

PERBANDINGAN MODEL DECISION TREE, RANDOM FOREST, DAN SVM PADA ANALISIS SENTIMEN BERBASIS ASPEK KOMENTAR FILM JUMBO

Heru Teguh Apriyanto , Ahmad Abdul Chamid, Rizkysari Meymaharani

Teknik Informatika, Universitas Muria Kudus

Jalan Lingkar Utara, Gondangmanis, Bae, Gondangmanis, Bae, Kabupaten Kudus

heruteguhapriyan@gmail.com

ABSTRAK

Industri perfilman Indonesia, khususnya genre animasi, mengalami perkembangan signifikan dengan kehadiran film "Jumbo" (2024). *YouTube* sebagai platform utama diskusi film menghasilkan ribuan komentar tidak terstruktur yang menyulitkan pemahaman objektif terhadap persepsi publik pada aspek-aspek spesifik film seperti cerita, visual, musik, dukungan, dan perbandingan. Penelitian ini bertujuan mengidentifikasi aspek yang paling banyak dibahas dan membandingkan performa tiga algoritma *machine learning* (*Decision Tree*, *Random Forest*, dan *SVM*) dalam mengklasifikasikan sentimen berbasis aspek. Menggunakan *Aspect-Based Sentiment Analysis* (ABSA) dengan pendekatan *lexicon-based* untuk pelabelan sentimen. Sebanyak 7.906 komentar dikumpulkan dari lima kanal *YouTube*, diproses melalui *preprocessing*, identifikasi aspek berdasarkan kata kunci, pelabelan sentimen, dan ekstraksi fitur *TF-IDF*. Tiga model klasifikasi dilatih dan dievaluasi pada 4.082 komentar berlabel. *Decision Tree* mencapai performa terbaik dengan rata-rata akurasi 91,0%, tertinggi pada aspek cerita_emosi (99,6%) dan terendah pada musik (78,9%). Aspek dukungan_apresiasi paling dominan (1.797 komentar, 88,7% sentimen positif), mengindikasikan respons positif publik. Penelitian ini memberikan wawasan objektif persepsi audiens dan membuktikan efektivitas ABSA untuk analisis ulasan film Indonesia.

Kata kunci: *Aspect-Based Sentiment Analysis; Machine Learning; Komentar YouTube; Film Jumbo; Animasi Indonesia*

1. PENDAHULUAN

Industri perfilman nasional mengalami perkembangan signifikan dalam beberapa tahun terakhir, termasuk genre animasi yang mulai mendapat tempat tersendiri. Film "Jumbo" yang dirilis tahun 2024 menjadi salah satu karya animasi Indonesia yang menarik perhatian publik dan memunculkan diskusi luas mengenai berbagai aspek substansi film, mulai dari kualitas cerita, visual, musik, hingga dukungan terhadap industri animasi nasional [1]. *YouTube* telah menjadi platform utama bagi audiens untuk menyampaikan opini terhadap produk sinematik. Perkembangan teknologi informasi telah mendorong perubahan signifikan dalam cara masyarakat mengakses informasi dan berinteraksi di dunia digital, menjadikan *YouTube* sebagai platform yang paling banyak diakses oleh pengguna internet di Indonesia dengan persentase mencapai 88% [2], [3]. Platform ini menyediakan fasilitas kolom komentar yang memungkinkan pengguna menyampaikan tanggapan, pendapat, atau opini terhadap konten video. Komentar pada konten video ini sering kali merepresentasikan opini publik yang autentik terhadap isu-isu sosial dan budaya [4]. *YouTube* memiliki peran strategis dalam menyebarkan informasi sekaligus menjadi ruang interaksi publik yang aktif dan merupakan sumber data potensial karena sifatnya yang terbuka dan interaktif [5], [6]. Analisis sentimen terhadap komentar *YouTube* telah menjadi area penelitian yang signifikan dalam memahami persepsi publik [7].

Komentar penonton menyimpan data penting tentang pandangan publik terhadap berbagai dimensi film, mulai dari cerita, visual, musik, hingga aspek

emosional. Namun, jumlah komentar yang *masif* dan tidak terstruktur menyebabkan analisis manual menjadi sangat kompleks dan memerlukan waktu yang lama [8]–[10]. Volume data teks yang besar ini tidak mungkin dianalisis secara manual satu per satu tanpa bantuan teknologi analisis data berbasis komputer. Selain itu, data teks yang dihasilkan oleh masyarakat melalui media sosial seringkali bersifat tidak terstruktur dan mengandung *noise* seperti singkatan, *typo*, dan penggunaan bahasa informal [11]. Analisis sentimen merupakan bidang *Natural Language Processing* (NLP) yang fokus pada identifikasi dan klasifikasi opini dalam teks untuk mengetahui pola emosi [12]. *Machine learning* telah terbukti efektif dalam membangun sistem deteksi dan klasifikasi otomatis untuk berbagai *domain* [13]. Analisis sentimen adalah studi komputasi yang bertujuan mengevaluasi dan mengenali sentimen atau opini secara otomatis yang terkandung dalam teks [14]–[16], dan umumnya terdiri dari tiga jenis sentimen: positif, negatif, atau netral [17]. Metode ini dapat membantu memahami opini publik terhadap produk atau layanan dan merupakan bidang riset yang berkembang pesat karena kemampuannya menggali makna *subjektif* dari data teks secara otomatis dan terstruktur [2]. Namun, pendekatan tradisional yang hanya mengklasifikasikan sentimen global (positif, negatif, netral) belum memadai untuk menghasilkan *insight* detail tentang komponen spesifik film yang mendapat apresiasi atau kritik [18]. Analisis sentimen yang umumnya dilakukan pada tingkat dokumen tidak menunjukkan secara spesifik apa yang disukai dan tidak disukai oleh pengguna.

