

PERBANDINGAN KINERJA ALGORITMA KLASIFIKASI PADA REVIEW PENGGUNA APLIKASI NETFLIX

Rauhulloh Ayatulloh Khomeini Noor Bintang^{1*}, Nova Tri Romadloni², Faisya Ramadhani³

^{1,2,3}Informatika, Universitas Muhammadiyah Karanganyar; Jl. Raya Solo-Tawangmangu No.KM. 12, Pandes, Papahan, Kec. Tasikmadu, Kabupaten Karanganyar, Jawa Tengah 57761; 08112801912

Received: 2 Maret 2025

Accepted: 27 Maret 2025

Published: 14 April 2025

Keywords:

Algoritma Klasifikasi;
Analisis Sentimen;
Google Play Store;
Machine Learning;
Netflix.

Correspondent Email:

rahullbintang@gmail.com

Abstrak. Penelitian ini bertujuan untuk menganalisis sentimen ulasan pengguna aplikasi Netflix yang diperoleh dari Google Play Store menggunakan metode web scraping dengan Python di Google Colab. Data ulasan diproses melalui tahap pembersihan teks, tokenisasi, penghapusan stopword, dan stemming, serta direpresentasikan menggunakan metode Term Frequency-Inverse Document Frequency (TF-IDF). Lima algoritma klasifikasi, yaitu Logistic Regression, Naive Bayes, Decision Tree, Random Forest, dan Support Vector Machine (SVM), dibandingkan untuk menentukan algoritma terbaik dalam klasifikasi sentimen positif, negatif, dan netral. Evaluasi dilakukan berdasarkan akurasi dengan pembagian data latih dan data uji sebesar 90:10. Hasil pengujian menunjukkan bahwa Logistic Regression dan Random Forest memiliki akurasi tertinggi sebesar 76%, diikuti oleh SVM sebesar 74%, Decision Tree sebesar 73%, dan Naive Bayes dengan akurasi terendah sebesar 71%. Temuan ini memberikan kontribusi bagi penelitian di bidang analisis sentimen serta dapat menjadi referensi bagi pengembang aplikasi dalam meningkatkan pengalaman pengguna berbasis data.

Abstract. This research aims to analyze the sentiment of user reviews of the Netflix application obtained from the Google Play Store using the web scraping method with Python in Google Colab. Review data is processed through the stages of text cleaning, tokenization, stopword removal, and stemming, and is represented using the Term Frequency-Inverse Document Frequency (TF-IDF) method. Five classification algorithms, namely Logistic Regression, Naive Bayes, Decision Tree, Random Forest, and Support Vector Machine (SVM), were compared to determine the best algorithm for classifying positive, negative, and neutral sentiment. Evaluation is carried out based on accuracy with a 90:10 division of training data and test data. The test results show that Logistic Regression and Random Forest have the highest accuracy at 76%, followed by SVM at 74%, Decision Tree at 73%, and Naive Bayes with the lowest accuracy at 71%. These findings contribute to research in the field of sentiment analysis and can be a reference for application developers in improving data-based user experiences.

1. PENDAHULUAN

Analisis sentimen merupakan metode yang digunakan untuk memahami opini dan persepsi pengguna terhadap suatu produk atau layanan [1]. Saat ini, ulasan pengguna yang tersedia di platform digital, seperti Google Play Store, menjadi sumber data yang kaya untuk menggali

wawasan terkait pengalaman pengguna [2]. Netflix, sebagai salah satu platform streaming terpopuler, sering menjadi objek penelitian dalam analisis sentimen. Studi ini berfokus pada pengelompokan sentimen dalam ulasan pengguna Netflix ke dalam kategori positif,

negatif, atau netral dengan menerapkan metode klasifikasi.

Pemilihan algoritma klasifikasi yang tepat menjadi faktor kunci dalam menghasilkan analisis yang akurat dan relevan. Dalam penelitian ini, lima algoritma klasifikasi dibandingkan, yaitu Logistic Regression, Naive Bayes, Decision Tree, Random Forest, dan Support Vector Machine (SVM). Logistic Regression sering digunakan dalam klasifikasi biner karena kemampuannya dalam memodelkan probabilitas berdasarkan atribut tertentu [3]. Naive Bayes, yang mengasumsikan independensi antar fitur, dikenal sebagai metode yang sederhana namun efisien dalam pengolahan teks [4]. Decision Tree bekerja dengan membagi data berdasarkan aturan keputusan yang logis, sementara Random Forest mengombinasikan beberapa pohon keputusan untuk meningkatkan akurasi dan mengurangi risiko overfitting [5]. Support Vector Machine menggunakan hyperplane untuk memisahkan kelas data dengan margin optimal, sehingga sering diterapkan dalam analisis data berskala besar dan kompleks [6].

Penelitian ini bertujuan untuk mengevaluasi kinerja masing-masing algoritma dengan memanfaatkan data ulasan pengguna aplikasi Netflix di Google Play Store. Evaluasi dilakukan berdasarkan tingkat akurasi model dalam mengklasifikasikan sentimen. Dengan pendekatan ini, diharapkan penelitian dapat memberikan wawasan yang lebih mendalam terkait keunggulan dan keterbatasan masing-masing algoritma dalam analisis sentimen. Selain itu, hasil penelitian ini diharapkan dapat memberikan kontribusi bagi pengembang aplikasi dalam pengambilan keputusan berbasis data serta memperkaya literatur akademik di bidang klasifikasi sentimen.

2. TINJAUAN PUSTAKA

Analisis sentimen adalah cabang dari pemrosesan bahasa alami (Natural Language Processing, NLP) yang berfokus pada ekstraksi [7], pengklasifikasian, dan interpretasi opini atau emosi yang terkandung dalam teks. Dalam definisi lainnya analisis sentimen sebagai upaya untuk memahami orientasi sentimen, apakah positif, negatif, atau netral, dari suatu ulasan atau opini [8]. Dalam berbagai penelitian,

analisis sentimen banyak digunakan untuk memahami pola perilaku konsumen, tren pasar, hingga pengambilan keputusan berbasis data [9]. Pada ulasan pengguna aplikasi seperti Netflix di Google Play Store, analisis sentimen menjadi alat yang penting untuk menggali pengalaman pengguna dan mengidentifikasi aspek yang perlu ditingkatkan [10].

