

HYBRID CNN-SVM UNTUK ANALISIS SENTIMEN KLUB TIM NASIONAL

Yudha Prastya Laksana Pratama^{1*}, Ronny Makhffudin Akbar², Yanuarini Nur Sukmaningtyas³

^{1,2,3}Universitas Islam Majapahit, Jl. Raya Jabon No.KM.0,7, Tambak Rejo, Gayaman, Kec. Mojoanyar, Kabupaten Mojokerto, Jawa Timur 61364, (0321) 399474

Keywords:

Analisis Sentimen, *Hybrid* CNN-SVM, Instagram, Tim Nasional Indonesia, *Text Mining*.

Correspondent Email:

prastyapratama48@gmail.com



Copyright © [JITET](http://www.jitet.org) (Jurnal Informatika dan Teknik Elektro Terapan). This article is an open access article distributed under terms and conditions of the Creative Commons Attribution (CC BY NC)

Abstrak. Media sosial dan kemajuan teknologi digital telah mengubah cara para penggemar mengekspresikan dukungan mereka kepada tim sepak bola, khususnya Tim Nasional Indonesia. Salah satu platform yang sering digunakan untuk mengekspresikan opini baik opini positif, negatif, maupun netral adalah Instagram. Tujuan dari penelitian ini adalah menerapkan teknik Hybrid CNN-SVM untuk menganalisis sentimen komentar pengguna pada unggahan akun Instagram resmi @timnasindonesia. Teknik ini menggabungkan kemampuan *Support Vector Machines* (SVM) untuk kategorisasi sentimen dengan *Convolutional Neural Networks* (CNN) untuk ekstraksi fitur. Untuk mengumpulkan data, 90 unggahan Instagram dari tanggal 18 Agustus 2023 hingga 16 November 2024, yang mencakup skenario menang, kalah, dan seri, di scrapping komentarnya. Sembilan ribu komentar dikumpulkan, dan mereka menjalani pra-pemrosesan, pelabelan, vektorisasi, pelatihan model, dan evaluasi kinerja dengan menggunakan metrik untuk akurasi, presisi, recall, dan F1-score. Bergantung pada rasio data pelatihan dan data uji, temuan evaluasi menunjukkan bahwa model tersebut bekerja pada kisaran akurasi 69% hingga 73%. Penelitian ini menunjukkan bahwa analisis sentimen pada data teks media sosial, terutama yang berkaitan dengan olahraga nasional, dapat dikatakan kurang baik menggunakan pendekatan hibrida CNN-SVM.

Abstract. Social media and advances in digital technology have changed the way fans express their support for soccer teams, especially the Indonesian National Team. One of the platforms that is often used to express opinions, whether positive, negative, or neutral, is Instagram. The purpose of this research is to apply the Hybrid CNN-SVM technique to analyze the sentiment of user comments on the official @timnasindonesia Instagram account uploads. This technique combines the capabilities of Support Vector Machines (SVM) for sentiment categorization with Convolutional Neural Networks (CNN) for feature extraction. To collect the data, 90 Instagram posts from August 18, 2023 to November 16, 2024, which included win, loss, and draw scenarios, had their comments scrapped. Nine thousand comments were collected, and they underwent pre-processing, labeling, vectorization, model training, and performance evaluation using metrics for accuracy, precision, recall, and F1-score. Depending on the ratio of training data and test data, the evaluation findings showed that the model performed at an accuracy range of 69% to 73%. This study demonstrates that sentiment analysis on social media text data, especially those related to national sports, can be underperformed using the CNN-SVM hybrid approach.

1. PENDAHULUAN

Sepak bola merupakan olahraga dengan jumlah penggemar terbesar di Indonesia. Antusiasme masyarakat terhadap Tim Nasional Indonesia sangat tinggi, terutama ketika tim berlaga dalam pertandingan internasional.

Platform media sosial seperti Instagram menjadi saluran utama bagi masyarakat untuk mengekspresikan opini, baik berupa dukungan maupun kritik terhadap performa tim. Unggahan dan komentar pada akun resmi @timnasindonesia memberikan data berharga

untuk menganalisis persepsi publik secara langsung dan real time [1].

Seiring meningkatnya aktivitas di media sosial, analisis sentimen menjadi alat penting dalam mengukur reaksi dan opini masyarakat. Sentimen positif, negatif, maupun netral terhadap tim nasional dapat memberikan wawasan strategis bagi pengelola tim, media, dan peneliti. Untuk itu, dibutuhkan metode yang efektif dalam mengklasifikasi sentimen dari komentar berbahasa informal dan variatif. Metode tradisional seperti Naive Bayes dan SVM telah banyak digunakan, tetapi memiliki keterbatasan dalam menangkap konteks dan struktur bahasa yang kompleks.

