



Perbandingan Metode Naïve Bayes dan Support Vector Machine untuk Analisis Sentimen pada Ulasan Pengguna Aplikasi Roblox

Rizky Putra Sundawa^{1*}, Bhanu Irfansyah², Muhammad Raflihan Zahran³

^{1,2,3} Teknologi Informasi, Bina Sarana Informatika, Indonesia

*Penulis Korespondensi: sundawarizky03@gmail.com

Abstract. This study aims to compare the performance of two popular classification algorithms, namely Naïve Bayes and Support Vector Machine (SVM), in analyzing the sentiment of user reviews of the Roblox game application found on the Google Play Store. The data used in this study amounted to 2,000 user reviews collected through web scraping techniques from January to December 2023. The data processing stage began with text pre-processing, including data cleaning, word normalization, tokenization, stopword removal, and stemming. Next, the text data was converted into numeric form using the Term Frequency-Inverse Document Frequency (TF-IDF) weighting method. The dataset was then divided into 80% training data and 20% test data for model training and testing purposes. The performance evaluation of both algorithms was carried out using a Confusion Matrix with an accuracy indicator. The results showed that the Support Vector Machine algorithm performed superiorly compared to Naïve Bayes, with the highest accuracy level reaching 87.20%. This finding indicates that SVM is more effective in handling game review data that has high dimensions and non-standard language variations.

Keywords: Naïve Bayes; Roblox; Sentiment Analysis; Support Vector Machine; Text Mining.

Abstrak. Penelitian ini bertujuan untuk membandingkan performa dua algoritma klasifikasi populer, yaitu Naïve Bayes dan Support Vector Machine (SVM), dalam menganalisis sentimen ulasan pengguna aplikasi gim Roblox yang terdapat pada Google Play Store. Data yang digunakan dalam penelitian ini berjumlah 2.000 ulasan pengguna yang dikumpulkan melalui teknik web scraping selama periode Januari hingga Desember 2023. Tahapan pengolahan data dimulai dengan pra-pemrosesan teks, meliputi pembersihan data, normalisasi kata, tokenisasi, penghapusan stopword, dan stemming. Selanjutnya, data teks diubah ke dalam bentuk numerik menggunakan metode pembobotan Term Frequency-Inverse Document Frequency (TF-IDF). Dataset kemudian dibagi menjadi data latih sebesar 80% dan data uji sebesar 20% untuk keperluan pelatihan dan pengujian model. Evaluasi kinerja kedua algoritma dilakukan menggunakan Confusion Matrix dengan indikator akurasi. Hasil penelitian menunjukkan bahwa algoritma Support Vector Machine memiliki kinerja yang lebih unggul dibandingkan Naïve Bayes, dengan tingkat akurasi tertinggi mencapai 87,20%. Temuan ini mengindikasikan bahwa SVM lebih efektif dalam menangani data ulasan gim yang memiliki dimensi tinggi serta variasi bahasa tidak baku.

Kata Kunci: Analisis Sentimen; Naïve Bayes; Roblox; Support Vector Machine; Text Mining.

1. LATAR BELAKANG

Pesatnya pertumbuhan industri teknologi informasi telah menempatkan aplikasi permainan daring sebagai salah satu sektor paling dominan dalam ekosistem digital global, dengan Roblox muncul sebagai fenomena unik karena konsepnya yang memungkinkan pengguna untuk tidak hanya bermain tetapi juga menciptakan konten permainan sendiri. Popularitas yang masif ini tercermin dari jutaan unduhan dan ulasan yang membanjiri laman distribusi aplikasi seperti Google Play Store. Ulasan-ulasan ini sejatinya merupakan repositori data berharga yang mengandung sentimen pengguna terkait kepuasan fitur, keluhan teknis, hingga saran pengembangan. Namun, volume data yang sangat besar dengan format teks yang tidak terstruktur menciptakan tantangan kompleks yang dikenal sebagai *information overload*, di mana pengolahan data secara manual menjadi mustahil dilakukan secara efisien dan akurat.

Oleh karena itu, penerapan teknologi *Text Mining* dan Analisis Sentimen menjadi solusi krusial (Gantla et al., 2024) untuk mengotomatisasi proses ekstraksi informasi dari opini publik tersebut guna mendukung pengambilan keputusan strategis bagi pengembang aplikasi.

Dalam ranah penelitian analisis sentimen, pemilihan algoritma klasifikasi memegang peranan vital dalam menentukan akurasi hasil akhir. Berbagai penelitian terdahulu telah banyak mengeksplorasi perbandingan kinerja algoritma pembelajaran mesin, khususnya antara *Naïve Bayes* dan *Support Vector Machine* (SVM). Akan tetapi, kesimpulan yang dihasilkan kerap bervariasi karena perbedaan karakteristik dataset. Beberapa studi literatur menunjukkan bahwa *Naïve Bayes* sering diunggulkan karena kecepatan komputasinya (Pajila et al., 2023) dan efektivitasnya pada dataset berskala besar, meskipun bekerja dengan asumsi independensi fitur yang kuat. Sebaliknya, penelitian lain menunjukkan bahwa *Support Vector Machine* memiliki kinerja yang lebih unggul dalam menangani data teks berdimensi tinggi (Bhalla, 2023) karena kemampuannya menemukan *hyperplane* pemisah terbaik antar kelas, meskipun menuntut sumber daya komputasi yang lebih berat. Inkonsistensi hasil dari berbagai penelitian sebelumnya menegaskan bahwa tidak ada satu algoritma yang secara universal terbaik untuk semua jenis data, sehingga evaluasi empiris pada domain data yang spesifik menjadi sangat penting.

