

ANALISIS PERBANDINGAN ALGORITMA *NAÏVE BAYES* SVM DAN *RANDOM FOREST* PADA KLASIFIKASI SENTIMEN KOMENTAR YOUTUBE "*CLASH OF* *CHAMPIONS*" RUANGGURU

Irlon.

Program Studi Teknik Informatika, FTI, Institut Teknologi Budi Utomo Jakarta,
dahil.irlon@gmail.com

Abstrak

Perkembangan platform digital seperti YouTube membuka peluang baru untuk menganalisis opini publik melalui komentar pengguna. Penelitian ini membandingkan kinerja tiga algoritma klasifikasi—*Naïve Bayes*, *Support Vector Machine (SVM)*, dan *Random Forest*—dalam mengklasifikasikan sentimen (positif, negatif, netral) pada komentar video “*Clash of Champions*” oleh Ruangguru. Data dikumpulkan dengan web crawling menggunakan YouTube API (Maret–Mei 2025), menghasilkan 15.423 komentar. Proses prapemrosesan mencakup *cleaning*, *casefolding*, normalisasi, tokenisasi, *stopword removal*, dan *stemming*. Pelabelan dilakukan dengan pendekatan leksikon dan pembobotan *TF-IDF* sebelum klasifikasi menggunakan Python (*scikit-learn*). Parameter: kernel linear (*SVM*), 100 estimator (*Random Forest*), model multinomial (*Naïve Bayes*). Evaluasi dengan akurasi, presisi, *recall*, dan *f1-score* menunjukkan *SVM* unggul dengan akurasi 83%, disusul *Random Forest* (82%) dan *Naïve Bayes* (68%). *SVM* menunjukkan kinerja paling seimbang dan direkomendasikan untuk klasifikasi sentimen komentar YouTube dalam studi ini.

Kata kunci: Analisis Sentimen, *Naïve Bayes*, *Support Vector Machine*, *Random Forest*, YouTube, *TF-IDF*.

1. PENDAHULUAN

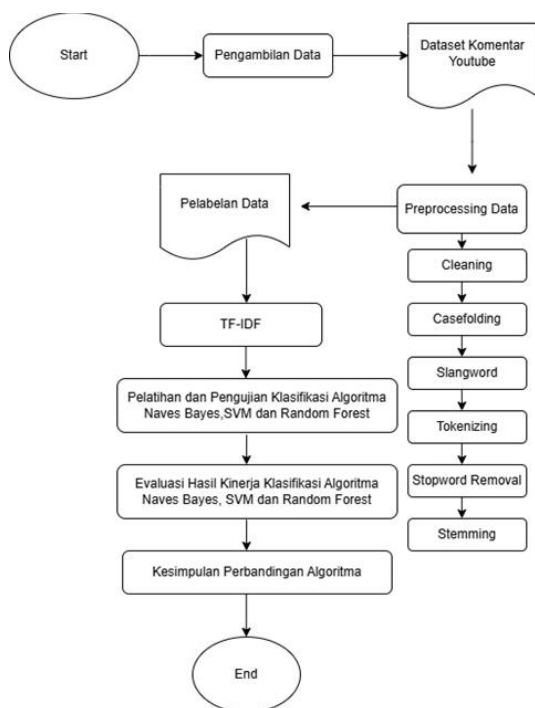
YouTube merupakan salah satu *platform* media sosial populer yang memungkinkan pengguna mengunggah dan menonton konten visual (Hendrawan & Sela, 2024). *Platform* ini dimanfaatkan oleh berbagai kalangan, termasuk perusahaan edukasi seperti Ruangguru yang memiliki video edukasi dengan tingkat keterlibatan tinggi dari pengguna (Ruangguru, 2024). Salah satu pendekatan untuk memahami respons audiens terhadap konten tersebut adalah dengan menganalisis sentimen komentar yang ditinggalkan penonton menggunakan metode analisis sentimen berbasis pembelajaran mesin (Purnamasari et al., 2023). Komentar pengguna pada video Ruangguru mencerminkan beragam opini publik, mulai dari dukungan, kritik, hingga saran, sehingga analisis sentimen dapat memberikan wawasan penting untuk mengevaluasi keberhasilan konten dan meningkatkan kualitas program mendatang.

