



Klasifikasi Sentimen Pada Dataset yang Terbatas Menggunakan Algoritma Convolutional Neural Network

M Ridho Saputra, Surya Agustian*, Jasril, Novriyanto

Fakultas Sains dan Teknologi, Program Studi Teknik Informatika, UIN Sultan Syarif Kasim Riau, Pekanbaru, Indonesia

Email: ¹12050116746@students.uin-suska.ac.id, ^{2,*}surya.agustian@uin-suska.ac.id, ³jasril@uin-suska.ac.id,

³novriyanto@uin-suska.ac.id

Email Penulis Korespondensi: surya.agustian@uin-suska.ac.id

Abstrak—Penelitian ini bertujuan untuk menganalisis respons publik terhadap penunjukan Kaesang Pangarep sebagai Ketua Umum Partai Solidaritas Indonesia (PSI) dengan pendekatan klasifikasi sentimen berbasis algoritma Convolutional Neural Network (CNN). Dataset utama berisi 300 tweet berbahasa Indonesia yang diklasifikasikan ke dalam tiga kategori sentimen: positif, negatif, dan netral. Keterbatasan jumlah data pelatihan menjadi tantangan utama karena dapat menghambat kemampuan model dalam melakukan generalisasi. Untuk mengatasi kendala ini, perluasan data dilakukan dengan menggabungkan data eksternal dengan topik Covid-19 dan Open Topic. Tahapan pra-pemrosesan mencakup pembersihan teks, normalisasi, dan tokenisasi. Model CNN yang dikembangkan menggunakan arsitektur berlapis dengan penerapan teknik regularisasi seperti L2 dan dropout guna mengurangi risiko overfitting. Metrik akurasi, F1-Score, Precision dan Recall digunakan untuk mengevaluasi performa model. Hasil pengujian menunjukkan bahwa F1-Score sebesar 0,62 pada data validasi dan 0,51 pada data uji, menunjukkan bahwa kinerja terbaik diperoleh ketika data Kaesang dan Covid-19 digabungkan. Temuan ini menunjukkan bahwa menambahkan data eksternal dapat meningkatkan akurasi klasifikasi, bahkan dalam kondisi data yang terbatas. Penelitian ini berkontribusi pada pengembangan metode klasifikasi sentimen berbasis deep learning untuk teks bahasa Indonesia.

Kata Kunci: Klasifikasi Sentimen; Convolutional Neural Network; Kaesang Pangarep; Dataset Terbatas; Partai Solidaritas Indonesia

Abstract—This study aims to analyze public responses to the appointment of Kaesang Pangarep as the Chairman of the Indonesian Solidarity Party (PSI) using a sentiment classification approach based on the Convolutional Neural Network (CNN) algorithm. The primary dataset consists of 300 Indonesian-language tweets categorized into three sentiment classes: positive, negative, and neutral. The limited size of the training data presents a major challenge, as it can hinder the model's ability to generalize. To address this issue, data augmentation was carried out by incorporating external datasets with Covid-19 and Open Topic themes. The preprocessing stages include text cleaning, normalization, and tokenization. The developed CNN model utilizes a layered architecture and applies regularization techniques such as L2 and dropout to reduce the risk of overfitting. Accuracy, F1-score, precision, and recall were used as performance evaluation metrics. Experimental results show that the best performance was achieved when the Kaesang and Covid-19 datasets were combined, yielding an F1-score of 0.62 on the validation set and 0.51 on the test set. These findings indicate that adding external data can improve classification accuracy, even under limited data conditions. This study contributes to the development of deep learning-based sentiment classification methods for Indonesian-language texts.

Keywords: Sentiment Classification; Convolutional Neural Network; Kaesang Pangarep; Limited Dataset; Indonesian Solidarity Party

1. PENDAHULUAN

Kemajuan teknologi informasi telah mengubah cara masyarakat berinteraksi dengan berita dan peristiwa publik, mengarah pada pergeseran aktivitas komunikasi ke ranah digital, khususnya media sosial. Twitter, yang sekarang dikenal sebagai X, adalah salah satu platform paling populer yang memungkinkan orang-orang dari seluruh dunia untuk bebas berbagi pemikiran mereka. Informasi yang dihasilkan dari percakapan di platform ini memiliki potensi besar untuk diteliti lebih lanjut melalui analisis opini guna memahami persepsi masyarakat terhadap isu tertentu [1], [2]. Twitter menjadi pusat diskusi dinamis, terutama dalam konteks isu-isu politik. Terpilihnya Kaesang Pangarep sebagai Ketua Umum Partai Solidaritas Indonesia (PSI) menjadi salah satu pokok bahasan yang banyak menyita perhatian. Partai Solidaritas Indonesia (PSI) yang berdiri pada tahun 2014, mengusung nilai-nilai seperti keberagaman, partisipasi generasi muda, dan pembaruan sistem politik. Masuknya Kaesang seorang figur muda yang dikenal publik ke dalam kepemimpinan PSI menimbulkan berbagai tanggapan yang tersebar luas di Twitter [3], [4]. Reaksi beragam ini mencerminkan ketertarikan dan perhatian publik, menjadikan fenomena tersebut relevan untuk dikaji secara sistematis.

Dalam konteks ini, analisis sentimen menjadi pendekatan yang tepat untuk mengungkap pandangan masyarakat terhadap peristiwa politik tersebut melalui teks yang dipublikasikan di media sosial. Penelitian ini melanjutkan inisiatif kolaboratif yang dilakukan oleh S. Agustian *et al* [5], berupa shared task yang bertujuan untuk mengklasifikasikan sentimen terhadap tokoh publik secara efisien. Salah satu tantangan utama dalam pendekatan ini adalah jumlah data pelatihan yang terbatas, sehingga diperlukan strategi untuk memperluas dataset menggunakan topik lain yang relevan guna meningkatkan kinerja model.

Beberapa studi sebelumnya menawarkan berbagai pendekatan untuk mengatasi keterbatasan data. Ravil *et al.* di dalam [6] menunjukkan bahwa F1-Score ditingkatkan secara signifikan dari 52,70% menjadi 72,96% dengan memasukkan lebih banyak data dari masalah COVID-19 ke dalam dataset Kaesang dan menggunakan metode *Naïve Bayes* yang dioptimalkan dengan *Particle Swam Optimization* (PSO). Sementara itu, Zahwa *et al.* di dalam [7] memperlihatkan bahwa penggabungan data eksternal dengan teknik *oversampling* dapat secara nyata meningkatkan performa model Bi-LSTM, yang menghasilkan f1-score hingga 67,77% dengan fitur *word2vec*. Penelitian lain oleh Saputra *et al.* [8] menggunakan metode TF-IDF bersama dengan algoritma *Support Vector Machine* (SVM), dan



