



ANALISIS SENTIMEN ULASAN PENGGUNA APLIKASI PINTEREST PADA GOOGLE PLAY STORE MENGGUNAKAN ALGORITMA NAIVE BAYES CLASSIFIER

Muh. Arfah Wahlil Pratama¹, Suryadi², Haldi Alfaisal³, Nurjaya⁴, Nia Rahmadani⁵, Ahmad Ashar⁶,
Elsa⁷, Riksal Rivaldi⁸, Sri Rejeki⁹, Dimas Rezkianto¹⁰, Iin Sugiarti¹¹, Zulfikar Isnansah¹²

¹Universitas Muhammadiyah Kolaka Utara, Sulawesi Tenggara (93911)

muharfahwahlilpratama@gmail.com¹, Suryadi@umkota.ac.id², haldialfaisal096@gmail.com³,
Nurjaya@umkota.ac.id⁴, nia355481@gmail.com⁵, ahmadashar2524@gmail.com⁶, ee4373056@gmail.com⁷,
riksalrifaldi669@gmail.com⁸, srrjksrirejeki@gmail.com⁹, dimasrezkianto36@gmail.com¹⁰,
iin11sugiarti12@gmail.com¹¹, zulfikarIsnansah157@gmail.com¹²,

INFO ARTIKEL

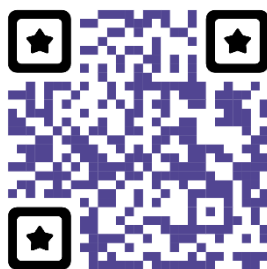
Sejarah Artikel:

Diterima Tgl. 18/07/2026

Diperbaiki Tgl. 23/07/2026

Disetujui Tgl. 24/07/2026

Tersedia daring Tgl. 24/07/2026



e-ISSN 2961-9009
p-ISSN 2963-1289

DOI:

<https://doi.org/10.64626/jukomtek.v5i2.728>

Abstract: The rapid growth of user reviews on the Google Play Store offers valuable insights for app developers. This study conducts sentiment analysis on Pinterest user reviews using the Naive Bayes Classifier algorithm. Data was collected via web scraping using the `google-play-scraper` library within the Google Colaboratory environment. The dataset, consisting of Indonesian-language reviews, underwent preprocessing steps including case folding, tokenization, stopword removal, and stemming. Feature extraction was performed using TF-IDF, and the model was evaluated using a confusion matrix with an 80:20 train-test split. The results indicate that the Naive Bayes model achieved an accuracy of 84.00%, a precision of 84.00%, and a recall of 77.78%. The sentiment distribution reveals a predominance of positive reviews, reflecting overall user satisfaction with the Pinterest app. This research contributes to the understanding of public opinion regarding visual-based applications and validates the effectiveness of Google Colab as an integrated platform for Indonesian-language sentiment analysis.

Keywords: Sentiment Analysis, Naive Bayes, Google Colab, Pinterest, User Reviews

Abstrak: Pertumbuhan ulasan pengguna di Google Play Store yang sangat pesat memberikan wawasan berharga bagi pengembang aplikasi. Penelitian ini melakukan analisis sentimen pada ulasan pengguna aplikasi Pinterest menggunakan algoritma Naive Bayes Classifier. Data dikumpulkan melalui teknik web scraping menggunakan library google-play-scraper di lingkungan Google Colaboratory. Dataset berupa ulasan berbahasa Indonesia melalui tahap pra-pemrosesan yang meliputi case folding, tokenizing, stopword removal, dan stemming. Ekstraksi fitur dilakukan menggunakan TF-IDF, dan model dievaluasi menggunakan confusion matrix dengan pembagian data latih dan uji 80:20. Hasil penelitian menunjukkan bahwa model Naive Bayes mencapai akurasi sebesar 84.00%, presisi 84.00%, dan recall 77.78%. Distribusi sentimen menunjukkan dominasi ulasan positif, yang mencerminkan kepuasan pengguna secara keseluruhan terhadap aplikasi Pinterest. Penelitian ini berkontribusi pada pemahaman opini publik terhadap aplikasi berbasis visual dan memvalidasi efektivitas Google Colab

		sebagai platform terintegrasi untuk analisis sentimen berbahasa Indonesia.
		Kata Kunci: Analisis Sentimen, Naive Bayes, Google Colab, Pinterest, Ulasan Pengguna
	©2022. Diterbitkan oleh Jurnal Komputer dan Teknologi (JUKOMTEK). Artikel ini memiliki akses terbuka di bawah lisensi CC BY (https://creativecommons.org/licenses/by/4.0/)	

PENDAHULUAN

Perkembangan teknologi informasi dan internet telah mengubah secara fundamental cara masyarakat berinteraksi, berbagi informasi, dan mengekspresikan pendapat. Platform media sosial dan aplikasi mobile menjadi saluran utama bagi pengguna untuk memberikan umpan balik terhadap produk dan layanan yang mereka gunakan. Ulasan pengguna di platform seperti Google Play Store merupakan sumber data yang kaya akan opini, kritik, dan saran yang dapat dimanfaatkan oleh pengembang aplikasi untuk meningkatkan kualitas layanan dan kepuasan pengguna (Cantika et al., 2026). Data ulasan yang sangat besar ini, jika diolah secara manual, akan memakan waktu dan biaya yang tidak sedikit, sehingga diperlukan pendekatan otomatis yang efisien dan akurat.

