



Analisis Sentimen tentang Penundaan Pengangkatan CPNS 2025 pada Platform X Menggunakan Metode IndoBERT

Muh. Hasbi Asshiddiq^{1*}, Arita Witanti²

^{1,2} Informatika, Universitas Mercu Buana Yogyakarta, Indonesia
Gg. Jemb. Merah No.84C, Soropadan, Condongcatur
E-Mail *: 211110051@student.mercubuana-yogya.ac.id

ABSTRACT

The postponement of the 2025 Civil Servant Candidate (CPNS) appointment drew significant public attention in Indonesia, triggering diverse reactions on social media. This study aims to analyze public sentiment toward the delay and evaluate the performance of the IndoBERT model in classifying Indonesian-language opinions on the issue. The research adopts a descriptive quantitative approach using data mining and Natural Language Processing (NLP) methods. Data collection begins by crawling tweets from Platform X using the Tweet Harvest tool, which is based on auth tokens and keywords related to the CPNS 2025 postponement issue. The data is then processed through pre-processing stages including cleaning, case folding, tokenization, and normalization. Subsequent steps include filtering, refinement, manual labeling, and splitting the data into three sets with a ratio of 80:10:10 for training, validation, and testing. The model is developed using the Indobert-base-p2 transformer, fine-tuned with the Adam optimizer and optimal configuration. Model performance is evaluated using accuracy, precision, recall, and F1-score metrics. Out of 3,079 collected tweets, 2,479 were balanced in terms of positive, negative, and neutral sentiment. The IndoBERT model achieved an accuracy of 84.27%, with an average precision, recall, and F1-score of around 84%. Sentiment analysis showed that negative reactions dominated early on but shifted positively following government clarification. Neutral sentiment mostly came from official and media accounts. These findings underscore social media's role in real-time opinion monitoring and the potential of NLP-based sentiment analysis for evaluating public communication and informing policy decisions.

Keywords: CPNS 2025, IndoBERT, Platform X, Sentiment Analysis

ABSTRAK

Penundaan pengangkatan Calon Pegawai Negeri Sipil (CPNS) 2025 menjadi isu yang menarik perhatian publik Indonesia dan memicu beragam reaksi, terutama di media sosial. Penelitian ini bertujuan menganalisis sentimen publik terhadap penundaan tersebut sekaligus mengevaluasi performa model IndoBERT dalam mengklasifikasikan opini masyarakat berbahasa Indonesia terkait isu ini. Penelitian menggunakan pendekatan kuantitatif deskriptif dengan metode data mining dan *Natural Language Processing* (NLP). Pengumpulan data dimulai dengan *crawling* data tweet dari Platform X menggunakan tools *Tweet Harvest* berbasis *auth token* dan kata kunci terkait isu penundaan CPNS 2025. Data kemudian diproses melalui tahap *pre-processing* yang meliputi *cleaning*, *case folding*, tokenisasi, dan normalisasi. Selanjutnya dilakukan *filtering*, *refinement*, pelabelan manual, serta pembagian data menjadi tiga set dengan rasio 80:10:10 untuk pelatihan, validasi, dan pengujian. Pemodelan dilakukan dengan menggunakan *transformer* Indobert-base-p2 yang di-*fine-tune* dengan *optimizer* Adam dan sesuai konfigurasi optimal. Evaluasi performa model dilakukan menggunakan metrik akurasi, presisi, *recall*, dan *F1-score*. Dari 3.079 tweet yang terkumpul, diperoleh 2.479 data seimbang dari segi sentimen positif, negatif, dan netral. Model IndoBERT berhasil mencapai akurasi sebesar 84,27% dengan rata-rata presisi, *recall*, dan *F1-score* sekitar 84%. Analisis sentimen menunjukkan dominasi sentimen negatif pada awal isu, yang kemudian berangsur berubah menjadi positif setelah klarifikasi dari pemerintah. Sentimen netral banyak berasal dari akun resmi dan media. Temuan ini menegaskan pentingnya media sosial sebagai sumber pemantauan opini publik secara *real-time* serta menunjukkan potensi besar analisis sentimen berbasis NLP sebagai alat evaluasi komunikasi publik dan pengambilan keputusan kebijakan pemerintah.

Kata Kunci : Analisis Sentimen, CPNS 2025, IndoBERT, Platform X

Naskah diterima 04 Juni 2025; Revisi 12 Juni 2025; Diterima 30 Juni 2025. Tanggal Publikasi 01 September 2025
Jurnal teknika berada pada lisensi *Creative Commons Attribution-ShareAlike 4.0 International License*
DOI: 10.30736/jt.v17i2.1448_Hal 97-108



1. PENDAHULUAN

Penundaan pengangkatan Calon Pegawai Negeri Sipil (CPNS) 2025 menjadi isu yang menarik perhatian masyarakat Indonesia. CPNS adalah tahapan awal sebelum seseorang resmi diangkat menjadi Pegawai Negeri Sipil (PNS), yang bertugas mendukung pelayanan publik dan pemerintahan (Zuhri et al., 2023). Pada tahun 2025, terjadi penundaan pengangkatan CPNS hasil seleksi tahun 2024 hingga Oktober 2025. Keputusan ini diumumkan oleh Kementerian Pendayagunaan Aparatur Negara dan Reformasi Birokrasi (PANRB), dan berdampak langsung terhadap peserta seleksi yang telah lulus namun belum diangkat secara resmi (Yulianti, 2025). Beberapa penyebab penundaan ini diantaranya adalah perlunya penyesuaian data formasi dan penempatan, penyeragaman tanggal mulai tugas (TMT), serta penyelesaian dan tenaga non-ASN (Kristi, 2024).

Penundaan ini memicu beragam reaksi dari masyarakat, terutama para peserta seleksi yang terdampak. Banyak dari mereka yang telah mengundurkan diri dari pekerjaan sebelumnya karena berharap segera diangkat sebagai CPNS, sehingga mengalami ketidakpastian ekonomi (Magfirah, 2025). Tak sedikit pula masyarakat yang mengungkapkan opini, keluhan, atau dukungan terhadap kebijakan ini melalui media sosial. Untuk mengetahui dan memahami reaksi tersebut secara lebih objektif, dapat dilakukan analisis sentimen guna memahami persepsi masyarakat secara umum terhadap kebijakan penundaan pengangkatan CPNS ini. Analisis sentimen dapat mengelompokkan opini menjadi positif, negatif, atau netral, dan membantu pemerintah serta pemangku kebijakan untuk merespons keresahan publik secara lebih tepat (Arham et al., 2022).

Data komentar masyarakat dapat dikumpulkan dari berbagai sosial media salah satunya dari Platform X (sebelumnya dikenal sebagai Twitter) yang dipilih untuk analisis sentimen karena berbagai keunggulan yang dimilikinya dibandingkan media sosial lain. Menurut Statista, pengguna aktif platform X di Indonesia mencapai 24,85 juta, menjadikannya negara dengan jumlah pengguna terbesar ke-4 secara global (Statista, 2024). Sebagai salah satu media sosial utama di Indonesia, platform ini menawarkan basis pengguna yang luas dan beragam, mencerminkan opini masyarakat dari berbagai latar belakang (Putriyetti et al., 2025). Aktivitas diskusi publik di platform X cenderung aktif, terutama dalam isu-isu terkini, sehingga menjadikannya media ideal untuk mengamati reaksi masyarakat terhadap penundaan pengangkatan CPNS. Fitur seperti hashtags, komentar, dan retweets mendukung pengumpulan data secara *real-time* dan terstruktur (Talalu & Valentine, 2021). Hal ini memperkuat relevansi platform X dalam analisis opini publik.

Untuk menganalisis data teks dari Platform X, berbagai metode pemrosesan bahasa alami (*Natural Language Processing/NLP*) yang merupakan cabang dari *Artificial Intelligence* (AI) telah digunakan.

Metode tersebut mencakup pendekatan berbasis *machine learning* tradisional (Putri et al., 2022), hingga teknik *deep learning* modern (Putra & Juanita, 2021).

