

Classification of Sentiment in Reddit Forum Comments using Fine-Tuned IndoBERT (Case Study: Free Nutritious Meal Program) [Klasifikasi Sentimen Komentar Forum Reddit menggunakan Fine-Tuned IndoBERT (Studi Kasus : Program Makan Bergizi Gratis)]

Bram Aji Saka Putra¹⁾, Suprianto ^{*,2)}

¹⁾Program Studi Teknik Informatika, Universitas Muhammadiyah Sidoarjo, Indonesia

²⁾ Program Studi Teknik Informatika, Universitas Muhammadiyah Sidoarjo, Indonesia

*Email Penulis Korespondensi: suprianto@umsida.ac.id

Abstract. This study analyzes public opinion on the Free Nutritious Meals Program (MBG) policy on the Reddit platform using the IndoBERT model. The study uses a dataset of 6,295 Reddit comments collected from 2024 to 2025. The preprocessing stage implemented a comprehensive suite of textual refinement procedures, encompassing normalization of colloquial language, emoji conversion, and translation into Indonesian to establish linguistic consistency across the corpus. Manual annotation was performed with meticulous care on a balanced dataset of 3,000 comments, evenly allocated among positive, negative, and neutral classes. An 80:20 train-test split was applied to enhance the reliability of model training and evaluation. A fine-tuned IndoBERT model exhibited outstanding performance, attaining 99.00% for accuracy, F1-score, and precision. Applied to the full dataset, the model predicted neutral sentiment as predominant (59.44%), followed by positive (34.11%) and negative (6.45%) sentiments, a distribution that suggests a measured and reflective public discourse on the topic. Model reliability was further supported by a mean confidence score of 0.8502 and a processing throughput of 137.93 samples per second, indicating strong potential for deployment as an effective tool for near real-time sentiment monitoring in public policy contexts.

Keywords - IndoBERT; sentiment analysis; Reddit; Free Nutritious Meal Program; natural language processing

Abstrak. Penelitian ini menganalisis opini publik terhadap kebijakan Program Makan Bergizi Gratis (MBG) di platform Reddit menggunakan model IndoBERT. Penelitian menggunakan dataset dengan 6.295 komentar Reddit yang dikumpulkan dari periode 2024-2025. Pra-pemrosesan data meliputi pembersihan teks, normalisasi slang, konversi emoji, dan penerjemahan ke bahasa Indonesia. Pelabelan manual dilakukan pada 3.000 komentar yang seimbang (positif, negatif, netral) dengan pembagian data 80:20. Model IndoBERT yang di-fine-tune mencapai akurasi 99,00%, F1-score 99,00%, dan presisi 99,00%. Hasil prediksi pada dataset lengkap menunjukkan 59,44% sentimen netral, 34,11% positif, dan 6,45% negatif, mengindikasikan diskusi publik yang informatif dan analitis. Model yang dikembangkan dalam penelitian ini menunjukkan tingkat keandalan yang signifikan, tercermin dari skor kepercayaan rata-rata sebesar 0,8502 serta kemampuan memproses data hingga 137,93 sampel per detik. Pencapaian tersebut menegaskan relevansi dan responsivitas model ini sebagai instrumen yang efektif untuk pemantauan kebijakan publik secara waktu nyata, khususnya dalam menangkap dinamika serta pergeseran opini masyarakat yang senantiasa berkembang.

Kata Kunci - IndoBERT; analisis sentimen; Reddit; Program Makan Bergizi Gratis; pemrosesan bahasa alami

I. PENDAHULUAN

Perkembangan pesat teknologi internet telah mengubah secara fundamental cara manusia berkomunikasi dan memperoleh informasi. Dalam beberapa tahun terakhir, banyak orang mulai beralih dari media berita tradisional ke platform digital, dengan media sosial menjadi tempat utama untuk mendapatkan informasi [1]. Lewat media sosial, pengguna tidak hanya bisa mendapatkan berita secara cepat dan terbuka, tetapi juga saling terhubung untuk berbagi pendapat dan berdiskusi bersama [15].

Reddit dipilih sebagai sumber data karena memungkinkan konten teks yang jauh lebih panjang (hingga 40.000 karakter untuk postingan dan 10.000 karakter untuk komentar) dibandingkan platform media sosial lain seperti Twitter yang terbatas 280 karakter, sehingga menghasilkan diskusi mendalam dengan kompleksitas linguistik dan nuansa sentimen yang kaya karakteristik yang krusial untuk analisis sentimen berbasis NLP [13][2]. Selain itu, struktur subreddit yang terorganisir berdasarkan topik, tingkat anonimitas yang mendorong ekspresi autentik, serta ketersediaan API untuk ekstraksi data skala besar menjadikan Reddit sumber data yang ideal untuk mengkaji opini publik terhadap kebijakan Program Makan Bergizi Gratis [6].

Program Makan Bergizi Gratis (MBG) dipilih sebagai fokus penelitian karena merupakan kebijakan publik berskala jangka panjang yang berlanjut lintas pemerintahan. Dengan menargetkan anak-anak, pelajar, dan ibu hamil

Copyright © Universitas Muhammadiyah Sidoarjo. This preprint is protected by copyright held by Universitas Muhammadiyah Sidoarjo and is distributed under the Creative Commons Attribution License (CC BY). Users may share, distribute, or reproduce the work as long as the original author(s) and copyright holder are credited, and the preprint server is cited per academic standards.

Authors retain the right to publish their work in academic journals where copyright remains with them. Any use, distribution, or reproduction that does not comply with these terms is not permitted.

untuk mengatasi stunting dan kelaparan, keberhasilan program ini memiliki dampak besar terhadap kualitas sumber daya manusia dan berpotensi menjadi dasar bagi kebijakan serupa di masa mendatang [3]. Penelitian ini juga relevan secara temporal karena data dikumpulkan sepanjang 2024–2025, meliputi fase pengumuman, tahap implementasi awal, hingga evaluasi lanjutan. Dengan demikian, hasil penelitian tidak hanya menggambarkan opini publik saat ini, tetapi juga berfungsi sebagai tolok ukur untuk memantau dinamika perubahan sentimen di masa depan.

Analisis sentimen dipilih sebagai pendekatan untuk memahami respons masyarakat terhadap Program Makan Bergizi Gratis (MBG). Dengan memanfaatkan teknik Natural Language Processing (NLP), metode ini memungkinkan pengklasifikasian opini publik ke dalam kategori positif, negatif, dan netral [4] sekaligus menyingkap nuansa emosional dan kecenderungan evaluatif yang terkandung dalam ungkapan warga. Mengingat kompleksitas ragam bahasa Indonesia, IndoBERT dipilih sebagai solusi yang efektif karena kemampuannya menangkap konteks lintas gaya penulisan, variasi tingkat formalitas, dan dinamika bahasa sehari-hari secara mendalam, sehingga mendukung interpretasi yang lebih sensitif terhadap makna sosial dan kultural di balik setiap komentar.

Penelitian ini secara esensial bertujuan untuk menilai tingkat akurasi sekaligus performa model IndoBERT dalam mengklasifikasikan sentimen publik terkait Program Makan Bergizi Gratis pada platform Reddit. Di samping itu, penelitian ini juga berupaya menelaah pola distribusi serta karakteristik sentimen yang terungkap melalui hasil prediksi model, sehingga pada akhirnya dapat memperluas dan memperdalam pemahaman kita bersama mengenai harapan serta aspirasi masyarakat terhadap kebijakan kesejahteraan anak yang tengah digulirkan.

Berdasarkan tinjauan terhadap studi-studi terdahulu, teridentifikasi kesenjangan penelitian yang substansial, yang menjadi dasar pengembangan penelitian ini untuk memperkaya wawasan tentang dinamika sentimen masyarakat.

Nurjoko dan Rahardi (2024) telah menerapkan IndoBERT untuk mengidentifikasi sentimen kekerasan verbal di Twitter dengan akurasi 72%, namun belum melakukan optimalisasi hiperparameter dan optimizer secara sistematis guna memaksimalkan potensi model dalam menangkap nuansa emosional masyarakat.

Nur, Wibisono, dan Megasari (2025) memanfaatkan fine-tuning IndoBERT bersama BERTopic untuk analisis sentimen merek teknologi di platform X; tetapi, ketidakadaan perbandingan performa antara model pre-trained dan fine-tuned secara terpisah menyulitkan isolasi efek optimasi yang sebenarnya.

Munir (2025) menganalisis sentimen Program MBG menggunakan SVM dan BERT generik pada platform X dengan dataset tidak seimbang; namun, absennya IndoBERT yang dioptimalkan serta karakteristik linguistik khas Reddit menjadi celah metodologis yang substansial dalam memahami dinamika wacana publik.

Pham et al. (2025) [2] membandingkan VADER, TextBlob, dan BERT generik untuk analisis sentimen di Reddit terkait AI, namun penggunaan BERT generik bukan IndoBERT yang dioptimalkan untuk bahasa Indonesia serta topik AI yang tidak relevan dengan opini publik kebijakan Indonesia menunjukkan perlunya adaptasi model dan domain.

