

draft-jurnal-kelompok-jen-revisi- final_1765804619783.pdf

By Turnitin Acc

Analisis Sentimen Ulasan Pengguna Aplikasi X (Twitter) pada Google Play Store Menggunakan Komparasi Algoritma Support Vector Machine, K-Nearest Neighbor, dan Naive Bayes

Jim Maxwell Manampiring¹, Jefiri Za², Jenianus Halawa³, Waeisul Bismi⁴, Ika Kurniawati⁵, Rizal Fahlap⁶

Abstrak

Transformasi *rebranding* media sosial Twitter menjadi X menimbulkan respons yang beragam dari pengguna global, termasuk di Indonesia. Ulasan pengguna pada platform Google Play Store memuat informasi berharga mengenai kepuasan dan keluhan pengguna, namun volume data yang besar dan tidak terstruktur menyulitkan analisis secara manual. Studi ini difokuskan pada penerapan analisis sentimen guna mengklasifikasi opini pengguna menjadi kategori positif dan negatif, serta membandingkan kinerja tiga algoritma *Machine Learning*, yaitu *Support Vector Machine* (SVM), *K-Nearest Neighbor* (KNN), dan *Naive Bayes*. Dataset yang digunakan berjumlah 10.000 data berbahasa Indonesia yang dikumpulkan melalui *scraping*. Melalui tahapan *preprocessing* yang meliputi *cleaning*, *tokenizing*, dan *stemming*, data dilatih dengan pembagian rasio 80:20. Hasil pengujian menunjukkan algoritma SVM menggunakan kernel linear menghasilkan kinerja terbaik dengan akurasi sebesar 84,8%, diikuti oleh Naive Bayes sebesar 83,1%, dan KNN sebesar 79,7%. Kesimpulan dari studi ini menegaskan bahwa SVM merupakan metode yang paling efisien guna menangani klasifikasi teks pada data ulasan aplikasi X yang memiliki dimensi tinggi, meskipun terdapat ketidakseimbangan kelas pada dataset.

Kata Kunci: Analisis Sentimen, X (Twitter), Support Vector Machine, Naive Bayes, KNN.

Abstract

The *rebranding* transformation of the social media platform Twitter to X has elicited diverse responses from global users, including those in Indonesia. User reviews on the Google Play Store contain valuable information regarding user satisfaction and complaints; however, the large volume of unstructured data complicates manual analysis. This study focuses on applying sentiment analysis to classify user opinions into positive and negative categories, as well as comparing the performance of three *Machine Learning* algorithms: *Support Vector Machine* (SVM), *K-Nearest Neighbor* (KNN), and *Naive Bayes*. The dataset consists of 10,000 Indonesian-language reviews collected via *scraping* techniques. Following preprocessing stages—which include cleaning, tokenizing, and stemming—the data was trained using an 80:20 split ratio. Test results indicate that the SVM algorithm with a linear kernel achieved the best performance with an accuracy of 84.8%, followed by Naive Bayes at 83.1%, and KNN at 79.7%. This study concludes that SVM is the most efficient method for handling text classification on high-dimensional review data of the X application, despite the presence of class imbalance in the dataset.

Keywords: Sentiment Analysis, X (Twitter), Support Vector Machine, Naive Bayes, KNN, Text Mining.

1. PENDAHULUAN

Transformasi fundamental media sosial Twitter menjadi "X" pada pertengahan tahun 2023 merupakan salah satu peristiwa *rebranding* terbesar dalam industri teknologi. Perubahan ini tidak hanya meliputi pergantian nama dan logo, tetapi juga perombakan kebijakan fitur yang memicu respons masif dari pengguna global, termasuk di Indonesia. Beberapa penelitian terdahulu mencatat bahwa *rebranding* ini menimbulkan polemik di kalangan pengguna setia, di mana sentimen negatif cenderung mendominasi akibat ketidaksiapan pengguna terhadap perubahan antarmuka dan fungsionalitas aplikasi [1], [2]. Reaksi publik ini terekam secara *real-time* melalui ulasan (*reviews*) pada kanal penyedia aplikasi seluler resmi berbasis Android, yakni Google Play Store [3], [4].

Ulasan di Google Play Store mengandung informasi vital yang dapat digunakan pengembang untuk mengevaluasi kepuasan pengguna dan mendeteksi permasalahan

