



PERBANDINGAN SENTIMEN ULASAN APLIKASI GEMINI PADA GOOGLE PLAY STORE MENGGUNAKAN METODE SVM DAN NAIVE BAYES

Dimas Pungky Eprianza¹, Mochamad Agung Laksono², Aria Farhan Qur'ani³,
Susliansyah⁴, Hendro Priyono⁵

¹19210719@bsi.ac.id, ²19211127@bsi.ac.id, ³19210739@bsi.ac.id,

⁴susliansyah.slx@bsi.ac.id, ⁵Hendro.hop@bsi.ac.id

^{1,2,3,4,5}Program Studi Sistem Informasi, Fakultas Teknik dan Informatika,
Universitas Bina Sarana Informatika

Abstrak

Teknologi kecerdasan buatan (AI) berkembang pesat dan telah diterapkan di berbagai aspek kehidupan masyarakat, termasuk *chatbot* seperti Gemini AI dari Google. Meskipun memiliki banyak fitur menarik, beberapa pengguna telah mengungkapkan ketidakpuasan terhadap fungsionalitas aplikasi tersebut. Berdasarkan ulasan dari *Google Play Store*, penelitian ini menyelidiki persepsi pengguna terhadap aplikasi Gemini. Metode *Support Vector Machine* dan *Naive Bayes* digunakan untuk menganalisis 1.000 ulasan yang dikumpulkan melalui proses *scraping*. Pra-pemrosesan, pembobotan TF-IDF, pemisahan data, implementasi algoritma, dan evaluasi model Google Colab merupakan bagian dari fase penelitian. Berdasarkan temuan penelitian ini, SVM mencapai akurasi 81,4%, sedangkan teknik *Naive Bayes* mencapai akurasi terbaik sebesar 82,4%. Hasil analisis mengindikasikan bahwa kedua metode memiliki kemampuan yang lebih baik dalam klasifikasi sentimen positif relatif terhadap sentimen negatif.

Kata kunci: *Artificial Intelligence, Analisis Sentimen, Gemini AI, SVM, Naive Bayes*

Abstract

Artificial intelligence (AI) technology is advancing rapidly and has been applied to various aspects of people's lives, including chatbots such as Google's Gemini AI. Although it offers many attractive features, some users have expressed dissatisfaction with the app's functionality. Based on reviews from the Google Play Store, this study investigates user perceptions of the Gemini app. The Support Vector Machine and Naive Bayes algorithms were used to analyze 1,000 reviews collected through a scraping process. Preprocessing, TF-IDF weighting, data splitting, algorithm implementation, and model evaluation on Google Colab were part of the research phases. Based on the findings of this study, SVM achieved an accuracy of 81.4%, while the Naive Bayes technique achieved the best accuracy of 82.4%.

Keywords: *Artificial Intelligence, Analisis Sentimen, Gemini AI, SVM, Naive Bayes*

1. Pendahuluan

Artificial Intelligence merupakan cabang ilmu komputer yang berfokus pada pengembangan sistem atau mesin yang mampu menjalankan tugas-tugas yang biasanya membutuhkan kecerdasan manusia. AI bekerja dengan memanfaatkan komputer dan sistem berbasis algoritma serta model matematika untuk mempelajari data, mengenali pola, dan mengambil keputusan secara otomatis[1].

Kehadiran *Artificial Intelligence* (AI) telah membawa perubahan signifikan di berbagai bidang kehidupan. Salah satu inovasi yang menarik dari kecerdasan buatan adalah kemampuan untuk berkomunikasi antara manusia dan mesin yang disebut *chatbot*. *Chatbot* adalah program komputer yang dapat berinteraksi dengan manusia melalui internet seperti *ChatGPT*, *Deepseek*, *ClaudeAI* dan *Gemini AI*[2].

Salah satu *chatbot* yang banyak dimanfaatkan ialah Gemini AI salah satu kecerdasan buatan yang banyak digunakan saat ini. Yang dikembangkan oleh *Google Deep Mind*, aplikasi ini dirancang dengan teknologi multimodal untuk memproses berbagai jenis data seperti teks, gambar, audio, video, dan kode program[3]. Meskipun Gemini AI menawarkan berbagai keunggulan, banyak pengguna mengeluhkan respons yang kurang akurat, khususnya pada pertanyaan teknis seperti pemrograman. Beberapa ulasan juga menyebutkan jawaban yang tidak relevan dengan pertanyaan utama. Hal ini menunjukkan adanya

kesenjangan antara ekspektasi dan pengalaman pengguna, sehingga analisis sentimen diperlukan untuk menilai terhadap kinerja Gemini AI.

