

ANALISIS SENTIMEN TERHADAP ACARA RUANGGURU CLASH OF CHAMPIONS (COC) DENGAN ALGORITMA NAÏVE BAYES CLASSIFIER

Muhammad Rizki Arya Rifai^{1*}, Betha Nurina Sari², Garno³

^{1,2,3}Univeristas Singaperbangsa Karawang, Jl. H.S. Ronggowaluyo, Puseurjaya, Telukjambe Timur, Karawang, Jawa Barat 41361; Telp: (0267) 641177; Fax: (0267) 641367

Keywords:

Analisis Sentimen;
Youtube;
Ruangguru *Clash of Champions*;
Naïve Bayes Classifier;

Correspondent Email:

mrizkiaryarifai@gmail.com

Abstrak. Pertumbuhan media sosial seperti YouTube menjadikannya sumber data penting untuk melihat opini publik terhadap suatu konten, termasuk program edukatif. Penelitian ini bertujuan untuk menganalisis sentimen publik terhadap acara Ruangguru Clash of Champions (COC) melalui komentar di YouTube dengan menggunakan algoritma Naïve Bayes Classifier. Proses analisis mengikuti tahapan Knowledge Discovery in Database (KDD) yang mencakup: pemilihan data, pra-pemrosesan, transformasi, data mining, dan interpretasi. Komentar yang digunakan berasal dari episode pertama acara COC, diambil dalam rentang waktu 29 Juni hingga 31 Desember 2024. Data diberi label sentimen positif, negatif, dan netral, menggunakan metode leksikon. Transformasi data dilakukan dengan metode TF-IDF yang sudah dituning untuk menghasilkan fitur yang relevan. Selanjutnya, klasifikasi sentimen dilakukan menggunakan algoritma Multinomial Naïve Bayes dan divalidasi menggunakan teknik 10-fold cross-validation. Hasil evaluasi melalui Confusion Matrix menunjukkan rata-rata dari 10 fold dengan nilai akurasi sebesar 70,55%, presisi 74,75%, recall 62,36%, dan f1-score 64,38%. Penelitian ini memberikan gambaran persepsi masyarakat terhadap acara COC dan dapat menjadi bahan evaluasi untuk meningkatkan kualitas program edukatif ke depan.

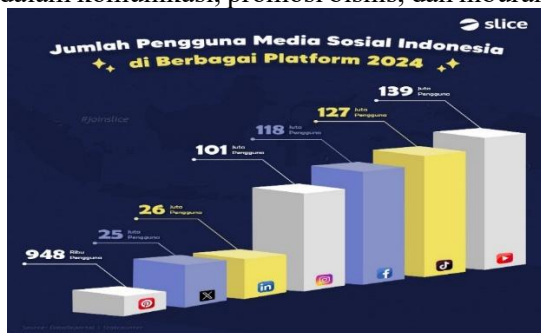


JITET is licensed under a Creative Commons Attribution-NonCommercial 4.0 International License.

Abstract. The growth of social media platforms like YouTube has made them important sources of data for understanding public opinion on various content, including educational programs. This study aims to analyze public sentiment toward the Ruangguru Clash of Champions (COC) event by examining comments on YouTube using the Naïve Bayes Classifier algorithm. The analysis follows the Knowledge Discovery in Database (KDD) process, which includes data selection, preprocessing, transformation, data mining, and interpretation. The comments were collected from the first episode of the COC event, within the period of June 29 to December 31, 2024. Sentiment labels positive, negative, and neutral, were assigned using the lexicon-based method. The data was transformed using a tuned TF-IDF method to generate more informative features. Sentiment classification was then performed using the Multinomial Naïve Bayes algorithm and validated with 10-fold cross-validation. Evaluation results using a Confusion Matrix showed average values across the 10 folds: accuracy of 70.55%, precision of 74.75%, recall of 62.36%, and F1-score of 64.38%. This study provides insights into public perception of the COC event and can serve as valuable input for improving the quality of future educational programs.

1. PENDAHULUAN

Meningkatnya akses internet dan penggunaan perangkat mobile mendorong pertumbuhan media sosial di Indonesia. Saat ini, lebih dari 215 juta warga terhubung ke internet, dan sebagian besar aktif di platform seperti *YouTube*, *TikTok*, *Instagram*, *Facebook*, dan lainnya. Pada Gambar 1.1 menunjukkan bahwa *YouTube* menjadi media sosial paling populer dengan 139 juta pengguna. Perluasan jaringan *4G* dan *5G* turut mempercepat adopsi digital, menjadikan media sosial sebagai bagian penting dalam kehidupan sehari-hari, termasuk dalam komunikasi, promosi bisnis, dan hiburan.



Gambar 1.1 Data Pengguna Aktif Media Sosial di Indonesia
(Sumber: blog.slice.id, 2024)

Seiring perkembangan digital, *YouTube* menjadi sumber penting untuk analisis sentimen. Selain sebagai platform berbagi video, Penelitian sebelumnya menunjukkan bahwa komentar *YouTube* bisa mencerminkan persepsi publik terhadap program Pendidikan [1]. Salah satu contohnya adalah acara Ruangguru *Clash of Champions* (COC), sebuah kompetisi edukatif dari Ruangguru yang bertujuan meningkatkan minat belajar pelajar dan mahasiswa. Acara ini tidak hanya menunjukkan kemampuan peserta, tetapi juga memberi kesempatan bagi penyelenggara untuk menerima masukan dari audiens.

Karena itu, penting untuk menganalisis sentimen masyarakat terhadap acara ini agar penyelenggara memahami respons audiens dan menilai efektivitasnya. Menurut [2], analisis sentimen juga bermanfaat untuk mengevaluasi program atau kebijakan pendidikan berdasarkan opini publik.