Wagay dan Jahiruddin (2025) [22] menunjukkan keunggulan SBERT-BiLSTM untuk klasifikasi mental illness dari Reddit posts, namun penggunaan model hybrid bukan IndoBERT tunggal dan tanpa evaluasi fine-tuning serta fokus distribusi sentimen tiga kelas pada domain kebijakan publik menjadi batasan yang belum teratasi.

Berangkat dari kesenjangan tersebut, penelitian ini secara khusus mengoptimalkan model IndoBERT melalui konfigurasi hyperparameter dan optimizer AdamW untuk analisis sentimen kebijakan publik Indonesia di platform Reddit, sekaligus menggali secara komprehensif distribusi serta karakteristik sentimen masyarakat terhadap Program Makan Bergizi Gratis dengan mempertimbangkan variasi bahasa sehari-hari, dialek regional, serta kebisingan kontekstual yang melekat pada data media sosial.

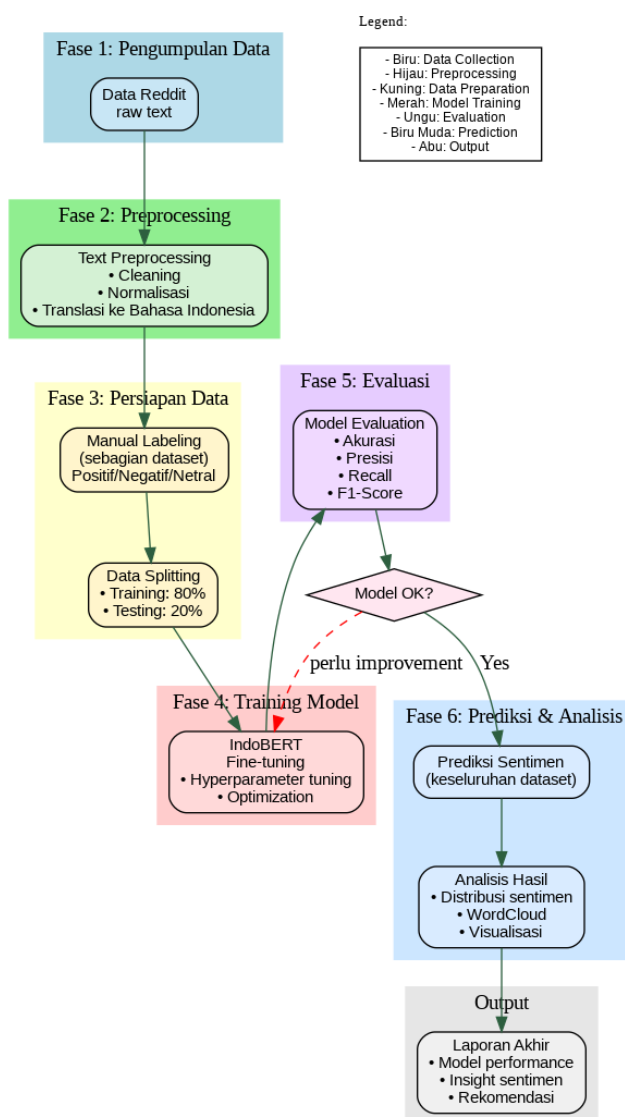
II. METODE

Penelitian ini menerapkan pendekatan kuantitatif melalui metode analisis sentimen berbasis deep learning, guna mengungkap wawasan mendalam tentang dinamika opini publik secara sistematis dan objektif. Penelitian mengadopsi paradigma positivisme dengan pendekatan eksperimental, di mana pengukuran dilakukan secara objektif menggunakan metrik kuantitatif untuk mengevaluasi performa model [5].

2.1 Alur Kerja Penelitian

Alur kerja penelitian dirancang secara bertahap melalui enam fase utama: (1) pengumpulan data menggunakan PRAW API dengan kata kunci "MBG" dan "Makan Bergizi Gratis"; (2) preprocessing yang mencakup normalisasi slang, konversi emoji, serta penerjemahan ke Bahasa Indonesia; (3) persiapan data melalui pelabelan manual parsial dan pembagian dataset rasio 80:20 untuk pelatihan model; (4) fine-tuning model IndoBERT dengan hyperparameter tuning dan optimizer AdamW; (5) evaluasi berbasis metrik akurasi, presisi, recall, serta F1-score; dan (6) prediksi serta Analisis secara menyeluruh terhadap keseluruhan dataset dilakukan guna mengungkap pola distribusi sentimen serta karakteristik opini publik yang terkandung di dalamnya. Sebagai penunjang pemahaman atas kerangka metodologis yang digunakan, Gambar 1. menyajikan ilustrasi alur penelitian secara komprehensif, sekaligus

memberikan landasan visual yang memudahkan pembaca dalam memahami tahapan-tahapan penelitian secara sistematis dan terstruktur.

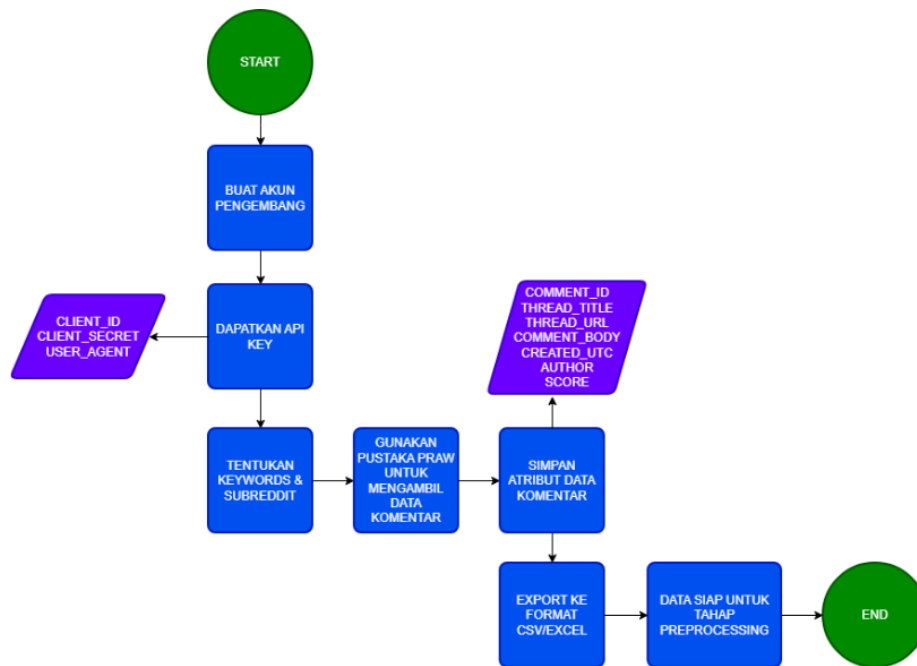


Gambar 1. Alur Penelitian

Pada Gambar 1. dijelaskan alur metodologi penelitian yang terdiri dari enam fase utama. Fase 1 (Pengumpulan Data) merupakan tahap akuisisi data mentah dari Reddit menggunakan API. Fase 2 (Preprocessing) meliputi pembersihan dan normalisasi teks. Fase 3 (Persiapan Data) mencakup pelabelan manual dan pembagian data training-testing. Fase 4 (Training Model) adalah proses fine-tuning IndoBERT dengan hyperparameter tuning. Fase 5 (Evaluasi) menggunakan metrik akurasi, presisi, recall, dan F1-score. Fase 6 (Prediksi dan Analisis) menerapkan model pada dataset lengkap untuk analisis distribusi sentimen. Diagram ini menunjukkan adanya mekanisme feedback loop (garis putus-putus merah) yang memungkinkan perbaikan model jika hasil evaluasi belum memenuhi kriteria.

2.2 Pengumpulan Data

Data penelitian bersumber dari forum diskusi Reddit, khususnya subreddit r/Indonesia. Reddit memiliki karakteristik diskursif yang mendukung debat mendalam dan tingkat anonimitas relatif tinggi, menjadikannya repositori data primer yang bernilai untuk kajian opini publik [12]. Pengumpulan data dilakukan melalui web scraping menggunakan Reddit API dan pustaka PRAW (Python Reddit API Wrapper) [6]. Gambar 2. menunjukkan flowchart proses pengumpulan data.



Gambar 2. Flowchart Pengumpulan Data

Gambar 2. mengilustrasikan tahapan sistematis pengumpulan data dari Reddit. Proses dimulai dengan pembuatan akun pengembang Reddit untuk memperoleh API key (client_id, client_secret, user_agent). Selanjutnya ditentukan keywords "MBG" dan "Makan Bergizi Gratis" serta subreddit target. Pustaka PRAW digunakan untuk mengambil data komentar dengan atribut yang disimpan meliputi: comment_id, thread_title, thread_url, comment_body, created_utc, author, dan score. Data kemudian diekspor ke format CSV/Excel untuk tahap preprocessing. Flowchart ini menekankan pentingnya autentikasi API dan penyimpanan atribut komprehensif untuk analisis lanjutan.