Salah satu pendekatan terbaru yang terbukti efektif dalam konteks Bahasa Indonesia adalah IndoBERT (Hakim et al., 2024), sebuah model *transformer* berbasis arsitektur BERT (*Bidirectional Encoder Representations from Transformers*) yang dikembangkan oleh (Wilie et al., 2020). IndoBERT dilatih secara khusus menggunakan korpus Bahasa Indonesia yang terdiri dari Wikipedia, berita daring (Kompas, Detik), dan data media sosial. Model ini mengadopsi arsitektur BERT-base dengan 12 lapisan *encoder*, 768 dimensi *hidden state*, 12 *self-attention heads*, dan sekitar 110 juta parameter (Wilie et al., 2020).

Keunggulan utama IndoBERT terletak pada kemampuannya menangkap konteks kalimat secara dua arah (*bidirectional*) serta memahami struktur dan nuansa emosional dalam teks Bahasa Indonesia. Hal ini membuatnya unggul dalam berbagai tugas NLP seperti klasifikasi teks, *Named Entity Recognition* (NER), dan analisis sentimen, serta menunjukkan performa lebih baik dibandingkan model NLP tradisional (Koto et al., 2020).

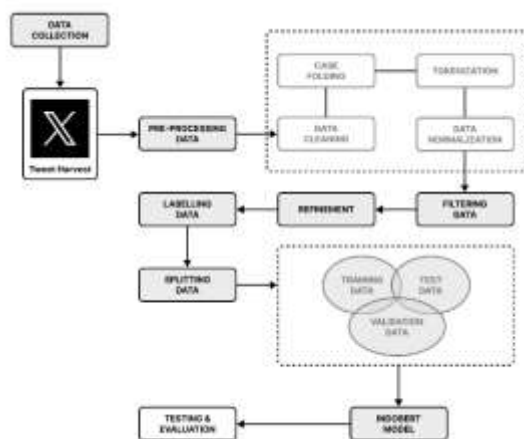
Sebagai bagian dari penerapan AI dalam bidang NLP, penggunaan IndoBERT pada teks berbahasa Indonesia telah memberikan peningkatan signifikan dalam akurasi analisis sentimen dibandingkan dengan metode seperti *Logistic Regression* dan *Long Short-Term Memory* (LSTM) (Khotimah & Sarno, 2019). Sebuah penelitian juga mencatat bahwa IndoBERT dapat mengurangi ambiguitas yang sering terjadi pada analisis sentimen dalam bahasa Indonesia karena pemahaman konteks yang lebih mendalam (Singgalen, 2025). Oleh karena itu, penggunaan IndoBERT dalam analisis sentimen terhadap data dari Platform X diharapkan dapat memberikan hasil yang lebih akurat dan representatif terhadap kondisi sebenarnya di masyarakat.

Penelitian ini akan berfokus pada analisis sentimen terhadap data teks yang dikumpulkan dari media sosial Platform X, guna mengetahui bagaimana respons publik terhadap kebijakan penundaan pengangkatan CPNS 2025. Penelitian ini bertujuan untuk mengidentifikasi kecenderungan sentimen masyarakat baik positif, negatif, maupun netral terhadap kebijakan yang dinilai menunda hak publik dan menjadi isu hangat dalam perbincangan daring. Selain menganalisis sentimen, penelitian ini juga akan mengevaluasi efektivitas model IndoBERT dalam memahami dan mengklasifikasikan opini publik berbahasa Indonesia dalam konteks isu kebijakan sosial yang kompleks. Pendekatan ini merupakan bagian dari penerapan *Artificial Intelligence* (AI) dalam bidang pemrosesan bahasa alami (NLP), yang memungkinkan analisis teks secara lebih kontekstual dan akurat.

2. METODE

2.1 Jenis dan Pendekatan Penelitian

Penelitian ini merupakan penelitian kuantitatif dengan pendekatan deskriptif yang memanfaatkan metode *data mining* dan *Natural Language Processing* (NLP) untuk menganalisis sentimen masyarakat terhadap penundaan pengangkatan CPNS 2025. Tujuan dari pendekatan ini adalah untuk memperoleh pemahaman mendalam mengenai kecenderungan opini publik melalui data teks yang diperoleh dari media sosial. Dalam hal ini, analisis dilakukan berdasarkan klasifikasi sentimen ke dalam tiga kategori utama: positif, negatif, dan netral, dengan memanfaatkan model IndoBERT sebagai alat klasifikasi berbasis kecerdasan buatan. Arsitektur sistem dalam penelitian ini terdiri atas beberapa tahapan utama yang ditampilkan dalam Diagram Alur pada Gambar 1.



Gambar 1. Diagram Alur Penelitian

2.2 Sumber dan Teknik Pengumpulan Data

Sumber data dalam penelitian ini berasal dari media sosial Platform X. Pengumpulan data dilakukan selama periode 3 Maret 2025 hingga 16 Mei 2025, dengan menggunakan kata kunci: "cpns 2025" dan "penundaan cpns". Teknik pengumpulan data dilakukan dengan *crawling* data menggunakan tools Tweet Harvest, yang memungkinkan ekstraksi tweet berdasarkan kata kunci dan rentang waktu tertentu tanpa perlu menggunakan API resmi Twitter. Tools ini memungkinkan pengguna mengunduh tweet secara massal dalam bentuk file terstruktur.

Untuk memulai proses pengambilan data, peneliti terlebih dahulu memasukkan *auth token* pribadi yang diperoleh dari akun Twitter masing-masing, yang berfungsi sebagai otorisasi untuk mengakses data publik di platform tersebut. Setelah otorisasi berhasil, langkah selanjutnya adalah menentukan kata kunci atau *hashtag* yang relevan dengan topik penelitian, yaitu "CPNS 2025" dan "penundaan CPNS", yang sering digunakan oleh pengguna Twitter saat membahas isu ini. Tweet yang sesuai kemudian dikumpulkan secara maksimal, mencakup teks utama serta metadata pendukung seperti waktu unggahan, jumlah retweet, jumlah suka,

dan sebagainya. Seluruh data disimpan dalam format .csv agar mudah diolah pada tahap selanjutnya.

Data yang diperoleh berisi teks-teks berbahasa Indonesia yang mencerminkan opini, reaksi, dan sentimen publik terhadap isu penundaan pengangkatan CPNS. Dalam konteks pemrosesan bahasa alami (NLP), data ini menjadi fokus utama untuk dianalisis. Sebelum masuk ke tahap pemodelan, data mentah ini akan diproses melalui tahap *pre-processing* yang mencakup *filtering*, *case folding*, *cleansing*, tokenisasi, hingga normalisasi. Proses ini penting untuk meningkatkan kualitas data dan memastikan hasil analisis sentimen yang lebih akurat.

2.3 Pre-Processing Data

Pre-processing merupakan tahap penting yang bertujuan untuk membersihkan dan mempersiapkan data mentah sebelum masuk ke proses klasifikasi. Pada tahap ini, dilakukan serangkaian proses pembersihan dan transformasi data untuk meningkatkan kualitas serta konsistensi data yang akan digunakan dalam analisis sentimen. Tahapan *pre-processing* dalam penelitian ini adalah sebagai berikut.

1) Cleaning Data

Pembersihan data (*cleaning*) merupakan proses menghilangkan elemen-elemen yang tidak relevan dari teks, seperti karakter spesial, emotikon, angka, tanda baca, dan URL. Elemen-elemen ini umumnya tidak memiliki kontribusi terhadap pemahaman semantik dalam analisis sentimen. *Cleaning* bertujuan untuk meningkatkan kualitas data dengan cara menyederhanakan input teks tanpa kehilangan makna utama. Contoh data pada proses *cleaning* dapat dilihat pada Tabel 1.