Aspect-Based Sentiment Analysis (ABSA) hadir sebagai solusi yang lebih tepat. ABSA merupakan salah satu jenis analisis sentimen yang merupakan proses di mana sentimen terhadap berbagai aspek atau fitur tertentu dari suatu entitas diekstraksi dan diklasifikasikan sehingga bisa memberikan sebuah wawasan mendalam [19], [20]. ABSA sangat penting dalam memahami opini pengguna secara detail dan *komprehensif* [21]. ABSA tidak hanya mengidentifikasi sentimen umum, tetapi juga mendeteksi aspek spesifik dalam komentar dan menetapkan sentimen untuk setiap aspek secara terpisah [22]. Penelitian Pranatawijaya et al.[22] menggunakan ABSA untuk analisis ulasan aplikasi *ride-hailing* dan berhasil mengungkap sentimen detail per aspek. Namun, penelitian tersebut belum membandingkan performa berbagai algoritma *machine learning*. Irawan et al.[15] melakukan studi perbandingan *Random Forest*, *Naive Bayes*, dan *SVM* pada analisis sentimen aplikasi *CapCut* dan menemukan *SVM* memberikan akurasi terbaik (86%), namun membutuhkan waktu pemrosesan yang lebih lama.

Berdasarkan gap penelitian di atas, penelitian ini bertujuan mengidentifikasi aspek yang paling banyak dibahas dalam komentar film Jumbo dan membandingkan performa *Decision Tree*, *Random Forest*, dan *SVM* dalam mengklasifikasikan sentimen berbasis aspek. Penelitian ini menggunakan metode ABSA dengan pendekatan *lexicon-based* untuk pelabelan sentimen otomatis, kemudian membandingkan tiga algoritma *machine learning* untuk menentukan model terbaik dalam konteks analisis sentimen film Indonesia. Dengan demikian, penelitian ini diharapkan dapat memberikan kontribusi berupa: (1) pemahaman objektif tentang persepsi audiens terhadap film Jumbo per aspek; (2) rekomendasi model klasifikasi terbaik untuk analisis sentimen berbasis aspek pada ulasan film; dan (3) wawasan untuk pengembangan industri film animasi Indonesia.

2. TINJAUAN PUSTAKA

Beberapa penelitian sebelumnya telah mengeksplorasi analisis sentimen pada berbagai *platform* media sosial, khususnya *YouTube*, dengan menggunakan berbagai algoritma *machine learning*. Berikut adalah tinjauan penelitian terdahulu yang relevan dengan penelitian ini:

Penelitian Muhayat melakukan analisis sentimen terhadap komentar video *YouTube* menggunakan *Support Vector Machine* (*SVM*). Penelitian tersebut berhasil memperoleh akurasi sebesar 86% dalam mengklasifikasikan sentimen komentar. *SVM* dipilih karena kemampuannya dalam menangani data berdimensi tinggi dan efektif dalam klasifikasi teks. Namun, penelitian tersebut masih berfokus pada penggunaan satu teknik pembobotan fitur yaitu *TF-IDF*, dan tidak melakukan perbandingan dengan algoritma lain. Keterbatasan lain adalah penelitian ini

menggunakan analisis sentimen tingkat dokumen tanpa mempertimbangkan aspek-aspek spesifik yang dibahas dalam komentar [5].

Ariyani melakukan *ekstraksi* pengetahuan dari ulasan aplikasi *CapCut* menggunakan metode *Aspect-Based Sentiment Analysis* (ABSA) dan klasifikasi. Penelitian ini menekankan bahwa *SVM* memberikan hasil yang baik dalam analisis sentimen berbasis aspek dengan akurasi yang kompetitif. Penelitian ini berhasil mengidentifikasi aspek-aspek spesifik yang dibahas pengguna dan memberikan insight detail tentang kelebihan dan kekurangan aplikasi. Namun, peneliti menyarankan agar di masa depan dilakukan perbandingan dengan algoritma lain seperti *Decision Tree* atau *Random Forest* untuk melihat metode mana yang lebih efektif dalam konteks ABSA. Keterbatasan lain adalah jumlah data yang digunakan masih terbatas sehingga perlu diperluas untuk generalisasi yang lebih baik [19].

Teguh Arlovin dan Kusri menganalisis sentimen review pengguna aplikasi *Fizzo Novel* di *Google Play* menggunakan algoritma *Naive Bayes*. Dalam penelitian tersebut, *Naive Bayes* menghasilkan akurasi sebesar 83% dalam mengklasifikasikan sentimen positif dan negatif. *Naive Bayes* dipilih karena kesederhanaan dan efisiensinya dalam pemrosesan teks. Meskipun akurasi yang diperoleh cukup baik, penelitian ini masih berfokus pada satu algoritma saja tanpa membandingkan dengan metode lain. Peneliti menyarankan penggunaan algoritma lain seperti *Random Forest* atau *SVM* untuk mendapatkan performa yang lebih optimal. Selain itu, penelitian ini juga belum menerapkan analisis berbasis aspek sehingga tidak dapat memberikan insight detail tentang aspek-aspek spesifik aplikasi yang disukai atau tidak disukai pengguna [11].

