerik 1–5. Setelah seluruh penyesuaian dilakukan, *dataset* dinyatakan bersih dan siap digunakan pada tahap transformasi berikutnya.

Data Transformation

Transformasi data dilakukan melalui beberapa langkah penting untuk menyiapkan dataset sebelum proses klasifikasi.

a. Penghitungan Skor Total Tiap Responden

Setiap responden memiliki n indikator (item kuesioner) dengan skala Likert 1–5.

Skor total dihitung dengan rumus:

$$\text{Skor Total}_i = \sum_{j=1}^n X_{ij}$$

di mana:

X_{ij} = nilai responden ke- i pada indikator ke- j

n = jumlah indikator kesadaran privasi

b. Pengelompokan Kelas Kesadaran

Pengelompokan kelas dilakukan berdasarkan skor total responden. Secara teoretis, *cut-off* dihitung dari nilai minimum dan maksimum skor Likert:

$$\text{CutOff} = \frac{\text{Skor Maksimum} + \text{Skor Minimum}}{2}$$

Dengan 13 indikator, diperoleh skor minimum 13 dan maksimum 65, sehingga *cut-off* teoretis adalah:

$$\text{CutOff} = \frac{65 + 13}{2} = 39$$

Dalam penelitian survei, pendekatan ini termasuk metode *mid-range classification*, yaitu teknik kategorisasi berdasarkan titik tengah rentang teoritis. Namun, jika distribusi data tidak seimbang, literatur menganjurkan penggunaan *distribution-based categorization*, yakni penentuan *cut-off* berdasarkan pola sebaran data aktual agar kategori lebih representatif.

Pada penelitian ini, lebih dari 90% responden memiliki skor di atas 39, sehingga *cut-off* teoretis menghasilkan ketidakseimbangan kelas yang ekstrem. Oleh karena itu digunakan *cut-off* empiris 50, yang memberikan pemisahan kelas lebih seimbang dan sesuai untuk proses klasifikasi. Kategori ditetapkan sebagai berikut:

$$\text{Kelas}_i = \begin{cases} \text{TINGGI}, & \text{Jika Skor Total}_i \geq 50 \\ \text{STANDAR}, & \text{Jika Skor Total}_i < 50 \end{cases}$$

c. Penyesuaian Format Dataset

Dataset disusun dalam format tabel yang memuat 13 variabel input, satu variabel turunan berupa (Total Skor), dan satu variabel kelas (TINGGI/STANDAR). Dataset yang telah tersusun dalam format CSV/Excel tersebut digunakan sebagai masukan pada Orange Data Mining untuk proses klasifikasi.

d. Ilustrasi Data Mentah dan Hasil Transformasi

Data mentah berisi jawaban skala Likert pada 13 indikator kesadaran privasi yang menjadi dasar perhitungan skor total responden.

Tabel 2. Contoh Data Mentah Responden

Pertanyaan	P1	P2	P3	P4	...	P13
R01	4	4	5	5	...	5
R02	3	4	4	3	...	4
R03	5	5	5	5	...	4
R04	4	4	4	3	...	5
R05	5	3	4	3	...	4

Setiap responden dihitung skor totalnya, kemudian diklasifikasikan ke dalam kategori TINGGI atau STANDAR menggunakan cut-off empiris 50.

Tabel 3. Contoh Transformasi Data

Pertanyaan	P1	P2	P3	P4	...	P13	Total Skor	Kelas
R01	4	4	5	5	...	5	58	TINGGI
R02	3	4	4	3	...	4	49	STANDAR
R03	5	5	5	5	...	4	65	TINGGI
R04	4	4	4	3	...	5	50	TINGGI
R05	5	3	4	3	...	4	47	STANDAR

Tabel 3 menunjukkan hasil transformasi ini menghasilkan dataset final yang telah dilengkapi dua atribut tambahan, yaitu (Total Skor) dan (Kelas). Struktur *dataset* semacam ini diperlukan agar pemodelan Naïve Bayes dapat berjalan optimal, karena setiap instance telah memiliki label kelas yang jelas sesuai tujuan klasifikasi.

