

ANALISIS SENTIMEN KOMENTAR MASYARAKAT TERHADAP RANGKA ESAF HONDA MENGGUNAKAN ALGORITMA NAIVE BAYES DAN SUPER VECTOR MACHINE

Gusti Andika Joni, Nuri Cahyono*, Anna Baita, Nur Aini
Program Studi S1 Informatika, Universitas Amikom Yogyakarta
nuricahyono@amikom.ac.id

ABSTRAK

Honda, sebagai salah satu merek motor terkemuka di Indonesia, memperkenalkan teknologi *Enhanced Smart Architecture Frame* (eSAF) untuk meningkatkan stabilitas dan kenyamanan pengendalian sepeda motor. Namun, teknologi ini mendapatkan berbagai tanggapan dari masyarakat terkait masalah keropos, karat, dan kerentanannya terhadap patah, yang banyak dibahas di media sosial, terutama Instagram. Penelitian ini bertujuan untuk menganalisis sentimen terhadap kerangka eSAF menggunakan metode *Multinomial Naive Bayes* (MNB) dan *Support Vector Machine* (SVM). Data dikumpulkan dari akun Instagram resmi PT. Astra Honda Motor (@welovehonda_id) antara 29 Juli 2023 hingga 11 September 2023, dengan total 4607 komentar, yang setelah *pre-processing* berkurang menjadi 4435 komentar. *Labeling* sentimen dilakukan dengan kamus VADER, menghasilkan distribusi sentimen yang seimbang: 50,69% positif dan 49,31% negatif. Evaluasi model menggunakan rasio data 90:10, 80:20, 70:30, 60:40, dan 50:50 menunjukkan bahwa SVM unggul dalam Precision, dengan nilai tertinggi 91% pada rasio 70:30. MNB memiliki Precision tertinggi 87% pada rasio 90:10 dan Accuracy 87% pada rasio yang sama. Secara keseluruhan, SVM menunjukkan kinerja lebih baik dibandingkan MNB, terutama dalam ketepatan prediksi (*Precision*), sehingga direkomendasikan untuk analisis sentimen di masa depan.

Kata kunci : *Sentiment Analysis, Multinomial Naive Bayes, Support Vector Machine, Machine Learning*

1. PENDAHULUAN

Honda, sebagai salah satu merek motor terkemuka di Indonesia, terus berkomitmen untuk menghadirkan inovasi dalam produk mereka guna memenuhi kebutuhan konsumen. Salah satu inovasi terbaru dari Honda adalah penerapan teknologi *Enhanced Smart Architecture Frame* (eSAF), yang dirancang untuk meningkatkan lincahnya serta stabilitas *handling* sepeda motor, sehingga memberikan kemudahan dalam pengendalian dan kenyamanan saat bermanuver [1].

Meskipun teknologi ini menawarkan berbagai keunggulan, terdapat sejumlah kritik yang menyoroti kelemahan signifikan dari kerangka eSAF, seperti masalah keropos, karat, dan kerentanannya terhadap patah. Isu-isu ini telah memicu berbagai tanggapan dari masyarakat, yang banyak dibagikan melalui platform media sosial, terutama Instagram [2].

Instagram, sebagai salah satu platform media sosial terpopuler, memainkan peran penting dalam berbagi informasi dan interaksi antar pengguna. Pada Juli 2020, Indonesia tercatat sebagai negara dengan pengguna Instagram terbanyak keempat di dunia, dengan 73 juta pengguna atau 27% dari total populasi [3].

Dengan begitu banyak komentar dan umpan balik yang tersedia di media sosial, analisis sentimen menjadi alat penting untuk memahami persepsi publik terhadap isu-isu tertentu. Analisis sentimen merupakan salah satu area penelitian dalam *text mining* yang digunakan untuk mengkategorikan

dokumen teks berupa opini berdasarkan jenis sentimennya [6].

Dalam konteks ini, analisis sentimen dapat memberikan wawasan tentang opini, evaluasi, dan emosi masyarakat terhadap produk atau isu tertentu [3].

Penelitian ini bertujuan untuk melakukan analisis sentimen terhadap komentar masyarakat mengenai kerangka eSAF Honda yang terdapat di Instagram, dengan menggunakan dua metode klasifikasi utama: *Multinomial Naive Bayes* dan *Support Vector Machine* (SVM). *Multinomial Naive Bayes* adalah metode probabilistik yang beroperasi berdasarkan teorema Bayes dan frekuensi kemunculan kata dalam dokumen [4].