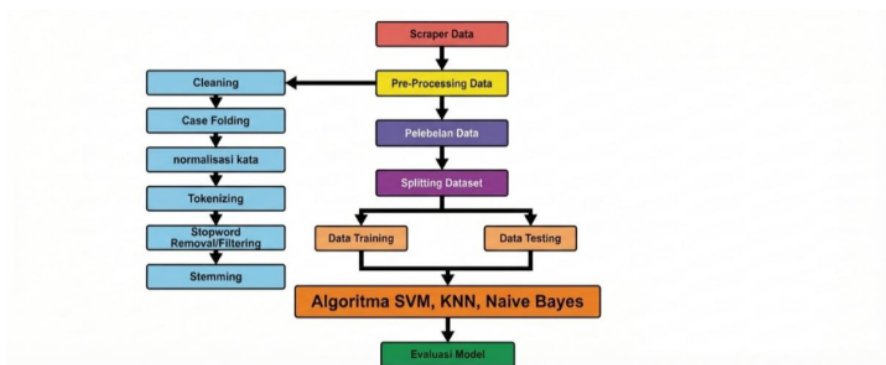
teknis (*bug*). Namun, dengan jutaan pengguna aktif, volume data ulasan yang masuk sangat besar dan tidak terstruktur. Analisis manual terhadap data berskala besar ini menjadi tidak efisien dan rentan terhadap subjektivitas [11]. Oleh karena itu, penerapan *Sentiment Analysis* (Analisis Sentimen) berbasis *Machine Learning* menjadi solusi krusial untuk mengotomatisasi klasifikasi opini publik ke dalam kategori positif dan negatif secara akurat dan efisien [7].

Dalam ranah klasifikasi teks, algoritma *Supervised Learning* seperti *Support Vector Machine* (SVM), *K-Nearest Neighbor* (KNN), dan *Naive Bayes* sering menjadi standar utama dalam studi komparasi performa [5], [6]. Ketiga algoritma ini memiliki karakteristik unik dalam menangani data teks. SVM dikenal memiliki performa yang *robust* dan akurasi tinggi pada data berdimensi besar sekalipun distribusinya tidak seimbang [8], [9]. Di sisi lain, *Naive Bayes* (khususnya varian Multinomial) diunggulkan karena efisiensi komputasinya dalam mengolah data berbasis frekuensi kata [10], sedangkan KNN menawarkan pendekatan berbasis jarak yang sederhana meskipun sering menghadapi tantangan komputasi pada dataset teks yang bersifat *sparse* [12].

Meskipun banyak studi telah membahas analisis sentimen, penelitian yang membandingkan secara langsung kinerja SVM, KNN, dan *Naive Bayes* secara spesifik pada kasus *rebranding* aplikasi X dengan dataset berbahasa Indonesia masih perlu diperdalam untuk menemukan model yang paling optimal. Merujuk pada permasalahan tersebut, riset ini bertujuan mengkomparasi performa ketiga algoritma menggunakan 10.000 data ulasan dari Google Play Store. Melalui studi ini, kontribusi yang ingin dicapai adalah memberikan rekomendasi metode klasifikasi terbaik yang adaptif terhadap karakteristik data ulasan aplikasi media sosial pasca-*rebranding*.

2. METODE PENELITIAN

Tahapan penelitian ini dirancang secara sistematis untuk memastikan validitas hasil analisis. Alur penelitian berawal dari tahap *collecting data*, *preprocessing*, pelabelan fitur pemisahan, serta penerapan model menggunakan algoritma Machine Learning, hingga evaluasi kinerja model.



Gambar 1. Tahapan Penelitian Analisis Sentimen Aplikasi X

6

2.1. Pengumpulan Data (Data Collection)

Data penelitian **dikumpulkan** menggunakan **teknik Web Scraping** pada platform Google Play Store. Metode ini umum digunakan untuk mengakuisisi ulasan pengguna secara massal guna analisis sentimen [5]. Objek penelitian adalah ulasan pengguna pada aplikasi "X" (sebelumnya Twitter) dengan paket aplikasi com.twitter.android. Jumlah dataset yang digunakan sebanyak 10.000 ulasan berbahasa Indonesia. Atribut data yang diambil meliputi *username*, tanggal ulasan (*date*), rating bintang (*score*), dan teks ulasan (*content*). Data disimpan dalam format CSV (*Comma Separated Values*) untuk pengolahan lebih lanjut.

Tabel 1. Sample Rating Data dari Google Play

	Review ID	Username	Rating	Review Text	Date
0	66d44f63-8907-44c8-b07a-06ebc2e6f913	Pengguna Google	2	X sakarang bener ² ga nyaman, setiap mau buka a...	2025-11-30 05:19:07
1	28cdb2d1-57ee-405d-a5bb-3802f1a68925	Pengguna Google	1	aplikasi taik tiba-tiba scroll GC bisa masuk n...	2025-11-30 05:05:54
2	13938a78-5053-47f4-a983-0e673c2f2e1b	Pengguna Google	1	apalah mau daftar dikit ² "kesalahan teknis" na...	2025-11-30 05:05:04
3	ab3dc132-8c4e-4e93-b950-13ce619b373e	Pengguna Google	1	gak bisa Download	2025-11-30 04:40:36
4	ee3cfae8-e81d-4609-85f0-b7eed29282a0	Pengguna Google	1	kenapa setiap login pasti gagal karena kesalah...	2025-11-30 04:37:14
5	2f053073-beba-4dff-877b-4730f96aaba0	Pengguna Google	5	terbaik	2025-11-30 04:30:21
6	fcbec9ef-973e-4c9f-b7ed-09c466861efa	Pengguna Google	1	gabisa masuk akunnya! JELEK BANGET	2025-11-30 04:15:13