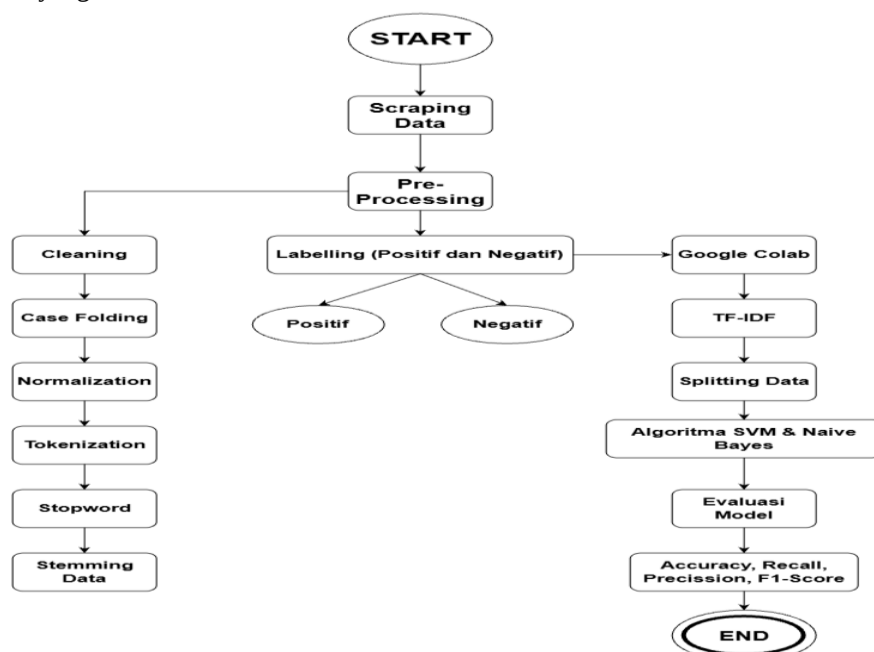
Penelitian ini berfokus pada penerapan analisis sentimen menggunakan pendekatan *machine learning*. Studi mengenai sentimen adalah sebuah pandangan, dan emosi yang tertuang dalam tulisan dikenal sebagai analisis sentimen [4]. Berdasarkan penelitian sebelumnya oleh Pamungkas & Cahyono menunjukkan bahwa algoritma SVM memiliki akurasi 80%, sedangkan algoritma KNN memiliki akurasi 71%[5]. Lalu oleh Najibulloh menunjukkan bahwa metode SVM sebesar 95,4%, *Decision Tree* sebesar 92,1% dan *Neural Network* (NN) memiliki hasil 95,5%[6].

Untuk menganalisis sentimen diatas dibutuhkan suatu metode yaitu menggunakan metode algoritma SVM dan Naive Bayes. Melalui pendekatan ini pada dua metode ini, penelitian diharapkan dapat memberikan kontribusi dalam bentuk evaluasi sistematis terhadap pandangan pengguna terhadap Gemini AI, sekaligus menjadi acuan bagi pengembang aplikasi untuk menjadi lebih baik kedepannya. Dengan membandingkan dua metode klasifikasi tersebut berdasarkan hasil evaluasi model[7]. Penelitian ini diharapkan dapat menentukan metode optimal dalam mengukur kepuasan pengguna dan efektivitas Gemini AI, sekaligus memberikan kontribusi akademik serta masukan praktis bagi pengembang untuk menjadikan aplikasinya menjadi lebih baik. Kebaruan pada hasil ini terletak pada penggunaan data ulasan terkini yang dikaji berdasarkan rentang waktu selama satu tahun, yaitu dari Juni 2024 hingga Juni 2025. Pendekatan ini memungkinkan diperolehnya gambaran sentimen pengguna yang lebih aktual dan relevan terhadap perkembangan fitur, performa, serta pengalaman penggunaan Gemini AI dalam periode terbaru.

Sebagian besar penelitian analisis sentimen sebelumnya berfokus pada aplikasi e-commerce, media sosial, dan layanan digital lainnya dengan dominasi penggunaan metode SVM yang sering menunjukkan performa terbaik. Sementara itu, penelitian yang bertujuan untuk menganalisis persepsi pengguna pada aplikasi berbasis *Artificial Intelligence* (AI), khususnya Gemini AI, masih sangat terbatas. Dari beberapa penelitian sebelumnya tidak terlalu banyak membahas dan membandingkan performa dari algoritma SVM dan Naive Bayes dari pada Gemini AI untuk mengetahui metode yang dominan serta efektif.

2. Metode

Dalam metode ini mengimplementasikan alur sistematis dengan mengkaji aplikasi Gemini pada *Google Play Store*. Tahapan dimulai dari *scraping* data, lalu dilakukannya proses *Preprocessing*. Selanjutnya, TF-IDF serta memberikan label pada setiap ulasan sentimen. Data kemudian dibagi dengan 80% data latih serta 20% untuk dilakukannya data pengujian untuk proses klasifikasi menggunakan SVM. Evaluasi performa model diukur berdasarkan hasil dari *confusion matrix*. Pada Gambar 1 menyajikan gambaran mengenai tahapan penelitian yang dilakukan.