Convolutional Neural Network (CNN) dikenal andal dalam mengekstraksi fitur dari data teks, sedangkan Support Vector Machine (SVM) unggul dalam klasifikasi. Kombinasi keduanya, yang disebut hybrid CNN-SVM, menawarkan pendekatan yang lebih akurat dalam mengolah data komentar media sosial. Namun, hingga saat ini belum banyak penelitian yang menerapkan pendekatan hybrid CNN-SVM secara spesifik pada data komentar media sosial Instagram yang menggunakan bahasa informal khas masyarakat Indonesia. Keterbatasan metode konvensional dalam menangkap konteks dan struktur teks yang tidak baku menjadi tantangan utama dalam klasifikasi sentimen. Oleh karena itu, penelitian ini dilakukan untuk menjawab pertanyaan: Bagaimana menggunakan pendekatan hybrid CNN-SVM untuk analisis sentimen pada akun media sosial Instagram Tim Sepak Bola Nasional?

2. TINJAUAN PUSTAKA

Dalam klasifikasi sentimen tingkat dokumen, setiap kalimat dianalisis menggunakan CNN sebagai detektor subjektivitas sebelum diklasifikasikan dengan SVM. Model ini dilatih dengan output dari CNN untuk menentukan polaritas sentimen suatu teks. Pendekatan ini memungkinkan model menangani negasi serta mengidentifikasi dan mengekstrak fitur subjektif yang penting untuk prediksi sentimen. Dengan memadukan kekuatan CNN dalam pembelajaran fitur dan SVM dalam klasifikasi, arsitektur yang diusulkan diharapkan mampu meningkatkan akurasi prediksi sentimen pada tingkat dokumen. Penelitian ini terbukti lebih unggul

dibandingkan metode klasifikasi lainnya dan dapat memperoleh hasil yang lebih baik. Tingkat recall dan F1-score juga menunjukkan bahwa model kami memiliki performa yang tinggi dengan presisi 97%, recall 96%, F1-score 96%, akurasi 97% penelitian tersebut bersumber dari jurnal yang berjudul “*Deep CNN With SVM-Hybrid Model for Sentence-based Document Level Sentiment Analysis Using Subjectivity Detection*” [2].

Analisis sebelumnya yang dilakukan oleh Piqih Aditiya, Ultach Enri, Iqbal Maulana pada tahun 2022 dengan judul “Analisis Sentimen Ulasan Pengguna Aplikasi MyIM3 pada Situs Google Play Menggunakan *Support Vector Machine*” memanfaatkan algoritma SVM dengan kernel linear dan RBF untuk melakukan klasifikasi sentimen ulasan. Data diperoleh melalui metode *web scraping*, kemudian diproses menggunakan tahapan *text mining* seperti data *cleaning*, *case folding*, *tokenizing*, *filtering*, dan *stemming* menggunakan *library* Sastrawi. Proses transformasi data dilakukan dengan pembobotan TF-IDF. Hasil penelitian menunjukkan bahwa kernel linear memberikan akurasi terbaik sebesar 87% pada skenario C (70:30), sementara kernel RBF mencapai akurasi serupa pada skenario A (90:10). Visualisasi *word cloud* menunjukkan kata “kuota” sebagai kata dominan dalam ulasan positif dan “sinyal” dalam ulasan negatif. Temuan ini menunjukkan efektivitas algoritma SVM dalam analisis sentimen terhadap ulasan aplikasi berbasis teks, [3].

Studi lain yang dilakukan oleh Safrizal Ahmad, Ahmad Maulid Ridwan, Gilang Dwi Setiawan pada tahun 2023 berjudul “Analisis Sentimen Product Tools & Home Menggunakan Metode CNN dan LSTM” membandingkan performa kedua metode dalam analisis sentimen. Hasilnya menunjukkan bahwa model LSTM memiliki akurasi lebih baik sebesar 60,74% dibandingkan CNN yang memperoleh 57,99%. Meskipun demikian, model CNN unggul dalam presisi dengan nilai 65,91%, sedangkan LSTM sebesar 60,22%. Untuk recall, model LSTM lebih tinggi dengan 60,74%, sedangkan CNN mencatatkan 57,99%. Selain itu, nilai F1-Score dari model LSTM sedikit lebih unggul dibandingkan CNN, masing-masing sebesar 60,47% dan 60,07%. Penelitian ini menunjukkan bahwa meskipun CNN lebih baik dalam mengklasifikasikan

kelas positif secara tepat (presisi), LSTM lebih andal dalam menemukan seluruh data kelas positif (recall) serta mempertahankan keseimbangan performa melalui F1-Score [4].

2.1 Analisis Sentimen

Analisis sentimen adalah proses mengidentifikasi dan mengkategorikan opini yang dinyatakan dalam teks untuk menentukan sikap penulis terhadap suatu topik [5]. Pendekatan ini digunakan untuk memahami opini publik, terutama dalam media sosial, dengan klasifikasi sentimen menjadi tiga kategori utama: positif, negatif, dan netral. Analisis ini banyak digunakan dalam bidang bisnis, politik, dan olahraga untuk memahami persepsi publik secara otomatis dan dalam skala besar.[6]

2.2 Convolutional Neural Network (CNN)

CNN adalah salah satu jenis arsitektur deep learning yang dikenal efektif dalam pengolahan data berbentuk gambar dan teks. CNN bekerja dengan cara mengekstraksi fitur dari input menggunakan lapisan konvolusi, pooling, dan fully connected. Dalam konteks teks, CNN mampu menangkap pola lokal seperti frasa penting dalam kalimat. CNN banyak digunakan dalam klasifikasi teks karena keunggulannya dalam mengenali pola sekuensial [7], [8].