Sementara itu, SVM adalah algoritma klasifikasi yang bekerja dengan mencari hyperplane terbaik untuk memisahkan data [5]. Dengan menerapkan kedua metode ini, diharapkan dapat diperoleh pemahaman yang mendalam mengenai persepsi masyarakat terhadap kerangka eSAF dan memberikan masukan yang konstruktif bagi Honda dalam upaya perbaikan dan pengembangan produk di masa depan.

2. TINJAUAN PUSTAKA

2.1. Text Mining

Text mining adalah proses memperoleh pengetahuan dengan menggunakan alat analisis, di mana pengguna berinteraksi secara kontinu dengan kumpulan dokumen [10].

Text mining bertujuan untuk menganalisis opini, sentimen, evaluasi, sikap, penilaian, dan emosi

seseorang, sehingga dapat mengungkapkan pandangan mereka terhadap suatu topik, layanan, organisasi, individu, atau kegiatan tertentu. Langkah-langkah yang umum dilakukan dalam *text mining* meliputi *text pre-processing* dan pemberian bobot [3].

2.2. Analisis Sentimen

Analisis sentimen merupakan salah satu area penelitian dalam *text mining* yang digunakan untuk mengkategorikan dokumen teks berupa opini berdasarkan jenis sentimennya [6].

Proses ini meliputi pemahaman, ekstraksi, dan pemrosesan data teks secara otomatis untuk menghasilkan informasi yang berguna. Selain itu, analisis sentimen juga mempelajari opini, sentimen, evaluasi, sikap, dan emosi terkait dengan suatu entitas tertentu [3].

Tujuan utama dari analisis sentimen adalah untuk mengklasifikasikan teks yang berasal dari dokumen, kalimat, atau fitur tertentu ke dalam kategori positif, negatif, atau netral [8].

2.3. Instagram

Instagram merupakan sebuah platform media sosial yang menyediakan ruang bagi penggunanya untuk berbagi konten, berinteraksi, dan berkomunikasi melalui fitur *like* dan komentar pada unggahan pengguna lain [2].

Pada Juli 2020, Indonesia berada di posisi keempat dalam daftar negara dengan jumlah pengguna Instagram terbanyak di dunia, dengan total pengguna mencapai 73 juta, atau sekitar 27% dari keseluruhan populasi di Indonesia [3].

Instagram memungkinkan penggunanya untuk mengunggah foto dan video yang disertai dengan teks, yang kemudian dapat dimanfaatkan sebagai sumber data untuk analisis sentimen.

2.4. Text Pre-processing

Dalam penelitian ini, *text pre-processing* dilakukan melalui serangkaian tahapan yang berurutan, yaitu *case folding*, *cleaning*, *tokenizing*, *stopword removal*, dan *stemming* adalah beberapa tahapan dalam pemrosesan teks. *Case folding* berfungsi untuk mengonversi semua huruf besar dalam dokumen menjadi huruf kecil. *Cleaning* bertujuan untuk menghilangkan karakter-karakter yang tidak penting, seperti "https:", "@", "#", dan URL. *Tokenizing* memecah kalimat dalam dokumen menjadi kata-kata individu. *Stopword removal* bertujuan untuk menghilangkan kata-kata yang mengandung sedikit informasi atau dianggap kurang signifikan. *Stemming* menghilangkan imbuhan, baik awalan maupun akhiran, sehingga kata-kata dikembalikan ke bentuk dasarnya. Tujuan utama dari *text pre-processing* adalah untuk mengubah data teks yang tidak terstruktur menjadi data yang lebih terstruktur dan siap untuk dianalisis [7].

2.5. VADER Sentiment

Valence Aware Dictionary and Sentiment Reasoner (VADER) adalah metode analisis sentimen yang menggunakan kamus *lexicon* untuk menilai intensitas emosional dari data teks. Metode ini memungkinkan penilaian sentimen pada frasa dan kalimat tanpa memerlukan analisis tambahan yang kompleks. Sentimen yang dihasilkan oleh VADER dapat dikategorikan sebagai negatif, netral, atau positif, dan juga dapat dinyatakan secara numerik dalam bentuk skor intensitas. Pendekatan leksikal VADER menganalisis setiap kata dalam sebuah kalimat dan menggabungkan atribut *compound* dari setiap kata tersebut untuk menentukan polaritas keseluruhan kalimat. Kategorisasi sentimen dilakukan dengan menetapkan nilai numerik: 1 untuk sentimen positif ($\text{compound} > 0,5$), 0 untuk sentimen netral ($-0,5 \leq \text{compound} \leq 0,5$), dan -1 untuk sentimen negatif ($\text{compound} < -0,5$) [12].