Irawan melakukan studi perbandingan algoritma *Random Forest*, *Naive Bayes*, dan *SVM* dalam analisis sentimen pada aplikasi *CapCut*. Penelitian ini menyimpulkan bahwa *SVM* adalah algoritma dengan akurasi terbaik yaitu 86% untuk analisis sentimen ulasan *CapCut*, diikuti oleh *Random Forest* dengan akurasi 84% dan *Naive Bayes* dengan akurasi 83%. Meskipun *SVM* memberikan akurasi tertinggi, penelitian ini menemukan bahwa *SVM* membutuhkan waktu pemrosesan yang cukup lama dibandingkan dengan kedua algoritma lainnya, terutama ketika bekerja dengan dataset yang besar. *Random Forest* menunjukkan keseimbangan yang baik antara akurasi dan efisiensi komputasi. Keterbatasan penelitian ini adalah belum menerapkan analisis berbasis aspek sehingga tidak dapat mengidentifikasi sentimen per aspek spesifik [15].

2.1. Analisis Sentimen

Analisis sentimen adalah studi komputasi yang bertujuan mengevaluasi dan mengenali sentimen atau opini secara otomatis yang terkandung dalam teks [14], [15]. Metode ini berkembang pesat karena

kemampuannya menggali makna subjektif dari data teks secara otomatis dan terstruktur [2].

Analisis sentimen bekerja dengan mengklasifikasikan teks berdasarkan *polaritas* emosi yang terkandung di dalamnya, umumnya terdiri dari tiga kategori utama: positif, negatif, atau netral [17]. Sentimen positif mengindikasikan apresiasi, kepuasan, atau pandangan yang menguntungkan terhadap suatu topik. Sentimen negatif menunjukkan ketidakpuasan, kritik, atau pandangan yang tidak menguntungkan. Sementara sentimen netral menunjukkan opini yang tidak memihak atau pernyataan faktual tanpa emosi yang jelas.

Dalam konteks praktis, analisis sentimen memiliki berbagai aplikasi penting. Perusahaan menggunakan analisis sentimen untuk memahami opini publik terhadap produk atau layanan mereka melalui analisis ulasan pelanggan, komentar media sosial, atau *feedback* survei. Di bidang politik, analisis sentimen digunakan untuk mengukur respon publik terhadap kebijakan atau tokoh politik. Dalam industri hiburan, khususnya film, analisis sentimen membantu produser dan distributor memahami persepsi audiens terhadap karya mereka.

2.2. Aspect-Based Sentiment Analysis

ABSA merupakan jenis analisis sentimen yang mengekstraksi dan mengklasifikasikan sentimen terhadap berbagai aspek spesifik dari suatu entitas [10], [11]. ABSA dapat digunakan untuk memperoleh opini dan polaritas sentimen dalam sebuah ulasan [23]. ABSA memberikan informasi sentimen yang lebih *komprehensif* mengenai hal yang disukai atau tidak disukai pengguna pada masing-masing aspek [12].

Perbedaan fundamental antara analisis sentimen tradisional dan ABSA terletak pada level *granularitas* analisis. Analisis sentimen tradisional hanya memberikan satu label sentimen untuk keseluruhan dokumen atau kalimat. Misalnya, ulasan "Film ini memiliki visual yang menakjubkan tetapi ceritanya membosankan" akan diklasifikasikan sebagai sentimen netral atau ambigu oleh analisis sentimen tradisional. Sebaliknya, ABSA dapat mengidentifikasi dua aspek berbeda dalam ulasan tersebut: aspek "visual" dengan sentimen positif dan aspek "cerita" dengan sentimen negatif. Kemampuan ini membuat ABSA sangat penting dalam memahami opini pengguna secara detail dan *komprehensif* [21].

ABSA terdiri dari dua tugas utama: (1) *Aspect Extraction* - mengidentifikasi aspek-aspek yang dibahas dalam teks, dan (2) *Aspect Sentiment Classification* - menentukan polaritas sentimen untuk setiap aspek yang telah diidentifikasi [22]. Dalam penelitian ini, *aspect extraction* dilakukan menggunakan metode keyword matching berbasis kamus kata kunci yang telah didefinisikan untuk setiap aspek. Sementara *aspect sentiment classification* menggunakan pendekatan *lexicon-based* untuk pelabelan otomatis, dilanjutkan dengan klasifikasi menggunakan algoritma machine learning.

2.3. Algoritma Machine Learning

Decision Tree adalah model klasifikasi berbasis pohon keputusan yang membuat serangkaian aturan untuk mengklasifikasikan data [16]. *Decision Tree* menonjol karena beberapa keunggulan: (1) Interpretabilitas tinggi - struktur pohon mudah divisualisasikan dan dipahami oleh pengguna non-teknis; (2) Efisiensi komputasi - proses *training* dan prediksi relatif cepat; (3) Kemampuan memproses variabel numerik dan kategorikal tanpa perlu *preprocessing* ekstensif; (4) Tidak memerlukan normalisasi data; (5) Dapat menangani *missing values* dengan baik [24]. Namun, *Decision Tree* memiliki kecenderungan untuk *overfitting* pada data *training*, terutama ketika pohon tumbuh terlalu dalam. Untuk mengatasi hal ini, teknik *pruning* atau pembatasan kedalaman pohon sering diterapkan.