	Review ID	Username	Rating	Review Text	Date
7	8a28b0f6-5c67-4ff9-85f0-c8c62dc2e59c	Pengguna Google	1	error terus ni aplikasi	2025-11-30 04:05:53
8	9965afb4-19e4-43c1-a6f5-8a1e726c706d	Pengguna Google	5	Aplikasi sangat bagus, Sangat membantu untuk m...	2025-11-30 04:01:06
9	a092e4bf-0060-4dc3-a5c1-f0efbdd7b819	Pengguna Google	1	kenapa akun saya di tangguhkan,sebab nya apa? ...	2025-11-30 03:14:26
10	75caec27-d485-47c4-b1c7-893ce61e6a16	Pengguna Google	5	ok MANTAP dah gan	2025-11-30 03:10:08

Implementasi sistem dilakukan menggunakan bahasa pemrograman Python 3.10 pada lingkungan Google Colab. Pustaka utama yang digunakan meliputi *Pandas* untuk manipulasi data, *NLTK* dan *Sastrawi* untuk pemrosesan teks, serta *Scikit-Learn* untuk pemodelan algoritma Machine Learning.

2.2. Pra-pemrosesan Data (*Data Preprocessing*)

Data mentah hasil *scraping* memiliki banyak *noise* dan format yang tidak terstruktur. Tahap ini bertujuan untuk membersihkan dan mempersiapkan data agar dapat diproses oleh algoritma. Tahapan *text mining* yang tepat sangat berpengaruh terhadap akurasi klasifikasi [10]. Langkah-langkahnya meliputi:

1. **Pembersihan Data (*Cleaning*):** Mengeliminasi elemen teks yang dianggap *noise* (gangguan) meliputi angka, puntuasi (tanda baca), serta simbol, *emoji*, *username* (@), *hashtag* (#), dan URL.
2. **Penghapusan Duplikat:** Menghapus data ulasan yang memiliki isi teks identik untuk menghindari bias pada model.
3. **Pengubahan Huruf (*Case Folding*):** mengonversi format teks dengan mentransformasi semua alfabet menjadi huruf kecil (*lowercase*) untuk menyeragamkan data.
4. **Normalisasi Kata:** Mengubah kata-kata non baku, singkatan, dan bahasa sehari-hari (*slang*) menjadi bentuk formal yang merujuk pada pedoman Bahasa Indonesia menggunakan kamus leksikon yang telah disusun.
5. **Stopword Removal:** Menghapus kata sambung atau kosakata frekuensi tinggi yang sering muncul namun tidak memiliki makna sentimen signifikan (seperti

"dan", "yang", "di", "dari") menggunakan pustaka NLTK (*Natural Language Toolkit*).

6. **Stemming:** Mengembalikan kata yang mengandung afiks (imbuhan) ke bentuk asalnya (kata dasar). menggunakan pustaka Sastrawi, yang dirancang khusus untuk morfologi Bahasa Indonesia.

14

2.3. Pelabelan Data (*Data Labeling*)

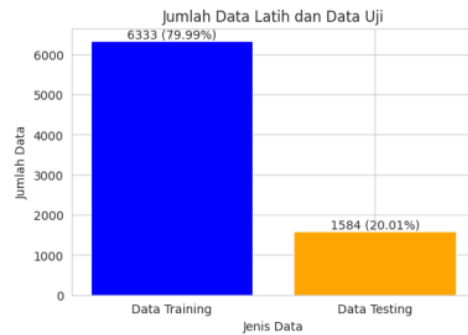
Pelabelan sentimen dilakukan secara otomatis dengan memanfaatkan atribut *Rating* (skala 1-5) yang diberikan pengguna [15]. Skema pelabelan ditentukan sebagai berikut:

- **Sentimen Negatif:** Ulasan dengan *Rating* 1, 2, dan 3. *Rating* 3 dikategorikan ke dalam kelas negatif karena dalam konteks kepuasan pengguna aplikasi, nilai 'cukup' atau 'netral' seringkali mengindikasikan adanya aspek yang belum memenuhi ekspektasi pengguna atau keluhan pasif yang memerlukan perbaikan, sehingga lebih relevan dikelompokkan sebagai sentimen yang tidak mendukung (non-positif).
- **Sentimen Positif:** Ulasan dengan *Rating* 4 dan 5. Kategori ini dipilih karena merepresentasikan tingkat kepuasan yang tinggi (*user satisfaction*) dan penerimaan yang baik terhadap fungsionalitas aplikasi, di mana pengguna cenderung memberikan apresiasi atau tidak memiliki keluhan signifikan yang mengganggu pengalaman penggunaan.

Proses ini mengubah masalah regresi menjadi klasifikasi biner (dua kelas).