Gambar 1. Kerangka Penelitian

2.1. Scrapping Data

Teknik *scrapping* digunakan untuk mendapatkan data atau informasi dari situs *web* tertentu. Ini dapat berupa teks, tautan, video, audio, atau dokumen[8]. Untuk *scrapping* ini dilakukan dalam batas-batas tertentu. Adapun batasan yang diambil pada penelitian ini menscraping data ulasan pada aplikasi Gemini AI dalam *scrapping* ini mengambil data berupa *Username*, *Rating*, *Review Text/Ulasan*, Tanggal ulasan. Dataset awal terdiri dari 1.000 data hasil *scrapping*. Setelah dilakukan proses penghapusan data duplikat, jumlahnya berkurang menjadi 999. Selanjutnya, pada tahap *preprocessing* dan pelabelan, Beberapa data tidak memenuhi kriteria, jadi dihapus. sehingga total data akhir yang digunakan menjadi 994.

2.2. Preprocessing Data

Langkah pertama dalam proses adalah prapemrosesan, yang mencakup mengubah data menjadi format yang siap dianalisis [9]. Pada preprocessing ini meliputi:

1. *Cleaning*: sebuah proses untuk menghapus elemen mengganggu agar data lebih bersih dan bisa dianalisis.
2. *Case folding*: merupakan bagian dari pemrosesan teks untuk menyamakan penulisan dengan mengubah semua huruf menjadi kecil.
3. *Normalization*: merupakan suatu proses transformasi data numerik ke dalam skala tertentu.
4. *Tokenization*: merupakan sebuah proses pemisah kalimat dan menjadi perkata di sebuah data.
5. *Stopword Removal*: merupakan tahap dalam proses analisis yang melibatkan penyaringan kata yang tidak penting.
6. *Stemming*: merupakan proses menghapus kata awal dan kata akhir yang nantinya menghasilkan kata dasar.

2.3. Labelling Sentimen

Pada proses *labelling* sentimen sangat penting untuk memperoleh data yang sesuai karena menggunakan label positif dan negatif dalam analisis sentimen, yang berguna untuk melatih untuk menghasilkan data yang relevan. [10]. Pada penelitian ini labelling untuk memudahkan klasifikasi dan analisis sentimen ulasan pada Aplikasi Gemini AI.

2.4. Term Frequency-Inverse Document Frequency

Pemrosesan TF-IDF dengan menentukan relevansi kata dalam sebuah dokumen dievaluasi menggunakan metode pemrosesan teks TF-IDF dengan membandingkannya dengan seluruh kumpulan dokumen. [11]. Adapun rumus pada TF-IDF sebagai berikut:

$$TF(t) = \frac{\text{jumlah kata } (t) \text{ dalam komentar}}{\text{jumlah total kata dalam komentar}} \quad (1)$$

$$IDF = \frac{1 + \text{total dokumen frekuensi } (N)}{1 + \text{jumlah total kata (DF)}} + 1 \quad (2)$$

Keterangan:

TF-IDF = TF × IDF

TF : *Term Frequency*

DF : *Document Frequency*

IDF : *Inverse Document Frequency*

(t) : Term atau sebuah kata dalam komentar

N : Jumlah total dokumen

2.5. Klasifikasi Support Vector Machine & Naive Bayes

Proses dengan metode *Support Vector Machine* (SVM) diawali dengan memasukkan data yang telah dipilih. Selanjutnya, data tersebut melalui tahap klasifikasi menggunakan SVM untuk memperoleh nilai akurasi. Tahapan awal dalam proses ini yaitu mengidentifikasi karakteristik data yang dijadikan parameter klasifikasi, kemudian dilanjutkan dengan perhitungan akurasi dari hasil klasifikasi tersebut[12]. Adapun rumus *Support Vector Machine (Linear)*:

$$f(x) = w \cdot x + b \quad (3)$$

Keterangan:

w: vektor bobot (*weight*)

x: vektor fitur input (hasil TF-IDF)

b: *bias*

f(x): fungsi keputusan

Naive Bayes termasuk metode klasifikasi berdasarkan peluang. Metode ini menentukan peluang suatu data pada kelas yang bersifat spesifik pada banyaknya kemunculan serta hubungan antar fitur pada data. Algoritma ini didasarkan pada *Teorema Bayes* bahwa setiap atribut tidak bergantung satu sama lain terhadap nilai pada variabel kelas [13]. Adapun rumus *Naive Bayes*:

$$P(H|X) = \frac{P(X|H) \cdot P(H)}{P(X)} \quad (4)$$

Keterangan:

X = Data yang belum diketahui label kelasnya

H = Hipotesis bahwa data X termasuk ke dalam kelas tertentu

P(H|X) = Probabilitas hipotesis H setelah diberikan informasi berupa data X (*Posteriori Probality*)

P(H) = Probabilitas awal hipotesis H sebelum mempertimbangkan data (*Prior Probability*)

P(H|X) = Probabilitas munculnya data X jika hipotesis H benar

P(X) = Probabilitas terjadinya data X

2.6. Evaluasi Model

Evaluasi model adalah untuk mengukur baik kinerja atau model dalam menyelesaikan tugas yang telah diterapkan, seperti klasifikasi, regresi, atau analisis sentimen. Evaluasi model dilakukan dengan menggunakan *Confussion Matrix*, yaitu *Confussion Matrix* merupakan hasil dari membandingkan hasil klasifikasi dengan label sebenarnya dan menghitung jumlah prediksi yang sesuai serta tidak sesuai pada setiap kelas [14]. *Confussion Matrix* terdiri dari beberapa komponen yaitu:

- True Positive* (TP): prediksi model sesuai dengan kondisi sebenarnya, yaitu keduanya bernilai benar.
- True Negative* (TN): prediksi model sesuai dengan kondisi sebenarnya, yaitu keduanya bernilai salah.
- False Positive* (FP): prediksi model menyatakan benar, tetapi kondisi sebenarnya bernilai salah.
- False Negative* (FN): nilai prediksi salah dan nilai sebenarnya benar prediksi model menyatakan salah, namun kondisi sebenarnya bernilai benar.