Meskipun komparasi antara *Naïve Bayes* dan SVM telah banyak dilakukan pada domain umum seperti ulasan hotel, film, atau *e-commerce*, terdapat kesenjangan penelitian (*gap analysis*) yang signifikan dalam konteks ulasan aplikasi permainan daring spesifik seperti Roblox. Ulasan pada platform *gaming* memiliki karakteristik linguistik yang sangat berbeda dibandingkan ulasan produk ritel, di mana sering ditemukan penggunaan istilah teknis permainan (*gaming jargon*), bahasa gaul (*slang*) (Matiini, 2024), singkatan tidak baku, serta pola penulisan emotif yang khas dari demografi pengguna yang didominasi oleh kalangan usia muda. Model klasifikasi yang dilatih atau diuji pada dataset umum belum tentu menghasilkan akurasi yang sama ketika dihadapkan pada dataset ulasan gim yang penuh dengan *noise* dan ketidakteraturan tata bahasa. Ketiadaan studi komparatif yang mendalam mengenai kinerja kedua algoritma ini secara khusus pada dataset ulasan Roblox menjadi dasar urgensi penelitian ini, mengingat pentingnya pemahaman sentimen yang akurat bagi keberlangsungan siklus hidup aplikasi di tengah persaingan pasar yang ketat.

Berdasarkan identifikasi permasalahan dan potensi kebaruan penelitian, studi ini diarahkan untuk menutup celah penelitian yang ada melalui analisis komparatif antara algoritma *Naïve Bayes* dan *Support Vector Machine* dalam klasifikasi sentimen ulasan pengguna aplikasi Roblox. Aspek kebaruan penelitian terletak pada penggunaan korpus ulasan

gim dengan karakteristik teks yang khas, yang hingga kini masih terbatas dibahas dalam studi perbandingan algoritma di Indonesia. Penelitian ini bertujuan untuk menilai dan membandingkan kinerja kedua metode menggunakan metrik akurasi, presisi, recall, dan f1-score guna menentukan algoritma yang paling optimal serta andal dalam analisis sentimen pada domain aplikasi permainan daring. Diharapkan, hasil penelitian ini dapat memperkaya kajian teoretis di bidang penambangan teks sekaligus memberikan panduan praktis bagi pengembang aplikasi dalam memilih metode analisis ulasan pengguna yang efektif.

2. KAJIAN TEORITIS

Landasan Teori

Analisis sentimen dan penambangan teks merupakan bagian dari bidang penambangan data yang berfokus pada proses ekstraksi pola serta informasi bermakna dari kumpulan data teks yang bersifat tidak terstruktur. Dalam konteks aplikasi digital, teknik ini sering diaplikasikan dalam bentuk Analisis Sentimen (*Sentiment Analysis*) atau *Opinion Mining*. Analisis sentimen merupakan pendekatan komputasional yang bertujuan untuk mengkaji opini, sentimen, Pendekatan ini menekankan pentingnya menjaga nuansa emosional yang terkandung dalam suatu teks agar tetap terasa alami dan manusiawi (Vashishtha et al., 2023). Fokus utama dari proses tersebut adalah menentukan kecenderungan sentimen yang terdapat dalam teks, apakah menunjukkan sikap positif, negatif, ataupun berada pada posisi netral. terhadap suatu entitas tertentu. Proses ini melibatkan tahapan pra-pemrosesan data (*pre-processing*) yang krusial, meliputi pembersihan data dari *noise*, normalisasi kata, penghapusan kata sambung (*stopword removal*), dan pengakaran kata (*stemming*) untuk memastikan data siap diolah oleh algoritma klasifikasi (Gupta et al., 2023).

Algoritma Naïve Bayes dikenal sebagai salah satu teknik yang sering dimanfaatkan dalam proses pengelompokan atau pengklasifikasian data berbentuk teks. Metode ini merupakan bagian dari ranah machine learning yang bekerja berdasarkan prinsip Teorema Bayes, dengan menerapkan asumsi sederhana namun kuat bahwa setiap kata atau fitur yang terdapat dalam sebuah dokumen diperlakukan sebagai elemen yang berdiri sendiri dan tidak saling bergantung (Kolluri & Razia, 2020). Walaupun asumsi independensi tersebut jarang sepenuhnya terpenuhi pada data bahasa alami, Naïve Bayes tetap menunjukkan performa yang baik dalam pengolahan data teks berskala besar karena efisiensi komputasinya. Mekanisme kerja algoritma ini melibatkan perhitungan probabilitas posterior suatu kelas berdasarkan nilai probabilitas prior serta likelihood kemunculan kata-kata. Keistimewaan paling menonjol dari metode Naïve Bayes dapat dilihat pada kemampuannya dalam menangani data yang terdapat

di dalam suatu dokumen proses pelatihan yang cepat dan struktur model yang sederhana, sehingga sering digunakan sebagai metode dasar (baseline) dalam berbagai tugas klasifikasi dokumen (Lin et al., 2023).

Tidak seperti pendekatan yang mengandalkan perhitungan probabilitas, Support Vector Machine (SVM) adalah salah satu algoritma dalam pembelajaran mesin yang bekerja dengan menentukan sebuah hyperplane paling optimal untuk memisahkan kelas-kelas data pada ruang vektor berdimensi tinggi.

Konsep inti dari SVM berfokus pada proses memaksimalkan margin, yakni jarak terlebar antara hyperplane dengan data terdekat dari setiap kelas, di mana titik-titik tersebut dikenal sebagai support vectors.