Tabel 1. *Cleaning Data*

No	Full_Text	Clean
1.	Harus demo dulukah biar penundaan pengangkatan cpns ini dibatalkan? @Gerindra @kempnrb @BKNgoid	Harus demo dulukah biar penundaan pengangkatan cpns ini dibatalkan
2.	Kuat2 buat sahabat aku yg uda resign dan udah nikahin pacarnya krn sesuai mimpi mereka mereka udah lulus cpns di instansi	Kuat buat sahabat aku yg uda resign dan udah nikahin pacarnya krn sesuai mimpi mereka mereka udah lulus cpns di instansi
3.	@dekade_08 Negara menjelang Bangkrut ditandai dgn penundaan pengangkatan CPNS	Negara menjelang Bangkrut ditandai dgn penundaan pengangkatan CPNS

2) Case folding

Case folding adalah proses mengubah seluruh huruf dalam teks menjadi huruf kecil (*lowercase*). Proses ini penting karena sistem komputer pada umumnya membedakan huruf besar dan kecil, sehingga kata seperti “CPNS” dan “cpns” akan dianggap berbeda. Dengan melakukan case folding, keseragaman representasi kata dapat dicapai, sehingga menghindari redundansi dalam fitur. Contoh data pada proses *case folding* dapat dilihat pada Tabel 2.

Tabel 2. Case Folding

No	Clean	Case Folding
1.	Harus demo dulukah biar penundaan pengangkatan cpns ini dibatalkan	harus demo dulukah biar penundaan pengangkatan cpns ini dibatalkan
2.	Kuat buat sahabat aku yg uda resign dan udah nikahin pacarnya krn sesuai mimpi mereka mereka udah lulus cpns di instansi	kuat buat sahabat aku yg uda resign dan udah nikahin pacarnya krn sesuai mimpi mereka mereka udah lulus cpns di instansi
3.	Negara menjelang Bangkrut ditandai dgn penundaan pengangkatan CPNS	negara menjelang bangkrut ditandai dgn penundaan pengangkatan cpns

3) Tokenization

Tokenisasi adalah proses memecah teks menjadi unit-unit kecil yang disebut token, biasanya berupa kata atau frasa. Proses ini sangat penting dalam NLP karena banyak algoritma pemrosesan bahasa membutuhkan input dalam bentuk token agar dapat dianalisis secara numerik. Tokenisasi memungkinkan setiap elemen bahasa diidentifikasi secara individual untuk dianalisis lebih lanjut. Contoh data pada proses *tokenization* dapat dilihat pada Tabel 3.

Tabel 3. Tokenization

No	Case Folding	Token
1.	harus demo dulukah biar penundaan pengangkatan cpns ini dibatalkan	harus,demo,dulukah, biar,penundaan, pengangkatan,cpns,ini, dibatalkan
2.	kuat buat sahabat aku yg uda resign dan udah nikahin pacarnya krn sesuai mimpi mereka mereka udah lulus cpns di instansi	kuat,buat,sahabat,aku, yg,uda,resign,dan,udah ,nikahin,pacarnya,krn, sesuai,mimpi,mereka, mereka,udah,lulus,cpn s, di,instansi
3.	negara menjelang bangkrut ditandai dgn penundaan pengangkatan cpns	negara,menjelang, bangkrut,ditandai,dgn , penundaan, pengangkatan,cpns

4) Normalization

Normalisasi teks bertujuan untuk menyamakan variasi bentuk kata yang memiliki arti sama. Dalam bahasa Indonesia, pengguna media sosial sering menggunakan singkatan atau bentuk tidak baku seperti “gm” untuk “bagaimana”, “tdk” untuk “tidak”, dan sebagainya. Proses ini dilakukan dengan merujuk pada kamus normalisasi yang berisi pasangan kata tidak baku dan bentuk bakunya. Normalisasi diperlukan untuk memastikan bahwa variasi penulisan tidak menghasilkan makna berbeda saat dilakukan analisis. Contoh data pada proses *normalization* dapat dilihat pada Tabel 4.

Tabel 4. Normalization

No	Token	Normalization
1.	harus,demo,dulukah, h, biar,penundaan, pengangkatan,cpns ,ini, dibatalkan	harus,demo,dulukah,b iar,penundaan,pengan gkatan,cpns,ini,dibatal kan
2.	kuat,buat,sahabat,a ku,yg,uda,resign,d an,udah,nikahin,pa carnya,krn,sesuai, mimpi,mereka, mereka,udah,lulus, cpns,di,instansi	kuat,buat,sahabat,aku, yang,sudah,resign,dan ,sudah,nikahin,pacarn ya,karena,sesuai,mim pi,mereka,mereka,sud ah,lulus,cpns,di, instansi
3.	negara,menjelang, bangkrut,ditandai,d gn, penundaan, pengangkatan,cpns	negara,menjelang,ba ngkrut,ditandai, dengan, penundaan,pengangk atan,cpns

Semua proses ini dilakukan dengan bantuan pustaka Python seperti Pandas untuk manipulasi data dan NLTK untuk proses tokenisasi. Dengan tahapan *pre-processing* yang menyeluruh, data menjadi lebih bersih, terstruktur, dan representatif untuk digunakan dalam pemodelan dan analisis sentimen.

2.4 Filtering Data

Filtering adalah proses yang bertujuan untuk menyeleksi data agar hanya menyisakan tweet yang relevan dengan tujuan analisis. Dalam konteks ini, tweet yang bersifat spam, promosi, tidak mengandung teks (kosong), atau bersifat duplikat dihapus dari dataset. Tujuan dari filtering adalah untuk menghindari distorsi hasil analisis akibat adanya *noise* dalam data. *Filtering* juga dapat meningkatkan efisiensi pemrosesan karena hanya data yang bernilai informatif yang digunakan.

2.5 Refinement

Refinement merupakan tahapan tambahan yang dilakukan setelah proses *pre-processing* untuk memastikan bahwa data yang telah dibersihkan benar-benar layak digunakan dalam proses pelabelan dan pemodelan. Pada tahap ini dilakukan koreksi manual terhadap hasil *pre-processing*, terutama untuk menangani kasus-kasus ambiguitas atau kesalahan konversi yang mungkin tidak terdeteksi oleh proses

otomatis sebelumnya. Misalnya, terdapat kata-kata slang atau istilah lokal yang tidak dikenali dalam proses normalisasi, sehingga perlu penyesuaian secara manual berdasarkan konteks penggunaannya. *Refinement* bertujuan untuk meningkatkan kualitas dan akurasi data teks, dengan mempertimbangkan konteks linguistik dan semantik dari setiap *tweet*.

2.6 Labelling Data

Labelling data merupakan tahap penting dalam pemodelan klasifikasi, di mana setiap *tweet* diberikan label sentimen berdasarkan isi atau makna yang terkandung dalam teks. Dalam penelitian ini, proses pelabelan dilakukan secara manual oleh peneliti, dengan mempertimbangkan makna eksplisit maupun implisit dari *tweet*. Label yang digunakan terdiri dari tiga kategori sentimen, yaitu: positif, negatif, dan netral.

Untuk mendukung proses pelabelan yang lebih akurat dan konsisten, peneliti menggunakan bantuan dari ChatGPT sebagai alat bantu dalam mempertimbangkan konteks semantik dan memberikan saran klasifikasi apabila terdapat keraguan terhadap sentimen suatu *tweet*. Penggunaan ChatGPT bukan sebagai alat pelabel otomatis, tetapi sebagai *decision support tool* untuk memperkuat validitas keputusan yang diambil secara manual. Seluruh data yang telah diberi label ini kemudian digunakan sebagai dataset terannotasi untuk proses pelatihan dan pengujian model klasifikasi sentimen pada tahap selanjutnya.

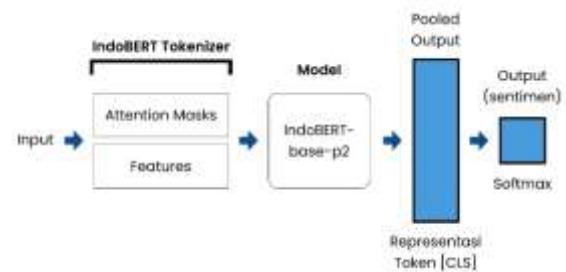
2.7 IndoBERT Model

Tahap pemodelan merupakan inti dari proses analisis dalam penelitian ini, di mana opini masyarakat terhadap isu penundaan CPNS 2025 dianalisis menggunakan model berbasis transformer, yaitu IndoBERT. Model yang digunakan adalah *indobert-base-p2*, salah satu varian dari IndoBERT yang dikembangkan oleh IndoNLU sebagai bagian dari inisiatif IndoBenchmark untuk evaluasi dan pengembangan model NLP berbahasa Indonesia.