Random Forest merupakan model *ensemble* yang menggunakan banyak pohon keputusan untuk meningkatkan akurasi [1]. Algoritma ini diperkenalkan oleh Leo Breiman dan bekerja berdasarkan prinsip "*wisdom of the crowd*" - kombinasi dari banyak model yang lemah dapat menghasilkan model yang kuat. Proses kerja *Random Forest* melibatkan beberapa tahap: (1) *Bootstrap aggregating (bagging)* - membuat *multiple subset* dari data *training* dengan metode *sampling with replacement*; (2) *Random feature selection* - pada setiap *split node*, hanya *subset* acak dari fitur yang dipertimbangkan; (3) *Majority voting* - untuk klasifikasi, prediksi akhir ditentukan berdasarkan *vote* mayoritas dari seluruh pohon. Pendekatan ini mengurangi varians model dan meningkatkan generalisasi. Keunggulan *Random Forest* meliputi: (1) *Robust* terhadap *overfitting* karena *averaging*; (2) Dapat menangani *dataset* dengan fitur berdimensi tinggi; (3) Memberikan estimasi *feature importance*; (4) Performa baik pada berbagai jenis data tanpa *tuning parameter* ekstensif. Kelemahan utama adalah kompleksitas model yang lebih tinggi dan waktu *training* yang lebih lama dibanding *Decision Tree* tunggal, serta interpretabilitas yang lebih rendah [15].

Support Vector Machine (SVM) adalah teknik klasifikasi yang menemukan *hyperplane* optimal untuk memisahkan kelas data dengan margin terbesar [10]. SVM bekerja berdasarkan prinsip *structural risk minimization*, yang bertujuan meminimalkan *error* generalisasi bukan hanya *error training*. Konsep dasar SVM meliputi: (1) *Hyperplane* - bidang pemisah berdimensi $n-1$ yang memisahkan data dalam ruang berdimensi n ; (2) *Support vectors* - titik data yang terletak paling dekat dengan *hyperplane* dan menentukan posisi *hyperplane*; (3) *Margin* - jarak antara *hyperplane* dengan *support vector* terdekat dari masing-masing kelas; (4) *Kernel trick* - transformasi data ke ruang berdimensi lebih tinggi untuk menangani data yang tidak *linearly separable*. SVM memiliki beberapa keunggulan: (1) Efektif dalam ruang berdimensi tinggi; (2) Masih efektif ketika jumlah dimensi lebih besar dari jumlah sampel; (3) *Memory*

efficient karena hanya menggunakan *subset* dari *training points* (*support vectors*); (4) *Versatile* karena dapat menggunakan berbagai *kernel function* (*linear*, *polynomial*, *RBF*, *sigmoid*). Namun, SVM memiliki kelemahan: (1) Sensitif terhadap pemilihan parameter dan *kernel*; (2) Waktu *training* yang lama pada *dataset* besar; (3) Tidak memberikan estimasi probabilitas secara langsung; (4) Sulit diinterpretasi dibanding *Decision Tree* [19].

3. METODE PENELITIAN

Metodologi penelitian dilakukan secara terstruktur menggunakan pendekatan *data mining* dengan tahapan: (1) pengumpulan data, (2) pra-pemrosesan, (3) ABSA 1 - ekstraksi aspek, (4) ABSA 2 - pelabelan sentimen, (5) pembagian data dan TF-IDF, (6) perbandingan model, dan (7) interpretasi hasil.

3.1. Pengumpulan Data

Dataset dikumpulkan menggunakan teknik *web scraping* dari 5 channel YouTube yang membahas film JUMBO menggunakan YouTube API V3 melalui Google Cloud, menghasilkan 7.906 komentar [8].

3.2. Pra-Pemrosesan

Tahapan pra-pemrosesan meliputi penghapusan duplikasi, *cleaning* (menghapus URL, emoji, karakter non-alfabet), *case folding*, normalisasi, tokenisasi, dan *stopword removal* menggunakan NLTK [14].

3.3. ABSA 1 - Ekstraksi Aspek

Mengidentifikasi lima aspek penting (cerita/emosi, visual, dukungan/apresiasi, perbandingan, musik) menggunakan *keyword matching* untuk mengekstrak aspek dalam komentar [10].

3.4. ABSA 2 - Pelabelan Sentimen

Pelabelan sentimen menggunakan metode *lexicon-based* dengan InSet Lexicon untuk menentukan polaritas (positif/negatif) setiap aspek [12]. Komentar masih belum berlabel sentimen, sehingga dilakukan pelabelan sentimen menggunakan metode berbasis leksikon. Proses pelabelan data yang baik sangat penting karena akan menghasilkan model klasifikasi yang lebih akurat [25].

3.5. Pembagian Data dan TF-IDF

Data dibagi menjadi 80% *training* dan 20% *testing*, kemudian diubah ke representasi vektor menggunakan TF-IDF (*Term Frequency-Inverse Document Frequency*) [8].