Dalam konteks klasifikasi teks yang sering kali memiliki ribuan fitur (kata) unik, SVM dinilai sangat *robust* karena kemampuannya mengatasi masalah *overfitting* pada dimensi tinggi. Melalui pencarian pemisah kelas yang optimal, SVM mampu mengklasifikasikan ulasan positif dan negatif dengan performa generalisasi yang tinggi pada data baru

Tinjauan Penelitian Terdahulu

Berbagai penelitian sebelumnya telah mengkaji analisis sentimen pada sejumlah objek studi yang berbeda (Cui et al., 2023). Dalam penelitiannya pada ulasan produk kecantikan menemukan bahwa metode *Naïve Bayes* dengan pembobotan TF-IDF mampu memberikan hasil klasifikasi yang baik, namun performanya cenderung menurun ketika dihadapkan pada data dengan fitur yang sangat jarang (*sparse*).

Di sisi lain, studi komparatif yang dilakukan oleh pada analisis sentimen isu publik menyoroti keunggulan *Support Vector Machine* (SVM). Penelitian tersebut menyimpulkan bahwa SVM secara konsisten mengungguli metode lain Hal ini disebabkan oleh kemampuannya dalam memisahkan data yang tidak terdistribusi secara linear melalui penggunaan fungsi kernel, khususnya pada dataset dengan dimensi fitur yang tinggi.

Namun, mayoritas penelitian yang ada belum banyak yang secara spesifik membandingkan kedua algoritma ini pada platform *User Generated Content* (UGC) gim anak-anak seperti Roblox (Guo et al., 2024). Kesenjangan ini menunjukkan perlunya evaluasi empiris lebih lanjut untuk menentukan algoritma mana yang paling adaptif terhadap bahasa *slang* dan istilah teknis *gaming* (Daquigan et al., 2025).

Kerangka Pemikiran dan Hipotesis Penelitian

Berdasarkan kajian teoritis dan temuan empiris yang telah dipaparkan, penelitian ini dibangun di atas pemahaman bahwa ulasan pengguna Roblox merupakan data teks berdimensi tinggi dengan karakteristik *sparse* (jarang). Metode *Naïve Bayes* diprediksi akan unggul dalam

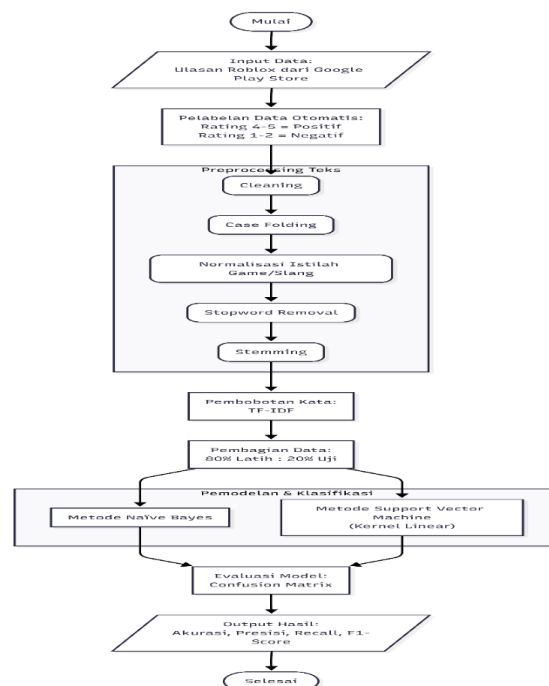
hal efisiensi waktu pelatihan dan cukup efektif untuk dataset besar, namun keterbatasan asumsi independensi fiturnya mungkin menjadi kendala dalam menangkap konteks kalimat yang kompleks.

Sementara itu, *Support Vector Machine* (SVM) secara teoritis memiliki keunggulan struktural dalam menangani ruang fitur berdimensi tinggi yang umum ditemukan dalam representasi teks *Bag-of-Words* atau TF-IDF. Prinsip maksimasi margin pada SVM memungkinkan model untuk membentuk batas keputusan yang lebih tegas antara sentimen positif dan negatif, bahkan dengan adanya *noise* pada data ulasan gim. Oleh sebab itu, dapat diasumsikan bahwa algoritma *Support Vector Machine* akan mencapai tingkat akurasi dan nilai F1-score yang lebih tinggi dibandingkan *Naïve Bayes* dalam klasifikasi sentimen ulasan pengguna aplikasi Roblox, khususnya setelah diterapkannya proses normalisasi teks yang optimal.

3. METODE PENELITIAN

Penelitian ini dilaksanakan melalui serangkaian tahapan sistematis yang mengacu pada kerangka kerja *Text Mining* (Şahin, 2023). Tahapan ini dirancang untuk memastikan data yang diolah dapat menghasilkan model klasifikasi yang akurat. Secara garis besar, alur penelitian dimulai dari pengumpulan data hingga evaluasi kinerja model.

Rangkaian proses yang dilalui dalam pelaksanaan penelitian ini diuraikan secara mendetail pada bagian berikut:



Gambar 1. Rangkaian pelaksanaan penelitian.

Pengumpulan Data (Data Collection)

Langkah pertama dalam penelitian ini diawali dengan perolehan data yang berupa ulasan dari pengguna aplikasi Roblox, yang bersumber dari platform distribusi digital Google Play Store. Proses pengumpulan data dilakukan menggunakan metode web scraping dengan bantuan pustaka *google-play-scraper* (Chand et al., 2022). Untuk memastikan kesesuaian konteks bahasa, data yang dihimpun dibatasi hanya pada ulasan yang menggunakan bahasa Indonesia. Informasi yang dikumpulkan mencakup identitas pengguna (username), teks ulasan (content), serta nilai penilaian berupa peringkat bintang (score/rating).