Sebelum dilakukan pelatihan, label sentimen yang terdiri dari tiga kategori: negatif, netral, dan positif, dikonversi ke dalam bentuk numerik, masing-masing direpresentasikan sebagai 0, 1, dan 2. Selanjutnya, dataset yang telah di *pre-processing* dibagi menjadi tiga bagian, yaitu data latih (*training*), data validasi (*validation*), dan data uji (*testing*). Data latih digunakan untuk membangun model, data validasi berfungsi untuk mengatur dan memantau performa model selama pelatihan serta mencegah *overfitting*, sementara data uji digunakan untuk mengevaluasi performa akhir model secara objektif.

Pembagian dataset dilakukan berdasarkan hasil studi terdahulu oleh (Bouchra et al., 2024), yang menunjukkan bahwa rasio pembagian 80:10:10 memberikan performa terbaik untuk tugas klasifikasi teks. Oleh karena itu, penelitian ini mengadopsi rasio tersebut untuk memastikan model dapat dilatih secara optimal dan hasil evaluasinya representatif.

Selama proses pelatihan, model dikembangkan menggunakan framework *TensorFlow* dan *PyTorch*, dengan memanfaatkan pustaka dari HuggingFace *Transformers*. Untuk memperoleh performa terbaik, dilakukan eksperimen terhadap beberapa kombinasi *hyperparameter* menggunakan *optimizer* Adam. Parameter yang diuji antara lain jumlah *epoch* (2, 3, dan 4), *batch size* (16 dan 32), serta *learning rate* (2e-5, 3e-5, dan 5e-5) (Devlin et al., 2019). Fungsi *loss* yang digunakan adalah *SparseCategoricalCrossentropy* dengan *from_logits=True*, yang sesuai untuk klasifikasi multikelas dengan label dalam bentuk numerik. Evaluasi performa model selama pelatihan menggunakan metrik *SparseCategoricalAccuracy*, yang mengukur akurasi prediksi terhadap tiga kelas sentimen (positif, negatif, dan netral). Setiap kombinasi diuji secara terpisah untuk mengevaluasi kinerjanya. Setelah pelatihan selesai, model diuji menggunakan data uji untuk menilai performa akhir dalam mengklasifikasikan sentimen, dan hasilnya dianalisis dengan berbagai metrik evaluasi. Visualisasi arsitektur model IndoBERT yang digunakan dalam penelitian ini dapat dilihat pada Gambar 2.



Gambar 2. Diagram Arsitektur Indo BERT

2.8 Evaluasi Model

Tahap evaluasi model dilakukan untuk mengukur kinerja model IndoBERT dalam mengklasifikasikan sentimen *tweet* terkait isu penundaan CPNS 2025. Tujuan utama dari evaluasi ini adalah untuk menilai seberapa baik model mampu mengenali dan mengelompokkan opini masyarakat ke dalam tiga kategori sentimen, yaitu positif, negatif, dan netral.

Model yang telah dilatih dan divalidasi sebelumnya diuji menggunakan data uji yang telah dipisahkan selama proses pembagian dataset. Evaluasi kinerja dilakukan dengan menggunakan metrik-metrik yang diperoleh dari *confusion matrix*, yaitu: Akurasi, Presisi, *Recall*, dan *F1-score*. Akurasi menunjukkan seberapa besar proporsi prediksi yang benar dibandingkan dengan total data uji. Presisi mengukur ketepatan model dalam mengklasifikasikan suatu kategori sentimen dengan benar, sedangkan *recall* menunjukkan sejauh mana model dapat mengenali seluruh instance dari kategori tersebut. *F1-score* digunakan untuk menyeimbangkan presisi dan *recall*, terutama ketika distribusi data tidak seimbang antar kelas sentimen.

Selain pengukuran numerik, hasil evaluasi juga disajikan dalam bentuk visualisasi hasil pelatihan dan validasi model, seperti kurva akurasi dan loss selama proses training. Visualisasi ini membantu dalam memantau performa model dan mendeteksi kemungkinan *overfitting* atau *underfitting*. Selain itu, *confusion matrix* ditampilkan untuk menggambarkan distribusi prediksi model terhadap kelas sentimen secara lebih rinci. Evaluasi juga dilengkapi dengan tabel classification report yang memuat nilai akurasi, presisi, *recall*, dan *F1-score* untuk masing-masing kelas sentimen. Penyajian ini bertujuan untuk memberikan gambaran menyeluruh mengenai kekuatan dan kelemahan model dalam mengklasifikasikan opini masyarakat ke dalam kategori positif, negatif, dan netral.

Hasil evaluasi ini menjadi dasar penting untuk memahami persepsi publik terhadap kebijakan pemerintah dalam isu penundaan CPNS 2025. Selain itu, hasil ini juga menunjukkan keunggulan pendekatan berbasis transformer, khususnya IndoBERT, dibandingkan metode klasifikasi konvensional. Hal ini dikarenakan IndoBERT mampu menangkap konteks kalimat dalam bahasa Indonesia secara lebih kompleks dan mendalam, sehingga memberikan performa klasifikasi yang lebih akurat.

2.9 Analisis dan Pelaporan

Tahap analisis dan pelaporan merupakan langkah akhir dalam proses penelitian yang berfokus pada pengolahan hasil klasifikasi sentimen serta penyusunan laporan secara sistematis dan komprehensif. Setelah model IndoBERT melakukan klasifikasi pada tweet terkait isu penundaan CPNS 2025, hasil yang diperoleh dianalisis untuk mengidentifikasi pola distribusi sentimen masyarakat. Proporsi sentimen positif, negatif, dan netral divisualisasikan menggunakan diagram batang atau diagram lingkaran, yang memudahkan interpretasi data secara kuantitatif. Untuk mendukung analisis visual, *tools* seperti Microsoft Excel digunakan dalam pembuatan grafik dan ringkasan data.

Selain analisis kuantitatif, peneliti melakukan interpretasi kualitatif dengan meninjau konten tweet bernuansa negatif guna mengungkap isu dan kekhawatiran utama masyarakat, seperti kekecewaan terhadap kebijakan pemerintah, ketidakpastian nasib pelamar CPNS, atau ketimpangan informasi publik. Hasil analisis ini kemudian dibandingkan dengan latar belakang dan tujuan penelitian yang telah dirumuskan sebelumnya, untuk menilai pencapaian target penelitian.

Seluruh temuan dan proses penelitian disusun dalam laporan ilmiah yang sistematis menggunakan Microsoft Word, dengan struktur yang meliputi pendahuluan, metodologi, hasil dan pembahasan, serta kesimpulan. Laporan ini tidak hanya menyajikan hasil klasifikasi secara teknis, tetapi juga menguraikan implikasi sosial dari temuan penelitian,

memberikan rekomendasi bagi pemangku kebijakan, serta membuka peluang pengembangan penelitian selanjutnya. Dengan demikian, tahap analisis dan pelaporan menjadi bagian krusial dalam menyampaikan hasil penelitian secara efektif dan ilmiah.

2.10 Alat Penelitian

Penelitian ini menggunakan kombinasi perangkat keras dan perangkat lunak untuk mendukung seluruh proses pengumpulan, pengolahan, serta analisis data. Dari sisi perangkat keras, digunakan komputer dengan spesifikasi prosesor AMD RYZEN 5 5600H, VGA NVIDIA RTX 3015 Ti, RAM 16 GB, dan koneksi internet stabil, yang memadai untuk menjalankan proses *crawling data*, *pre-processing*, dan pelatihan model NLP.