Proses pengumpulan data dilaksanakan dengan memanfaatkan kata kunci "MBG" dan "Makan Bergizi Gratis" dalam rentang waktu Januari 2024 hingga Agustus 2025. Seluruh data tersebut selanjutnya menjalani proses seleksi dan pembersihan secara ketat, sehingga diperoleh dataset akhir yang lebih representatif. Sebagai pelengkap, Tabel 1. menyajikan rangkuman atribut-atribut data yang berhasil dikumpulkan, sekaligus memberikan gambaran yang lebih utuh mengenai karakteristik dan keragaman sumber opini masyarakat yang menjadi objek kajian.

Tabel 1. Atribut Komentar Reddit

Atribut	Detail
comment_id	ID unik komentar.
Keyword_pencarian	Kata kunci yang digunakan saat proses pengumpulan data
Subreddit	Komunitas pada forum reddit
thread_title	Judul topik/thread.
thread_url	Alamat laman topik/thread.
comment_body	Teks komentar.
created_utc	Waktu pembuatan komentar.
author	Nama pengguna.
score	Jumlah upvote/downvote komentar.

Tabel 1. merinci delapan atribut utama yang diekstraksi dari setiap komentar Reddit. Atribut *comment_id* berfungsi sebagai pengenal unik untuk mencegah duplikasi data; *keyword_pencarian* mencatat metode pengambilan guna validasi relevansi; *subreddit* mengidentifikasi komunitas sumber diskusi; sementara *thread_title* dan *thread_url* menyediakan konteks percakapan yang kaya. *Comment_body* menjadi inti data untuk analisis sentimen, *created_utc* memungkinkan analisis temporal, sedangkan *author* dan *score* memberikan metadata keterlibatan yang berharga untuk eksplorasi lebih lanjut.

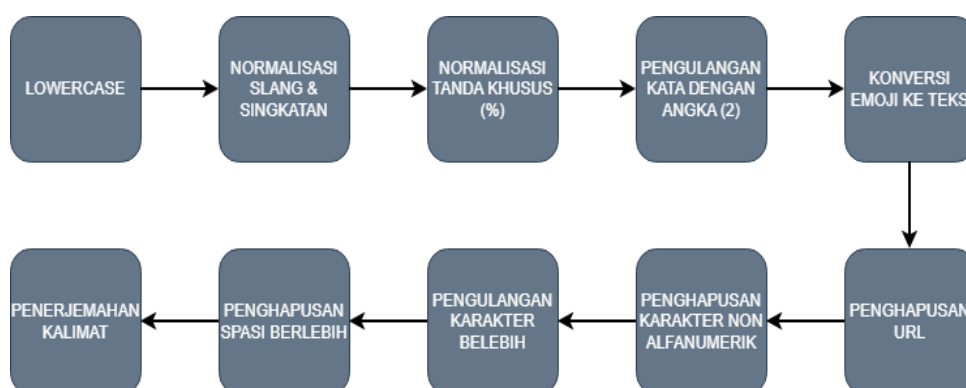
2.3 Preprocessing Data

Tahap preprocessing data merupakan langkah krusial dalam pipeline NLP untuk membersihkan dan menormalisasi teks mentah dari media sosial [11]. Tabel 2. menunjukkan contoh hasil preprocessing.

Tabel 2. Contoh Hasil Preprocessing	
Raw Comment	Preprocessing
Woyy program MBG keren bgt nih!! 🔥🔥 Program makan bergizi gratis 100% bakal ngebantu bgt buat anak2 sekolah di daerah terpencil 🙏🙏 Link lengkapnya di sini: https://reddit.com/r/indonesia/comments/xyz123	hei program makan bergizi gratis hebat banget ini fire fire program makan bergizi gratis seratus persen akan membantu banget untuk anak anak sekolah di daerah terpencil crying face crying face saya sih setuju seratus persen sama program ini tapi ada yang bilang kualitasnya kurang bagus ya bagaimana lagi ya kalau budgetnya terbatas smiling face with tears smiling face with tears ngomongo ngomong menurut kalian bagaimana layak tidak sih thinking face
Gue sih setuju 100% sama program ini, tapi ada yg bilang kualitasnya kurang bagus... Ya gimanaaa lagi ya kalo budgetnya terbatas 😊😊	
Btw, menurut kalian gimana? Worth it gak sih? 🤔	
#MBG #MakanBergiziGratis #IndonesiaEmas2024	

Tabel 2. mendemonstrasikan empat contoh transformasi preprocessing pada komentar Reddit. Kolom Raw Comment menunjukkan teks asli dengan karakteristik media sosial: slang ("woyy", "bgt", "ngebantu"), singkatan ("MBG", "2"), emoji, URL, dan campuran bahasa Indonesia-Inggris. Kolom Preprocessing Result menunjukkan hasil normalisasi: konversi slang ke bentuk standar, ekspansi singkatan, konversi emoji ke deskripsi teks [14] dalam bracket, penghapusan URL, dan standarisasi bahasa. Transformasi ini memastikan konsistensi korpus untuk input model IndoBERT.

Serangkaian proses pembersihan sistematis dan normalisasi adaptif diperlukan untuk mentransformasikan data mentah menjadi representasi terstruktur yang koheren [17].



Gambar 3. Urutan Tahap Preprocessing

Gambar 3. menggambarkan pipeline preprocessing dalam 10 tahapan berurutan. Urutan dimulai dengan lowercase untuk uniformitas, diikuti normalisasi slang dan singkatan menggunakan kamus khusus. Normalisasi tanda khusus (%) dan pengulangan kata dengan angka (2) menangani pola khas bahasa sosial media. Konversi emoji ke teks mempertahankan nuansa emosional. Penghapusan URL dan karakter non-alfanumerik mengurangi noise. Pengurangan pengulangan karakter berlebih serta penghapusan spasi redundan dilakukan untuk menstandarisasi format teks secara konsisten. Tahap penerjemahan ke bahasa Indonesia ditempatkan sebagai langkah akhir guna menjamin homogenitas linguistik yang optimal. Urutan proses ini dirancang mengikuti praktik terbaik (best practices), demi memaksimalkan kualitas data input yang mendukung analisis sentimen secara akurat dan bermakna.

2.4 Pelabelan dan Pembagian Data

Pelabelan manual diprioritaskan karena superioritasnya dalam menghasilkan ground truth berkualitas tinggi [7]. Selain label sentimen primer tiga kelas, anotator menyertakan metadata tambahan konteks afektif untuk memperkaya dataset [8] [18]. Pelabelan manual dilakukan pada sebagian komentar dengan distribusi seimbang yakni 1:1:1. Tabel 3. menunjukkan skema pelabelan yang digunakan.

Tabel 3. Skema Pelabelan

Label	Kode Numerik	Deskripsi	Kriteria
Positif	1	Komentar yang mendukung, memuji atau optimis terhadap program	Mengandung kata-kata pujian, Ekspresi harapan positif, Dukungan eksplisit, Testimoni positif
Negatif	0	Komentar yang mengkritik, menolak, atau pesimis terhadap program	Kritik langsung, Ekspresi kekecewaan, Sarkasme/sinisme, Penolakan eksplisit
Netral	2	Komentar informatif, bertanya atau tidak memiliki stance jelas	Bersifat deskriptif, Pertanyaan factual, Informasi tanpa opini, Balance positif-negatif

Tabel 3. mempresentasikan kerangka klasifikasi tiga kelas sentimen yang dikembangkan khusus untuk konteks kebijakan publik. Setiap kelas didefinisikan berdasarkan dimensi afektif (dukungan/kritik/informasi), kode numerik untuk pemrosesan komputasi, dan kriteria operasional yang dapat diobservasi. Kriteria multi-layer mencakup indikator leksikal (kata pujian/kritik), konteks kalimat, dan intensitas emosi. Pendekatan ini memastikan konsistensi antar-annotator dan reproducibility proses pelabelan.

Pembagian dataset dilakukan dengan metode hold-out split menggunakan rasio 80:20, di mana 80% data dialokasikan sebagai training set untuk fine-tuning model IndoBERT dan 20% sebagai test set untuk evaluasi independen [5]. Pendekatan stratified splitting diterapkan untuk mempertahankan proporsi kelas sentimen yang identik pada kedua subset, menghindari bias evaluasi akibat ketidakseimbangan distribusi. Metode ini dipilih karena ukuran dataset yang memadai sehingga pembagian statis sudah cukup representatif tanpa memerlukan teknik seperti k-fold cross-validation. Dengan demikian, test set berfungsi sebagai data yang benar-benar tidak terlihat oleh model selama proses pelatihan, sehingga hasil evaluasi mencerminkan kemampuan generalisasi model secara objektif.

