

Analisis Sentimen Suporter terhadap Performa Persija Jepara Menggunakan Metode Naïve Bayes

DOI: <http://dx.doi.org/10.35889/jutisi.v15i4.3883>

Creative Commons License 4.0 (CC BY – NC)

M. Rikza Tamyiz^{1*}, Esti Wijayanti², Ahmad Jazuli³

Teknik Informatika, Universitas Muria Kudus, Kudus, Indonesia

*Email Corresponding Author: 202251023@std.umk.ac.id

Abstract

Social media has given football fans an open channel to voice unfiltered reactions to their club's performance, and the Instagram account of Persija Jepara now attracts a large volume of such comments that has never been examined in a structured way. This study applies a Machine Learning technique, specifically the Multinomial Naïve Bayes algorithm, to classify the sentiment expressed by these supporters. A total of 5,493 distinct Instagram comments were collected, cleaned through a text-preprocessing pipeline, and converted into numerical features using TF-IDF weighting. Testing the resulting model produced an accuracy of 74.32%, alongside a precision score of 0.76, a recall score of 0.72, and an F1-score of 0.74.

Keywords: Sentiment Analysis; Machine Learning; Naïve Bayes; TF-IDF; Persija Jepara.

Abstrak

Kemajuan media sosial membuka ruang bagi suporter sepak bola untuk menuliskan pendapat mereka secara bebas mengenai performa klub. Akun Instagram resmi Persija Jepara menerima volume komentar yang tinggi, tetapi data teks tersebut belum pernah diolah dan dinilai secara terstruktur. Riset ini berupaya mengukur kecenderungan sentimen suporter dengan memanfaatkan teknik Machine Learning berupa algoritma Multinomial Naïve Bayes. Sebanyak 5.493 komentar unik dikumpulkan dari Instagram, kemudian dibersihkan melalui tahap preprocessing teks serta diubah menjadi fitur numerik menggunakan pembobotan TF-IDF. Hasil pengujian model menunjukkan akurasi 74,32%, presisi 0,76, recall 0,72, serta F1-score 0,74.

Kata kunci: Analisis Sentimen; Machine Learning; Naïve Bayes; TF-IDF; Persija Jepara

1. Pendahuluan

Laju kemajuan teknologi informasi dan komunikasi telah mengubah cara masyarakat menyuarakan pandangan sekaligus mencari informasi. Ruang publik kini banyak berpindah ke media sosial, tempat masyarakat berinteraksi, berbagi pengalaman, dan menanggapi berbagai persoalan, termasuk isu-isu olahraga. Kegiatan tersebut memunculkan tumpukan data teks yang terus bertambah setiap harinya, menjadikannya sumber informasi berharga untuk membaca opini publik. Menelaah opini yang tersebar di media sosial menjadi hal yang penting sebab hasilnya dapat menggambarkan respons, kepuasan, maupun kepercayaan masyarakat terhadap suatu hal secara lebih cepat dan objektif dibanding metode konvensional. Atas dasar itu, analisis sentimen yang ditopang *Machine Learning* menjadi pilihan yang tepat untuk mengolah opini secara otomatis, sehingga hasilnya dapat dipakai sebagai landasan pengambilan keputusan yang lebih tepat.

Sebagai klub yang menjadi kebanggaan warga Kabupaten Jepara, Persija Jepara memiliki basis suporter yang cukup aktif menanggapi performa tim lewat kanal media sosial, terutama Instagram. Tanggapan tersebut muncul dalam bentuk komentar pada unggahan resmi klub yang mengulas hasil laga, performa pemain, taktik pelatih, sampai kebijakan manajemen. Derasnya interaksi ini menghasilkan tumpukan data teks yang sebenarnya merekam persepsi suporter terhadap perkembangan tim, namun data tersebut belum digarap secara sistematis.

Selama ini penilaian performa tim cenderung hanya mengandalkan statistik pertandingan seperti jumlah gol, penguasaan bola, dan posisi di klasemen, sedangkan persepsi dan sentimen suporter sebagai salah satu *stakeholder* penting klub belum pernah diukur secara kuantitatif. Akibatnya, pihak manajemen belum memiliki gambaran terstruktur tentang kecenderungan opini suporter yang bisa dijadikan bahan pertimbangan dalam evaluasi dan pengambilan keputusan.

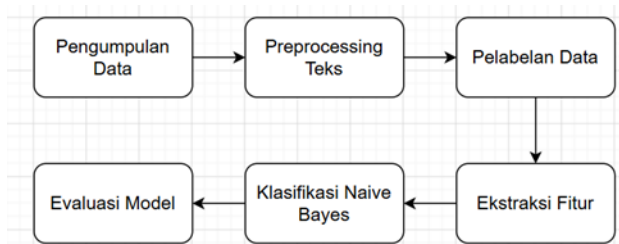
Sejumlah studi sebelumnya telah memanfaatkan analisis sentimen untuk membaca persepsi masyarakat terhadap berbagai peristiwa olahraga di media sosial. Mulya dkk. [1], misalnya, mengkaji sentimen publik seputar tren olahraga sepanjang masa pandemi COVID-19 dengan memanfaatkan data Twitter dan algoritma Naïve Bayes Classifier; alur penelitiannya diawali dari penelusuran topik olahraga yang tengah tren lewat Google Trends, dilanjutkan dengan crawling data Twitter, tahap preprocessing (case folding, cleansing, normalisasi, stopword removal, stemming, convert negation, dan tokenizing), pelabelan data, klasifikasi Naïve Bayes, hingga evaluasi memakai confusion matrix. Studi lain oleh Kholilullah dkk. [2] mengulas sentimen pengguna Twitter (X) terkait penyelenggaraan Piala Dunia U-17 di Indonesia, juga dengan algoritma Naïve Bayes, melalui rangkaian tahapan pengumpulan data via web scraping, preprocessing, pelabelan sentimen, pembobotan TF-IDF, klasifikasi Naïve Bayes, dan evaluasi performa model.