Dari sisi perangkat lunak, penelitian ini mengandalkan model pra-latih IndoBERT yang dirancang khusus untuk bahasa Indonesia sebagai inti analisis sentimen. Berbagai pustaka Python seperti Pandas dan NumPy digunakan untuk manipulasi dan analisis data, serta Matplotlib untuk visualisasi hasil selama proses pelatihan dan evaluasi model. Sebelumnya, untuk visualisasi data awal, seperti diagram batang dan diagram lingkaran yang menampilkan total data tweet, digunakan Microsoft Excel guna memudahkan pemahaman distribusi data sebelum masuk tahap pemodelan.

Proses pengembangan dan eksperimen kode dilakukan pada platform Google Colab, yang menyediakan lingkungan interaktif dan akses komputasi berbasis cloud. Kombinasi alat-alat ini memastikan kelancaran dan efisiensi dalam menjalankan tahapan penelitian, mulai dari pengumpulan data hingga interpretasi hasil.

3. PEMBAHASAN

3.1 Dataset

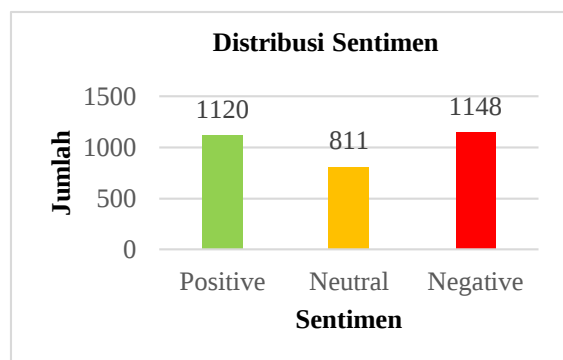
Data yang digunakan diperoleh dari platform X dengan menggunakan kata kunci yang berkaitan dengan penundaan pengangkatan CPNS tahun 2025. Proses pengumpulan data dilakukan selama periode tertentu dengan bantuan *tools Tweet Harvest*. Data yang dikumpulkan kemudian dilakukan *pre-processing*. Proses *preprocessing* ini mencakup penghapusan karakter khusus, tautan, emotikon, konversi teks menjadi huruf kecil, tokenisasi, serta normalisasi kata tidak baku. Tahapan ini dilakukan untuk memastikan bahwa data yang dianalisis bersih dan dapat diolah secara optimal oleh model IndoBERT. Setelah proses *preprocessing*, data yang valid adalah sebanyak 3079 data tweet, dengan distribusi sentimen yang dapat dilihat pada Tabel 5.

Tabel 5. Dataset Sentimen

Sentimen	Jumlah
Positive	1120
Neutral	811
Negative	1148

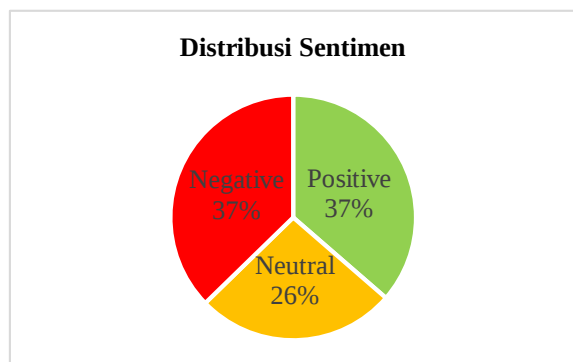
Sentimen	Jumlah
Total	3079

Berdasarkan Tabel 4, total data yang dikumpulkan adalah 3079 data tweet. Dengan 1120 kategori positif, 811 kategori netral, dan 1148 kategori negatif. Untuk memperjelas persebaran data sentimen, dilakukan visualisasi dalam bentuk diagram batang dan diagram lingkaran. Visualisasi ini bertujuan untuk memberikan gambaran yang lebih intuitif terkait proporsi masing-masing kategori sentimen. Diagram batang ini digunakan untuk menggambarkan jumlah absolut dari masing-masing kategori sentimen yang telah diidentifikasi, yaitu Positif, Netral, dan Negatif. Setiap batang mewakili total tweet dalam kategori tersebut, ditunjukkan pada Gambar 3.



Gambar 3. Distribusi Sentimen

Pada Gambar 3 di atas, terlihat bahwa jumlah tweet terbanyak berada pada kategori Negatif (1148 tweet), disusul oleh kategori Positif (1120 tweet) dan Netral (811 tweet). Selanjutnya Diagram lingkaran atau pie chart digunakan untuk menggambarkan proporsi atau persentase dari masing-masing kategori sentimen terhadap keseluruhan dataset. Visualisasi ini memberikan gambaran persentase yang lebih mudah dipahami secara sekilas yang ditunjukkan pada Gambar 4.



Gambar 4. Proporsi Sentimen

Berdasarkan Gambar 4, terlihat bahwa sentimen Negatif dan Positif masing-masing memiliki proporsi sebesar 37% dari total keseluruhan data, sementara sentimen Netral hanya mencakup 26%. Perbandingan ini menunjukkan bahwa opini publik terkait isu penundaan CPNS 2025 cenderung terbagi dua antara

yang bersifat mendukung dan menolak, sementara opini netral lebih sedikit dan umumnya berupa informasi deskriptif atau berita tanpa ekspresi emosional yang kuat.

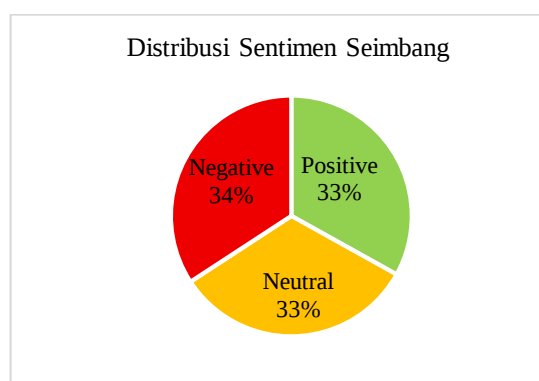
Distribusi data yang tidak merata ini menimbulkan permasalahan ketidakseimbangan kelas (*class imbalance*), yaitu kondisi di mana jumlah data dari satu atau dua kelas secara signifikan lebih besar dibandingkan kelas lainnya. Ketidakseimbangan kelas dapat berdampak negatif terhadap performa model klasifikasi karena model cenderung lebih mudah belajar dari kelas mayoritas dan mengabaikan kelas minoritas (Yutika et al., 2021). Hal ini dapat menyebabkan penurunan akurasi khususnya dalam mengklasifikasikan data dari kelas dengan jumlah yang lebih sedikit, serta menghasilkan bias klasifikasi.

Untuk memastikan bahwa model IndoBERT dapat belajar secara adil dari setiap kategori sentimen, dilakukan proses penyeimbangan data (*data balancing*). Teknik yang digunakan dalam penelitian ini adalah undersampling, yaitu dengan mengurangi jumlah data pada kelas mayoritas agar sebanding dengan jumlah data pada kelas minoritas (Rahma & Suadaa, 2023). Dalam konteks ini, kelas Positif dan Negatif yang memiliki jumlah lebih besar dibandingkan kelas Netral dikurangi masing-masing sebanyak 300 data secara acak. Tujuannya adalah untuk menghindari dominasi kelas mayoritas selama proses pelatihan model. Distribusi data setelah dilakukan undersampling ditampilkan pada Tabel 6.

Tabel 6. Dataset Sentimen yang Seimbang

Sentimen	Jumlah
Positive	820
Neutral	811
Negative	848
Total	2479

Berdasarkan Tabel 6. Dapat dilihat pada total data telah seimbang dengan 2479 data tweet. Untuk melihat lebih jelas pembagian data dapat dilihat pada diagram lingkaran. Visualisasi dalam bentuk diagram lingkaran bertujuan untuk memperlihatkan hasil dari proses balancing dataset secara proporsional. Ini penting untuk menilai apakah distribusi antar kelas sudah cukup setara dan layak untuk dijadikan input pelatihan model, yang ditunjukkan pada Gambar 5.



Gambar 5. Proporsi Sentimen Seimbang

Pada Gambar 5, terlihat bahwa proporsi sentimen telah berhasil diseimbangkan. Sentimen negatif mendominasi sebesar 34%, diikuti oleh positif dan netral masing-masing sebesar 33%. Proporsi ini menunjukkan bahwa dataset telah berada dalam kondisi yang relatif seimbang dan siap digunakan untuk tahap selanjutnya, yaitu pelatihan model analisis sentimen menggunakan IndoBERT.