Penelitian ini membandingkan kinerja tiga algoritma klasifikasi sentimen, yaitu *Naïve Bayes*, *Support Vector Machine (SVM)*, dan *Random Forest*, yang banyak digunakan dalam penelitian analisis teks karena performanya yang andal dalam mengenali pola data (Adrian et al., 2021). Untuk itu, diperlukan dataset yang berkualitas dan proses pra-pemrosesan teks mendalam, seperti *cleaning*, *casefolding*, *tokenizing*, *stopword removal*, dan *stemming* (Purnamasari et al., 2023). Tujuan penelitian ini tidak hanya menganalisis sentimen publik terhadap video “*Clash of Champions*” Ruangguru, tetapi juga mengevaluasi efektivitas ketiga algoritma tersebut dalam mendeteksi sentimen komentar baru di YouTube. Dengan demikian, penelitian ini diharapkan memberi kontribusi signifikan baik secara teoritis maupun praktis, khususnya dalam penerapan analisis sentimen pada konten edukasi digital.

2. METODOLOGI

2.1 Kerangka Pemikiran

Untuk menjelaskan alur logis dari pelaksanaan penelitian ini, diperlukan suatu kerangka berpikir yang menggambarkan proses penelitian secara sistematis, mulai dari pengumpulan data hingga evaluasi model klasifikasi. Kerangka pemikiran ini disusun untuk memperjelas tahapan-tahapan utama yang dilalui peneliti dalam melakukan analisis sentimen komentar YouTube menggunakan tiga algoritma machine learning, yaitu *Naïve Bayes*, *Support Vector Machine* (SVM), dan *Random Forest*. Dengan kerangka ini, hubungan antar proses mulai dari *preprocessing* data hingga evaluasi performa model dapat digambarkan secara runtut dan terpadu.



Gambar 1 Kerangka Pikir
Sumber : Pola pikir mandiri

Berikut adalah uraian singkat alur kerangka pemikiran dalam penelitian ini yang berfokus pada analisis sentimen komentar YouTube menggunakan algoritma *Naïve Bayes*, *Support Vector Machine* (SVM), dan *Random Forest*:

1. Pengumpulan data dilakukan dengan teknik *crawling* menggunakan YouTube API untuk mengambil komentar pada

video "*Clash of Champions*" milik Ruangguru.

2. Pembersihan data awal dilakukan dengan menghapus komentar duplikat agar data yang digunakan unik dan valid.
3. Pra-pemrosesan data (*preprocessing*) mencakup tahapan *cleaning*, *casefolding*, normalisasi kata slang, tokenisasi, *stopword removal*, dan *stemming* untuk mengubah data mentah menjadi format yang siap dianalisis.
4. Pelabelan data dilakukan menggunakan pendekatan leksikon untuk pendekatan leksikon untuk mengklasifikasikan setiap komentar ke dalam kategori sentimen: positif, negatif, atau netral.
5. Pelabelan data dilakukan menggunakan pendekatan leksikon untuk mengklasifikasikan setiap komentar ke dalam kategori sentimen: positif, negatif, atau netral.
6. Pembagian data ke dalam data latih dan data uji bertujuan untuk melatih model klasifikasi dan mengevaluasi performanya secara terpisah.
7. Transformasi data menggunakan metode *Term Frequency-Inverse Document Frequency* (TF-IDF) agar teks dapat direpresentasikan dalam bentuk numerik.
8. Klasifikasi sentimen dilakukan menggunakan tiga algoritma, yaitu *Naïve Bayes*, *SVM*, dan *Random Forest* untuk memprediksi sentimen berdasarkan model yang telah dilatih.

2.2 Metode Analisa Data

Metode analisis data dilakukan untuk mengolah komentar YouTube agar dapat diklasifikasikan ke dalam sentimen positif, negatif, dan netral. Tahapan meliputi pra-pemrosesan teks, pelabelan sentimen, transformasi data menggunakan TF-IDF, lalu klasifikasi menggunakan algoritma *Naïve Bayes*, *SVM*, dan *Random Forest*. Hasil klasifikasi kemudian dievaluasi dengan metrik akurasi, presisi, *recall*, dan *f1-score*.