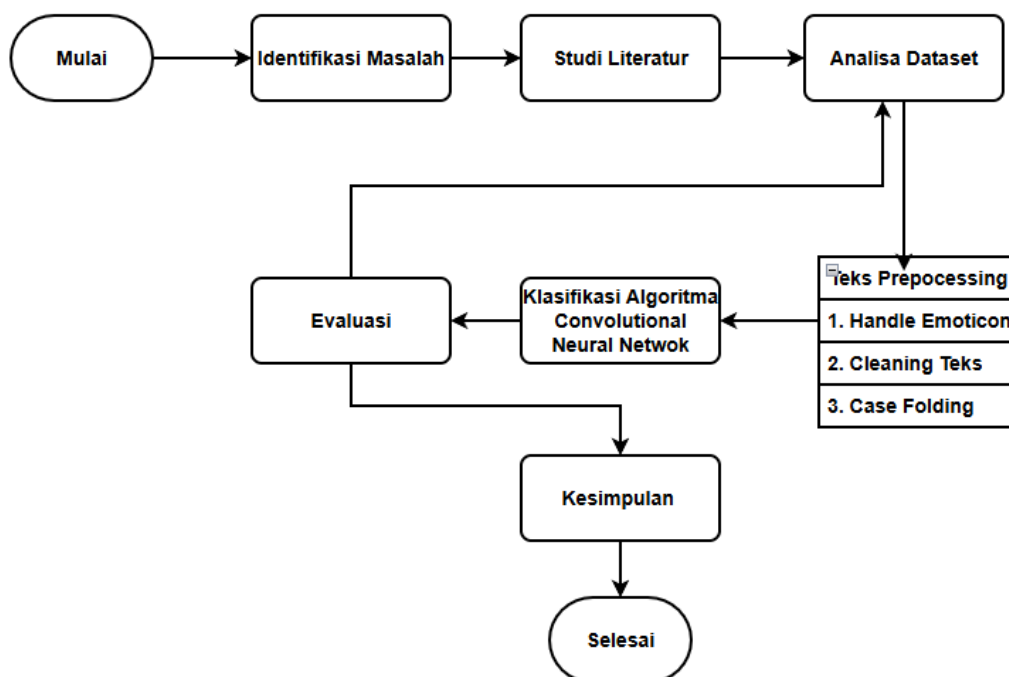
mencapai F1-Score sebesar 0,51. Namun, sebagian besar pendekatan tersebut masih bergantung pada penambahan data dan belum sepenuhnya mengatasi tantangan keterbatasan data secara optimal. Di sisi lain, pendekatan berbasis *Convolutional Neural Network* (CNN) menunjukkan performa yang lebih menjanjikan. Studi oleh [9] mengungkap bahwa peningkatan jumlah lapisan konvolusi dalam arsitektur CNN dapat meningkatkan akurasi klasifikasi hingga mencapai 90%. CNN dikenal memiliki keunggulan dalam mengenali pola penting dalam data teks tanpa proses ekstraksi fitur secara manual. Algoritma ini telah banyak digunakan dalam bidang *Natural Language Processing* (NLP), termasuk untuk tugas analisis sentimen, identifikasi emosi, dan pembuatan ringkasan teks. Dibandingkan metode konvensional seperti SVM, Naïve Bayes, atau Logistic Regression, CNN sering kali menghasilkan kinerja yang lebih baik. Beberapa arsitektur lanjutan seperti *DoubleMax* CNN bahkan menunjukkan peningkatan akurasi klasifikasi rata-rata sebesar 17% [10], [11]. Meskipun beberapa studi telah mengatasi keterbatasan data melalui perluasan dataset, pendekatan tersebut belum sepenuhnya optimal untuk klasifikasi sentimen pada data terbatas. Oleh karena itu, diperlukan kajian lebih lanjut mengenai efektivitas arsitektur CNN berlapis dengan regularisasi dalam menangani klasifikasi sentimen pada dataset terbatas, untuk dapat melihat perbandingan metode *machine learning* dan *deep learning* dalam mengatasi tantangan keterbatasan data.

Dengan mempertimbangkan hal-hal tersebut, penelitian ini bertujuan untuk menilai opini publik terhadap penunjukan Kaesang Pangarep sebagai Ketua Umum Partai Solidaritas Indonesia (PSI) melalui penerapan algoritma *Convolutional Neural Network* (CNN). Dataset terdiri dari 300 tweet yang dianalisis dikumpulkan menggunakan kata kunci terkait isu Kaesang dan diberi label positif, negatif, atau netral.. Pendekatan ini diharapkan dapat memberikan analisis yang mendalam serta menunjukkan efektivitas CNN dalam melakukan klasifikasi sentimen, meskipun dengan jumlah data yang terbatas.

2. METODOLOGI PENELITIAN

2.1. Tahapan Penelitian

Pendekatan yang diterapkan dalam penelitian ini menggambarkan tahapan-tahapan yang harus dilalui untuk menyelesaikan studi, yang disajikan pada Gambar 1 di bawah ini.



Gambar 1. Tahapan Penelitian

Penelitian dimulai dengan identifikasi masalah yang menjadi tahap awal. Setelahnya, dilakukan kajian pustaka untuk memperoleh wawasan yang lebih mendalam mengenai topik yang sedang dianalisis. Langkah selanjutnya adalah analisis dataset, yang mencakup pengumpulan dan seleksi data relevan untuk analisis sentimen. Setelah dataset tersedia, tahap berikutnya adalah *Text Preprocessing*, yang mencakup langkah-langkah penting seperti *Handle Emoticon* (penanganan emoji) sebagai tahap awal. Tahap selanjutnya di *Text Preprocessing* yaitu, *Cleaning Text* (pembersihan kata) dengan melakukan penghapusan URL, *mention*, angka. Tahap terakhir pada *Text Preprocessing* yaitu, *case folding*. Setelah *Text Preprocessing* selesai, data akan dikelompokkan dengan menggunakan algoritma *Convolutional Neural Network* (CNN) sesuai kategori yang tepat. Setelah model dikembangkan, evaluasi dilakukan untuk mengukur kinerja model dengan menggunakan metrik yang telah ditentukan seperti *F1-score*, *accuracy*, *precision*, dan *recall*. Hasil evaluasi ini kemudian dianalisis dan disajikan dalam bagian hasil dan kesimpulan, sebelum penelitian dinyatakan selesai.



2.2. Identifikasi Masalah

Penelitian ini berfokus pada upaya peningkatan kinerja klasifikasi sentimen dalam kondisi keterbatasan data dengan memanfaatkan data eksternal, seperti data terkait COVID-19 dan Open Topic, melalui penerapan algoritma *Convolutional Neural Network* (CNN). Proses penelitian diawali dengan identifikasi permasalahan yang kemudian dilanjutkan dengan studi pendahuluan untuk memperkuat kerangka konsep dan pendekatan penelitian. Fokus utama terletak pada penerapan CNN untuk mengklasifikasikan sentimen dalam teks berbahasa Indonesia. Penelitian ini bertujuan untuk meningkatkan efisiensi klasifikasi sentimen, khususnya dalam menghadapi tantangan keterbatasan jumlah data yang tersedia.

2.3. Studi Literatur

Kajian pustaka ini bertujuan untuk meningkatkan pemahaman mengenai topik penelitian serta mengidentifikasi teori, konsep, dan metodologi yang telah digunakan dalam penelitian-penelitian sebelumnya yang relevan. Kajian pustaka ini akan difokuskan pada klasifikasi sentimen menggunakan algoritma *Convolutional Neural Network* (CNN). Selain itu, pemanfaatan data eksternal seperti data COVID-19 juga akan dibahas untuk menilai bagaimana penambahan data tersebut dapat memperkaya dataset dan meningkatkan performa model klasifikasi sentimen, terutama pada dataset yang terbatas. Studi literatur ini mengumpulkan informasi dari berbagai sumber, seperti buku, jurnal ilmiah, artikel internasional, dan referensi lainnya.

