

Analisis Sentimen Opini Netizen Terhadap Pergantian Pelatih Timnas Indonesia pada Platform X Menggunakan Metode TF-IDF dan Naïve Bayes

Muhammad Raihan Rabbani^{1*)}, Luthfi Alfian²⁾, Riska Dwi Handayani³⁾

^{1*, 2)} Teknik Informatika, STMIK Bina Patria

³⁾ Manajemen Informatika, STMIK Bina Patria

^{1*)} rabbaniraihan7@gmail.com, ²⁾ luthfialfian399@gmail.com, ³⁾ riska@stmikbinapatria.ac.id

ABSTRACT

This study aims to apply text mining techniques for sentiment analysis on netizen opinions regarding national team coach appointments on X platform. The research stages include data collection, text preprocessing, feature extraction using TF-IDF, classification model development using Multinomial Naive Bayes, and performance evaluation. Text preprocessing is carried out through case folding, data cleaning, tokenization, stopword removal, and stemming to improve data quality. The results indicate that the model achieves an accuracy of 84.87%. However, further analysis reveals an accuracy paradox caused by significant class imbalance in the dataset. While the model performs exceptionally well in predicting the majority class (Neutral), it fails to identify positive and negative sentiments, resulting in a precision, recall, and F1-score of 0.00 for these classes. This study concludes that while Multinomial Naive Bayes provides high global accuracy, alternative approaches such as Support Vector Machine (SVM) or data balancing techniques are required to improve sensitivity toward minority classes.

Keywords: Sentiment Analysis, X, Multinomial Naive Bayes, Accuracy Paradox, Class Imbalance.

I. PENDAHULUAN

Media sosial saat ini menjadi salah satu alat penting bagi masyarakat dalam berbagi informasi, menyampaikan pendapat, serta membentuk opini umum mengenai berbagai isu penting, termasuk keputusan pergantian pelatih di klub maupun tim nasional di berbagai negara. Opini dari masyarakat online bisa menunjukkan bagaimana masyarakat umum merespons suatu keputusan atau peristiwa karena informasi bisa menyebar dengan cepat (Kurniawan & Setiawan, 2022). Namun, jumlah komentar yang sangat banyak membuat proses analisis secara manual menjadi tidak efisien, memakan waktu, dan berisiko tinggi terhadap bias subjektif. Sebab itu, dibutuhkan sistem otomatis yang bisa mengelompokkan komentar sesuai dengan sikap atau perasaan orang-orang secara teratur.

Metode text mining berhasil mengolah data teks dari media sosial. TF-IDF, yang memberikan nilai pada kata-kata penting dalam dokumen atau komentar, adalah salah satu teknik yang sering digunakan dalam text mining. Selanjutnya, komentar dapat dikategorikan ke dalam kategori tertentu menggunakan algoritma *Naive Bayes* (Pratama & Lestari, 2021). Kategori ini dapat positif, negatif, atau netral. Dengan menggabungkan TF-IDF dan *Naive Bayes*, klasifikasi opini masyarakat dapat dilakukan dengan cepat, efektif, dan akurat. Oleh karena itu, penelitian ini dapat menghasilkan wawasan tentang persepsi publik terhadap pergantian pelatih di beberapa tim di seluruh dunia.

Namun demikian, sebagian besar penelitian sebelumnya cenderung menggunakan dataset dengan distribusi kelas yang seimbang. Belum banyak studi yang secara spesifik mengeksplorasi efektivitas *Naive Bayes* pada opini terkait penunjukan pelatih Tim Nasional Indonesia di platform X, platform X dipilih sebagai sumber data karena merupakan salah satu media sosial yang paling aktif digunakan oleh masyarakat Indonesia untuk menyampaikan opini secara terbuka dan *real-time* mengenai isu-isu populer, termasuk olahraga (Wahyudi & Jannah, 2020). Terdapat tantangan berupa dominasi opini netral yang

sangat tinggi dibandingkan opini positif atau negatif. Ketidakseimbangan data (*class imbalance*) ini krusial karena dapat menyebabkan model mengalami bias dan gagal mengenali sentimen ekstrem secara akurat. Oleh karena itu, penelitian ini bertujuan untuk mengimplementasikan algoritma *Naive Bayes* dan fitur TF-IDF untuk mengklasifikasikan opini netizen terkait penunjukan pelatih Tim Nasional Indonesia, sekaligus mengevaluasi performa model dalam menghadapi fenomena *mode collapse* pada dataset yang didominasi oleh kelas tertentu.