2.3 Support Vector Machine (SVM)

Support Vector Machine (SVM) adalah algoritma klasifikasi yang bekerja dengan mencari hyperplane terbaik untuk memisahkan data ke dalam dua kelas atau lebih. SVM berusaha memaksimalkan margin antara data dari kelas berbeda, sehingga membuat model menjadi lebih general dan tidak overfitting. SVM sangat efektif untuk dataset berdimensi tinggi dan tetap bekerja optimal meskipun jumlah fitur lebih banyak daripada jumlah sampel. Dalam konteks analisis sentimen, SVM digunakan untuk mengklasifikasikan opini menjadi kategori sentimen seperti positif, negatif, atau netral. Keunggulan SVM terletak pada kestabilannya dalam menangani data yang kompleks dan tidak linear, terutama dengan penggunaan kernel trick [9].

2.4 Hybrid CNN-SVM

Hybrid CNN-SVM adalah pendekatan gabungan yang memanfaatkan keunggulan CNN dalam mengekstraksi fitur dari teks dan keandalan SVM dalam melakukan klasifikasi. CNN bertugas mengenali pola-pola lokal dalam data teks melalui lapisan konvolusi dan pooling, menghasilkan vektor fitur yang merepresentasikan konten semantik dari komentar. Vektor fitur ini kemudian digunakan sebagai input oleh SVM untuk menentukan kategori sentimen. Pendekatan ini dianggap efektif karena CNN dapat menangkap fitur non-linear dan kompleks dari teks, sementara SVM memberikan margin pemisah yang optimal antar kelas. Kombinasi ini sering digunakan dalam tugas klasifikasi teks untuk meningkatkan akurasi dan mengurangi overfitting dibandingkan jika keduanya digunakan secara terpisah.

2.5 Term Frequency-Inverse Document Frequency (TF-IDF)

Untuk memproses dan mengevaluasi data teks, data yang telah diproses sebelumnya harus diubah menjadi bentuk numerik. Algoritme prediktif digunakan untuk memproses data ini setelah diubah menjadi vektor dan diberi bobot dan nilai untuk setiap kata. Teknik *Term-Frequency Times Inverse Document Frequency* (TF-IDF) digunakan selama tahap pembobotan kata.[10]

Metode ini dipilih karena efektivitas, kemudahan penerapan, dan ketepatan hasilnya. Prosesnya melibatkan perhitungan *Term Frequency* (TF) dan *Inverse Document Frequency* (IDF) untuk setiap kata dalam dokumen [11]. Intinya, teknik ini menghitung seberapa sering kata atau istilah tertentu muncul dalam sebuah dokumen. Frekuensi Istilah (TF) merujuk pada seberapa sering suatu kata atau frasa muncul dalam dokumen, di mana frekuensi istilah yang lebih tinggi (TF tinggi) berkaitan dengan bobot dokumen yang lebih besar, yang menunjukkan nilai kesamaan yang lebih tinggi. Rumus TF-IDF sebagai berikut.

$$TF - IDF_{(t,d,D)} = TF_{(t,d)} \times IDF_{(t,D)} \quad 2.1$$

Keterangan:

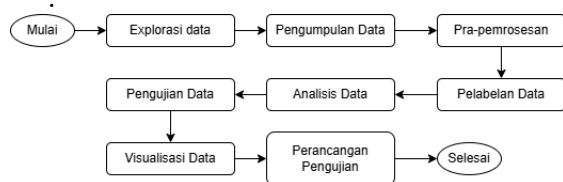
$TF - IDF_{(t,d,D)}$ = pembobotan TF-IDF

$TF_{(t,d)}$ = Jumlah kemunculan term dalam dokumen

$IDF_{(t,d)}$ = bobot inverse dalam nilai df

3. METODE PENELITIAN

Serangkaian tindakan sistematis yang bertujuan untuk mengarahkan proses penelitian dan menghasilkan hasil yang diinginkan dikenal sebagai metodologi penelitian. Gambar 3.1 menunjukkan seluruh prosedur metodologi penelitian.



Gambar 3. 1 Diagram Alir Analisis Sentimen

Pemeriksaan data pada akun Instagram tim nasional Indonesia merupakan langkah pertama dalam penelitian ini. Plugin IG *Comment Export* di Google Chrome digunakan untuk *scraping* komentar untuk mengumpulkan data. Data tersebut kemudian melalui tahapan pra-pemrosesan seperti lemmatization, cleaning, case folding, filtering, tokenizing, dan stemming.

3.1 Pengumpulan Data

Data dikumpulkan dari akun Instagram resmi @timnasindonesia menggunakan ekstensi Chrome IG Comment Export. Total 90 unggahan dipilih berdasarkan tiga kategori hasil pertandingan (menang, kalah, seri), dengan masing-masing 30 unggahan. Dari unggahan ini terkumpul 9.000 komentar.