2.4. Pembagian Data (*Data Splitting*)

Dataset yang telah bersih dan berlabel dibagi menjadi dua bagian menggunakan teknik *Hold-out Validation*. Dataset dipartisi dengan proporsi 80:20, rinciannya 80% dialokasikan sebagai himpunan latih (*Training Set*) guna pembelajaran model, sementara 20% sisanya difungsikan sebagai himpunan uji (*Testing Set*) untuk evaluasi akurasi. Pembagian dilakukan secara *stratified* untuk memelihara rasio label data sentimen yang seimbang di dalam *training set* dan *testing set*. [13]



Gambar 2. Hasil Data Splitting

2.5. Ekstraksi Fitur (*Feature Extraction*)

Sebelum masuk ke tahap pemodelan, informasi tekstual dikonversi ke dalam format angka (vektor) melalui pendekatan *Bag of Words* (BoW) atau *Count Vectorizer*. Metode ini menghitung frekuensi kemunculan setiap kata dalam dokumen untuk membentuk matriks fitur yang dapat diproses oleh mesin.

Secara matematis, jika dimisalkan D adalah sebuah dokumen dan $V=\{w_1, w_2, \dots, w_n\}$ adalah daftar kosakata (*vocabulary*) unik dari seluruh korpus data, maka representasi vektor x untuk dokumen D didefinisikan pada Persamaan:

$$x_D = [c(w_1), c(w_2), \dots, c(w_n)]$$

Di mana $c(w_i)$ merepresentasikan frekuensi kemunculan kata ke- i (w_i) di dalam dokumen D . Jika kata tersebut tidak muncul dalam dokumen, maka nilainya adalah 0.

2.6. Metode Klasifikasi

Penelitian ini membandingkan kinerja tiga algoritma *Supervised Machine Learning*:

1. **Support Vector Machine (SVM)**: Menggunakan *kernel Linear*. Mekanisme utama algoritma ini adalah menemukan *hyperplane* ideal yang berfungsi untuk memisahkan dua kelas sentimen dengan margin maksimal [16].

Fungsi keputusan pada *hyperplane* linear didefinisikan dalam Persamaan:

$$f(x) = w^T x + b$$

Dengan w merepresentasikan vektor bobot (*weight*), x sebagai input vektor, dan b merupakan nilai bias. Kelas diprediksi berdasarkan tanda positif atau negatif dari hasil fungsi tersebut. *Kernel linear* digunakan karena efisiensinya yang efektif saat menghadapi data tekstual berdimensi kompleks [4].

2. **K-Nearest Neighbor (KNN)**: Menggunakan parameter $k=5$. Algoritma ini mengklasifikasikan input baru merujuk pada dominasi label dari 5 tetangga terdekat (*nearest neighbors*) di dalam ruang vektor [12]. Perhitungan jarak

kedekatan antar data menggunakan rumus *Euclidean Distance* seperti pada Persamaan:

$$d(x, y) = \sqrt{\sum_{i=1}^n (x_i - y_i)^2}$$

Di mana $d(x, y)$ adalah jarak antara data uji x dan data latih y , sedangkan n adalah jumlah fitur.

3. **Naive Bayes:** Menggunakan varian *Multinomial Naive Bayes*. Metode ini berlandaskan prinsip Teorema Bayes yang memegang asumsi independensi antar fitur, yang sangat cocok untuk data berupa frekuensi kata diskrit [2]. Klasifikasi dilakukan dengan menghitung probabilitas tertinggi [14], [20] (Maximum A Posteriori) menggunakan Persamaan:

$$P(c|x) = \frac{P(x|c) \cdot P(c)}{P(x)}$$

Di mana $P(c|x)$ adalah peluang kelas c (Positif/Negatif) jika diberikan fitur x , $P(x|c)$ adalah *likelihood*, dan $P(c)$ adalah probabilitas prior kelas.

2.7. Evaluasi Kinerja (Evaluation)

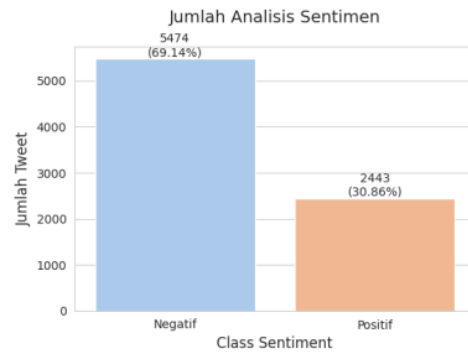
Kinerja model dievaluasi menggunakan Confusion Matrix untuk menghitung empat metrik utama [17]. **Akurasi** digunakan untuk mengukur rasio prediksi benar terhadap keseluruhan data. **Presisi** mengukur ketepatan prediksi positif dibandingkan dengan total data yang diprediksi positif, sedangkan **Recall** mengukur kemampuan sistem mengenali seluruh data positif yang sebenarnya. Terakhir, **F1-Score** digunakan sebagai rata-rata harmonik antara presisi dan recall untuk memberikan gambaran performa yang lebih objektif pada dataset yang tidak seimbang (imbalanced) [19].