2.6. TF-IDF

Term Frequency-Inverse Document Frequency (TF-IDF) adalah metode pemodelan yang digunakan untuk mengevaluasi relevansi antara istilah dan data dengan memberikan bobot pada setiap kata atau istilah [13]. Bobot kata akan lebih tinggi jika kata tersebut sering muncul dalam satu dokumen, tetapi akan lebih rendah jika kata tersebut muncul di banyak dokumen [14]. Metode ini dapat meningkatkan efisiensi dan akurasi dalam berbagai tugas pemrosesan teks [11]. Rumus untuk menghitung bobot kata menggunakan TF-IDF adalah sebagai berikut

$$w_{t,d} = tf_{t,d} \times idf_t = tf_{t,d} \times \log \frac{N}{df_t} \quad (1)$$

$w_{t,d}$ = Bobot TF-IDF

$tf_{t,d}$ = Jumlah frekuensi kata

idf_t = Jumlah *inverse* frekuensi dokumen tiap kata

df_t = Jumlah frekuensi dokumen tiap kata

N = Jumlah total dokumen

2.7. Multinomial Naive Bayes

Multinomial Naive Bayes adalah metode berbasis probabilitas yang mengacu pada teorema Bayes dan sering digunakan dalam NLP. Algoritma ini bergantung pada konsep frekuensi kata, yaitu jumlah kemunculan kata dalam suatu dokumen. Model ini mempertimbangkan dua faktor utama: keberadaan kata dalam dokumen dan frekuensi kemunculannya. Persamaan untuk MNB dituliskan sebagai berikut [14]:

$$P(p|n) \propto P(p) \prod_{1 \leq k \leq nd} P(t_k|p) \quad (2)$$

$P(t_k|p)$ = probabilitas munculnya dokumen text (tk)

n = jumlah dokumen dan

p = polaritas

Selanjutnya, untuk menghitung probabilitas kemunculan kata dalam dokumen dengan polaritas tertentu, dapat digunakan rumus berikut:

$$P(t_k|p) = \frac{\text{count}(t_k|p)+1}{\text{count}(t_p)+|V|} \quad (3)$$

(tk|p) = Jumlah kemunculan kata (tk) dalam dokumen yang memiliki polaritas p. count(tp) merujuk pada jumlah total token dalam dokumen dengan polaritas p.

2.8. Support Vector Machine

Support Vector Machine (SVM) adalah metode yang sering digunakan dalam analisis sentimen. Algoritma ini berfungsi dengan menentukan *hyperplane* terbaik untuk membagi data ke dalam dua kelompok berbeda, yaitu kelompok +1 dan kelompok -1, di mana setiap kelompok memiliki pola yang khas [5].

Algoritma ini bertujuan untuk menemukan *hyperplane* yang memisahkan kelas-kelas tersebut dengan margin terbesar. Model SVM dapat dirumuskan dengan persamaan berikut:

$$f(x) = \sum_{i=1}^{\infty} a_i y_i K(x, x') + b \quad (5)$$

a_i = Koefisien Lagrange untuk data latih ke- i

y_i = Kelas dari data latih ke- i

M = Jumlah total data latih

$K(x, x')$ = Fungsi *kernel* yang digunakan, di mana

x = Data uji

x_i = Data latih ke- i

b = Bias

2.9. Evaluasi Model

Confusion matrix berfungsi untuk menggambarkan jumlah prediksi yang akurat dan tidak akurat dari model [15].

Dalam matriks ini, setiap baris mewakili kelas sebenarnya dari data, sementara setiap kolom mewakili kelas yang diprediksi oleh model (atau sebaliknya) [9]. Tabel berikut menunjukkan hasil evaluasi model:

Tabel 1. Confusion Matrix

	Predicted Negative	Predicted Positive
Actual Negative	True Negative (TN)	False Positive (FP)
Actual Positive	False Negative (FN)	True Positive (TP)

Berbagai metrik evaluasi model dapat dihitung dari *confusion matrix*:

- Accuracy* mengukur seberapa sering model benar dalam mengklasifikasikan data secara keseluruhan. Formula untuk menghitung *accuracy* adalah:

$$\frac{TP+TN}{Total} \quad (6)$$

- Precision* mengukur seberapa akurat model mengklasifikasikan hasil sebagai positif. Formula untuk menghitung *precision* adalah:

$$\frac{TP}{FP+TP} \quad (7)$$

- Recall* mengukur seberapa sering model memprediksi positif ketika kelas sebenarnya adalah positif. Formula untuk menghitung *recall* adalah:

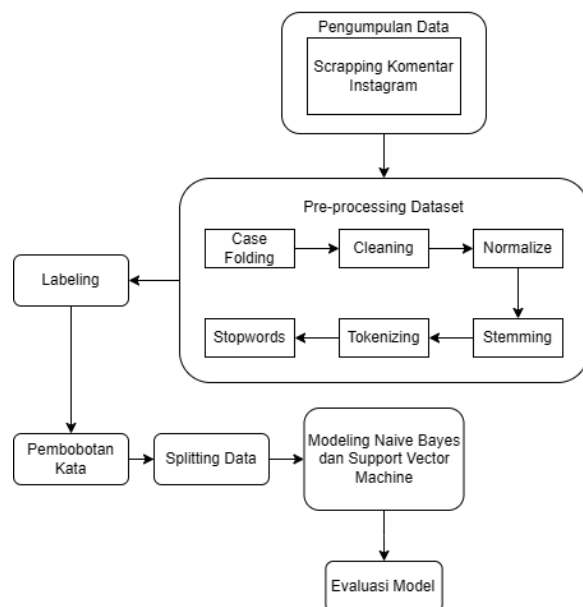
$$\frac{TP}{FN+TP} \quad (8)$$

- F1-Score* adalah ukuran gabungan dari *precision* dan *recall* yang memberikan keseimbangan antara keduanya.. Formula untuk menghitung *F1-Score* adalah:

$$2 * \frac{precision * recall}{precision + recall} \quad (9)$$

3. METODE PENELITIAN

Penelitian ini dilengkapi dengan *flowchart* yang mewakili berbagai tahapan penelitian untuk memudahkan pemahaman dengan menunjukkan langkah - langkah yang harus dilakukan.



Gambar 1. Alur Penelitian

Gambar 1 menunjukkan proses penelitian. Diagram alur ini memvisualisasikan langkah-langkah penelitian Anda dan membantu Anda memahami tahapan yang perlu dilakukan.

3.1. Pengumpulan Data

Data dikumpulkan melalui proses scrapping dari Instagram menggunakan PhantomBuster. Scrapping dilakukan pada akun Instagram resmi PT. Astra Honda Motor, @welovehonda_id, dari 29 Juli 2023 hingga 11 September 2023. Fitur Post Commenters Export digunakan untuk mengumpulkan komentar dari setiap postingan, yang kemudian diunduh dalam format CSV untuk analisis lebih lanjut.

3.2. Pre-processing Data

Proses pre-processing meliputi:

- Case Folding:** Semua teks diubah menjadi huruf kecil untuk konsistensi.
- Cleaning:** Menghapus elemen yang tidak relevan seperti tanda baca, angka, dan simbol.
- Normalization:** Meratakan variasi kata menjadi bentuk standar.
- Stemming:** Mengubah kata ke bentuk dasarnya untuk menyederhanakan variasi kata.
- Tokenizing:** Memecah teks menjadi token (kata atau frasa) untuk analisis.
- Stopwords Removing:** Menghapus kata-kata umum yang tidak signifikan untuk meningkatkan kualitas data.

3.3. Labeling

Labeling dilakukan menggunakan *library* VADER. VADER menganalisis sentimen dengan menggunakan kamus lexicon untuk menentukan intensitas emosional dari teks dan mengategorikannya sebagai positif, negatif, atau netral berdasarkan nilai compound[12].

3.4. Pembobotan Kata

Pembobotan kata dengan TF-IDF mengevaluasi peranan kata dalam dokumen berdasarkan keseluruhan kumpulan dokumen. Metode ini menghitung berapa kali sebuah kata muncul dalam dokumen (TF) dan berapa kali muncul di seluruh dokumen (IDF) dan memberikan bobot yang mencerminkan relevansi kata tersebut.

3.5. Splitting Data

Data dipisahkan menjadi bagian pelatihan dan bagian pengujian dengan rasio 90:10, 80:20, 70:30, 60:40, dan 50:50. Pembagian ini bertujuan untuk menemukan rasio yang menghasilkan akurasi tertinggi pada model dan memastikan bahwa model dapat diandalkan dalam melakukan generalisasi data.