3.2 Pra-pemrosesan

Berfungsi sebagai awal mengikuti prosedur berbagi data dalam fase ini. Fase pra-pemrosesan teks ini akan menjelaskan langkah-langkah dalam menangani data sebagai persiapan untuk tahap berikutnya [12]. Berikut adalah tahapan yang terlibat dalam tahap pra-pemrosesan teks.

a. Cleaning

Cleaning Proses ini melibatkan penghapusan komponen yang tidak relevan dari sebuah dokumen, termasuk nama pengguna, tagar, angka, tanda baca, dan tautan. Tabel 3.1 adalah contoh dari prosedur ini [13].

Tabel 3. 1 Cleaning

Sebelum	Sesudah
@roni_sullivan49 Timnas Qatar menangis	Timnas Qatar menangis
@rdhostwww AFC nya yg gk bener aturan mainnya	AFC nya yg gk bener aturan mainnya

b. Case folding

Case folding Proses ini melibatkan penggantian semua karakter menjadi huruf kecil untuk menstandarkan format teks, sehingga data menjadi seragam dan siap untuk analisis lebih lanjut [14].

Tabel 3. 2 Case Folding

Sebelum	Sesudah
Timnas Qatar menangis	timnas qatar menangis
AFC nya yg gk bener aturan mainnya	afc nya yg gk bener aturan mainnya

c. Filtering

Filtering adalah proses menyaring kata-kata dalam teks dengan menghapus *stopwords*, adalah istilah yang sering digunakan seperti kata ganti, kata sambung, dan kata seru yang tidak memengaruhi cara dokumen dikategorikan.

Tabel 3. 3 Filtering

Sebelum	Sesudah
pelatihnya siapa si ini gils bgt	[pelatihnya, si, gils, bgt]
afc nya yg gk bener aturan mainnya	[afc, nya, yg, gk, bener, aturan, mainnya]

d. Tokenizing

Metode tokenisasi melibatkan pembagian teks menjadi kata-kata diskrit atau token, dipisahkan oleh spasi, yang kemudian digunakan dalam proses pembuatan matriks dokumen [15].

Tabel 3. 4 Tokenizing

Sebelum	Sesudah
pelatihnya siapa si ini gils bgt	[pelatihnya,siapa,si,ini,gils,bgt]
afc nya yg gk bener aturan mainnya	[afc,nya,yg,gk,bener,aturan,main nya]

e. *Stemming*

Stemming adalah proses pengurangan kata-kata penting ke bentuk dasarnya berdasarkan kriteria Kamus Besar Bahasa Indonesia (KBBI). Pada tahap ini, perpustakaan Python Sastrawi digunakan untuk melaksanakan prosedur stemming.

Tabel 3. 5 Steaming

Sebelum	Sesudah
afc nya yg gk benar aturan mainnya	afc nya yg gk benar atur main
ini lah yang di namakan hilang nya seni sepakbola	nama hilang nya seni sepakbola

3.3 Pelabelan Data

Setelah tahap pra-pemrosesan, prosedur pelabelan kelas sentimen dilakukan. Untuk mendapatkan representasi korpus yang sesuai dengan kebutuhan analisis, langkah ini sama pentingnya. Dengan menentukan nilai penilaian menggunakan kamus leksikal, pelabelan sentimen secara otomatis selesai. Sentimen positif, netral, dan negatif adalah tiga kategori sentimen yang secara umum dipisahkan dalam prosedur kategorisasi ini. Daftar istilah bahasa Indonesia yang mencakup kata-kata positif dan negatif digunakan untuk mengevaluasi sebuah teks dan menentukan apakah teks tersebut termasuk dalam kategori positif atau negatif. Kelas sentimen setiap dokumen akan ditentukan oleh nilai skor dari istilah-istilah ini.

Menghitung jumlah kata positif dan negatif dalam setiap kalimat ulasan adalah dasar untuk klasifikasi otomatis. Ulasan diklasifikasikan sebagai positif jika skor kata positif lebih tinggi, negatif jika skor negatif lebih tinggi, dan netral jika jumlahnya sama.

3.4 Vektorisasi Data

Dalam penelitian ini, vektorisasi data teks dilakukan menggunakan metode *Term Frequency-Inverse Document Frequency* (TF-IDF) untuk mengubah dokumen teks menjadi vektor numerik yang dapat diolah oleh

algoritma *machine learning*, yaitu *Hybrid CNN-SVM*. TF-IDF dipilih karena kemampuannya menonjolkan kata-kata penting dalam dokumen tertentu sambil mengurangi bobot kata-kata umum yang muncul di banyak dokumen dalam *corpus*. Proses vektorisasi dimulai dengan pra-pemrosesan teks, meliputi tokenisasi untuk memecah teks menjadi kata-kata, pembersihan tanda baca dan karakter khusus, pengubahan ke huruf kecil, serta penghapusan *stop words* seperti “dan” atau “di” menggunakan pustaka seperti Sastrawi untuk bahasa Indonesia. Selanjutnya, semua kata unik dari *corpus* dikumpulkan untuk membentuk vokabuler. Bobot TF-IDF dihitung berdasarkan frekuensi kata dalam dokumen dan kelangkaannya di seluruh *corpus*, menghasilkan matriks TF-IDF di mana setiap dokumen direpresentasikan sebagai vektor.

