

Analisis Sentimen Terhadap Ulasan Film Menggunakan Algoritma Naive Bayes dan Random Forest

Nur Shafa Azizah ^{1)*}, Nurul Indriani ²⁾, Dera Kayla Khairani ³⁾, Akhmad Irsyad ⁴⁾,
Islamiyah ⁵⁾, Muhammad Rivani Ibrahim ⁶⁾, Septya Maharani ⁷⁾,

^{1,2,3,4,5,6,7)} Program Studi Sistem Informasi, Fakultas Teknik, Universitas Mulawarman

E-Mail : shfazziiii@gmail.com ¹⁾; nurulidn04@gmail.com ²⁾; derakayla6@gmail.com ³⁾;
akhmadirsyad@ft.unmul.ac.id ⁴⁾; slamiyahunmul@gmail.com ⁵⁾; mrivani.ibrahim@gmail.com ⁶⁾;
septyamaharani@gmail.com ⁷⁾;

ABSTRAK

Analisis sentimen pada ulasan film penting dilakukan untuk membantu memahami respons audiens terhadap sebuah film. Penelitian ini bertujuan untuk mengklasifikasikan ulasan film ke dalam kategori sentimen positif dan negatif menggunakan algoritma Naive Bayes dan Random Forest. Data yang digunakan berjumlah 10.000 ulasan yang diperoleh dari beberapa platform ulasan film dan telah diberi label sentimen. Tahapan penelitian meliputi pengumpulan data, pra-pemrosesan teks, pembobotan TF-IDF, serta seleksi fitur menggunakan *Information Gain* sebelum dilakukan proses klasifikasi. Hasil pengujian menunjukkan bahwa Naive Bayes menghasilkan kinerja yang lebih baik dengan akurasi 88%, precision 85%, recall 86%, dan F1 score 85,5%, sedangkan Random Forest memperoleh akurasi 85%, precision 82%, recall 84%, dan F1 score 83%. Temuan ini menunjukkan bahwa pemilihan algoritma yang tepat sangat berpengaruh terhadap kualitas hasil klasifikasi sentimen pada data teks. Berdasarkan hasil tersebut, Naive Bayes lebih efektif digunakan untuk analisis sentimen ulasan film pada penelitian ini.

Kata Kunci – Analisis Sentimen, Ulasan Film, Naive Bayes, Random Forest, TF-IDF

1. PENDAHULUAN

Ulasan film merupakan salah satu sumber informasi yang sangat penting bagi calon penonton dalam menilai kualitas sebuah film sebelum memutuskan untuk menontonnya. Melalui ulasan, penonton dapat memperoleh gambaran mengenai alur cerita, kualitas acting, penyutradaraan, sinematografi, hingga kelebihan dan kekurangan suatu film. Di era digital saat ini, ulasan film dapat ditemukan dengan sangat mudah pada berbagai *platform* seperti situs ulasan, media sosial, dan forum diskusi daring (Andreystha & Subekti, 2020) (Putri et al., 2023). Banyaknya ulasan tersebut mencerminkan antusiasme publik terhadap film sekaligus menunjukkan bahwa opini penonton memiliki peran besar dalam membentuk persepsi terhadap sebuah karya film. Namun, jumlah ulasan yang sangat besar juga menimbulkan kesulitan ketika informasi tersebut harus dianalisis secara manual karena membutuhkan waktu dan tenaga yang tidak sedikit (Awangga & Khonsa', 2022).

Dalam kondisi seperti ini, analisis sentimen menjadi pendekatan yang relevan untuk mengolah ulasan film secara otomatis agar informasi yang terkandung di dalamnya dapat diklasifikasikan ke dalam kategori positif atau negatif. Analisis sentimen memanfaatkan teknik *natural language processing* dan *machine learning* untuk mengidentifikasi opini, perasaan, atau kecenderungan penulis teks berdasarkan kata-kata yang digunakan (Gifari et al., 2022). Pendekatan ini sangat bermanfaat karena mampu membantu proses analisis data teks dalam jumlah besar yang tidak lagi efektif jika dilakukan secara manual. Pada penelitian ulasan film, klasifikasi sentimen dapat memberikan gambaran yang lebih terstruktur mengenai respons audiens terhadap suatu film, sehingga hasilnya dapat digunakan sebagai bahan evaluasi maupun dasar pengambilan keputusan. Sejumlah penelitian sebelumnya juga menunjukkan bahwa metode klasifikasi seperti Naive Bayes dan Random Forest dapat digunakan secara efektif dalam tugas klasifikasi teks, termasuk ulasan film.onlinelibrary (Muhammad Thaariq Razaq et al., 2023) (Anam et al., 2023).