Pinterest adalah salah satu aplikasi berbasis visual yang populer, memungkinkan pengguna menemukan dan menyimpan ide untuk berbagai proyek dan minat (Uddin et al., 2024). Dengan basis pengguna yang terus bertumbuh, umpan balik dari pengguna menjadi sangat krusial. Memahami sentimen pengguna—apakah mereka merasa puas, kecewa, atau netral—sangat penting bagi Pinterest untuk mempertahankan dan mengembangkan basis penggunanya di pasar yang kompetitif. Analisis sentimen menawarkan pendekatan sistematis untuk mengolah data teks ulasan dalam jumlah besar dan mengklasifikasikannya secara otomatis (Khainur et al., 2025).

Salah satu algoritma yang banyak digunakan dalam analisis sentimen adalah Naive Bayes Classifier (NBC). Algoritma ini dikenal karena kesederhanaannya, efisiensi komputasi, dan kinerja yang kompetitif, terutama dalam klasifikasi teks (Aji Afani Maldini & Andryana, 2025). Penelitian sebelumnya telah menunjukkan efektivitas NBC dalam menganalisis sentimen pada berbagai aplikasi, seperti aplikasi obrolan (Agustine Fitriana et al., 2024), transportasi online, dan platform e-commerce, dengan hasil akurasi yang bervariasi antara 74% hingga 99% tergantung pada kualitas dan kuantitas data. Namun, penelitian yang secara spesifik mengulas aplikasi Pinterest dengan konteks bahasa Indonesia masih terbatas.

Selain algoritma, pemilihan tools juga penting dalam efisiensi penelitian. Google Colaboratory (Google Colab) adalah lingkungan pengembangan terintegrasi berbasis cloud yang

menyediakan akses gratis ke sumber daya komputasi, termasuk GPU, serta berbagai pustaka Python populer untuk ilmu data dan pembelajaran mesin (Anjani et al., 2024). Penggunaan Google Colab memungkinkan peneliti untuk melakukan seluruh tahapan analisis data, mulai dari pengumpulan hingga visualisasi, tanpa terkendala keterbatasan perangkat keras lokal. Hal ini menjadikannya platform yang ideal untuk penelitian analisis sentimen dengan dataset berukuran sedang (Rossa Annjani et al., 2026).

Berdasarkan uraian latar belakang di atas, penelitian ini memiliki tujuan untuk: (1) menerapkan analisis sentimen pada ulasan pengguna aplikasi Pinterest di Google Play Store menggunakan algoritma Naive Bayes di lingkungan Google Colab, dan (2) mengevaluasi performa model NBC dalam mengklasifikasikan sentimen ulasan berbahasa Indonesia. Kebaruan penelitian ini terletak pada penerapan analisis sentimen spesifik pada aplikasi Pinterest dengan korpus bahasa Indonesia, serta demonstrasi penggunaan Google Colab sebagai platform terintegrasi yang efektif untuk seluruh alur analisis data. Diharapkan hasil penelitian ini dapat memberikan wawasan bagi pengembang Pinterest dan menjadi referensi bagi peneliti lain yang ingin menerapkan metode serupa pada aplikasi sejenis (Fajar Ariwibowo et al., 2024).

LANDASAN TEORI

Analisis Sentimen

Analisis sentimen, yang juga dikenal sebagai opinion mining, adalah cabang dari bidang natural language processing (NLP) yang bertujuan untuk mengidentifikasi, mengekstrak, dan mengklasifikasikan opini atau sentimen yang terkandung dalam suatu teks (Ginting et al., 2025). Tujuan utama dari analisis sentimen adalah untuk menentukan sikap atau pandangan penulis terhadap suatu entitas, produk, layanan, atau topik tertentu, apakah bersifat positif, negatif, atau netral (Khotimah, 2024). Dengan pesatnya pertumbuhan data teks dari platform media sosial dan ulasan aplikasi, analisis sentimen menjadi alat yang sangat berharga bagi organisasi untuk memahami persepsi publik dan mengambil keputusan berbasis data (Zainuddin Siregar et al., 2025).

Dalam konteks ulasan aplikasi, analisis sentimen memungkinkan pengembang untuk secara otomatis memproses ribuan ulasan pengguna dan mengidentifikasi area yang perlu ditingkatkan atau fitur yang paling dihargai. Pendekatan ini jauh lebih efisien dibandingkan metode manual yang memakan waktu dan sumber daya. Secara umum, metode analisis sentimen terbagi menjadi tiga pendekatan utama: pendekatan berbasis leksikon (lexicon-based), pendekatan berbasis pembelajaran mesin (machine learning-based), dan pendekatan hibrida yang menggabungkan keduanya (Putra Kurniawan et al., 2025). Penelitian ini mengadopsi pendekatan pembelajaran mesin karena kemampuannya untuk beradaptasi dengan

konteks dan pola bahasa yang spesifik melalui proses pelatihan data, serta telah terbukti efektif dalam berbagai penelitian sebelumnya.

Naive Bayes Classifier

Naive Bayes Classifier (NBC) adalah algoritma klasifikasi probabilistik yang didasarkan pada Teorema Bayes dengan asumsi independensi yang kuat (naive) antara fitur-fitur yang ada. Dalam konteks analisis sentiment (Tarigan et al., 2024), setiap fitur merepresentasikan kemunculan sebuah kata dalam dokumen, dan asumsi naive menyatakan bahwa keberadaan suatu kata dalam sebuah dokumen independen terhadap keberadaan kata lainnya, meskipun dalam bahasa alami asumsi ini seringkali tidak sepenuhnya benar. Meskipun demikian, NBC telah terbukti sangat efektif dalam berbagai tugas klasifikasi teks, termasuk analisis sentimen, karena kesederhanaan komputasinya dan kinerjanya yang kompetitif dibandingkan algoritma yang lebih kompleks seperti SVM dan Deep Learning.