Kendati penelitian-penelitian tersebut terbukti berhasil menerapkan Naïve Bayes untuk analisis sentimen di ranah olahraga, fokus objeknya masih berkutat pada turnamen berskala nasional maupun internasional dengan sumber data Twitter (X). Sampai sejauh ini, kajian yang secara spesifik menyoroti sentimen suporter terhadap performa klub sepak bola daerah — khususnya Persija Jepara — dengan memanfaatkan komentar Instagram masih jarang ditemukan. Penelitian ini hadir untuk menutup celah tersebut dengan menerapkan algoritma Naïve Bayes guna mengklasifikasikan sentimen komentar suporter Persija Jepara di Instagram, sehingga hasilnya dapat menjadi sumber informasi pendukung evaluasi performa tim dan pengambilan keputusan oleh manajemen klub.

Guna menjawab kesenjangan tersebut, penelitian ini mengusulkan pemanfaatan algoritma Naïve Bayes untuk mengelompokkan sentimen komentar suporter Persija Jepara di Instagram ke dalam tiga kategori: positif, negatif, dan netral. Pilihan jatuh pada Naïve Bayes karena algoritma ini unggul dalam mengolah data teks, tetap efisien meski menangani data berukuran besar [3], dan telah terbukti memberikan performa klasifikasi yang baik pada berbagai riset analisis sentimen, tidak terkecuali pada topik olahraga. Rangkaian tahapan analisis — mulai dari *web scraping*, *preprocessing*, pembobotan TF-IDF, hingga klasifikasi Naïve Bayes — dinilai mampu menghasilkan model yang akurat dalam mengenali kecenderungan opini suporter. Kebaruan riset ini terletak pada penggunaan komentar Instagram, bukan Twitter (X) [4] seperti kebanyakan penelitian terdahulu, sebagai sumber data untuk menganalisis sentimen suporter Persija Jepara — topik yang selama ini jarang disentuh karena penelitian sejenis umumnya berfokus pada turnamen sepak bola berskala nasional atau internasional. Diharapkan, temuan penelitian ini tidak sekadar memotret persepsi suporter terhadap performa tim, tetapi juga menjadi masukan berbasis data bagi manajemen Persija Jepara dalam mengevaluasi tim dan mengambil keputusan.

2. Metodologi

Metode yang digunakan dalam penelitian ini mencakup serangkaian tahapan sistematis, mulai dari pengolahan data teks media sosial hingga terbentuknya model klasifikasi yang dapat mengenali sentimen suporter Persija Jepara.



Gambar 1. Tahapan Penelitian

2.1. Pengumpulan Data

Penelitian ini memakai data primer berupa opini atau komentar suporter yang diambil langsung dari akun media sosial resmi Persija Jepara (Instagram). Proses scraping dilakukan memakai ekstensi bernama Instant Data Scraper. Cakupan pengambilan data meliputi komentar pada unggahan hasil pertandingan Persija, baik saat menang, kalah, maupun seri. Dari proses ini terkumpul 5.943 data mentah, yang setelah melalui pembersihan menyusut menjadi 2.500 data. Data tersebut selanjutnya dikelompokkan ke dalam kelas sentimen melalui pelabelan, dengan ketentuan Positif=1, Negatif=-1, dan Netral=0

2.2. Preprocessing Teks

Tahap ini dirancang untuk merapikan data teks mentah agar lebih terstruktur [5] *noise* — misalnya emoji dan simbol — dihapus, sementara kata berimbuhan seperti 'permainan' dikembalikan ke bentuk dasarnya, 'main', menggunakan *library Sastrawi*. Pra-pemrosesan menjadi langkah awal yang wajib dilalui agar data siap diolah lebih lanjut [6] serta siap dibaca oleh algoritma. Adapun tahapannya mencakup:

- 1) *Case Folding*
Menyeragamkan seluruh huruf pada teks menjadi huruf kecil (*lowercase*).
- 2) *Cleaning*
Membuang unsur tidak relevan seperti URL, username, tanda baca, angka, dan simbol khusus.
- 3) *Tokenizing*
Memisahkan kalimat menjadi satuan kata atau token tersendiri.
- 4) *Stopword Removal*
Membuang kata-kata umum yang kerap muncul namun tidak membawa makna sentimen berarti (misalnya: "dan", "yang", "di")
- 5) *Stemming*
Mengembalikan kata berimbuhan ke bentuk dasarnya guna menekan variasi kata.