Data Mining (Algoritma Naïve Bayes)

Proses klasifikasi dilakukan menggunakan algoritma Naïve Bayes dengan tahapan berikut: Naïve Bayes menggunakan Teorema Bayes sebagai dasar perhitungan probabilitas:

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

Keterangan:

- C = kelas yang ingin diprediksi (TINGGI atau STANDAR)
- X = kumpulan atribut atau fitur responden
- $P(C|X)$ = probabilitas responden masuk ke kelas C setelah melihat datanya
- $P(X|C)$ = probabilitas kemunculan atribut X jika diketahui kelasnya C
- $P(C)$ = probabilitas awal (*prior*) kelas C
- $P(X)$ = probabilitas total dari data (dapat dianggap sebagai faktor normalisasi)

Dengan mengasumsikan bahwa seluruh atribut bersifat independen:

$$P(X|C) = \prod_{i=1}^n P(X_i | C)$$

Keterangan Rumus:

- x_i = nilai fitur ke-i (misalnya skor pengetahuan, sikap, perilaku, dll.)
- n = jumlah fitur
- $\prod$ = perkalian seluruh probabilitas fitur

Rumus ini berarti setiap fitur dianggap saling bebas, sehingga probabilitas digabung melalui perkalian.

Model memilih kelas dengan probabilitas posterior terbesar:

$$C^* = \arg \max \left[P(C) \prod_{i=1}^n P(X_i | C) \right]$$

Keterangan Rumus:

- C^* = kelas hasil prediksi
 - $\arg \max$ = operator yang memilih nilai dengan probabilitas terbesar
- Model akan memilih apakah seorang responden masuk kategori TINGGI atau STANDAR berdasarkan nilai probabilitas akhir.

Interpretation / Evaluation

Evaluasi model dilakukan untuk mengetahui performa klasifikasi Naïve Bayes:

- a. Akurasi
Akurasi mengukur persentase prediksi yang benar terhadap seluruh data uji.

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

Keterangan:

- TP (*True Positive*) : prediksi TINGGI yang benar
- TN (*True Negative*) : prediksi STANDAR yang benar



FP (*False Positive*) : diprediksi TINGGI tetapi sebenarnya STANDAR
 FN (*False Negative*) : diprediksi STANDAR tetapi sebenarnya TINGGI

Akurasi menunjukkan seberapa besar proporsi prediksi model yang tepat

b. *Precision*

Precision mengukur seberapa banyak prediksi positif (TINGGI) yang benar-benar termasuk kategori tersebut.

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

Keterangan:

Semakin tinggi *precision*, semakin baik model menghindari kesalahan dalam memprediksi kelas TINGGI.

c. *Recall*

Recall mengukur kemampuan model dalam menangkap seluruh data aktual yang termasuk kategori TINGGI.

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

Keterangan:

Semakin tinggi *recall*, semakin sedikit data TINGGI yang salah diklasifikasikan menjadi STANDAR.

d. *F1-Score*

F1-Score merupakan rata-rata harmonik dari *precision* dan *recall*. Metrik ini penting jika distribusi kelas tidak seimbang.

$$F1 = 2 \times \frac{Precision \times Recall}{Precision + Recall}$$

Keterangan:

F1-Score memberikan penilaian menyeluruh terhadap performa model, terutama pada kasus FN dan FP yang sensitif seperti klasifikasi tingkat kesadaran.

e. *Confusion Matrix*

Confusion Matrix menampilkan distribusi prediksi model untuk setiap kelas (TINGGI dan STANDAR). Tabel ini berfungsi untuk melihat pola kesalahan model secara rinci, termasuk:

- Berapa banyak mahasiswa dengan kesadaran TINGGI diprediksi benar
- Berapa banyak mahasiswa dengan kesadaran STANDAR diprediksi salah
- Tipe kesalahan apa yang paling sering terjadi (FP atau FN)

3. HASIL DAN PEMBAHASAN

Gambaran Umum Data Penelitian

Data penelitian diperoleh dari 56 mahasiswa yang mengisi 13 indikator kesadaran privasi menggunakan skala Likert 1–5. Setiap indikator menggambarkan dimensi pengetahuan, sikap, dan perilaku dalam menjaga data pribadi di platform digital. Setelah dilakukan transformasi data, seluruh responden memiliki (Total Skor) sebagai representasi tingkat kesadaran serta sebuah label kelas yang membedakan kategori TINGGI dan STANDAR. Distribusi awal menunjukkan bahwa sebagian besar responden berada pada kategori TINGGI, memberikan indikasi awal bahwa kelompok mahasiswa memiliki pemahaman relatif baik terhadap isu privasi data.

Hasil Tahapan KDD

Data Selection

Seluruh 56 data responden dinyatakan lengkap dan valid untuk digunakan. Tidak ada data yang tereliminasi karena setiap responden mengisi seluruh indikator secara penuh. Variabel yang tidak relevan seperti *timestamp* dikeluarkan sehingga analisis hanya berfokus pada indikator kesadaran privasi.

Preprocessing

Tahap pembersihan data menunjukkan tidak adanya duplikasi maupun *missing value*. Skala Likert telah dikonversi menjadi nilai numerik 1–5 dan nama variabel distandarkan agar seragam. Dengan demikian, *dataset* dinyatakan siap untuk tahap transformasi tanpa memerlukan perlakuan tambahan.



Data Transformation

Transformasi meliputi penghitungan (Total Skor) per responden dan pengelompokan kelas kesadaran menggunakan *cut-off* empiris 50. *Cut-off* ini dipilih karena pembagian teoretis (39) menghasilkan ketidakseimbangan kelas yang ekstrem, sedangkan *cut-off* 50 menghasilkan distribusi yang lebih representatif:

- TINGGI : 46 responden
- STANDAR : 10 responden

Hasil transformasi menunjukkan bahwa mayoritas mahasiswa memiliki skor yang cukup tinggi, sehingga model perlu mampu membedakan pola halus antara dua kategori tersebut.

Data Mining Menggunakan Naïve Bayes

Model klasifikasi dibangun menggunakan algoritma Naïve Bayes. *Dataset* final dengan 13 fitur, total skor, dan kelas *output* dimasukkan ke dalam proses pelatihan. Data kemudian diuji untuk menghasilkan prediksi kategori kesadaran. Skenario pengujian mengikuti mekanisme evaluasi bawaan algoritma, yaitu pembagian data ke dalam proses pelatihan–pengujian dan pengukuran performa menggunakan berbagai metrik. Hasil pemodelan disajikan pada bagian berikut.

Hasil Evaluasi Model

Evaluasi dilakukan menggunakan *confusion matrix* dan metrik performa seperti akurasi, *precision*, *recall*, *F1-score*, AUC, dan MCC. Penilaian ini bertujuan memastikan model mampu membedakan kelas secara konsisten dan tidak bias terhadap salah satu kategori.

Confusion Matrix

Tabel 4. *Confusion Matrix*

Actual / Prediction	STANDAR	TINGGI	Total
STANDAR	9	1	10
TINGGI	4	42	46
Total	13	43	56

Pada Tabel 4, model menghasilkan 51 prediksi benar dari 56 klasifikasi, atau akurasi sekitar 91.1%. Sebagian besar prediksi benar terdapat pada kelas TINGGI (42 TP). Kesalahan klasifikasi paling banyak terdapat pada FN, yakni 4 mahasiswa yang sebenarnya berkategori TINGGI namun diprediksi STANDAR.