2.2.1. Text Preprocessing

Suatu proses Guna mentransformasi data tidak terstruktur ke dalam format yang lebih terorganisir dan dapat dianalisis secara efektif, Tujuan *preprocessing* adalah untuk menghasilkan sebuah set kata indeks yang dapat digunakan untuk menggambarkan data. Sebagai bagian dari proses *preprocessing* data, baris duplikat dan baris kosong dihapus dari dataframe. kemudian melakukan *preprocessing* seperti:

1. *Cleaning* : Melakukan pembersihan teks, termasuk menghapus elemen-elemen seperti mention, hashtag, retweet, tautan, angka, tanda baca, dan spasi di awal maupun akhir kalimat.
2. *Casefolding* Seluruh karakter dalam teks dikonversi menjadi huruf kecil oleh fungsi ini guna memastikan konsistensi data.
3. *Slangword* : Fungsi ini digunakan untuk mengubah kalimat tidak baku menjadi kalimat baku.
4. *Tokenizing* : Proses pembagian teks menjadi bagian yang lebih kecil. untuk mempermudah analisis teks lebih lanjut.
5. *Stopword Removal*: Melakukan penghapusan kata-kata umum seperti "yang", "dengan", "dan", dan "atau" yang tidak memiliki makna kontekstual atau analitis.
6. *Stemming* : Dalam fungsi ini, stemming digunakan pada teks, yang berarti mengubah kata-kata ke bentuk aslinya. untuk stemming dalam bahasa Indonesia, Anda dapat menggunakan literatur sastra.

2.2.2 Labeling Data

Suatu proses *Labeling* memberikan label pada teks untuk sentimen, topik, atau kategori lainnya selama analisis teks. Ini dilakukan dengan menggunakan metode Lexicon Based pada setiap komentar. Labeling data merupakan langkah krusial agar model klasifikasi dapat dilatih dengan benar serta dievaluasi secara akurat.

2.2.3 Pembobotan Kata (TF-IDF)

Tahap ini mengubah data teks menjadi bentuk numerik dengan menghitung bobot tiap kata menggunakan metode TF-IDF (*Term Frequency–Inverse Document Frequency*). Bobot ditentukan berdasarkan frekuensi kemunculan kata dalam dokumen serta seberapa jarang kata tersebut muncul di keseluruhan dataset.

2.2.4 Klasifikasi Naïve Bayes

Setelah data diproses dan diberi bobot TF-IDF, dilakukan klasifikasi menggunakan algoritma Naïve Bayes. Algoritma ini mengelompokkan komentar ke dalam sentimen positif, negatif, atau netral berdasarkan probabilitas tertinggi dari setiap kata dalam data latih.

2.2.5 Klasifikasi *Support Vector Machine* (SVM)

SVM mengklasifikasikan data dengan mencari *hyperplane* terbaik yang memisahkan tiap kelas sentimen. Data hasil *preprocessing* dan TF-IDF ditransformasikan ke ruang fitur berdimensi tinggi, lalu diuji untuk mendapatkan hasil klasifikasi optimal.

2.2.6 Klasifikasi *Random Forest*

Random Forest memanfaatkan banyak *decision tree* untuk memprediksi sentimen. Setiap *tree* memberikan hasil klasifikasi, dan kelas akhir ditentukan berdasarkan voting mayoritas, sehingga meningkatkan akurasi dan mengurangi risiko *overfitting*.

2.3 Metode Pembahasan Hasil

Dalam analisis data dan *machine learning*, *confusion matrix* digunakan untuk mengevaluasi efektivitas model klasifikasi dengan membandingkan hasil prediksi dan data sebenarnya. Matriks ini terdiri dari empat komponen utama: *True Positive* (TP), prediksi positif yang benar; *True Negative* (TN), prediksi negatif yang benar; *False Positive* (FP), prediksi positif yang salah; dan *False Negative* (FN), prediksi negatif yang salah. Dari keempat nilai ini, dapat dihitung metrik evaluasi

seperti akurasi, presisi, *recall*, dan *f1-score* untuk menilai kinerja model.