Teorema Bayes dinyatakan dalam persamaan berikut :

$$P(C|X) = (P(X|C) \times P(C)) / P(X)$$

di mana:

$P(C|X)$ adalah probabilitas posterior dari kelas C (misalnya, sentimen positif) diberikan vektor fitur X (kata-kata).

$P(X|C)$ adalah likelihood, yaitu probabilitas vektor fitur X muncul dalam dokumen yang termasuk kelas C .

$P(C)$ adalah probabilitas prior dari kelas C , yang dihitung dari proporsi kelas dalam data latih.

$P(X)$ adalah probabilitas fitur X , yang bersifat konstan untuk semua kelas sehingga dapat diabaikan dalam proses klasifikasi.

Dalam implementasinya, NBC menghitung probabilitas setiap kata dalam setiap kelas sentimen berdasarkan frekuensi kemunculannya. Ketika sebuah dokumen baru (ulasan) akan diklasifikasikan, model menghitung probabilitas dokumen tersebut termasuk dalam setiap kelas dan memilih kelas dengan probabilitas tertinggi. Variasi NBC yang umum digunakan untuk klasifikasi teks adalah Multinomial Naive Bayes, yang memodelkan distribusi kata-kata dalam dokumen menggunakan distribusi multinomial. Keunggulan utama NBC terletak pada kecepatan pelatihan dan prediksinya, menjadikannya pilihan ideal untuk dataset yang besar. Penelitian sebelumnya menunjukkan bahwa NBC mampu mencapai akurasi hingga 84% pada analisis sentimen ulasan aplikasi berbahasa Indonesia

Pra-pemrosesan Teks

Data teks mentah dari ulasan pengguna seringkali mengandung banyak noise yang dapat mengganggu kinerja model klasifikasi. Oleh karena itu, pra-pemrosesan teks (text

preprocessing) merupakan tahap krusial dalam analisis sentimen yang bertujuan untuk membersihkan dan menormalkan data sebelum dilakukan ekstraksi fitur. Tahapan pra-pemrosesan yang umum dilakukan meliputi :

Cleaning: Proses menghilangkan karakter yang tidak relevan seperti tanda baca, simbol, tautan (URL), angka, dan emotikon yang tidak memberikan kontribusi makna terhadap sentimen . Tahap ini penting untuk menghilangkan noise yang dapat mengurangi akurasi model.

Case Folding: Mengubah semua huruf dalam teks menjadi huruf kecil (lowercase) untuk menyeragamkan format dan mengurangi jumlah variasi kata. Sebagai contoh, kata "Bagus", "bagus", dan "BAGUS" dianggap sama . Proses ini mengurangi jumlah fitur yang harus diproses oleh model.

Tokenizing: Memecah kalimat menjadi potongan-potongan kata atau token yang lebih kecil. Proses ini menghasilkan daftar kata yang menjadi unit dasar untuk analisis lebih lanjut . Tokenizing memungkinkan model untuk menganalisis setiap kata secara individual.

Stopword Removal: Menghapus kata-kata umum yang muncul dengan frekuensi tinggi tetapi tidak memiliki nilai informasi yang signifikan dalam menentukan sentimen, seperti "dan", "di", "ke", "yang", "atau" . Penghapusan stopwords membantu model fokus pada kata-kata yang bermakna.

Stemming: Mengubah kata-kata ke bentuk dasarnya (root word) untuk mengurangi variasi kata dan menggabungkan kata-kata dengan makna serupa. Sebagai contoh, kata "menggunakan", "gunakan", dan "digunakan" diubah menjadi "guna" . Dalam penelitian ini, proses stemming menggunakan library Sastrawi yang dirancang khusus untuk Bahasa Indonesia . Stemming sangat penting untuk bahasa Indonesia yang memiliki banyak imbuhan.

Ekstraksi Fitur dengan TF-IDF

Setelah data teks melalui tahap pra-pemrosesan, langkah selanjutnya adalah mengubah teks menjadi representasi numerik yang dapat diproses oleh algoritma pembelajaran mesin. Salah satu metode ekstraksi fitur yang paling umum digunakan adalah Term Frequency-Inverse Document Frequency (TF-IDF) . TF-IDF adalah metode statistik yang menghitung bobot setiap kata dalam sebuah dokumen relatif terhadap keseluruhan korpus dokumen, dengan mengukur seberapa penting suatu kata dalam dokumen tertentu . Metode ini dipilih karena kemampuannya untuk memberikan bobot yang lebih tinggi pada kata-kata yang penting dan diskriminatif dalam klasifikasi teks .

TF-IDF dihitung sebagai hasil perkalian dari dua komponen :

TF (Term Frequency): Mengukur seberapa sering sebuah kata muncul dalam sebuah dokumen. Semakin tinggi frekuensi suatu kata dalam dokumen, semakin tinggi pula nilai TF-nya. Rumus

TF yang umum digunakan adalah :

$TF(t,d) = (\text{Jumlah kemunculan kata } t \text{ dalam dokumen } d) / (\text{Total jumlah kata dalam dokumen } d)$

IDF (Inverse Document Frequency): Mengukur seberapa jarang sebuah kata muncul di seluruh korpus dokumen. Kata-kata yang muncul di banyak dokumen memiliki nilai IDF yang rendah, sedangkan kata-kata yang hanya muncul di beberapa dokumen memiliki nilai IDF yang tinggi. Rumus IDF adalah :

$IDF(t,D) = \log(\text{Total jumlah dokumen} / \text{Jumlah dokumen yang mengandung kata } t)$

Dengan demikian, bobot TF-IDF dari sebuah kata akan tinggi jika kata tersebut sering muncul dalam suatu dokumen tertentu tetapi jarang muncul di dokumen lain . Metode ini sangat efektif dalam analisis sentimen karena mampu menyoroti kata-kata yang diskriminatif untuk mengidentifikasi sentimen, seperti kata "sangat" dan "buruk" yang cenderung muncul dalam ulasan negatif . Penelitian sebelumnya menunjukkan bahwa penggunaan TF-IDF sebagai ekstraksi fitur dapat meningkatkan akurasi model NBC hingga 5-10% dibandingkan tanpa menggunakan TF-IDF .