Analisis sentimen tak hanya mengukur opini publik, tetapi juga mengungkap faktor-faktor yang memengaruhi persepsi mereka, seperti isi acara, cara penyampaian, dan interaksi peserta. Memahami hal ini membantu

penyelenggara meningkatkan kualitas acara agar lebih menarik. Menurut [3], memahami perasaan audiens penting untuk merancang program yang sesuai dengan harapan masyarakat.

Penelitian oleh [4] menganalisis sentimen terhadap kebijakan Kampus Merdeka menggunakan algoritma *Naïve Bayes* pada komentar *YouTube* kanal Mendikbud. Proses analisis menggunakan TF-IDF untuk pembobotan kata dan validasi data dengan *K-Fold Cross Validation* ($k=10$). Hasil terbaik dicapai pada iterasi ke-4 dengan akurasi 97%, dan rata-rata akurasi keseluruhan sebesar 91,8%. Nilai *precision*, *recall*, dan *f-measure* masing-masing adalah 90,35%, 93,6%, dan 91,25%.

Penelitian oleh [5] melakukan analisis sentimen terhadap kenaikan biaya haji 2023 menggunakan data komentar *YouTube* dari salah satu *channel* berita nasional. Mereka membandingkan tiga algoritma klasifikasi: *Naïve Bayes Classifier*, *Decision Tree*, dan *Random Forest*. Hasilnya, *Naïve Bayes* mencatat akurasi tertinggi sebesar 90%, disusul *Random Forest* 87%, dan *Decision Tree* 83%.

Penelitian ini bertujuan menganalisis sentimen opini publik di *YouTube* terhadap acara Ruangguru *Clash of Champions* (COC) dengan menggunakan algoritma *Naïve Bayes Classifier* untuk mengklasifikasikan sentimen menjadi positif, negatif, atau netral. Data diperoleh dari komentar *YouTube* menggunakan teknik *crawling data*, lalu diproses melalui tahapan *preprocessing*, seperti penghapusan karakter khusus, normalisasi teks, *stopword removal*, dan *stemming*.

Penelitian ini diharapkan dapat memberikan wawasan tentang sentimen publik terhadap acara Ruangguru *Clash of Champions* (COC), apakah bersifat positif, negatif, atau netral. Hasil analisis ini bisa menjadi acuan bagi tim Ruangguru, peserta, institusi pendidikan, maupun peneliti untuk meningkatkan kualitas acara, menyusun strategi pengembangan layanan pendidikan digital, serta memahami kebutuhan dan preferensi masyarakat terhadap program pendidikan.

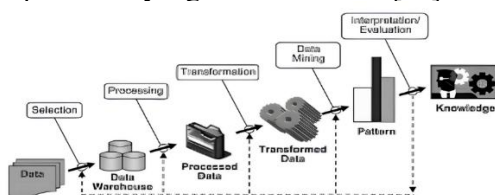
2. TINJAUAN PUSTAKA

2.1. Analisis Sentimen

Analisis sentimen adalah proses mengelompokkan opini menjadi positif, negatif, atau netral dengan bantuan teknik *Natural Language Processing* (NLP) untuk memahami makna dan konteks teks [6]. Menurut [7], metode ini berguna untuk melihat tanggapan masyarakat terhadap isu tertentu di platform seperti YouTube. Analisis ini penting dalam berbagai bidang, termasuk pendidikan, sosial, dan bisnis, karena membantu perusahaan memahami persepsi pengguna serta menemukan area yang perlu diperbaiki [8].

2.2. Knowledge Discovery In Database (KDD)

Knowledge Discovery in Database (KDD) adalah proses untuk menemukan informasi dan wawasan bernilai dari data, yang dapat digunakan sebagai dasar pengambilan keputusan [9]. Dalam *text mining*, KDD membantu mengenali pola dari data berbasis teks. Menurut [10], proses KDD mencakup pemilihan data (*data selection*), pembersihan (*cleansing*), transformasi (*transformation*), penggalian data (*data mining*), dan evaluasi (*interpretation*). Dalam penelitian ini, KDD digunakan untuk menggali pola sentimen pengguna media sosial, dengan tujuan mengubah data mentah menjadi pengetahuan yang relevan dan bermanfaat bagi berbagai bidang seperti pendidikan, bisnis, dan penelitian. Pada gambar 2.1 menunjukkan tahapan KDD yang bersumber dari [11].



Gambar 2.1 Tahapan KDD
(Sumber: [11])

2.2.1. Data Selection

Tahap awal KDD di mana data dipilih dari berbagai sumber sesuai dengan tujuan penelitian. Misalnya, dalam analisis sentimen, data dapat dikumpulkan dari platform seperti *YouTube*, *Twitter*, atau *Google Play Store* menggunakan teknik *crawling*.

2.2.2. Preprocessing

Tahapan untuk membersihkan dan menyiapkan data sebelum dianalisis. Proses *preprocessing* terdiri dari *cleansing*, *case folding*, *tokenizing*, *spelling normalization*, *filtering*, dan *stemming*. Ini mencakup penghapusan duplikat, perbaikan kesalahan, serta penanganan data yang hilang atau tidak konsisten agar data siap digunakan pada proses berikutnya.

2.2.3. Transformation

Data diubah ke dalam format yang sesuai untuk analisis, seperti mengubah data teks ke angka atau melakukan normalisasi. Transformasi bertujuan menyesuaikan data agar kompatibel dengan metode data mining yang akan digunakan.

2.2.4. Data Mining

Inti dari proses KDD, yaitu menggali pola atau informasi penting dari data menggunakan algoritma tertentu, seperti klasifikasi, *clustering*, atau asosiasi.