1. Akurasi

Akurasi mengukur seberapa banyak prediksi yang dilakukan oleh model sesuai dengan nilai sebenarnya, dibandingkan dengan total seluruh prediksi.

$$\text{akurasi} = \frac{TP+TN}{TP+TN+FP+FN}$$

2. Presisi

Menunjukkan proporsi prediksi positif yang benar dibandingkan dengan seluruh prediksi positif yang dihasilkan oleh model

$$\text{presisi} = \frac{TP}{TP+FP}$$

3. Recall

Menghitung jumlah data positif yang berhasil dikenali atau diprediksi oleh model.

$$\text{recall} = \frac{TP}{TP+FN}$$

4. F1-Score

Evaluasi yang digunakan untuk mengevaluasi seberapa baik model klasifikasi berfungsi, terutama dalam hal masalah klasifikasi yang tidak seimbang. Untuk menilai keseimbangan antara dua metrik utama, Precision dan Recall, F1-Score menggabungkan keduanya menjadi satu nilai.

$$F1 = \frac{2 \times \text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}}$$

5. Area Under Curve (AUC)

AUC (*Area Under Curve*) adalah metrik evaluasi untuk mengukur kemampuan model klasifikasi dalam membedakan kelas sentimen menggunakan kurva ROC (*Receiver Operating Characteristic*). Nilainya berkisar 0–1; semakin mendekati 1 menunjukkan klasifikasi semakin baik, sedangkan 0,5 setara prediksi acak. AUC lebih kuat

dibanding metrik berbasis confusion matrix karena mempertimbangkan probabilitas, tahan terhadap ambang klasifikasi, dan efektif pada data tidak seimbang. Dalam penelitian ini, AUC digunakan untuk menilai seberapa baik model memisahkan kelas sentimen positif, negatif, dan netral

6. HASIL DAN PEMBAHASAN

6.1 Crawling Data

Data yang digunakan dalam penelitian ini adalah data primer yang diperoleh dari komentar pada channel youtube Ruangguru, khususnya pada serial "*Clash of Champions*". Pengambilan data di peroleh dengan teknik crawling.

Proses pengumpulan data dilakukan melalui *crawling* menggunakan YouTube API untuk mengambil komentar dari video "*Clash of Champions*" milik Ruangguru. Komentar yang berhasil dikumpulkan mencakup informasi tanggal publikasi, nama pengguna, isi komentar, dan jumlah like. Data tersebut kemudian disusun dalam bentuk dataframe sehingga siap untuk dianalisis. Hasil *crawling* ini ditampilkan pada Gambar 2 sebagai gambaran data mentah sebelum tahap pra-pemrosesan.

```
import pandas as pd
df = pd.DataFrame(comments, columns=['publishedAt', 'authorDisplayName', 'textDisplay', 'likeCount'])
df
```

	publishedAt	authorDisplayName	textDisplay	likeCount
0	2024-07-03T06:39:10Z	@Ruangguru	Yuk, konsultasi dan kiamat diskonnya sekarang d...	2100
1	2024-07-03T07:16:05Z	@BayuFocus	Emoh	82
2	2024-07-03T08:46:15Z	@FazaZa-kidz	👍👍👍	26
3	2024-07-03T11:38:14Z	@Mangputi	Nonton video clash of champions, lalu nonton v...	0
4	2024-07-03T12:14:35Z	@ulaula3577	mantap min	12
...	...	...	...	...
15419	2024-06-28T11:32:53Z	@rijwaaf	besok kak...	0
15420	2024-06-28T11:44:50Z	@fa_nny	yaallah kukira skrg Uda tanggal 29@rijwaaf	1
15421	2024-06-28T23:00:07Z	@desyrisanta4775	@rijwaafjam brpa ya	0
15422	2024-06-28T11:02:47Z	@SaprianPereira	Gak sabar nggu besok yaa...	28
15423	2024-06-28T10:57:40Z	@relFaanel	Wow	0