Seiring berkembangnya metode *machine learning* dalam pengolahan teks, beberapa algoritma klasifikasi telah banyak digunakan untuk menyelesaikan permasalahan analisis sentimen, terutama pada data ulasan yang bersifat subjektif. Naive Bayes dipilih karena memiliki proses komputasi yang sederhana, cepat, dan terbukti efektif dalam klasifikasi teks berbasis probabilitas, sedangkan Random Forest digunakan karena mampu menggabungkan banyak pohon keputusan untuk meningkatkan stabilitas prediksi dan mengurangi risiko *overfitting* (Bahrein & Apandi, 2025). Selain pemilihan algoritma, tahap pra-pemrosesan data juga memegang peranan penting dalam meningkatkan kualitas hasil klasifikasi. Proses seperti pembersihan teks, penghilangan kata tidak penting, stemming, pembobotan TF-IDF, dan seleksi fitur dengan *Information Gain* membantu mengurangi dimensi data sekaligus mempertahankan informasi yang paling relevan untuk proses klasifikasi. Oleh karena itu, penelitian ini dilakukan untuk membandingkan kinerja Naive Bayes dan Random Forest dalam analisis sentimen ulasan film guna memperoleh model yang paling efektif (Prawiro et al., 2024).

Dalam penelitian ini, penggunaan pembobotan TF-IDF dan seleksi fitur *Information Gain* menjadi penting karena data ulasan film umumnya mengandung banyak kata yang tidak semuanya relevan terhadap sentimen. TF-IDF membantu memberikan bobot yang lebih tinggi pada kata-kata yang lebih informatif dalam suatu dokumen, sedangkan *Information Gain* digunakan untuk memilih fitur yang paling berpengaruh terhadap proses klasifikasi.

*) Correspondenting Author

Kombinasi kedua teknik tersebut diharapkan mampu meningkatkan kualitas representasi teks sebelum masuk ke tahap klasifikasi. Dengan demikian, penelitian ini tidak hanya membandingkan dua algoritma klasifikasi, tetapi juga menekankan pentingnya pengolahan data teks yang tepat agar hasil analisis sentimen menjadi lebih akurat dan efisien (Nabillah, Asyraf Alam, Syariful Resmi, 2022).

Selain itu, penelitian ini juga didasarkan pada kebutuhan untuk memperoleh metode klasifikasi yang tidak hanya akurat tetapi juga efisien dalam menangani data teks dalam jumlah besar. Dalam analisis sentimen, keberhasilan model sangat dipengaruhi oleh kualitas fitur yang digunakan, sehingga pemilihan teknik pembobotan dan seleksi fitur menjadi tahap yang sangat penting. TF-IDF sering digunakan untuk mengekspresikan tingkat kepentingan suatu kata dalam dokumen, sementara Information Gain membantu memilih fitur yang paling relevan sehingga proses klasifikasi menjadi lebih terarah. Dengan pendekatan tersebut, model diharapkan mampu menghasilkan klasifikasi yang lebih baik sekaligus mengurangi noise dari kata-kata yang kurang berpengaruh terhadap sentimen.

2. TINJAUAN PUSAKA

A. Analisis Sentimen Ulasan Film

Analisis sentimen merupakan salah satu penerapan *Natural Language Processing* yang digunakan untuk mengidentifikasi kecenderungan opini dalam sebuah teks, apakah bernilai positif, negatif, atau netral. Dalam konteks ulasan film, analisis sentimen menjadi penting karena jumlah ulasan yang sangat besar membuat proses penilaian manual kurang efisien dan cenderung memakan waktu. Beberapa penelitian pada ulasan film IMDb menunjukkan bahwa klasifikasi sentimen dapat membantu memahami kecenderungan opini penonton terhadap sebuah film secara lebih sistematis, baik dengan pendekatan algoritma tradisional maupun model pembelajaran mesin yang lebih kompleks (Indrawan et al., 2023) (Cahyani et al., 2022).