Algoritma klasifikasi merupakan inti dari proses analisis sentimen, yang bertugas memetakan data teks ke dalam kategori tertentu. Beberapa algoritma yang umum digunakan dalam analisis sentimen yaitu Logistic Regression, Naive Bayes, Decision Tree, Random Forest, dan Support Vector Machine (SVM).

Dalam analisis sentimen, langkah pemrosesan data teks melibatkan pembersihan data (data cleaning), tokenisasi, stemming, dan penghapusan stop words. Preprocessing data teks sangat penting untuk meningkatkan akurasi model klasifikasi. Pendekatan berbasis term frequency-inverse document frequency (TF-IDF) untuk merepresentasikan teks dalam format numerik yang dapat diproses oleh algoritma.

Penelitian sebelumnya banyak membahas analisis sentimen. Dalam beberapa tahun terakhir, analisis sentimen telah menjadi pendekatan yang banyak digunakan untuk mengevaluasi kepuasan pengguna terhadap berbagai layanan digital, termasuk platform streaming. Analisis ini memungkinkan pengembang aplikasi untuk memahami opini pengguna berdasarkan ulasan yang diberikan, baik dalam bentuk komentar positif, netral, maupun negatif. Beberapa penelitian sebelumnya telah membahas metode dan algoritma yang digunakan dalam analisis sentimen untuk berbagai aplikasi layanan digital dan platform streaming.

Dalam penelitiannya tentang analisis sentimen kepuasan pengguna terhadap layanan streaming Mola menggunakan algoritma Random Forest menemukan bahwa algoritma tersebut mampu mengklasifikasikan ulasan pengguna ke dalam kategori positif, negatif, dan netral dengan tingkat akurasi yang tinggi, mencapai 98% [11].

Hasil ini menunjukkan bahwa metode berbasis ensemble learning, seperti Random Forest, memiliki keunggulan dalam menangani teks ulasan yang kompleks. Selain itu, analisis

sentimen ulasan aplikasi pendidikan "Sister for Students UNEJ" dengan menggunakan algoritma Random Forest yang dioptimalkan dengan teknik Synthetic Minority Over-sampling Technique (SMOTE). Penelitian ini menunjukkan bahwa dengan penyeimbangan dataset, akurasi model meningkat secara signifikan hingga 98.9% [12].

Selain penggunaan algoritma Random Forest, beberapa penelitian telah membandingkan metode klasifikasi lain. Misalnya, menganalisis sentimen aplikasi NEWSAKPOLE menggunakan Naïve Bayes menunjukkan bahwa algoritma ini memiliki akurasi sebesar 88% setelah diuji dengan K-Fold validation [13].

Sementara itu, yang melakukan analisis sentimen ulasan e-commerce Shopee menunjukkan bahwa Naïve Bayes lebih unggul dibandingkan Support Vector Machine (SVM) dengan tingkat akurasi 85% [14]. Namun, dalam penelitian lain terkait aplikasi sentimen setelah pilpres 2024, algoritma Logistic Regression memberikan hasil terbaik dalam klasifikasi sentimen dengan akurasi mencapai 84,39% [15].

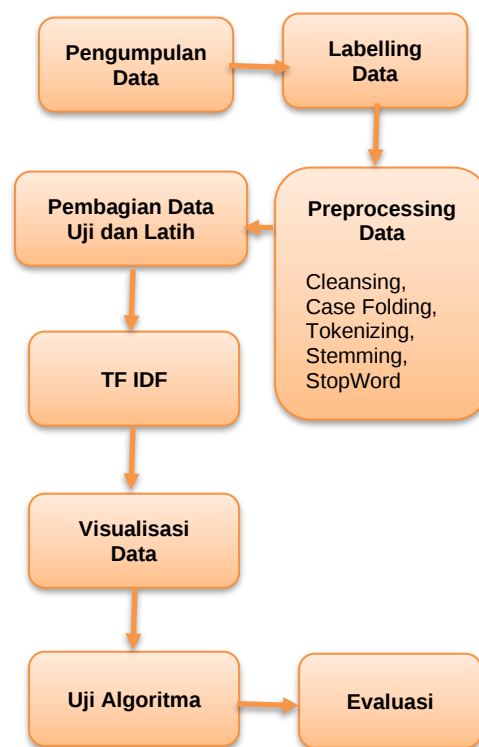
Berdasarkan hasil penelitian sebelumnya, dapat disimpulkan bahwa metode Random Forest, Naïve Bayes, dan Support Vector Machine sering digunakan dalam analisis sentimen ulasan pengguna aplikasi digital. Namun, masing-masing metode memiliki kelebihan dan keterbatasan tergantung pada karakteristik dataset yang digunakan. Oleh karena itu, penelitian ini bertujuan untuk menganalisis sentimen ulasan pengguna aplikasi Netflix dengan membandingkan beberapa algoritma klasifikasi, termasuk Logistic Regression, Naïve Bayes, Decision Tree, Random Forest, dan Support Vector Machine (SVM). Evaluasi model akan dilakukan untuk menentukan algoritma terbaik dalam mengklasifikasikan sentimen positif, negatif, dan netral pada ulasan pengguna Netflix, sehingga dapat memberikan wawasan lebih mendalam dalam memahami opini pelanggan dan meningkatkan kualitas layanan platform streaming ini.