3.5 Pembagian Dataset

Dalam penelitian ini, pembagian data untuk pelatihan dan pengujian model dilakukan dengan 5 rasio yaitu 90% : 10%, 80% : 20%, 70% : 30%, 60% : 40%, 50% : 50%. Pendekatan ini umum digunakan dalam machine learning untuk memastikan model dapat belajar dari sebagian besar data sekaligus dievaluasi pada data yang belum pernah dilihat sebelumnya, sehingga dapat mengukur kemampuan generalisasi model.

3.6 Pelatihan Model

a. Pelatihan Model Menggunakan CNN

Penelitian ini menggunakan Convolutional Neural Network (CNN) untuk mengekstraksi fitur dari data teks. CNN dipilih karena kemampuannya dalam mengenali pola-pola lokal dalam kalimat, seperti frasa atau hubungan antar kata, yang penting dalam analisis sentimen. Proses pelatihan melibatkan pengubahan teks menjadi representasi vektor dan dilanjutkan dengan pendeteksian fitur penting untuk membedakan makna dalam konteks kalimat.

Hasil pelatihan menunjukkan bahwa CNN mampu menghasilkan representasi teks yang lebih kaya dan kontekstual dibandingkan metode tradisional seperti TF-IDF. Dengan

demikian, model ini lebih efektif dalam mengidentifikasi sentimen positif maupun negatif dari teks, karena mampu menangkap pola-pola khusus yang mencerminkan emosi atau opini dalam bahasa.

b. Penggabungan antara TF-IDF dengan CNN

Penelitian ini menggabungkan dua pendekatan representasi teks, yaitu TF-IDF dan CNN, untuk memperoleh hasil yang lebih komprehensif dalam analisis sentimen. TF-IDF dikenal mampu merepresentasikan seberapa penting suatu kata dalam dokumen berdasarkan frekuensinya, sedangkan CNN unggul dalam menangkap pola kontekstual dan hubungan antar kata dalam kalimat. Dengan menggabungkan kedua pendekatan ini, diharapkan model dapat memahami teks dari dua sisi: makna kata secara statistik dan makna dalam konteks.

Hasil penggabungan kedua fitur ini menciptakan representasi yang lebih kaya dan informatif. Pendekatan ini terbukti meningkatkan kemampuan model dalam membedakan sentimen positif dan negatif karena memanfaatkan keunggulan masing-masing metode. Namun, peningkatan kompleksitas data juga membawa tantangan tersendiri, seperti kebutuhan komputasi yang lebih besar dan potensi terjadinya overfitting jika tidak ditangani dengan baik.

c. Klasifikasi menggunakan SVM

Tahap akhir dalam proses analisis sentimen dilakukan dengan menggunakan metode klasifikasi Support Vector Machine (SVM) berbasis kernel linear. SVM dipilih karena kemampuannya dalam menangani data berdimensi tinggi, seperti fitur gabungan yang berasal dari metode TF-IDF dan CNN. Dalam konteks ini, SVM berperan sebagai alat untuk memisahkan data ke dalam

kategori sentimen positif dan negatif secara optimal.

Model SVM dilatih menggunakan data yang telah diperkaya dengan kombinasi dua jenis fitur tersebut, sehingga mampu mengenali pola sentimen dengan lebih baik. Pendekatan ini dikenal sebagai Hybrid CNN-SVM, yang menggabungkan keunggulan CNN dalam mengekstraksi informasi kontekstual dari teks dan keandalan SVM dalam melakukan klasifikasi. Hasilnya, model yang dihasilkan diharapkan memiliki performa yang lebih akurat dan stabil dalam membedakan jenis-jenis sentimen yang terkandung dalam data teks.

4. Hasil Evaluasi model

Tabel 4. 1 Evaluasi keseluruhan

Rasio Train : Test	Akurasi	0 (P/R/F1)	1 (P/R/F1)	2 (P/R/F1)
90:10	73%	0.76 / 0.79 / 0.78	0.70 / 0.74 / 0.72	0.72 / 0.60 / 0.65
80:20	71%	0.75 / 0.78 / 0.77	0.67 / 0.72 / 0.70	0.72 / 0.55 / 0.63
70:30	70%	0.74 / 0.75 / 0.75	0.65 / 0.71 / 0.68	0.71 / 0.57 / 0.63
60:40	69%	0.71 / 0.76 / 0.73	0.67 / 0.68 / 0.68	0.70 / 0.56 / 0.52
50:50	69%	0.71 / 0.76 / 0.73	0.66 / 0.68 / 0.67	0.72 / 0.56 / 0.63

Hasil evaluasi kinerja model klasifikasi berdasarkan lima skenario pembagian data latih dan uji 90:10, 80:20, 70:30, 60:40, dan 50:50 ditampilkan pada Tabel 4.2. Untuk setiap kelas (kelas 0, kelas 1, dan kelas 2), evaluasi dilakukan dengan menggunakan metrik akurasi, presisi, *recall*, dan *f1-score*. Selain itu, nilai rata-rata *makro* dan rata-rata tertimbang dihitung sebagai ukuran kinerja rata-rata model secara keseluruhan.