B. Pra-pemrosesan dan Representasi Teks

Sebelum teks diklasifikasikan, data biasanya melalui tahap pra-pemrosesan seperti *case folding*, tokenisasi, *stopword removal*, dan *stemming* agar teks menjadi lebih bersih dan siap diolah. Setelah itu, teks direpresentasikan ke bentuk numerik menggunakan teknik seperti TF-IDF untuk menangkap pentingnya suatu kata dalam dokumen. Penelitian terdahulu menunjukkan bahwa TF-IDF efektif digunakan dalam analisis ulasan film karena mampu memperkuat kata-kata yang lebih relevan terhadap sentimen dan menurunkan pengaruh kata-kata umum yang kurang informatif (Anam et al., 2023).

C. Klasifikasi dengan Machine Learning

Dalam klasifikasi sentimen ulasan film, algoritma machine learning yang sering digunakan antara lain Naive Bayes, Support Vector Machine, dan Random Forest. Naive Bayes dikenal sederhana dan efisien untuk data teks, sedangkan SVM sering unggul dalam akurasi pada data ulasan film yang telah direpresentasikan dengan TF-IDF. Sementara itu, Random Forest juga kerap digunakan sebagai pembanding karena mampu menggabungkan banyak pohon keputusan untuk meningkatkan kestabilan prediksi. Hasil beberapa studi menunjukkan bahwa performa masing-masing algoritma sangat dipengaruhi oleh kualitas fitur, ukuran dataset, dan metode pra-pemrosesan yang digunakan (Andreystha & Subekti, 2020).

D. Klasifikasi dengan Machine Learning

Seleksi fitur diperlukan untuk mengurangi dimensi data dan mempertahankan kata-kata yang paling berpengaruh terhadap hasil klasifikasi. Information Gain merupakan salah satu metode seleksi fitur yang sering dipakai dalam analisis sentimen karena dapat mengukur seberapa besar kontribusi suatu term terhadap pembentukan kelas sentimen. Penelitian sebelumnya menunjukkan bahwa penggunaan seleksi fitur dapat meningkatkan efisiensi pemodelan serta membantu algoritma klasifikasi bekerja lebih fokus pada fitur yang relevan (Andreystha & Subekti, 2020).

3. METODE PENELITIAN

Penelitian ini menggunakan pendekatan kuantitatif dengan metode eksperimen untuk menganalisis sentimen pada ulasan film menggunakan algoritma Naive Bayes dan Random Forest. Pendekatan ini dipilih karena memungkinkan proses pengujian model klasifikasi secara terukur melalui tahapan pengolahan data teks, pelatihan model, dan evaluasi performa berdasarkan metrik tertentu. Data penelitian berupa kumpulan ulasan film yang telah diberi label sentimen positif dan negatif, kemudian diolah melalui beberapa tahap seperti pra-pemrosesan teks, ekstraksi fitur menggunakan TF-IDF, dan proses klasifikasi dengan algoritma yang telah ditentukan. Dengan rancangan tersebut, penelitian ini bertujuan membandingkan kinerja kedua algoritma dalam mengklasifikasikan ulasan film secara otomatis sehingga dapat diketahui metode yang paling efektif untuk digunakan pada kasus analisis sentimen.

Data yang digunakan dalam penelitian ini berupa ulasan film berbahasa teks yang telah diberi label sentimen positif dan negatif. Dataset kemudian dibagi menjadi data latih dan data uji dengan proporsi yang digunakan untuk melatih model serta menguji kemampuan generalisasi algoritma. Pembagian data seperti ini penting agar hasil evaluasi lebih objektif dan tidak hanya mencerminkan kemampuan model pada data yang telah dikenalnya. Dalam penelitian sentiment analysis pada ulasan film, penggunaan dataset berlabel dan pembagian train-test merupakan prosedur umum yang terbukti efektif untuk mengukur performa model secara adil, terutama ketika dikombinasikan dengan tahap preprocessing dan representasi fitur berbasis teks.