Evaluasi Performa Model

Tabel 5. Evaluasi Model

Model	Akurasi	Presisi	Recall	F1 - Score	Auc	MCC
Naive Bayes	0.911	0.926	0.911	0.915	0.976	0.738

Pada Tabel 5, nilai AUC yang tinggi (0.976) menunjukkan kemampuan diskriminasi model yang sangat baik terhadap kedua kelas, sementara MCC sebesar 0.738 mengindikasikan korelasi prediksi yang kuat meskipun data tidak seimbang.

Pembahasan

Interpretasi *Confusion Matrix*

Tingginya jumlah *True Positive* (42) menunjukkan bahwa pola jawaban mahasiswa dengan kesadaran privasi tinggi cenderung stabil sehingga mudah dikenali oleh model. Sebaliknya, munculnya 4 *False Negative* mengindikasikan bahwa sebagian mahasiswa dengan total skor tinggi memberikan respons yang cenderung mirip dengan kategori STANDAR pada beberapa indikator spesifik. Fenomena ini sejalan dengan literatur mengenai *privacy paradox*, yaitu kondisi ketika pemahaman terhadap privasi tidak selalu konsisten dengan perilaku yang ditunjukkan pengguna.

Kesalahan *False Positive* yang hanya berjumlah satu menunjukkan bahwa model cukup konservatif dalam menetapkan label TINGGI, sehingga tingkat *precision* yang tinggi dapat dicapai.

Analisis Metrik Kinerja

Nilai akurasi 0.911 dan *F1-score* 0.915 menunjukkan bahwa model bekerja sangat baik untuk kasus klasifikasi biner ini. *Precision* yang tinggi (0.926) menandakan bahwa hampir semua prediksi mahasiswa berkategori TINGGI benar adanya. Hal ini penting dalam konteks penelitian karena prediksi yang keliru sebagai TINGGI dapat menyebabkan institusi salah menilai tingkat kesiapan mahasiswa dalam memahami privasi.

Nilai AUC 0.976 menegaskan bahwa model mampu membedakan kelas dengan sangat stabil di berbagai ambang batas klasifikasi. Temuan ini konsisten dengan teori bahwa Naïve Bayes memberikan performa optimal pada *dataset* yang relatif kecil dan memiliki fitur yang bersifat independen secara kondisi.

MCC sebesar 0.738 mengonfirmasi bahwa model memiliki kemampuan prediksi yang baik meskipun terdapat ketidakseimbangan jumlah antara kelas STANDAR dan TINGGI. Dengan demikian, model tidak menunjukkan bias yang signifikan.

Keterkaitan dengan Penelitian Sebelumnya

Temuan penelitian ini sejalan dengan studi-studi sebelumnya yang menunjukkan bahwa Naïve Bayes merupakan algoritma yang efektif untuk data survei berbasis Likert dan ukuran sampel kecil hingga menengah. Penelitian terkait kesadaran privasi juga menegaskan bahwa responden muda umumnya memiliki tingkat kesadaran privasi lebih tinggi, yang sesuai dengan dominasi kelas TINGGI dalam penelitian ini. Pada sisi lain, kemunculan FN juga konsisten dengan penelitian mengenai *attitude-behavior gap* dalam perilaku digital, di mana pemahaman tinggi tidak selalu selaras dengan tindakan yang diambil.

Keterbatasan dan Implikasi

Sampel penelitian yang relatif kecil dan diperoleh melalui *voluntary response sampling* berpotensi menimbulkan bias karena responden yang mengisi survei cenderung berasal dari kelompok yang memiliki literasi digital lebih tinggi. Oleh sebab itu, generalisasi temuan harus dilakukan secara hati-hati. Selain itu, pemilihan *cut-off* empiris meskipun efektif dalam menyeimbangkan kelas, tetap memerlukan pengujian lanjutan menggunakan pendekatan statistik lain untuk memastikan keandalannya.