II. TINJAUAN PUSTAKA

2.1 Literatur Review

Penelitian tentang analisis opini publik di media sosial telah menjadi bagian penting dari pemahaman kita tentang bagaimana masyarakat bertindak terhadap peristiwa secara real-time. Untuk mengukur opini publik, Kurniawan dan Setiawan (2022) menggunakan metode TF-IDF dan *Naive Bayes*. Hasilnya menunjukkan bahwa metode ini efektif dalam klasifikasi teks dengan tepat. Sebuah penelitian yang dilakukan oleh Pratama dan Lestari (2021) menunjukkan bahwa metode otomatis berbasis teks mining memproses data yang sangat besar dengan lebih cepat dan konsisten daripada metode manual. Namun, masalah utama yang sering muncul dalam penelitian sejenis adalah ketergantungan model pada distribusi data: jika dataset tidak seimbang, model cenderung memberikan hasil yang bias pada kelas mayoritas.

2.2 Text Mining dan TF-IDF

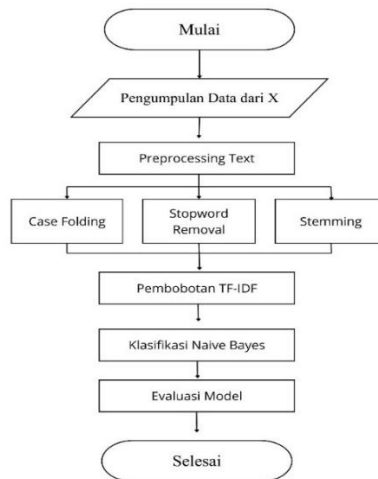
Text mining adalah proses menemukan pola atau informasi penting dari data teks tidak terstruktur melalui fase pra-pemrosesan, ekstraksi fitur, dan klasifikasi. Salah satu teknik ekstraksi fitur yang paling penting adalah *Term Frequency-Inverse Document Frequency* (TF-IDF). Teknik ini memberikan bobot pada setiap kata berdasarkan frekuensi yang muncul dalam sebuah dokumen (*Term Frequency*) dan perbedaan yang dimilikinya dibandingkan dengan frekuensi seluruh dokumen dalam korpus. TF-IDF membantu algoritma klasifikasi membedakan polaritas sentimen dengan memberi nilai tinggi pada kata yang spesifik dan bermakna.

2.3 Algoritma Naïve Bayes

Naive Bayes adalah algoritma klasifikasi probabilistik yang menggunakan teori Bayes dan menganggap bahwa setiap fitur independen. Berdasarkan kata-kata dalam teks, algoritma ini menghitung kemungkinan masuknya teks ke dalam kategori tertentu, seperti positif, negatif, atau netral. Karena kecepatan komputasinya, *Naive Bayes* sering digunakan dalam analisis sentimen media sosial, meskipun memiliki asumsi yang sederhana. Namun, kinerja algoritma ini sangat tergantung pada komposisi dataset. Ketidakseimbangan jumlah data antara kelas dapat menyebabkan fenomena *mode collapse*, di mana model hanya dapat mengidentifikasi kelas dominan.

III. METODE PENELITIAN

Flowchart pada Gambar 1 menunjukkan langkah-langkah metode penelitian yang digunakan dalam penelitian ini. Penelitian dimulai dengan proses pengumpulan data berupa tweet dari platform X yang membahas penunjukan pelatih Tim Nasional Indonesia. Setelah itu, data tersebut diproses melalui tahap preprocessing yang mencakup proses *case folding*, penghapusan stopwords, dan *stemming*. Selanjutnya, data diubah menjadi bentuk representasi menggunakan metode TF-IDF, kemudian diklasifikasikan dengan algoritma *Naive Bayes*, dan terakhir dinilai untuk mengetahui seberapa baik performa model klasifikasi tersebut.