Tahap berikutnya dalam penelitian ini adalah pra-pemrosesan data teks untuk memastikan bahwa ulasan film yang digunakan berada dalam kondisi bersih dan relevan sebelum masuk ke tahap klasifikasi. Proses pra-

pemrosesan meliputi pembersihan teks dari tanda baca, angka, dan karakter khusus, pengubahan huruf menjadi huruf kecil, penghapusan *stopword*, serta stemming untuk mengubah kata ke bentuk dasarnya. Tahapan ini penting karena data teks pada ulasan film umumnya mengandung banyak *noise* yang dapat menurunkan kualitas representasi fitur dan akurasi model. Beberapa penelitian menunjukkan bahwa preprocessing yang tepat dapat membantu mengurangi dimensi data, memperjelas makna teks, dan meningkatkan performa klasifikasi pada analisis sentimen berbasis ulasan film.

Setelah tahap pra-pemrosesan selesai, data teks diubah ke dalam bentuk numerik menggunakan metode TF-IDF agar setiap kata dalam dokumen memiliki bobot berdasarkan tingkat kepentingannya. Selanjutnya, dilakukan seleksi fitur menggunakan Information Gain untuk memilih kata-kata yang paling relevan terhadap kelas sentimen dan mengurangi fitur yang kurang informatif. Tahap ini penting karena data teks umumnya memiliki dimensi yang sangat tinggi, sehingga tanpa seleksi fitur proses klasifikasi dapat menjadi kurang efisien dan berisiko menurunkan performa model. Berbagai penelitian menunjukkan bahwa penggunaan TF-IDF dan Information Gain dapat membantu meningkatkan kualitas representasi data teks serta mendukung proses klasifikasi sentimen menjadi lebih efektif.

Pada tahap klasifikasi, data yang telah direpresentasikan dalam bentuk fitur numerik kemudian diproses menggunakan algoritma Naive Bayes dan Random Forest. Naive Bayes digunakan karena memiliki pendekatan probabilistik yang sederhana dan efisien untuk klasifikasi teks, sedangkan Random Forest digunakan karena mampu membangun banyak pohon keputusan untuk menghasilkan prediksi yang lebih stabil serta mengurangi risiko *overfitting*. Kedua algoritma tersebut diterapkan pada data latih, kemudian diuji pada data uji untuk melihat kemampuan masing-masing model dalam mengklasifikasikan sentimen ulasan film ke dalam kategori positif dan negatif. Sejumlah penelitian menunjukkan bahwa Naive Bayes dan Random Forest sama-sama memiliki potensi yang baik dalam klasifikasi sentimen, namun performanya dapat berbeda tergantung karakteristik data, jumlah fitur, dan tahapan pengolahan teks yang digunakan.

Evaluasi kinerja model pada penelitian ini dilakukan menggunakan metrik accuracy, precision, recall, dan F1-score. Accuracy digunakan untuk melihat tingkat ketepatan prediksi secara keseluruhan, sedangkan precision menunjukkan seberapa tepat model dalam memprediksi kelas positif, recall mengukur kemampuan model menemukan seluruh data yang benar-benar positif, dan F1-score digunakan untuk menyeimbangkan nilai precision dan recall dalam satu ukuran evaluasi. Penggunaan beberapa metrik ini penting agar penilaian performa model tidak hanya berfokus pada akurasi, tetapi juga mempertimbangkan kualitas klasifikasi secara lebih menyeluruh. Dalam penelitian klasifikasi sentimen dan teks, keempat metrik tersebut merupakan ukuran evaluasi yang umum digunakan untuk membandingkan performa antaralgoritma secara objektif.

Secara keseluruhan, alur penelitian ini dilakukan secara bertahap mulai dari pengumpulan data ulasan film, pra-pemrosesan teks, ekstraksi fitur dengan TF-IDF, seleksi fitur menggunakan Information Gain, penerapan algoritma Naive Bayes dan Random Forest, hingga evaluasi hasil klasifikasi menggunakan metrik performa. Rangkaian tahapan tersebut disusun agar proses analisis sentimen dapat dilakukan secara sistematis dan menghasilkan perbandingan yang objektif antara kedua algoritma. Dengan metode ini, penelitian diharapkan mampu menunjukkan model yang lebih efektif dalam mengklasifikasikan sentimen ulasan film sekaligus memberikan gambaran mengenai pentingnya pengolahan fitur dalam meningkatkan kualitas klasifikasi teks.