2.4. Analisis Dataset

Dataset dari penelitian Agustian et al. digunakan dalam penelitian ini sebagaimana tercantum dalam[5]. Sumber data Berbagai ungkapan sikap publik terkait Covid-19, penunjukan Kaesang sebagai Ketua Umum Partai Solidaritas Indonesia (PSI), dan isu luas lainnya (*open topics*). Dataset utama dikumpulkan melalui proses crawling dari media sosial Twitter, dengan total 2.309 tweet yang berhasil dihimpun selama periode 25 September hingga 3 Oktober 2023. Pengambilan data dilakukan menggunakan kata kunci “Kaesang PSI”, dan proses pelabelan sentimen dilakukan oleh beberapa anotator melalui pendekatan crowdsourcing. Dari keseluruhan data yang terkumpul, hanya 300 tweet yang digunakan sebagai data pelatihan, dengan distribusi seimbang untuk masing-masing kelas sentimen, yaitu positif, netral, dan negatif masing-masing sebanyak 100 tweet. Dataset yang tersedia dalam penelitian ini dapat dilihat pada Tabel 1.

Tabel 1. Dataset Penelitian yang tersedia [5]

No	Dataset	Penggunaan	Jumlah Sampel Data	Distribusi Kelas		
				Positif	Negatif	Netral
1	Kaesang V1	Training	300	100	100	100
2	Kaesang V2	Training	300	100	100	100
3	Covid-19	Training	8000	463	873	6664
4	Open Topic	Training	7569	1505	2656	3408
5	Kaesang	Testing	924			

Berdasarkan Tabel 1, penelitian ini menggunakan dua jenis data, yaitu data pelatihan (training) dan data pengujian (testing). Data pelatihan terdiri dari Dataset Kaesang V1, Kaesang V2, serta Dataset Covid-19 dan Open Topic yang digunakan untuk membangun model klasifikasi sentimen. Sementara itu, data pengujian berupa Dataset Kaesang yang terdiri dari 924 tweet digunakan untuk mengevaluasi performa model. Evaluasi dilakukan melalui sistem *leaderboard* guna membandingkan hasil prediksi antar model secara objektif.

Pada penelitian ini, Dataset Kaesang V1 dan V2 digabungkan menggunakan fungsi `pd.concat` dari library *pandas*. Gabungan dataset tersebut kemudian dibagi menjadi 80% data pelatihan dan 20% data validasi menggunakan fungsi `train_test_split` dari *scikit-learn*, dengan parameter `random_state=42` untuk menjamin reproduktibilitas. Selanjutnya, bagian data pelatihan hasil pembagian ini digabungkan kembali dengan Dataset Covid-19 dan *Open Topic*, juga menggunakan `pd.concat`, guna memperluas jumlah data yang digunakan dalam proses pelatihan. Kemudian, dataset uji Kaesang digunakan sebagai data pengujian untuk mendapatkan skor evaluasi pada sistem peringkat *leaderboard*.

2.5. Text Preprocessing

Text Preprocessing adalah langkah pertama dalam memodifikasi teks asli dengan prosedur tertentu untuk menghapus atau mengubah elemen-elemen yang tidak relevan, sehingga mempermudah pengolahan data lebih lanjut. *Text Preprocessing* bertujuan untuk mengonsistenkan bentuk kata dalam indeks agar lebih seragam. Indeks yang dimaksud adalah representasi dari konten dokumen yang digunakan dalam proses pencarian. Kebanyakan mesin tidak dapat membedakan antara huruf besar, huruf kecil, tanda baca, dan elemen lainnya, sehingga *Text Preprocessing* menjadi penting untuk menjembatani kesenjangan ini[12]. Langkah-langkah yang dilakukan dalam *Text preprocessing* meliputi:

a. Handle Emoticon

Pada proses *Text preprocessing*, penanganan emoji dilakukan dengan memanfaatkan fungsi `remove_emoji()` dari pustaka Python *emoji*. Tindakan ini bertujuan untuk menghilangkan gangguan dari elemen non-verbal (emoji) yang berpotensi menurunkan performa model dalam tahap tokenisasi dan klasifikasi sentimen.

b. Cleaning Text

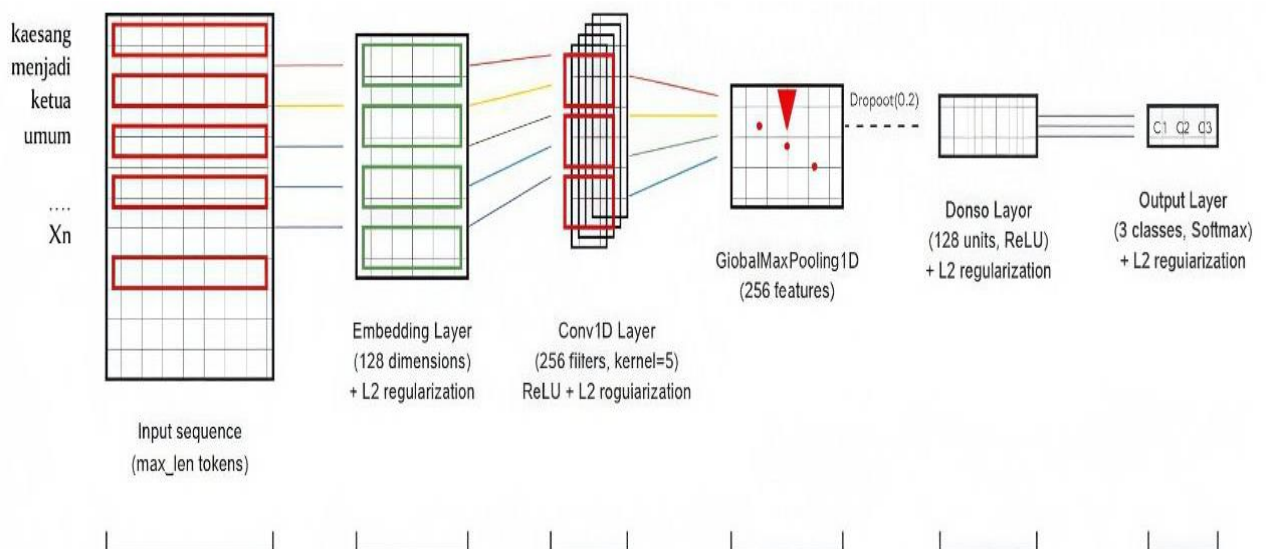
Proses *Cleaning Text* pembersihan data memanfaatkan metode *Regular Expression (regex)* untuk menghapus komponen non-teks seperti tautan, *mention*, tagar, angka, simbol, serta karakter non-alfabet guna menjaga konsistensi dan kebersihan data sebelum memasuki tahap tokenisasi[13].

c. Case Folding

Case folding dilakukan dengan mengonversi seluruh teks menjadi huruf kecil menggunakan fungsi `.lower()`. Tujuannya adalah untuk menyeragamkan bentuk kata serta mengurangi perbedaan yang timbul akibat variasi kapitalisasi sebelum proses tokenisasi dan klasifikasi.

