

Penerapan BERT dalam Memetakan Opini Pengguna Instagram Terhadap Program Makan Bergizi Gratis

Rika Septiana¹, Ali Ibrahim^{2*}, Endang Lestari¹, Yadi Utama¹

¹Sistem Informasi, Ilmu Komputer, Universitas Sriwijaya, Indralaya, Indonesia

²Ilmu Komputer, Ilmu Komputer, Universitas Sriwijaya, Indralaya, Indonesia

Email: ¹septiana.rika09@email.com, ²aliibrahim@unsri.ac.id, ³endanglestari@unsri.ac.id, ⁴yadiutama@unsri.ac.id

Email Penulis Korespondensi: aliibrahim@unsri.ac.id

Submitted 04-12-2025; Accepted 01-01-2026; Published 19-02-2026

Abstrak

Makan Siang Gratis adalah program baru yang dibuat pemerintah untuk meningkatkan asupan gizi serta fokus belajar siswa selama mengikuti kegiatan pendidikan di sekolah. Kebijakan ini tidak hanya ditujukan untuk mendukung perkembangan kemampuan berpikir anak, tetapi juga menjadi langkah dalam mencegah masalah kesehatan seperti stunting dan meningkatkan kesejahteraan masyarakat. Seiring penerapannya, muncul berbagai respons dari masyarakat melalui media sosial, khususnya Instagram, baik berupa dukungan terhadap manfaat program maupun kritik mengenai pelaksanaannya. Untuk mengetahui sentimen masyarakat terhadap program tersebut, khususnya pada unggahan yang berasal dari wilayah Sumatera Selatan dengan menggunakan tagar #mbgsumsel. Data komentar yang berhasil dikumpulkan kemudian melalui rangkaian proses analisis, mulai dari pembersihan data, tahap pra-pemrosesan, pembagian data ke dalam set pelatihan, validasi, serta pengujian, fine-tuning hingga tahap evaluasi kinerja model. Model BERT dipilih karena memiliki kemampuan memahami konteks bahasa yang kompleks dan hubungan antarkata secara dua arah, sehingga memberikan hasil klasifikasi yang lebih akurat dibanding metode konvensional seperti Naïve Bayes, K-Nearest Neighbors, maupun Support Vector Machine. Hasil pengujian menunjukkan tingkat akurasi model mencapai 88%, dengan persebaran sentimen terdiri atas 751 komentar positif, 233 komentar netral, dan 257 komentar negatif. Semua nilai metrik tertinggi diperoleh pada kelas positif, yang menandakan bahwa model lebih efektif dalam mengidentifikasi dukungan publik. Melalui temuan ini, penelitian memberikan gambaran kuantitatif mengenai persepsi masyarakat di media sosial terhadap pelaksanaan program pemerintah tersebut.

Kata Kunci: Makan Siang Gratis; BERT; Analisis Sentimen; Media Sosial; Instagram

Abstract

The Free Lunch Program is a new government initiative designed to improve students' nutritional intake and enhance their focus during school activities. Its implementation has gained significant public attention, generating a variety of reactions across social media platforms, particularly Instagram. This study aims to examine public sentiment toward the program's rollout in South Sumatra by collecting user comments from posts using the hashtag #mbgsumsel. The collected comments were processed through several stages, including data cleaning, text pre-processing, dataset partitioning, fine-tuning and model evaluation. The BERT model was employed due to its strong capability in capturing contextual meaning within text, making it more effective than conventional classification approaches. Experimental results indicate that the model achieved an accuracy of 88% in classifying sentiments. The test dataset consisted of 751 positive comments, 233 neutral comments, and 257 negative comments. Overall, this study provides a quantitative overview of how the public perceives the Free Lunch Program based on their social media expressions.