Dengan berbagai pendekatan dan metode yang telah disebutkan, penelitian ini memberikan landasan yang kuat untuk mengeksplorasi perbandingan algoritma klasifikasi dalam analisis sentimen. Kajian

literatur ini juga menunjukkan relevansi algoritma yang dipilih dengan konteks analisis ulasan pengguna aplikasi Netflix, sehingga memberikan kontribusi signifikan dalam bidang pengolahan data dan pengembangan aplikasi berbasis data.

3. METODE PENELITIAN

Penelitian ini menganalisis ulasan pengguna aplikasi Netflix di google play store menggunakan algoritma logistic regression, naïve bayes, decision tree, random forest dan support vector machines (SVM). Berikut terdapat gambar 1 yang merupakan tahapan alur penelitian yang dilakukan.



Gambar 1. Alur Penelitian

3.1. Pengumpulan Data

Dalam penelitian ini pengambilan data ulasan pengguna aplikasi Netflix dilakukan melalui metode web scraping menggunakan bahasa Python di Google Colab dengan mengidentifikasi ID aplikasi Netflix, data ulasan diambil langsung dari Google Play Store dengan pengaturan parameter tertentu, seperti bahasa (`lang='id'`) dan wilayah (`country='id'`), untuk mengatur hanya ulasan berbahasa Indonesia dari pengguna di Indonesia yang dikumpulkan. Selain itu, ulasan dipilih

berdasarkan waktu terbaru agar hasilnya lebih relevan dengan opini pengguna terkini.

Setiap ulasan yang diambil berisi informasi penting seperti isi ulasan, skor rating (1-5 bintang), tanggal ulasan, nama pengguna, dan versi aplikasi yang digunakan. Pengambilan data dibatasi hingga 1000 ulasan setiap permintaan. Setelah data diperoleh, hasil scraping disimpan dalam format CSV untuk mempermudah pengelolaan dan analisis lebih lanjut.

3.2. Labeling Data

Tahap berikutnya adalah pelabelan, yang bertujuan untuk membagi data ke dalam berbagai kelas sentimen yang akan digunakan dalam penelitian. Tiga kelas sentimen yang paling umum digunakan adalah negatif, netral, dan positif. Tujuan dari proses pelabelan ini adalah membagi dataset menjadi dua bagian yaitu data latih dan data uji. Data pelatihan digunakan untuk mengajarkan sistem agar mengenali pola yang dicari, sedangkan data pengujian digunakan untuk menguji hasil pelatihan yang telah dilakukan.

3.3. Preprocessing Data

Preprocessing data merupakan tahapan yang dilakukan setelah pelabelan data. Pada tahap ini, data disiapkan agar siap untuk dianalisis. Beberapa tahapan yang dilakukan dalam preprocessing data yaitu tahap pertama adalah cleaning, yaitu menghapus simbol, angka, URL, dan karakter tidak relevan yang dapat mengganggu analisis. Selanjutnya, case folding dilakukan dengan mengubah seluruh teks menjadi huruf kecil guna menjaga konsistensi data. Proses berikutnya adalah tokenizing, yaitu memecah teks menjadi kata-kata atau unit terkecil. Setelah itu, stopword removal diterapkan untuk menghilangkan kata-kata umum yang tidak memiliki makna penting, seperti "dan" atau "yang". Tahap terakhir adalah stemming, yang bertujuan mengembalikan kata-kata ke bentuk dasarnya untuk mengurangi variasi bentuk kata yang sama. Dengan melakukan langkah-langkah ini, data teks menjadi lebih bersih, terstruktur, dan siap digunakan dalam proses analisis sentimen, sehingga dapat meningkatkan akurasi model yang diterapkan.

3.4. Pembagian Data

Pembagian data latih (training data) dan data uji (testing data) merupakan langkah krusial dalam pengembangan model machine learning, yang bertujuan untuk memastikan bahwa model dapat menggeneralisasi dengan baik terhadap data yang belum pernah dilihat sebelumnya. Data latih digunakan untuk melatih model agar dapat mengenali pola dan hubungan dalam dataset, sementara data uji berfungsi untuk menguji kinerja model setelah proses pelatihan. Metode pembagian yang umum digunakan meliputi pembagian acak, di mana data dibagi menjadi dua subset dengan proporsi yang sering kali 90% untuk data latih dan 10% untuk data uji. Pembagian yang tepat sangat penting untuk menghindari overfitting, di mana model terlalu cocok dengan data latih dan tidak dapat berfungsi dengan baik pada data baru.

Dataset biasanya dibagi menjadi set pelatihan dan set pengujian, dengan pembagian umum 70% untuk pelatihan dan 30% untuk pengujian [16]. Pembagian ini sangat penting untuk mengevaluasi kinerja model pada data yang tidak terlihat." Dengan demikian, pemilihan metode pembagian yang tepat dan pemahaman tentang pentingnya pembagian ini akan membantu dalam membangun model yang lebih robust dan efektif dalam aplikasi dunia nyata.

3.5. Pembobotan

Metode TF-IDF (Term Frequency-Inverse Document Frequency) adalah teknik dalam pengolahan teks yang digunakan untuk menilai seberapa penting sebuah kata dalam sebuah dokumen relatif terhadap kumpulan dokumen lainnya. TF-IDF terdiri dari dua komponen utama yaitu TF (Term Frequency) yang mengukur seberapa sering sebuah kata muncul dalam dokumen tertentu, dan IDF (Inverse Document Frequency) yang menghitung tingkat kelangkaan kata tersebut di seluruh dokumen. Kombinasi kedua komponen ini membantu mengurangi pengaruh kata-kata umum yang muncul di hampir semua dokumen, seperti "dan", "ke", "di", "yang", dan lainnya. sehingga memberikan bobot lebih besar pada kata-kata yang signifikan untuk analisis.

3.6. Visualisasi Data

Visualisasi kata pada tahap ini ditampilkan menggunakan Wordcloud, sebuah representasi

visual dari teks di mana kata-kata yang sering muncul memiliki ukuran lebih besar, sementara kata-kata yang jarang muncul memiliki ukuran lebih kecil.