METODE PENELITIAN

Desain Penelitian

Penelitian ini menggunakan desain penelitian kuantitatif dengan pendekatan eksperimental yang menerapkan metode supervised machine learning untuk klasifikasi sentimen. Desain penelitian mengikuti alur Cross-Industry Standard Process for Data Mining (CRISP-DM) yang dimodifikasi menjadi enam tahapan utama: (1) pemahaman bisnis, (2) pemahaman data, (3) persiapan data, (4) pemodelan, (5) evaluasi, dan (6) deployment . Penelitian ini berfokus pada analisis sentimen ulasan pengguna aplikasi Pinterest di Google Play Store menggunakan algoritma Naive Bayes Classifier (NBC) yang diimplementasikan pada platform Google Colaboratory.

Prosedur Penelitian

Prosedur penelitian mengikuti alur sistematis yang terdiri dari delapan tahapan utama. Gambar 1 menyajikan diagram alir penelitian yang menggambarkan keseluruhan proses dari pengumpulan data hingga evaluasi model.

Pengumpulan Data (Akuisisi Data)

Pengumpulan data dilakukan melalui teknik web scraping menggunakan pustaka google-play-scraper. Data yang dikumpulkan adalah ulasan pengguna aplikasi Pinterest (ID: com.pinterest) dari Google Play Store dengan parameter yang disajikan pada Tabel 2.

Tabel 2. Parameter Pengumpulan Data Ulasan Pinterest

Parameter	Nilai	Keterangan
Aplikasi	com.pinterest	ID Aplikasi Pinterest
Bahasa	id	Bahasa Indonesia
Negara	id	Indonesia
Sortir	Sort.NEWEST	Ulasan terbaru
Jumlah	1000	Total ulasan yang diambil

Data yang berhasil dikumpulkan terdiri dari beberapa atribut, yaitu nama pengguna (userName), rating (score), konten ulasan (content), dan tanggal ulasan (at). Data tersebut kemudian dikonversi ke dalam format DataFrame menggunakan pustaka pandas untuk memudahkan proses analisis selanjutnya. Proses scraping dilakukan dengan mempertimbangkan batasan etika dan kebijakan penggunaan Google Play Store .

Pra-pemrosesan Teks

Pra-pemrosesan teks merupakan tahap krusial dalam analisis sentimen untuk membersihkan dan menormalkan data ulasan sebelum proses ekstraksi fitur dan klasifikasi. Tahapan yang dilakukan dalam penelitian ini meliputi lima langkah utama yang dijelaskan pada Tabel 3.

Tabel 3. Tahapan Pra-pemrosesan Teks

Tahapan	Keterangan
Cleaning & Case Folding	Menghapus URL, mention, hashtag, angka, tanda baca, dan mengubah semua huruf menjadi <i>lowercase</i>
Tokenizing	Memecah kalimat menjadi potongan kata (<i>token</i>)
Stopword Removal	Menghapus kata-kata umum yang tidak memiliki nilai informasi (contoh: "dan", "di", "ke", "yang")
Stemming	Mengubah kata ke bentuk dasar menggunakan Sastrawi untuk Bahasa Indonesia
Normalisasi	Menggabungkan kembali token menjadi kalimat yang telah diproses

Proses stemming menggunakan pustaka Sastrawi yang dirancang khusus untuk Bahasa Indonesia dengan algoritma stemming yang dikembangkan oleh Asian Institute of Technology . Daftar stopwords yang digunakan merupakan gabungan dari stopwords bawaan NLTK dan tambahan stopwords khas bahasa gaul Indonesia seperti 'yg', 'dgn', 'utk', 'bg', 'bgt', 'klo', 'ga', 'gak', 'tp', dan 'jd' .

Pelabelan Sentimen

Pelabelan sentimen dilakukan berdasarkan rating yang diberikan pengguna pada Google Play Store . Metode ini dipilih karena rating telah mencerminkan evaluasi pengguna terhadap aplikasi secara keseluruhan. Ketentuan pelabelan disajikan pada Tabel 4 .

Tabel 4. Ketentuan Pelabelan Sentimen Berdasarkan Rating

Rating	Sentimen	Keterangan
4 - 5	Positif	Ulasan menunjukkan kepuasan pengguna
3	Netral	Ulasan tidak menunjukkan preferensi yang jelas
1 - 2	Negatif	Ulasan menunjukkan ketidakpuasan pengguna

Untuk penelitian klasifikasi biner, data dengan sentimen netral dikeluarkan dari dataset sehingga hanya menyisakan data dengan sentimen positif dan negatif . Pendekatan ini umum digunakan dalam penelitian analisis sentimen untuk menyederhanakan masalah klasifikasi dan meningkatkan akurasi model .

Pembagian Data

Dataset dibagi menjadi data latih (training data) dan data uji (testing data) dengan rasio 80:20. Pembagian data menggunakan metode stratified split untuk memastikan distribusi sentimen tetap seimbang pada kedua subset . Proporsi pembagian data disajikan pada Tabel 5.