Dari temuan ini, terbukti bahwa rasio 90:10, atau 73%, menghasilkan akurasi terbaik. Akurasi model menurun seiring dengan berkurangnya persentase data pelatihan,

mencapai 71% pada 80:20, 70% pada 70:30, dan 69% pada 60:40 dan 50:50. Hal ini menunjukkan bahwa kapasitas model untuk menggeneralisasi data uji sangat bergantung pada *volume* data pelatihan yang digunakan, dengan semakin besar kumpulan data pelatihan, semakin baik kinerja model.

Kelas 0 secara teratur memiliki metrik yang lebih besar daripada kelas lainnya, menurut tinjauan kinerja untuk setiap kelas. Nilai akurasi, *recall*, dan *f1-score* kelas 0 berada di antara 0.71 dan 0.78, menunjukkan bahwa model dapat secara akurat mengidentifikasi dan meramalkan kelas ini. Dengan nilai *f1-score* berkisar antara 0,67 hingga 0,72, performa model di kelas 1 cenderung bervariasi secara signifikan. Kelas 2 memiliki kinerja terburuk, terutama dalam hal *recall*, yang hanya berkisar antara 0,55 dan 0,60. Meskipun nilai presisi kelas 2 relatif tinggi (sekitar 0,70-0,72), nilai *recall* yang rendah menunjukkan bahwa model tidak dapat mengidentifikasi sebagian besar data aktual kelas tersebut, yang menghasilkan nilai *f1-score* yang rendah (sekitar 0,52-0,65). Mengingat kelas 2 memiliki data yang jauh lebih sedikit dibandingkan kelas lainnya, distribusi dataset yang tidak merata dapat menjadi penyebab kinerja yang buruk.

Model ini juga secara konsisten berkinerja lebih buruk ketika rasio data pelatihan turun jika dilihat dari nilai rata-rata makro dan rata-rata tertimbang. Pada rasio 60:40 dan 50:50, rata-rata makro turun dari 0.72 pada 90:10 menjadi 0.68. Karena memperhitungkan jumlah data di setiap kelas, rata-rata tertimbang cenderung sedikit lebih tinggi, tetapi juga menunjukkan tren penurunan yang serupa. Hal ini lebih lanjut menunjukkan bagaimana kinerja model dipengaruhi oleh ketidakseimbangan kelas dan sangat sensitif terhadap kuantitas data pelatihan.

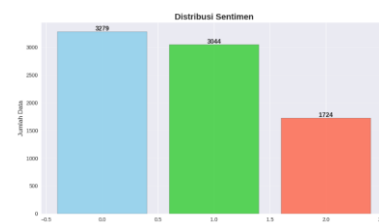
Akurasi metode hybrid CNN-SVM yang diteliti berkisar antara 69% hingga 73% pada variasi pembagian data (90:10, 80:20, 70:30, 60:40, 50:50) bisa disebabkan oleh beberapa faktor berikut:

- Ukuran dan kualitas dataset: Model mungkin tidak dapat mengidentifikasi pola terbaik jika dataset tidak cukup besar atau representatif. Performa juga dapat dipengaruhi oleh ketidakseimbangan data, dan akurasi akan terhambat jika tidak ada metode

pengambilan sampel yang berlebihan seperti SMOTE atau pengoptimalan fitur.

- CNN dan SVM secara bersama-sama bukanlah kombinasi terbaik: Penyesuaian yang cermat diperlukan saat menggabungkan SVM untuk klasifikasi dengan CNN untuk ekstraksi fitur. Agar CNN dapat menghasilkan fitur yang representatif, pelatihan yang memadai dan prapemrosesan yang baik diperlukan. Jika tidak, teknik tunggal yang lebih baik dapat mengungguli kinerja gabungan.
- Distribusi data pelatihan dan pengujian: Performa dapat berubah jika rasio distribusi data pelatihan dan pengujian diubah. Karena model dilatih dengan lebih banyak data, pembagian 90:10 biasanya menghasilkan akurasi yang lebih tinggi; sebaliknya, pembagian yang sama (50:50) dapat menghasilkan akurasi yang lebih rendah karena ada lebih sedikit data pelatihan (penurunan yang dilaporkan dari 73% menjadi 69%).
- Perbedaan kontekstual dan prapemrosesan data: Analisis sentimen sangat sensitif terhadap representasi teks (word2vec, TF-IDF) dan prapemrosesan teks (tokenisasi, stopwords, stemming). Akurasi dapat menurun jika metode yang digunakan kurang efisien