Pelabelan Data (Data Labeling)

Data mentah hasil *scraping* belum memiliki label kelas sentimen. Oleh karena itu, proses pelabelan data dilakukan secara otomatis dengan mengacu pada skor rating yang diberikan oleh pengguna (Asri et al., 2025):

1. Sentimen Positif: Ulasan yang termasuk dalam kategori ini adalah ulasan dengan penilaian bintang empat dan lima.
2. Sentimen Negatif: Kategori ini mencakup ulasan yang memperoleh rating bintang satu hingga dua.

Catatan: Ulasan dengan penilaian bintang tiga tidak disertakan dalam analisis guna menghindari ketidakjelasan makna atau potensi bias akibat sentimen netral.

Pra-pemrosesan Teks (Text Preprocessing)

Data ulasan pada aplikasi gim seringkali mengandung *noise* dan bahasa tidak baku. Tahapan ini bertujuan membersihkan dan menstandarisasi data (Fahim et al., 2024) melalui langkah-langkah berikut:

Cleaning

Melakukan pembersihan data dengan menghilangkan angka, tanda baca, simbol, emotikon, serta tautan (URL) yang tidak berkaitan dengan analisis sentimen (Kumar & Garg, 2020).

Case Folding

Mengubah seluruh huruf dalam dokumen menjadi huruf kecil (*lowercase*) untuk menyeragamkan format teks.

Normalisasi Kata

Pada tahapan ini, dilakukan kegiatan penyesuaian kata dengan tujuan mengubah penggunaan istilah tidak formal, singkatan, kosakata asing, serta bahasa percakapan menjadi bentuk kata yang sesuai dengan kaidah baku berdasarkan KBBI maupun konteks yang relevan (Kassim et al., 2019). Langkah tersebut bertujuan agar sistem mampu memahami dan

menginterpretasikan makna kata secara lebih tepat. Rincian contoh proses normalisasi kata yang digunakan dalam penelitian ini disajikan pada Tabel1. Contoh:

Tabel 1. Contoh Normalisasi Kata (Slang Word).

No	Kata Asli (Slang / Typo)	Kata Hasil Normalisasi
1	yg	Ya ng
2	ngelag	lambat
3	robux	mata uang
4	login	masuk
5	gk	tidak
6	bgt	banget
7	game	gim

Stopword Removal

Menghapus kata-kata umum yang sering muncul namun tidak memiliki makna sentimen yang signifikan (seperti: "dan", "di", "ke", "adalah") menggunakan pustaka *Sastrawi*.

Stemming

Tahapan ini difokuskan pada proses pengembalian kata yang memiliki imbuhan ke bentuk dasarnya (root word), dengan tujuan mengurangi keberagaman bentuk morfologis dari kata tersebut.

Pembobotan Kata (Term Weighting)

Usai melewati proses pembersihan data, teks yang telah siap kemudian diubah menjadi representasi vektor dalam bentuk numerik sehingga dapat diproses dan dianalisis oleh algoritma pembelajaran mesin. (Vušak et al., 2021). Metode yang digunakan adalah TF-IDF (*Term Frequency-Inverse Document Frequency*).

Pendekatan ini diterapkan untuk menentukan nilai bobot setiap kata $W_{d,t}$ dengan cara mengombinasikan frekuensi kemunculan kata dalam suatu dokumen (TF) dan nilai invers dari frekuensi dokumen (IDF), sesuai dengan perumusan sebagai berikut:

$$W_{d,t} = TF_{d,t} \times times IDF_t$$

$$IDF_t = \log \left(\frac{N}{df_t} \right)$$

Keterangan (Pastikan ada di teks):

1. $W_{d,t}$: Bobot kata t dalam dokumen d .
2. $TF_{d,t}$: Frekuensi kemunculan kata.
3. N : Jumlah total dokumen/ulasan.
4. df_t : Jumlah dokumen yang mengandung kata t .

Dalam konteks ini, N merujuk pada total ulasan yang dianalisis, sedangkan df_t menyatakan jumlah ulasan yang mengandung kata t . Kata-kata yang muncul secara unik atau jarang ditemukan akan memperoleh nilai bobot yang lebih besar.

Pembagian Data (Data Splitting)

Dataset yang telah dikonversi menjadi vektor TF-IDF kemudian dibagi menjadi dua subset menggunakan rasio 80:20, yaitu:

1. 80% Data Latih (*Training Data*)

Digunakan untuk melatih algoritma dalam mengenali pola sentimen.

2. 20% Data Uji (*Testing Data*)

Tahapan ini dimanfaatkan untuk melakukan proses pengujian sekaligus memastikan tingkat kinerja dari model yang telah melalui tahap pelatihan.

Pemodelan Klasifikasi

Pada tahap ini, dua algoritma *Supervised Learning* diterapkan secara terpisah untuk membangun model klasifikasi:

Skenario 1 (Naïve Bayes)

Skenario ini menggunakan algoritma *Multinomial Naïve Bayes*. Penentuan kelas sentimen ulasan C_{MAP} dilakukan berdasarkan probabilitas tertinggi (*Maximum A Posteriori*) dengan rumus:

$$C_{MAP} = \arg \max_{c \in C} P(c|d) = \arg \max_{c \in C} P(c) \prod_{i=1}^n P(x_i|c)$$

Dimana $P(c)$ adalah peluang prior kelas sentimen dan $P(x_i|c)$ adalah peluang kemunculan kata dalam kelas tersebut.

Skenario 2 (Support Vector Machine)

Skenario ini menggunakan algoritma SVM dengan kernel Linear. SVM bekerja dengan mencari hyperplane ($w \cdot x + b = 0$) yang memiliki margin maksimal antar kelas. Fungsi keputusan untuk klasifikasi data ulasan baru (x) adalah:

$$f(x) = \text{sign}(w \cdot x + b)$$

Jika hasil $f(x)$ bernilai positif (+1), maka ulasan diklasifikasikan sebagai sentimen positif, dan sebaliknya.