3.2 Hasil Klasifikasi Sentimen

Setelah proses *preprocessing* dan penyeimbangan data, dilakukan pembagian data (*splitting*) dengan rasio 80:10:10 untuk data latih, data validasi, dan data uji. Tahap selanjutnya adalah proses *fine-tuning* terhadap model IndoBERT. Model yang digunakan dalam penelitian ini adalah *indobert-base-p2*, salah satu varian IndoBERT dari IndoNLU yang efektif untuk klasifikasi teks bahasa Indonesia.

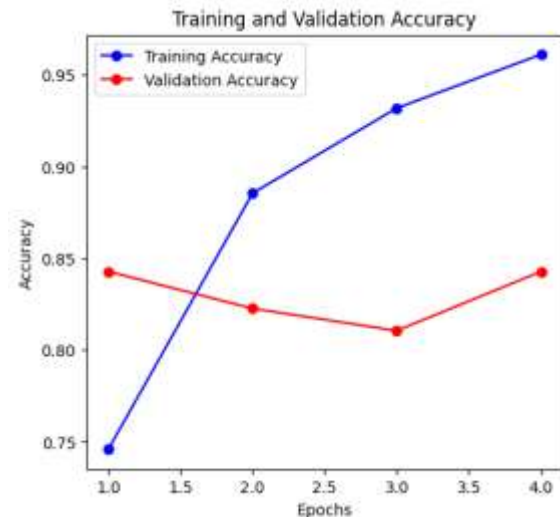
Proses *fine-tuning* dilakukan dengan melakukan eksplorasi terhadap beberapa konfigurasi *hyperparameter*, guna memperoleh performa model yang optimal. Adapun variasi *hyperparameter* yang diuji antara lain jumlah *epoch* 2, 3, dan 4, dengan 16 dan 32 *batch size*, dan $2e-5$, $3e-5$, dan $5e-5$ untuk *learning rate*. dengan total 18 kombinasi model. Seluruh model dilatih menggunakan *optimizer* Adam dan fungsi *loss* *SparseCategoricalCrossentropy*, sebagaimana dijelaskan pada bagian metode. Evaluasi dilakukan dengan metrik akurasi, *precision*, *recall*, dan *f1-score* dengan hasil perbandingan pada Tabel 7.

Tabel 7. Hasil Pengujian *Hyperparameter*

<i>Epoch</i>	<i>Batch Size</i>	<i>Learning Rate</i>	<i>Accuracy</i>	<i>Precision</i>	<i>Recall</i>	<i>F1-Score</i>
2	16	$2e-5$	82,66%	84%	83%	83%
		$3e-5$	79,03%	81%	79%	79%
		$5e-5$	78,63%	81%	79%	79%
		$2e-5$	81,85%	82%	82%	82%
	32	$3e-5$	79,84%	80%	80%	80%
		$5e-5$	83,46%	84%	84%	83%
		$2e-5$	79,03%	81%	80%	80%
		$3e-5$	82,66%	84%	83%	83%
3	16	$5e-5$	79,03%	80%	79%	79%
		$2e-5$	82,26%	86%	82%	83%
		$3e-5$	83,47%	84%	83%	83%
		$5e-5$	82,66%	83%	83%	83%
	32	$2e-5$	80,64%	81%	81%	80%
		$3e-5$	82,66%	83%	82%	83%
		$5e-5$	81,85%	83%	82%	82%
		$2e-5$	84,27%	84%	84%	84%
4	32	$3e-5$	81,85%	82%	82%	82%
		$5e-5$	82,26%	82%	82%	82%

Berdasarkan hasil pada Tabel 7, kombinasi *hyperparameter* terbaik diperoleh pada 4 *epoch*, *batch size* 32, dan *learning rate* $2e-5$, dengan hasil akurasi mencapai 84,27%, presisi 84%, *recall* 84%, dan *f1-score* sebesar 84%. Nilai-nilai metrik tersebut menunjukkan bahwa model mampu mengklasifikasikan data sentimen dengan baik dan seimbang antar ketiga kelas. Untuk mengevaluasi stabilitas proses pelatihan, grafik akurasi terhadap data pelatihan dan validasi disajikan pada Gambar 6. Grafik tersebut memperlihatkan konvergensi model yang

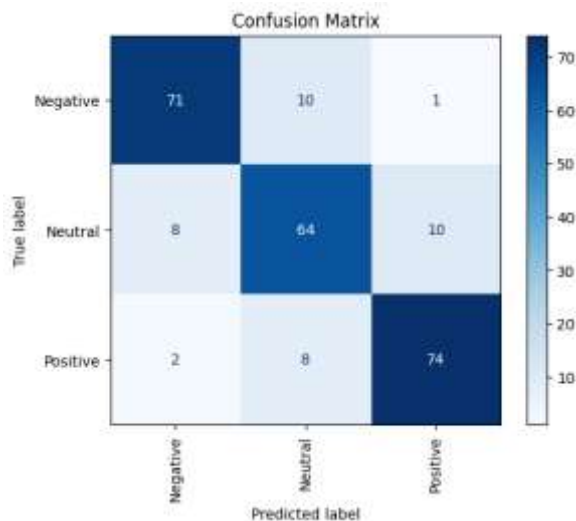
baik, ditandai dengan peningkatan akurasi pelatihan dan kestabilan akurasi validasi selama empat *epoch fine-tuning*.



Gambar 6. Training & Validation History

Gambar 6 menunjukkan grafik akurasi pelatihan dan akurasi validasi dari model IndoBERT selama proses *fine-tuning* selama empat *epoch*. Sumbu horizontal (x) merepresentasikan jumlah *epoch*, sementara sumbu vertikal (y) menunjukkan nilai akurasi. Dari grafik tersebut dapat dilihat bahwa akurasi pada data pelatihan (*training accuracy*) meningkat secara signifikan dari *epoch* ke-1 hingga *epoch* ke-4, dimulai dari 75% hingga mencapai lebih dari 95%. Hal ini menunjukkan bahwa model mampu mempelajari pola dalam data pelatihan dengan sangat baik seiring bertambahnya jumlah *epoch*. Sementara itu, akurasi validasi (*validation accuracy*) terlihat cukup stabil meskipun tidak meningkat secara signifikan. Nilai akurasi validasi berada di kisaran diatas 80–84% sepanjang empat *epoch*. Pada *epoch* ke-2 terlihat sedikit penurunan hingga *epoch* ke-3, namun kembali meningkat di *epoch* ke-4. Hal ini mengindikasikan bahwa model tidak mengalami *overfitting* secara drastis, namun terdapat gap performa antara pelatihan dan validasi yang perlu dicermati. Perbedaan tren antara akurasi pelatihan dan validasi ini menggarisbawahi pentingnya pemilihan *epoch* yang tepat agar model tidak hanya menghafal data pelatihan, tetapi juga dapat melakukan generalisasi yang baik terhadap data baru.

Setelah dilakukan pelatihan model menggunakan konfigurasi *hyperparameter* terbaik, evaluasi performa klasifikasi dilakukan menggunakan *confusion matrix* yang digunakan untuk menggambarkan kinerja model klasifikasi dengan membandingkan antara label asli (*true label*) dan label yang diprediksi oleh model (*predicted label*). Visualisasi *confusion matrix* dapat membantu melihat secara rinci seberapa baik model mengenali masing-masing kelas sentimen, yaitu Negatif, Netral dan Positif. Hasil *confusion matrix* dapat dilihat pada Gambar 7.