Keywords: Free Lunch Program; BERT; Sentiment Analysis; Social Media; Instagram

1. PENDAHULUAN

Program Makan Siang Gratis (MBG) adalah program baru pemerintah berdasarkan janji kampanye calon presiden dan wakil presiden 2024 [1]. MBG adalah program yang dibuat pemerintah untuk meningkatkan kemampuan anak dalam berkonsentrasi pada pelajaran, memberikan dukungan kesehatan melalui makanan bergizi, mencegah stunting, dan memanfaatkan sumber daya pangan lokal [2]. Selain itu, program diharapkan mendorong pertumbuhan ekonomi karena makanan akan dibeli dari petani lokal [3]. Kebijakan yang optimal dan berkelanjutan mengenai program makan siang gratis dapat memberikan dukungan yang signifikan bagi ketahanan pangan nasional [4]. Namun, program ini menuai berbagai opini di masyarakat, hal ini karena masyarakat khawatir tentang berbagai masalah yang akan timbul, terutama mengenai anggaran yang sangat besar mencapai Rp71 triliun [5]. Selain itu dari sisi kualitas gizi, distribusi dan masalah lainnya yang dikhawatirkan masyarakat untuk kedepannya. Untuk mengetahui pasti pendapat pengguna instagram di Sumatera Selatan terhadap program MBG maka dilakukanlah analisis sentimen.

Analisis sentimen adalah cara untuk dapat mengetahui opini seseorang terhadap sesuatu [6]. Analisis sentimen dapat menunjukkan makna secara eksplisit maupun implisit [7]. Analisis sentimen merupakan tugas dasar *Natural Language Processing* (NLP) yang mengkategorikan sentimen berdasarkan positif, negatif, atau netral [8]. Untuk menganalisis sentimen diperlukannya suatu metode untuk memahami makna teks. Dalam penelitian ini digunakan model *Bidirectional Encoder Representations from Transformers* (BERT). BERT adalah pemrosesan bahasa alami NLP yang memungkinkan komputer memahami teks secara mendalam layaknya manusia [9]. BERT memanfaatkan arsitektur transformer untuk memahami konteks secara dua arah [10]. Pengelompokan sentimen bisa dilakukan menggunakan metode tradisional seperti *Naïve Bayes*, *Support Vector Machine* (SVM), *Logistic Regression*, dan *Decision Tree*. Namun, seiring berkembangnya teknologi pembelajaran mendalam, metode modern seperti *Recurrent Neural Network* (RNN), *Convolutional Neural Network* (CNN), hingga model berbasis Transformer seperti BERT mulai banyak digunakan karena mampu memahami konteks bahasa secara lebih baik. BERT menggunakan jaringan saraf berlapis untuk membuat

representasi kata yang kontekstual dan fleksibel sesuai teks, sehingga berguna untuk kata yang memiliki banyak makna, serta memakai pra-pelatihan diikuti penyetalan untuk tugas tertentu [11]. BERT efektif untuk menganalisis sentimen di media sosial karena memiliki pemahaman yang baik tentang bahasa secara keseluruhan dan dapat mencapai akurasi hingga 84% [12]. Selanjutnya, metode BERT dipilih karena melampaui metode tradisional seperti *Naïve Bayes*, *Support Vector Machine*, dan *K-Nearest Neighbors* (KNN) dengan akurasi mencapai 84,69% [13]. Metode seperti *Naïve Bayes* dan *Long Short-Term Memory* (LSTM) memang bekerja cukup baik untuk mengklasifikasikan teks, tetapi memiliki keterbatasan dalam memahami konteks secara umum [14]. Pada penelitian mengenai analisis ulasan film, BERT mendapatkan nilai akurasi 96% [15]. Penelitian lain juga menunjukkan bahwa BERT mencapai akurasi, presisi, recall, dan skor F1 sebesar stabil sebesar 86% untuk analisis sentimen pengguna game [16]. Keunggulan BERT tersebut membuatnya banyak digunakan untuk analisis sentimen. Namun, dalam menganalisis sentimen biasanya dibutuhkan sumber data. Program MBG tentunya banyak didiskusikan, terutama di media sosial yang sering digunakan saat ini.