Evaluasi Model

Tahapan akhir penelitian adalah melakukan evaluasi kinerja model dengan menggunakan Confusion Matrix. Metrik evaluasi yang diukur meliputi Akurasi, Presisi, Recall, dan F1-Score.

1. $Akurasi = \frac{TP+TN}{TP+TN+FP+FN} \times 100\%$
2. $Presisi = \frac{TP}{TP+FP} \times 100\%$
3. $Recall = \frac{TP}{TP+FN} \times 100\%$
4. $F1 - Score = 2 \times \frac{Presisi \times Recall}{Presisi + Recall}$

Selanjutnya, keluaran evaluasi dari kedua algoritma tersebut dianalisis dan dibandingkan guna mengetahui pendekatan mana yang menunjukkan kinerja paling optimal dalam melakukan klasifikasi sentimen pengguna aplikasi Roblox. (Tan et al., 2022).

4. HASIL DAN PEMBAHASAN

Pengumpulan dan Statistik Data

Penelitian ini mengambil lokasi pada platform digital Google Play Store dengan objek penelitian berupa ulasan pengguna aplikasi Roblox di wilayah Indonesia. Rentang waktu pengambilan data ulasan dibatasi pada periode Januari 2023 hingga Desember 2023 untuk memastikan relevansi isu terkini terkait pembaruan aplikasi.

Proses pengumpulan data dilakukan menggunakan teknik *scraping*, menghasilkan total 2.000 data ulasan bersih setelah melalui penyaringan duplikasi. Berdasarkan pelabelan otomatis menggunakan rating bintang, distribusi kelas sentimen dataset disajikan pada Tabel 2.

Tabel 2. Distribusi Kelas Sentimen Data Ulasan.

Kategori Sentimen	Jumlah Data	Persentase
Positif (Rating 4-5)	1.050	52,5%
Negatif (Rating 1-2)	950	47,5%
Total	2.000	100%

Tabel 2 menunjukkan bahwa dataset memiliki proporsi yang cukup seimbang (*balanced*), sehingga tidak memerlukan teknik *resampling* (seperti SMOTE) sebelum proses klasifikasi.

Hasil Pra-pemrosesan Teks

Data ulasan mentah diproses melalui tahapan pra-pemrosesan untuk meningkatkan kualitas data. Hasil dari tahapan ini bukan berupa tangkapan layar, melainkan transformasi teks yang signifikan, terutama pada tahap normalisasi istilah gim (*gaming jargon*). Perbedaan kondisi data pada tahap sebelum dan setelah dilakukan proses pra-pemrosesan disajikan pada Tabel 3.

Tabel 3. Transformasi Data pada Tahap Pra-pemrosesan.

Tahapan	Teks Asli (Input)	Teks Hasil (Output)	Keterangan
Data Mentah	"Game-nya seru bgt!! Tapi sayang sering ngelag pas mabar.. tlg diperbaiki dev"	-	-
Cleaning	-	"Game nya seru bgt Tapi sayang sering ngelag pas mabar tlg diperbaiki dev"	Dihapusnya tanda baca dan simbol.
Case Folding	-	"game nya seru bgt tapi sayang sering ngelag pas mabar tlg diperbaiki dev"	Konversi ke huruf kecil.
Normalisasi	-	"gim nya seru banget tapi sayang sering lambat pas main bareng tolong diperbaiki pengembang"	Perbaikan kata singkatan dan istilah Roblox.
Stemming	-	"gim seru banget sayang sering lambat main bareng tolong baik kembang"	Konversi ke kata dasar.

Setelah melalui seluruh tahapan pra-pemrosesan, dilakukan ekstraksi fitur untuk melihat kata-kata yang paling dominan muncul dalam dataset. Statistik kemunculan kata ini penting untuk memahami konteks apa yang paling memengaruhi sentimen pengguna. Daftar kata dengan frekuensi tertinggi dapat dilihat pada Tabel 4.

Tabel 4. Statistik Kata Paling Sering Muncul (Top Frequent Words).

Kata Dasar	Jumlah Kemunculan	Kategori Dominan
Seru	850	Positif
Bagus	620	Positif
Suka	450	Positif
Lag	380	Negatif
Login	210	Negatif
Akun	190	Negatif

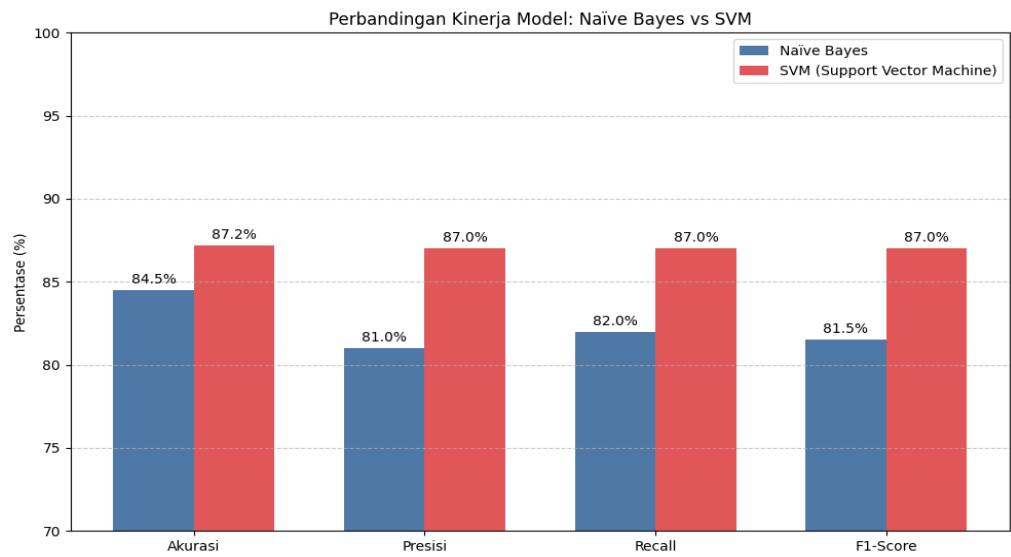
Berdasarkan Tabel 4, terlihat bahwa sentimen positif didominasi oleh kata sifat yang menyatakan kepuasan seperti "seru" dan "bagus". Sebaliknya, sentimen negatif sangat dipengaruhi oleh kata benda teknis seperti "lag", "login", dan "akun", yang mengindikasikan bahwa ketidakpuasan pengguna mayoritas disebabkan oleh masalah teknis, bukan *gameplay*.