Tabel 1. *Confussion Matrix*

	<i>Predicted Positive</i>	<i>Predicted Negative</i>
<i>Actual Positive</i>	TP	FN
<i>Actual Negative</i>	FP	TN

Confussion Matrix bisa dihitung matrix evaluasi seperti tingkat akurasi, presisi, *recall*, dan *F1 score* yang nantinya menunjukkan kinerja pada model klasifikasi yang ditentukan.

- Accuracy* adalah rasio prediksi yang benar terhadap seluruh jumlah data yang diuji.

$$Accuracy = \frac{TP+TN}{TP+TN+FP+FN} \quad (5)$$

- Recall* menghitung tingkat keberhasilan model dalam mengenali data positif yang sebenarnya ada.

$$Recall \text{ (Kelas Positif)} = \frac{TP}{TP+FN} \quad (6)$$

$$Recall \text{ (Kelas Negatif)} = \frac{TN}{TN+FP} \quad (7)$$

3. *Precision* menghitung seberapa banyak prediksi positif yang benar positif.

$$Precision \text{ (Kelas Positif)} = \frac{TP}{TP + FP} \quad (8)$$

$$Precision \text{ (Kelas Negatif)} = \frac{TN}{TN + FN} \quad (9)$$

4. *F1 Score* merupakan rata-rata harmonis dari *precision* dan *recall*. Metrik ini digunakan untuk menyeimbangkan keduanya.

$$F1 \text{ Score (Positif)} = 2X \frac{Precision(Positif) \times Recall(Positif)}{Precision(Positif) + Recall(Positif)} \quad (10)$$

$$F1 \text{ Score (Negatif)} = 2X \frac{Precision(Negatif) \times Recall(Negatif)}{Precision(Negatif) + Recall(Negatif)} \quad (11)$$

3. Hasil dan Pembahasan

Penelitian ini diawali dengan proses *scrapping data* di *Google Colab* untuk aplikasi Gemini AI. Data tersebut kemudian diolah melalui serangkaian langkah prapemrosesan dan diberi label sentimen positif dan negatif. Fitur TF-IDF kemudian diterapkan pada dataset yang telah diberi label, yang selanjutnya dibagi menjadi uji pelatihan dan uji pengujian. Fokus utama studi ini adalah klasifikasi sentimen menggunakan algoritma SVM dan *Naive Bayes*. Kinerja dievaluasi menggunakan skor *Confusion Matrix*.

Berdasarkan hasil yang diperoleh, *Naive Bayes* memiliki tingkat kinerja yang baik dibandingkan SVM karena mampu memanfaatkan probabilitas kemunculan kata dalam setiap kelas untuk mengidentifikasi pola teks secara lebih efektif. Namun, kedua model masih kurang optimal dalam mengenali sebuah sentimen karena data positif lebih banyak dibandingkan negatif. Dengan kondisi tersebut, model mampu mengkategorikan sentimen positif dengan tingkat ketepatan yang lebih tinggi karena lebih banyak dipelajari selama proses pelatihan. Meskipun dari sebagian penelitian yang lain SVM sering menghasilkan akurasi lebih tinggi, pada penelitian ini *Naive Bayes* lebih unggul, yang kemungkinan dipengaruhi oleh karakteristik data dan tahapan pengolahan yang digunakan.

3.1. Pengumpulan Data

Untuk pengumpulan data ini dilakukan dengan teknik *scraping* yang menggunakan bahasa *Python* dengan menggunakan *library googleplayscraper*. Untuk pengambilan data set ini diambil pada tanggal Juni 2024 –Juni 2025 sebanyak 1000 data. Dari data tersebut akan menampilkan hasil dari *scraping* yang berisi berbagai ulasan yang diberikan oleh user dari aplikasi Gemini AI yang ada di *Playstore* yang ditunjukkan pada Tabel 2.

Table 2. Hasil *Scrapping Data*

No	Nama	Rating	Ulasan	Tanggal Ulasan
1.	Hahan Dirdja Subagdja	5	sajauh ini bagus, semoga terus bagus biar bayak penggunaanya	11/06/2025 17:34:19
2.	Arman William	5	Sangat Berguna dan bermanfaat. mantap. Dua jempol untuk developer,	24/05/2025 22:37:56
3.	Galuh Ika Permadi	4	mudah di pakai dan sangat membantu.	11/06/2025 13:42:29
...	...	...	...	...
1000	Muh Amar	5	sangat membantu dan tentunya bikin bangga.	15/04/2025 23:59:57

Dari Tabel 2 ini dapat hasil dari scraping data yaitu ulasan mentah yang diperoleh dari *Google Play Store*. Dari data itu masih belum bersih ditunjukkan pada Tabel 2. Dalam ini menunjukkan bahwa hasil *scraping* diperlukan proses preprocessing agar data tersebut bisa diolah.