2.3. Pelabelan Data

Proses pelabelan dilakukan secara manual menggunakan Microsoft Excel setiap komentar diberi skor sentimen; skor > 1 dikategorikan positif, skor < -1 dikategorikan negatif, skor < 0 dikategorikan netral.

2.4. Ekstraksi Fitur

Tahap ini mengubah teks yang telah melalui *preprocessing* menjadi representasi numerik agar dapat dihitung secara matematis oleh komputer. Pembobotan kata dilakukan menggunakan metode TF-IDF (*Term Frequency - Inverse Document Frequency*), yang menilai setiap kata berdasarkan seberapa sering kata tersebut muncul pada satu dokumen dibanding pada keseluruhan dataset.

2.5. Klasifikasi Naïve Bayes

Pelatihan model dilakukan dengan algoritma Naïve Bayes, yang menghitung peluang setiap kata terhadap masing-masing kategori sentimen (positif, negatif, atau netral) berdasarkan teorema Bayes. Algoritma ini dipilih karena hemat memori dan cepat dalam mengolah data teks berskala besar.

Efektivitas Naïve Bayes pada klasifikasi teks juga pernah dibuktikan oleh penelitian terdahulu yang memakai seleksi fitur *Relevance Frequency* guna mempertajam akurasi klasifikasi sentimen *publik* [5]. Walau tergolong metode klasifikasi konvensional, kesederhanaan perhitungannya menjadikan algoritma ini pilihan ideal sebagai model dasar (*baseline*) sebelum beranjak ke arsitektur yang lebih rumit.

$$P(C|X) = \sum [P(X|C) \times P(C)] / P(X) \quad (1)$$

C pada rumus tersebut menyatakan kelas sentimen, sedangkan X menyatakan fitur teks. Kelas sentimen akhir ditetapkan berdasarkan nilai probabilitas tertinggi yang dihasilkan algoritma ini. Algoritma ini tetap unggul dari sisi efisiensi memori dan kecepatan pemrosesan data teks berskala besar.

2.6. Evaluasi Model

Kinerja model dalam mengklasifikasikan sentimen diukur menggunakan beberapa metrik evaluasi berikut:

- 1) *Confusion Matrix*: Membandingkan hasil prediksi model dengan data aktual.
- 2) *Accuracy, Precision, Recall, and F1-Score*: Mengukur tingkat ketepatan klasifikasi pada tiap kelas sentimen.

3. Hasil dan Pembahasan

3.1. Pengumpulan Dataset

Data penelitian dihimpun melalui proses *scraping* pada kolom komentar akun Instagram resmi Persija Jepara. Eksekusi proses ini menghasilkan 3 berkas CSV terpisah yang kemudian digabungkan (*merging*). Cakupannya meliputi komentar pada unggahan hasil pertandingan Persija, baik saat menang, kalah, maupun seri. Total 5.943 data mentah berhasil dihimpun, yang setelah dibersihkan tersisa 2.500 data. Data ini kemudian dikelompokkan ke kelas sentimen melalui pelabelan, dengan ketentuan Positif=1, Negatif=-1, dan Netral=0

- 1) Dataset Mentah: sebelum dibersihkan, data yang terkumpul (5.943) masih banyak mengandung duplikasi.
- 2) Dataset Olah: setelah pembersihan, tersisa 2.500 data siap pakai.

```
-----
BERHASIL! Total data unik: 5493 baris.
File baru tersimpan: dataset_gabungan_siap_label.csv
```

Gambar 2. Jumlah Seluruh Data

3.2. Tahapan *Preprocessing* Data

Guna mendapatkan akurasi optimal, sebanyak 5.493 baris dataset mentah dijalankan melalui fungsi *Preprocessing* lengkap. Perubahan data sebelum dan sesudah proses pembersihan disajikan pada Tabel 1:

Tabel 1. Hasil *Preprocessing* Data

Proses	Sebelum	Sesudah
Cleaning	Pelatihnya minim taktik, ganti saja! #PersijapDay	Pelatihnya minim taktik ganti saja
Case Folding	Pelatihnya minim taktik, ganti saja! #PersijapDay	pelatihnya minim taktik, ganti saja
Tokenizing	pelatihnya minim taktik, ganti saja! #PersijapDay	(pelatih, minim, taktik, ganti, saja,)
Stopword	(pelatih, minim, taktik, ganti, saja,)	pelatih minim taktik ganti saja
Stemming	Pelatihnya minim taktik, ganti saja! #PersijapDay	pelatih minim taktik, ganti saja!