Gambar 1. Metode Penelitian

3.1 Pengumpulan Data (*Data Crawling*)

Penelitian ini menggunakan teknik web scraping pada platform X menggunakan instrumen Tweet-Harvest, yang berbasis Node.js, untuk mengumpulkan data dari halaman pencarian X melalui browser. Untuk autentikasi sistem, Twitter Auth Token digunakan untuk mengumpulkan data. Untuk menentukan kriteria pencarian, kata kunci khusus yang terkait dengan pergantian pelatih Tim Nasional Indonesia digunakan. Selain itu, parameter lang:id digunakan untuk membatasi data dalam bahasa Indonesia.

Untuk memudahkan pengolahan lanjutan, data yang ditarik secara otomatis disimpan ke dalam format.csv. Setelah proses *crawling* selesai, data dimuat ke dalam lingkungan pemrograman menggunakan pustaka Pandas pada Python untuk memverifikasi jumlah dan disiapkan untuk tahap pra-pemrosesan. Untuk menangkap dinamika pendapat netizen secara real-time, penelitian menggunakan skema ini untuk mendapatkan dataset yang objektif dan sesuai dengan periode waktu yang ditetapkan.

3.2 Pelabelan Data (*Labeling*)

Untuk menentukan kebenaran atau kelas target dari setiap opini, tahap pelabelan data dilakukan setelah proses pengumpulan data selesai. Peneliti menggunakan pelabelan manual (*manual labeling*) berdasarkan konteks kalimat dalam tweet. Semua data dimasukkan ke dalam kategori sentimen positif, negatif, dan netral.

Sentimen Negatif diberikan pada komentar yang menunjukkan kritik, kekecewaan, atau penolakan terhadap pemilihan pelatih, seperti dengan kata kunci "kecewa" atau "gagal". Sentimen Positif diberikan pada komentar yang menunjukkan dukungan, kepuasan, atau apresiasi terhadap pemilihan pelatih. Meskipun demikian, sentimen Netral ditunjukkan pada tweet yang informatif, menantang, atau tidak menunjukkan polaritas emosi yang jelas. Untuk

memastikan bahwa model *Naive Bayes* memiliki referensi yang akurat dalam mempelajari pola kata sebelum melakukan klasifikasi otomatis pada data uji.

3.3 Pra-pemrosesan Teks (*Text Preprocessing*)

Untuk mengubah data mentah yang dihasilkan dari *crawling* menjadi data terstruktur yang siap untuk diproses oleh algoritma klasifikasi, tahapan pra-pemrosesan teks adalah langkah penting. Untuk menyeragamkan data dan menghindari duplikasi kata karena perbedaan kapitalisasi, proses ini diawali dengan *Case Folding*, yang mengubah semua karakter teks menjadi huruf kecil. Selanjutnya, proses pembersihan dilakukan untuk menghilangkan gangguan (*noise*) dari data teks, termasuk angka, simbol, tanda baca, tautan URL, dan atribut media sosial seperti *username* (@) dan *hashtag* (#).

Tokenizing adalah proses yang memecah kalimat menjadi potongan kata tunggal atau token setelah data bebas simbol. Langkah berikutnya adalah penghapusan kata-kata umum, seperti kata hubung atau preposisi, yang tidak memiliki makna dalam Bahasa Indonesia. Untuk mengubah setiap kata berimbuhan menjadi kata dasarnya, *stemming* adalah tahap akhir dari pra-pemrosesan. Sebelum masuk ke tahap ekstraksi fitur TF-IDF, proses *stemming* menggunakan algoritma Nazief & Adriani untuk memastikan setiap kata memiliki representasi makna yang konsisten.

3.4 Ekstraksi Fitur TF-IDF

Tujuan dari tahapan ekstraksi fitur adalah untuk mengubah data teks yang telah diproses sebelumnya menjadi representasi numerik yang dapat diolah oleh algoritma klasifikasi. Metode TF-IDF digunakan untuk memberikan bobot pada setiap kata berdasarkan frekuensi kemunculannya dalam dokumen dan kelangkaannya di seluruh dataset. Penggunaan metode ini terbukti efektif dalam menonjolkan kata-kata kunci yang memiliki nilai informasi tinggi untuk proses klasifikasi teks (Hidayatullah & Ma'arif, 2017). *Term Frequency* (TF) mengukur seberapa sering suatu kata muncul dalam satu tweet, sementara *Inverse Document Frequency* (IDF) mengurangi bobot kata-kata yang terlalu sering muncul dalam banyak dokumen yang sama.