2.6. Klasifikasi dengan Algoritma *Convolutional Neural Network*

Penelitian ini menggunakan model *Convolutional Neural Network (CNN)* dengan arsitektur berlapis untuk klasifikasi sentimen dalam teks berbahasa Indonesia. Dengan menerapkan filter pada berbagai bagian teks, CNN mampu menangkap pola dan ketergantungan lokal, seperti kombinasi kata yang sering muncul dan menunjukkan sentimen tertentu [14]. Struktur model ditampilkan pada Gambar 2 yang terdiri dari enam lapisan utama yang tersusun secara berurutan.



Gambar 2. Arsitektur CNN

Gambar 2 menampilkan arsitektur *Convolutional Neural Network (CNN)* yang diimplementasikan dalam penelitian ini. Diagram tersebut menggambarkan tahapan pemrosesan data teks, dimulai dari input hingga prediksi kelas sentimen. Proses diawali dengan lapisan *Embedding* yang mengubah setiap token input ($X_1, X_2, \dots, X_n$) menjadi vektor berdimensi 128, yang merepresentasikan makna semantik dari kata-kata tersebut dalam ruang vektor. Regularisasi L2 diterapkan untuk meminimalkan kemungkinan *overfitting* pada tahap awal ini[15].

Selanjutnya, representasi vektor dari *embedding* diproses oleh lapisan *Conv1D*, yang memiliki 256 filter dengan ukuran *kernel* 5. Lapisan ini berperan dalam mengekstraksi pola lokal dari sekuens kata. Fungsi aktivasi *ReLU* digunakan untuk menambahkan non-linearitas ke dalam proses pembelajaran, sementara regularisasi L2 turut diterapkan untuk menjaga kestabilan bobot.

Fitur yang dihasilkan dari lapisan konvolusional kemudian dipadatkan menggunakan *GlobalMaxPooling1D*, yang memilih nilai maksimum dari setiap feature map, sehingga menghasilkan representasi fitur yang ringkas dan paling signifikan. Proses ini diikuti oleh lapisan *Dropout* dengan tingkat dropout 0.2, yang berfungsi sebagai mekanisme regularisasi tambahan dengan menonaktifkan sebagian neuron selama pelatihan untuk mencegah ketergantungan berlebihan terhadap fitur tertentu. Hasil *pooling* selanjutnya diteruskan ke lapisan *Dense* dengan 128 neuron dan aktivasi *ReLU*, yang berfungsi menggabungkan dan menyaring kembali informasi dari fitur sebelumnya. Regularisasi L2 sebesar 0.001 diterapkan untuk memperkuat generalisasi model. Terakhir, lapisan *output* berupa *Dense* dengan 3 neuron dan aktivasi *softmax* digunakan untuk menghasilkan distribusi probabilitas terhadap tiga kelas sentimen: positif, negatif, dan netral.

Arsitektur CNN ini dirancang untuk mengidentifikasi fitur semantik penting dalam teks secara efisien, bahkan ketika data pelatihan terbatas. Kombinasi lapisan konvolusional dan *pooling* membantu menangkap pola-pola lokal yang relevan terhadap sentimen, sedangkan teknik *dropout* dan regularisasi memberikan perlindungan terhadap *overfitting*. Dengan demikian, model mampu melakukan klasifikasi sentimen secara efektif pada data teks tidak terstruktur seperti komentar di media sosial berbahasa Indonesia.

2.7. Evaluasi

Pada tahap ini, teknik analisis utama yang digunakan untuk menilai kinerja model adalah confusion matrix, dengan jumlah baris yang dimodifikasi sesuai dengan jumlah model yang dibangun pada langkah sebelumnya. Confusion matrix digunakan untuk menggambarkan performa klasifikasi model terhadap tiga label sentimen: positif, negatif, dan netral.



Setiap baris menunjukkan jumlah data aktual dalam suatu kelas, sedangkan kolom menunjukkan prediksi yang dibuat oleh model. Dapat dilihat pada Tabel 2 yang menampilkan *multiclass confusion matrix*.

Tabel 2. *Multiclass Confusion Matrix*

		Prediksi		
		Positif	Negatif	Netral
Aktual	Positif	TPos	FPosNeg	FPosNet
	Negatif	FNegPos	TNeg	FNegNet
	Netral	FNetPos	FNetNeg	TNet

Hasil klasifikasi model ditunjukkan pada Tabel 2 terkait dengan tiga kategori sentimen: Positif (Pos), Negatif (Neg), dan Netral (Net). Label yang diproyeksikan adalah hasil dari model yang diterapkan, sedangkan label aktual adalah anotasi manual yang berfungsi sebagai referensi utama. Jika hasil yang diprediksi cocok dengan label aktual, klasifikasi dianggap benar (True/T) jika tidak, dianggap salah (False/F).

Hasil evaluasi ditampilkan sebagai tabel dengan metrik evaluasi termasuk *accuracy*, *precision*, *recall*, dan *F1-score* [16]. *Precision* mengukur persentase prediksi positif yang benar terhadap seluruh prediksi positif yang dihasilkan oleh model [17], sedangkan *recall* merupakan metrik yang digunakan untuk menilai sejauh mana model mampu mengenali seluruh data yang benar-benar termasuk dalam suatu kategori [18]. Persamaan yang digunakan untuk menentukan *recall* and *precision* untuk setiap kelas ditunjukkan dalam penelitian Yohana. *et al* [19]. Nilai *True* dan *False* untuk masing-masing kelas (Positif, Negatif, dan Netral) diperoleh dari tabel *confusion matrix*. Adapun nilai *accuracy* diperoleh dari rasio jumlah prediksi yang benar terhadap total keseluruhan data uji. Akurasi setiap kelas dihitung menggunakan Persamaan (1),

$$Akurasi = \frac{True\ Positif + True\ Negatif + True\ Netral}{Seluruh\ Data} \quad (1)$$

Rumus tersebut menunjukkan bahwa akurasi dihitung dengan menjumlahkan seluruh prediksi yang benar yaitu jumlah data yang benar diklasifikasikan sebagai positif (*True Positive*), negatif (*True Negative*), dan netral (*True Netral*), kemudian dibagi dengan total jumlah seluruh data uji. Nilai akurasi mencerminkan seberapa besar proporsi prediksi model yang tepat dibandingkan seluruh data yang dievaluasi. *F1-score* dihitung menggunakan Persamaan (2).

$$F1 - Score = 2 \cdot \frac{Precision \times Recall}{Precision + Recall} \quad (2)$$

Rumus tersebut menunjukkan bahwa *F1-Score* menghitung rata-rata harmonik dari *precision* dan *recall*. *F1-Score* memberikan ukuran keseimbangan antara keduanya, terutama saat terdapat ketidakseimbangan data antar kelas. Penelitian ini menggunakan *F1-score* sebagai metrik utama untuk evaluasi performa model ketika baik ketepatan prediksi (*precision*) maupun kemampuan menangkap seluruh kasus relevan (*recall*) sama-sama penting, yang kemudian dihitung menggunakan pendekatan *F1-macro average* pada masing-masing kelas, sebagaimana dijelaskan dalam Persamaan (3).

$$F1 - macro\ avg = \frac{F1(pos) + F1(net) + F1(neg)}{3} \quad (3)$$

Rumus tersebut menunjukkan bahwa *F1-macro average* menghitung rata-rata *F1-score* dari kelas positif, netral, dan negatif. Metode ini cocok digunakan saat ingin mengevaluasi performa model secara adil pada semua kelas.