15424 rows x 4 columns

Gambar 2. Hasil *crawling*
Sumber : Olahan Data Mandiri

Langkah selanjutnya adalah mengekstrak data nama pengguna dan komentar, lalu mengganti nama kolom data untuk penyesuaian. seperti pada gambar 3.

```
df = df[['authorDisplayName', 'textDisplay']]
df.head()
```

	authorDisplayName	textDisplay
0	@Ruangguru	Yuk, konsultasi dan klaim diskonnya sekarang d...
1	@BayuFocus	Emoh
2	@FazaZa-k8fr	🔥🔥🔥
3	@Mangputt	Nonton video clash of champions, lalu nonton v...
4	@juliaulia3577	mntap min

```
df = df.rename(columns={'authorDisplayName': 'Nama Pengguna', 'textDisplay': 'Komentar'})
df.head()
```

	Nama Pengguna	Komentar
0	@Ruangguru	Yuk, konsultasi dan klaim diskonnya sekarang d...
1	@BayuFocus	Emoh
2	@FazaZa-k8fr	🔥🔥🔥
3	@Mangputt	Nonton video clash of champions, lalu nonton v...
4	@juliaulia3577	mntap min

Gambar 3. Dataset analisis
Sumber : Olahan Data Mandiri

Data komentar yang telah tersusun dalam bentuk dataframe kemudian disimpan ke dalam format berkas CSV.

Nama Pengguna	Komentar
@Ruangguru	Yuk, konsultasi dan klaim diskonnya sekarang d...
@BayuFocus	Emoh
@FazaZa-k8fr	🔥🔥🔥
@Mangputt	Nonton video clash of champions, lalu nonton v...
@juliaulia3577	mntap min

Gambar 4. Dataset CSV
Sumber : Olahan Data Mandiri

3.2 Analisis Klasifikasi Sentimen

Analisis klasifikasi sentimen dilakukan untuk mengevaluasi performa tiga algoritma yang digunakan, yaitu *Naïve Bayes*, *Support Vector Machine* (SVM), dan *Random Forest*, dalam mengelompokkan komentar YouTube ke dalam kategori positif, negatif, dan netral. Proses evaluasi dilakukan dengan menggunakan metrik utama seperti akurasi, presisi, *recall*, dan *f1-score* yang diperoleh dari hasil pengujian pada data uji. Selain itu, kurva ROC dan nilai AUC juga digunakan untuk menilai kemampuan diskriminatif masing-masing model dalam membedakan kelas sentimen pada berbagai *threshold*. Hasil analisis ini memberikan gambaran komprehensif mengenai algoritma yang paling efektif dan stabil dalam mengklasifikasikan komentar serta menjadi dasar rekomendasi pemilihan model terbaik untuk studi serupa di masa mendatang.

3.2.1 Preprocessing

Tahap ini merupakan proses transformasi data tidak terstruktur menjadi data terstruktur agar dapat dianalisis secara sistematis. Tujuan dari *preprocessing* adalah untuk

menghasilkan himpunan term index yang dapat merepresentasikan data secara terstruktur dan dapat diproses oleh algoritma.

3.2.1.1 Cleaning

Tahap *cleaning* dilakukan untuk membersihkan data komentar YouTube dari elemen yang tidak relevan sehingga siap dianalisis. Proses ini meliputi penghapusan tanda baca, angka, tautan, simbol, serta karakter khusus seperti emoji yang tidak memiliki makna kontekstual dalam analisis sentimen. Selain itu, spasi berlebih dan teks duplikat juga dihilangkan untuk memastikan kualitas dan keunikan data. Hasil akhir tahap ini menghasilkan komentar dengan format teks yang lebih rapi dan konsisten sebagai dasar untuk proses pra-pemrosesan berikutnya.

Tabel 1. Hasil proses cleaning

No	Komentar	Cleaning
1	Yuk konsultasi dan klaim diskonnya sekarang	Yuk konsultasi dan klaim diskonnya sekarang
2	Gaga ini kalo ada yang menang 🏆🏆🏆	Gaga ini kalo ada yang menang
3	@juliaulia3577 1. Buka video ruangguru, nonton Clash Champions Bu...	Buka video ruangguru nonton Clash Champions Bu
4	tim yg nonton ulang coc d 2025	tim nonton ulang coc
5	ahref="https://www.youtube.com/watch?v=dY_3acmFFRQ∓=406">6:46 alfie kamu peramal ya 🤔	alfie kamu peramal

Sumber : Olahan Data Mandiri

3.2.1.2 Casefolding

Tahap *casefolding* dilakukan dengan mengubah seluruh huruf pada komentar menjadi huruf kecil untuk menyamakan format teks. Langkah ini penting agar kata dengan huruf kapital dan huruf kecil dianggap sama oleh sistem, sehingga proses analisis selanjutnya seperti tokenisasi dan pembobotan kata dapat dilakukan secara konsisten tanpa adanya duplikasi makna akibat perbedaan penulisan huruf.