3.6. Perbandingan Model

Melakukan perbandingan tiga model klasifikasi (*Decision Tree*, *Random Forest*, *SVM*) untuk menentukan model terbaik. Setiap model dievaluasi menggunakan *confusion matrix* dengan metrik *accuracy*, *precision*, *recall*, dan *F1-score*.

4. HASIL DAN PEMBAHASAN

4.1. Hasil Pengumpulan dan Preprocessing

Tabel 1. Penghapusan duplikasi

Keterangan	Jumlah
Data awal	7906
Komentar duplikat dan kosong	274
Data setelah dibersihkan	7632

Dari tabel 1 dapat dilihat bahwa dari 7.906 komentar awal, setelah penghapusan duplikasi dan data kosong tersisa 7.632 komentar yang siap diproses lebih lanjut.

4.2. Ekstraksi Aspek

Tabel 2. Distribusi Aspek

Aspek	Jumlah	Persentase
Dukungan/apresiasi	2050	34,4%
Cerita/emosi	1612	27,1%
Visual	1432	24,0%
Perbandingan	667	11,2%
Musik	195	3,3%
Total teridentifikasi	5956	78,0%

Dari tabel 2 terlihat bahwa sebanyak 5.956 komentar (78,0%) berhasil diidentifikasi aspeknya, dengan aspek Dukungan/Apresiasi mendominasi.

4.3. Pelabelan Sentimen

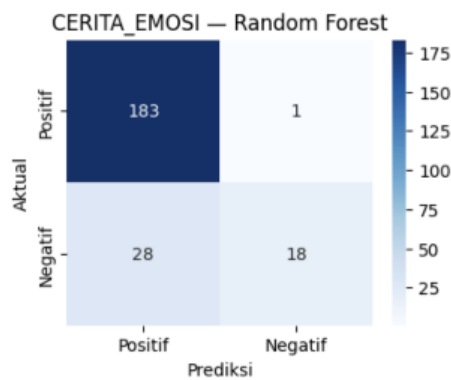
Tabel 3. Penghapusan duplikasi

Keterangan	Jumlah
Data dari ABSA 1	5956
Baris tanpa ekspresi sentimen	1874
Data akhir untuk dianalisis	4082

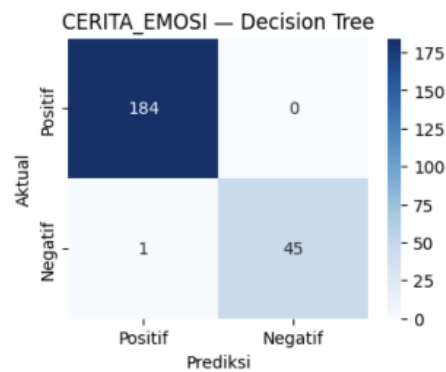
Dari tabel 3 dapat dilihat bahwa sebanyak 1.874 baris dihapus karena tidak memiliki ekspresi sentimen yang jelas (netral atau ambigu). Dengan demikian, dataset final yang digunakan untuk pelatihan dan pengujian model berjumlah 4.082 komentar dengan sentimen positif atau negatif yang jelas. Dataset ini kemudian dibagi menjadi data *training* (80%) dan data *testing* (20%) untuk setiap aspek secara terpisah. Pembagian ini dilakukan secara acak dengan mempertahankan proporsi sentimen di setiap aspek. Aspek Dukungan/Apresiasi memiliki jumlah data terbesar (1.797 komentar atau 44,0% dari total), sementara aspek Musik memiliki jumlah data terkecil (93 komentar atau 2,3% dari total).

4.4. Evaluasi Model dengan Confusion Matrix

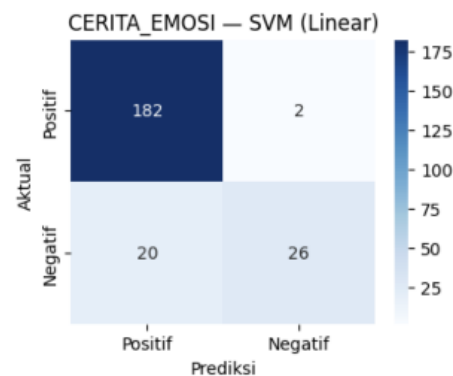
Setelah data dibagi, dilakukan pelatihan model menggunakan tiga algoritma: *Decision Tree*, *Random Forest*, dan *SVM* (Linear). Hasil evaluasi divisualisasikan menggunakan *confusion matrix* untuk setiap kombinasi model dan aspek.



Gambar 1. Confusion matrix random forest untuk aspek cerita/emosi



Gambar 3. Confusion matrix decision tree untuk aspek cerita/emosi



Gambar 2. Confusion matrix SVM untuk aspek cerita/emosi

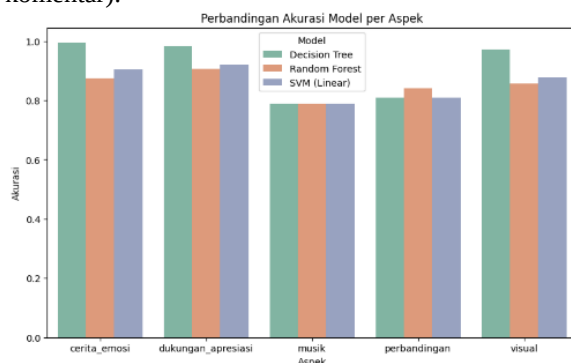
Dari gambar 1, 2, dan 3 dapat dilihat bahwa ketiga model dapat melakukan analisis sentimen dengan baik untuk aspek Cerita/Emosi, dengan *Decision Tree* menunjukkan performa terbaik.