3. HASIL DAN PEMBAHASAN

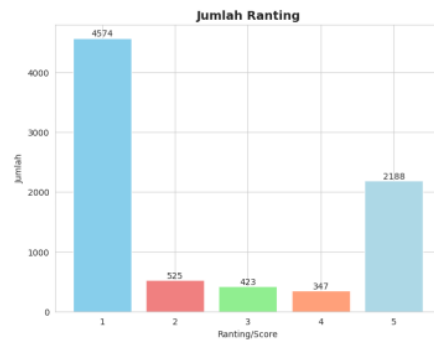
Segmen ini mendeskripsikan luaran eksperimen, bermula dari tahapan statistik deskriptif dataset hingga evaluasi kinerja algoritma Support Vector Machine (SVM), K-Nearest Neighbor (KNN), dan Naive Bayes.

3.1. Hasil Preprocessing

Berdasarkan hasil pelabelan otomatis menggunakan rating bintang, ditemukan bahwa distribusi sentimen pada dataset ulasan aplikasi X cenderung tidak seimbang (imbalanced) [15].



Gambar 3. Distribusi Persentase Kelas Sentimen Negatif Dan Positif

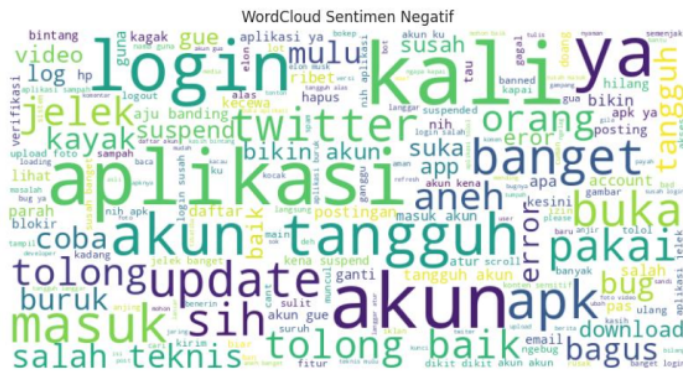


Gambar 4. Distribusi Jumlah Ulasan Berdasarkan Rating

Dari total 10.000 data ulasan yang dikumpulkan, sentimen negatif mendominasi secara signifikan. Hal ini mengindikasikan bahwa pasca-rebranding dari Twitter menjadi X, pengguna cenderung lebih aktif menyuarakan ketidakpuasan atau keluhan dibandingkan apresiasi. ¹⁹ Temuan ini sejalan dengan penelitian terdahulu yang menyatakan bahwa perubahan drastis pada identitas aplikasi sering memicu resistensi pengguna [2]. Analisis visual menggunakan Word Cloud dilakukan untuk mengidentifikasi topik dominan pada setiap kelas sentiment [10].



Gambar 5. Wordcloud Sentimen positif



Gambar 6. Wordcloud Sentimen negatif

Dalam visualisasi *Word Cloud* kategori negatif, kata kunci yang mendominasi meliputi "suspend", "susah", "akun", "tangguh", dan "aplikasi". Hal ini menunjukkan bahwa sumber utama sentimen negatif bukan sekadar perubahan *rebranding*, melainkan permasalahan fungsionalitas aplikasi. Pengguna banyak mengeluhkan ketidakstabilan sistem yang menyulitkan akses masuk (*login*), serta penegakan aturan yang ketat berupa penangguhan akun (*suspend*) yang dinilai mengganggu pengalaman pengguna. Sebaliknya, pada sentimen positif, kata kunci yang muncul cenderung bersifat umum seperti "bagus", "keren", "suka", dan "mantap", yang menunjukkan kepuasan pengguna bersifat general tanpa menyoroti fitur spesifik.

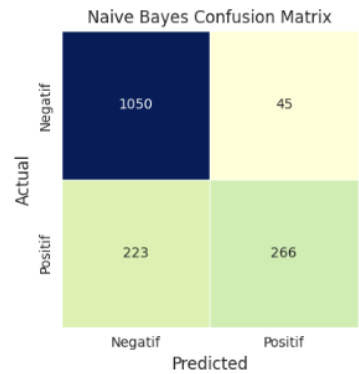
3.2. Evaluasi Kinerja Model

Pengujian dilakukan menggunakan data uji (testing set) sebesar 20% dari total dataset. Evaluasi kinerja ketiga algoritma (Naive Bayes, K-Nearest Neighbor, dan Support Vector Machine) diukur menggunakan *Confusion Matrix*. Metode ini memberikan representasi visual dan kuantitatif mengenai detail prediksi yang benar (*True*) dan salah (*False*) untuk kelas Positif dan Negatif [17]. Total jumlah keseluruhan sampel testing yang dipakai dalam evaluasi ini berjumlah 1.584 data.