Tabel 2 Hasil proses *casefolding*

No	Cleaning	Casefolding
1	Yuk konsultasi dan klaim diskonnya sekarang	yuk konsultasi dan klaim diskonnya sekarang
2	Gaga ini kalo ada yang menang	gaga ini kalo ada yang menang
3	SIAPA YG MASIH NONTON DI	siapa yg masih nonton di
4	KANGEN COC	kangen coc
5	alfie kamu peramal	alfie kamu peramal

Sumber : Olahan Data Mandiri

3.2.1.2 Slangword

Tahap *slangword* dilakukan untuk mengganti kata tidak baku atau bahasa gaul pada komentar menjadi kata baku sesuai kamus bahasa Indonesia. Proses ini penting agar istilah yang memiliki arti sama namun ditulis dengan cara berbeda dapat diproses secara konsisten oleh algoritma. Dengan normalisasi ini, kualitas representasi data meningkat dan analisis sentimen menjadi lebih akurat.

Tabel 3 Hasil proses *Slangword*

No	Cleaning	Slangword
1	mntap min	mantap admin
2	tim nonton ulang coc	tim menonton ulang coc
3	doang nih dri	doang nih dari
4	salfok dua jari semua	salah fokus dua jari semua
5	shakiraa kmu keren	shakiraa kamu keren

Sumber : Olahan Data Mandiri

3.2.1.3 Tokenizing

Tahap tokenizing dilakukan dengan memecah teks komentar menjadi potongan-potongan kata atau token. Proses ini bertujuan untuk mempermudah analisis pada tingkat kata sehingga setiap kata dapat diproses secara terpisah dalam tahapan selanjutnya, seperti penghapusan kata tidak penting (*stopword removal*) dan pembobotan kata dengan TF-IDF. Dengan tokenisasi, teks yang semula berupa kalimat utuh menjadi lebih terstruktur dan siap untuk dianalisis oleh algoritma klasifikasi.

Tabel 4 Hasil proses Tokenizing

No	Slangword	Tokenize
1	mantap admin	['mantap', 'admin']
2	tim menonton ulang coc	['tim', 'menonton', 'ulang', 'coc']
3	doang nih dari	['doang', 'nih', 'dari']
4	salah fokus dua jari semua	['salah', 'fokus', 'dua', 'jari', 'semua']
5	shakiraa kamu keren	['shakiraa', 'kamu', 'keren']

Sumber : Olahan Data Mandiri

3.2.1.4 Stopword Removal

Tahap *stopword removal* dilakukan untuk menghapus kata-kata umum yang tidak memiliki makna penting dalam analisis sentimen, seperti “dan”, “atau”, serta “yang”. Penghapusan kata ini bertujuan mengurangi gangguan pada proses analisis sehingga hanya kata bermakna utama yang dipertahankan. Dengan demikian, hasil representasi data menjadi lebih ringkas dan fokus pada kata-kata yang benar-benar berpengaruh terhadap klasifikasi sentiment

Tabel 5 Hasil proses Stopword Removal

No	Tokenize	Stopword Removal
1	['yuk', 'konsultasi', 'dan', 'klaim', 'diskonnya', 'sekarang']	['yuk', 'konsultasi', 'klaim', 'diskonnya']
2	['shakira', 'dia', 'yang', 'buka', 'game', 'dia', 'juga', 'yang', 'tutup', 'gamenya']	['shakira', 'buka', 'game', 'tutup', 'gamenya']
3	['siapa', 'yang', 'menonton', 'ulang']	['menonton', 'ulang']
4	['iya', 'lagi']	['iya']
5	['kangen', 'jadi', 'menonton', 'lagi']	['kangen', 'menonton']

Sumber : Olahan Data Mandiri

3.2.1.5 Stemming

Tahap *stemming* dilakukan untuk mengubah kata berimbuhan menjadi bentuk dasar atau kata akar. Proses ini penting agar variasi kata dengan makna sama, seperti “bermain”, “memainkan”, dan “mainan”, diperlakukan sebagai satu kata yang sama dalam analisis. Dengan penyederhanaan ini, jumlah fitur kata berkurang sehingga model klasifikasi dapat bekerja lebih efisien dan akurat.