2.5 Konfigurasi Model dan Training

Model yang digunakan adalah IndoBERT-base (indobenchmark/indobert-base-p1) [13] dengan arsitektur 12-layer transformer, 768 hidden size, dan 12 attention heads. Penggunaan IndoBERT sebagai model backbone telah terbukti efektif untuk berbagai tugas NLP bahasa Indonesia, termasuk deteksi sumber stres dalam teks media [9]. Tabel 4. menunjukkan konfigurasi hyperparameter.

Tabel 4. Konfigurasi Hyperparameter

Parameter	Nilai	Justifikasi
Epochs	5	Cukup untuk konvergensi tanpa overfitting
Batch Size (Train)	16	Sesuai kapasitas GPU, menjaga stabilitas gradien
Batch Size (Eval)	32	Lebih besar karena hanya forward pass
Learning Rate	2×10^{-5}	Optimal untuk fine-tuning BERT
Warmup Steps	500	Stabilisasi learning di awal training
Weight Decay	0.01	Regularisasi L2 untuk mencegah overfitting
Gradient Accumulation	2 steps	Simulasi batch size lebih besar
Mixed Precision	FP16	Percepatan komputasi, hemat memori
Evaluation Strategy	Per epoch	Monitoring performa berkala

Parameter	Nilai	Justifikasi
Save Strategy	Best F1-score	Simpan model terbaik
Early Stopping	Patience 2	Hentikan jika tidak ada improvement
Random Seed	42	Memastikan hasil dapat direproduksi

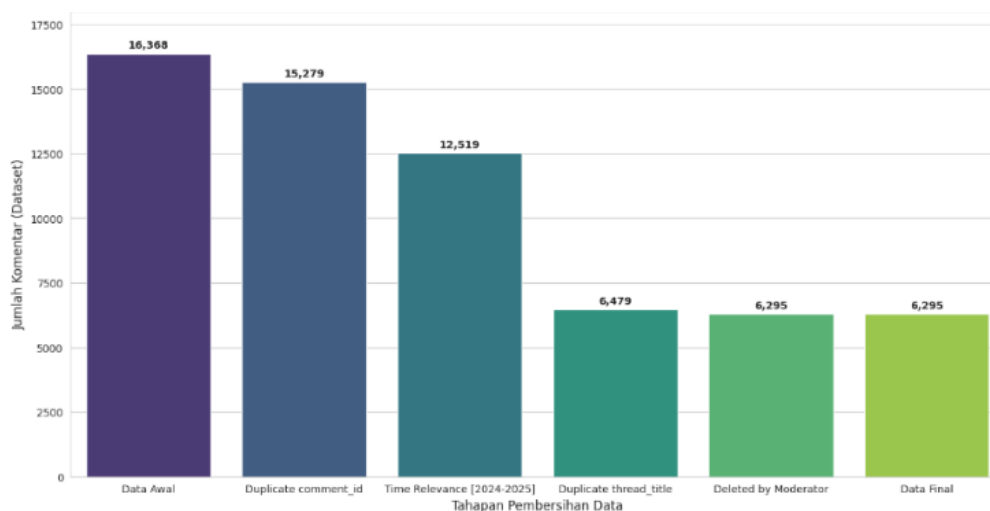
Tabel 4. menyajikan secara rinci 12 hiperparameter yang dioptimalkan dalam proses fine-tuning IndoBERT. Parameter-parameter tersebut diorganisasikan ke dalam tiga kelompok utama yang saling melengkapi. Kelompok pertama mencakup parameter pelatihan dasar seperti jumlah epochs, ukuran batch, dan learning rate yang dirancang secara hati-hati untuk menjaga keseimbangan antara konvergensi model dan efisiensi komputasi. Kelompok kedua berisi strategi optimasi, termasuk warmup, weight decay, gradient accumulation, dan mixed precision, yang berfungsi memastikan stabilitas proses pelatihan sekaligus mempercepat kinerja komputasi. Sementara itu, kelompok ketiga berkaitan dengan strategi evaluasi dan penyimpanan model, yang menjamin bahwa model dengan performa terbaik dapat terdokumentasi secara sistematis dan terstruktur.

Penetapan nilai setiap hiperparameter didasarkan pada kajian literatur mengenai fine-tuning IndoBERT serta disesuaikan dengan karakteristik khusus dari dataset yang digunakan. Dengan konfigurasi tersebut, diharapkan proses pelatihan dapat mencapai konvergensi yang optimal, sehingga meningkatkan kemampuan model dalam menangkap nuansa serta kedalaman sentimen publik secara lebih akurat dan bermakna.

III. HASIL DAN PEMBAHASAN

3.1 Hasil Pengumpulan

Proses pengumpulan berhasil mengumpulkan 16.268 komentar yang kemudian melalui tahapan pembersihan sistematis. Gambar 4. menunjukkan visualisasi penyaringan dataset.



Gambar 4. Visualisasi Penyaringan Dataset

Gambar 4. menampilkan proses reduksi data dari 16.268 komentar awal menjadi 6.295 data final melalui lima tahapan filtering. Tahap 1: Eliminasi duplikat berdasarkan comment_id mengurangi 989 komentar (6.1%). Tahap 2: Filter temporal 2024-2025 menghapus 2.760 komentar (18%) di luar rentang waktu. Tahap 3: Eliminasi duplikat thread_title mengurangi 6.040 komentar (48.2%) untuk menghindari bias topik. Tahap 4: Penghapusan komentar yang dihapus moderator (184 komentar, 2.8%). Hasil akhir 6.295 komentar (38.7% dari data awal) merepresentasikan dataset berkualitas tinggi untuk analisis. Grafik ini menunjukkan pentingnya kurasi data untuk validitas penelitian.

Dataset mentah yang diperoleh dari proses scraping disimpan dalam format CSV (.csv) dengan struktur yang terorganisir. Gambar 5. menunjukkan tampilan dataset mentah pada Microsoft Excel.

Gambar 5. Dataset Mentah pada Microsoft Excel

Gambar 5. menampilkan dataset mentah yang terbuka di Microsoft Excel. Terlihat kolom-kolom utama: comment_id (A), keyword_pencarian (B), subreddit (C), thread_title (D), thread_url (E), comment_author (F), comment_body (G), tanggal_komentar (H), skor_komentar (I). Dataset ini berisi 16.268 baris data komentar Reddit yang berhasil diekstrak. Format tersebut memudahkan proses kurasi manual sebelum tahap preprocessing otomatis.

3.2 Hasil Preprocessing Data

Setelah tahap preprocessing menggunakan Python (Pandas), hasilnya diekspor kembali ke format Excel untuk dokumentasi dan validasi. Gambar 6. menunjukkan dataset hasil preprocessing.

Gambar 6. Dataset Hasil Preprocessing

Gambar 6. menampilkan dataset hasil preprocessing komentar dari platform Reddit yang telah melalui beberapa tahapan pembersihan dan transformasi data. Dataset ini berisi komentar-komentar pengguna yang diproses untuk analisis lebih lanjut yakni untuk keperluan klasifikasi sentimen.

3.3 Hasil Pelabelan dan Pembagian Data

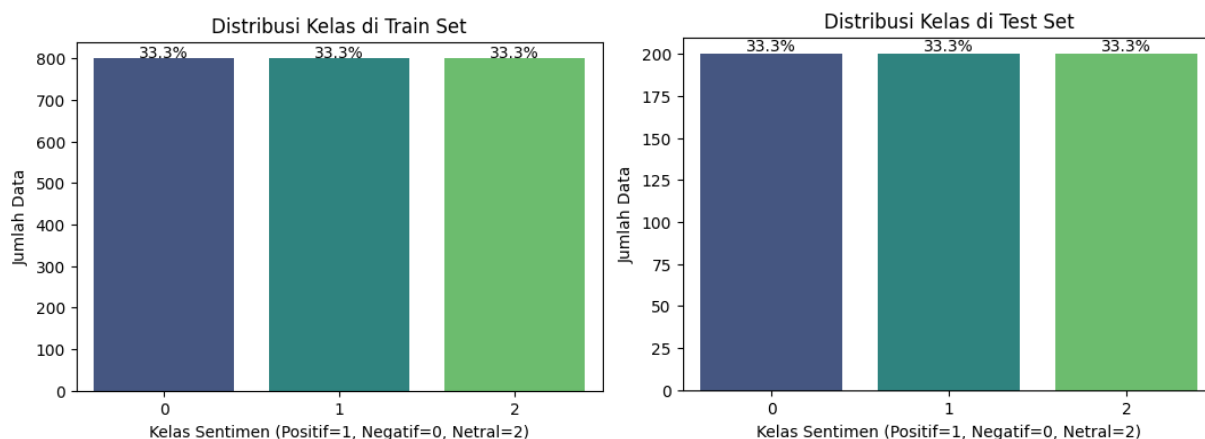
Proses pelabelan sentimen yang terdiri dari positif, negatif dan netral. terdapat juga kolom sentimen notes yang digunakan untuk mempermudah dalam menentukan sentimen dan sebagai acuan tambahan untuk model mempelajari sentimen. Pelabelan manual dilakukan pada 3.000 komentar dengan distribusi seimbang: 1.000 positif, 1.000 negatif, dan 1.000 netral. Gambar 7. menunjukkan hasil proses pelabelan.