Tabel 3. Hasil *Drop Duplicates*

No	Nama	Rating	Ulasan	Tanggal Ulasan
1.	Hahan Dirdja Subagdja	5	sajauh ini bagus, semoga terus bagus biar bayak penggunanya	11/06/2025 17:34:19
2.	Yans Lombok	3	berlangganan ke advance dapet 3 generate video per 24 jam wkwk 	11/06/2025 14:53:03
3.	Galuh Ika Permadi	4	mudah di pakai dan sangat membantu.	11/06/2025 13:42:29
...	...	...	...	...
999	Metta Yana	5	sangat membantu buat ngerjain tugas.	11/04/2025 12:18:06

Drop duplicates dilakukan untuk menghilangkan data ulasan yang semula 1.000 menjadi 999 agar data ulasan tidak menjadi duplicate yang ditunjukkan pada Tabel 3.

3.2. Preprocessing

Pada tahap *preprocessing* bertujuan untuk menganalisis data ataupun pemodelan dengan tujuan agar data lebih terstruktur. Maka dari itu, memerlukan beberapa tahapan proses seperti *Cleaning* membersihkan teks dari karakter yang tidak diperlukan seperti angka, tanda baca, URL, dan simbol agar data lebih bersih dan terstruktur, *Case Folding* mengubah seluruh huruf menjadi huruf kecil (*lowercase*) untuk menyamakan format penulisan kata. *Normalization* mengubah kata tidak baku atau singkatan ke bentuk baku agar konsisten. *Tokenization* memecah teks menjadi potongan kecil (token) seperti kata atau frasa untuk memudahkan analisis. *Stopword Removal* menghapus kata umum yang tidak berpengaruh terhadap makna. *Stemming* mengubah kata ke bentuk dasar dengan menghilangkan imbuhan.

Tabel 4. Hasil *Preprocessing Text*

Ulasan	Cleaning	Case folding	Normalization	Tokenization	Stopword Removal	Stemming
sajauh ini bagus, semoga terus bagus biar bayak penggunaanya	sajauh ini bagus semoga terus bagus biar bayak penggunaan ya	sajauh ini bagus semoga terus bagus biar bayak penggunaan ya	sajauh ini bagus semoga terus bagus biar bayak penggunaanya	['sajauh', 'ini', 'bagus', 'semoga', 'terus', 'bagus', 'biar', 'bayak', 'penggunaanya', 'ya']	['sajauh', 'bagus', 'semoga', 'bagus', 'biar', 'bayak', 'penggunaan', 'y a']	sajauh bagus moga bagus biar bayak guna
Sangat Berguna dan bermanfaat. Dua jempol untuk developer,	Sangat Berguna dan bermanfaat mantap Dua jempol untuk developer	sangat berguna dan bermanfaat mantap dua jempol untuk developer	sangat berguna dan bermanfaat mantap dua jempol untuk developer	['sangat', 'berguna', 'dan', 'bermanfaat', 'jempol', 'untuk', 'developer', 'dua']	['berguna', 'bermanfaat', 'mantap', 'jempol', 'developer']	guna manfaat mantap jempol develope r
Selalu memberikan Info yang tidak akurat	Selalu memberikan Info yang tidak akurat	selalu memberikan info yang tidak akurat	selalu memberikan info yang tidak akurat	['selalu', 'memberikan', 'info', 'yang', 'tidak', 'akurat']	['info', 'akurat']	info akurat

3.3. Hasil *Labelling Sentimen*

Pada proses *labelling* sentimen ini dilakukan berdasarkan *rating* dimana *rating* dua dikategorikan sebagai label “Negatif” dan *rating* selain daripada dua maka akan dikategorikan sebagai label “Positif” dimana proses *labelling* ini sangat penting untuk menentukan isi dari *review text* berdasarkan *rating* yang disajikan di Tabel 5.

Tabel 5. Hasil *Labelling Sentimen*

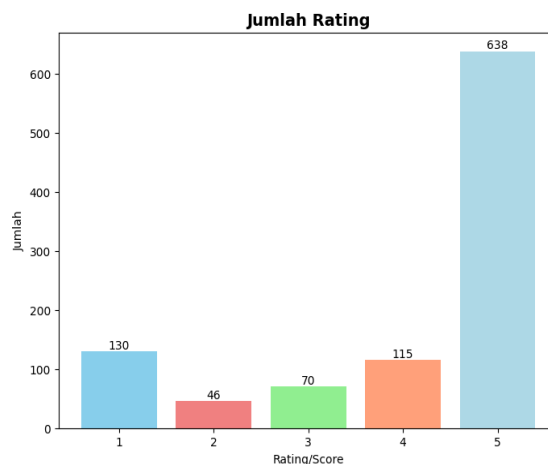
Ulasan	Rating	Sentimen
gua suka bug bug ya bikin kesal woi	2	Negatif
wow aplikasi bantu tugas request apa pokok keren	5	Positif
responsif relevan sesuai	5	Positif