3.3. Analisis Frekuensi Kata

Tabel 2. Frekuensi Kata

Peringkat	Kata	Frekuensi
1.	jap	439
2.	Alhamdulillah	205
3.	main	168
4.	lki	146
5.	wes	115
6.	latih	105
7.	menang	94
8.	persijap	93
9.	lawan	85
10.	liga	82

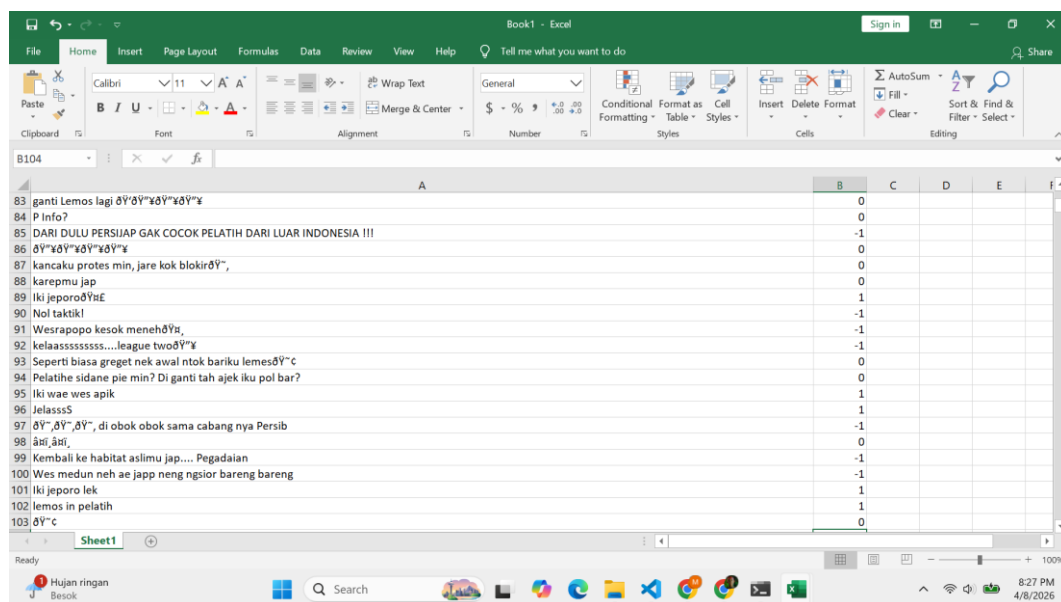
Hasil analisis frekuensi kata memperlihatkan bahwa kata "jap" paling banyak muncul, yakni sebanyak 439 kali, karena kerap dipakai suporter sebagai singkatan dari Persija. Kata "alhamdulillah" tercatat muncul 205 kali, mencerminkan banyaknya komentar bernuansa syukur atau apresiasi atas hasil laga. Sementara itu, kemunculan kata "main" (168 kali) dan "latih" (105 kali) menandakan bahwa perbincangan suporter banyak menyoroti performa pemain serta strategi pelatih.

Kata-kata lain seperti "menang" (94 kali), "persijap" (93 kali), "lawan" (85 kali), dan "liga" (82 kali) turut mengindikasikan bahwa sebagian besar komentar memang berkisar pada hasil pertandingan dan perjalanan Persija Jepara di kompetisi liga.

3.4. Distribusi Sentimen

Setelah tahap preprocessing rampung, tiap komentar diberi label sentimen untuk menentukan kategori komen pada setiap postingan. Pelabelan diawali dengan menilai apakah setiap komen cenderung bernada positif, negatif, atau netral terhadap pertandingan Persija.

Mengingat volume data yang besar, proses pelabelan tetap dikerjakan secara manual menggunakan Microsoft Excel setiap komentar diberi skor sentimen; skor > 1 dikategorikan positif, skor < -1 dikategorikan negatif, skor < 0 dikategorikan netral.



Gambar 3. Hasil Labelling Microsoft Excel

3.5. Hasil Pembobotan (TF-IDF)

Seluruh komentar yang telah melewati tahapan *preprocessing* selanjutnya diubah menjadi representasi numerik lewat metode **Term Frequency-Inverse Document Frequency (TF-IDF)**, dengan proses pembobotan dijalankan memakai kelas **TfidfVectorizer** dari pustaka *Scikit-learn*. Beberapa kata umum yang dinilai minim kontribusi informasi terhadap klasifikasi, seperti *info*, *pemain*, *baru*, *ada*, *yang*, *dan*, *di*, dan *dari*, dimasukkan ke parameter **stop_words** sehingga tidak ikut dijadikan fitur. Pembobotan TF-IDF baru dijalankan setelah dataset dipecah menjadi **80% data latih** dan **20% data uji** lewat fungsi **train_test_split** dengan **random_state = 42**. Fungsi **fit_transform()** kemudian diterapkan pada data latih untuk membangun kosakata (*vocabulary*) sekaligus bobot TF-IDF-nya, sementara data uji cukup diproses lewat fungsi **transform()** agar tetap mengacu pada kosakata yang sama dengan data latih. Melalui proses ini, setiap komentar berubah menjadi satu vektor numerik, dengan tiap baris mewakili satu komentar dan tiap kolom mewakili satu kata unik (*feature*) yang diperoleh dari data latih.

Tabel 3. Hasil Pembobotan TF-IDF pada Komentar

Dokumen	jap	menang	liga	latih	degradasi
D1	0.521	0.314	0	0	0
D2	0	0	0.463	0.392	0
D3	0	0.488	0	0	0
D4	0.317	0	0	0	0.541
D5	0	0	0.285	0	0

Angka-angka pada tabel tersebut menggambarkan bobot TF-IDF setiap kata pada masing-masing komentar.