$$W_{d,t} = tf_{d,t} \times \log \left(\frac{N}{df_t} \right) \quad (1)$$

Keterangan:

- $W_{d,t}$: Bobot kata t dalam dokumen d .
- $Tf_{d,t}$: Frekuensi kemunculan kata t dalam dokumen d .
- N : Total jumlah dokumen dalam dataset.
- Df_t : Jumlah dokumen yang mengandung kata t .

3.5 Klasifikasi Multinomial Naïve Bayes

Metode ini digunakan untuk mengklasifikasikan sentimen berdasarkan probabilitas kemunculan kata dalam dokumen. Multinomial *Naive Bayes* sangat efektif untuk data teks karena menghitung peluang setiap kelas secara efisien berdasarkan frekuensi fitur yang ada.

$$P(C | X) = \frac{P(X|C)P(C)}{P(X)} \quad (2)$$

Keterangan:

- $P(C|X)$: Peluang kelas C jika diketahui fitur X .
- $P(X|C)$: Peluang fitur X muncul pada kelas C .
- $P(C)$: Peluang awal sebuah kelas.
- $P(X)$: Peluang kemunculan fitur X .

3.6 Parameter dan Evaluasi Model

Studi ini menetapkan parameter untuk membagi dataset dengan rasio 60:40. Berdasarkan 723 entri data yang diperoleh dari proses *crawling* menggunakan instrumen Tweet-Harvest, data dibagi menjadi 434 data untuk tahap pelatihan (training) dan 289 data untuk tahap pengujian (testing). Metode ini dipilih karena memberikan data latihan yang cukup bagi algoritma Multinomial *Naive Bayes* untuk mengenali pola kata, serta jumlah data uji yang cukup untuk memvalidasi secara objektif kinerja model.

Tahap evaluasi dilakukan untuk mengevaluasi efektivitas klasifikasi dengan membandingkan hasil prediksi sistem dengan label yang sebenarnya menggunakan Matriks Konflik. Kinerja model dievaluasi dengan menggunakan metrik Akurasi untuk mengukur persentase prediksi yang tepat, serta metrik Ketepatan, Recall, dan F1-Score untuk memberikan penilaian yang lebih mendalam. Penggunaan metrik ini sangat penting mengingat adanya ketidakseimbangan data dalam dataset. Ini dibutuhkan untuk mengukur kemampuan model untuk menentukan kelas positif, negatif, dan netral.

IV. HASIL DAN PEMBAHASAN

4.1 Dataset

Dataset yang digunakan dalam penelitian ini berasal dari platform X, berupa komentar dan tweet yang membicarakan penunjukan pelatih di beberapa klub dan juga tim nasional di seluruh dunia. Total data yang dikumpulkan sebanyak 723 baris, semuanya ditulis dalam bahasa Indonesia. Sebelum diproses, teks dalam dataset dilakukan pembersihan sederhana dengan menghilangkan mention (@username) dan simbol *hashtag* (#).

Selanjutnya, dataset dibagi menjadi dua bagian, yaitu data latih dan data uji dengan rasio 60:40. Data latih terdiri dari 434 komentar dan data uji berisi 289 komentar. Dataset tersebut digunakan untuk melatih model klasifikasi teks dengan menggunakan TF-IDF sebagai cara mengubah kata-kata menjadi fitur dan *Naive Bayes* sebagai metode untuk mengklasifikasikan sentimen. Struktur dataset ini membantu model mempelajari pola kata dan bisa membedakan opini positif, negatif, atau netral terhadap penunjukan pelatih secara otomatis.