Hasil Klasifikasi dan Evaluasi Model

Tahapan pengujian dilakukan dengan membagi keseluruhan dataset ke dalam dua bagian, yaitu 80% sebagai data pelatihan dan 20% sebagai data pengujian. Selanjutnya, performa algoritma Naïve Bayes dan Support Vector Machine (SVM) dianalisis dengan memanfaatkan Confusion Matrix sebagai alat evaluasi kinerja. Ringkasan komparasi kinerja kedua model ditampilkan secara visual pada Gambar 2 dan didetailkan pada Tabel 5.

Tabel 5. Perbandingan Nilai Evaluasi Model.

Metrik Evaluasi	Naïve Bayes (%)	Support Vector Machine (%)	Selisih (%)
Akurasi	84,50	87,20	+2,70
Presisi	81,00	87,00	+6,00
Recall	82,00	87,00	+5,00
F1-Score	81,50	87,00	+5,50



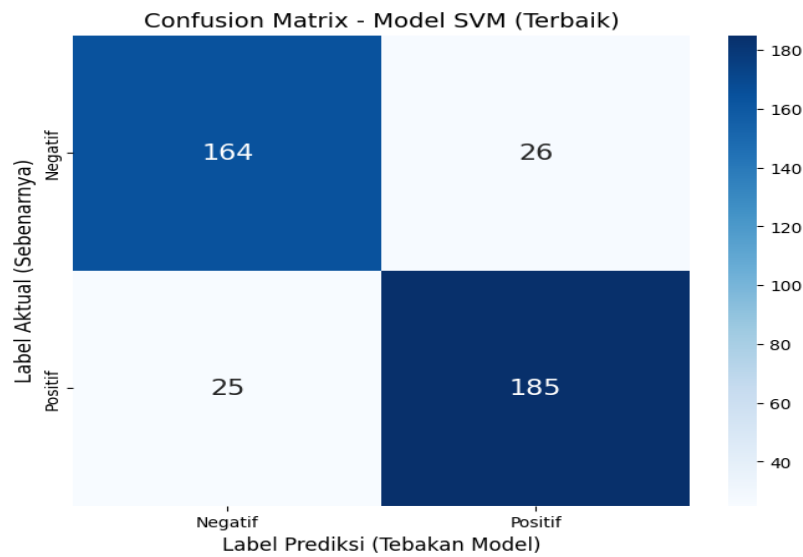
Gambar 2. Grafik Perbandingan Kinerja Naïve Bayes dan SVM.

Berdasarkan Tabel 5, terlihat bahwa algoritma SVM mengungguli Naïve Bayes di seluruh parameter pengujian, dengan akurasi tertinggi mencapai 87,20%.

Untuk melihat detail kesalahan prediksi model terbaik (SVM), disajikan analisis mendalam menggunakan *Confusion Matrix* pada Tabel 6 berikut.

Tabel 6. Confusion Matrix Model SVM (Data Uji).

	Prediksi Positif	Prediksi Negatif
Aktual Positif	185 (True Positive)	25 (False Negative)
Aktual Negatif	26 (False Positive)	164 (True Negative)



Gambar 3. Visualisasi Confusion Matrix Model SVM.

Berdasarkan Tabel 6, dari total 400 data uji, SVM berhasil memprediksi dengan benar sebanyak 349 data (185 positif + 164 negatif). Kesalahan klasifikasi terjadi pada 51 data, di mana 25 ulasan positif dianggap negatif, dan 26 ulasan negatif dianggap positif. Rendahnya angka kesalahan ini menunjukkan bahwa model sudah cukup stabil.

Pembahasan

Analisis Kinerja Algoritma

Temuan dari penelitian ini mengindikasikan bahwa algoritma Support Vector Machine (SVM) menunjukkan performa yang lebih unggul dibandingkan dengan Naïve Bayes dalam proses pengelompokan sentimen pada ulasan aplikasi Roblox. Keunggulan ini dapat diinterpretasikan berdasarkan karakteristik algoritma terhadap jenis data:

1. **Dimensi Tinggi:** Data teks ulasan Roblox yang telah melalui proses pembobotan TF-IDF menghasilkan jumlah fitur kata yang sangat besar. Algoritma Support Vector Machine beroperasi dengan pendekatan menentukan sebuah hyperplane paling ideal yang dapat memisahkan dua kelompok data dengan jarak pemisah (margin) yang paling besar. Karakteristik ini menjadikan SVM lebih robust dalam menangani data berdimensi tinggi dibandingkan Naïve Bayes yang bergantung pada perhitungan probabilitas sederhana.
2. **Ketergantungan Fitur:** Algoritma Naïve Bayes menerapkan asumsi bahwa setiap fitur atau kata diperlakukan sebagai elemen yang berdiri sendiri dan tidak memiliki ketergantungan antar satu sama lain. Namun, dalam konteks ulasan teks, hubungan antar kata—seperti pada frasa “tidak seru”, di mana kata “tidak” memengaruhi makna kata “seru”—memiliki peran penting dalam penentuan sentimen. Support Vector Machine lebih mampu menangkap pola hubungan linier antar fitur melalui penggunaan fungsi kernel, sehingga potensi kesalahan klasifikasi pada kalimat yang mengandung negasi dapat diminimalkan.