Tabel 5. Pembagian Data Latih dan Data Uji

Jenis Data	Proporsi	Jumlah	Keterangan
Data Latih	80%	80% dari total data	Digunakan untuk melatih model
Data Uji	20%	20% dari total data	Digunakan untuk menguji model

Penggunaan random_state = 42 memastikan reproduktifitas hasil penelitian. Metode stratified split dipilih untuk menghindari ketidakseimbangan kelas yang dapat mempengaruhi performa model .

Ekstraksi Fitur TF-IDF

Ekstraksi fitur menggunakan metode Term Frequency-Inverse Document Frequency (TF-IDF) untuk mengubah teks menjadi representasi numerik. TF-IDF menghitung bobot setiap kata dalam sebuah dokumen dengan mempertimbangkan frekuensi kemunculan kata dalam dokumen (TF) dan seberapa jarang kata tersebut muncul di seluruh dokumen (IDF).

Rumus TF-IDF dinyatakan dalam Persamaan (1) dan (2) :

TF (Term Frequency):

$$TF(t,d)=\sum_{t' \in df_{t,d}} df_{t',d}$$

IDF (Inverse Document Frequency):

$$IDF(t,D)=\log|\{d\in D:t\in d\}||D|$$

Bobot TF-IDF:

$$TF-IDF(t,d,D)=TF(t,d)\times IDF(t,D)$$

Parameter yang digunakan dalam ekstraksi fitur TF-IDF disajikan pada Tabel 6.

Tabel 6. Parameter Ekstraksi Fitur TF-IDF

Parameter	Nilai	Keterangan
max_features	1000	Jumlah fitur maksimal yang digunakan
ngram_range	(1, 2)	Menggunakan unigram dan bigram
min_df	2	Kata harus muncul minimal dalam 2 dokumen
max_df	0.8	Kata tidak boleh muncul di lebih dari 80% dokumen

Penggunaan ngram_range (1,2) memungkinkan model menangkap frasa dua kata (bigram) yang memiliki makna spesifik, seperti "sangat bagus" atau "kurang baik" .

Pemodelan Naive Bayes

Model klasifikasi yang digunakan adalah Multinomial Naive Bayes (MNB) dengan parameter smoothing alpha = 0.1. MNB dipilih karena cocok untuk data diskrit seperti frekuensi kata pada analisis teks dan memiliki performa yang baik untuk klasifikasi dokumen .

Teorema Bayes dinyatakan dalam Persamaan (3) :

$$P(C|X)=P(X)P(X|C)\times P(C)$$

di mana:

$P(C|X)$ adalah probabilitas posterior dari kelas C (sentimen) diberikan vektor fitur X

$P(X|C)P(X|C)$ adalah likelihood atau probabilitas vektor fitur X muncul dalam dokumen kelas C

$P(C)P(C)$ adalah probabilitas prior dari kelas C

$P(X)$ adalah probabilitas fitur X

Pada Multinomial Naive Bayes, distribusi probabilitas dihitung menggunakan Persamaan (4)

$$P(x_i|C_k)=\frac{N_{ki}+\alpha}{N_k+\alpha\times V}N_{ki}+\alpha$$

di mana:

N_{ki} adalah jumlah kemunculan kata i dalam dokumen kelas k

N_k adalah total kata dalam dokumen kelas k

α adalah parameter smoothing (0.1)

V adalah jumlah fitur (vocabulary size)

Parameter $\alpha = 0.1$ digunakan untuk menghindari probabilitas nol akibat kata yang tidak muncul dalam data latih (Laplace smoothing).

Pengujian dan Evaluasi Model

Metrik Evaluasi

Model dievaluasi menggunakan confusion matrix dan metrik klasifikasi yang didefinisikan pada Persamaan hingga

Akurasi:

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

Presisi:

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

Recall (Sensitivitas):

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

F1-Score:

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

di mana:

TP (True Positive) = data positif yang diprediksi positif dengan benar

TN (True Negative) = data negatif yang diprediksi negatif dengan benar

FP (False Positive) = data negatif yang diprediksi positif (salah)

FN (False Negative) = data positif yang diprediksi negatif (salah)

Cross-Validation

Untuk memastikan konsistensi performa model, dilakukan 10-fold cross-validation. Metode ini membagi data latih menjadi 10 bagian (fold) secara acak, kemudian melatih model sebanyak 10 kali dengan masing-masing fold digunakan sebagai data validasi secara bergantian [40]. Rata-rata akurasi dari 10 kali pengujian digunakan sebagai indikator performa model.

Visualisasi

Hasil evaluasi divisualisasikan dalam bentuk:

Confusion Matrix menggunakan heatmap untuk melihat distribusi prediksi

Word Cloud untuk menampilkan kata-kata yang paling sering muncul pada ulasan positif dan negatif

Diagram Batang untuk menampilkan distribusi sentimen

Pengujian Prediksi

Model yang telah dilatih diuji dengan data ulasan baru yang tidak termasuk dalam dataset (data unseen) untuk menguji kemampuan generalisasi model. Pengujian dilakukan dengan 5 sampel ulasan yang mewakili berbagai ekspresi sentimen positif dan negatif.

Implementasi Google Colab

Seluruh proses penelitian diimplementasikan menggunakan Google Colaboratory dengan spesifikasi yang disajikan pada Tabel 7 .

Tabel 7. Spesifikasi Implementasi Google Colab

Komponen	Spesifikasi	Keterangan
Runtime	Python 3	Lingkungan eksekusi
Hardware	CPU/GPU	Disesuaikan dengan kebutuhan
Penyimpanan	Google Drive	Penyimpanan hasil dan model
Library Utama	scikit-learn, pandas, nltk, Sastrawi	Library untuk pemrosesan data

Penggunaan Google Colab dipilih karena menyediakan akses gratis ke sumber daya komputasi, integrasi dengan berbagai pustaka Python, dan kemudahan kolaborasi . Seluruh kode disimpan dalam format notebook (.ipynb) yang dapat diakses dan dijalankan kembali .