Meskipun demikian, hasil penelitian memberikan kontribusi penting dalam pemetaan kesadaran privasi mahasiswa. Model yang dibangun dapat digunakan oleh lembaga pendidikan sebagai alat awal untuk menilai kebutuhan program literasi keamanan data dan strategi penguatan kesadaran privasi di lingkungan kampus.

Ringkasan Temuan

Penelitian menunjukkan bahwa model Naïve Bayes mampu melakukan klasifikasi tingkat kesadaran privasi mahasiswa dengan akurasi tinggi. Mayoritas mahasiswa berada pada kategori TINGGI, namun terdapat variasi dalam pola respon yang menyebabkan sebagian mahasiswa dengan skor tinggi diklasifikasikan keliru. Model memiliki performa yang kuat pada seluruh metrik evaluasi, terutama AUC. Temuan ini menguatkan literatur bahwa Naïve Bayes cocok diterapkan dalam analisis perilaku digital dan dapat menjadi dasar bagi pengembangan kebijakan literasi privasi di lingkungan akademik.

4. SIMPULAN

Berdasarkan hasil analisis dan evaluasi yang telah dilakukan, penelitian ini menyimpulkan bahwa algoritma Naïve Bayes mampu mengklasifikasikan tingkat kesadaran privasi mahasiswa secara efektif meskipun menggunakan ukuran sampel yang relatif kecil. Model menunjukkan performa tinggi dengan akurasi 91.1%, *precision* 92.6%, *recall* 91.1%, *F1-score* 91.5%, serta AUC 0.976 yang mengindikasikan kemampuan diskriminasi yang sangat baik. Struktur *Confusion Matrix* memperlihatkan bahwa model dapat mengenali mayoritas mahasiswa berkategori Tinggi dengan konsisten, meskipun masih terdapat sejumlah kecil *False Negative* yang berkaitan dengan variasi respons pada beberapa indikator. Temuan ini sejalan dengan literatur yang menyatakan bahwa Naïve Bayes cocok digunakan untuk data survei berbasis Likert dan konteks klasifikasi biner.

Penelitian ini juga menegaskan bahwa sebagian besar mahasiswa memiliki tingkat kesadaran privasi yang tinggi, namun tetap terdapat kelompok kecil yang berada pada kategori STANDAR. Hal ini memberikan implikasi bahwa institusi pendidikan perlu meningkatkan program literasi privasi secara lebih terarah, terutama untuk mahasiswa yang menunjukkan pola respons kurang konsisten. Keterbatasan penelitian mencakup penggunaan teknik *voluntary response sampling* dan ukuran sampel yang terbatas, sehingga generalisasi hasil perlu dilakukan secara hati-hati. Penelitian selanjutnya dianjurkan untuk menggunakan sampel lebih besar, metode pengelompokan kelas berbasis pendekatan statistik alternatif, serta membandingkan performa Naïve Bayes dengan algoritma klasifikasi lainnya.

DAFTAR PUSTAKA

- [1] A. M. Baqis and M. I. P. Nasution, "Pentingnya Privasi dan Keamanan Data Pribadi di Era Digital," *Jurnal Manajemen dan Pendidikan Agama Islam*, vol. 3, no. 3, pp. 396–404, 2025, doi: 10.61132/jmpai.v3i3.1150.
- [2] S. Stefanni, Z. Zulfachmi, Z. Zulkipli, and A. Saputra, "Analisis Sentimen Pengguna X Terhadap Kebocoran Data Pribadi Menggunakan Algoritma Naïve Bayes Classifier," *Jurnal Bangkit Indonesia*, vol. 14, no. 1, pp. 32–40, 2025, doi: 10.52771/bangkitindonesia.v14i1.434.
- [3] R. Firmansyah and D. Darmawan, "Kesadaran Masyarakat terhadap Perlindungan Data Pribadi di Era Digital," *Jurnal Ilmiah Komputer dan Informatika*, vol. 8, no. 2, pp. 123–130, 2020, doi: <https://doi.org/10.5281/zenodo.12608751>.
- [4] Nopriadi, "Menjaga Privasi Digital: Studi Tentang Kesadaran Mahasiswa dalam Perlindungan Data Pribadi di Media Sosial," *Polygon : Jurnal Ilmu Komputer dan Ilmu Pengetahuan Alam*, vol. 2, no. 6, pp. 87–97, 2024, doi: 10.62383/polygon.v2i6.297.