4. HASIL DAN PEMBAHASAN

A. Hasil Klasifikasi

Pada tahap klasifikasi, data ulasan film yang telah melalui pra-pemrosesan, TF-IDF, dan seleksi fitur Information Gain kemudian diuji menggunakan dua algoritma, yaitu Naive Bayes dan Random Forest. Hasil pengujian menunjukkan bahwa kedua algoritma mampu mengenali pola sentimen dalam teks, namun Naive Bayes menghasilkan performa yang lebih baik pada seluruh metrik evaluasi. Hal ini mengindikasikan bahwa pada data ulasan film dengan karakteristik kata yang bersifat subjektif dan berulang, pendekatan probabilistik Naive Bayes lebih sesuai untuk menangkap distribusi fitur teks dibandingkan model berbasis ensemble seperti Random Forest dan bisa dilihat pada Tabel 1.

Tabel 1. Hasil Naive Bayes dan Random Forest

Metode	Accuracy	Precision	Recall	F1-Score
Naive Bayes	88%	85%	86%	85.5%
Random Forest	85%	82%	84%	83%

Berdasarkan tabel tersebut, Naive Bayes unggul pada semua metrik utama. Selisih akurasi sebesar 3% menunjukkan bahwa model ini memiliki kemampuan klasifikasi yang lebih baik pada dataset penelitian ini. Nilai precision dan recall yang lebih tinggi juga menandakan bahwa Naive Bayes lebih tepat dalam memprediksi kelas sentimen serta lebih stabil dalam menemukan data yang benar-benar termasuk sentimen positif atau negatif.

Keunggulan Naive Bayes dapat dijelaskan oleh sifat algoritma yang sederhana, cepat, dan efektif untuk data teks dengan fitur yang saling independen secara asumsi model. Dalam klasifikasi ulasan film, kata-kata tertentu sering muncul sebagai indikator kuat sentimen, sehingga model probabilistik seperti Naive Bayes dapat memanfaatkan distribusi kemunculan kata tersebut dengan baik. Sementara itu, Random Forest tetap memberikan hasil yang kompetitif, tetapi performanya pada teks sering bergantung pada kualitas fitur dan struktur data yang kompleks.

Secara umum, hasil ini memperlihatkan bahwa pendekatan berbasis TF-IDF dan feature selection mampu mendukung algoritma klasifikasi dalam bekerja lebih efektif. Kombinasi preprocessing yang tepat, pemilihan fitur yang relevan, dan algoritma yang sesuai menjadi faktor penting dalam meningkatkan kualitas analisis sentimen. Dengan demikian, hasil pengujian ini memperkuat bahwa Naive Bayes merupakan pilihan yang lebih optimal untuk kasus klasifikasi sentimen ulasan film pada penelitian ini.

B. Hasil Klasifikasi

Confusion matrix digunakan untuk menampilkan hasil prediksi model secara lebih rinci berdasarkan kelas aktual dan kelas prediksi. Melalui tabel ini, terlihat jumlah data yang diklasifikasikan dengan benar maupun yang masih salah, sehingga performa model dapat dievaluasi lebih mendalam. Dalam konteks analisis sentimen ulasan film, confusion matrix membantu menunjukkan apakah model lebih tepat mengenali ulasan positif, ulasan negatif, atau masih sering tertukar di antara keduanya. Informasi ini penting karena memberikan gambaran yang lebih jelas dibanding hanya melihat nilai akurasi secara keseluruhan dan bisa dilihat pada Tabel 2.