4.5. Hasil Evaluasi Model per Aspek

Tabel 4. Tabel Model Per Aspek

Aspek	Model terbaik	Accuracy	Precision	Recall	F1-Score
Cerita/emosi	Decision Tree	0.995	0.994	1.000	0.997
Dukungan/apresiasi	Decision Tree	0.983	0.984	0.996	0.990
Visual	Decision Tree	0.972	0.990	0.973	0.981
Perbandingan	Random Forest	0.841	0.833	0.978	0.900
Musik	Decision Tree	0.789	0.833	0.833	0.833

Dari tabel 4 terlihat bahwa *Decision Tree* secara konsisten memberikan performa terbaik untuk sebagian besar aspek dengan akurasi rata-rata 91,0%. Akurasi tertinggi dicapai pada aspek cerita/emosi (99,5%) dan terendah pada aspek musik (78,9%), yang kemungkinan disebabkan keterbatasan data (hanya 93 komentar).



Gambar 4. Grafik perbandingan akurasi model untuk semua aspek

Dari gambar 4 dapat dilihat perbandingan akurasi ketiga model untuk setiap aspek secara visual. Grafik menunjukkan bahwa *Decision Tree* mendominasi performa terbaik di 4 dari 5 aspek dengan akurasi di atas 90% untuk aspek Cerita/Emosi, Dukungan/Apresiasi, dan Visual. *Random Forest* unggul pada aspek Perbandingan, kemungkinan karena kemampuan ensemble method dalam menangani variasi pola sentimen yang lebih kompleks. Terdapat kesenjangan performa yang signifikan antara aspek dengan data banyak (Cerita/Emosi: 1.149 komentar, Dukungan/Apresiasi: 1.797 komentar, Visual: 730 komentar) yang mencapai akurasi di atas 97% dengan *Decision Tree*, dan aspek dengan data terbatas (Musik: 93 komentar) yang hanya mencapai akurasi 78,9%. Hal ini menunjukkan adanya korelasi positif antara jumlah data dengan akurasi model, mengindikasikan pentingnya keseimbangan distribusi data dalam analisis sentimen.

4.6. Distribusi Sentimen

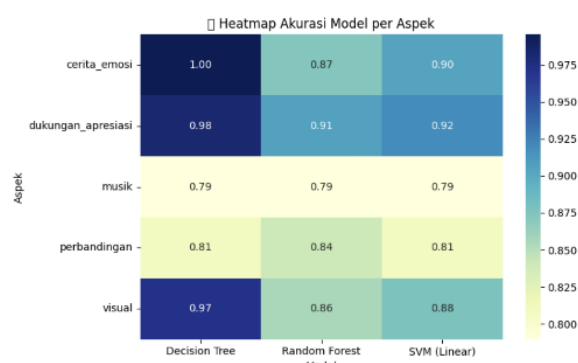
Tabel 5. Distribusi Aspek

Aspek	Positif	Negatif
Dukungan/apresiasi	C	11,3%
Cerita/emosi	80,0%	20,0%
Visual	75,9%	24,1%
Perbandingan	72,5%	27,5%
Musik	61,3%	38,7%

Dari tabel 5 dapat dilihat pola sentimen dominan untuk setiap aspek. Aspek Dukungan/Apresiasi memiliki persentase sentimen positif tertinggi 88,7%, menunjukkan bahwa penonton sangat mendukung dan mengapresiasi film JUMBO secara keseluruhan. Hanya 11,3% komentar yang bersifat negatif pada aspek ini. Aspek Cerita/Emosi mendapat sentimen positif yang cukup tinggi 80,0%, mengindikasikan bahwa alur cerita dan muatan emosional film berhasil diterima dengan baik oleh mayoritas penonton. Aspek Visual memperoleh 75,9% sentimen positif, menunjukkan apresiasi terhadap kualitas sinematografi dan aspek visual film. Aspek Perbandingan memiliki 72,5% sentimen positif, namun dengan 27,5% sentimen negatif, ini menjadi salah satu aspek dengan kritik yang cukup signifikan, kemungkinan terkait dengan ekspektasi penonton yang membandingkan dengan film lain. Aspek Musik memiliki persentase sentimen positif terendah 61,3% dengan sentimen negatif tertinggi 38,7%, mengindikasikan bahwa musik/soundtrack film menjadi aspek yang paling banyak mendapat kritik dari penonton.

Secara keseluruhan, film JUMBO mendapat sambutan positif dari penonton dengan rasio positif-negatif sebesar 4,6:1 (82,1% positif berbanding 17,9% negatif). Aspek Dukungan/Apresiasi, Cerita/Emosi, dan Visual menjadi kekuatan utama film, sementara aspek Musik menjadi area yang perlu mendapat perhatian lebih untuk produksi film berikutnya.