Tabel 6 Hasil proses Stemming

No	Stopword Removal	Stemming
1	['menonton', 'ulang']	tonton ulang
2	['pintar', 'menonton', 'orang', 'pintar']	pintar tonton orang pintar
3	['merekomendasikan', 'acara', 'muridmurid', 'biar', 'menonton']	rekomendasi acara muridmurid biar tonton
4	['berkompetisi', 'mahasiswa', 'berprestasi', 'dulunya', 'pakai', 'aplikasi', 'ruang', 'guru', 'ya']	kompetisi mahasiswa prestasi dulunya pakai aplikasi ruang guru ya
5	['menikmati', 'acara']	nikmat acara

Sumber : Olahan Data Mandiri

3.2.1 Pelabelan Data

Tahap pelabelan data dilakukan untuk mengklasifikasikan setiap komentar ke dalam kategori sentimen positif, negatif, atau netral. Proses ini menggunakan pendekatan berbasis leksikon, di mana daftar kata berpolaritas digunakan sebagai acuan penentuan label. Dengan pelabelan ini, data komentar yang sebelumnya tidak berlabel dapat digunakan sebagai data latih dan data uji untuk proses klasifikasi pada algoritma machine learning.

Tabel 7 Score polarity

No	Komentar	Score Polarity	Polarity
1	yuk konsultasi klaim diskon	7	Positif
2	rewatch kangen	0	Netral
3	pintar manfaat	5	Positif
4	kak maxwell kak chindita kayak pisah	-3	Negatif
5	jir gua tegang gua tonton doang	-5	Negatif

Sumber : Olahan Data Mandiri

3.2.1 Pembobotan kata TF-IDF

Tahap pembobotan kata menggunakan metode Term Frequency-Inverse Document Frequency (TF-IDF) dilakukan untuk mengubah teks komentar menjadi representasi numerik. Metode ini menghitung seberapa sering suatu kata muncul dalam sebuah komentar (TF) dan membandingkannya dengan seberapa jarang kata tersebut muncul di seluruh komentar (IDF). Dengan pembobotan ini, kata yang memiliki makna penting akan

mendapatkan bobot lebih tinggi, sedangkan kata umum akan berbobot rendah, sehingga memudahkan algoritma klasifikasi dalam mengenali pola sentimen.

Tabel 8 Hasil TF IDF

Term	TF			IDF TF-IDF		
	D1	D2	D3	D1	D2	D3
admin	0	0	1	1.693147	0.0	0.000000
betapa	0	1	0	1.693147	0.0	0.213201
champions	0	1	0	1.693147	0.0	0.213201
channel	0	1	0	1.693147	0.0	0.213201
clash	0	1	0	1.693147	0.0	0.213201
diskon	1	0	0	1.693147	0.5	0.000000
dumbnya	0	1	0	1.693147	0.0	0.213201
faham	0	1	0	1.693147	0.0	0.213201
klaim	1	0	0	1.693147	0.5	0.000000
konsultasi	1	0	0	1.693147	0.5	0.000000
langsung	0	1	0	1.693147	0.0	0.213201
mantap	0	0	1	1.693147	0.0	0.000000
melalui rekaman	0	1	0	1.693147	0.0	0.213201
tonton	0	1	0	1.693147	0.0	0.213201
video	0	1	0	1.693147	0.0	0.213201
yotube	0	1	0	1.693147	0.0	0.213201
yuk	1	0	0	1.693147	0.5	0.000000

Sumber : Olahan Data Mandiri

3.3 Hasil Perbandingan Algoritma

Tabel 9 Perbandingan Kalifikasi Algoritma

Klasifikasi	Akurasi Latih	Akurasi Uji	Precision	Recall	F1-score
Naïve Bayes	81%	68%	75%	65%	69%
SVM	97%	83%	83%	84%	84%
Random Forest	99%	82%	83%	83%	83%