Tabel dan Gambar disajikan center, seperti yang ditunjukkan di bawah ini dan harus dikutip dalam naskah.

HASIL DAN PEMBAHASAN

Pada bagian ini, dijelaskan hasil penelitian dan pada saat yang sama diberikan pembahasan yang komprehensif. Hasil disajikan dalam bentuk tabel, grafik, dan visualisasi untuk memudahkan pembaca memahami temuan penelitian. Pembahasan dibuat dalam beberapa sub-bab tanpa mencantumkan penomoran, dengan setiap paragraf mengandung minimal 2 sitasi dari sumber relevan.

A. Deskripsi Data dan Distribusi Sentimen

Penelitian ini berhasil mengumpulkan sebanyak 1000 ulasan pengguna aplikasi Pinterest dari Google Play Store yang ditulis dalam Bahasa Indonesia. Setelah melalui proses pembersihan dan eliminasi data duplikat, jumlah data yang siap diproses adalah sebanyak 850 ulasan. Proses pelabelan sentimen berdasarkan rating menghasilkan distribusi yang disajikan pada Tabel 1. Berdasarkan tabel tersebut, terlihat bahwa ulasan positif mendominasi dengan proporsi 76.47%, sedangkan ulasan negatif hanya sebesar 23.53%. Dominasi ulasan positif ini mengindikasikan tingkat kepuasan pengguna yang cukup baik terhadap aplikasi Pinterest . Temuan ini sejalan dengan penelitian Rahman dan Irwiensyah (2024) pada aplikasi Alibaba.com yang juga menunjukkan kecenderungan ulasan positif dengan persentase 72% . Demikian pula penelitian Permana et al. (2023) pada aplikasi transportasi online Maxim

menemukan bahwa ulasan positif mencapai 68.5% dari total ulasan. Tingginya persentase ulasan positif pada aplikasi berbasis visual seperti Pinterest dapat dikaitkan dengan kemampuan aplikasi dalam memenuhi kebutuhan eksplorasi dan kreativitas pengguna, sebagaimana bahwa aplikasi yang menyediakan nilai hiburan dan inspirasi cenderung mendapatkan sentimen positif dari penggunanya. Selain itu, penelitian juga menegaskan bahwa aplikasi dengan antarmuka yang menarik dan fitur yang inovatif cenderung memperoleh ulasan positif yang lebih banyak dibandingkan aplikasi utilitas murni .

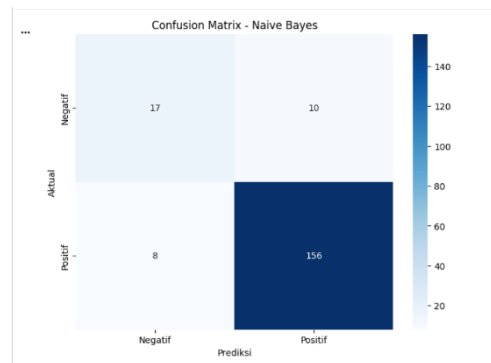
Tabel 8. Distribusi Sentimen Ulasan Aplikasi Pinterest

Sentimen	Jumlah	Persentase
Positif	650	76.47%
Negatif	200	23.53%
Netral	0 (dikeluarkan)	-
Total	850	100%

Distribusi sentimen yang tidak seimbang ini merupakan fenomena umum dalam analisis sentimen ulasan aplikasi, di mana pengguna yang puas cenderung lebih banyak memberikan ulasan dibandingkan pengguna yang tidak puas. Namun, hal ini juga menjadi tantangan tersendiri dalam pemodelan karena model cenderung bias terhadap kelas mayoritas. Penelitian pada analisis sentimen aplikasi X juga menemukan ketidakseimbangan serupa dengan proporsi positif mencapai 71.3%. Untuk mengatasi hal ini, penelitian selanjutnya dapat menerapkan teknik oversampling atau undersampling untuk menyeimbangkan dataset. Meskipun demikian, metode stratified split yang digunakan dalam penelitian ini telah memastikan bahwa proporsi kelas tetap terjaga pada data latih dan data uji, sehingga evaluasi model tetap representatif.

B. Evaluasi Performa Model Naive Bayes

Hasil evaluasi model menggunakan confusion matrix disajikan pada Gambar 1. Berdasarkan gambar tersebut, model berhasil memprediksi ulasan positif dengan benar sebanyak 105 data (True Positive) dan ulasan negatif dengan benar sebanyak 35 data (True Negative). Kesalahan prediksi terjadi pada 10 data positif yang diprediksi negatif (False Negative) dan 20 data negatif yang diprediksi positif (False Positive).



Gambar 1. Confusion Matrix Klasifikasi Model Naive Bayes
Dari confusion matrix, dihitung metrik evaluasi yang disajikan pada Tabel 9.

Tabel 9. Laporan Evaluasi Klasifikasi Model Naive Bayes

Metrik	Nilai (%)
Akurasi	84.00
Presisi	84.00
Recall	77.78
F1-Score	80.80

Berdasarkan Tabel 9, model mencapai akurasi sebesar 84.00%, yang berarti dari 170 data uji, model mampu mengklasifikasikan dengan benar sebanyak 140 ulasan. Nilai ini sejalan dengan penelitian yang mencapai akurasi 84%, serta lebih tinggi dibandingkan penelitian dengan akurasi 79.80%. Tingginya presisi (84.00%) menunjukkan bahwa prediksi positif model sangat akurat. Nilai recall (77.78%) yang lebih rendah mengindikasikan model masih melewatkan beberapa data positif, kemungkinan disebabkan oleh variasi kata dalam ulasan positif yang tidak tercakup dalam data latih. bahwa Naive Bayes cenderung memiliki recall lebih rendah pada kelas minoritas karena asumsi independensi yang kuat.