J	K	L	M	N
comment cleaned	comment translated	sentiment manual	sentiment notes	
tidak semua anjing berbent	tidak semua anjing berbentuk	negatif	sindiran	
pegawai negri sipil nya ya j	pegawai negri sipil nya ya jug	netral	informasi	
tipe tipe orang belum pemat	tipe tipe orang belum pernah c	negatif	ofensif	
goblok banget masih pegav	goblok banget masih pegawai	negatif	sindiran	
tuh pemuja program makar	tuh pemuja program makan b	netral	opini	
the regime knows best [sm	rezim tahu yang terbaik smilir	netral	ambigu	
si paling hak asasi manusi	si paling hak asasi manusia	negatif	sindiran	
sudah ketauan belum sih i	sudah ketauan belum sih iden	netral	pertanyaan	
this bastard needs a reality	bajingan ini perlu sadar diri	negatif	sindiran	
ya tidak heran kalau besar	ya tidak heran kalau besarnya	netral	opini	
anjing geter tangan saya m	anjing geter tangan saya mau	negatif	ofensif	
his apology is pretty funny	permintaan maafnya sangat lu	negatif	sindiran	
kamu masih anak kecil w	kamu masih anak kecil wah b	negatif	sindiran	
bule hitam in every sense	bule hitam dalam segala hal y	negatif	ejekan	
berapa gaji pegawai negri	berapa gaji pegawai negri sipil	negatif	sindiran	
negara kerajaan dengan se	negara kerajaan dengan sega	negatif	sindiran	
menolak makan bergizi gra	menolak makan bergizi gratis	netral	opini	
what the fuck	apa apaan	netral	ambigu	
aih kakak ini tidak ada ot	aih kakak ini tidak ada otak k	negatif	sindiran	
yang tendang itu orang pap	yang tendang itu orang papua	netral	informasi	
saya kira tendang kecil se	saya kira tendang kecil sepe	negatif	kata kasar	
untung ada yang rekam kel	untung ada yang rekam kelak	negatif	ejekan	
itu kenapa ditendang memi	itu kenapa ditendang memanc	netral	pertanyaan	
sekolah menengah atas ya	sekolah menengah atas ya se	netral	pertanyaan	
was gonna say wah bisa j	hanya mau bilang wah bisa j	netral	opini	

Gambar 7. Dataset Hasil Proses Pelabelan

Gambar 7. menampilkan hasil tahapan pelabelan sentimen manual yang dilakukan setelah proses preprocessing data komentar. Dataset ini menunjukkan transformasi dari teks komentar yang telah dibersihkan dan diterjemahkan, yang kemudian melalui proses pelabelan sentimen yang terstruktur, siap digunakan untuk pelatihan model klasifikasi sentimen.

Pembagian data dilakukan dengan rasio 80:20 melalui teknik stratified splitting, guna mempertahankan proporsi kelas secara proporsional antarset pelatihan dan pengujian. Gambar 8. mengilustrasikan distribusi kelas pada kedua set data tersebut, yang memastikan representasi aspirasi publik secara merata dalam analisis sentimen.



Gambar 8. Distribusi Kelas pada Training Set dan Test Set

Gambar 8. menampilkan visualisasi distribusi kelas sentimen yang identik antara training set (kiri) dan test set (kanan). Kedua subset menunjukkan proporsi 33.3% untuk masing-masing kelas (positif, netral, negatif), mengkonfirmasi penggunaan metode stratified splitting [19] untuk proses training model menggunakan dataset lain. Training set berisi 2.400 data (800 per kelas) dan test set 600 data (200 per kelas). Distribusi seimbang antar kelas penting untuk menghindari bias pada model yang dapat muncul apabila salah satu kategori lebih dominan, berisiko membuat model hanya mempelajari pola tertentu [10].

3.4 Hasil Training Model

Model mencapai performa optimal pada epoch kelima dengan metrik yang sangat memuaskan. Gambar 9. menunjukkan hasil tiap epoch pada training set dan test set.

Epoch	Training Loss	Validation Loss	Accuracy	F1	Precision	Recall
1	1.080000	0.648165	0.880000	0.880287	0.882936	0.880000
2	0.104800	0.062956	0.980000	0.980054	0.980462	0.980000
3	0.041700	0.057665	0.986667	0.986661	0.986831	0.986667
4	0.029400	0.070039	0.983333	0.983331	0.983636	0.983333
5	0.026500	0.059211	0.990000	0.990000	0.990024	0.990000

Gambar 9. Hasil Training Model IndoBERT per Epoch

Gambar 9. mengilustrasikan dinamika training loss, validation loss, dan akurasi selama 5 epoch. Epoch 1 menunjukkan training loss 1.0800 dan validation loss 0.6482 dengan akurasi 88.00%, mengindikasikan tahap pembelajaran awal. Epoch 2 menunjukkan penurunan drastis training loss menjadi 0.1048 dan validation loss 0.0630, dengan akurasi melonjak ke 98.00% - bukti efektivitas transfer learning. Epoch 3 mencapai akurasi tertinggi sementara 98.67%. Epoch 4 menunjukkan tanda awal overfitting dengan validation loss naik ke 0.0700. Epoch 5 mencapai performa puncak: akurasi 99.00%, F1-score 99.00%, dengan training loss 0.0265 dan validation loss 0.0592. Pola konvergensi cepat dan stabil menunjukkan keunggulan IndoBERT untuk domain spesifik.

3.5 Evaluasi Model

Evaluasi menyeluruh pada test set berisi 600 data menghasilkan performa yang sangat memuaskan, mencerminkan keandalan model dalam menangkap dinamika sentimen publik. Tabel 5. menyajikan metrik evaluasi secara keseluruhan, yang memberikan gambaran komprehensif tentang kemampuan model untuk memahami nuansa aspirasi masyarakat.

Tabel 5. Hasil Evaluasi Model pada Test Set

Metrik	Nilai
Evaluation Loss	0.0592
Accuracy	0.9900 (99.00%)
F1-Score	0.9900 (99.00%)
Precision	0.9900 (99.00%)
Recall	0.9900 (99.00%)
Evaluation Runtime	4.35 detik
Samples per Second	137.93

Tabel 5. menyajikan tujuh metrik evaluasi yang mengkonfirmasi performa model yang unggul. Evaluation loss 0.0592 yang konsisten dengan validation loss menunjukkan tidak adanya overfitting. Accuracy 99.00% berarti hanya 6 dari 600 data test yang salah klasifikasi. F1-score, precision, dan recall [16] yang identik (99.00%) mengindikasikan keseimbangan sempurna antara ketepatan prediksi positif dan kemampuan deteksi seluruh kasus positif. Efisiensi komputasi ditunjukkan oleh throughput 137.93 sampel/detik, memungkinkan pemrosesan 6.295 komentar dalam ~46 detik - karakteristik kritis untuk aplikasi real-time.

```

=== Classification Report ===
              precision    recall  f1-score   support

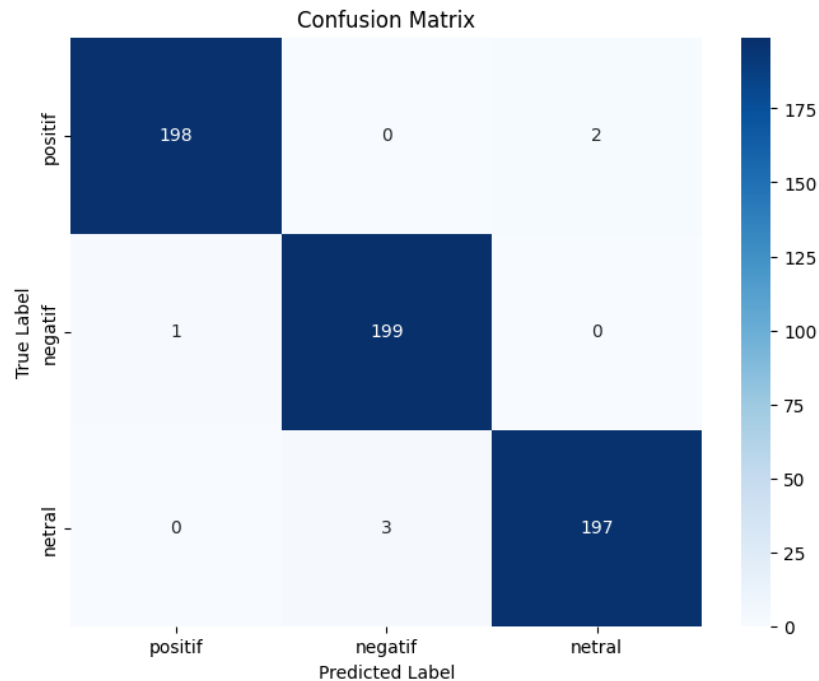
   positif         0.99         0.99         0.99         200
   negatif         0.99         0.99         0.99         200
    netral         0.99         0.98         0.99         200

 accuracy                   0.99         600
 macro avg         0.99         0.99         0.99         600
 weighted avg         0.99         0.99         0.99         600

```

Gambar 10. Classification Report

Gambar 10. Classification Report menyajikan precision, recall, F1-score, dan support untuk setiap kelas, dengan ketiga kelas mencapai nilai 0.99 pada precision dan F1-score, serta recall sebesar 0.98 untuk kelas netral dan 0.99 untuk kelas positif serta negatif. Macro average dan weighted average keduanya identik pada 0.99, yang menegaskan tidak adanya bias antar-kelas.