4.7. Visualisasi Performa Model



Gambar 5. Heatmap akurasi model per aspek

Dari gambar 5 dapat dilihat akurasi ketiga model (Decision Tree, Random Forest, dan SVM) untuk kelima aspek yang dianalisis. Intensitas warna yang

lebih gelap menunjukkan akurasi yang lebih tinggi. Decision Tree memiliki akurasi tertinggi untuk aspek Cerita/Emosi, Dukungan/Apresiasi, Visual, dan Musik, sedangkan Random Forest menunjukkan performa terbaik untuk aspek Perbandingan. Semua model menunjukkan performa yang konsisten rendah pada aspek Musik, mengindikasikan keterbatasan data atau kompleksitas sentimen pada aspek tersebut.

4.8. Pembahasan

Penelitian ini berhasil mengidentifikasi bahwa Decision Tree merupakan model paling efektif untuk klasifikasi sentimen berbasis aspek pada komentar film Jumbo. Keunggulan Decision Tree terletak pada kemampuannya membuat aturan klasifikasi yang jelas dan mudah diinterpretasi. Random Forest unggul pada aspek Perbandingan yang memiliki pola sentimen lebih kompleks.

Aspek Dukungan/Apresiasi yang mendominasi (44,0% dari total komentar) menunjukkan bahwa penonton lebih fokus memberikan apresiasi keseluruhan dibanding kritik detail. Sentimen positif yang tinggi (82,1%) mengindikasikan kesuksesan film dalam menarik perhatian publik, meskipun aspek Musik menjadi area yang perlu perbaikan. Keterbatasan data pada aspek tertentu mempengaruhi performa model, menunjukkan pentingnya keseimbangan distribusi data dalam analisis sentimen.

5. KESIMPULAN DAN SARAN

Memuat kesimpulan yang diperoleh dan saran-saran untuk penelitian selanjutnya (jika ada). Tuliskan Kesimpulan dan saran dalam 1 paragraf.

Penelitian ini berhasil mengimplementasikan *Aspect-Based Sentiment Analysis* pada 7.906 komentar YouTube film Jumbo dengan mengidentifikasi lima aspek utama. *Decision Tree* terbukti sebagai model terbaik dengan akurasi rata-rata 91,0%, mencapai 99,6% pada aspek cerita/emosi. Aspek dukungan/apresiasi mendominasi dengan 44,0% dari total komentar dan 88,7% sentimen positif, mengindikasikan respons publik yang sangat baik terhadap film. Penelitian ini memberikan wawasan objektif tentang persepsi audiens dan membuktikan efektivitas ABSA untuk analisis ulasan film Indonesia. Penelitian selanjutnya dapat memanfaatkan algoritma *Deep Learning* (BERT, IndoBERT), memperluas sumber data ke platform lain (Twitter/X, Instagram), menambah aspek analisis (dialog, karakter, *pacing*), dan menerapkan teknik *oversampling* untuk mengatasi keterbatasan data pada aspek tertentu.

DAFTAR PUSTAKA

- [1] J. Emarapenta, B. Sinulingga, H. Cesar, and K. Sitorus, "Analisis Sentimen Masyarakat terhadap Film Horor Indonesia Menggunakan Metode SVM dan TF-IDF Sentiment Analysis of Public towards Indonesian Horror Films Using SVM and TF-IDF Methods," vol. 14, no. April, pp. 42–53, 2024.