Berikut adalah rincian hasil evaluasi untuk masing-masing algoritma:

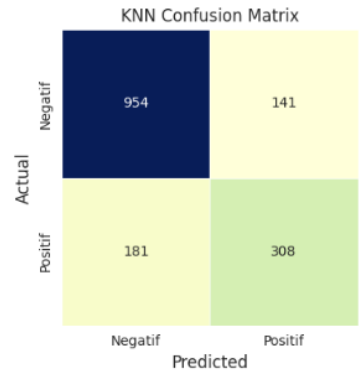
3.2.1. Naive Bayes Merujuk pada data hasil eksperimen, algoritma Naive Bayes menunjukkan kemampuan yang sangat baik dalam mendeteksi kelas Negatif. Model ini berhasil memprediksi 1.050 data secara tepat sebagai Negatif (True Negative) dan 266 data sebagai Positif (True Positive). Akan tetapi, terjadi misklasifikasi di mana 45 sampel negatif salah dikenali menjadi positif (*False Positive*), serta sejumlah besar ulasan positif yang terdeteksi sebagai negatif, yaitu sebanyak 223 data (False Negative). Hal ini

menunjukkan Naive Bayes cenderung konservatif dalam memprediksi kelas positif pada dataset [11].



Gambar 7. Hasil Confusion Matrix Naïve Bayes

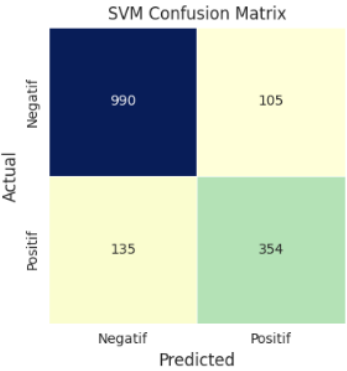
3.2.2. K-Nearest Neighbor (KNN) Pada algoritma KNN, hasil Confusion Matrix menunjukkan sedikit penurunan pada prediksi kelas negatif dibandingkan Naive Bayes. KNN memprediksi 954 data dengan benar sebagai Negatif (True Negative) dan 308 data sebagai Positif (True Positive). Tingkat kesalahan pada model ini terlihat dari adanya 141 data yang salah diprediksi sebagai positif (False Positive) dan 181 data yang salah diprediksi sebagai negatif (False Negative). Meskipun jumlah True Positive lebih tinggi daripada Naive Bayes, tingginya False Positive mengindikasikan model ini memiliki tingkat "alarm palsu" yang lebih tinggi [12].



Gambar 8. Hasil Confusion Matrix KNN

3.2.3. Support Vector Machine (SVM) Algoritma SVM menunjukkan kinerja yang paling seimbang dan optimal di antara ketiga model. SVM mencatatkan 990 data True Negative dan 354 data True Positive. Keunggulan utama SVM terlihat pada kemampuannya meminimalkan kesalahan, dengan hanya 105 data False Positive dan 135 data False

Negative. Jumlah False Negative yang paling rendah dibandingkan dua model lainnya (135 banding 181 dan 223) menunjukkan bahwa SVM adalah model yang paling sensitif dan akurat dalam mengenali sentimen positif tanpa mengorbankan terlalu banyak akurasi pada sentimen negative [1].



Gambar 9. Hasil Confusion Matrix SVM

3.3. Analisis Classification Report

Selain menggunakan Confusion Matrix, evaluasi kinerja model dilakukan secara mendalam menggunakan Classification Report untuk mengukur nilai Precision, Recall, dan F1-Score pada setiap kelas sentimen. Hal ini penting untuk melihat apakah model bekerja dengan baik pada kedua kelas (Positif dan Negatif) atau cenderung bias ke salah satu kelas saja akibat ketidakseimbangan data (imbalanced dataset) [8].

3.3.1. Algoritma Support Vector Machine (SVM) SVM dengan kernel Linear menunjukkan performa paling optimal dan seimbang.

Akurasi Tertinggi: Model ini mencapai akurasi global sebesar 84,8%.

Keseimbangan (Robustness): Capaian metrik F1-Score sebesar 0.89 didapatkan pada klasifikasi sentimen negatif, diikuti oleh kelas positif 0.75. Meskipun terdapat ketimpangan jumlah data (Support: 1095 data negatif vs 489 data positif), SVM mampu mempertahankan Recall yang cukup tinggi pada kelas minoritas (Positif) sebesar 0.724. Ini menunjukkan SVM tidak hanya "menebak" kelas mayoritas saja, sesuai dengan karakteristik SVM yang tangguh pada dimensi tinggi [4] [16].

Tabel 2. Classification Report for SVM:

	precision	recall	f1-score	support
Negatif	0.880	0.904	0.892	1095.000

	29			
	precision	recall	f1-score	support
Positif	0.771	0.724	0.747	489.000
accuracy	0.848	0.848	0.848	0.848
macro avg	0.826	0.814	0.819	1584.000
weighted avg	0.846	0.848	0.847	1584.000

3.3.2. Algoritma *Multinomial Naive Bayes* menempati peringkat kedua dengan perolehan akurasi sebesar 83,1%, namun menunjukkan adanya anomali signifikan berupa bias terhadap kelas mayoritas. Hal ini terlihat dari ketimpangan nilai *Recall* di mana model mampu mendeteksi sentimen negatif dengan sangat baik (0,959), tetapi memiliki kinerja yang buruk pada *Recall* sentimen positif yang hanya mencapai 0,544. Kondisi ini mengindikasikan bahwa model cenderung agresif dalam memprediksi kelas negatif, sehingga menyebabkan tingginya angka *False Negative* di mana hampir separuh ulasan positif gagal dideteksi dan salah diklasifikasikan sebagai negatif. Oleh karena itu, model ini dinilai kurang direkomendasikan apabila tujuan utama sistem adalah untuk menangkap *feedback* positif pengguna secara akurat [5].