Sumber : Olahan Data Mandiri

Tabel 9 perbandingan menunjukkan kinerja tiga algoritma klasifikasi, yaitu *Naïve Bayes*, *SVM*, dan *Random Forest*, berdasarkan lima metrik utama: akurasi latih, akurasi uji, *precision*, *recall*, dan *f1-score*. *Naïve Bayes* memiliki akurasi latih 81% dan akurasi uji 68%, dengan *precision* 75%, *recall* 65%, dan *f1-score* 69%. Hasil ini menandakan model cukup baik dalam mendeteksi sentimen tertentu, namun kesulitan dalam mengenali sentimen netral sehingga performanya menurun saat diuji pada data baru.

Sementara itu, *SVM* menunjukkan hasil terbaik dengan akurasi latih 97% dan akurasi uji 84%, disertai *precision*, *recall*, dan *f1-score* yang seimbang (masing-masing 83–84%). *Random Forest* memiliki akurasi latih tertinggi 99% namun akurasi uji sedikit menurun menjadi 82%, dengan *precision* dan *recall* 83%. Secara keseluruhan, *SVM* menjadi algoritma

paling stabil dan unggul, disusul Random Forest, sementara Naïve Bayes berada di posisi terendah terutama dalam menangani sentimen netral

3.4 Evaluasi Algoritma

Hasil evaluasi menunjukkan bahwa SVM memberikan performa terbaik di antara tiga algoritma klasifikasi, dengan akurasi latih 97%, akurasi uji 84%, serta *precision*, *recall*, dan *f1-score* yang seimbang di angka 83–84%, disertai AUC tertinggi 95% untuk kelas positif dan negatif. Random Forest menempati posisi kedua dengan akurasi latih 99% dan akurasi uji 82%, serta AUC stabil 93–94% pada semua kelas. Sementara itu, Naïve Bayes memiliki akurasi latih 81% dan akurasi uji 68%, dengan *precision* 75%, *recall* 65%, *f1-score* 69%, serta AUC terendah terutama pada kelas netral (87%), sehingga kurang optimal dibanding dua algoritma lainnya.

4. KESIMPULAN

Proses analisis sentimen diawali dengan pra-pemrosesan komentar video YouTube “*Clash of Champions*” milik Ruangguru, mencakup *cleaning*, *casefolding*, *tokenizing*, *stopword removal*, dan *stemming* untuk menyiapkan teks tidak terstruktur menjadi data siap analisis. Data kemudian diberi label sentimen (positif, negatif, netral) berbasis leksikon dan dibobot menggunakan TF-IDF sebelum diklasifikasikan menggunakan tiga algoritma: *Naïve Bayes*, *Support Vector Machine (SVM)*, dan *Random Forest*. Hasil evaluasi menunjukkan *Naïve Bayes* memiliki akurasi uji terendah (68,54%) dan kesulitan mendeteksi sentimen netral, *Random Forest* stabil dengan akurasi uji 82,86%, sementara SVM unggul dengan akurasi uji 83,79% serta *precision* dan *recall* yang seimbang. Secara keseluruhan, SVM menjadi algoritma paling efektif dan akurat untuk klasifikasi sentimen komentar pada penelitian ini.

DAFTAR PUSTAKA

A. Hendrawan and E. I. Sela, “Analisis Sentimen Komentar Youtube Tentang Resesi Global 2023 Menggunakan LSTM,” J. Indones. Manaj. Inform. dan Komun., vol. 5, no. 1, pp. 587–

593, 2024, doi:
10.35870/jimik.v5i1.526.

D. Purnamasari et al., *Pengantar Metode Analisis Sentimen*. 2023.

M. R. Adrian, M. P. Putra, M. H. Rafialdy, and N. A. Rakhmawati, “Perbandingan Metode Klasifikasi Random Forest dan SVM Pada Analisis Sentimen PSBB,” J. Inform. Upgris, vol. 7, no. 1, pp. 36–40, 2021, doi:
10.26877/jiu.v7i1.7099.

Ruangguru, “Tentang Ruangguru,” 2024. Accessed: Apr. 18, 2025. [Online]. Available:
<https://www.ruangguru.com/about-us>