3. HASIL DAN PEMBAHASAN

3.1. Dataset

Dataset awal terdiri dari beberapa sumber data yang berbeda, yaitu dataset utama Kaesang dan dataset eksternal Covid-19. Dataset Kaesang memuat tweet yang telah diklasifikasikan ke dalam tiga kategori sentimen: negatif, netral, dan positif.

Tabel 3. Data Kaesang

No	Kalimat	Label
1	@kurawa @psi_id Selamat untuk Mas Kaesang, semoga dpt membesarkan PSI. Salam dr JOKOWERS Mataram, Lombok. NTB. https://t.co/ffquj2oRz7	Positif
2	@kurawa @psi_id @kaesangp Lumayan buat kak Sangpisang mendapatkan panggung untuk salurkan bakatnya sbg stand-up komedi, Kaesang Jokowi emang jossðŸ˜ƒ #BuatLuculucuanAja #PSITambahLucu #PSISemangkinLucu.	Negatif
3	@_missGe @Metro_TV @kaesangp @jokowi @psi_id Gibran lebih conservative dibanding kaesang,...seperti bapaknya hati hati dalam melangkah, karna gibran bermain panjang...kaesang belum Tau ini masih baru buat di analisa menurut saya.	Netral
4	Salut buat mas Kaesang, PSI masuk Senayan Indonesia jaya! ðŸ™ŒðŸ™Œ #NinjakaesangPSI #SangKetum #SangKaesang Anak Mudaa PSI Untuk Indonesiaa https://t.co/Ps6FLdXn9A	Positif



Tabel 3 menyajikan distribusi tweet yang telah diklasifikasikan ke dalam kategori sentimen sebagai bagian dari analisis klasifikasi sentimen. Opini negatif menunjukkan ketidakpuasan atau kritik, sentimen positif menunjukkan dukungan atau pandangan optimis terhadap Kaesang sebagai Ketua Umum PSI. Sikap yang tidak positif maupun negatif ditunjukkan oleh kategori netral.

Table 4. Data Covid-19

No	Kalimat	Label
1	Alhamdulillah sudah vaksin ..Terima Kasih @utycampus #vaksāncovid19 #covid_19 #health #healthy #healthylifestyleâ€¦ https://t.co/mUe66BTZu3	Positif
2	@collegemenfess Belum bisa deh kayaknya nderðŸœŒ jangan2 nanti pas kampusmu offline kamu sendirian online gara2 ga vaksin hayoloðŸœŒ	Netral
3	BPOM sendiri aja bilang GMP dari pembuatan vaksin ini masih belum terkalibrasi dengan benar kok ya kepala batu bgtâ€¦ https://t.co/Ycx5NNacGq .	Negatif
4	@JSuryoP1 @BPOM_RI persoalan vaksin kok jd urusan dukung mendukung...sekalian aja di buat pemilu utk vaksin itu.	Negatif

Tabel 4 memperlihatkan sebaran tweet yang berkaitan dengan Covid-19, yang telah dikelompokkan ke dalam tiga kategori sentimen. Kategori positif mencerminkan dukungan atau pandangan optimis terhadap Covid-19, sementara kategori negatif merepresentasikan penolakan atau sikap kritis. Di sisi lain, sentimen netral mencerminkan pandangan yang tidak menunjukkan keberpihakan, baik mendukung maupun menolak.

Table 5. Data Open Topic

No	Kalimat	Label
1	"Innalillahi wa inna ilaihi raji'un, segenap keluarga besar Pemprov DKI Jakarta mengucapkan turut berduka cita sedalam-dalamnya atas berpulangnyā Bapak Muhammad Harmawan, Camat Kelapa Gading. Semoga husnul khotimah dan mendapatkan tempat terbaik di sisiNya. Aamiin...â€¦" https://t.co/0uFDvvn3Do	Positif
2	Portofolio masing2 pelajaran 2, ada yang 3 juga, belum lagi uprak yang ampir semua ribet Cape parah gile Gatau rasanya tah napa ku makin ga peduli nilai selain un karna itu merupakan harga diri sampe nanti Yang lainnya yang penting selesai doang targetnya	Netral
3	@BPPT_RI Sekalian pantau ke Techno Park Pondok Pusaka, kab kaur, Prov Bengkulu. Banyak sekali fasilitas produksi dan Laboratorium sdh dibangun Megah, tapi tidak jalan optimal. sayang uang negara habis tidak dilanjutkan pemda @rbtv_bengkulu @BappenasRI	Negatif
4	@IndoPluralitas @aniesbaswedan daannnn dengan anggaran yg 80 Trilyun dan semua habis terpakai, defisit malah, dia tidak bisa mengendalikan banjir... hebatnya hanya tata kata, lalu kerugian warga DKI siapa yg harus tanggung jawab pak?	Negatif

Tabel 5 menampilkan distribusi tweet yang membahas berbagai topik umum (*open topic*) yang tidak secara khusus berkaitan dengan tokoh atau isu tertentu. Tweet-tweet tersebut telah diklasifikasikan ke dalam tiga kategori sentimen, yaitu positif, negatif, dan netral. Kategori positif mencerminkan ekspresi empati, dukungan, atau doa yang bersifat membangun, seperti ucapan belasungkawa atau harapan baik. Sentimen negatif ditunjukkan melalui kritik, keluhan, atau ketidakpuasan terhadap kebijakan atau situasi tertentu. Sementara itu, sentimen netral menggambarkan opini atau pernyataan yang bersifat informatif tanpa muatan emosi yang kuat ke arah positif maupun negatif.

3.2. Penerapan Text Preprocessing

Setelah dilakukan proses penghapusan emoji. Kemudian, dilakukan pembersihan teks pada dataset yang mencakup sejumlah prosedur, seperti penghapusan URL, penghilangan spasi berlebih, penghapusan hashtag, serta penghapusan *username*. Setelah itu, penerapan teknik *case folding*. Hasil dari rangkaian proses pembersihan dan fitur *preprocessing* tersebut disajikan pada Tabel 6, yang memperlihatkan output akhir dari tahap *text preprocessing*.