Gambar 7. Confusion Matrix

Gambar 7 menyajikan *confusion matrix* hasil prediksi model IndoBERT terhadap data uji. Setiap baris pada matriks mewakili label sebenarnya (*true label*), sedangkan setiap kolom mewakili label yang diprediksi oleh model. Pada kelas Negatif, dari total 82 data, sebanyak 71 data berhasil diprediksi dengan benar, sedangkan 10 data diklasifikasikan sebagai Netral dan 1 data sebagai Positif. Untuk kelas Netral, dari 82 data, model memprediksi 64 data dengan benar, sementara 8 data diklasifikasikan sebagai Negatif dan 10 data sebagai Positif. Adapun kelas Positif, dari 84 data, sebanyak 74 data diklasifikasikan secara tepat, dengan 2 data salah diklasifikasikan sebagai Negatif dan 8 data sebagai Netral.

Secara keseluruhan, model menunjukkan performa klasifikasi yang cukup baik, dengan tingkat keberhasilan tinggi untuk kelas Negatif dan Positif, serta performa yang relatif lebih menantang pada kelas Netral. Model memiliki kecenderungan untuk memprediksi ke kelas Netral ketika tidak yakin, yang merupakan indikasi umum dalam klasifikasi multikelas ketika terdapat kemiripan semantik antar kelas. Fenomena serupa juga ditemukan dalam sebuah studi yang melaporkan bahwa model IndoBERT cenderung memilih kelas Netral saat menghadapi ekspresi sentimen yang ambigu atau tidak eksplisit (Wafda et al., 2025).

Untuk melengkapi evaluasi model selain menggunakan *confusion matrix*, digunakan juga *classification report* yang memberikan gambaran kuantitatif dari performa model terhadap masing-masing kelas sentimen: Negative (0), Neutral (1), dan Positive (2). Laporan ini mencakup metrik evaluasi utama seperti *precision*, *recall*, *f1-score* dan *accuracy*, serta jumlah data tiap kelas (*support*), yang memberikan pemahaman menyeluruh terhadap kekuatan dan kelemahan model dalam mengklasifikasikan data. Hasil *classification report* dapat dilihat pada Gambar 8.

	precision	recall	f1-score	support
0	0.88	0.87	0.87	82
1	0.78	0.78	0.78	82
2	0.87	0.88	0.88	84
accuracy			0.84	248
macro avg	0.84	0.84	0.84	248
weighted avg	0.84	0.84	0.84	248

Gambar 8. Classification Report

Berdasarkan Gambar 8, diperoleh hasil evaluasi performa model klasifikasi sentimen yang menunjukkan perbedaan performa antar kelas. Pada kelas Negatif (label 0), model memiliki *precision* sebesar 0.88, yang berarti dari seluruh prediksi terhadap kelas Negatif, sebanyak 88% di antaranya benar. *Recall*-nya sebesar 0.87 menunjukkan bahwa model mampu mengenali 87% dari seluruh data yang benar-benar termasuk dalam kelas Negatif. Kombinasi *precision* dan *recall* ini menghasilkan nilai *f1-score* sebesar 0.87. Sementara itu, pada kelas Netral (label 1), *precision* model relatif lebih rendah, yaitu sebesar 0.78, dengan *recall* yang juga hanya mencapai 0.78. Hal ini menunjukkan bahwa model sering memprediksi kelas Netral, meskipun tidak selalu tepat, sehingga menghasilkan *f1-score* sebesar 0.78. Adapun pada kelas Positif (label 2), model menunjukkan performa yang cukup baik dan seimbang, dengan *precision* sebesar 0.87 dan *recall* sebesar 0.88, menghasilkan *f1-score tertinggi* dibandingkan kelas lainnya, yaitu 0.88.

Secara keseluruhan, akurasi model mencapai 84%. Rata-rata performa model yang ditunjukkan juga sama-sama berada di angka 84% untuk *precision*, *recall*, dan *f1-score*, mencerminkan kestabilan model dalam menangani distribusi data yang relatif seimbang antar kelas. Nilai-nilai ini menunjukkan bahwa model mampu melakukan klasifikasi sentimen dengan baik secara umum, meskipun masih ada tantangan dalam mengenali kelas Netral secara lebih akurat. Hasil ini sejalan dengan temuan yang menunjukkan bahwa model IndoBERT mampu memberikan performa tinggi dalam tugas klasifikasi sentimen pada bahasa Indonesia, terutama ketika digunakan dengan *fine-tuning* yang tepat (Nuryadi et al., 2025).

Sebagai perbandingan, sebuah penelitian oleh (Khairunissabina & Kurniawan, 2025) menggunakan pendekatan *machine learning* klasik dengan *Support Vector Machine* (SVM) juga mengklasifikasikan sentimen terkait penerimaan CPNS 2024. Meskipun model SVM yang digunakan berhasil mencapai akurasi 82,5%, performa tersebut masih berada sedikit di bawah model IndoBERT-base-p2 dalam penelitian ini. Hal ini menunjukkan keunggulan model berbasis *transformer*, terutama dalam menangkap konteks linguistik bahasa Indonesia yang kompleks, dibandingkan pendekatan berbasis fitur statistik seperti TF-IDF.

Dengan performa model yang andal tersebut, analisis selanjutnya dilakukan terhadap data tweet terkait isu penundaan pengangkatan CPNS 2025 untuk mengetahui distribusi sentimen masyarakat seiring

waktu. Hasil klasifikasi menunjukkan bahwa sentimen negatif mendominasi pada fase awal munculnya isu tersebut. Hal ini mencerminkan kekhawatiran dan kekecewaan masyarakat terhadap ketidakjelasan informasi serta ketidakpastian nasib para pelamar. Seiring berjalannya waktu, terjadi pergeseran sentimen ke arah yang lebih positif, terutama setelah pemerintah memberikan klarifikasi dan mengumumkan penjadwalan ulang yang dimajukan. Sementara itu, sentimen netral didominasi oleh akun resmi dan media yang menyampaikan informasi tanpa ekspresi opini.

Temuan ini menunjukkan bahwa persepsi publik terhadap kebijakan sangat dinamis, dan media sosial menjadi kanal penting untuk memantau perubahan opini masyarakat secara real-time. Hal ini sejalan dengan penegasan mengenai pentingnya analisis sentimen media sosial sebagai alat pemantauan persepsi publik terhadap kebijakan pemerintah secara real-time (Habibi et al., 2022). Selain itu, hasil ini juga memperlihatkan potensi besar analisis sentimen berbasis NLP sebagai bahan evaluasi komunikasi publik bagi pembuat kebijakan.

4. KESIMPULAN

Penelitian ini berhasil mengumpulkan dan mengolah 3.079 data tweet terkait isu penundaan pengangkatan CPNS 2025, yang setelah dilakukan penyeimbangan menjadi 2.479 data dengan distribusi sentimen relatif seimbang yang terdiri dari 1120 data Positif, 1148 data Negatif, dan 811 data Netral. *Fine-tuning* model IndoBERT dengan konfigurasi *hyperparameter* terbaik (4 *epoch*, *batch size* 32, *learning rate* 2e-5) menghasilkan performa klasifikasi yang baik, dengan akurasi sebesar 84,27% serta nilai *precision*, *recall*, dan *f1-score* rata-rata 84%. Meskipun model mampu mengklasifikasikan sentimen dengan andal, tantangan masih muncul dalam mengenali kelas Netral yang bersifat ambigu.

Analisis sentimen menunjukkan dominasi sentimen negatif pada fase awal isu, yang mencerminkan kekhawatiran dan ketidakpastian masyarakat. Namun, setelah klarifikasi dan penjadwalan ulang oleh pemerintah, sentimen berangsur positif, sementara sentimen netral banyak berasal dari akun resmi dan media yang menyampaikan informasi tanpa opini.

Temuan ini menegaskan dinamika persepsi publik yang sangat dipengaruhi oleh informasi terkini dan pentingnya media sosial sebagai kanal pemantauan opini masyarakat secara real-time. Selain itu, penelitian ini menggarisbawahi potensi besar analisis sentimen berbasis NLP sebagai alat evaluasi komunikasi publik yang mendukung pengambilan keputusan kebijakan pemerintah.