Tabel 3. Classification Report for Naive Bayes:

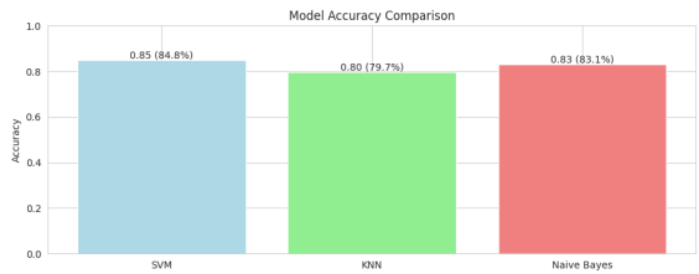
	precision	recall	f1-score	support
Negatif	0.825	0.959	0.887	1095.000
Positif	0.855	0.544	0.665	489.000
accuracy	0.831	0.831	0.831	0.831
macro avg	0.840	0.751	0.776	1584.000
weighted avg	0.834	0.831	0.818	1584.000

3.3.3. Algoritma K-Nearest Neighbor (KNN) dengan parameter $k=5$ menghasilkan kinerja terendah dengan akurasi 79,7%. **Kelemahan pada Dimensi Tinggi:** Rendahnya nilai *Precision* dan *Recall* pada kedua kelas (dibandingkan dua model lainnya) mengonfirmasi bahwa pendekatan berbasis jarak (*distance-based*) seperti KNN kurang

efektif menangani data teks yang memiliki dimensi fitur sangat besar (*sparse matrix*) [9], [18].

Tabel 4. Classification Report for KNN:

	precision	recall	f1-score	support
Negatif	0.841	0.871	0.856	1095.000
Positif	0.686	0.630	0.657	489.000
accuracy	0.797	0.797	0.797	0.797
macro avg	0.763	0.751	0.756	1584.000
weighted avg	0.793	0.797	0.794	1584.000



Gambar 10 Hasil akurasi model masing-masing algoritma

4. SIMPULAN

Berdasarkan hasil penelitian mengenai analisis sentimen terhadap *rebranding* aplikasi X, dapat disimpulkan bahwa respons publik didominasi oleh sentimen negatif yang mencerminkan tingginya ketidakpuasan pengguna terhadap stabilitas sistem, perubahan fitur, dan kendala teknis pasca-pembaruan. Dari sisi komparasi algoritma, metode *Support Vector Machine* (SVM) berbasis kernel Linear terbukti menjadi model klasifikasi yang paling efektif dan stabil dibandingkan *Naive Bayes* maupun *K-Nearest Neighbor* (KNN). Keunggulan SVM terletak pada ketahanannya (*robustness*) dalam menangani karakteristik data teks berdimensi tinggi serta kemampuannya menjaga keseimbangan prediksi pada dataset yang tidak seimbang, sedangkan algoritma *Naive Bayes* cenderung bias ke arah kelas mayoritas dan KNN memiliki keterbatasan dalam memetakan fitur yang jarang (*sparse*). Oleh karena itu, SVM direkomendasikan sebagai pendekatan terbaik

untuk menyelesaikan permasalahan klasifikasi opini publik pada karakteristik data ulasan aplikasi media sosial.

ORIGINALITY REPORT

12%

SIMILARITY INDEX

PRIMARY SOURCES

1	greenpub.org Internet	73 words — 2%
2	ejournal.undip.ac.id Internet	28 words — 1%
3	ejournal.polbeng.ac.id Internet	25 words — 1%
4	journal.ubpkarawang.ac.id Internet	20 words — 1%
5	Raisya Nadzira Zahirah Shafa, Asti Herliana. "ANALISIS SENTIMEN PENGGUNA APLIKASI SAPAWARGA - JABAR SUPER APPS MENGGUNAKAN ALGORITMA SUPPORT VECTOR MACHINE", Djtechno: Jurnal Teknologi Informasi, 2025 Crossref	17 words — 1%
6	ojs.stmik-banjarbaru.ac.id Internet	17 words — 1%
7	Shazifa Azhari, Nining Rahaningsih, Raditya Danar Dana, Mulyawan .. "PENINGKATAN AKURASI ANALISIS SENTIMEN PADA APLIKASI LOKLOK DENGAN METODE NAÏVE BAYES", Jurnal Informatika dan Teknik Elektro Terapan, 2025 Crossref	15 words — < 1%