2.2.5. Interpretation

Tahap evaluasi hasil untuk menilai performa analisis. Jika hasilnya belum optimal, peneliti dapat mengulang tahapan sebelumnya untuk meningkatkan akurasi dan kualitas temuan.

2.3. Algoritma Naïve Bayes Classifier

Menurut [12], Algoritma *Naïve Bayes* adalah metode klasifikasi berbasis *Teorema Bayes*, yang menghitung probabilitas suatu peristiwa berdasarkan bukti baru. Algoritma ini mengasumsikan bahwa setiap fitur dalam data bersifat independen satu sama lain, meskipun dalam kenyataannya asumsi ini tidak selalu sepenuhnya akurat. Algoritma *Naïve Bayes* didasarkan pada *Teorema Bayes*, yang memiliki rumus sebagai berikut [13]:

$$P(H|X) = \frac{P(X|H) \cdot P(H)}{P(X)} \quad (2.1)$$

Keterangan:

- X: Bukti atau data yang belum diketahui labelnya
- H: Hipotesis yang menyatakan data X masuk dalam kategori khusus
- P(X): Peluang data sampel X
- P(H): Probabilitas awal dari hipotesis H (probabilitas prior)

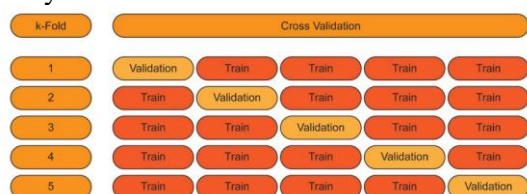
- $P(H|X)$: Untuk bukti (X), peluang hipotesis (H) benar. (probabilitas posterior)
- $P(X|H)$: Potensi data sampel (X) jika hipotesis (H) benar

2.4. Term Frequency-Inverse Document Frequency (TF-IDF)

Menurut [14], TF-IDF adalah metode untuk mengubah data teks menjadi data numerik dengan memberi bobot pada setiap kata. Term Frequency (TF) menunjukkan seberapa sering kata muncul dalam sebuah dokumen, sedangkan Inverse Document Frequency (IDF) menilai seberapa penting kata tersebut dengan melihat seberapa banyak dokumen yang mengandung kata tersebut. TF-IDF membantu menonjolkan kata-kata yang relevan dalam analisis teks.

2.5. K-Fold Cross Validation

K-Fold Cross Validation adalah metode evaluasi algoritma dengan membagi data secara acak menjadi K kelompok. Setiap kelompok secara bergantian digunakan sebagai data uji, sementara sisanya digunakan untuk pelatihan, sehingga hasil evaluasi lebih akurat dan merata [15]. Menurut [16], *Cross Validation* adalah teknik dalam data mining untuk meningkatkan akurasi model dengan membagi data menjadi data latih dan data uji. Proses ini dilakukan berulang sebanyak K kali, sehingga setiap bagian data digunakan bergantian sebagai data uji. Metode ini membantu menghasilkan evaluasi model yang lebih optimal dan menyeluruh.



Gambar 2.2 Ilustrasi Proses 5-Fold Cross Validation
(Sumber; [16])

2.6. Confusion Matrix

Confusion matrix adalah alat evaluasi untuk sistem klasifikasi, yang menyajikan tabel perbandingan antara hasil prediksi model dan kelas sebenarnya dari data. Matriks ini membantu menilai seberapa akurat model dalam mengklasifikasikan data [17]. Rumus *confusion matrix* ditunjukkan pada table 2.1.

Tabel 2.1 Struktur Tabel *Confusion Matrix*
(Sumber: [18])

Kelas	Terklarifikasi Positif	Terklarifikasi Negatif
Positif	True Positif (TP)	False Negatif (FN)
Negatif	False Positif (FP)	True Negatif (TN)

Keterangan:

- True Positif (TP): Data yang benar dan diklasifikasikan sebagai benar oleh sistem.
- False Positif (FP): Data yang salah tetapi diklasifikasikan sebagai benar oleh sistem.
- False Negatif (FN): Data yang benar tetapi diklasifikasikan sebagai salah oleh sistem.
- True Negatif (TN): Data yang salah dan diklasifikasikan sebagai salah oleh sistem.

Berikut merupakan rumus untuk menghitung *accuracy*, *precision*, *recall*, *f-Measure*.

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN} \quad (2.2)$$

$$Precision = \frac{TP}{TP + FP} \times 100\% \quad (2.3)$$

$$Recall = \frac{TP}{TP + FN} \quad (2.4)$$

$$F-Measure = 2 \times \frac{Precision \times Recall}{Precision + Recall} \quad (2.5)$$

2.7. Metode Lexicon

Metode Lexicon adalah pendekatan berbasis kamus kosakata yang memberi informasi kategori atau bobot sentimen pada setiap kata. Menurut [19], metode ini digunakan untuk mengukur sentimen dalam dataset dengan bantuan pustaka seperti VADER Sentiment, yang menghitung skor polaritas tiap kata. Skor tersebut digunakan untuk mengklasifikasikan sentimen:

- Positif jika skor > 0
- Negatif jika skor < 0
- Netral jika skor = 0

3. METODE PENELITIAN

Penelitian ini menggunakan metode *Knowledge Discovery in Database (KDD)* untuk melakukan analisis sentimen terhadap acara Ruangguru *Clash of Champions (COC)* dengan menganalisis kolom komentar pada

video episode pertama dari tanggal 29 Juni sampai 31 Desember 2024 yang diunggah di kanal YouTube Ruangguru. Data dikumpulkan melalui teknik *web crawling* pada platform YouTube.