Sebagai ilustrasi, fitur-fitur yang dihasilkan *TfidfVectorizer* berupa kumpulan kata unik hasil preprocessing, seperti dicontohkan berikut.

Tabel 4. Fitur (Kata Unik) Hasil *TfidfVectorizer*

No	Fitur
1	alhamdulillah
2	apik
3	degradasi
4	jap
5	jeporo
6	kalah
7	latih
8	liga
9	main
10	menang

Melalui proses pembentukan fitur berbasis TF-IDF, seluruh komentar berhasil diubah menjadi bentuk numerik dan siap diproses oleh algoritma **Multinomial Naïve Bayes**. TF-IDF sendiri memberi bobot berbeda pada tiap kata sesuai perannya dalam membedakan satu dokumen dari yang lain — kata yang sering muncul pada komentar tertentu namun jarang ditemui pada komentar lain akan diberi bobot lebih besar.

Pada penelitian ini, fitur yang terbentuk didominasi oleh istilah-istilah seputar pertandingan Persija Jepara, seperti **jap**, **persijap**, **menang**, **liga**, **latih**, **kalah**, dan **degradasi**— menandakan bahwa perbincangan suporter memang lebih banyak menyentuh performa tim, strategi pelatih, hasil laga, dan posisi klub di kompetisi. Keseluruhan fitur hasil pembobotan TF-IDF ini selanjutnya dipakai sebagai input (*input feature*) pada tahap pelatihan dan pengujian model klasifikasi sentimen berbasis **Multinomial Naïve Bayes**.

3.6. Implementasi Naïve Bayes

Pembangunan model klasifikasi mengandalkan algoritma *Multinomial Naïve Bayes*, dengan pembagian dataset sebagai berikut:

- 80% data latih (4.394 data), dipakai untuk melatih model mengenali pola kata.
- 20% data uji (1.099 data), dipakai untuk menguji keakuratan prediksi model.

Model yang sudah dilatih kemudian diuji dengan data uji untuk melihat performa klasifikasinya, dengan evaluasi memakai metrik Accuracy, Precision, Recall, dan F1-Score yang diambil dari fungsi `classification_report()`. Hasil pengujian menunjukkan akurasi sebesar 74,32%, yang berarti model mampu mengklasifikasikan secara tepat sekitar tiga dari empat data uji.

AKURASI SETELAH STEMMING: 74.32%

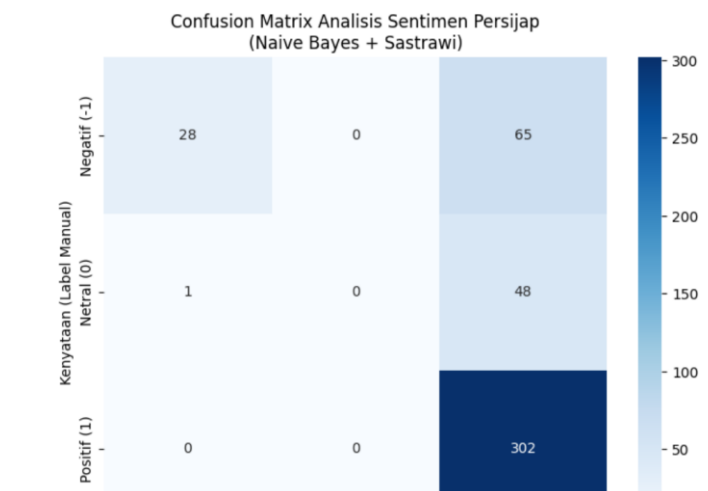
Laporan Klasifikasi:

	precision	recall	f1-score	support
-1.0	0.97	0.30	0.46	93
0.0	0.00	0.00	0.00	49
1.0	0.73	1.00	0.84	302
accuracy			0.74	444
macro avg	0.56	0.43	0.43	444
weighted avg	0.70	0.74	0.67	444

Gambar 4. Laporan Klasifikasi (Classification Report) Model Multinomial Naïve Bayes

3.7. Evaluasi Kinerja Model (*Confusion Matrix*)

Setelah tahap klasifikasi dengan Multinomial Naïve Bayes rampung, langkah berikutnya adalah menilai performa model melalui Confusion Matrix, yang berguna untuk melihat berapa banyak prediksi yang tepat maupun keliru pada tiap kelas sentimen — negatif, netral, dan positif. Penilaian ini dilakukan terhadap hasil prediksi model atas data uji, lalu divisualisasikan dalam bentuk heatmap agar distribusi hasil klasifikasi pada setiap kelas lebih mudah diamati.



Gambar 5. Confusion Matrix

Ketimpangan Data: dari 5.493 baris dataset mentah, sentimen positif mendominasi secara signifikan, sehingga model menjadi lebih terampil mengenali komentar dukungan dibanding kritik. Persoalan ketimpangan data (*imbalanced data*) semacam ini sebenarnya bisa diatasi dengan teknik khusus, sebagaimana dicontohkan oleh [7], yang memadukan algoritma *Long Short-Term Memory* dengan *Synthetic Minority Over-sampling Technique* (SMOTE) guna menyeimbangkan distribusi kelas.