3.7. Pengujian Algoritma

Pengujian algoritma machine learning bertujuan untuk mengevaluasi kinerja berbagai metode dalam menyelesaikan masalah tertentu. Setiap algoritma memiliki pendekatan matematis unik yang memengaruhi cara mereka belajar dari data dan menghasilkan prediksi. Dengan memahami teori di balik algoritma, pemilihan metode yang tepat dapat dilakukan untuk meningkatkan akurasi model. Berikut adalah algoritma yang digunakan pada analisis kali ini.

Logistic Regression adalah algoritma klasifikasi biner yang memodelkan hubungan antara variabel input dan probabilitas hasil melalui fungsi logit. Algoritma ini memberikan interpretasi probabilistik yang mudah dipahami dan cocok untuk data linear secara logistik. Sedangkan, Naïve Bayes merupakan algoritma probabilistik sederhana yang menghitung probabilitas suatu kelas dengan asumsi independensi antar fitur. Meski asumsinya kuat, algoritma ini efisien dan sering digunakan pada klasifikasi teks.

Decision Tree algoritma berbasis pohon yang membagi data berdasarkan atribut yang mengurangi ketidakpastian secara maksimal. Fleksibilitasnya membuat algoritma ini ideal untuk data numerik maupun kategorikal. Kemudian Random Forest adalah metode ensemble yang menggabungkan banyak Decision Tree untuk meningkatkan akurasi dan mengurangi overfitting. Prediksi akhir diperoleh melalui rata-rata atau voting mayoritas.

Support Vector Model (SVM) algoritma yang mencari hyperplane terbaik untuk memisahkan kelas dengan margin terbesar. Algoritma ini efektif untuk data berdimensi tinggi dan dapat menggunakan kernel untuk data yang tidak linear.

4. HASIL DAN PEMBAHASAN

Bagian ini menyajikan hasil dari proses analisis sentimen pada ulasan pengguna aplikasi Netflix yang diperoleh melalui metode web scraping. Analisis dilakukan dengan menerapkan serangkaian tahapan pemrosesan

data, mulai dari pengumpulan dan pelabelan data, preprocessing teks, pembobotan menggunakan TF-IDF, hingga implementasi berbagai algoritma klasifikasi.

Evaluasi model klasifikasi dilakukan dengan menggunakan metrik akurasi untuk menilai kinerja masing-masing algoritma dalam mengklasifikasikan sentimen positif, negatif, dan netral. Perbandingan hasil pengujian antar algoritma disajikan dalam bentuk tabel dan visualisasi guna memberikan wawasan yang lebih mendalam terkait efektivitas setiap metode yang digunakan.

Selain itu, dilakukan analisis terhadap distribusi data dan korelasi antara panjang ulasan dengan skor yang diberikan oleh pengguna. Visualisasi seperti word cloud dan diagram korelasi turut digunakan untuk memahami pola sentimen dalam ulasan yang dikumpulkan.

Pada subbagian berikut, akan dibahas secara rinci mengenai tahapan pengumpulan data, preprocessing, pembobotan teks, serta evaluasi performa algoritma klasifikasi yang digunakan dalam penelitian ini.

4.1. Pengumpulan Data

Hasil pengumpulan data menunjukkan cuplikan hasil scraping yang telah dilakukan. Data ini mencakup berbagai jenis sentimen, baik positif, negatif, maupun netral. Sebagai contoh, ulasan dengan skor 5 umumnya berisi komentar positif, seperti pujian terhadap kualitas video dan acara yang tersedia. Sebaliknya, ulasan dengan skor 1 mengindikasikan ketidakpuasan, seperti masalah login dan biaya berlangganan yang dianggap mahal.

	content	score	thumbsUpCount	\
0	Sangat bagus...videonya HD.. mahal tapi bulana...	5	0	
1	Saya beri bintang 5 karena acara2 dinetflix se...	5	0	
2	Ko saya gabisa masuk aja ke aplikasinya ya ken...	5	0	
3	Pilihan film nya sedikit. Jelek	1	0	
4	Jelek amat, udh bayar tapih gk bisa masuk lagi...	1	0	
..	...	...	...	
995	Sangat bagus	5	0	
996	GUA UDAH MELAKUKAN BERLANGGANAN TAPI GABISA MA...	1	1	
997	Susah banget loginnya	1	0	
998	njir lah bayar 🤔🤔🤔🤔🤔🤔🤔🤔	2	2	
999	Error terus	1	0	

Gambar 2. Hasil Scrapping pada Google Colab

Data yang diperoleh sebanyak 906 ulasan selanjutnya diunduh dalam format .csv yang akan melalui proses labelling data yang

mengubah teks ke dalam bentuk yang lebih terstruktur agar siap untuk tahap analisis sentimen lebih lanjut.