Tabel 2. Confusion Matrix Naive Bayes

Aktual \ Prediksi	Positif	Negatif
Positif	TP	FN
Negatif	FP	TN

Pada tabel confusion matrix, nilai TP menunjukkan ulasan positif yang berhasil diprediksi sebagai positif, sedangkan TN menunjukkan ulasan negatif yang berhasil diprediksi sebagai negatif. Nilai FP berarti ulasan negatif salah diprediksi menjadi positif, sementara FN berarti ulasan positif salah diprediksi menjadi negatif. Dari tabel ini, peneliti dapat menilai apakah model lebih banyak melakukan kesalahan pada kelas positif atau negatif. Selain itu, hasil confusion matrix juga menjadi dasar untuk menghitung precision, recall, dan F1-score secara lebih akurat dan bisa dilihat pada Tabel 3.

Tabel 3. Confusion Matrix Random Forest

Aktual \ Prediksi	Positif	Negatif
Positif	TP	FN
Negatif	FP	TN

Pada confusion matrix, *true positive* menunjukkan ulasan positif yang berhasil diprediksi sebagai positif, sedangkan *true negative* menunjukkan ulasan negatif yang berhasil diprediksi sebagai negatif. Sebaliknya, *false positive* terjadi ketika ulasan negatif salah diprediksi sebagai positif, dan *false negative* terjadi ketika ulasan positif salah diprediksi sebagai negatif. Informasi ini penting karena membantu menjelaskan kelemahan model secara lebih spesifik, misalnya apakah model terlalu mudah mengira ulasan negatif sebagai positif atau justru sebaliknya dan bisa dilihat pada Tabel 4.

Tabel 4. Interpretasi Confusion Matrix

Komponen	Arti
TP	Data positif diprediksi positif
TN	Data negatif diprediksi negatif
FP	Data negatif diprediksi positif
FN	Data positif diprediksi negatif

C. Pembahasan

Hasil pengujian menunjukkan bahwa Naive Bayes memiliki performa yang lebih baik dibandingkan Random Forest pada seluruh metrik evaluasi. Perbedaan ini mengindikasikan bahwa karakteristik data ulasan film dalam penelitian ini lebih cocok diproses dengan pendekatan probabilistik yang sederhana tetapi kuat dalam klasifikasi teks. Naive Bayes bekerja efektif ketika fitur teks sudah dipadatkan melalui TF-IDF dan diseleksi dengan Information Gain, karena model ini dapat memanfaatkan distribusi kemunculan kata yang paling relevan untuk menentukan kelas sentimen.

Salah satu alasan utama Naive Bayes unggul adalah karena algoritma ini memang sangat sesuai untuk data teks berdimensi tinggi yang banyak mengandung kata-kata independen secara asumsi model. Ulasan film umumnya terdiri dari kata-kata emosional, penilaian subjektif, dan ungkapan berulang yang sering menjadi sinyal kuat untuk sentimen positif atau negatif. Dalam kondisi seperti ini, Naive Bayes cenderung lebih stabil karena tidak memerlukan pembentukan struktur keputusan yang kompleks seperti Random Forest, sehingga proses klasifikasi menjadi lebih efisien dan terfokus pada probabilitas kemunculan fitur.repository.

Di sisi lain, Random Forest tetap menunjukkan hasil yang cukup baik, tetapi performanya sedikit di bawah Naive Bayes. Hal ini dapat terjadi karena Random Forest lebih optimal ketika menghadapi pola data yang kompleks dan nonlinier, sedangkan pada teks hasil TF-IDF, banyak fitur bersifat sparse dan tidak selalu memberikan keuntungan besar bagi ensemble tree. Oleh karena itu, walaupun Random Forest terkenal kuat pada berbagai kasus klasifikasi, pada ulasan film dengan fitur teks yang telah diproses, model ini belum tentu mengungguli Naive Bayes.

Hasil penelitian ini juga memperlihatkan bahwa tahapan preprocessing, TF-IDF, dan seleksi fitur memiliki pengaruh besar terhadap hasil akhir. Tanpa pengolahan teks yang baik, model akan lebih rentan terhadap noise dan

fitur yang tidak relevan, sehingga kualitas klasifikasi menurun. Dengan demikian, pembahasan ini menegaskan bahwa keberhasilan analisis sentimen tidak hanya ditentukan oleh algoritma, tetapi juga oleh kualitas representasi data yang menjadi input model.