- [2] A. Ardiansyah, C. Agustina, I. Maryani, and D. Pribadi, "Analisis Sentimen pada Komentar YouTube terkait Pembahasan eSIM Menggunakan Metode Naive Bayes dan Random Forest," vol. 11, no. 1, pp. 7–14, 2025.
- [3] S. M. Harahap *et al.*, "Analisis Sentimen Komentar Youtube terhadap Food Vlogger dengan Menggunakan Metode Naive Bayes Sulistia Maharani Harahap," vol. 9, no. 1, pp. 87–96, 2024.
- [4] N. V. Loca, D. Abdullah, T. Informatika, and U. M. Bengkulu, "Analisis Sentimen terhadap Dedy Mulyadi Berdasarkan Komentar Youtube Menggunakan Metode Naive Bayes," pp. 874–884, 2025.
- [5] T. Muhyat, A. Fauzi, and J. Indra, "Analisis Sentimen Terhadap Komentar Video Youtube Menggunakan Support Vector Machines," pp. 231–240, 2022.
- [6] N. A. Laia and S. P. Barus, "E ISSN : 2809-4069 Analisis Sentimen YouTube : " Di Balik Ambisi Jokowi dalam IKN ", vol. 5, no. 1, pp. 7–12, 2025.
- [7] W. Irmayani, "PERSEPSI PUBLIK TERHADAP KENAIKAN PPN 12%: PENDEKATAN SENTIMEN PADA KOMENTAR YOUTUBE," vol. 12, no. 2, pp. 112–118, 2024.
- [8] A. Hermawan and I. Jowensen, "Implementasi Text-Mining untuk Analisis Sentimen pada Twitter dengan Algoritma Support Vector Machine," vol. 12, no. 1, pp. 129–137, 2023.
- [9] R. N. Mauliza and Y. R. Sipayung, "Penerapan Text Mining Dalam Menganalisis Pendapat Masyarakat Terhadap Pemilu 2024 Pada Media Sosial X Menggunakan Metode Naive Bayes," vol. 9, no. 1, pp. 1–16, 2024.
- [10] M. R. Herdiansyah, A. Yuliana, T. Informatika, N. Bayes, and K. Sentimen, "NAIVE BAYES BERDASARKAN KOMENTAR PADA YOUTUBE," vol. 8, no. 6, pp. 12454–12459, 2024.
- [11] K. Teguh Arlovin, Kusri, "ANALISIS SENTIMEN REVIEW PENGGUNA APLIKASI FIZZO NOVEL DI GOOGLE PLAY MENGGUNAKAN ALGORITMA NAIVE BAYES," vol. 6, no. 1, pp. 65–70, 2024.
- [12] A. N. Syafia, M. F. Hidayattullah, and W. Suteddy, "Studi Komparasi Algoritma SVM Dan Random Forest Pada Analisis Sentimen Komentar Youtube BTS," vol. 8, no. 3, pp. 207–212, 2023.
- [13] A. P. Hendrawan *et al.*, "Design of an Exam Cheating Detection System Application Based on Machine Learning with the Computer Vision Method Abstrak PENDAHULUAN institusi pendidikan di seluruh dunia . Kecurangan ini tidak hanya merugikan para peserta untuk masalah ini . Computer vision memungkinkan komputer untuk menganalisis video dan pengawasan , tetapi juga meningkatkan akurasi dalam mendeteksi perilaku mencurigakan .," vol. 11, no. 2, pp. 509–521, 2025.
- [14] N. E. Oktaviana *et al.*, "ANALISIS SENTIMEN TERHADAP KEBIJAKAN KULIAH DARING SELAMA PANDEMI MENGGUNAKAN PENDEKATAN LEXICON BASED FEATURES DAN SENTIMENT ANALYSIS AGAINST ONLINE LECTURE POLICY DURING THE PANDEMIC USING LEXICON BASED FEATURES APPROACH AND SUPPORT," vol. 9, no. 2, 2022, doi: 10.25126/jtiik.202295625.
- [15] I. Irawan, M. H. Wathan, and M. B. Prayogi, "Studi Perbandingan : Algoritma Random Forest , Naive Bayes Dan Support Vector Machine Dalam Analisis Sentimen Pada Aplikasi Capcut," vol. 5, no. 4, pp. 189–201, 2024.
- [16] A. Karimah *et al.*, "ANALISIS SENTIMEN KOMENTAR VIDEO MOBIL LISTRIK DI PLATFORM," vol. 8, no. 1, pp. 767–773, 2024.
- [17] A. D. Hananto, A. M. Erfiana, B. Lexiani, and P. Putri, "Algoritma Machine Learning Naive Bayes pada Analisis Sentimen Kesepakatan Polri dan GNPF-MUI pada Aksi Bela Islam III ' 212 ' Naive Bayes Machine Learning Algorithm on Sentiment Analysis of Police and GNPF-MUI Agreement on ' 212 ' Islamic Defense Action III," vol. 4, no. 2, pp. 151–160, 2023.
- [18] A. R. Putra and D. E. Ratnawati, "ANALISIS SENTIMEN BERBASIS ASPEK PADA APLIKASI MOBILE MENGGUNAKAN NAIVE BAYES BERDASARKAN ULASAN PENGGUNA PLAYSTORE (STUDI KASUS : JCONNECT MOBILE) ASPECT-BASED SENTIMENT ANALYSIS ON MOBILE APPLICATIONS USING NAIVE BAYES BASED ON PLAYSTORE USER REVIEWS (CASE STUDY : JCONNECT MOBILE)," vol. 12, no. 2, pp. 293–300, 2025, doi: 10.25126/jtiik.2025127556.
- [19] I. P. Ariyani, K. D. Tania, A. Wedhasmara, and A. Meiriza, "Ekstraksi Pengetahuan dari Ulasan Aplikasi CapCut Menggunakan Metode Aspect-Based Sentiment Analysis dan Klasifikasi," vol. 16, no. April, pp. 80–92, 2025.
- [20] I. A. Atika Putri, Surya Agustian, Jasril, "Eksplorasi fitur fasttext, tf-idf dan indobert pada metode k-nearest neighbor untuk klasifikasi sentimen 1,2,3,4," vol. 7, no. 1, pp. 49–60, 2025.
- [21] A. Jazuli, A. A. Chamid, and R. Kusumaningrum, "Transformer-based semantic indexing for aspect-based sentiment analysis using an enhanced index generation algorithm with BERT," vol. 12, no. 127, 2025.
- [22] V. H. Pranatawijaya, N. Noor, and K. Sari, "Unveiling User Sentiment : Aspect-Based Analysis and Topic Modeling of Ride-Hailing and Google Play App Reviews," vol. 10, no. 3, pp. 328–339, 2024.

- [23] A. A. Chamid, “Graph-Based Semi-Supervised Deep Learning for Indonesian Aspect-Based Sentiment Analysis,” 2023.
- [24] I. H. Ariyanto, A. A. Chamid, and R. Fiati, “Demand Prediction and Apparel Production Management Using AI-Based Decision Tree,” vol. 8, no. 1, 2025, doi: 10.32877/bt.v8i1.2325.
- [25] A. A. Chamid and R. Kusumaningrum, “Ingénierie des Systèmes d ’ Information Labeling Consistency Test of Multi-Label Data for Aspect and Sentiment Classification Using the Cohen Kappa Method,” vol. 29, no. 1, pp. 161–167, 2024.