3.8. Analisis Hasil Prediksi

Model juga diujicobakan untuk memprediksi data baru yang belum pernah muncul dalam dataset. Hasil uji fungsi prediksi sentimen tersebut ditampilkan pada Gambar 6.

```

=== HASIL PENGUJIAN 3 SENTIMEN INPUT BARU ===
1. Komentar : 'Persijap main bagus banget hari ini, bangga sekali menang!'
   Hasil AI : POSITIF

2. Komentar : 'Kecewa sekali dengan permainan hari ini kalah terus, latihannya gimana sih'
   Hasil AI : POSITIF

3. Komentar : 'Info jadwal pertandingan persijap berikutnya tanggal berapa min?'
   Hasil AI : POSITIF

```

Gambar 6. Hasil Prediksi Data Baru

3.9. Pembahasan

Riset ini ditujukan untuk mengklasifikasikan sentimen suporter Persija Jepara berdasarkan komentar pada akun Instagram resminya, ke dalam tiga kategori — positif, negatif, netral — menggunakan algoritma Multinomial Naïve Bayes dengan fitur hasil ekstraksi Term Frequency–Inverse Document Frequency (TF-IDF). Model yang diuji memperoleh akurasi 74,32%, menandakan bahwa mayoritas komentar berhasil diklasifikasikan sesuai label sebenarnya.

Laporan klasifikasi memperlihatkan performa terbaik model justru muncul pada kelas positif, dengan precision 0,73, recall 1,00, dan F1-score 0,84 — artinya hampir seluruh komentar positif berhasil dikenali dengan tepat. Sebaliknya, model masih kesulitan mengenali komentar negatif, apalagi netral. Kondisi ini tampak dari recall kelas negatif yang hanya 0,30, serta precision, recall, dan F1-score kelas netral yang seluruhnya 0,00. Temuan ini mengonfirmasi kecenderungan model untuk memprediksi komentar ke kelas dengan jumlah data paling dominan [8].

Dengan capaian tersebut, tujuan riset untuk membangun model klasifikasi sentimen komentar suporter Persija Jepara dapat dikatakan tercapai. Model sudah mampu menangkap kecenderungan sentimen suporter secara umum, meski performanya pada kelas minoritas masih perlu ditingkatkan.

Dibandingkan sejumlah riset terdahulu yang juga memakai Naïve Bayes, hasil penelitian ini tergolong lebih unggul — sebagai perbandingan, riset A. Sholekhah hanya memperoleh akurasi 66,34% [9], sehingga menegaskan bahwa model yang dibangun pada penelitian ini mampu bekerja lebih baik pada dataset yang digunakan.

Temuan ini juga selaras dengan prinsip dasar Multinomial Naïve Bayes. Algoritma ini bekerja berdasarkan peluang kemunculan kata pada tiap kelas dan dikenal efektif untuk mengklasifikasikan dokumen teks yang telah direpresentasikan lewat TF-IDF atau representasi berbasis frekuensi kata [10].

Di sisi lain, pemakaian TF-IDF sebagai metode pembobotan fitur turut membantu model membedakan kata yang bernilai informasi tinggi dari kata yang sekadar sering muncul tanpa banyak makna.

TF-IDF memberi bobot lebih tinggi pada istilah yang bersifat diskriminatif, sehingga kualitas fitur yang dipakai dalam klasifikasi teks pun meningkat [11]. Dalam penelitian ini, kata-kata terkait performa tim, hasil pertandingan, dan pelatih menjadi fitur utama yang ikut menentukan hasil klasifikasi sentimen.

Temuan ini memperkuat sejumlah riset sebelumnya yang menyimpulkan bahwa kombinasi TF-IDF dan Naïve Bayes cukup efektif untuk analisis sentimen di media sosial. Naufal dkk. [12] menganalisis sentimen tim nasional Indonesia pada Piala Asia 2023 di media sosial X dengan pembobotan TF-IDF dan Multinomial Naïve Bayes, dan memperoleh akurasi berkisar 79%–81% pada berbagai kategori sentimen (kekalahan, kemenangan, dan keseluruhan). Capaian tersebut sejalan dengan hasil penelitian ini, yang juga berhasil mengklasifikasikan komentar suporter Persija Jepara dengan akurasi yang cukup memadai.

Sejalan dengan itu, Khairunnisa dkk. [13] dalam kajian sentimen media sosial Twitter pada masa pandemi COVID-19 melaporkan bahwa kualitas serta kombinasi tahapan text preprocessing sangat berpengaruh terhadap performa model klasifikasi. Hal serupa juga tampak pada penelitian ini, di mana tahapan case folding, cleaning, tokenizing, stopword removal, dan stemming menghasilkan data yang lebih bersih sehingga pembentukan fitur TF-IDF menjadi lebih optimal.

Penelitian ini turut memberi kontribusi pada bidang text mining dan analisis sentimen olahraga, khususnya dalam mengkaji opini suporter terhadap performa klub sepak bola lokal. Jika sebagian besar penelitian sejenis memakai objek berupa ulasan produk atau layanan digital, penelitian ini justru menjadikan komentar media sosial sebagai sumber data untuk mengukur persepsi suporter terhadap Persija Jepara.