4.1. Distribusi Sentimen

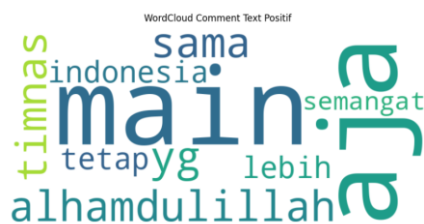


Gambar 4. 1 Distribusi Sentimen

Gambar 4.1 menunjukkan distribusi sentimen pada data komentar akun Instagram Tim Nasional Indonesia ditampilkan dalam bentuk grafik batang. Hasil klasifikasi menunjukkan bahwa kategori sentimen negatif, yang ditandai dengan warna biru, merupakan kategori dengan jumlah terbanyak yaitu sebanyak 3.279 komentar. Disusul oleh

sentimen netral berwarna hijau dengan jumlah 3.044 komentar. Sementara itu, kategori sentimen positif yang ditandai dengan warna merah menjadi kategori dengan jumlah paling sedikit, yakni sebanyak 1.724 komentar. Berdasarkan distribusi ini dapat disimpulkan bahwa mayoritas komentar yang diterima akun Instagram Tim Nasional Indonesia selama periode pengamatan didominasi oleh komentar bernada negatif, diikuti oleh komentar netral, dan paling sedikit berupa komentar positif.

4.2. Wordcloud Positif



Gambar 4. 1 Wordcloud Positif

Sikap yang mendukung Grafik *WordCloud* menunjukkan kumpulan istilah yang paling sering digunakan dalam komentar dukungan yang dibuat oleh penggemar tim nasional Indonesia di Instagram. Di antara kata-kata yang paling sering digunakan adalah “bagus”, “mantap”, “timnas juara”, “main apik”, dan “semangat”. Kemunculan kata-kata ini menunjukkan bahwa sebagian besar komentar positif berisi apresiasi, pujian, serta dukungan moral terhadap performa Timnas, baik saat meraih kemenangan maupun ketika pertandingan berlangsung.

4.3. Wordcloud Negatif



Gambar 4. 2 Wordcloud Negatif

Gambar *WordCloud* sentimen negatif memperlihatkan kata-kata yang paling sering muncul dalam komentar yang bernada kritik, kekecewaan, atau keluhan terhadap Timnas Indonesia. Kata-kata seperti “kalah”, “main”, “ga”, “timnas”, dan “apa” tampak mendominasi *WordCloud* ini. Hal tersebut menunjukkan

adanya ketidakpuasan dari sebagian pendukung terhadap performa tim, keputusan wasit, maupun hasil pertandingan.

4.4. Wordcloud Netral



Gambar 4. 3 Wordcloud Netral

Gambar *WordCloud* sentimen netral menampilkan kata-kata yang paling sering ditemukan dalam komentar yang tidak secara langsung mengandung opini positif maupun negatif. Kata-kata yang muncul didominasi oleh istilah umum seperti “timnas”, “indonesia”, “main”, dan “garuda”. Komentar-komentar ini umumnya bersifat deskriptif, menyebut fakta pertandingan, atau sekadar menyampaikan informasi tanpa menunjukkan ekspresi emosional.

5. Kesimpulan

Berdasarkan hasil penelitian yang telah dilakukan terhadap analisis sentimen komentar pengguna Instagram pada akun resmi @timnasindonesia dengan menggunakan metode Hybrid CNN-SVM, maka diperoleh beberapa kesimpulan sebagai berikut:

1. Metode *Hybrid* CNN-SVM dapat digunakan untuk melakukan klasifikasi sentimen pada komentar media sosial, khususnya dalam mengidentifikasi komentar positif, negatif, dan netral. Model ini bekerja dengan cara memanfaatkan CNN untuk mengekstrak fitur penting dari data teks, kemudian SVM digunakan sebagai algoritma klasifikasi.
2. Akurasi model yang dihasilkan dalam penelitian ini berada pada rentang 69% hingga 73%, dengan hasil terbaik pada skenario pembagian data latih dan uji 90:10. Nilai akurasi tersebut menunjukkan bahwa model dapat bekerja cukup baik, namun belum optimal untuk keperluan klasifikasi yang sangat presisi. Akurasi yang

sedang ini kemungkinan disebabkan oleh karakteristik data komentar Instagram yang tidak terstruktur, mengandung bahasa tidak baku, serta banyak menggunakan slang dan sarkasme.

3. Kinerja model menunjukkan bahwa semakin besar proporsi data latih, maka semakin tinggi akurasi yang dapat dicapai. Akan tetapi, peningkatannya tidak signifikan dan cenderung stagnan setelah titik tertentu. Hal ini mengindikasikan bahwa perlu dilakukan optimasi lanjutan, baik dari sisi pre-processing, arsitektur model, maupun pemilihan dan penyempurnaan data latih.
4. Secara umum, metode *hybrid* CNN-SVM mampu memberikan hasil klasifikasi yang cukup baik untuk kebutuhan awal analisis sentimen di media sosial. Namun, untuk mencapai performa yang lebih andal dan akurat, perlu dilakukan pengembangan lebih lanjut pada aspek pemrosesan teks, pemilihan fitur, serta teknik balancing data dan tuning *hyperparameter*.