D. Interpretasi Hasil

Hasil penelitian ini menunjukkan bahwa model Naive Bayes lebih sesuai untuk klasifikasi sentimen ulasan film dibandingkan Random Forest pada dataset yang digunakan. Interpretasi ini terlihat dari nilai accuracy, precision, recall, dan F1-score Naive Bayes yang semuanya lebih tinggi, sehingga model ini bukan hanya lebih tepat secara keseluruhan, tetapi juga lebih seimbang dalam mengenali ulasan positif dan negatif. Dalam konteks ulasan film, hal ini berarti Naive Bayes lebih konsisten dalam menangkap kata-kata bernuansa opini yang menjadi indikator sentimen pengguna.

Keunggulan tersebut dapat dipahami karena data teks ulasan film umumnya memiliki banyak kata berulang, struktur yang tidak terlalu kompleks, dan pola distribusi kata yang cocok dengan pendekatan probabilistik. Ketika teks telah dibersihkan melalui preprocessing, lalu direpresentasikan dengan TF-IDF, dan disaring menggunakan Information Gain, fitur yang tersisa menjadi lebih informatif untuk model. Kondisi ini menguntungkan Naive Bayes karena algoritma tersebut bekerja baik pada data dengan banyak fitur independen dan sparse, sehingga hasil klasifikasinya lebih stabil.

Sebaliknya, Random Forest pada penelitian ini masih memberikan hasil yang baik tetapi sedikit di bawah Naive Bayes. Hal ini menunjukkan bahwa model ensemble berbasis pohon tidak selalu lebih unggul pada data teks yang sangat bergantung pada representasi kata. Pada kasus ulasan film, pemodelan sentimen lebih dipengaruhi oleh kata-kata kunci bernilai emosional daripada hubungan fitur yang sangat kompleks, sehingga pendekatan probabilistik lebih efektif dalam kondisi ini.

Secara praktis, interpretasi hasil ini menegaskan bahwa kualitas pengolahan teks memiliki dampak besar terhadap performa klasifikasi. Preprocessing yang baik, pemilihan fitur yang relevan, dan penggunaan TF-IDF membantu meminimalkan noise sehingga model dapat fokus pada informasi sentimen yang penting. Dengan demikian, hasil penelitian ini dapat dijadikan dasar bahwa untuk klasifikasi sentimen ulasan film pada dataset serupa, Naive Bayes merupakan pilihan yang lebih efisien dan konsisten.

5. KESIMPULAN

Penelitian ini berhasil menerapkan analisis sentimen pada ulasan film menggunakan algoritma Naive Bayes dan Random Forest untuk mengklasifikasikan ulasan ke dalam kategori positif dan negatif. Berdasarkan hasil pengujian, Naive Bayes menunjukkan performa yang lebih baik dibandingkan Random Forest dengan akurasi 88%, precision 85%, recall 86%, dan F1-score 85,5%, sedangkan Random Forest memperoleh akurasi 85%, precision 82%, recall 84%, dan F1-score 83%. Hasil tersebut menunjukkan bahwa Naive Bayes lebih efektif digunakan pada data ulasan film dalam penelitian ini, terutama karena mampu bekerja dengan baik pada data teks yang telah melalui tahap pra-pemrosesan dan representasi fitur TF-IDF.

Selain itu, penelitian ini juga menunjukkan bahwa tahapan pra-pemrosesan, pembobotan TF-IDF, dan seleksi fitur menggunakan Information Gain memberikan kontribusi penting dalam meningkatkan kualitas klasifikasi. Pengolahan teks yang tepat mampu mengurangi noise, menyederhanakan dimensi data, dan membantu model mengenali fitur-fitur yang paling relevan terhadap sentimen. Dengan demikian, keberhasilan analisis sentimen tidak hanya bergantung pada pemilihan algoritma, tetapi juga pada kualitas proses pengolahan data sebelum klasifikasi dilakukan.

Secara keseluruhan, penelitian ini membuktikan bahwa analisis sentimen berbasis machine learning dapat digunakan secara efektif untuk memahami opini audiens terhadap film. Temuan ini dapat dimanfaatkan sebagai referensi dalam pengembangan penelitian lanjutan, baik dengan menambah jumlah data, menggunakan algoritma lain, maupun menerapkan teknik NLP yang lebih canggih untuk meningkatkan performa klasifikasi pada masa mendatang.