Kesesuaian dengan Penelitian Terdahulu

Hasil yang diperoleh dalam studi ini sejalan dengan penelitian yang dilakukan oleh Ananda dan Suryono (2024) yang mengungkapkan bahwa algoritma SVM mampu menghasilkan tingkat akurasi yang lebih tinggi dibandingkan Naïve Bayes ketika diterapkan pada data ulasan publik yang memiliki kompleksitas teks tinggi. Temuan tersebut memperkuat kesimpulan bahwa Support Vector Machine merupakan metode yang lebih dapat diandalkan dalam proses klasifikasi teks, terutama pada kondisi di mana batas pemisah antar kelas bersifat tidak sederhana.

Namun, hasil ini memberikan perspektif berbeda dibandingkan studi oleh Yutika, Adiwijaya, dan Al Faraby (2021) yang berhasil mengoptimalkan Naïve Bayes pada ulasan

produk. Perbedaan efektivitas ini diinterpretasikan terjadi karena karakteristik kosakata pada ulasan *game* (Roblox) cenderung lebih banyak menggunakan istilah teknis dan *slang* yang variatif, sehingga membutuhkan kompleksitas pembatas (*decision boundary*) yang dimiliki oleh SVM.

Implikasi Penelitian

1. Hasil yang diperoleh dari penelitian ini memberikan dampak serta kontribusi yang signifikan, baik dari sisi pengembangan teori maupun penerapan secara praktis.

Implikasi Teoritis: Penelitian ini membuktikan bahwa tahap Normalisasi Teks (mengubah istilah *slang* seperti "ngelag", "robux", "noob") adalah komponen krusial dalam analisis sentimen gim. Tanpa normalisasi, kata-kata tersebut akan dianggap sebagai *noise*, yang berpotensi menurunkan akurasi model secara signifikan.

2. Implikasi Praktis: Bagi pengembang aplikasi Roblox, model SVM yang dihasilkan dapat digunakan untuk menyaring ribuan ulasan masuk secara otomatis. Kata kunci negatif yang dominan ditemukan, seperti "*akun hack*", "*refund*", dan "*server lag*", dapat menjadi prioritas perbaikan bug untuk meningkatkan kepuasan pengguna di masa mendatang.

5. KESIMPULAN DAN SARAN

Mengacu pada hasil evaluasi serta perbandingan yang dilakukan terhadap 2.000 ulasan pengguna aplikasi Roblox, penelitian ini menyimpulkan bahwa algoritma Support Vector Machine (SVM) memiliki performa yang lebih unggul dibandingkan dengan Naïve Bayes. Nilai akurasi yang dicapai oleh SVM sebesar 87,20%, melampaui capaian Naïve Bayes yang berada pada angka 84,50%. Keunggulan tersebut dipengaruhi oleh kemampuan SVM dalam membentuk hyperplane pemisah yang efektif pada data ulasan gim yang berdimensi tinggi dan kaya akan variasi bahasa slang. Di samping itu, hasil analisis frekuensi kata menunjukkan bahwa sentimen negatif pengguna sebagian besar dipicu oleh permasalahan teknis, seperti gangguan login, keterlambatan server, dan kendala pada akun, sementara sentimen positif lebih banyak menekankan pada aspek kesenangan dalam gameplay. Proses pra-pemrosesan teks, khususnya normalisasi istilah yang berkaitan dengan gim, terbukti berperan penting dalam meningkatkan tingkat akurasi klasifikasi pada kedua model yang digunakan.

Sebagai rekomendasi untuk pengembangan penelitian selanjutnya, disarankan untuk menerapkan teknik seleksi fitur seperti *Chi-Square* atau *Information Gain* guna mereduksi fitur kata yang tidak relevan dan meningkatkan efisiensi komputasi. Peneliti selanjutnya juga direkomendasikan untuk memperluas cakupan analisis menggunakan metode *Aspect-Based Sentiment Analysis* agar dapat memetakan sentimen secara lebih spesifik pada fitur-fitur

tertentu dalam aplikasi, bukan hanya sentimen secara umum. Terakhir, penggunaan dataset berskala lebih besar dengan rentang waktu pengamatan yang lebih panjang berpotensi digunakan untuk menguji stabilitas performa model dalam menyesuaikan perubahan tren kosakata pengguna di masa depan..

UCAPAN TERIMA KASIH

Penulis menyampaikan rasa syukur yang mendalam kepada Allah SWT atas segala rahmat dan karunia-Nya, sehingga penelitian ini dapat diselesaikan dengan lancar. Penghargaan juga diberikan kepada Universitas Bina Sarana Informatika atas fasilitas serta bekal keilmuan yang telah diberikan selama masa perkuliahan. Karya ilmiah ini merupakan hasil pengembangan dari skripsi yang disusun penulis pada Program Studi Teknologi Informasi. Ucapan terima kasih selanjutnya ditujukan kepada dosen pembimbing yang telah memberikan arahan, masukan, serta saran yang bersifat membangun dalam penyempurnaan naskah dan metode penelitian. Tidak lupa, penulis mengucapkan apresiasi kepada rekan-rekan seperjuangan atas dukungan dan motivasi yang diberikan sepanjang proses penelitian.