Tabel 1. Contoh Dataset Opini

Content	Label	Keyword
Pelatih baru ini strategi timnya mantap	Positif	mantap
Pemilihan pelatih ini sangat bagus	Positif	bagus
<i>Track record</i> nya dia waktu jadi pelatih bagus	Positif	bagus
Banyak pemain kecewa dengan pelatih baru	Negatif	kecewa
No 1 nya jadi pelatih timnas ya bang	netral	-
Saya belum bisa menilai pelatih ini	Netral	-
Tidak ada perubahan signifikan	Netral	-
Pelatih gagal menyesuaikan strategi	Negatif	gagal

4.2 Preprocessing Text

Selanjutnya, dataset komentar dan tweet yang digunakan dalam penelitian ini diproses melalui tahap preprocessing. Tahap ini bertujuan untuk membersihkan data dari kata atau simbol yang tidak relevan, sehingga pola opini positif, negatif, dan netral dapat dikenali dengan lebih baik. Pada penelitian ini, preprocessing dilakukan melalui tiga teknik utama:

a. Case Folding

Tahap berikutnya adalah mengubah semua huruf dalam komentar menjadi huruf kecil. Tujuannya adalah agar perbedaan antara huruf besar dan kecil tidak mengganggu proses analisis opini. Contoh penerapan teknik *case folding* bisa dilihat pada Tabel 2.

Tabel 2. Contoh Penerapan *Case Folding*

content	case folding
Pelatih baru ini strategi timnya mantap	pelatih baru ini strategi timnya mantap
Pemilihan pelatih ini sangat bagus	pemilihan pelatih ini sangat bagus
<i>Track record</i> nya dia waktu jadi pelatih bagus	<i>track record</i> nya dia waktu jadi pelatih bagus
Banyak pemain kecewa dengan pelatih baru	banyak pemain kecewa dengan pelatih baru
No 1 nya jadi pelatih timnas ya bang	no 1 nya jadi pelatih timnas ya bang
Saya belum bisa menilai pelatih ini	saya belum bisa menilai pelatih ini
Tidak ada perubahan signifikan	tidak ada perubahan signifikan
Pelatih gagal menyesuaikan strategi	pelatih gagal menyesuaikan strategi

b. *Stopword Removal*

Stopword adalah kata-kata yang tidak terkait langsung dengan pendapat utama dalam sebuah komentar, meskipun kata-kata tersebut sering muncul. Kata-kata ini tidak berpengaruh besar dalam menentukan jenis pendapat tersebut. Contoh penerapan penghapusan *stopword* dapat dilihat pada Tabel 3.

Tabel 3. Contoh Penerapan *Stopword Removal*

content	stopword
Pelatih baru ini strategi timnya mantap	pelatih baru strategi timnya mantap
Pemilihan pelatih ini sangat bagus	pemilihan pelatih sangat bagus
<i>Track record</i> nya dia waktu jadi pelatih bagus	<i>track record</i> nya waktu jadi pelatih bagus
Banyak pemain kecewa dengan pelatih baru	banyak pemain kecewa pelatih baru
No 1 nya jadi pelatih timnas ya bang	no 1 nya jadi pelatih timnas bang
Saya belum bisa menilai pelatih ini	saya bisa menilai pelatih ini
Tidak ada perubahan signifikan	tidak perubahan signifikan
Pelatih gagal menyesuaikan strategi	pelatih gagal menyesuaikan strategi

c. *Stemming*

Stemming adalah proses mengubah kata-kata dalam kalimat menjadi bentuk akar (*root word*). Tujuannya adalah untuk menyamakan makna kata agar pemrosesan lebih mudah. Algoritma yang digunakan dalam penelitian ini adalah algoritma Nazief & Adriani. Contoh penerapan *stemming* dapat dilihat pada Tabel 4.

Tabel 4. Contoh Penerapan *Stemming*

content	stemmed
Pelatih baru ini strategi timnya mantap	latih baru strategi tim mantap
Pemilihan pelatih ini sangat bagus	pilih latih sangat bagus
<i>Track record</i> nya dia waktu jadi pelatih bagus	<i>track record</i> nya waktu jadi latih bagus
Banyak pemain kecewa dengan pelatih baru	banyak main kecewa latih baru
No 1 nya jadi pelatih timnas ya bang	no 1 nya jadi latih timnas bang
Saya belum bisa menilai pelatih ini	saya bisa nilai latih ini
Tidak ada perubahan signifikan	tidak ubah signifikan
Pelatih gagal menyesuaikan strategi	latih gagal sesuai strategi

4.3 TF-IDF

Setelah proses pra-pemrosesan selesai, data teks diubah menjadi bentuk angka menggunakan metode *Term Frequency-Inverse Document Frequency* (TF-IDF). TF-IDF digunakan untuk memberi nilai penting pada setiap kata berdasarkan sejauh mana kata tersebut relevan dalam sebuah dokumen. *Term Frequency* (TF) menunjukkan seberapa sering suatu kata muncul dalam dokumen tertentu, sedangkan *Inverse Document Frequency* (IDF) menunjukkan seberapa jarang kata tersebut muncul di semua dokumen.