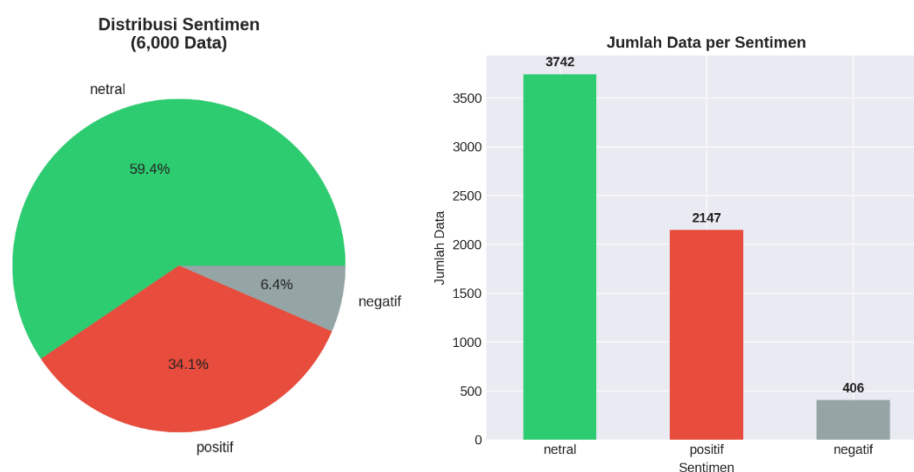


Gambar 11. Confusion Matrix

Gambar 11. Confusion Matrix menyajikan heatmap 3×3 yang mengkuantifikasi true positives masing-masing 198 (positif), 199 (negatif), dan 197 (netral), dengan hanya 6 kesalahan klasifikasi: 2 positif terdeteksi sebagai netral, 1 negatif sebagai positif, serta 3 netral sebagai negatif. Pola kesalahan ini mencerminkan ambiguitas alami pada batas antar-kelas, terutama netral-positif dan netral-negatif, yang menuntut pemahaman lebih dalam terhadap kompleksitas ekspresi publik

3.6 Hasil Prediksi pada Dataset Lengkap

Model diterapkan pada 6.295 komentar untuk analisis sentimen publik. Gambar 12. menunjukkan distribusi sentimen hasil prediksi.



Gambar 12. Distribusi Sentimen Hasil Prediksi

Gambar 12. Pie chart (kiri) merepresentasikan distribusi sentimen publik terhadap Program MBG di Reddit. Dominasi sentimen netral (59.44%) mengindikasikan karakteristik diskusi yang informatif dan analitis, di mana pengguna lebih banyak berbagi informasi faktual daripada opini eksplisit. Proporsi positif signifikan (34.11%) menunjukkan dukungan substansial terhadap program. Sentimen negatif minoritas (6.45%) dapat merefleksikan penerimaan umum atau adanya kritik implisit yang terdeteksi sebagai netral. Bar chart (kanan) memperjelas perbandingan sepenuhnya: 3.742 netral, 2.147 positif, dan 406 negatif. Visualisasi ini secara efektif mengkomunikasikan pola sentimen yang bebas dari bias, di mana mayoritas diskusi bersifat objektif namun tetap mencerminkan dukungan yang substansial terhadap Program MBG, sehingga memperkaya pemahaman kita akan aspirasi kolektif masyarakat.

J	K	L	M	N	O	P	Q
comment_cleaned	comment_translated	predicted_sentiment	confidence_score	prob_positif	prob_negatif	prob_netral	prob_diff
tidak semua anjing bertidak semua anjing bertidak	netral	0.996797979	0.001332977	0.001869099	0.996797979	0.99492888	
pegawai negri sipil nya pegawai negri sipil nya	netral	0.996615827	0.002238161	0.001146015	0.996615827	0.994377666	
tipe tipe orang belum p tipe tipe orang belum p	negatif	0.702251554	0.025141621	0.702251554	0.272606879	0.429644674	
goblok banget masih pegoblok banget masih pe	negatif	0.964005709	0.004417186	0.964005709	0.031577043	0.932428665	
tuh pemuja program mi tuh pemuja program mi	netral	0.551222324	0.435541451	0.013236193	0.551222324	0.115680873	
the regime knows best rezim tahu yang terbaik	netral	0.991795003	0.00317472	0.005030277	0.991795003	0.986764727	
si paling hak asasi manusia si paling hak asasi manusia	netral	0.994898975	0.002844419	0.002256655	0.994898975	0.992054556	
sudah ketahuan belum si sudah ketahuan belum si	netral	0.997010946	0.001757625	0.001231429	0.997010946	0.995253321	
this bastard needs a reabjangan ini perlu sadar	netral	0.834585249	0.103731878	0.06168285	0.834585249	0.730853371	
ya tidak heran kalau beuya tidak heran kalau beu	positif	0.9380216	0.9380216	0.009630748	0.052342615	0.885673985	
anjing geter tangan sayi anjing geter tangan sayi	negatif	0.970774472	0.007832997	0.970774472	0.021392467	0.949382005	
his apology is pretty fur permintaan maafnya sai	positif	0.476003289	0.476003289	0.426888764	0.09710791	0.049114525	
kamu masih anak kecil kamu masih anak kecil	netral	0.911578238	0.020361504	0.068060242	0.911578238	0.843517996	
bule hitam in every sen bule hitam dalam segali	netral	0.489782453	0.228814706	0.281402767	0.489782453	0.208379686	
berapa gaji pegawai ne berapa gaji pegawai ne	netral	0.995951056	0.002023969	0.002024917	0.995951056	0.993926139	
negara kerajaan dengari keadaan kerajaan dengi	netral	0.995735407	0.002201657	0.00206287	0.995735407	0.99353375	
menolak makan bergizi menolak makan bergizi	negatif	0.965476573	0.017822923	0.965476573	0.016700443	0.947653649	
what the fuck apa apan	negatif	0.939988971	0.027709411	0.939988971	0.032301661	0.90768731	
aih kakak ini tidak ada o aih kakak ini tidak ada o	negatif	0.784316599	0.123596601	0.784316599	0.09208677	0.660719998	
yang tendang itu orang yang tendang itu orang	netral	0.4827438	0.339879185	0.177376986	0.4827438	0.142864615	
saya kira tendang kecil saya kira tendang kecil	netral	0.995604396	0.002887208	0.001508477	0.995604396	0.992717188	
untung ada yang rekam untung ada yang rekam	negatif	0.790052772	0.018405233	0.790052772	0.191542029	0.598510742	
itu kenapa ditendang mi itu kenapa ditendang mi	netral	0.993892312	0.004090228	0.002017446	0.993892312	0.989802084	

Gambar 13. Hasil Proses Tahap Prediksi

Gambar 13. menampilkan output akhir dari proses prediksi sentimen menggunakan model machine learning yang diterapkan pada seluruh dataset (6.295 komentar). Dataset ini merepresentasikan tahapan inferensi di mana model yang telah dilatih membuat prediksi sentimen secara otomatis beserta tingkat kepercayaannya.

3.7 Analisis Confidence Score

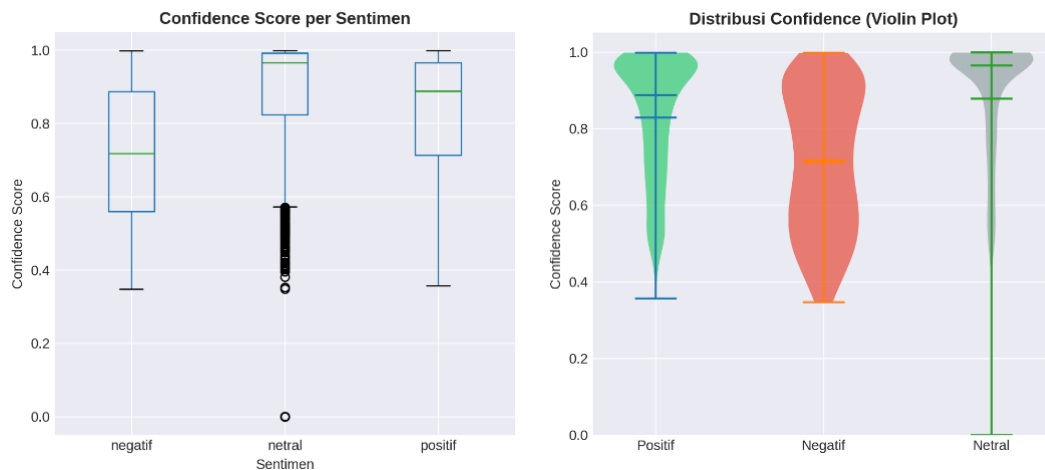
Confidence score mengindikasikan tingkat kepercayaan model terhadap prediksi. Tabel 6 menunjukkan statistik confidence score.