Media sosial merupakan sarana yang menyediakan kemudahan dalam berkomunikasi, berbagi informasi serta berinteraksi aktif [17]. Hal tersebut menjadi alasan utama mengapa media sosial banyak digunakan saat ini. Setiap media sosial pastinya memiliki karakteristik masing-masing untuk membedakannya dengan media sosial lain dan agar menarik lebih banyak pengguna. Media sosial tiktok memiliki karakteristik seperti short video, youtube dengan video panjang yang biasa digunakan untuk menonton film, X dengan tweet nya berupa teks yang sering menjadi ruang diskusi, instagram dengan karakteristik visualnya menjadikannya sering dipakai dalam hal *personal branding* dan media sosial lainnya. Di Indonesia sendiri pengguna dengan banyak 103 juta pada awal tahun 2025 adalah Instagram, yang menjadikannya sebagai salah satu media sosial paling populer [18]. Nama Instagram diambil dari kata 'insta', yang berdasarkan konsep fotografi gaya Polaroid [19]. Platform ini digunakan oleh berbagai kelompok seperti pemerintah, media, dan masyarakat umum yang semakin memperluas ruang diskusi dengan perspektif yang berbeda. Tingkat interaksi di Instagram tinggi karena suatu postingan menarik secara visual, yang dapat memengaruhi cara orang merespons isu-isu tertentu [20]. Instagram menyediakan fitur interaktif seperti *like*, komentar, *story*, dan lainnya. Konten Instagram memainkan peran penting dalam pembentukan opini publik [21]. Instagram berperan sebagai sarana naratif yang kuat untuk menyampaikan [22]. Pada penelitian ini pengambilan data dari Instagram berdasarkan hastag *mbgsumsel*. Penggunaan hastag ini karena suatu tren yang paling populer diciptakan oleh pengguna melalui hastag [23]. Jadi, instagram bukan hanya media sosial dengan karakteristik visual saja namun bisa menjadi ruang diskusi aktif yang tidak kalah dari platform media sosial lainnya.

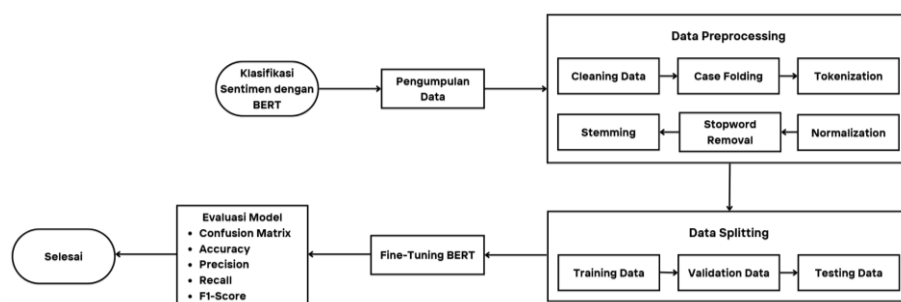
Tujuan dari penelitian adalah untuk mengetahui opini pengguna instagram terutama di Sumatera Selatan apakah positif, netral atau negatif dan seberapa banyak hasil komentar tersebut. Selain itu tujuan penelitian ini untuk mengukur seberapa baik model BERT melakukan analisis sentimen terhadap program makan siang gratis dengan menggunakan instagram yang jarang digunakan dalam analisis sentimen. Setelah mengetahui hasilnya, diharapkan dapat membantu pemerintah untuk mengetahui opini masyarakat berdasarkan data *real-time* dan menjadi masukan untuk program tersebut kedepannya.

Berdasarkan penelitian-penelitian sebelumnya, sebagian besar studi hanya berfokus pada analisis sentimen media sosial secara umum tanpa melakukan pemetaan opini secara spesifik terhadap Program MBG pada platform Instagram. Selain itu, mayoritas penelitian analisis sentimen lebih banyak memanfaatkan platform X sebagai sumber data karena kemudahan akses dan karakteristik teks yang lebih sederhana, sementara penggunaan Instagram masih relatif jarang dilakukan. Padahal, kolom komentar pada Instagram memiliki konteks opini yang lebih beragam dan dapat menggambarkan pandangan pengguna secara lebih mendalam terhadap suatu kebijakan. Oleh karena itu, penelitian ini berkontribusi dengan memanfaatkan komentar pengguna Instagram untuk menganalisis dan memetakan opini terhadap Program MBG menggunakan model BERT, sehingga diharapkan mampu memberikan sudut pandang yang berbeda dan lebih kontekstual dibandingkan penelitian sebelumnya.

2. METODOLOGI PENELITIAN

2.1 Tahapan Penelitian

Metodologi yang akan digunakan adalah metode BERT (*Bidirectional Encoder Representations from Transformers*). BERT mampu menganalisis konteks teks secara mendalam untuk klasifikasi teks (Kelana, 2025). Tahapan penelitian ini akan dijelaskan pada gambar di bawah ini.