Dengan metode TF-IDF, kata-kata yang muncul sering tapi tidak bermakna akan memiliki nilai rendah, sedangkan kata-kata yang spesifik dan bermakna akan memiliki nilai lebih tinggi. Representasi ini membantu algoritma klasifikasi dalam membedakan pendapat positif, negatif, dan netral secara lebih baik.

Tabel 5. Tabel Perhitungan TF-IDF

Term (<i>Stemming</i>)	Bobot TF-IDF D1	Bobot TF-IDF D2
pelatih	0,4123	0,2874
ganti	0,3561	0
tim	0,2987	0,3152
strategi	0	0,3018
performa	0	0,2649
Jumlah Bobot	1,0671	1,1693

4.4 Hasil Klasifikasi *Naive Bayes*

Setelah proses preprocessing dan pembobotan kata menggunakan TF-IDF, tahap selanjutnya adalah melakukan klasifikasi teks menggunakan algoritma *Naive Bayes*. *Naive Bayes* merupakan algoritma klasifikasi berbasis probabilistik yang bekerja berdasarkan Teorema Bayes dengan asumsi bahwa setiap fitur (kata) bersifat independen satu sama lain.

Tabel 6. *Confusion Matrix* Hasil Klasifikasi *Naive Bayes*

Kelas Aktual \ Prediksi	Negatif	Positif	Netral
Negatif	0	0	13
Positif	0	0	31
Netral	0	0	247

Pada titik ini, model yang telah dilatih digunakan untuk menguji 289 data uji. Tabel 6 menunjukkan fenomena di mana model cenderung memprediksi seluruh data ke dalam kategori Netral, menunjukkan bahwa model tidak dapat mengidentifikasi karakteristik khusus dari opini Positif dan Negatif. Secara teknis, algoritma *Naive Bayes* tidak dapat mempelajari pola di kelas minoritas karena distribusi data yang tidak seimbang. Akibatnya, model mengalami bias terhadap kelas mayoritas.

4.5 Evaluasi Kinerja Model

Tabel 7. Hasil Evaluasi Klasifikasi *Naive Bayes*

Kelas	Precision	Recall	F1-Score	Support
Negatif	0.00	0.00	0.00	13
Positif	0.00	0.00	0.00	31
Netral	0.85	1.00	0.92	247
Akurasi			0.85	291
Macro Avg	0.28	0.33	0.31	291
Weighted Avg	0.72	0.85	0.78	291

Tabel 7 menunjukkan bahwa model mencatat nilai akurasi sebesar 0,85. Terlepas dari kenyataan bahwa angka-angka ini sangat tinggi, angka-angka ini hanya menunjukkan keberhasilan model dalam memprediksi kelas mayoritas (Netral), bukan kemampuan klasifikasi sebenarnya. Nilai Precision, Recall, dan F1-Score untuk kelas Positif dan Negatif mencapai 0,00. Secara ilmiah, nilai akurasi yang tinggi pada kondisi ini disebut *Accuracy Paradox*. Fenomena ini sejalan dengan penelitian Sari & Sabur (2019) yang menyatakan bahwa pada dataset yang tidak seimbang, akurasi tinggi sering kali bersifat semu karena model cenderung memiliki bias terhadap kelas dengan jumlah sampel terbanyak. Di sini, model terlihat baik secara statistik global, tetapi tidak dapat menemukan sentimen ekstrem secara efektif. Model tampaknya "bermain aman" dengan memasukkan semua prediksi ke kelas Netral untuk menjaga tingkat akurasi yang tinggi karena banyaknya data Netral (247 data) dibandingkan dengan data Positif (31) dan Negatif (13).