UCAPAN TERIMA KASIH

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

DAFTAR PUSTAKA

- [1] A. M. Jon, "Analisis Sentimen Pada Media Sosial Instagram Klub Persija Jakarta Menggunakan Metode Naive Bayes," *Automata*, vol. 958, pp. 1–8, 2022.
- [2] S. M. V. Raviya, K., "Deep CNN With SVM-Hybrid Model for Sentence-Based Document Level Sentiment Analysis Using Subjectivity Detection," *ICTACT J. Soft Comput.*, vol. 11, no. 03, pp. 2344–2352, 2021, doi: 10.21917/ijsc.2021.0335.
- [3] P. Aditiya, U. Enri, and I. Maulana, "Analisis Sentimen Ulasan Pengguna Aplikasi Myim3 Pada Situs Google Play Menggunakan Support Vector Machine," *JURIKOM (Jurnal Ris. Komputer)*, vol. 9, no. 4, p. 1020, 2022, doi: 10.30865/jurikom.v9i4.4673.
- [4] S. Ahmad, A. M. Ridwan, and G. D. Setyawan, "Analisis Sentimen Product Tools & Home Menggunakan Metode Cnn Dan Lstm," *Teknokon*, vol. 6, no. 2, pp. 133–140, 2023, doi: 10.31943/teknokon.v6i2.154.
- [5] D. Duei Putri, G. F. Nama, and W. E. Sulistiono, "Analisis Sentimen Kinerja Dewan Perwakilan Rakyat (DPR) Pada Twitter Menggunakan Metode Naive Bayes Classifier," *J. Inform. dan Tek. Elektro Terap.*, vol. 10, no. 1, pp. 34–40, 2022, doi: 10.23960/jitet.v10i1.2262.
- [6] B. W. Sari and F. F. Haranto, "Implementasi Support Vector Machine Untuk Analisis Sentimen Pengguna Twitter Terhadap Pelayanan Telkom Dan Biznet," *J. Pilar Nusa Mandiri*, vol. 15, no. 2, pp. 171–176, 2019, doi: 10.33480/pilar.v15i2.699.
- [7] A. F. Hidayatullah and R. A. N. Nayoan, "Analisis Sentimen Berbasis Fitur pada Ulasan Tempat Wisata Menggunakan Metode Convolutional Neural Network(CNN)," *Univ. Islam Indones.*, no. 2019-12-04, pp. 1–7, 2019, [Online]. Available: www.cnet.com.
- [8] R. A. Tilasefana and R. E. Putra, "Penerapan Metode Deep Learning Menggunakan Algoritma CNN Dengan Arsitektur VGG NET Untuk Pengenalan Cuaca," *J. Informatics Comput. Sci.*, vol. 05, no. 1, pp. 48–57, 2023.
- [9] R. W. Pratiwi, S. F. H., D. Dairoh, D. I. Af'idah, Q. R. A., and A. G. F., "Analisis Sentimen Pada Review Skincare Female Daily Menggunakan Metode Support Vector Machine (SVM)," *J. Informatics, Inf. Syst. Softw. Eng. Appl.*, vol. 4, no. 1, pp. 40–46, 2021, doi: 10.20895/inista.v4i1.387.
- [10] A. T. Ni'mah and A. Z. Arifin, "Perbandingan Metode Term Weighting terhadap Hasil Klasifikasi Teks pada Dataset Terjemahan Kitab Hadis," *Rekayasa*, vol. 13, no. 2, pp. 172–180, 2020, doi: 10.21107/rekayasa.v13i2.6412.
- [11] S. Lestari and M. Febryanti, "Analisis Sentimen Mengenai Produk Inovasi Invisible Tv Menggunakan Algoritma Naïve Bayes Sentiment Analysis Using the Naïve Bayes Algorithm for Invisible Tv Innovation Product," *J. Inf. Technol. Comput. Sci.*, vol. 7, no. 5, 2024, [Online]. Available: <https://databoks.katadata.co.id/>
- [12] H. Harnelia, "Analisis Sentimen Review Skincare Skintific Dengan Algoritma Support Vector Machine (Svm)," *J. Inform. dan Tek. Elektro Terap.*, vol. 12, no. 2, 2024, doi: 10.23960/jitet.v12i2.4095.
- [13] R. C. Larasati, C. Dewi, and H. J. Christanto, "Analisis Sentimen Produk Kecantikan Jenis Moisturizer Di Twitter Menggunakan Algoritma Super Vector Machine," *J. Tek.*

- Inf. dan Komput.*, vol. 7, no. 1, p. 124, 2024, doi: 10.37600/tekinkom.v7i1.1243.
- [14] A. Liawati, R. Narasati, D. Solihudin, C. Lukman Rohmat, and S. Eka Permana, "Analisis Sentimen Komentar Politik Di Media Sosial X Dengan Pendekatan Deep Learning," *JATI (Jurnal Mhs. Tek. Inform.*, vol. 7, no. 6, pp. 3557–3563, 2024, doi: 10.36040/jati.v7i6.8248.
- [15] J. W. Iskandar and Y. Nataliani, "Perbandingan Naïve Bayes, SVM, dan k-NN untuk Analisis Sentimen Gadget Berbasis Aspek," *J. RESTI (Rekayasa Sist. dan Teknol. Informasi)*, vol. 5, no. 6, pp. 1120–1126, 2021, doi: 10.29207/resti.v5i6.3588.