Metodologi KDD dipilih karena mendukung proses analisis yang terstruktur, yang terdiri dari lima tahapan utama: pengambilan data (*data selection*), pembersihan data (*preprocessing*), transformasi data (*transformation*), klasifikasi data (*data mining*) menggunakan algoritma *Naïve Bayes Classifier*, dan evaluasi hasil (*interpretation*).

3.1. Data Selection

Penelitian ini menggunakan teknik *crawling* di platform YouTube dengan bantuan Python di Google Colaboratory untuk mengumpulkan komentar pada episode pertama Ruangguru *Clash of Champions* (COC). Data dikumpulkan dalam rentang waktu 29 Juni 2024 hingga 31 Desember 2024.

3.2. Preprocessing

Setelah data dikumpulkan, Langkah selanjutnya adalah *pre-processing teks* untuk menyiapkan data agar siap diolah. Tujuan utamanya adalah mengubah data yang tidak terstruktur menjadi format yang lebih terstruktur sesuai kebutuhan penelitian. Proses ini mencakup enam tahap, yaitu:

3.2.1. Cleansing

Pada tahap ini, data dibersihkan menggunakan. Proses ini mencakup penghapusan *emoticon*, *hashtag*, *URL*, *username*, dan simbol-simbol tidak relevan untuk mempermudah proses data mining. Contohnya, komentar seperti “Membuat malas belajar!!” akan diubah menjadi “Membuat malas belajar”.

3.2.2. Case Folding

Pada tahap ini dilakukan *case folding*, yaitu mengubah semua huruf kapital menjadi huruf kecil. Tujuannya agar teks seragam dan lebih mudah diolah dalam proses data mining. Misalnya, komentar “Sangat Membantu” diubah menjadi “sangat membantu”.

3.2.3. Tokenizing

Tahap selanjutnya adalah *tokenizing*, yaitu proses memecah kalimat menjadi kata per kata (*tokens*). Tujuannya untuk mempermudah proses stemming pada

tahap berikutnya. Contoh: kalimat “sangat seru acara ini” akan diubah menjadi: “sangat”, “seru”, “acara”, “ini”.

3.2.4. Spelling Normalization

Tahap ini adalah normalisasi, yaitu menyamakan ejaan kata agar seragam. Hal ini penting karena komentar di media sosial seperti YouTube sering menggunakan bahasa gaul atau singkatan. Contohnya: kata seperti “gk”, “ga”, “gak”, dan “ngk” akan diubah menjadi “tidak”.

3.2.5. Filtering

Tahap ini disebut *filtering*, yaitu proses menyaring atau membuang kata-kata yang jarang muncul. Tujuannya agar data yang dianalisis lebih relevan dan optimal saat masuk ke tahap data mining.

3.2.6. Stemming

Tahap ini adalah *stemming*, yaitu proses menghilangkan imbuhan pada setiap kata. Setelah tokenizing, kata-kata akan distem agar menjadi bentuk dasar, sehingga lebih mudah diproses saat data mining. Contoh: “berjalan” → “jalan”, “keren” → “keren”, “aplikasi” → “aplikasi”, “lancar” → “lancar”

3.3. Transformation

Setelah *pre-processing* selesai, langkah selanjutnya adalah transformasi data, yaitu mengubah teks menjadi bentuk numerik. Proses ini menggunakan metode TF-IDF untuk memberikan bobot pada setiap kata berdasarkan seberapa sering kata tersebut muncul dalam satu dokumen dibandingkan dengan dokumen lainnya. Skor TF-IDF menunjukkan seberapa penting sebuah kata dalam konteks analisis.

3.4. Data Mining

Langkah berikutnya adalah data mining, dengan menggunakan teknik *K-Fold Cross-Validation*. Data dibagi menjadi 10 bagian (*fold*) berukuran sama. Setiap kali, satu *fold* digunakan sebagai data uji, dan sisanya sebagai data latih. Proses ini diulang sebanyak 10 kali, lalu hasilnya dirata-rata untuk memperoleh estimasi performa model yang lebih akurat dan mengurangi bias. Algoritma yang digunakan dalam tahap ini adalah *Naïve Bayes Classifier*.

3.5. Interpretation

Tahap Interpretation atau evaluasi data dilakukan untuk menilai kinerja model. Pada tahap ini digunakan *Confusion Matrix* untuk membandingkan hasil prediksi dengan kondisi sebenarnya. Evaluasi menghasilkan metrik seperti *accuracy*, *precision*, *recall*, dan *f-measure*, yang digunakan untuk menilai seberapa baik dan akurat model dalam melakukan klasifikasi.

1. *Accuracy*: Mengukur seberapa dekat prediksi model dengan nilai sebenarnya secara keseluruhan. Semakin tinggi akurasi, semakin banyak prediksi yang benar.
2. *Precision*: Menilai seberapa tepat model dalam memprediksi data positif. *Precision* penting saat *False Positive* (positif palsu) tinggi.
3. *Recall*: Menilai seberapa banyak data yang benar-benar positif berhasil dikenali oleh model.
4. *F-Measure*: Merupakan rata-rata harmonis antara *precision* dan *recall*, berguna saat dibutuhkan keseimbangan antara keduanya.

4. HASIL DAN PEMBAHASAN

Berikut ini hasil dan pembahasan analisis sentiment menggunakan algoritma *Naive Bayes Classifier* sebagai berikut.

4.1. Data Selection

Data Selection merupakan tahap awal dalam metode penelitian, yaitu mengumpulkan dataset yang sesuai dengan penelitian. Pengambilan data dilakukan dengan teknik *Crawling* menggunakan *Youtube API v3* pada *Google Cloud*. Hasil *crawling* yang dikumpulkan berupa komentar *Youtube* pada acara Ruangguru *Clash of Champions* (COC) episode pertama, dengan rentang waktu 29 Juni hingga 31 Desember 2024 dataset yang didapat terdiri dari 4 kolom yaitu: Tanggal, Username, Komentar, dan Like. Total data yang diperoleh sebanyak 15.432. Contoh hasil dari pengumpulan data dapat dilihat pada tabel 4.1 sebagai berikut.