V. KESIMPULAN DAN SARAN

5.1 Kesimpulan

Berdasarkan hasil penelitian yang telah dilakukan, dapat disimpulkan bahwa penggunaan algoritma Multinomial *Naive Bayes* untuk menganalisis sentimen opini netizen terhadap penunjukan pelatih Tim Nasional Indonesia menghasilkan tingkat akurasi sebesar 85%. Namun, nilai ini bersifat semu, atau paradoks akurasi, karena model hanya dapat memprediksi kelas mayoritas, yaitu kategori Netral, dengan nilai recall 1,00.

Sebaliknya, model gagal mengkategorikan sentimen ke kategori Positif dan Negatif secara keseluruhan, yang ditunjukkan dengan nilai F1-Score 0,00. Dalam kumpulan data ini, terdapat ketidakseimbangan data kelas yang signifikan, di mana jumlah data netral jauh lebih banyak daripada data sentimen lainnya. Akibatnya, algoritma *Naive Bayes* mengalami bias karena mempertahankan probabilitas tertinggi pada kelas mayoritas dengan mengabaikan karakteristik kata kelas minoritas.

5.2 SARAN

Berdasarkan keterbatasan penelitian ini, peneliti menyarankan beberapa topik untuk pengembangan lebih lanjut. Untuk mendapatkan volume data yang lebih besar dan variasi pendapat yang lebih beragam tentang kelas positif dan negatif, peneliti berikutnya harus melakukan *crawling*, atau memperpanjang periode waktu pengumpulan data. Selain itu, sangat disarankan untuk menggunakan algoritma klasifikasi alternatif seperti Support Vector Machine (SVM) karena dikenal memiliki kinerja yang lebih stabil dalam menangani dataset yang kompleks dibandingkan dengan *Naive Bayes*. Peneliti selanjutnya juga dapat mempertimbangkan penggunaan teknik pembobotan kelas atau pemilihan fitur yang lebih spesifik untuk meningkatkan kemampuan model dalam mendeteksi sentimen ekstrem secara lebih akurat.

DAFTAR PUSTAKA

- Hidayatullah, A. F., & Ma'arif, M. R. (2017). "Penerapan Term Frequency Inverse Document Frequency (TF-IDF) untuk Penentuan Nilai Penting Kata pada Klasifikasi Teks". *Jurnal Teknologi Informasi*.
- Kusuma, R., & Fajar, A., 2021, Analisis Sentimen Publik di Media Sosial, *Jurnal Teknologi dan Informasi*, No.3, Vol.9, hal.34–42.
- Kurniawan, D., & Setiawan, E., 2022, Analisis Opini Publik di Media Sosial Menggunakan TF-IDF dan *Naive Bayes*, *Jurnal Teknologi Informasi Terapan*, No.2, Vol.10, hal.45–53.
- Mustafa, T. F., & Alfianti, H., 2020, Klasifikasi Berita dan Opini Berbahasa Indonesia Menggunakan Algoritma *Naive Bayes*, *Jurnal Sains Informatika Terapan*, No.2, Vol.3, hal.45–52.
- Pratama, A., & Lestari, N., 2021, Penerapan Text Mining untuk Analisis Sentimen Netizen pada Isu Olahraga, *Jurnal Sistem Informasi dan Teknologi*, No.1, Vol.8, hal.23–31.
- Sari, B. W., & Sabur, F. (2019). Analisis Sentimen Menggunakan *Naive Bayes* Classifier Dengan Data Imbalanced. *Jurnal Teknologi Informasi dan Ilmu Komputer (JTIK)*, 6(6), 621-626.
- Sari, R. P., & Handoko, M., 2020, Pemanfaatan Analisis Opini Netizen dalam Evaluasi Kinerja Pelatih Olahraga, *Jurnal Manajemen dan Teknologi Informasi*, No.2, Vol.5, hal.12–20.
- Sriyano, C. S., & Setiawan, E., 2021, Pendeteksian Opini Netizen Menggunakan *Naive Bayes* dan Fitur TF-IDF, *e-Proceeding of Engineering*, No.5, Vol.8, hal.5678–5685.
- Wahyudi, T., & Jannah, M. (2020). "Analisis Sentimen Twitter Menggunakan Algoritma *Naive Bayes*". *Jurnal Informatika*.