Karena model konvensional seperti Multinomial Naïve Bayes masih terbatas dalam mengenali kelas minoritas maupun pola bahasa rumit seperti sarkasme, beberapa pendekatan berbasis pembelajaran mendalam berikut layak dipertimbangkan sebagai arah pengembangan selanjutnya. Salah satu opsi adalah beralih ke jaringan saraf tiruan mendalam atau Deep Learning [14]. Arsitektur Long Short-Term Memory (LSTM), misalnya, cukup populer karena mampu mengingat konteks jangka panjang dalam teks dan sering dipakai pula untuk peramalan data deret waktu secara akurat [15]. Untuk hasil yang lebih tinggi, arsitektur multi-class

classification berbasis Bidirectional LSTM (Bi-LSTM) juga terbukti tangguh menangani teks tidak terstruktur [16]. Selain rumpun LSTM, model Transformer berbasis IndoBERT tengah menjadi tren karena pemahamannya terhadap semantik bahasa Indonesia jauh lebih mendalam, cocok untuk analisis sentimen berbasis aspek (Aspect-Based Sentiment Analysis) [17]. Algoritma generator indeks semantik berbasis BERT pun terbukti mampu mendongkrak akurasi ekstraksi fitur secara nyata [18] serta dapat diperkuat lagi lewat kombinasi Bidirectional Autoencoder. Terakhir, untuk mengatasi kendala pelabelan manual dalam jumlah besar yang rawan bias subjektif, penelitian mendatang dapat mempertimbangkan metode pelabelan otomatis (Auto Labeling) berbasis K-Nearest Neighbors (KNN) guna mempercepat proses anotasi dataset [19].

4. Simpulan

Berdasarkan analisis sentimen terhadap 5.493 data opini suporter Persija Jepara memakai metode Multinomial *Naïve Bayes*, dapat disimpulkan bahwa algoritma ini cukup efektif memetakan persepsi publik digital, dengan akurasi 74,32%. Opini suporter didominasi oleh Sentimen Positif (68%) yang mencerminkan dukungan moral, disusul Sentimen Negatif (18%) berupa kritik performa, dan Sentimen Netral (14%) yang umumnya berisi informasi logistik. Capaian akurasi ini tidak lepas dari keberhasilan tahap *Preprocessing (cleaning, case folding, dan pembobotan TF-IDF)* dalam mereduksi noise teks, meski model masih terbatas dalam mengenali struktur kalimat sarkasme dan kosakata kritik yang spesifik akibat ketimpangan data (imbalanced data).

Penelitian ini masih menyisakan sejumlah keterbatasan. Pertama, data hanya diambil dari komentar pada akun Instagram resmi Persija Jepara, sehingga belum mewakili keseluruhan opini suporter yang sebetulnya juga aktif di platform lain seperti Facebook, X (Twitter), atau TikTok. Kedua, jumlah data pada tiap kelas sentimen belum seimbang, yang turut memengaruhi kemampuan model dalam mengenali kelas negatif dan netral. Ketiga, penelitian ini baru menerapkan satu algoritma, yaitu Multinomial *Naïve Bayes*, sehingga belum ada perbandingan performa dengan algoritma lain seperti *Support Vector Machine (SVM)*, *Random Forest*, atau pendekatan berbasis Deep Learning.

Untuk penelitian ke depan, sumber data dapat diperluas ke berbagai platform media sosial lain, disertai penerapan teknik penyeimbangan data seperti SMOTE atau Random Oversampling, serta eksplorasi pendekatan *Deep Learning* dan model bahasa terkini agar model lebih akurat dalam mengenali kelas minoritas dan memahami konteks sarkasme.

Daftar Referensi

- [1] B. C. Nbc, S. Mulya, and H. Sujaini, "Analisis Sentimen Tren Olahraga di Masa Pandemi COVID-19 pada Twitter dengan Metode Naïve," *JEPIN (Jurnal Edukasi dan Penelit. Inform.*, vol. 8, no. 2, pp. 284–291, 2022.
- [2] M. Kholilullah, U. Hayati, T. Informatika, M. Informatika, and K. Cirebon, "Analisis Sentimen Pengguna Twitter (X) Tentang Piala Dunia Usia 17 Menggunakan Metode Naive Bayes," *JATI (Jurnal Mhs. Tek. Inform.*, vol. 8, no. 1, pp. 392–398, 2024.
- [3] H. P. Jelita and M. I. Saad, "Penerapan Algoritma Naïve Bayes Dalam Analisis Sentimen Masyarakat Terhadap STMIK Widya Cipta Dharma," *Bull. Inf. Technol.*, vol. 6, no. 2, pp. 148–160, 2025, doi: 10.47065/bit.v5i2.2029.
- [4] R. H. Nufus and U. Surapati, "Analisis Sentimen Persepsi Masyarakat Terhadap Timnas Indonesia U-23 dalam AFC-23 Asian Cup 2024 Pada Media Sosial X Menggunakan Metode Naïve Bayes Classifier," *J. Indones. Manaj. Inform. dan Komun.*, vol. 5, no. 3, pp. 2647–2657, 2024.
- [5] K. Verena, S. Toy, Y. A. Sari, and I. Cholissodin, "Analisis Sentimen Twitter menggunakan Metode Naive Bayes dengan Relevance Frequency Feature Selection (Studi Kasus : Opini Masyarakat mengenai Kebijakan New Normal)," *J. Pengemb. Teknol. Inf. dan Ilmu Komput.*, vol. 5, no. 11, pp. 5068–5074, 2021.
- [6] M. Dwiky, C. Fardani, E. Wijayanti, and A. A. Chamid, "Application of Naïve Bayes for Sentiment Analysis of Shopee App User Comments," *Bit-Tech (Binary Digit. - Technol.*, vol. 8, no. 2, 2025, doi: 10.32877/bt.v8i2.2854.
- [7] A. Ahmad, "Sentimen Analisis dengan Long Short-Term Memory dan Synthetic Minority Over Sampling Technic Pada Aplikasi Digital Perbankan," *JTIK (Jurnal Teknol. Inf. dan Komunikasi)*, vol. 8, no. 4, 2024.