Tabel 4.1 Hasil Crawling Data

Tanggal	Username	Komentar	Like
2024-07-03T06:39:10Z	@Ruangguru	Yuk, konsultasi dan klaim diskonnya sekarang di <a href="https://bit.ly/KlaimDiskonCOC" https://bit.ly/KlaimDiskonCOC	2107
2024-07-03T07:18:05Z	@BayuFocus	Emoh	81

4.2. Preprocessing

Pada tahapan ini terdiri dari enam proses yaitu *cleansing*, *case folding*, *tokenizing*, *spelling normalization*, *filtering*, *stemming*. Berikut merupakan contoh hasil dari preprocessing text pada tabel 4.2.

Tabel 4.2 Hasil Tahapan Preprocessing

Praproses	Teks
Dataset	Yuk, konsultasi dan klaim diskonnya sekarang di <a href="https://bit.ly/KlaimDiskonCOC" https://bit.ly/KlaimDiskonCOC
Cleansing	Yuk konsultasi dan klaim diskonnya sekarang di a
Case Folding	yuk konsultasi dan klaim diskonnya sekarang di a
Tokenizing	[yuk, konsultasi, dan, klaim, diskonnya, sekarang, di, a]
Spelling Normalization	[ayo, konsultasi, dan, klaim, diskonnya, sekarang, di, a]
Filtering	[ayo, konsultasi, klaim, diskonnya, sekarang]
Stemming	[ayo, konsultasi, klaim, diskon, sekarang]

Setelah seluruh tahap *preprocessing* selesai, token kata digabung kembali menjadi satu kalimat utuh dan disimpan dalam kolom baru bernama 'komentar_bersih'. Sebelum melanjutkan, dilakukan pembersihan tambahan,

seperti: Menghapus duplikat, Menghapus data kosong, Menghapus komentar pendek (kurang dari 3 kata), Normalisasi spasi berlebih, dan Reset index. Langkah ini bertujuan menghasilkan data yang lebih bersih dan siap untuk tahap data transformation.

Jumlah data awal sebanyak 15.432 komentar, dan setelah preprocessing tersisa 10.716 komentar. Gambar 4.1 menampilkan contoh hasil akhir dari proses preprocessing tersebut.

komentar_bersih	
0	ayo konsultasi klaim diskon sekarang
1	tonton bentrok juara lalu tonton youtube sendi...
2	buka ruangguru tonton bentrok juara buka sendi...
3	kangen clash of champions
4	universitas perang elit league keren akhir sek...

Gambar 4.1 Sampel data hasil akhir Preprocessing

4.3. Transformation

Sebelum melakukan transformasi data, dilakukan terlebih dahulu proses pembobotan sentimen untuk memberikan label (positif, negatif, atau netral) pada setiap komentar. Pelabelan ini penting sebagai dasar dalam proses pembobotan kata menggunakan TF-IDF.

Prosesnya dilakukan dengan menghitung jumlah kata positif, negatif, dan netral dalam setiap kalimat, berdasarkan kamus sentimen (lexicon). Dengan metode ini, pelabelan dilakukan secara otomatis tanpa input manual.

Kamus sentimen yang digunakan untuk pelabelan diperoleh dari file khusus berisi daftar kata-kata bermakna positif dan negatif. Penilaian dilakukan dengan memberi skor 1 untuk kalimat positif, dan skor -1 untuk kalimat negatif. Tabel 4.3 menampilkan contoh isi file tersebut, yang berisi kumpulan kata berdasarkan makna sentimennya.

Tabel 4.3 Sampel file Kamus Skor Sentimen

KAMUS SKOR SENTIMEN	
Kata Positif (+1)	Kata Negatif (-1)
Adil	Ancaman
Amanah	Bajingan
Bagus	Beban
Bahagia	Ceroboooh
Cerdas	Dendam
Efisien	Egois

Setelah menambahkan kamus sentimen yang berisi skor positif dan negatif, langkah

selanjutnya adalah menghitung skor sentimen pada setiap komentar hasil *preprocessing*. Proses ini dilakukan dengan menganalisis polaritas kalimat menggunakan kamus tersebut sebagai acuan.

- Skor > 0 → Sentimen Positif
- Skor < 0 → Sentimen Negatif
- Skor = 0 → Sentimen Netral

Proses ini membantu mengelompokkan komentar berdasarkan sentimennya. Gambar 4.2 menunjukkan hasil akhir dari pembobotan sentimen tersebut.

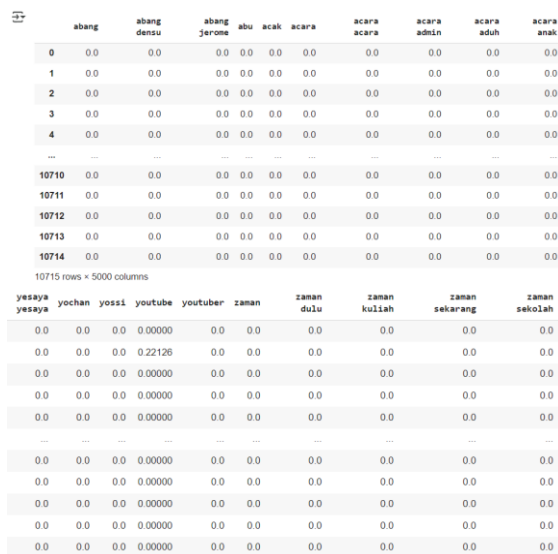
	komentar_bersih	polaritas	sentimen
0	ayo konsultasi klaim diskon sekarang	0	netral
1	tonton bentrok juara lalu tonton youtube sendi...	0	netral
2	buka ruangguru tonton bentrok juara buka sendi...	0	netral
3	kangen clash of champions	0	netral
4	universitas perang elit league keren akhir sek...	2	positif
5	ugm tunggu iya moga partisipasi	0	netral
6	iya akhir swafoto hari terima kasih kamu clash...	1	positif
7	kesini gara-gara gara-gara lihat aoc episode k...	0	netral
8	sini nonton ulang	-1	negatif
9	nonton ulang lagi kangen episode bulan yang lalu	-1	negatif