Table 6. Preprocessing Teks

No	Sebelum Preprocessing	Sesudah Preprocessing
1	@CNNIndonesia Masuk Rekor Muri Belum seminggu gabung langsung otw jadi ketum PSI + ketum termuda di indonesia Mantap ,master senior di dunia politik harus belajar dari mas kaesang ðŸœŒ	masuk rekor muri belum seminggu gabung langsung otw jadi ketum psi ketum termuda di indonesia mantap master senior di dunia politik harus belajar dari mas kaesang
2	Panda Nababan merespon diangkatnya Kaesang jadi Ketua Umum PSI @psi_id. Menurutnya hanyalah ajang lucu-lucuan semata. Kalau menurut Anda? #Kaesang #PSI #Politik https://t.co/OdOVBJ2ss8	panda nababan merespon diangkatnya kaesang jadi ketua umum psi menurutnya hanyalah ajang luculucuan semata kalau menurut anda
3	Tak hanya itu, Ketua MPR Bambang Soesatyo, dan Ketua Umum Partai Solidaritas Indonesia (PSI) Kaesang	tak hanya itu ketua mpr bambang soesatyo dan ketua umum partai solidaritas indonesia psi



No	Sebelum Preprocessing	Sesudah Preprocessing
	Pangarep juga terlihat hadir. #BersamaPrabowo #MenataMasaDepan Prabowo Subianto	kaesang pangarep juga terlihat hadir prabowo subianto

Tabel 6 memperlihatkan perbandingan antara kalimat sebelum dan sesudah dilakukan *preprocessing*, yang menunjukkan bahwa teks hasil pembersihan menjadi lebih bersih, sederhana, dan siap digunakan pada tahap selanjutnya.

3.3. Model Baseline CNN

Model Baseline CNN menjadi tahap awal untuk menguji kinerja model *Convolutional Neural Network* (CNN) sebelum dilakukan penambahan data eksternal. Pada tahap ini, model dilatih menggunakan gabungan dataset Kaesang V1 dan V2 yang terdiri dari 600 tweet, dengan distribusi yang seimbang untuk tiga kelas sentimen: positif, negatif, dan netral. Data kemudian dikonversi ke dalam representasi vektor menggunakan *embedding layer*. Arsitektur CNN yang digunakan terdiri dari lapisan *embedding*, *Conv1D* dengan 256 filter, *GlobalMaxPooling1D*, *dropout*, dan dua lapisan *Dense*, dengan regularisasi L2 untuk mencegah *overfitting*. Hasil evaluasi pada data validasi menunjukkan bahwa model mampu mencapai *F1-score* sebesar 0,60, *accuracy* 0,59, *precision* 0,63, dan *recall* 0,60. Nilai ini menjadi acuan dasar untuk membandingkan peningkatan performa pada proses optimasi model yang melibatkan data eksternal.

3.4. Proses Optimasi Model

Pada proses optimasi model, model selanjutnya dievaluasi dengan mengintegrasikan data eksternal ke dalam dataset pelatihan Kaesang. Dataset pelatihan utama merupakan data asli Kaesang, sementara data eksternal yang ditambahkan secara bertahap berasal dari dataset Covid-19 dan *Open Topic*. Evaluasi performa model dilakukan menggunakan data validasi Kaesang V1+V2. Hasil eksperimen dari penambahan data eksternal terhadap Kaesang V1+V2 disajikan pada Tabel 7.

Tabel 7. Eksperimen Setup

Dataset	ID Run	Jumlah Distribusi Kelas			Evaluasi terhadap data validasi			
		Positif	Negatif	Netral	F1-Score	Akurasi	Recall	Precision
Kaesang V1+V2	R1	200	200	200	0.60	0.59	0.60	0.63
Kaesang V1+V2 dan Covid-19	R2	200	200	200	0.57	0.57	0.58	0.57
	R3	400	400	400	0.62	0.61	0.61	0.62
	R4	450	650	1100	0.62	0.62	0.62	0.68
Kaesang V1+V2 dan Open Topic	R5	200	200	200	0.52	0.54	0.58	0.57
	R6	400	400	400	0.61	0.61	0.61	0.64
	R7	450	650	1100	0.53	0.53	0.53	0.56
Kaesang V1+V2 dan Covid-19 + Open Topic	R8	200	200	200	0.50	0.50	0.51	0.51
	R9	400	400	400	0.61	0.60	0.61	0.60
	R10	450	650	1100	0.57	0.57	0.56	0.61

Tabel 7 menunjukkan performa model CNN pada berbagai kombinasi dataset yang menggabungkan data utama Kaesang V1+V2 dengan data eksternal dari domain Covid-19 dan *Open Topic*. Dua konfigurasi memberikan hasil evaluasi terbaik dengan distribusi kelas yang berbeda.

Pada konfigurasi Kaesang V1+V2 dan Covid-19 dengan distribusi kelas sebanyak 450 positif, 650 negatif, dan 1100 netral, model mencapai *F1-Score* tertinggi sebesar 0,62, dengan *accuracy*, *recall*, dan *precision* masing-masing sebesar 0,62, 0,62, dan 0,68. Hasil ini mengindikasikan bahwa penambahan data Covid-19 yang memiliki distribusi kelas lebih besar dan bervariasi mampu meningkatkan performa klasifikasi, khususnya dalam mengenali kelas mayoritas (netral). Nilai *precision* yang relatif tinggi menunjukkan kemampuan model dalam meminimalkan prediksi positif yang salah.

Sementara itu, pada kombinasi Kaesang V1+V2 dan *Open Topic* dengan distribusi kelas seimbang sebanyak 400 data per kelas, model menunjukkan performa yang cukup baik dengan *F1-Score* 0,61, *accuracy* 0,61, *recall* 0,61, dan *precision* 0,64. Distribusi data yang seimbang ini membantu model mempelajari pola secara adil dari setiap kelas, sehingga menghasilkan prediksi yang lebih konsisten dan tidak memihak terhadap kelas tertentu. Dengan demikian, hasil ini menegaskan bahwa dalam konteks pemodelan dengan data terbatas. Model CNN terbukti cukup baik dalam memproses data teks pendek seperti tweet, selama distribusi data dijaga dengan baik dan sesuai dengan konteks klasifikasi yang ditargetkan. Secara keseluruhan, kedua konfigurasi tersebut menegaskan pentingnya pemilihan dan pengaturan proporsi data tambahan untuk meningkatkan kemampuan generalisasi model, khususnya pada dataset dengan jumlah terbatas.

3.5. Pengujian Pada Data Test

Setelah tahap Eksperimen Setup diselesaikan, langkah selanjutnya adalah pengujian menggunakan 924 data uji berupa tweet. Evaluasi dilakukan terhadap model klasifikasi berbasis *Convolutional Neural Network* (CNN) untuk mengukur akurasi dalam klasifikasi sentimen, yang dinilai melalui sistem *leaderboard*. Model yang digunakan merupakan model dengan performa terbaik berdasarkan hasil pada Tabel 7. Pengujian pertama menggunakan dataset Kaesang V1 ditambah



V2. Pengujian kedua menggabungkan dataset Kaesang V1 dan V2 dengan data *Open Topic* yang memiliki distribusi kelas positif sebanyak 400, negatif 400, dan netral 400. Pengujian ketiga mengombinasikan dataset Kaesang V1 dan V2 dengan data Covid-19 dengan distribusi kelas positif 450, negatif 650, dan netral 1100. Hasil evaluasi dari ketiga pengujian tersebut ditampilkan pada tabel 8.