3.6. Pembuatan Model

Model dibuat menggunakan dua metode: *Multinomial Naive Bayes* (MNB) dan *Support Vector Machine* (SVM). MNB adalah metode probabilistik yang efektif untuk klasifikasi teks berdasarkan frekuensi kemunculan kata, sementara SVM mencari batas pemisah terbaik untuk membagi data ke dalam kategori yang berbeda dengan margin maksimal.

3.7. Evaluasi Model

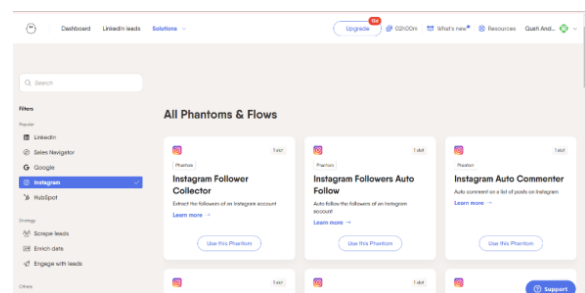
Evaluasi model dilakukan menggunakan Confusion Matrix. Matrix ini menyajikan perbandingan antara kelas sebenarnya dan kelas prediksi, memungkinkan analisis mendalam tentang akurasi, precision, recall, dan F1-score. Ini membantu dalam memahami kekuatan dan kelemahan model dalam mengklasifikasikan data.

4. HASIL DAN PEMBAHASAN

Hasil dan pembahasan akan dibagi menjadi beberapa subbab untuk penjelasan yang lebih rinci.

4.1. Pengumpulan Data

Penelitian ini memanfaatkan scrapping data dari Instagram menggunakan platform PhantomBuster untuk mengumpulkan komentar pada postingan dari akun resmi PT. Astra Honda Motor (@welovehonda_id). Proses scrapping dilakukan dari 29 Juli 2023 hingga 11 September 2023, menghasilkan 4607 komentar. Data dikumpulkan dengan fitur Post Commenters Export dari PhantomBuster, yang mengumpulkan komentar secara otomatis dari tautan postingan yang diberikan. Data yang diperoleh disimpan dalam format CSV untuk analisis lebih lanjut, memastikan data yang digunakan adalah interaksi nyata pengguna Instagram dengan konten PT. Astra Honda Motor.



Gambar 2. Website PhantomBuster

4.2. Pre-processing Data

Data komentar yang dikumpulkan dari akun Instagram @welovehonda_id sebanyak 4607 komentar, diproses melalui beberapa tahapan pre-processing:

- Case folding** mengubah seluruh teks menjadi huruf kecil. untuk menghindari perbedaan antara huruf kapital dan kecil. Misalnya, "Honda" menjadi "honda". Hal ini penting untuk konsistensi analisis teks.
- Cleaning** menghapus unsur yang tidak penting seperti simbol dan angka. Misalnya, "Honda ES@F is the best!" menjadi "Honda ESF is the best".
- Normalization** mengubah kata ke bentuk standar untuk mengurangi variasi bentuk kata. Misalnya, "mobil" dan "mobilnya" dinormalisasi menjadi "mobil".
- Stemming** mengubah kata ke bentuk dasarnya. Misalnya, "berlari" dan "pelari" menjadi "lari".
- Tokenizing** memecah teks menjadi unit yang lebih kecil seperti kata. Misalnya, "Honda adalah mobil yang bagus" menjadi ["Honda", "adalah", "mobil", "yang", "bagus"].
- Stopwords removing** menghapus istilah umum yang tidak menyumbang informasi signifikan dalam analisis. Misalnya, kata "dan" dan "yang" dihapus dari teks.

Tabel 2. Hasil Preprocessing

Proses	Sebelum	Sesudah
Case Folding	Pengen banget punya genio nih min @welovehonda_id ❤️❤️	pengen banget punya genio nih min @welovehonda_id ❤️❤️
Cleaning	pengen banget punya genio nih min @welovehonda_id ❤️❤️	pengen banget punya genio nih min
Normalize	pengen banget punya genio nih min	ingin banget punya genio nih min
Stemming	ingin banget punya genio nih min	ingin banget punya genio nih min
Tokenizing	ingin banget punya genio nih min	['ingin', 'banget', 'punya', 'genio', 'nih', 'min']
Stopwords Removing	['ingin', 'banget', 'punya', 'genio', 'nih', 'min']	['banget', 'genio', 'nih', 'min']