	A	B	C	D	E	F	G	H	I	J	K	L
1	reviewId	userId	userName	content	score	thumbsUp	thumbsDown	reviewCreatedAt	replyCount	replyAt	appVersion	
2	915945c	Dani		https://p/ Adakah kode negara tidak bisa di rubah	1	0	0	10/11/2024 5:54			8.135.1 build 750902	
3	98021264	Als Wilhel		https://p/ Nonton film itu gak perlu cari sesenya po	4	0	0	10/11/2024 5:30			8.135.1 build 750902	
4	3846021C	Cahya Ayu		https://p/ Mahal, ribet, bolak balik bayar tetap gak bi	1	0	0	10/11/2024 4:49			8.135.1 build 750902	
5	04679342	Aranda G		https://p/ Mau daftar atau login pake disuruh bertang	1	0	0	10/11/2024 2:29			8.135.1 build 750902	
6	44c0514	Dafra Aeri		https://p/ Bagus filmnya	4	0	0	10/11/2024 1:53			8.135.1 build 750902	
7	00204611	Zainal Anif		https://p/ Baru juga daftar plus bayar malah gk bisa d	1	0	0	10/11/2024 1:34			8.135.1 build 750902	
8	7050020F	claudia d		https://p/ saya gk bisa masuk ke aplikasi nya,tolong d	1	0	0	10/11/2024 1:06			8.135.1 build 750902	
9	68355005	Alnur Rafli		https://p/ Menghbur	5	0	0	10/10/2024 13:32			8.135.1 build 750902	
10	cc489555	Aldira Put		https://p/ Kng gk bisa di buka?	5	0	0	10/10/2024 17:32			8.135.1 build 750902	
11	1414ac32	Ira Puaput		https://p/ Serang	1	0	0	10/10/2024 16:58			7.110.0 build 633594	
12	13855a10	Avilla Awi		https://p/ Please release S3 fate the wins saga RVRB	1	0	0	10/10/2024 16:43			8.126.0 build 950771	
13	115fa13	Beben Infi		https://p/ Netflix 822406/29/2024	5	0	0	10/10/2024 16:41				
14	8c12ab10	Adnan Via		https://p/ Rf Rf Rf	5	0	0	10/10/2024 15:31				
15	9f2d63a	Dafa Fa		https://p/ Kayak anj ga bisa daftar	1	0	0	10/10/2024 15:05			8.135.1 build 750902	
16	ba302a5c	ERWINDA		https://p/ Ditengah2 episode selalu tak ada subtitle	1	0	0	10/10/2024 14:41			8.135.1 build 750902	
17	661647dc	Agil Wirat		https://p/ Sangat seru dengan pilihan film bioskop terl	5	0	0	10/10/2024 13:14			8.135.1 build 750902	
18	38f974ee	Adun Tico		https://p/ bagus	5	0	0	10/10/2024 12:31				
19	ca51115f	Bryan Kha		https://p/ Bagus	5	0	0	10/10/2024 11:23				
20	c1f1069c	UKIN BOT		https://p/ nob	1	0	0	10/10/2024 11:05			8.127.1 build 1050788	
21	1539962c	Yana Sriya		https://p/ Terbaik	5	0	0	10/10/2024 8:31			8.135.1 build 750902	
22	08a62882	Muh Rafli		https://p/ Masa 1x liat video 54rb parah banget si..k	1	0	0	10/10/2024 6:12			8.135.1 build 750902	
23	902915ca	Misa asha		https://p/ Sulit login	1	0	0	10/10/2024 5:11				
24	05469a3c	Adnan cha		https://p/ Film nya lengkap	5	0	0	10/10/2024 5:08				
25	29fcd87f	Agungfati		https://p/ Mantap	5	0	0	10/10/2024 5:08			8.135.1 build 750902	

Gambar 3. Hasil Scrapping dengan Format File .csv

4.2. Hasil Pemberian Label

Setiap ulasan dikategorikan ke dalam tiga label sentimen, yaitu positif, negatif, dan netral. Pemberian label ini didasarkan pada skor yang diberikan oleh pengguna dalam ulasannya.

Sentimen Positif dengan skor 4 dan 5 dianggap memiliki sentimen positif. Ulasan dalam kategori ini biasanya berisi komentar yang memuji layanan Netflix, seperti kualitas video yang baik, variasi konten yang menarik, dan pengalaman pengguna yang memuaskan.

Sentimen Netral dengan skor 3 dikategorikan sebagai netral, di mana pengguna mungkin memberikan komentar yang mengandung kelebihan dan kekurangan tanpa kecenderungan yang kuat ke arah positif atau negatif.

Sentimen Negatif: Ulasan dengan skor 1 dan 2 dikategorikan sebagai negatif. Ulasan dalam kategori ini umumnya berisi keluhan tentang masalah teknis, harga yang dianggap mahal, atau pengalaman buruk dalam menggunakan layanan Netflix.

Gambar 4. Dibawah ini hasil pelabelan kategori Positif, Negatif, dan Netral menunjukkan distribusi data setelah proses pelabelan dilakukan. Kategori ini akan digunakan sebagai dasar untuk analisis lebih lanjut dalam proses klasifikasi menggunakan algoritma machine learning.

	score	sentiment_label
0	5	Positif
1	5	Positif
2	5	Positif
3	1	Negatif
4	1	Negatif
..	...	...
994	1	Negatif
996	1	Negatif
997	1	Negatif
998	2	Negatif
999	1	Negatif

[906 rows x 2 columns]

Gambar 4. Hasil Pelabelan Kategori Positif, Negative, dan Netral

4.3. Hasil Cleaning dan Case Folding

Setelah data ulasan pengguna dikategorikan ke dalam sentimen positif, netral, dan negatif, tahap selanjutnya adalah proses cleaning dan case folding. Tahapan ini bertujuan untuk membersihkan teks dari elemen-elemen yang tidak relevan serta menyamakan format teks agar lebih mudah diolah dalam tahap analisis selanjutnya.

	content
0	sangat bagus videonya hd mahal tapi bulananya
1	saya beri bintang 5 karena acara2 dinetflix se...
2	ko saya gabisa masuk aja ke aplikasinya ya ken...
3	pilihan film nya sedikit jelek
4	jelek amat udh bayar tapih gk bisa masuk lagi ...
..	...
994	saya sudah menginstal aplikasi ini tapi pas pe...
996	gua udah melakukan berlangganan tapi gabisa ma...
997	susah banget loginnya
998	njir lah bayar
999	eror terus

[906 rows x 1 columns]

Gambar 5. Hasil dari tahap Cleaning dan Case Folding

Gambar 5. Hasil dari tahap Cleaning dan Case Folding menunjukkan perbedaan antara teks ulasan sebelum dan sesudah dilakukan pembersihan. Setelah tahap ini selesai, teks akan lebih bersih dan siap untuk diproses dalam tahap tokenisasi serta analisis lebih lanjut.