Tabel 6. Statistik Confidence Score pada Dataset Lengkap

Metrik	Nilai
Rata-rata (Mean)	0.8502
Median	0.9332
Standar Deviasi	0.1777
Minimum	0.0000*
Maximum	0.9985

*Catatan: Nilai minimum 0.0000 mengindikasikan adanya kasus di mana model memberikan probabilitas hampir sama untuk semua kelas (ambiguous cases)

Tabel 6. mengkuantifikasi reliabilitas prediksi model. Mean 0.8502 dan median 0.9332 yang lebih tinggi mengindikasikan distribusi right-skewed, di mana mayoritas prediksi memiliki confidence tinggi namun ada subset dengan confidence rendah yang menarik mean ke bawah. Standar deviasi 0.1777 menunjukkan variabilitas moderat. Rentang skor kepercayaan dari 0,0000 hingga 0,9985 meliputi spektrum lengkap mulai dari ketidakpastian total (ambiguitas ekstrem) hingga keyakinan yang nyaris mutlak. Interpretasi distribusi menunjukkan 57,27% prediksi dengan kepercayaan $\geq 0,9$ (sangat andal), 22,47% dalam rentang 0,7-0,9 (cukup andal), dan 20,25% di bawah 0,7 (memerlukan verifikasi manual), sehingga memberikan landasan pengambilan keputusan yang bijaksana dalam menangkap nuansa sentimen masyarakat.



Gambar 14. Confidence Score per Kategori Sentimen

Gambar 14. terdiri dari box plot (kiri) dan violin plot (kanan) yang membandingkan distribusi confidence antar-kelas. Box plot menunjukkan: netral memiliki median tertinggi (0.97) dan IQR (Interquartile Range) sempit, positif median 0.89 dengan variabilitas moderat, negatif median 0.72 dengan rentang terluas dan banyak outlier rendah. Violin plot mengkonfirmasi: pola distribusi kepercayaan secara jelas: kelas netral terkonsentrasi di atas 0,9 (mudah diklasifikasikan), kelas positif menampilkan bimodalitas dengan puncak pada 0,9 dan 0,6, sementara kelas negatif memiliki penyebaran paling luas dengan densitas signifikan di bawah 0,6. Temuan ini menggarisbawahi bahwa sentimen negatif khususnya sarkasme dan kritik tersirat merupakan tantangan utama model, yang sekaligus mencerminkan kekayaan nuansa dalam ekspresi aspirasi masyarakat.

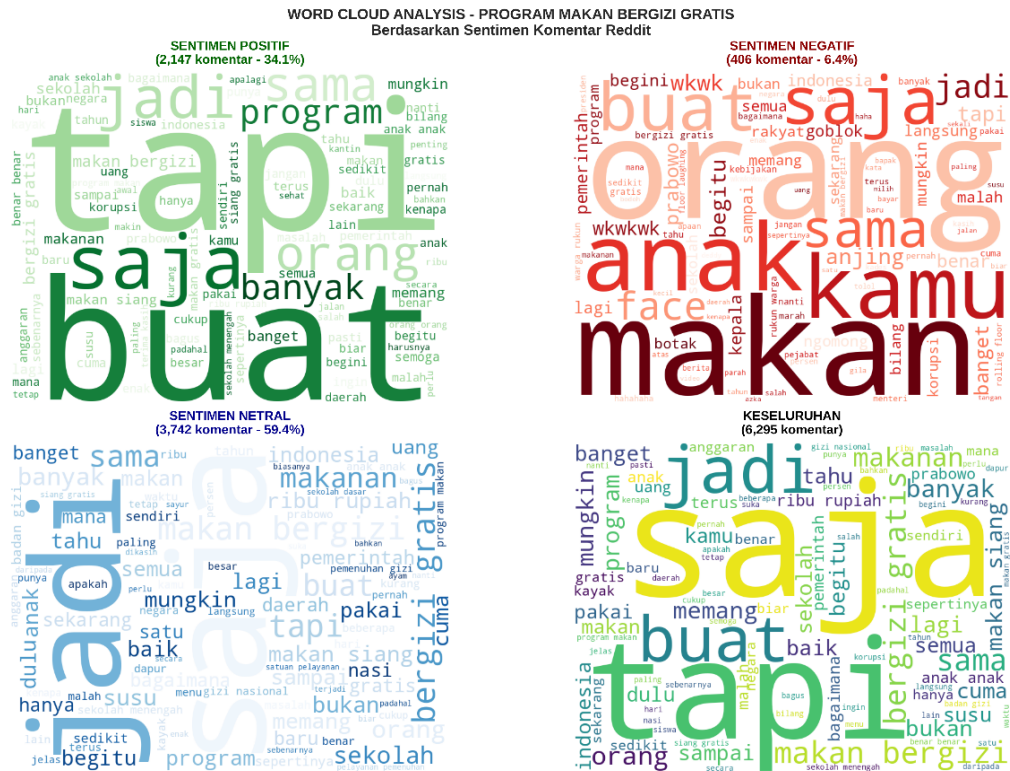
3.8 Analisis Kata Kunci dan Pola Linguistik

Analisis kata kunci mengungkap karakteristik linguistik per kategori sentimen. Tabel 7. menunjukkan top 10 kata kunci per sentimen.

Tabel 7. Top 10 Kata Kunci per Kategori Sentimen

Ranking	Sentimen Positif (Freq.)	Sentimen Negatif (Freq.)	Sentimen Netral (Freq.)
1	makan (866)	makan (54)	makan (1.132)
2	buat (587)	orang (45)	gratis (768)
3	gratis (538)	face (37)	sekolah (693)
4	sekolah (535)	anak (34)	saja (647)
5	tapi (534)	kamu (33)	bergizi (559)
6	saja (532)	saja (31)	jadi (548)
7	orang (505)	buat (31)	makanan (524)
8	jadi (468)	sama (30)	tapi (494)
9	anak (467)	jadi (27)	anak (492)
10	program (376)	anjing (25)	sama (487)

Tabel 7. membandingkan frekuensi leksikal dominan antar-tiga kategori sentimen. Kata "makan" mendominasi di semua kategori namun dengan frekuensi berbeda: 1.132 (netral), 866 (positif), 54 (negatif) - mencerminkan volume diskusi. Kategori positif ditandai kata kerja operasional ("buat", "jadi") dan istilah program. Kategori negatif menunjukkan karakteristik unik: frekuensi rendah, kehadiran "anjing" (vulgarisme), "face" (emoji negatif), dan "kamu" (konfrontasi). Kategori netral didominasi terminologi teknis ("bergizi", "makanan", "sekolah") dan konektor diskursif ("saja", "tapi", "jadi"). Analisis linguistik dilakukan secara sistematis untuk membedah pola bahasa intrinsik, tema diskursif emergent, serta struktur semantik yang mewarnai komentar Reddit [20].



Gambar 15. Word Cloud Analysis - Program Makan Bergizi Gratis

Gambar 15. menyajikan empat word cloud yang kaya informasi: sentimen positif (kiri atas, hijau), negatif (kanan atas, merah), netral (kiri bawah, biru), dan keseluruhan (kanan bawah, multi-warna), dengan ukuran font yang sebanding dengan frekuensi kemunculan kata. Word cloud positif menampilkan kata-kata fungsional seperti "buat", "tapi", "saja", "banyak" yang membentuk argumentasi konstruktif; negatif didominasi "orang", "anak", "kamu", "makan", "saja" dengan warna merah yang mencerminkan intensitas emosi; sementara netral mengandung pola informatif "saja", "tapi", "jadi", "gratis", "sekolah". Visualisasi keseluruhan mengintegrasikan tema-tema ini melalui "saja", "buat", "tapi", "jadi" sebagai pusat bahasa, memfasilitasi pemahaman intuitif terhadap tema dominan dan memperdalam wawasan kita akan dinamika percakapan masyarakat.

Tabel 8. Statistik Panjang Teks per Sentimen

Sentimen	Rata-rata (kata)	Median (kata)	Min (kata)	Max (kata)
Positif	31.4	18	1	355
Negatif	12.6	7	1	275
Netral	25.4	15	1	763

Tabel 8. mengkarakterisasi pola komunikasi antar-sentimen berdasarkan kompleksitas teks. Sentimen positif memiliki rata-rata tertinggi (31.4 kata) dengan disparitas mean-median signifikan, mengindikasikan adanya outlier komentar panjang berisi testimoni detail atau analisis komprehensif. Sentimen negatif paling ringkas (12.6 kata, median 7), mencerminkan ekspresi impulsif-emosional atau kritik langsung. Netral menempati posisi menengah (25.4 kata) dengan maksimum ekstrem (763 kata), merefleksikan keragaman dari pertanyaan singkat hingga analisis teknis mendalam. Pola ini konsisten dengan karakteristik platform: Reddit memfasilitasi diskusi panjang untuk dukungan dan informasi, namun reaksi negatif cenderung singkat.