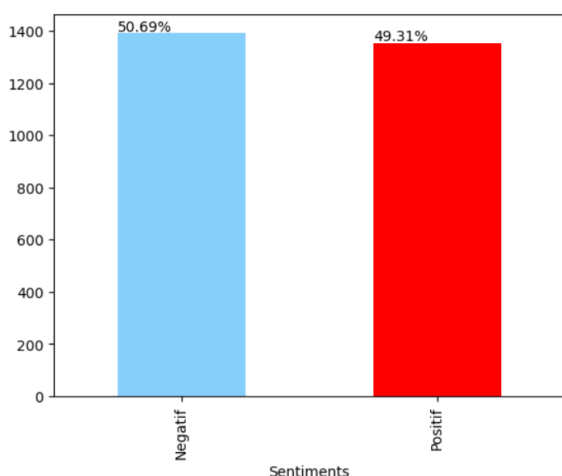
4.3. Labeling

Labeling menggunakan *library* VADER untuk menganalisis sentimen. VADER menilai intensitas emosional dari teks dengan menggunakan kamus lexicon yang mengkategorikan kata sebagai positif, negatif, atau netral dan memberikan skor numerik. Teks diberi label sebagai positif jika nilai compound > 0, dan negatif jika nilai compound ≤ 0.

Tabel 3. Hasil Labeling

Compound Score	Sentiments	Comment
0.9345	Positif	genio really hopefully win...
-0.0276	Negatif	min sparepart honda difficult...

Hasil labeling ini menunjukkan distribusi sentimen dengan 50,69% positif dan 49,31% negatif.



Gambar 3. Grafik Persebaran Data

Distribusi ini hampir seimbang, yang penting untuk memastikan bahwa model dapat memprediksi sentimen secara akurat tanpa bias terhadap salah satu kategori.

4.4. Pembobotan Kata

Metode pembobotan kata TF-IDF menilai seberapa sering kata muncul dalam dokumen dan seberapa umum kata tersebut di seluruh kumpulan dokumen untuk mengevaluasi relevansi kata dalam dokumen dibandingkan dengan kumpulan dokumen secara keseluruhan.

Tabel 4. Pembobotan Kata

term	Rank
broken	170.013864
frame	156.217910
honda	152.043988
skeleton	135.934088

4.5. Splitting Data

Dalam penelitian ini, data akan dibagi dengan empat rasio berbeda: 90:10, 80:20, 70:30, dan 60:40. Proses ini, membagi dataset menjadi *training data* untuk melatih model dan *testing data* untuk menilai kinerjanya. Pembagian ini bertujuan untuk menemukan rasio yang menghasilkan akurasi tertinggi pada model dan memastikan bahwa model dapat diandalkan dalam melakukan generalisasi data.

Tabel 5. Rasio Pembagian Data

Data Train	Data Test
90%	10%
80%	20%
70%	30%
60%	40%

4.6. Multinomial Naive Bayes

Model Multinomial Naive Bayes menunjukkan hasil yang relatif stabil pada berbagai rasio pembagian data. Hasil yang signifikan terlihat pada rasio 90:10 dengan nilai precision 87%, recall 84%, F1-Score 85%, dan accuracy 87%. Berikut tabel hasil performa MNB:

Tabel 6. Hasil Model MNB

Rasio	Precision	Recall	F1-score	Acc
90:10	90:10	90:10	90:10	90:10
80:20	80:20	80:20	80:20	80:20
70:30	70:30	70:30	70:30	70:30
60:40	60:40	60:40	60:40	60:40
50:50	50:50	50:50	50:50	50:50

4.7. Support Vector Machine Learning

Model SVM menunjukkan kinerja yang konsisten dengan nilai precision dan recall yang tinggi. Terutama pada rasio 90:10, SVM mencapai precision 89%, recall 89%, F1-Score 89%, dan accuracy 87%. Berikut tabel hasil performa SVM:

Tabel 7. Hasil Model SVM

Rasio	Precision	Recall	F1-score	Acc
90:10	90:10	90:10	90:10	90:10
80:20	80:20	80:20	80:20	80:20
70:30	70:30	70:30	70:30	70:30
60:40	60:40	60:40	60:40	60:40
50:50	50:50	50:50	50:50	50:50

4.8. Evaluasi

Penilaian model dilakukan dengan menerapkan metrik Confusion Matrix, Accuracy, Precision, Recall, dan F1-Score untuk menilai kinerja MNB dan SVM.