- [5] P. R. Sihombing and I. F. Yuliati, "Penerapan Metode Machine Learning dalam Klasifikasi Risiko Kejadian Berat Badan Lahir Rendah di Indonesia," *MATRIK : Jurnal Manajemen, Teknik Informatika dan Rekayasa Komputer*, vol. 20, no. 2, pp. 417–426, 2021, doi: 10.30812/matrik.v20i2.1174.
- [6] M. Siddik, Y. Desnelita, and Gustientiedina, "Penerapan Naïve Bayes untuk Memprediksi Tingkat Kepuasan Mahasiswa Terhadap Pelayanan Akademis," *Jurnal Infomedia*, vol. 2, no. 4, pp. 89–93, 2019, doi: 10.30811/jim.v4i2.1892.
- [7] A. Nurian, M. S. Ma'arif, I. N. Amalia, and C. Rozikin, "Analisis Sentimen Pengguna Aplikasi Shopee Pada Situs Google Play Menggunakan Naive Bayes Classifier," *Jurnal Informatika dan Teknik Elektro Terapan*, vol. 12, no. 1, pp. 97–105, 2024, doi: 10.23960/jitet.v12i1.3631.
- [8] D. Nasien *et al.*, "Perbandingan Implementasi Machine Learning Menggunakan Metode KNN , Naive Bayes , Dan Logistik Regression Untuk Mengklasifikasi Penyakit Diabetes," *JEKIN - Jurnal Teknik Informatika*, vol. 4, no. 1, pp. 10–17, 2024, doi: 10.58794/jekin.v4i1.640.
- [9] D. Nurfadlillah, "Sentiment Analysis About Indonesian People'S Awareness About Cyber Security in Determining Data Leakage Using Naïve Bayes Classifier Algorithm," *Jurnal Elektro Luceat*, vol. 9, no. 1, pp. 64–72, 2023, doi: 10.32531/jelekn.v9i1.598.
- [10] A. Zy and W. Hadikristanto, "Implementasi Algoritma Metode Naive Bayes dan Support Vector Machine Tentang Pembobolan dan Kebocoran Data di Twitter," *Bulletin of Information Technology (BIT)*, vol. 4, no. 1, pp. 49–56, 2023, doi: 10.47065/bit.v4i1.493.
- [11] A. Turmudi Zy, A. Nugroho, A. Rivaldi, and I. Afriantoro, "Analisis Sentimen Terhadap Pembobolan Data pada Twitter dengan Algoritma Naive Bayes," *Jurnal Teknologi Informatika dan Komputer*, vol. 8, no. 2, pp. 202–213, 2022, doi: 10.37012/jtik.v8i2.1240.
- [12] M. Muharrom, "Analisis Komparasi Algoritma Data Mining Naive Bayes, K-Nearest Neighbors dan Regresi Linier Dalam Prediksi Harga Emas," *Bulletin of Information Technology (BIT)*, vol. 4, no. 4, pp. 430–438, 2023, doi: 10.47065/bit.v4i4.986.
- [13] H. Syahputra, L. Mayola, and D. Guswandi, "Jurnal KomtekInfo Clustering Tingkat Penjualan Menu (Food and Beverage)," *KomtekInfo*, vol. 9, no. 1, pp. 29–33, 2022, doi: 10.35134/komtekinfo.v9i1.274.