4.4. Tokenizing, Stopword, dan Stemming

Tahap selanjutnya dalam reprocessing adalah tokenizing, stopwords removal, dan stemming. Tahapan ini bertujuan untuk menyederhanakan teks sehingga lebih optimal dalam analisis sentimen.

```

                                processed_content
0          [bagusvideonya, hd, mahal, bulananya]
1  [bintang, 5, acara2, dinetflix, seru, netflix,...
2  [ko, gabisa, masuk, aja, aplikasi, ya, udah, u...
3          [pilih, film, nya, jelek]
4  [jelek, udh, bayar, tapih, gk, masuk, netflix,...
..
994 [menginstal, aplikasi, pas, daftar, suruh, lan...
996 [gua, udah, langgan, gabisa, masuk, busett, du...
997          [susah, banget, loginnya]
998          [njir, bayar]
999          [eror]

[906 rows x 2 columns]

```

Gambar 6. Hasil Tokenizing, Stopword, dan Stemming

Gambar diatas menunjukkan perubahan teks ulasan sebelum dan sesudah melalui tahapan ini. Setelah preprocessing selesai, teks akan siap untuk tahap selanjutnya, yaitu pembobotan dengan TF-IDF (Term Frequency - Inverse Document Frequency) sebelum diterapkan ke model klasifikasi.

4.5. Pembobotan TF-IDF

Semua teks ulasan yang telah melalui tahap tokenisasi, stopword removal, dan stemming dikonversi ke dalam representasi numerik menggunakan TF-IDF. Kata-kata yang memiliki bobot TF-IDF tinggi menunjukkan bahwa kata tersebut memiliki pengaruh besar terhadap klasifikasi sentimen.

Data hasil pembobotan ini akan digunakan sebagai masukan (input features) dalam model machine learning untuk analisis sentimen lebih lanjut.

```

Shape of TF-IDF matrix: (906, 1931)
TF-IDF features: ['01' '10' '100' ... 'zaman' 'zetvideo' 'zoooonk']
[[0. 0. 0. ... 0. 0. 0.]]

```

Gambar 7. Hasil Pembobotan TF-IDF

Hasil Pembobotan TF-IDF menunjukkan matriks TF-IDF yang dihasilkan, di mana setiap baris mewakili satu ulasan dan setiap kolom mewakili kata-kata unik yang telah diproses.

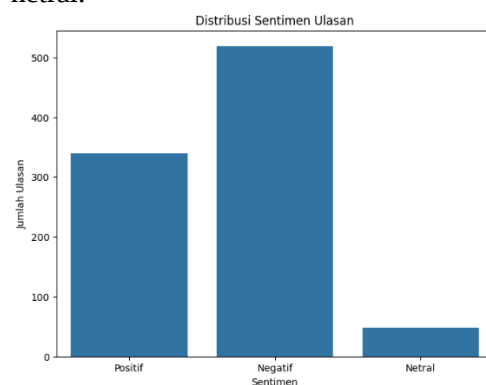
4.6. Visualisasi Data

Visualisasi data untuk memahami distribusi dan karakteristik dari dataset ulasan yang telah dikumpulkan. Visualisasi ini bertujuan untuk menggambarkan pola sentimen pengguna terhadap layanan Netflix serta membantu dalam analisis lebih lanjut sebelum proses klasifikasi.

Gambar 8 menghasilkan ulasan dengan sentimen negatif memiliki jumlah tertinggi dibandingkan kategori lainnya, dengan lebih dari 500 ulasan. Hal ini menunjukkan bahwa sebagian besar pengguna memberikan ulasan yang kurang puas terhadap layanan Netflix.

Kemudian, sentimen positif berjumlah sekitar 300 ulasan, yang menunjukkan bahwa masih ada banyak pengguna yang memberikan apresiasi terhadap layanan Netflix. Hal ini bisa mencerminkan kepuasan pengguna terhadap fitur tertentu seperti kualitas video, pilihan film, atau pengalaman streaming.

Selanjutnya, ulasan dengan sentimen netral memiliki jumlah yang paling sedikit, kurang dari 50 ulasan. Ini menunjukkan bahwa kebanyakan pengguna cenderung memiliki opini yang jelas, baik positif maupun negatif, dibandingkan memberikan ulasan yang bersifat netral.



Gambar 8. Hasil Visualisasi Berdasarkan Label

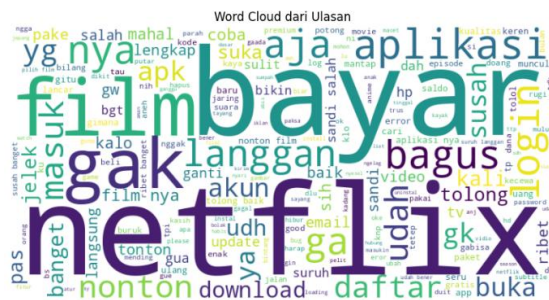
Dari hasil distribusi sentimen ini, dapat dilihat bahwa layanan Netflix mendapat lebih banyak ulasan negatif daripada positif. Analisis lebih lanjut dapat dilakukan untuk mengidentifikasi faktor-faktor utama yang menyebabkan sentimen negatif mendominasi, seperti harga berlangganan, kualitas konten, atau masalah teknis dalam aplikasi.

Hasil Visualisasi WordCloud menampilkan kumpulan kata yang paling sering muncul dalam ulasan pengguna terhadap layanan Netflix. WordCloud adalah teknik visualisasi teks yang menunjukkan frekuensi kemunculan kata dalam suatu kumpulan dokumen. Kata yang muncul lebih sering akan ditampilkan dengan ukuran yang lebih besar, sedangkan kata yang jarang muncul akan berukuran lebih kecil.