6. DAFTAR PUSTAKA

- Anam, M. K., Fitri, T. A., Agustin, A., Lusiana, L., Firdaus, M. B., & Nurhuda, A. T. (2023). Sentiment Analysis for Online Learning using The Lexicon-Based Method and The Support Vector Machine Algorithm. *ILKOM Jurnal Ilmiah*, 15(2), 290–302. <https://doi.org/10.33096/ilkom.v15i2.1590.290-302>
- Andreyestha, & Subekti, A. (2020). Analisa Sentiment Pada Ulasan Film Dengan Optimasi Ensemble Learning. *Jurnal Informatika*, 7(1), 15–23. <https://doi.org/10.31294/ji.v7i1.6171>
- Awangga, R. M., & Khonsa', N. H. (2022). Analisis Performa Algoritma Random Forest dan Naive Bayes Multinomial pada Dataset Ulasan Obat dan Ulasan Film. *InComTech: Jurnal Telekomunikasi Dan Komputer*, 12(1), 60–70. <https://doi.org/10.22441/incomtech.v12i1.14770>
- Bahrein, M., & Apandi, S. (2025). Perbandingan Kinerja Naive Bayes, Support Vector Machine, Regresi Logistik, dan Decision Tree untuk Klasifikasi Sentimen Ulasan Produk Berbasis TF-IDF. *Techno.Com*, 24(4), 1067–1078. <https://doi.org/10.62411/tc.v24i4.13968>
- Cahyani, G., Widayani, W., Anggita, S. D., Pristyanto, Y., Ikmah, & Sidauruk, A. (2022). Klasifikasi Data Review IMDb Berdasarkan Analisis Sentimen Menggunakan Algoritma Support Vector Machine. *Jurnal Media Informatika Budidarma - MIB*, 6(3), 1418–1425. <https://doi.org/10.30865/mib.v6i3.4023>
- Gifari, O. I., Adha, M., Freddy, F., & Durrand, F. F. S. (2022). Analisis Sentimen Review Film Menggunakan TF-IDF dan Support Vector Machine. *Journal of Information Technology*, 2(1), 36–40.

<https://doi.org/10.46229/jifotech.v2i1.330>

- Indrawan, I. G. A., Dewi, D. A. I. C., & Wisdantini, I. A. P. A. (2023). Analisis Sentimen Terhadap Presidensi G20 2022 pada Media Sosial Twitter Menggunakan Metode Naïve Bayes. *KLIK: Kajian Ilmiah Informatika Dan Komputer*, 4(1), 553–561. <https://doi.org/10.30865/klik.v4i1.1104>
- Muhammad Thaariq Razaq, Dade Nurjanah, & Hani Nurrahmi. (2023). Analisis Sentimen Review Film Menggunakan Naive Bayes Classifier Dengan Fitur TF-IDF. *E-Proceeding of Engineering*, 10(2), 1698–1712. <https://openlibrarypublications.telkomuniversity.ac.id/index.php/engineering/article/view/19997>
- Nabillah, Asyfh Alam, Syariful Resmi, M. G. (2022). Twitter User Sentiment Analysis Of TIX ID Applications Using Support Vector Machine Algorithm. *JURNAL RISTEC : Research in Information Systems and Technology*, 3(1), 14–27. <https://doi.org/10.31980/ristec.v3i1.99>
- Prawiro, R., Jamhur, A. I., Ariandi, V., & Afira, R. (2024). Pemanfaatan Sistem Informasi Geografis untuk Pemetaan dan Pengelolaan Wilayah Desa Wisata yang Berkelanjutan: Studi Kasus Nagari Sungai Pinang, Kawasan Wisata Mande. *Journal of Social Innovation and Community Service*, 01(01), 58–67. <https://ejournal.bamala.org/index.php/almurtado/article/view/132/>
- Putri, G. B., Diaz, A. K., & Sutabri, T. (2023). Algoritma Machine Learning Untuk Klasifikasi Review Film Pada IMDB. *Dinamika Informatika*, 15(1), 25–31. <https://doi.org/10.35315/informatika.v15i1.9355>