Secara praktis, analisis sentimen berbasis NLP dapat menjadi alat efektif bagi pemerintah dan pembuat kebijakan dalam memantau opini publik dan merespons perubahan persepsi masyarakat dengan tepat. Media dan institusi informasi juga dapat memanfaatkan hasil ini untuk meningkatkan kualitas komunikasi publik yang lebih transparan dan

interaktif. Model IndoBERT yang digunakan memiliki potensi besar untuk diterapkan dalam berbagai konteks analisis opini bahasa Indonesia, bahkan dapat menjadi acuan bagi industri lain yang ingin membangun sistem analisis sentimen di bidangnya masing-masing.

Penelitian ini memiliki keterbatasan pada pengumpulan data yang hanya berasal dari satu platform media sosial, yaitu platform X, serta dalam rentang waktu tertentu, sehingga hasilnya belum sepenuhnya merepresentasikan persepsi publik secara menyeluruh. Selain itu, kombinasi *hyperparameter* yang diuji juga masih terbatas, sehingga potensi optimalisasi model belum sepenuhnya tergalai.

Sebagai tindak lanjut, penelitian selanjutnya disarankan untuk memperluas pengumpulan data dari berbagai platform media sosial dan periode waktu yang lebih panjang agar dapat memperoleh gambaran sentimen publik yang lebih komprehensif dan representatif. Selain itu, eksplorasi teknik penyeimbangan data lain, seperti *oversampling* atau *augmentasi data*, perlu dilakukan untuk mengatasi masalah ketidakseimbangan kelas sekaligus meningkatkan performa model. Pemantauan sentimen secara berkelanjutan juga dianjurkan guna mendukung evaluasi komunikasi publik secara real-time yang lebih akurat dan responsif.

PUSTAKA

- Arham, A., Swedia, E. R., Cahyanti, M., & Septian, M. R. D. (2022). Implementasi Sentiment Analysis Pada Opini Masyarakat Indonesia Di Twitter Terhadap Virus Covid-19 Varian Omicron Dengan Algoritma Naïve Bayes, Decision Tree, Dan Support Vector Machine. *Sebatik*, 26(2), 565–572.
- Bouchra, F., Suarjaya, I. M. A. D., & Rusjyanthi, N. K. D. (2024). Analisis Sentimen Masyarakat terhadap Tayangan Televisi Nasional menggunakan Metode Deep Learning . *Jurnal Buana Informatika*, 15(2), 89–99.
- Habibi, M., Ma'arif, M. R., & Subekti, D. (2022). The Development of Social Media Intelligence System for Citizen Opinion and Perception Analysis over Government Policy. *Telematika: Jurnal Informatika Dan Teknologi Informasi*, 19(1), 31–46.
- Hakim, G., Fatyanosa, T. N., & Widodo, A. W. (2024). Analisis Sentimen Masyarakat terhadap Kereta Cepat Whoosh pada Platform X menggunakan IndoBERT. *Jurnal Pengembangan Teknologi Informasi Dan Ilmu Komputer*, 8(10).
- Khairunissabina, K., & Kurniawan, R. (2025). Analisis Sentimen Mengenai Penerimaan Calon Pegawai Negeri Sipil Tahun 2024 Menggunakan Support Vector Machine. *Progresif: Jurnal Ilmiah Komputer*, 21(1), 132–143.
- Khotimah, D. A. K., & Sarno, R. (2019). Sentiment Analysis of Hotel Aspect Using Probabilistic

- Latent Semantic Analysis, Word Embedding and LSTM. *International Journal of Intelligent Engineering and Systems*, 12(4), 275–290.
- Koto, F., Rahimi, A., Lau, J. H., & Baldwin, T. (2020). *IndoLEM and IndoBERT: A Benchmark Dataset and Pre-trained Language Model for Indonesian NLP*. 757–770. <https://doi.org/https://doi.org/10.48550/arXiv.2011.00677>
- Kristi, A. K. (2024). *Pengangkatan CPNS 2024 Mundur ke Oktober 2025, Apa Alasannya?*
- Magfirah, C. E. (2025). *Pelamar CPNS 2024 dan PPPK 2024 Bingung! Resign Sudah Diajukan, Pengangkatan Malah Ditunda* Artikel ini telah tayang di *Tribungayo.com* dengan judul *Pelamar CPNS 2024 dan PPPK 2024 Bingung! Resign Sudah Diajukan, Pengangkatan Malah Ditunda*, <https://gayo.tribunnews.com/2025/03/08/pelamar-cpns-2024-dan-pppk-2024-bingung-resign-sudah-diajukan-pengangkatan-malah-ditunda>. Penulis: Cut Eva Magfirah | Editor: Sri Widya Rahma.
- Nuryadi, D., Metandi, F., Hadiwijaya, N. A., Rohman, M. Z., Hartanto, S., Syafrizal, A., & Yadie, E. (2025). Fine Tuning Indobert Untuk Analisis Sentimen Pada Ulasan Pengguna Aplikasi Tiket .com digoogle Play Store. *JATI (Jurnal Mahasiswa Teknik Informatika)*, 9(2), 3577–3583.
- Putra, A. D. A., & Juanita, S. (2021). Analisis Sentimen pada Ulasan pengguna Aplikasi Bibit Dan Bareksa dengan Algoritma KNN. *JATISI (Jurnal Teknik Informatika Dan Sistem Informasi)*, 8(2), 636–646.
- Putri, D. D., Nama, G. F., & Sulistiono, W. E. (2022). Analisis Sentimen Kinerja Dewan Perwakilan Rakyat (DPR) Pada Twitter Menggunakan Metode Naive Bayes Classifier. *Jurnal Informatika Dan Teknik Elektro Terapan*, 10(1), 34–40.
- Putriyeki, A., Sulistiawati, D., Khairani, N., & Kusuma, R. R. A. (2025). Analisis Multidimensional Sentimen Masyarakat terhadap Program Makan Bergizi Gratis pada Media Sosial X. *Integrative Perspectives of Social and Science Journal*, 2(1), 1004–1024.
- Rahma, I. A., & Suadaa, L. H. (2023). Penerapan Text Augmentation untuk Mengatasi Data yang Tidak Seimbang pada Klasifikasi Teks Berbahasa Indonesia. *Jurnal Teknologi Informasi Dan Ilmu Komputer*, 10(6), 1329–1340.
- Singgalen, Y. A. (2025). Performance Analysis of IndoBERT for Sentiment Classification in Indonesian Hotel Review Data. *Journal of Information System Research (JOSH)*, 6(2), 976–986.
- Statista. (2024). *Leading countries based on number of X (formerly Twitter) users as of April 2024*.
- Talalu, T. R., & Valentine, F. (2021). Interaksi Pendengar dan Promosi Program Siaran Radio “Polemik Trijaya” di Twitter. *Jurnal Dakwah Dan Komunikasi*, 6(2), 237–254.
- Wafda, A., Fudholi, D. H., & Nugraha, J. (2025). Aspect-Based Sentiment Analysis On Twitter Tweets About The Merdeka Curriculum Using IndoBERT. *JITK (Jurnal Ilmu Pengetahuan Dan Teknologi Komputer)*, 10(3), 586–599.
- Wilie, B., Vincentio, K., Winata, G. I., Cahyawijaya, S., Li, X., Lim, Z. Y., Soleman, S., Mahendra, R., Fung, P., Bahar, S., & Purwarianti, A. (2020). *IndoNLU: Benchmark and resources for evaluating Indonesian natural language understanding*. <https://doi.org/https://doi.org/10.48550/arXiv.2009.05387>
- Yulianti, C. (2025). *Pengangkatan CPNS 2024 Diundur Jadi Oktober 2025, PPPK Kapan?* DetikEdu.
- Yutika, C. H., Adiwijaya, & Faraby, S. Al. (2021). Analisis Sentimen Berbasis Aspek pada Review Female Daily Menggunakan TF-IDF dan Naïve Bayes. *J. Media Inform. Budidarma*, 5(2), 422–430.
- Zuhri, H., Sawitri, D., & Muawanah, U. (2023). Implementasi dan Implikasi Kebijakan Moratorium CPNS di Lingkungan Pemerintah Kabupaten Mojokerto. *Jurnal Ekonomi Dan Manajemen*, 24(2), 22–38.

HALAMAN IINI SENGAJA DI KOSONGKAN