Dari hasil visualisasi wordcloud dibawah ini, beberapa kata yang dominan muncul dalam

ulasan pengguna kata-kata paling sering muncul yaitu kata "Netflix", "film", "bayar", dan "langganan" tampak memiliki ukuran terbesar, yang menunjukkan bahwa topik ini paling sering disebut dalam ulasan. Hal ini wajar karena ulasan berfokus pada pengalaman pengguna terhadap layanan Netflix, termasuk sistem berlangganan dan konten film yang disediakan.

Kata-kata yang mengindikasikan keluhan pengguna seperti kata "gak", "susah", "error", "mahal", dan "log" juga cukup besar, yang menunjukkan bahwa banyak pengguna mengalami masalah saat menggunakan layanan netflix. kemungkinan besar, ulasan negatif terkait dengan kendala login, harga berlangganan yang dianggap mahal, serta error dalam aplikasi. Sedangkan kata-kata yang mengindikasikan apresiasi pengguna seperti kata "bagus", "nonton", dan "aplikasi" muncul dalam ukuran yang cukup besar, yang berarti ada juga ulasan positif terkait kualitas film atau aplikasi netflix secara umum.



Gambar 9. Hasil Visualiasi WordCloud

Sedangkan untuk Gambar 10 skor ulasan berkisar dari 1 hingga 5, di mana skor 5 menunjukkan ulasan sangat positif, dan skor 1 menunjukkan ulasan sangat negatif. Sebagian besar ulasan tersebar di skor 1, 2, dan 5, menunjukkan adanya dominasi ulasan sangat positif dan sangat negatif, sementara skor tengah (3 dan 4) lebih jarang diberikan. Ulasan panjang lebih sering muncul pada skor 1 dan 5, yang menunjukkan bahwa pengguna yang sangat puas atau sangat kecewa lebih cenderung memberikan ulasan panjang.

Dari gambar 10 dibawah ini hasil analisis ini, dapat dijelaskan bahwa pengguna Netflix cenderung memberikan ulasan yang singkat, terutama jika mereka memberikan skor menengah (2–4). Sementara itu, pengguna yang sangat puas (skor 5) atau sangat kecewa (skor

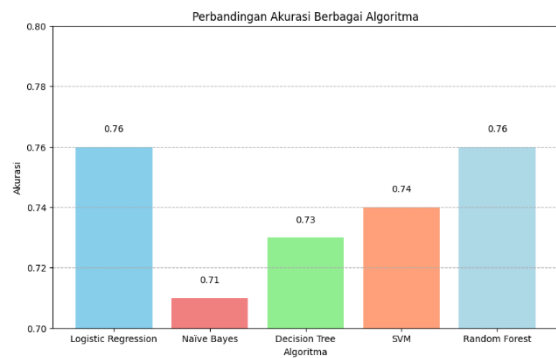
1) lebih cenderung menuliskan ulasan yang lebih panjang untuk mengekspresikan pendapat mereka.



Gambar 10. Korelasi Panjang Ulasan dan Skor

4.7. Hasil Pengujian Algoritma

Berikut ini hasil akurasi dari berbagai algoritma klasifikasi yang digunakan untuk ulasan Netflix.



Gambar 11. Perbandingan Hasil Pengujian

Algoritma

Analisis hasil pengujian Logistic Regression dan Random Forest memiliki akurasi tertinggi sebesar 76% Ini menunjukkan bahwa kedua algoritma ini memiliki performa terbaik dalam mengklasifikasikan sentimen ulasan. SVM (Support Vector Machine) memiliki akurasi 74% yang sedikit lebih rendah dibandingkan Logistic Regression dan Random Forest, tetapi masih cukup kompetitif. Decision Tree mendapatkan akurasi 73% yang lebih rendah dibandingkan SVM. Ini mengindikasikan bahwa model pohon keputusan mungkin mengalami overfitting atau kurang optimal dalam menangani fitur dari ulasan. Naïve Bayes memiliki akurasi terendah sebesar 71% yang menunjukkan bahwa metode

ini kurang efektif dalam menangani data ulasan Netflix, kemungkinan karena asumsi independensi antar fitur yang kurang sesuai dengan data teks yang sering memiliki keterkaitan antar kata.

Logistic Regression dan Random Forest menjadi pilihan terbaik untuk klasifikasi sentimen dalam penelitian ini karena memberikan akurasi tertinggi. Naïve Bayes memiliki performa paling rendah, yang mungkin disebabkan oleh ketergantungan antar kata dalam teks yang tidak dapat ditangani dengan baik oleh asumsi algoritma ini. SVM dan Decision Tree masih menunjukkan performa yang cukup baik, tetapi tidak lebih unggul dibandingkan Logistic Regression dan Random Forest.

5. KESIMPULAN

Berdasarkan hasil penelitian yang telah dilakukan, analisis terhadap ulasan pengguna aplikasi Netflix menunjukkan bahwa mayoritas ulasan bersentimen negatif, diikuti oleh ulasan positif, dan hanya sedikit yang bersifat netral.

Berdasarkan pengujian yang telah dilakukan, hasil menunjukkan bahwa algoritma terbaik dengan akurasi tertinggi untuk klasifikasi ulasan pengguna Netflix adalah Random Forest, dengan akurasi mencapai 0.76. Algoritma ini menunjukkan performa lebih optimal dibandingkan Logistic Regression, Naïve Bayes, SVM, dan Decision Tree dalam menganalisis data ulasan yang telah diproses. Oleh karena itu, Random Forest direkomendasikan sebagai model terbaik dalam analisis sentimen ulasan pengguna karena kemampuannya dalam menangani data teks dengan karakteristik yang kompleks. Model ini dapat diterapkan dalam sistem analitik otomatis untuk memahami opini pelanggan secara lebih mendalam, sehingga pengembang aplikasi dapat menyesuaikan strategi layanan yang lebih sesuai dengan kebutuhan pengguna.