IV. KESIMPULAN

Penelitian ini berhasil mengimplementasikan fine-tuning IndoBERT untuk klasifikasi sentimen berbahasa Indonesia dengan capaian performa yang sangat tinggi, ditunjukkan melalui akurasi, F1-score, dan presisi masing-masing sebesar 99,00% tanpa adanya indikasi bias klasifikasi. Analisis terhadap 6.295 komentar di platform Reddit

memperlihatkan dominasi sentimen netral sebesar 59,44%, diikuti oleh sentimen positif 34,11% dan sentimen negatif 6,45%. Pola distribusi ini mencerminkan karakter wacana publik mengenai Program Makan Bergizi Gratis yang cenderung informatif sekaligus analitis. Selain itu, model menunjukkan reliabilitas yang kuat dengan skor kepercayaan rata-rata 0,8502 serta kecepatan pemrosesan mencapai 137,93 sampel per detik, sehingga layak diposisikan sebagai instrumen pemantauan kebijakan publik secara waktu nyata. Kendati demikian, tantangan dalam mendeteksi sarkasme dan kritik implisit masih menyisakan ruang pengembangan lebih lanjut bagi penyempurnaan model di masa mendatang.

V. PENELITIAN SELANJUTNYA

Berdasarkan temuan serta keterbatasan penelitian, sejumlah rekomendasi pengembangan ke depan dapat diajukan. Pertama, perluasan sumber data ke berbagai platform digital untuk memperkaya representasi opini publik. Kedua, integrasi IndoBERT dengan sistem berbasis aturan guna meningkatkan kemampuan deteksi sarkasme. Ketiga, penerapan Aspect-Based Sentiment Analysis (ABSA) untuk mengungkap sentimen pada aspek spesifik Program Makan Bergizi Gratis. Keempat, pengembangan dashboard pemantauan secara waktu nyata agar hasil analisis dapat diakses secara praktis. Kelima, pelaksanaan penelitian longitudinal berbasis time-series untuk menelusuri evolusi opini publik sepanjang siklus kebijakan, sehingga dapat mendukung perumusan kebijakan yang lebih adaptif dan responsif.

UCAPAN TERIMA KASIH

Penulis menyampaikan ucapan terima kasih yang sebesar-besarnya kepada rekan-rekan mahasiswa yang telah memberikan dukungan moral serta diskusi ilmiah yang sangat berharga sepanjang proses penelitian klasifikasi sentimen ini. Ucapan terima kasih yang tulus juga disampaikan kepada para dosen dan ahli yang dengan murah hati berbagi masukan mendalam untuk memperkaya analisis dan metodologi. Dukungan praktis dari rekan-rekan dalam pelaksanaan eksperimen, disertai motivasi tak kenal lelah dari orang tua, menjadi pondasi utama keberhasilan penelitian ini.

REFERENSI

- [1] B. Auxier and M. Anderson, "Social Media Use in 2021," Pew Research Center, 2021. [Online]. Available: www.pewresearch.org.
- [2] K. Pham, K. C. Rao Kathala, and S. Palakurthi, "Reddit Sentiment Analysis on the Impact of AI Using VADER, TextBlob, and BERT," *Procedia Computer Science*, vol. 258, pp. 85-94, 2025.
- [3] A. Kiftiyah et al., "Program Makan Bergizi Gratis (MBG) dalam Perspektif Keadilan Sosial dan Dinamika Sosial-Politik," *Pancasila: Jurnal Keindonesiaan*, vol. 5, no. 1, pp. 45-62, 2025.
- [4] R. A. Munir, "Analisis Sentimen Cuitan di Media Sosial X tentang Program Makan Bergizi Gratis dengan Metode NLP," *Jurnal Informatika Dan Teknik Elektro Terapan*, vol. 13, no. 1, pp. 123-135, 2025.
- [5] V. Agustina, A. Herliana, and E. P. Korespondensi, "Analisis Sentimen Publik atas Kebijakan Efisiensi Anggaran 2025 dengan Text Mining dan Natural Language Processing," *Jurnal Sistem Informasi dan Teknologi*, vol. 7, no. 2, pp. 78-89, 2025.
- [6] B. Boe, "PRAW: The Python Reddit API Wrapper," GitHub Documentation, 2023. [Online]. Available: <https://praw.readthedocs.io/en/stable/index.html>
- [7] M. Rayhan Nur, Y. Wibisono, and R. Megasari, "Analisis Sentimen dan Pemodelan Topik pada Post tentang Merek Teknologi di X Menggunakan Fine-tuning IndoBERT dan BERTopic," *Jurnal Komputer Teknologi Informasi Sistem Informasi (JUKTISI)*, vol. 4, no. 2, pp. 45-58, 2025.
- [8] Babanejad, N., Agrawal, A., An, A., & Papagelis, M. (2020). A Comprehensive Analysis of Preprocessing for Word Representation Learning in Affective Tasks. <https://github.com/NastaranBa/>
- [9] M. Faza and M. Taufik, "Penerapan Model IndoBERT Untuk Deteksi Potensi Sumber Stres Dalam Teks Media," *Jurnal Teknologi Informasi dan Ilmu Komputer*, vol. 12, no. 3, pp. 567-578, 2025.
- [10] A. Kunaefi, Z. Abidin, and R. Kusumawati, "Klasifikasi Berita Hoaks Bahasa Indonesia Menggunakan IndoBERT Fine-Tuning Dengan Pendekatan Focal Loss pada Data Tidak Seimbang," *JUPI (Jurnal Ilmiah Penelitian Dan Pembelajaran Informatika)*, vol. 10, no. 2, pp. 89-102, 2025.
- [11] D. Prasetya et al., "Analisis Sentimen Pengguna Aplikasi MyBluebird Dengan Algoritma Naïve Bayes Di Playstore," *Jurnal Informatika Dan Teknik Elektro Terapan*, vol. 13, no. 2, pp. 145-156, 2025.

- [12] N. Proferes et al., "Studying Reddit: A Systematic Overview of Disciplines, Approaches, Methods, and Ethics," *Social Media and Society*, vol. 7, no. 2, pp. 1-14, 2021.
- [13] W. Wang et al., "IndoBERT: Pre-trained Model for Bahasa Indonesia," IndoNLP Research Group, 2020. [Online]. Available: <https://huggingface.co/indobenchmark>
- [14] Z. Liu et al., "Improving Sentiment Analysis Accuracy with Emoji Embedding," *Journal of Safety Science and Resilience*, vol. 2, no. 4, pp. 289-298, 2021.
- [15] M. Arif Afandi and I. Suri, "Effectiveness of Government Communication in Delivering Public Policy on Social Media," *Jurnal Komunikasi Pemerintah*, vol. 9, no. 1, pp. 23-38, 2025.
- [16] N. Ratnaswari, N. C. Wibowo, and D. S. Y. Kartika, "Analisis Sentimen Menggunakan Metode Lexicon-Based Dan Support Vector Machine pada Presiden dan Wakil Presiden Indonesia Periode 2024–2029," *Jurnal Informatika Dan Teknik Elektro Terapan*, vol. 13, no. 1, pp. 67-78, 2025.
- [17] S. S. Berutu et al., "Data preprocessing approach for machine learning-based sentiment classification," *JURNAL INFOTEL*, vol. 15, no. 4, pp. 317-325, 2023.
- [18] H. T. Duong and T. A. Nguyen-Thi, "Preprocessing techniques and data augmentation for sentiment analysis," *Computational Social Networks*, vol. 8, no. 1, pp. 1-15, 2021.
- [19] R. Srinivasan and C. N. Subalalitha, "Sentimental analysis from imbalanced code-mixed data using machine learning approaches," *Distributed and Parallel Databases*, vol. 41, no. 1-2, pp. 37-52, 2023.
- [20] A. Ari Putra Wibowo and Nurul Hidayat, "Eksplorasi Linguistik Komputasional dalam Analisis Bahasa Alami untuk Mengungkap Evolusi Dialek Digital di Era Media Sosial Global," *Journal of New Trends in Sciences*, vol. 1, no. 3, pp. 45-52, 2023.
- [21] Nurjoko and A. Rahardi, "Model Indo-BERT untuk Identifikasi Sentimen Kekerasan Verbal di Twitter," *IJCCS (Indonesian Journal of Computer and Cybernetics)*, vol. 18, no. 2, pp. 156-168, 2024.
- [22] F. A. Wagay and Jahiruddin, "Classification of Mental Illnesses from Reddit Posts Using Sentence-BERT Embeddings and Neural Networks," *Procedia Computer Science*, vol. 258, pp. 234-243, 2025.

Conflict of Interest Statement:

The author declares that the research was conducted in the absence of any commercial or financial relationships that could be construed as a potential conflict of interest.