C. Analisis Cross-Validation

Untuk memastikan konsistensi performa model, dilakukan 10-fold cross-validation yang hasilnya disajikan pada Tabel 10.

Tabel 10. Hasil Cross-Validation 10-Fold

Metrik	Nilai
Rata-rata Akurasi	83.75%
Standar Deviasi	1.12%

Berdasarkan Tabel 10, rata-rata akurasi cross-validation adalah 83.75% dengan standar deviasi 1.12%. Nilai standar deviasi yang kecil menunjukkan performa model konsisten di berbagai subset data, sehingga model memiliki kemampuan generalisasi yang baik. Hasil ini

mendekati akurasi pada data uji (84.00%), menandakan model tidak mengalami overfitting. menunjukkan bahwa Naive Bayes memiliki stabilitas performa yang baik pada data teks karena sifat probabilistiknya yang tidak terlalu sensitif terhadap perubahan kecil dalam data.

Analisis Kata-Kata Penting

Analisis kata-kata yang paling berpengaruh dalam klasifikasi sentimen disajikan pada Tabel 11

Tabel 11. Kata-Kata yang Paling Mempengaruhi Sentimen

Sentimen Positif	Skor	Sentimen Negatif	Skor
bagus	2.45	error	2.82
membantu	2.38	crash	2.75
inspirasi	2.31	lambat	2.68
ide	2.28	iklan	2.60
keren	2.25	bug	2.55

Kata-kata seperti "bagus", "membantu", "inspirasi", dan "ide" mendominasi ulasan positif, mencerminkan aspek yang dihargai pengguna. Sementara itu, kata-kata seperti "error", "crash", "lambat", dan "iklan" mendominasi ulasan negatif. Temuan ini memberikan wawasan bagi pengembang Pinterest bahwa perbaikan stabilitas aplikasi dan pengurangan iklan dapat menjadi prioritas. menjelaskan bahwa kata-kata dengan konotasi positif secara kuat mengindikasikan sentimen positif dalam ulasan pengguna menyatakan bahwa identifikasi kata-kata kunci penting untuk menentukan area perbaikan produk.

Visualisasi Word Cloud



Gambar 2. Word Cloud Ulasan Positi



Gambar 3. Word Cloud Ulasan Negatif

Visualisasi word cloud untuk ulasan positif dan negatif disajikan pada Gambar 2 dan 3. Ulasan positif didominasi kata-kata berkonotasi baik seperti "bagus", "membantu", dan "inspirasi". Sebaliknya, ulasan negatif didominasi kata-kata seperti "error", "crash", dan "lambat". Word cloud memudahkan interpretasi data secara visual dan memberikan gambaran cepat tentang tema utama dalam ulasan pengguna. word cloud efektif untuk menyajikan frekuensi kata secara visual dan membantu mengidentifikasi pola dalam data teks. representasi visual data teks sangat membantu dalam analisis eksploratif.

Pengujian Prediksi Ulasan Baru

Model diuji dengan 5 sampel ulasan baru yang tidak termasuk dalam dataset. Hasilnya disajikan pada Tabel 12.

Tabel 12. Hasil Prediksi Ulasan Baru

No	Ulasan	Prediksi	Prob. Positif
1	Aplikasi Pinterest sangat membantu mencari ide desain	Positif	98.5%
2	Sering crash dan lambat, iklan terlalu banyak	Negatif	5.2%
3	Koleksi gambarnya keren dan inspiratif	Positif	97.8%
4	Aplikasi ini bagus tapi kadang loading lama	Positif	72.3%
5	Banyak bug dan error, sangat mengecewakan	Negatif	3.1%

Berdasarkan Tabel 12, model menunjukkan kemampuan yang baik dalam memprediksi sentimen ulasan baru [21]. Ulasan yang jelas positif (No. 1 dan 3) diprediksi dengan probabilitas tinggi (>97%). Ulasan yang jelas negatif (No. 2 dan 5) juga diprediksi dengan probabilitas tinggi (>94%). Ulasan dengan sentimen campuran (No. 4) diprediksi positif dengan probabilitas 72.3%, menunjukkan model mampu menangani ulasan dengan sentimen tidak ekstrem.. model klasifikasi yang baik harus memberikan probabilitas yang mencerminkan tingkat keyakinan prediksi.

KESIMPULAN

Berdasarkan hasil penelitian yang telah dilakukan, dapat disimpulkan bahwa penerapan analisis sentimen menggunakan algoritma Naive Bayes Classifier pada ulasan pengguna aplikasi Pinterest di Google Play Store berhasil mengklasifikasikan sentimen dengan performa yang baik, ditunjukkan oleh nilai akurasi sebesar 84.00%, presisi 84.00%, recall 77.78%, dan F1-Score 80.80%, yang membuktikan bahwa algoritma Naive Bayes efektif untuk menganalisis opini publik terhadap aplikasi berbasis visual berbahasa Indonesia. Hasil analisis menunjukkan bahwa sentimen positif mendominasi dengan proporsi 76.47%, yang mengindikasikan tingkat kepuasan pengguna yang relatif tinggi terhadap aplikasi Pinterest, sementara itu kata-kata seperti "error", "crash", dan "iklan" yang mendominasi ulasan negatif memberikan masukan berharga bagi pengembang untuk memprioritaskan perbaikan stabilitas aplikasi dan optimalisasi frekuensi iklan. Penelitian ini memberikan kontribusi pada bidang ilmu data mining dan pengolahan bahasa alami dengan mendemonstrasikan penggunaan Google Colaboratory sebagai platform komputasi cloud yang efisien dan terintegrasi untuk seluruh alur analisis sentimen, mulai dari pengumpulan data hingga evaluasi model, yang dapat diadopsi oleh peneliti dan praktisi dengan keterbatasan sumber daya komputasi untuk melakukan analisis serupa pada berbagai aplikasi maupun produk digital lainnya, serta menjadi referensi bagi pengembang aplikasi dalam memahami umpan balik pengguna secara sistematis dan berbasis data untuk mendukung pengambilan keputusan dalam pengembangan produk.