Gambar 4.2 Pembobotan Sentimen

Selanjutnya masuk ke dalam proses Transformasi data dengan menggunakan *Term Frequency-Inverse Document Frequency* (TF-IDF) tuned dengan mengatur

- *ngram_range* = (1, 2) menjadi Unigram (kata tunggal) seperti saya, suka, makan, dan sebagainya dan Bigram (dua kata) seperti saya suka, suka makan, makan nasi, dan sebagainya.
- *min_df* = 3 untuk menampilkan minimal tiga komentar.
- *max_df* = 0.9 menampilkan lebih dari 90% data komentar tujuannya untuk membuang kata-kata yang jarang muncul dan menghindari kata-kata yang tidak relevan.
- *max_features* = 5000 untuk mempertahankan maksimal 5000 kata atau frasa yang paling informatif berdasarkan skor TF-IDF.

Hasil proses pembobotan TF-IDF dapat dilihat dibawah ini.



10715 rows x 5000 columns

Gambar 4.3 Hasil proses pembobotan TF-IDF

Hasil proses pembobotan TF-IDF ini menghasilkan 5000 fitur kata yang mewakili data komentar dapat dilihat pada. Beberapa kata dengan bobot TF-IDF rata-rata tertinggi antara lain pada Tabel 4.4.

Tabel 4.4 Top 10 kata dengan rata-rata tertinggi TF-IDF

NO.	KATA/FRASA	RATA-RATA TF-IDF
1	Keren	0.028866
2	Iya	0.026795
3	Acara	0.024417
4	Tonton	0.022177
5	Pintar	0.022168
6	Orang	0.020698
7	Begini	0.018055
8	Anak	0.017597
9	Buat	0.015160
10	Lihat	0.014912

4.4. Data Mining

Proses data mining dilakukan menggunakan algoritma *Multinomial Naïve Bayes*. Terdapat dua tahap utama: pembagian data dan klasifikasi. Pembagian data menggunakan metode *K-Fold Cross Validation* sebanyak 10 kali, di mana data dibagi menjadi 10 bagian dan diuji secara bergantian.

Setelah proses klasifikasi, hasil prediksi dievaluasi menggunakan *confusion matrix* dari 10 kali pengujian. Tujuan dari validasi silang ini adalah untuk menilai kemampuan model dalam menggeneralisasi data baru dan menghindari overfitting.

Hasil akurasi dari total 10 *fold* dan rata-rata dari akurasi hasil akhir pengujian model *naive bayes* dapat dilihat pada tabel 4.5.

Tabel 4.5 Hasil accuracy dan avg_accuracy dari 10 fold

K-Fold	Accuracy
Fold-1	0.70
Fold-2	0.70
Fold-3	0.71
Fold-4	0.70
Fold-5	0.70
Fold-6	0.68
Fold-7	0.72
Fold-8	0.70
Fold-9	0.71
Fold-10	0.70
Rata-rata	0.70

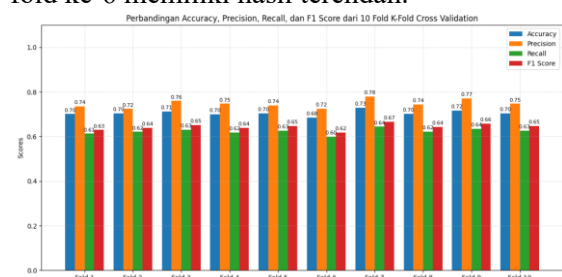
Berdasarkan hasil evaluasi pada Tabel 4.5 dari 10 fold, diperoleh rata-rata akurasi sebesar 70.55%. Hasil accuracy tertinggi terlihat pada fold-7 sebesar 72.85%, accuracy terendah pada fold-6 sebesar 68.47%.

4.5. Interpretation

Pada tahap Interpretation, evaluasi model dilakukan menggunakan Confusion Matrix dan K-Fold Cross Validation sebanyak 10 kali untuk memastikan stabilitas dan konsistensi performa model.

Setiap fold menggunakan data latih dan data uji yang berbeda, sehingga hasil evaluasi menjadi lebih menyeluruh. Nilai accuracy, precision, recall, dan F1-score dari setiap fold kemudian dirata-ratakan untuk mewakili performa model secara keseluruhan.

Tabel 4.6 dan Gambar 4.4 menunjukkan hasil evaluasi dari 10 fold pengujian, di mana fold ke-7 mencatat performa terbaik, sedangkan fold ke-6 memiliki hasil terendah.