Tabel 8. Pengujian Data Test

ID Run	Data Validasi		Data Test	
	F1-Score	Akurasi	F1-Score	Akurasi
R1	0.60	0.59	0.43	0.48
R6	0.60	0.61	0.48	0.55
R4	0.62	0.62	0.51	0.61

Tabel 8 memperlihatkan hasil evaluasi kinerja model pada data uji yang terdiri dari 924 tweet. Model yang dilatih hanya dengan dataset Kaesang V1 dan V2 menghasilkan *F1-Score* sebesar 0,43 dan akurasi 0,48 pada pengujian pertama dengan *ID Run* R1, menunjukkan kemampuan generalisasi yang rendah karena hanya menggunakan data dari satu topik.

Pada pengujian kedua dengan *ID Run* R6, model menggunakan kombinasi dataset Kaesang V1 dan V2 bersama data *Open Topic* dengan distribusi kelas distribusi kelas positif sebanyak 400, negatif 400, dan netral 400. Hasilnya menunjukkan peningkatan performa, yakni *F1-Score* sebesar 0.48 dan *accuracy* 0.55. Hal ini mengindikasikan bahwa keberagaman topik pada data pelatihan dapat membantu model dalam mengenali beragam pola ekspresi sentimen pada data uji.

Pengujian ketiga dengan *ID Run* R4, menunjukkan hasil terbaik dengan *F1-Score* sebesar 0.51 dan *accuracy* 0.61. Model ini dilatih dengan menggabungkan dataset Kaesang V1 dan V2 serta data Covid-19 yang memiliki dengan distribusi kelas positif 450, negatif 650, dan netral 1100. Peningkatan jumlah serta kompleksitas data pelatihan berkontribusi baik terhadap kemampuan model dalam mengklasifikasikan sentimen secara lebih tepat pada data uji. Hasil pada Tabel 8 menegaskan bahwa penambahan data eksternal ke dalam dataset yang terbatas berdampak signifikan terhadap peningkatan performa model dalam klasifikasi sentimen.

3.6. Diskusi

Setelah proses pengujian dilakukan terhadap data uji melalui sistem *leaderboard*, hasil yang diperoleh kemudian dibandingkan dengan beberapa eksperimen lain yang juga diikutsertakan dalam sistem tersebut. Evaluasi kinerja model dilakukan dengan menghitung metrik *F1-Score*, *accuracy*, *precision*, dan *recall*. Di antara metrik tersebut, *F1-Score* merupakan acuan utama untuk mengevaluasi performa model, karena metrik ini merupakan gabungan dari *precision* dan *recall*, sehingga mampu memberikan penilaian yang seimbang terhadap kedua aspek tersebut. Perbandingan hasil dari beberapa penelitian ditampilkan pada Tabel 9.

Tabel 9. Perbandingan Hasil *Leaderboard*

Tim	Metode	F1-Score	Akurasi	Precision	Recall
Ridho Ilahi[13]	BidirectionalLSTM	0.59	0.67	0.58	0.66
	IndoBERT				
Joni Pranata[20]	IndoBERT dan RF	0.54	0.63	0.58	0.63
Penelitian Ini	CNN	0.51	0.61	0.53	0.56
Yoga El Saputra[8]	SVM TF-IDF	0.51	0.61	0.52	0.59
Muhammad Ravil[6]	Naive Bayes PSO	0.50	0.59	0.52	0.58
Admin	Baseline	0.40	0.45	0.49	0.48

Hasil evaluasi menunjukkan bahwa model berbasis Bidirectional LSTM dan IndoBERT memperoleh performa tertinggi dengan *F1-Score* sebesar 0.59 dan akurasi 0.67, menandakan efektivitas embedding kontekstual dalam memahami teks pendek. Model lain yang juga menggunakan IndoBERT, namun dikombinasikan dengan Random Forest, mencatat *F1-Score* 0.54 dan akurasi 0.63. Model CNN dalam penelitian ini menghasilkan *F1-Score* 0.51 dan akurasi 0.61, setara dengan model berbasis SVM dan TF-IDF. Meskipun nilai *recall* model SVM sedikit lebih tinggi (0.59 vs 0.56), CNN memiliki *precision* yang lebih baik, menunjukkan keseimbangan performa. Sementara itu, model *Naive Bayes* dengan optimasi PSO menghasilkan *F1-Score* 0.50, masih kompetitif namun sedikit di bawah CNN dan SVM. Secara keseluruhan, model yang memanfaatkan *embedding* kontekstual dari *transformer-based models* seperti IndoBERT menunjukkan hasil terbaik, namun CNN tetap menjadi alternatif efisien yang kompetitif dalam kondisi keterbatasan data.

4. KESIMPULAN

Penelitian ini bertujuan untuk mengevaluasi sejauh mana algoritma *Convolutional Neural Network* (CNN) efektif dalam mengklasifikasikan sentimen pada kondisi data yang terbatas, dengan studi kasus persepsi publik terhadap Kaesang Pangarep sebagai Ketua Umum Partai Solidaritas Indonesia (PSI). Dataset utama terdiri dari 300 tweet yang telah dikategorikan ke dalam tiga label sentimen: positif, negatif, dan netral yang kemudian di gunakan sebagai eksperimen dasar dengan hasil evaluasi *F1-score* sebesar 0,60, *accuracy* 0,59, *precision* 0,63, dan *recall* 0,60. Nilai ini menjadi acuan



dasar untuk membandingkan peningkatan performa model, pada eksperimen setup dengan data eksternal bertopik Covid-19 dan *Open Topic*. Hasil eksperimen menunjukkan bahwa penambahan data eksternal mampu meningkatkan performa model, baik pada data validasi maupun data uji. Konfigurasi terbaik dicapai pada kombinasi data Kaesang dengan data Covid-19, terutama saat distribusi kelas diperbesar hingga 450 data positif, 650 data negatif, dan 1100 data netral. Konfigurasi ini menghasilkan *F1-score* sebesar 0,62 pada data validasi dan 0,51 pada data uji. Strategi perluasan data melalui penambahan dataset eksternal yang relevan dan proporsional sangat penting dalam konteks pelatihan model dengan sumber data terbatas selama *preprocessing text* dilakukan dengan benar dan data pelatihan dikelola secara optimal. Hasil evaluasi tertinggi berdasarkan perbandingan *Leaderboard* menunjukkan bahwa model berbasis *Bidirectional LSTM* dan *IndoBERT* memperoleh performa tertinggi dengan *F1-Score* sebesar 0.59 dan akurasi 0.67. Sehingga, mendapatkan gambaran umum bahwa model yang menggunakan *embedding* dari *transformer-based*, seperti *IndoBERT*, berhasil mencapai performa tertinggi. Meskipun demikian, CNN tetap menjadi pilihan yang efisien dan kompetitif, khususnya ketika menghadapi keterbatasan data. Oleh karena itu, penelitian selanjutnya disarankan untuk mengeksplorasi arsitektur *deep learning* yang lebih kompleks seperti GRU atau yang lainnya, serta mempertimbangkan penggunaan metode *embedding* berbasis konteks seperti *IndoRoBERTa* untuk mengevaluasi pengaruh representasi kata statis dan kontekstual terhadap performa klasifikasi sentimen.