Berdasarkan nilai penilaian yang diberikan oleh pengguna, Tabel 5 menunjukkan hasil pelabelan sentimen. Nilai penilaian 2 dilabelkan negatif, sedangkan nilai 5 dilabelkan positif. Hasil pelabelan menunjukkan bahwa ulasan yang mengandung apresiasi, kepuasan, atau penilaian baik dikategorikan sebagai sentimen positif. Pada penelitian, proses *labelling* menghasilkan data berlabel, yang merupakan dasar untuk proses analisis sentimen dan pembuatan model klasifikasi.

```
Sentiment
Positif    818
Negatif    176
Name: count, dtype: int64
```

Gambar 2. Jumlah *Labelling Sentimen*

Berikut ini adalah hasil dari jumlah data positif dan negatif yang dimana data positif berjumlah 818 dan data negatif berjumlah 176, langkah ini penting dalam analisis sentimen selanjutnya.



Gambar 3. Jumlah Rating Sentimen

Berdasarkan grafik di atas, mayoritas pengguna memberikan *rating* 3-5, yang mengindikasikan kepuasan terhadap aplikasi Gemini AI. Meskipun demikian, masih terdapat sebagian pengguna yang memberikan rating 1-2, menunjukkan adanya keluhan atau ketidakpuasan. Secara umum, sentimen positif pada ulasan lebih mendominasi daripada sentimen negatif

3.4. Hasil *Term Frequency-Inverse Document Frequency*

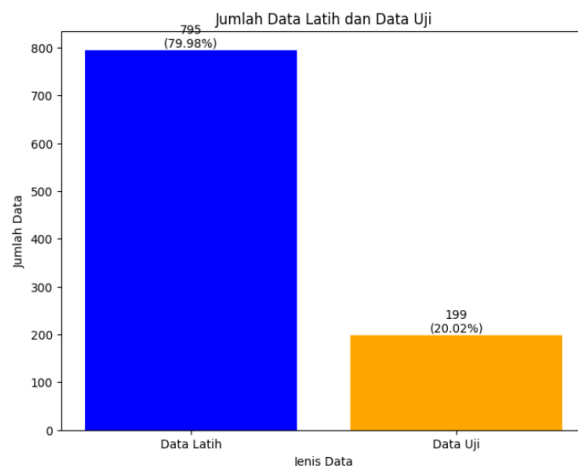
Setiap hasil bobot kata dibuat secara manual memanfaatkan metode TF dan TF-IDF. Tujuannya menilai signifikansi sebuah kata dalam suatu ulasan, dibandingkan dengan ulasan lain. Salah satu contoh ulasan dari *review text* yang akan digunakan untuk pembobotan kata “aplikasi bagus banget update ulang ya”.

Tabel 6. Hasil Perhitungan TF-IDF

Term	TF	DF	IDF	TF-IDF
aplikasi	$1/6 = 0,16$	193	5,231	0,836
bagus	$1/6 = 0,16$	285	3,870	0,619
banget	$1/6 = 0,16$	118	7,899	1,263
update	$1/6 = 0,16$	39	21.52	3,443
ulang	$1/6 = 0,16$	18	44.21	7,073
ya	$1/6 = 0,16$	167	5,886	0,941

Hasil pembobotan TF-IDF berdasarla memperlihatkan bahwa kata “ulang” memperoleh nilai tertinggi sebesar 7,073. Hal ini menunjukkan bahwa meskipun jarang muncul pada keseluruhan korpus kata tersebut memiliki tingkat kepentingan yang tinggi dalam 1 ulasan tertentu. Sebaliknya, kata “bagus” memiliki nilai lebih rendah (0,619) karena frekuensinya yang tinggi pada berbagai dokumen, sehingga tingkat kekhususannya menjadi rendah. Adapun kata “update” (3,443) dan “banget” (1,263) juga menunjukkan nilai TF-IDF yang relatif tinggi, menandakan kemunculan yang lebih jarang namun signifikan terhadap makna ulasan.