Gambar 4.4 Grafik perbandingan dari 10 fold

Tabel 4.6 Hasil Grafik dari perbandingan ke 10 fold

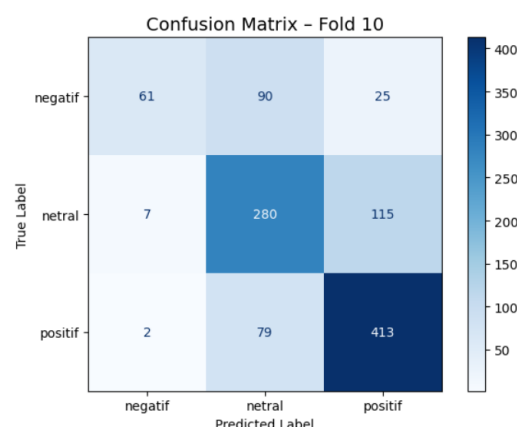
K-fold	Accuracy	Precision	Recall	F1-Score
Fold-1	0.70	0.73	0.61	0.62
Fold-2	0.70	0.72	0.62	0.63
Fold-3	0.71	0.75	0.63	0.65
Fold-4	0.70	0.74	0.61	0.63
Fold-5	0.70	0.73	0.62	0.64
Fold-6	0.68	0.72	0.59	0.61
Fold-7	0.72	0.78	0.64	0.66
Fold-8	0.70	0.74	0.62	0.64
Fold-9	0.71	0.77	0.63	0.65
Fold-10	0.70	0.74	0.62	0.64
Rata-rata	0.70	0.74	0.62	0.64

Rata-rata Evaluasi 10-Fold:
Accuracy 0.705559
Precision 0.747502
Recall 0.623666
F1-Score 0.643827
dtype: float64

Gambar 4.5 Hasil rata-rata evaluasi dari 10 fold

Gambar 4.5 menampilkan total rata-rata dari hasil evaluasi 10 fold yang menunjukkan rata-rata dari accuracy sebesar 70.55%, precision sebesar 74.75%, recall sebesar 62.36%, dan f1-score sebesar 64.38%.

Untuk memahami lebih dalam performa model, dilakukan analisis confusion matrix pada fold ke-10 dari proses K-Fold Cross Validation. Gambar 4.6 memperlihatkan perbandingan antara hasil prediksi model dan label asli pada tiga kelas sentimen: positif, netral, dan negatif.

**Gambar 4.6 Confusion Matrix dari fold ke-10**

Model menunjukkan performa terbaik dalam mengenali komentar positif, dengan 413 dari 494 komentar berhasil diklasifikasikan dengan benar. Sebaliknya, kelas negatif menjadi yang paling sulit dikenali, hanya 61 dari 176 komentar yang terklasifikasi tepat. Sebagian besar komentar negatif justru diklasifikasikan sebagai netral (90) atau positif (25).

Hasil ini menunjukkan bahwa model cenderung bias terhadap kelas positif, terutama saat menghadapi komentar netral atau ambigu. Hal ini bisa terjadi karena komentar netral sering mengandung kata-kata bernuansa positif seperti “terima kasih” atau “oke”, sehingga membingungkan model dalam klasifikasi.

5. KESIMPULAN

- Penelitian ini melakukan analisis sentimen terhadap opini masyarakat pengguna Youtube terhadap acara RuangRang Clash of Champions (COC) menggunakan algoritma Naïve Bayes dan pendekatan Knowledge Discovery in Database (KDD). Proses mencakup tahapan: Data Selection, Preprocessing, Transformation, Data Mining, dan Interpretation. Data dikumpulkan melalui Teknik crawling menggunakan Youtube API v3, menghasilkan 15.432 komentar dan balasan, yang setelah dibersihkan menjadi 10.722 data. Dengan hasil sentimen: positif (4946), negatif (1753), dan netral (4021). Hasil menunjukkan bahwa sentimen positif mendominasi, mengindikasikan respons positif masyarakat terhadap acara COC termasuk kelas sentimen positif.

- b. Model diuji menggunakan K-fold Cross Validation (10 fold) dengan Multinomial Naive Bayes. Evaluasi dilakukan menggunakan confusion matrix, menghasilkan:

- Accuracy: 70.55%
- Precision: 74.75%
- Recall: 62.36%
- F1-Score: 64.38%

Nilai-nilai ini menunjukkan bahwa model cukup efektif dalam mengklasifikasikan sentimen positif komentar Youtube ke dalam kategori positif, negatif, dan netral, dengan tingkat akurasi yang memadai

UCAPAN TERIMA KASIH

Penulis mengucapkan terima kasih kepada pihak-pihak terkait yang telah memberi dukungan terhadap penelitian ini.