8	academic-accelerator.com Internet	15 words — < 1%
9	jurnal.polibatam.ac.id Internet	15 words — < 1%
10	id.scribd.com Internet	14 words — < 1%
11	Dwi Nanda Agustia, Ryan Randy Suryono. "Comparison of Naïve Bayes, Random Forest, and Logistic Regression Algorithms for Sentiment Analysis Online Gambling", INOVTEK Polbeng - Seri Informatika, 2025 Crossref	13 words — < 1%
12	jutif.if.unsoed.ac.id Internet	12 words — < 1%
13	ijcs.net Internet	11 words — < 1%
14	Sebastian Ubaidillah Royan, Nana Suarna, Irfan Ali, Dodi Solihudin. "ANALISIS SENTIMEN ULASAN PRODUK SKINCARE DI SHOPEE UNTUK MENINGKATKAN KUALITAS PRODUK MENGGUNAKAN METODE SUPPORT VECTOR MACHINE", Jurnal Informasi dan Komputer, 2025 Crossref	10 words — < 1%
15	repository.uinsu.ac.id Internet	10 words — < 1%
16	Mohammad Bayu Anggara. "Comparison of Naïve Bayes and SVM Methods in Sentiment Analysis of User Reviews on the RSUD AL IHSAN Mobile Application", Competitive, 2025 Crossref	9 words — < 1%

17 Windy Livia Azzahra, Evangs Mailoa. "Analisis Sentimen terhadap RSUD Salatiga Menggunakan SVM dan TF-IDF", Jurnal Indonesia : Manajemen Informatika dan Komunikasi, 2025 9 words — < 1%
Crossref

18 jurnal.instiki.ac.id 9 words — < 1%
Internet

19 purnomowati.blogspot.com 9 words — < 1%
Internet

20 Abu Thaariq Nur Diwongso, Ubaid Fauzan, Yusritsan Ghinan Elfatih, Sopiyan Apandi. "Analisis Klasterisasi Indikator Makroekonomi Indonesia Menggunakan K-Means", RIGGS: Journal of Artificial Intelligence and Digital Business, 2025 8 words — < 1%
Crossref

21 Adryan Syah, Firman Nurdiyansyah, Aviv Yuniar Rahman. "ANALISIS SENTIMEN APLIKASI SHOPEE, TOKOPEDIA, LAZADA DAN BLIBLI MENGGUNAKAN LEKSIKON DAN RANDOM FOREST", Jurnal Informatika dan Teknik Elektro Terapan, 2024 8 words — < 1%
Crossref

22 Dany Pratmanto, Fabriyan Fandi Dwi Imaniawan. "Analisis Sentimen Terhadap Aplikasi Canva Menggunakan Algoritma Naive Bayes Dan K-Nearest Neighbors", Computer Science (CO-SCIENCE), 2023 8 words — < 1%
Crossref

23 Mhd Arief Hasan, Novia Bimby. "Analisis Sentimen Publik Terhadap Kenaikan Pajak PPN di Indonesia Tahun 2024 Menggunakan Algoritma Machine Learning", JURNAL FASILKOM, 2025 8 words — < 1%
Crossref

24 Rizki Aulia Putra, Rice Novita, Tengku Khairil Ahsyar, Zarnelly. "Implementation of Classification Algorithm for Sentiment Analysis: Measuring App User Satisfaction", Teknika, 2024

8 words — < 1%

Crossref

25 ejurnal.ung.ac.id

Internet

8 words — < 1%

26 Ahmad Fathan Hidayatullah, Aufa Aulia Fadila, Kiki Purnama Juwairi, Royan Abida Nayoan. Jurnal Linguistik Komputasional (JLK), 2019

7 words — < 1%

Crossref

27 IIP IMRON DAIPAH, RINI ASTUTI, WILLY PRIHARTONO. "PREDIKSI CHURN PELANGGAN PADA LAYANAN DESAIN GRAFIS HOME DESAIN MENGGUNAKAN ALGORITMA NAÏVE BAYES", Jurnal Informatika dan Teknik Elektro Terapan, 2025

7 words — < 1%

Crossref

28 Fastabiquul Khusna, Khothibul Umam, Siti Nur'aini, Maya Rini Handayani. "EVALUASI EFEKTIVITAS SUPPORT VECTOR MACHINE DAN RANDOM FOREST DALAM KLASIFIKASI ULASAN PENGGUNA APLIKASI STREAMING VIDIO", JSil (Jurnal Sistem Informasi), 2025

6 words — < 1%

Crossref

29 Muhammad Arif Arkan, Farhan Nul Hakim, Ghiyats Al Robbani, Meta Windyawati Windyawati, Agun Afiransah Afriansah. "ANALISIS DISPARITAS PEMBANGUNAN MELALUI INTEGRASI MACHINE LEARNING PADA DATA SPASIAL DAN TEMPORAL", JUTECH : Journal Education and Technology, 2025

6 words — < 1%

Crossref

EXCLUDE QUOTES	OFF	EXCLUDE SOURCES	OFF
EXCLUDE BIBLIOGRAPHY	OFF	EXCLUDE MATCHES	OFF