3.5. Splitting Data

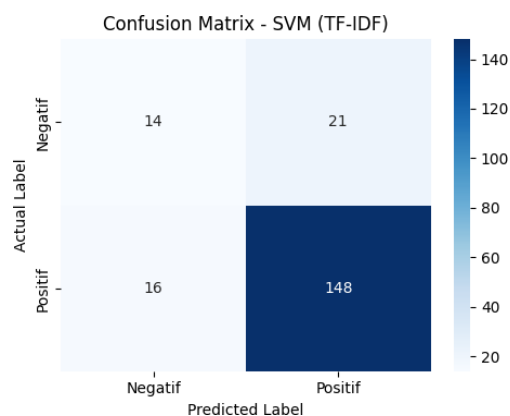


Gambar 4. Hasil Splitting Data

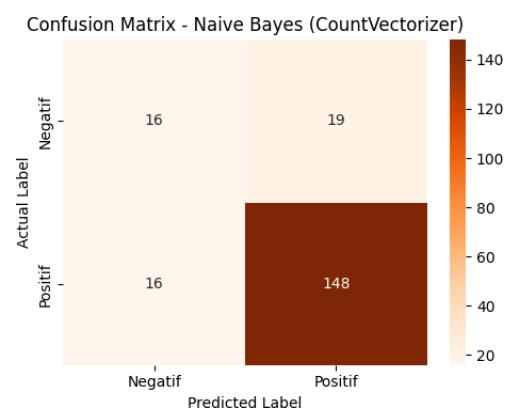
Grafik di atas menunjukkan hasil bahwa dataset telah dibagi menjadi dua data, yaitu data latih sebanyak 795 data (79,98%) yang ditampilkan dengan warna biru, dan data uji sebanyak 199 data (20,02%) yang ditampilkan dengan warna oranye. Visualisasi ini bertujuan untuk memastikan proporsi pembagian data sesuai dengan parameter yang telah ditentukan, yaitu 80% pelatihan dan 20% pengujian.

3.6. Hasil Klasifikasi SVM & Naive Bayes

Dilakukannya proses dengan evaluasi model untuk mengukur kinerja klasifikasi sentimen yang dihasilkan oleh algoritma SVM & Nave Bayes. Dengan memanfaatkan hasil dari confusion matrix dan *F1-Score* sebagai indikator utama. *Confusion Matrix* disajikan untuk melihat evaluasi yang dilakukan dengan membandingkan hasil prediksi terhadap data aktual.



Gambar 5. Confusion Matrix SVM



Gambar 6. Confusion Matrix Naive Bayes

Perbandingan antara *Confusion Matrix* SVM dan *Naive Bayes* yang ditunjukkan pada Gambar 5 dan Gambar 6 menunjukkan bahwa SVM Dari 199 data uji, sebanyak 148 data positif berhasil diklasifikasikan dengan benar, sedangkan 16 data positif salah diklasifikasikan sebagai negatif. Untuk data negatif, hanya 14 yang diprediksi dengan benar, sementara 21 data negatif salah diklasifikasikan sebagai positif.

Hasil ini mengindikasikan model lebih efektif dalam kemampuan klasifikasi sentimen positif daripada sentimen negatif, meskipun nilai akurasi model mencapai sekitar 81,41%. Sedangkan pada Naive Bayes Model berhasil mengidentifikasi 148 data positif dengan tepat (*True Positive*) dan 16 data negatif dengan benar (*True Negative*). Tetapi, terdapat kesalahan klasifikasi pada 19 data negatif yang terdeteksi sebagai positif (*False Positive*) dan 16 data positif yang terbaca sebagai negatif (*False Negative*).

Confusion Matrix:
[[14 21]
[16 148]]

Classification Report:

	precision	recall	f1-score	support
Negatif	0.47	0.40	0.43	35
Positif	0.88	0.90	0.89	164
accuracy			0.81	199
macro avg	0.67	0.65	0.66	199
weighted avg	0.80	0.81	0.81	199

Accuracy: 0.8140703517587939

Gambar 7. Evaluasi Kinerja Model SVM

Confusion Matrix:
[[16 19]
[16 148]]

Classification Report:

	precision	recall	f1-score	support
Negatif	0.50	0.46	0.48	35
Positif	0.89	0.90	0.89	164
accuracy			0.82	199
macro avg	0.69	0.68	0.69	199
weighted avg	0.82	0.82	0.82	199

Accuracy: 0.8241206030150754

Gambar 8. Evaluasi Kinerja Model Naive Bayes

Perbandingan evaluasi kinerja pada model SVM dan *Naive Bayes* yang ditunjukkan pada Gambar 7 dan Gambar 8 menunjukkan bahwa hasil evaluasi model SVM menunjukkan akurasi model sebesar 81,4%. Model memiliki kinerja sangat baik dalam mendeteksi ulasan *positif* *precision* 88%, *recall* 90,2%, *F1-score* 0,90, tetapi kurang optimal pada ulasan *negatif* *precision* 47%, *recall* 40%, *F1-score* 0,43. Hal ini mengindikasikan ketidakseimbangan performa, di mana model lebih dominan mengenali pola sentimen positif dibandingkan negatif. Sedangkan hasil evaluasi pada model *Naive Bayes* menghasilkan akurasi 82,4%. Pada kelas positif, performanya tinggi dengan *precision* 88%, *recall* 90,2%, dan *F1-score* 0,89, menunjukkan konsistensi dalam mengenali ulasan positif. Namun, di kelas negatif, *precision* 50%, *recall* 45%, dan *F1-score* 0,47 menandakan kesulitan dalam mengidentifikasi ulasan negatif. Secara keseluruhan, model lebih optimal menangani sentimen positif dibandingkan negatif.