DAFTAR PUSTAKA

- [1] R. F. Alhujaili and W. M. S. Yafooz, "Sentiment Analysis for Youtube Videos with User Comments: Review," *Proc. - Int. Conf. Artif. Intell. Smart Syst. ICAIS 2021*, pp. 814–820, 2021, doi: 10.1109/ICAIS50930.2021.9396049.
- [2] L. Lesmana, Mukrodin, and F. Nabyala, "Analisis Sentimen Pengguna Twitter PPDB Menggunakan Algoritma Multinomial Naive Bayes," *J. Sist. Inf. dan Teknol. Perad.*, vol. 1, no. 1, 2020, [Online]. Available: <https://journal.peradaban.ac.id/index.php/jsitp/article/view/604>
- [3] V. A. Fitri, R. Andreswari, and M. A. Hasibuan, "Sentiment analysis of social media Twitter with case of Anti-LGBT campaign in Indonesia using Naïve Bayes, decision tree, and random forest algorithm," *Procedia Comput. Sci.*, vol. 161, pp. 765–772, 2019, doi: 10.1016/j.procs.2019.11.181.
- [4] D. F. Zhafira, B. Rahayudi, and I. Indriati, "Analisis Sentimen Kebijakan Kampus Merdeka Menggunakan Naive Bayes dan Pembobotan TF-IDF Berdasarkan Komentar pada Youtube," *J. Sist. Informasi, Teknol. Informasi, dan Edukasi Sist. Inf.*, vol. 2, no. 1, pp. 55–63, 2021, doi: 10.25126/justsi.v2i1.24.
- [5] M. Yasir and R. Suraji, "Perbandingan Metode Klasifikasi Naive Bayes, Decision, Tree, Random Forest Terhadap Analisis Sentimen Kenaikan Biaya Haji 2023 pada Media Sosial Youtube," *J. Cahaya Mandalika*, vol. 3, no. 2, pp. 180–192, 2023.
- [6] P. Nandwani and R. Verma, "A review on sentiment analysis and emotion detection from text," *Soc. Netw. Anal. Min.*, vol. 11, no. 1, pp. 1–19, 2021, doi: 10.1007/s13278-021-00776-6.
- [7] Q. A. Chairunnisa, Y. Herdiyeni, M. Kusuma, D. Hardhienata, and J. Adisantoso, "Analisis Sentimen Pengguna Twitter Terhadap Program Vaksinasi Covid-19 di Indonesia Menggunakan Algoritme Support Vector Machine Sentiment Analysis of Twitter Users on COVID-19 Vaccination Program in Indonesia using Support Vector Machine Algorithm," *J. Ilmu Komput. dan Agri-Komputer*, vol. 9, no. 1, p. 79, 2022, [Online]. Available: <http://journal.ipb.ac.id/index.php/jika>
- [8] A. Khusnul Khotimah, "Analisis Sentimen Terhadap Kualitas Pelayanan," *JATI (Jurnal Mhs. Tek. Inform.)*, vol. 8, no. 3, pp. 3044–3048, 2024, doi: 10.36040/jati.v8i3.9520.
- [9] W. Astuti, R. Kurniawan, and Y. A. Wijaya, "Analisis Data Sentimen Ulasan Aplikasi Dana di Google Play Store Menggunakan Algoritma Naive Bayes," *J. Inform. dan Rekayasa Perangkat Lunak*, vol. 6, no. 1, pp. 158–163, 2024, doi: <http://dx.doi.org/10.36499/jinrpl.v6i1.10272>.
- [10] Wartumi, R. Kurniawan, and A. Y. Wijaya, "Analisis Data Sentimen Ulasan Pengguna Aplikasi Shopee di Google Play Store dengan Klasifikasi Algoritma Naive Bayes," *J. Inform. dan Rekayasa Perangkat Lunak*, vol. 6, no. 1, pp. 164–170, 2024.
- [11] I. Y. Yudo Bismo Utomo, Iin Kurniasari, "Penerapan Knowledge Discovery in Database," *J. Tek. Inform. Kaputama*, vol. 7, no. 1, 2023.
- [12] L. Muhamad, A. Arrozi, A. Y. Rahman, R. P. Putra, U. W. Malang, and K. Malang, "Klasifikasi tingkat tutur bahasa sasak berbasis teks menggunakan naïve bayes," vol. 12, no. 3, 2024.
- [13] F. S. Pamungkas and I. Kharisudin, "Analisis Sentimen dengan SVM, NAIVE BAYES dan KNN untuk Studi Tanggapan Masyarakat Indonesia Terhadap Pandemi Covid-19 pada Media Sosial Twitter," *Pros. Semin. Nas. Mat.*, vol. 4, pp. 1–7, 2021, [Online]. Available: <https://journal.unnes.ac.id/sju/prisma/article/view/45038>
- [14] D. Atika, A. Ari Aldino, S. Informasi, J. Pagar Alam No, L. Ratu, and K. Kedaton, "Term Frequency-Inverse Document Frequency Support Vector Machine Untuk Analisis Sentimen Opini Masyarakat Terhadap Tekanan Mental Pada Media Sosial Twitter," *J. Teknol. dan Sist. Inf.*, vol. 3, no. 4, p. page-page, 2022, [Online]. Available: <http://jim.teknokrat.ac.id/index.php/JTISI>

- [15] C. A. Salsabila, F. Yulianto, and T. A. Y. Siswa, "IMPLEMENTASI METODE NAIVE BAYES UNTUK KLASIFIKASI KECELAKAAN LALU LINTAS DI KOTA SAMARINDA," *J. Inform. dan Tek. Elektro Terap.*, vol. 13, no. 1, Jan. 2025, doi: 10.23960/jitet.v13i1.5890.
- [16] R. R. R. Arisandi, B. Warsito, and A. R. Hakim, "Aplikasi Naïve Bayes Classifier (Nbc) Pada Klasifikasi Status Gizi Balita Stunting Dengan Pengujian K-Fold Cross Validation," *J. Gaussian*, vol. 11, no. 1, pp. 130–139, 2022, doi: 10.14710/j.gauss.v11i1.33991.
- [17] Z. A. Sriyanti, D. S. Y. Kartika, and A. R. E. Najaf, "Implementasi Model Bert Pada Analisis Sentimen Pengguna Twitter Terhadap Aksi Boikot Produk Israel," *J. Inform. dan Tek. Elektro Terap.*, vol. 12, no. 3, pp. 2335–2342, 2024, doi: 10.23960/jitet.v12i3.4743.
- [18] E. A. Pratama, C. M. Hellyana, and N. I. Fadlilah, "Perbandingan 3 Algoritma Klasifikasi Data Mining Dalam Pro-Kontra Bahaya Rokok Elektrik," *J. Teknoinfo*, vol. 16, no. 1, p. 93, 2022, doi: 10.33365/jti.v16i1.1534.
- [19] M. Jamil, H. Hadiyanto, and R. Sanjaya, "Sentiment Analysis: Classifying Public Comments on YouTube in Disaster Management Simulation in Indonesia Using Naïve Bayes and Support Vector Machine," *Ingénierie des systèmes d Inf.*, vol. 29, no. 2, pp. 437–446, Apr. 2024, doi: 10.18280/isi.290205.