Tabel 8. Perbandingan Kinerja Multinomial Naive Bayes dan Support Vector Machine

Metrik	MNB (Rasio)	SVM (Rasio)
90:10	Precision: 87%, Recall: 84%, F1-Score: 85%, Accuracy: 87%	Precision: 89%, Recall: 89%, F1-Score: 89%, Accuracy: 87%
80:20	Precision: 84%, Recall: 83%, F1-Score: 83%, Accuracy: 84%	Precision: 90%, Recall: 87%, F1-Score: 89%, Accuracy: 84%
70:30	Precision: 85%, Recall: 84%, F1-Score: 85%, Accuracy: 84%	Precision: 91%, Recall: 87%, F1-Score: 89%, Accuracy: 84%
60:40	Precision: 84%, Recall: 84%, F1-Score: 84%, Accuracy: 84%	Precision: 90%, Recall: 87%, F1-Score: 88%, Accuracy: 84%
50:50	Precision: 84%, Recall: 85%, F1-Score: 85%, Accuracy: 84%	Precision: 90%, Recall: 84%, F1-Score: 87%, Accuracy: 84%

Pada rasio 90:10, baik model Multinomial Naive Bayes (MNB) maupun Support Vector Machine (SVM) menunjukkan performa yang sangat baik. Namun, SVM menunjukkan keunggulan dengan nilai precision, recall, dan F1-Score yang sedikit lebih tinggi dibandingkan MNB, meskipun kedua model memiliki nilai accuracy yang sama yaitu 87%. Pada rasio 80:20, SVM tetap unggul dengan precision yang lebih tinggi, tetapi recall mengalami penurunan. Di sisi lain, MNB mengalami penurunan kinerja di semua metrik pada rasio ini. Pada rasio 70:30 dan 60:40, SVM menunjukkan performa yang lebih baik secara keseluruhan, terutama dalam precision, sementara MNB tetap stabil dengan nilai yang sedikit lebih rendah. Pada rasio 50:50, kinerja kedua algoritma mendekati satu sama lain, namun SVM tetap unggul dalam precision dan F1-Score.

5. KESIMPULAN DAN SARAN

Analisis kinerja algoritma Multinomial Naive Bayes (MNB) dan Support Vector Machine (SVM) menunjukkan bahwa SVM unggul dalam hal precision dan memberikan performa yang lebih baik secara keseluruhan. SVM menunjukkan precision tertinggi pada rasio 70:30 dengan nilai 91%, meskipun recall

menurun dari 89% pada rasio 90:10 menjadi 84% pada rasio 50:50. F1-Score dan accuracy SVM juga konsisten tinggi, dengan accuracy mencapai 87% pada rasio 90:10. Sebaliknya, MNB menunjukkan stabilitas yang baik dengan precision antara 84% hingga 87% dan accuracy tetap pada 84%, kecuali pada rasio 90:10 yang mencapai 87%. Meskipun SVM merupakan algoritma yang lebih unggul dalam analisis ini, eksplorasi algoritma machine learning lain seperti Random Forest, Gradient Boosting, atau Deep Learning bisa bermanfaat untuk menemukan model yang mungkin menawarkan performa lebih baik. Selain itu, memperluas sumber data dengan menambahkan komentar dari platform lain seperti Twitter, Facebook, atau forum otomotif dapat memberikan perspektif yang lebih luas dan memperkaya dataset, yang pada akhirnya dapat meningkatkan kualitas analisis sentimen.