UCAPAN TERIMA KASIH

Penulis menyampaikan apresiasi kepada semua pihak yang telah memberikan dukungan dalam pelaksanaan penelitian ini. Semoga segala bantuan yang telah diberikan mendapatkan balasan yang setimpal, dan hasil penelitian ini dapat memberikan manfaat bagi pengembangan ilmu pengetahuan serta aplikasi di bidang yang relevan.

DAFTAR PUSTAKA

- [1] N. T. Romadloni and W. Supriyanti, "Analisis Sentimen Penggunaan Teknologi Pada Pendidikan Anak Usia Dini," *J. Ilm. SINUS*, vol. 21, no. 2, p. 101, 2023, doi: 10.30646/sinus.v21i2.759.
- [2] M. D. Hendriyanto, A. A. Ridha, and U. Enri, "Analisis Sentimen Ulasan Aplikasi Mola Pada Google Play Store Menggunakan Algoritma Support Vector Machine," *INTECOMS J. Inf. Technol. Comput. Sci.*, vol. 5, no. 1, pp. 1–7, 2022, doi: 10.31539/intecomsv5i1.3708.
- [3] B. Ramadhani and R. R. Suryono, "Komparasi Algoritma Naïve Bayes dan Logistic Regression Untuk Analisis Sentimen Metaverse," *J. Media Inform. Budidarma*, vol. 8, no. 2, p. 714, 2024, doi: 10.30865/mib.v8i2.7458.
- [4] N. T. Romadloni and Hilman F Pardede, "Seleksi Fitur Berbasis Pearson Correlation Untuk Optimasi Opinion Mining Review Pelanggan," *J. RESTI (Rekayasa Sist. dan Teknol. Informasi)*, vol. 3, no. 3, pp. 505–510, 2019, doi: 10.29207/resti.v3i3.1189.
- [5] G. Kanugrahan, V. Hafizh, C. Putra, and Y. Ramdhani, "Analisis Sentimen Aplikasi Gojek Menggunakan SVM, Random Forest dan Decision Tree," vol. 6, no. 2, 2024.
- [6] P. Verma, A. Dumka, A. Bhardwaj, and A. Ashok, "Product Review-Based Customer Sentiment Analysis Using an Ensemble of mRMR and Forest Optimization Algorithm (FOA)," *Int. J. Appl. Metaheuristic Comput.*, vol. 13, no. 1, pp. 1–21, 2022, doi: 10.4018/ijamc.2022010107.
- [7] H. Rahmawati and A. Sudrajat, "MAHASISWA BARU DI POLITEKNIK TEDC BANDUNG," vol. 13, no. 1, 2025.
- [8] R. Maulana, A. Voutama, and T. Ridwan, "Analisis Sentimen Ulasan Aplikasi MyPertamina pada Google Play Store menggunakan Algoritma NBC," *J. Teknol. Terpadu*, vol. 9, no. 1, pp. 42–48, 2023, doi: 10.54914/jtt.v9i1.609.
- [9] D. Ardiansyah, A. Saepudin, R. Aryanti, E. Fitriani, and Royadi, "Analisis Sentimen Review Pada Aplikasi Media Sosial Tiktok Menggunakan Algoritma K-Nn Dan Svm Berbasis Pso," *J. Inform. Kaputama*, vol. 7, no. 2, pp. 233–241, 2023, doi: 10.59697/jik.v7i2.148.
- [10] I. Aida Sapitri and M. Fikry, "Pengklasifikasian Sentimen Ulasan Aplikasi Whatsapp Pada Google Play Store Menggunakan Support Vector Machine," *J. TEKINKOM*, vol. 6, no. 1, pp. 1–7, 2023, doi: 10.37600/tekinkom.v6i1.773.

- [11] S. Nanda, D. Mualfah, and D. A. Fitri, "Analisis Sentimen Kepuasan Pengguna Terhadap Layanan Streaming Mola Menggunakan Algoritma Random Forest," *J. Apl. Teknol. Inf. dan Manaj.*, vol. 3, no. 2, pp. 210–219, 2022, doi: 10.31102/jatim.v3i2.1592.
- [12] A. F. Anjani, D. Anggraeni, and I. M. Tirta, "Implementasi Random Forest Menggunakan SMOTE untuk Analisis Sentimen Ulasan Aplikasi Sister for Students UNEJ," *J. Nas. Teknol. dan Sist. Inf.*, vol. 9, no. 2, pp. 163–172, 2023, doi: 10.25077/teknosi.v9i2.2023.163-172.
- [13] E. B. Susanto, Paminto Agung Christianto, Mohammad Reza Maulana, and Sattriedi Wahyu Binabar, "Analisis Kinerja Algoritma Naïve Bayes Pada Dataset Sentimen Masyarakat Aplikasi NEWSAKPOLE Samsat Jawa Tengah," *J. CoSciTech (Computer Sci. Inf. Technol.)*, vol. 3, no. 3, pp. 234–241, 2022, doi: 10.37859/coscitech.v3i3.4343.
- [14] Tania Puspa Rahayu Sanjaya, Ahmad Fauzi, and Anis Fitri Nur Masruriyah, "Analisis sentimen ulasan pada e-commerce shopee menggunakan algoritma naive bayes dan support vector machine," *INFOTECH J. Inform. Teknol.*, vol. 4, no. 1, pp. 16–26, 2023, doi: 10.37373/infotech.v4i1.422.
- [15] I. Syahrohim, S. D. Saputra, R. W. Saputra, V. H. Pranatawijaya, and R. Priskila, "Perbandingan Analisis Sentimen Setelah Pilpres 2024 Di Twitter Menggunakan Algoritma Machine Learning," *J. Inform. dan Tek. Elektro Terap.*, vol. 12, no. 2, 2024, doi: 10.23960/jitet.v12i2.4249.
- [16] H. Hakim, D. Kamil, and B. Alatas, "Pendekatan Machine Learning untuk Estimasi Harga Rumah dengan Regresi Linier," vol. 1, no. 1, pp. 18–22, 2025.