- [8] K. Kowsari, K. Jafari Meimandi, M. Heidarysafa, S. Mendu, L. Barnes, and D. Brown, "Text Classification Algorithms: A Survey," *Information*, vol. 10, no. 4, p. 150, Apr. 2019, doi: 10.3390/info10040150.
- [9] A. Sholekhah, "Comparison of Naïve Bayes and Support Vector Machine in Sentiment Analysis of Confiscation of Corrupt Assets on Twitter Perbandingan Naïve Bayes dan Support Vector Machine Dalam Analisa Sentimen Tentang Penyitaan Aset Koruptor di Twitter," *MALCOM Indones. J. Mach. Learn. Comput. Sci.*, vol. 5, no. July, pp. 981–989, 2025.
- [10] S. Xu, Y. Li, and Z. Wang, "Bayesian Multinomial Naïve Bayes Classifier to Text Classification," *Inst. Sci. Tech. Inf. China*, pp. 347–352, 2017, doi: 10.1007/978-981-10-5041-1_57.
- [11] Z. Jiang, B. Gao, Y. He, Y. Han, P. Doyle, and Q. Zhu, "Text Classification Using Novel Term Weighting Scheme-Based Improved TF-IDF for Internet Media Reports," *Math. Probl. Eng.*, vol. 2021, pp. 1–30, Mar. 2021, doi: 10.1155/2021/6619088.
- [12] M. Naufal, B. Arifitama, and S. D. H. Permana, "Sentimen Analisis Tim Sepak Bola Di Piala Asia 2023 Berdasarkan Pengguna Twitter (X) Menggunakan Algoritma Multinomial Naïve Bayes," *Kurawal - J. Teknol. Inf. dan Ind.*, vol. 7, no. 2, pp. 12–21, Oct. 2024, doi: 10.33479/kurawal.v7i2.1098.
- [13] S. Khairunnisa, A. Adiwijaya, and S. Al Faraby, "Pengaruh Text Preprocessing terhadap Analisis Sentimen Komentar Masyarakat pada Media Sosial Twitter (Studi Kasus Pandemi COVID-19)," *J. MEDIA Inform. BUDIDARMA*, vol. 5, no. 2, p. 406, Apr. 2021, doi: 10.30865/mib.v5i2.2835.
- [14] B. Yanuargi, E. Utami, Kusriani, and A. A. Parikesit, "Deep Learning Neural Dalam Analisis Sentimen : Sebuah Studi Literatur," *JITSI J. Ilm. Teknol. Sist. Inf.*, vol. 6, no. 3, pp. 321–328, Sep. 2025, doi: 10.62527/jitsi.6.3.499.
- [15] U. Mahdiyah and R. Wulanningrum, "Implementasi Long Short Term Memory (LSTM) Untuk Forecasting Jumlah Penjualan Daging Ayam," *INOTEK*, vol. 9, pp. 546–553.
- [16] A. A. Chamid, R. R. Isnanto, A. Jazuli, W. H. Sugiharto, A. D. Widianoro, and Sumaji, "Deep Learning Approach for Multi-Class Classification of Textual Data via Bi-LSTM," in *2025 IEEE 11th Information Technology International Seminar (ITIS)*, IEEE, Oct. 2025, pp. 74–79. doi: 10.1109/ITIS67966.2025.11308963.
- [17] 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," *Int. J. Adv. Technol. Eng. Explor.*, vol. 12, no. 127, 2025.
- [18] A. Jazuli, W. Widowati, R. Kusumaningrum, and T. N. Khusna, "Enhancing Aspect-Based Sentiment Analysis in Student Reviews Using Bidirectional Autoencoder and Index Generator Algorithm," *TEM J.*, vol. 14, no. 4, pp. 3412–3426, 2025, doi: 10.18421/TEM144.
- [19] A. Jazuli, Widowati, and R. Kusumaningrum, "Auto Labeling to Increase Aspect-Based Sentiment Analysis Using K-Nearest Neighbors Method," *E3S Web Conf.*, vol. 359, p. 05001, Oct. 2022, doi: 10.1051/e3sconf/202235905001.