DAFTAR REFERENSI

- Adebiyi, M. O., & Yasmin, N. (2023). Text classification on customer review dataset using support vector machine. In *Proceedings of the International Conference on Intelligent Systems*. Springer. https://doi.org/10.1007/978-981-19-7663-6_38
- Arpita, Kumar, P., & Garg, K. (2020, February 1). *Data cleaning of raw tweets for sentiment analysis*. IEEE. <https://doi.org/10.1109/Indo-TaiwanICAN48429.2020.9181326>
- Asri, Y., Kuswardani, D., Suliyanti, W. N., Manullang, Y. O., & Ansyari, A. R. (2025). Sentiment analysis based on Indonesian language lexicon and IndoBERT on user reviews PLN mobile application. *Indonesian Journal of Electrical Engineering and Computer Science*, 38(1), 677–688. <https://doi.org/10.11591/ijeecs.v38.i1.pp677-688>
- Bhalla, A. (2023, August 3). *Comparative analysis of text classification using SVM, Naïve Bayes, and KNN models*. IEEE. <https://doi.org/10.1109/ICIRCA57980.2023.10220728>
- Chand, R., Khan, S. U. R., Hussain, S., & Wang, W. (2022, December 1). *TTAG+R: A dataset of Google Play Store's top trending Android games and user reviews*. IEEE. <https://doi.org/10.1109/QRS-C57518.2022.00093>
- Cui, J., Wang, Z., Ho, S.-B., & Cambria, E. (2023). Survey on sentiment analysis: Evolution of research methods and topics. *Artificial Intelligence Review*. <https://doi.org/10.1007/s10462-022-10386-z>
- Daquigan, J. M., Marbella, G. K. G., Dioses, R. M., Co, J. D. C., Centeno, C. J., & Mata, K. E. (2025). Enhancement of profanity filtering and hate speech detection algorithm applied in Minecraft chats. *Technoarete Transactions on Advances in Computer Applications*. <https://doi.org/10.36647/TTACA/04.02.A001>

- Fahim, S. F., Sounok, S. A., Shaeed, N., Orpa, M. H., & Niloy, N. T. (2024). *Analyzing user sentiment in Google Play Store reviews: A natural language processing approach*. IEEE. <https://doi.org/10.1109/ICCCNT61001.2024.10725684>
- Gantla, H. R., Mamodiya, U., A, K., Prasad, K. R., Vijayasanthi, M., & Srinivas, V. K. S. (2024, December 20). *Text mining and sentiment analysis using big data techniques*. IEEE. <https://doi.org/10.1109/ARIIA63345.2024.11051875>
- Guo, K., Utkarsh, A., Ding, W., Ondracek, I., Zhao, Z., Freeman, G., Vishwamitra, N., & Hu, H. (2024). *Moderating illicit online image promotion for unsafe user-generated content games using large vision-language models*. arXiv. <https://doi.org/10.48550/arXiv.2403.18957>
- Gupta, G., Raja, K., Gupta, M., Jan, T., Whiteside, S. T., & Prasad, M. (2024). A comprehensive review of deepfake detection using advanced machine learning and fusion methods. *Electronics*, 13(1), 95. <https://doi.org/10.3390/electronics13010095>
- Kassim, M. N., Mat Jali, S. H., Maarof, M. A., Zainal, A., & Abdul Wahab, A. (2020). Enhanced text stemmer with noisy text normalization for Malay texts. In *Proceedings of the International Conference on Soft Computing and Data Engineering*. Springer. https://doi.org/10.1007/978-981-15-0077-0_44
- Kolluri, J., & Razia, S. (2020). Text classification using Naïve Bayes classifier. *Materials Today: Proceedings*. <https://doi.org/10.1016/j.matpr.2020.10.058>
- Lin, Y.-C., Chen, S.-A., Liu, J.-J., & Lin, C.-J. (2023). Linear classifier: An often-forgotten baseline for text classification. *Proceedings of the Annual Meeting of the Association for Computational Linguistics*. <https://doi.org/10.18653/v1/2023.acl-short.160>
- Matiini, G. (2024). Varieties of neologism used in online gaming conversation. *Language Literacy*, 8(1). <https://doi.org/10.30743/ll.v8i1.9190>
- Pajila, P. J. B., Sheena, B. G., Gayathri, A., Aswini, J., Nalini, M., & R, S. (2023, September 20). *A comprehensive survey on Naïve Bayes algorithm: Advantages, limitations and applications*. IEEE. <https://doi.org/10.1109/ICOSEC58147.2023.10276274>
- Şahin, D. (2023). The use of text mining in new media studies: A systematic review. *Yeni Medya Çalışmaları*. <https://doi.org/10.55609/yenimedya.1340590>
- Tan, J. Y., Chow, A., & Tan, C. W. (2022). A comparative study of machine learning algorithms for sentiment analysis of game reviews. *Journal of the Institution of Engineers, Malaysia*, 82(3). <https://doi.org/10.54552/v82i3.101>
- Vashishtha, S., Gupta, V., & Mittal, M. (2023). Sentiment analysis using fuzzy logic: A comprehensive literature review. *Wiley Interdisciplinary Reviews: Data Mining and Knowledge Discovery*, 13(5), e1509. <https://doi.org/10.1002/widm.1509>
- Vusak, E., Kuzina, V., & Jovic, A. (2021, September 27). *A survey of word embedding algorithms for textual data information extraction*. IEEE. <https://doi.org/10.23919/MIPRO52101.2021.9597076>