DAFTAR PUSTAKA

- [1] L. Nardo and B. Prasetyo, "PENGARUH KUALITAS PRODUK DAN PERSEPSI HARGA TERHADAP KEPUTUSAN PEMBELIAN SEPEDA MOTOR HONDA BEAT PADA DEALER CV. SUPRA JAYA MOTOR CIANJUR", *Transekonomika*, vol. 2, no. 5, pp. 433-448, Jul. 2022
- [2] N. Anisah, M. Sartika, and H. Kurniawan, "Penggunaan Media Sosial Instagram dalam Meningkatkan Literasi Kesehatan pada Mahasiswa Universitas Syiah Kuala," *Jurnal Peurawi: Media Kajian Komunikasi Islam*, vol. 04, no. 02, pp. 94-107, 2021. E-ISSN: 2598-6031, ISSN: 2598-6032.
- [3] M. F. Asshiddiqi and K. M. Lhaksana, "Perbandingan Metode Decision Tree dan Support Vector Machine untuk Analisis Sentimen pada Instagram Mengenai Kinerja PSSI," *e-Proceeding of Engineering*, vol. 7, no. 3, pp. 9936-9948, Dec. 2020. ISSN: 2355-9365.
- [4] Yuyun, Nurul Hidayah, and Supriadi Sahibu, "Algoritma Multinomial Naïve Bayes Untuk Klasifikasi Sentimen Pemerintah Terhadap Penanganan Covid-19 Menggunakan Data Twitter", *J. RESTI (Rekayasa Sist. Teknol. Inf.)*, vol. 5, no. 4, pp. 820 - 826, Aug. 2021
- [5] A. Zy and Wahyu Hadikristanto, "Implementasi Algoritma Metode Naive Bayes dan Support Vector Machine Tentang Pembobolan dan Kebocoran Data di Twitter", *bit*, vol. 4, no. 1, pp. 49 - 56, Mar. 2023.
- [6] F. Zamachsari, Gabriel Vangeran Saragih, Susafa'ati, and Windu Gata, "Analysis of Sentiment of Moving a National Capital with Feature Selection Naive Bayes Algorithm and Support Vector Machine", *J. RESTI (Rekayasa Sist. Teknol. Inf.)*, vol. 4, no. 3, pp. 504 - 512, Jun. 2020.
- [7] P. Arsi, R. Wahyudi, and R. Waluyo, "Optimasi SVM Berbasis PSO pada Analisis Sentimen

- Wacana Pindah Ibu Kota Indonesia”, *J. RESTI (Rekayasa Sist. Teknol. Inf.)*, vol. 5, no. 2, pp. 231–237, Apr. 2021.
- [8] R. S. A. Al-Zaelani, Y. R. Ramadhan, and M. A. Komara, "Analisis Sentimen Review Produk Motor Honda PCX dan Yamaha N-Max pada Twitter Menggunakan Algoritma Naïve Bayes," *JATI (Jurnal Mahasiswa Teknik Informatika)*, vol. 7, no. 3, pp. 1812, Jun. 2023.
- [9] M. M. Fajri and I. M. Karo Karo, "Implementasi Algoritma Support Vector Machine Untuk Analisis Sentimen Aplikasi Easycash di Playstore", *Scientica*, vol. 1, no. 3, pp. 145–152, Dec. 2023
- [10] D. Oktavia, Y. R. Ramadhan, and M. Minarto, "Analisis Sentimen Terhadap Penerapan Sistem E-Tilang Pada Media Sosial Twitter Menggunakan Algoritma Support Vector Machine (SVM)," *KLIK: Kajian Ilmiah Informatika dan Komputer*, vol. 4, no. 1, pp. 407–417, Aug. 2023, doi: 10.30865/klik.v4i1.1040.
- [11] F. Baehaqi and N. Cahyono, "Analisis Sentimen Terhadap Cyberbullying Pada Komentar di Instagram Menggunakan Algoritma Naïve Bayes," **Indonesian Journal of Computer Science**, vol. 13, no. 1, pp. 1051–1063, Jan. 2024.
- [12] D. Abimanyu, E. Budianita, E. P. Cynthia, F. Yanto, and Yusra, "Analisis Sentimen Akun Twitter Apex Legends Menggunakan VADER," **Jurnal Nasional Komputasi dan Teknologi Informasi**, vol. 5, no. 3, pp. 423–431, Jun. 2022.
- [13] A. D. Pangestu, I. E. S.Kom., M.Si., and N. Chamidah S.Kom, M.Kom., "Analisis Sentimen Terhadap PPKM Darurat Pada Media Sosial Twitter Menggunakan Metode Naïve Bayes Dengan Seleksi Fitur Information Gain," in *Seminar Nasional Mahasiswa Ilmu Komputer dan Aplikasinya (SENAMIKA)*, Jakarta, Indonesia, 20 Aug. 2022, pp. 662.
- [14] J. A. Septian, T. M. Fachrudin, and A. Nugroho, "Analisis Sentimen Pengguna Twitter Terhadap Polemik Persepakbolaan Indonesia Menggunakan Pembobotan TF-IDF dan K-Nearest Neighbor", *INSYST*, vol. 1, no. 1, pp. 43–49, Aug. 2019.
- [15] I. H. Kusuma and N. Cahyono, "Analisis Sentimen Masyarakat Terhadap Penggunaan E-Commerce Menggunakan Algoritma K-Nearest Neighbor," *Jurnal Informatika: Jurnal Pengembangan IT (JPIT)*, vol. 8, no. 3, pp. 302–307, Sep. 2023.