REFERENCES

- [1] T. S. Sabrila, V. R. Sari, and A. E. Minarno, "Analisis Sentimen Pada Tweet Tentang Penanganan Covid-19 Menggunakan Word Embedding Pada Algoritma Support Vector Machine Dan K-Nearest Neighbor," *Fountain of Informatics Journal*, vol. 6, no. 2, p. 69, Jul. 2021, doi: 10.21111/fij.v6i2.5536.
- [2] T. Spinde, E. Richter, M. Wessel, J. Kulshrestha, and K. Donnay, "What do Twitter comments tell about news article bias? Assessing the impact of news article bias on its perception on Twitter," *Online Soc Netw Media*, vol. 37–38, Sep. 2023, doi: 10.1016/j.osnem.2023.100264.
- [3] N. Rohman, "Peran Partai Solidaritas Indonesia (PSI) dalam Pemilihan Presiden 2024: Analisis Terhadap Pemilih Pemula," *JPW (Jurnal Politik Walisongo)*, vol. 5, no. 1, pp. 85–102, 2023, doi: 10.21580/jpw.v5i1.18330.
- [4] K. Tsabitah and I. Suryawati, "Analisis Wacana Kritis Pidato Pertama Kaesang Pangarep sebagai Ketua Umum Partai Solidaritas Indonesia," *CARAKA : Indonesia Journal of Communication*, vol. 5, no. 1, pp. 27–38, 2024, doi: 10.25008/caraka.v5i1.109.
- [5] S. Agustian *et al.*, "Arah Baru Penelitian Klasifikasi Teks: Memaksimalkan Kinerja Klasifikasi Sentimen dari Data Terbatas," *arXiv preprint*, pp. 1–9, Jul. 2024, doi: <https://doi.org/10.48550/arXiv.2407.05627>.
- [6] M. Ravil, S. Agustian, M. Fikry, and F. Insani, "Peningkatan Performa Klasifikasi Sentimen Tweet Kaesang Menggunakan Naïve Bayes dengan PSO pada Dataset Kecil," *KLIK: Kajian Ilmiah Informatika dan Komputer*, vol. 4, no. 6, pp. 2909–2917, 2024, doi: 10.30865/klik.v4i6.1939.
- [7] P. Zahwa, S. Agustian, F. Yanto, and S. Baru, "Implementasi Bi-Directional Long Short Term Memory Terhadap Klasifikasi Sentimen di Twitter Pada Dataset Terbatas," *ZONasi: Jurnal Sistem Informasi*, vol. 7, no. 1, pp. 11–24, 2023, doi: <https://doi.org/10.31849>.
- [8] Y. El Saputra, S. Agustian, and S. Ramadhani, "Klasifikasi Sentimen SVM Dengan Dataset yang Kecil Pada Kasus Kaesang Sebagai Ketua Umum PSI," *KLIK: Kajian Ilmiah Informatika dan Komputer*, vol. 4, no. 6, pp. 2902–2908, 2024, doi: 10.30865/klik.v4i6.1944.
- [9] S. N. Listyarini and D. A. Anggoro, "Analisis Sentimen Pilkada di Tengah Pandemi Covid-19 Menggunakan Convolution Neural Network (CNN)," *Jurnal Pendidikan dan Teknologi Indonesia*, vol. 1, no. 7, pp. 261–268, Jul. 2021, doi: 10.52436/1.jpti.60.
- [10] B. A. Yuniarossy *et al.*, "Analisis Sentimen Terhadap Isu Feminisme di Twitter Menggunakan Model Convolutional Neural Network (Cnn)," *Lebesgue: Jurnal Ilmiah Pendidikan Matematika, Matematika dan Statistika*, vol. 5, no. 1, 2024, doi: 10.46306/lb.v5i1.
- [11] F. A. Irawan and D. A. Rochmah, "Penerapan Algoritma CNN Untuk Mengetahui Sentimen Masyarakat Terhadap Kebijakan Vaksin Covid-19," *JURNAL INFORMATIKA*, vol. 9, no. 2, 2022, doi: <https://doi.org/10.31294>.
- [12] N. Satya Marga, A. Rahman Isnain, and D. Alita, "Sentimen Analisis Tentang Kebijakan Pemerintah Terhadap Kasus Corona Menggunakan Metode Naive Bayes," *Jurnal Informatika dan Rekayasa Perangkat Lunak (JATIKA)*, vol. 4, no. 4, pp. 453–463, 2021, doi: <https://doi.org/10.33365>.
- [13] R. Illahi, S. Agustian, Jasril, and F. Yanto, "Klasifikasi Sentimen Menggunakan Bidirectional Lstm dan Indobert Dengan Dataset Terbatas," *ZONasi: Jurnal Sistem Informasi*, vol. 7, no. 1, 2025, doi: <https://doi.org/10.31849>.
- [14] L. Ashbaugh and Y. Zhang, "A Comparative Study of Sentiment Analysis on Customer Reviews Using Machine Learning and Deep Learning," *Computers*, vol. 13, no. 12, Dec. 2024, doi: 10.3390/computers13120340.
- [15] E. Setia Budi, A. Nofriyaldi Chan, P. Priscillia Alda, and M. Arif Fauzi Idris, "Optimasi Model Machine Learning untuk Klasifikasi dan Prediksi Citra Menggunakan Algoritma Convolutional Neural Network," *RESOLUSI: Rekayasa Teknik Informatika dan Informasi*, vol. 4, no. 5, pp. 502–509, 2024, doi: <https://doi.org/10.30865>.
- [16] M. Haikal and U. Hayati, "Analisis Sentimen Terhadap Penggunaan Aplikasi Game Online Pubg Mobile Menggunakan Algoritma Naive Bayes," *Jurnal Mahasiswa Teknik Informatika*, vol. 7, no. 6, 2023, doi: <https://doi.org/10.36040>.
- [17] I. Habib Kusuma and N. Cahyono, "Analisis Sentimen Masyarakat Terhadap Penggunaan E-Commerce Menggunakan Algoritma K-Nearest Neighbor," *Jurnal pengembangan IT (JPIT)*, vol. 8, no. 3, 2023, doi: <https://doi.org/10.30591>.
- [18] S. BAYAT and G. İŞIK, "Evaluating the Effectiveness of Different Machine Learning Approaches for Sentiment Classification," *Iğdır Üniversitesi Fen Bilimleri Enstitüsü Dergisi*, vol. 13, no. 3, pp. 1496–1510, Sep. 2023, doi: 10.21597/jist.1292050.
- [19] P. Yohana, S. Agustian, and S. Kurnia Gusti, "Klasifikasi Sentimen Masyarakat terhadap Kebijakan Vaksin Covid-19 pada Twitter dengan Imbalance Classes Menggunakan Naive Bayes," *Seminar Nasional Teknologi Informasi, Komunikasi dan Industri (SNTIKI)*, vol. 26, pp. 69–80, Oct. 2022, [Online]. Available: <https://ejournal.uin-suska.ac.id/index.php/SNTIKI/article/view/19012>



- [20] J. Pranata, S. Agustian, J. Jasril, and E. Haerani, “Penggunaan Model Bahasa indoBERT pada metode Random Forest untuk Klasifikasi Sentimen dengan Dataset Terbatas,” *Building of Informatics, Technology and Science (BITS)*, vol. 6, no. 3, pp. 1668–1676, Dec. 2024, doi: 10.47065/bits.v6i3.6335.