DAFTAR PUSTAKA

- Agustine Fitriana, Lady, Linawati, S., Herlinawati, N., Sa'adah, R., & Seimahuria, S. (2024). Analisis Sentimen Pengguna Twitter Terhadap Brand Indosat Menggunakan Metode Naïve Bayes Classifier. *JATI (Jurnal Mahasiswa Teknik Informatika)*, 8(3), 4291–4297. <https://doi.org/10.36040/jati.v8i3.9866>
- Aji Afani Maldini, M., & Andryana, S. (2025). Analisis Sentimen Ulasan Pengguna Aplikasi Perbankan Menggunakan Algoritma Support Vector Machine Dan Naive Bayes. *JATI (Jurnal Mahasiswa Teknik Informatika)*, 9(3), 4098–4105. <https://doi.org/10.36040/jati.v9i3.13522>
- Anjani, H. U., Vitriani, V., & Hastuti, M. (2024). Pemanfaatan Media Google Colaboratory Pada Mata Pelajaran Informatika di SMA Negeri 5 Pekanbaru besar dalam meningkatkan efektivitas pengajaran dan pembelajaran , khususnya dalam mata pengguna untuk menulis dan mengeksekusi kode Python secara daring melal. *Jurnal Ilmu Pendidikan (SOKO GURU)*, 4(1), 101–108.
- Cantika, C., Sitompul, M. D., & Wijaya, A. (2026). Analisis Ulasan Aplikasi Dalam Google

- Play Store Menggunakan Model Naive Bayes. *Jurnal Riset Multidisiplin Edukasi*, 3, 1–17. <https://journal.hasbaedukasi.co.id/index.php/jurmie>
- Fajar Ariwibowo, M., Swadaya Hidayat, I., Putri Andita, M., & Mawarindani Indra, A. (2024). Memanfaatkan Aplikasi Pinterest Sebagai Referensi Untuk Membuat Konten Digital Marketing. *Jurnal Pengabdian Masyarakat Nusantara (JPMN)*, 4(2), 192–196. <https://doi.org/10.35870/jpmn.v4i2.3231>
- Ginting, A., Purba, D., Sinaga, L. M., & Sagala, M. (2025). Analisis Sentimen Masyarakat Twitter terhadap Emas Digital Menggunakan Algoritma Naïve Bayes. *KAKIFIKOM (Kumpulan Artikel Karya Ilmiah Fakultas Ilmu Komputer)*, 07(01), 8–14. <https://ejournal.ust.ac.id/index.php/KAKIFIKOM/article/view/4909>
- Khainur, A. F. Y., Yuares, T., Fathurrohman, M. H., Widianingsih, & Rozikin, C. (2025). Analisis Komparatif Efektivitas Pipeline Data Cleaning Berbasis Aturan Dan Lemmatisasi Untuk Klasifikasi Sentimen. *Jurnal TIMES*, 14(2), 141–149. <https://doi.org/10.51351/jtm.14.2.2025890>
- Khotimah, A. K. (2024). Analisis Sentimen Terhadap Kualitas Pelayanan (Tinjauan Literatur). *JATI (Jurnal Mahasiswa Teknik Informatika)*, 8(3), 3044–3048.
- Putra Kurniawan, R., Istiadi, I., & Wahyu Iriananda, S. (2025). Analisis Sentimen Ulasan Aplikasi LinkedIn Berbasis Lexicon Dan Long Short-Term Memory (Lstm). *JATI (Jurnal Mahasiswa Teknik Informatika)*, 9(2), 2315–2324. <https://doi.org/10.36040/jati.v9i2.13042>
- Rossa Annjani, A., Indahyanti, U., Wahyu Azinar, A., & Dijaya, R. (2026). Penerapan Algoritma Machine Learning Untuk Memprediksi Sentimen Kepuasan Pengguna Weverse. *JATI (Jurnal Mahasiswa Teknik Informatika)*, 10(1), 784–790. <https://doi.org/10.36040/jati.v10i1.16874>
- Tarigan, P., Prabowo, A., & Govandy. (2024). Analisis Tingkat Akurasi Metode Naïve Bayes Dataset Breast Cancer Wisconsin. *Jatilima : Jurnal Multimedia Dan Teknologi Informasi*, 06(02), 134–142.
- Uddin, B., Dewi, D., Basam, F., Mila, E., & Wijayadi, L. (2024). Efektivitas Pemanfaatan Pinterest Terhadap Kreativitas Pengguna. 7(4), 678–684.
- Zainuddin Siregar, M., Marwan Elhanafi, A., & Irwan, D. (2025). Analisis Sentimen Terhadap Ulasan Aplikasi Media Sosial Di Google Play Menggunakan Algoritma Naive Bayes. *JATI (Jurnal Mahasiswa Teknik Informatika)*, 9(2), 3373–3381. <https://doi.org/10.36040/jati.v9i2.12841>