4. Kesimpulan dan Saran

4.1 Kesimpulan

Penelitian ini menggunakan teknik *Support Vector Machines* dan *Naive Bayes* untuk menganalisis sentimen pengguna terhadap Gemini AI berdasarkan 1.000 ulasan di *Google Play Store*. Tahapan yang dilakukan meliputi prapemrosesan, pembobotan TF-IDF, serta pemisahan data pelatihan dan data uji. Berdasarkan hasil evaluasi, SVM mencapai akurasi 81,4%, sedangkan *Naive Bayes* mencapai akurasi maksimum 82,4%. Kedua pendekatan tersebut cenderung mendeteksi sentimen positif dengan lebih akurat daripada sentimen negatif selama pengujian. Meski demikian, *Naive Bayes* tetap menjadi yang paling unggul dalam penelitian ini.

4.2 Saran

Untuk penelitian selanjutnya, diharapkan menambah dari algoritma yang lain seperti *Random Forest*, *Logistic Regression*, LSTM, atau BERT untuk perbandingan. Selain itu, penerapan teknik *balancing* data seperti SMOTE atau *undersampling* diperlukan agar performa model terhadap kelas minoritas, khususnya sentimen negatif, dapat ditingkatkan.

Daftar Pustaka

- [1] Eriana, E. S., & Zein, D. A. (2023). Artificial Intelligence. *Angewandte Chemie International Edition*, 6(11), 1.
- [2] A. D. Pratama and H. Hendry, "Analisa Sentimen Masyarakat Terhadap Penggunaan Chatgpt Menggunakan Metode Support Vector Machine (Svm)," *JIPi (Jurnal Ilm. Penelit. dan Pembelajaran Inform.*, vol. 9, no. 1, pp. 327–338, 2024, doi: 10.29100/jipi.v9i1.4285.
- [3] R. Maulida, "Komparasi Respons ChatGPT dan Gemini terhadap Command Pattern Identik dengan Metode Black Box," vol. X, no. 2, pp. 2442–2445, 2024.

-
- [4] O. Manullang and C. Prianto, “Analisis Sentimen dalam Memprediksi Hasil Pemilu Presiden dan Wakil Presiden : Systematic Literature Review,” *J. Inform. Dan Teknol. Komput.*, vol. 4, no. 2, pp. 104–113, 2023, [Online]. Available: <https://ejurnalunsam.id/index.php/jicom/>.
 - [5] A. S. Pamungkas and N. Cahyono, “Analisis Sentimen Review ChatGPT di Play Store menggunakan Support Vector Machine dan K-Nearest Neighbor,” *Edumatic J. Pendidik. Inform.*, vol. 8, no. 1, pp. 1–10, 2024, doi: 10.29408/edumatic.v8i1.24114.
 - [6] I. K. Najibulloh, D. Intan, and S. Saputra, “Edumatic : Jurnal Pendidikan Informatika Analisis Sentimen Ulasan Co-Pilot Google Play dengan SVM , Neural Network , dan Decision Tree,” vol. 9, no. 1, pp. 275–283, 2025, doi: 10.29408/edumatic.v9i1.29673.
 - [7] A. Of, H. Disease, P. Using, A. C. Study, and M. L. Algorithms, “Analisis Prediksi Penyakit Jantung Menggunakan Perbandingan Algoritma *Machine Learning Analysis Of Heart Disease Prediction Using A Comparative Study Of*,” vol. 4, no. 2, pp. 262–269, 2025.
 - [8] S. Kusumo, “Penerapan Web Scraping Deskripsi Produk Menggunakan Selenium Python Dan Framework Laravel,” *JATISI (Jurnal Tek. Inform. dan Sist. Informasi)*, vol. 9, no. 4, pp. 3426–3435, 2022, doi: 10.35957/jatisi.v9i4.2727.
 - [9] H. H. Mubaroroh, H. Yasin, and A. Rusgiyono, “Analisis Sentimen Data Ulasan Aplikasi Ruangguru Pada Situs Google Play Menggunakan Algoritma Naive Bayes Classifier Dengan Normalisasi Kata Levenshtein Distance,” *J. Gaussian*, vol. 11, no. 2, pp. 248–257, 2022, doi: 10.14710/j.gauss.v11i2.35472.
 - [10] D. Moldovan, “A majority voting framework for reliable sentiment analysis of product reviews,” *PeerJ Comput. Sci.*, vol. 11, pp. 1–31, 2025, doi: 10.7717/peerj-cs.2738.
 - [11] D. Septiani and I. Isabela, “Analisis Term Frequency Inverse Document Frequency (TF-IDF) Dalam Temu Kembali Informasi Pada Dokumen Teks,” *SINTESIA J. Sist. dan Teknol. Inf. Indones.*, vol. 1, no. 2, pp. 81–88, 2023.
 - [12] O. Muhammad *et al.*, “Analisa Algoritma Support Vector Machine Pada Data Bunga Iris,” vol. 30, no. 1, pp. 477–487, 2022.
 - [13] M. Iqbal, “Penerapan Metode Naive Bayes Pada Ulasan Aplikasi Reelshort,” vol. 10, no. 1, pp. 49–59, 2025.
 - [14] T. Gori, A. Sunyoto, and H. Al Fatta, “Preprocessing Data dan Klasifikasi untuk Prediksi Kinerja Akademik Siswa,” *J. Teknol. Inf. dan Ilmu Komput.*, vol. 11, no. 1, pp. 215–224, 2024, doi: 10.25126/jtiik.20241118074.