Gambar 1. Tahapan Penelitian Klasifikasi Sentimen Model BERT

Gambar 1 merupakan tahapan penelitian yang akan dilakukan pada penelitian ini, dimulai dari pengumpulan data dari komentar instagram, *pre-processing* data, *splitting* data, *fine-tuning* BERT dan evaluasi model

- a. Pengumpulan Data: data akan dikumpulkan dengan teknik web *scraping*, yaitu data diambil secara otomatis dari platform digital. Dalam penelitian ini, data dikumpulkan dengan tagar #mbgsumsel di komentar Instagram yang menyebutkan program makan siang gratis di Sumatra Selatan. Postingan yang dikumpulkan berisi opini pengguna instagram mengenai implementasi program, opini publik, dan reaksi masyarakat terhadap program tersebut. Teknik ini dipilih karena memungkinkan pengumpulan data secara langsung dari pengguna media sosial. Data yang dikumpulkan dari komentar akan diproses dalam tahap selanjutnya *pre-processing*.
- b. *Pre-processing* Data: dilakukan agar dataset telah siap digunakan untuk tahap selanjutnya. *Pre-processing* data terdiri dari enam tahap yaitu, *cleaing*, *case folding*, tokenisasi, normalisasi, *stopword removal*, dan *stemming* [24]. Tahap ini dilakukan agar data siap digunakan dalam pemodelan.
 1. *Cleaning*, digunakan untuk menghilangkan karakter yang tidak valid.
 2. *Case folding*, menyeragamkan seluruh huruf menjadi huruf kecil.
 3. Tokenisasi, memecah kalimat menjadi unit kata yang lebih kecil.
 4. Normalisasi, mengubah kata tidak baku menjadi bentuk sesuai kaidah bahasa Indonesia.
 5. *Stopword removal*, kata yang tidak memiliki makna penting dihilangkan.
 6. *Stemming*, mengembalikan kata ke bentuk dasar untuk menyederhanakan analisis.
 Tujuan dilakukan tahapan ini agar membantu model memahami teks dengan lebih baik dan membuat prediksi sentimen menjadi lebih akurat.
- c. *Splitting* Data: setelah data selesai di *pre-processing*, data dibagi menjadi 70:20:10 untuk data latih, validasi, dan uji. Pembagian ini bertujuan melatih model, memantau kemungkinan *overfitting*, dan mengukur kemampuan model pada data baru. Dengan pembagian ini, model dapat menghasilkan hasil yang baik.
- d. *Fine-tuning* BERT: pada tahap ini, data komentar Instagram yang telah melalui proses *pre-processing* dan pembagian data dimasukkan ke dalam model BERT untuk dilakukan pelatihan ulang (*fine-tuning*). Pada proses ini, setiap komentar terlebih dahulu diubah ke dalam bentuk token melalui tokenizer BERT agar dapat diproses oleh model. Selanjutnya, model mempelajari pola teks pada data latih dan menyesuaikan parameternya sehingga mampu mengenali kecenderungan sentimen pada komentar. Hasil dari proses *fine-tuning* ini digunakan untuk memprediksi kategori sentimen pada data uji, yaitu positif, netral, atau negatif.
- e. Evaluasi Model: setelah proses pelatihan, validasi, dan pengujian selesai, kinerja model dievaluasi menggunakan beberapa metrik seperti akurasi, presisi, *recall*, skor f1 dan *confusion matrix*. Metrik-metrik tersebut memberikan informasi terkait kemampuan model dalam membedakan kelas sentimen, tingkat kesalahan yang terjadi, serta efektivitas dalam mengidentifikasi komentar positif, netral, maupun negatif.
 1. Akurasi digunakan untuk menunjukkan seberapa akurat prediksi sesuai dengan data uji. Metrik ini menilai seberapa banyak prediksi benar dibandingkan dengan total data uji. Semakin tinggi nilai presisi, semakin efektif model dalam mengenali pola data dan prediksi yang akurat.

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN} \quad (1)$$

2. Presisi mengukur seberapa tepat prediksi pada suatu kelas dibandingkan seluruh prediksi untuk kelas tersebut. Semakin tinggi nilainya, semakin sedikit kesalahan yang dilakukan model saat melabeli kelas tersebut.

$$Precision = \frac{TP}{TP + FP} \quad (2)$$

3. Recall menunjukkan kemampuan suatu model untuk mengidentifikasi semua label yang benar untuk suatu kelas dengan membandingkan prediksi yang benar dengan semua label aktual dalam kelas tersebut. Semakin tinggi nilai recall, semakin efektif model tersebut dalam mengambil semua data kelas tersebut.

$$Recall = \frac{TP}{TP + FN} \quad (3)$$

4. Skor f1 adalah rata-rata dari presisi dan recall. Semakin besar nilai *Skor f1* maka model dapat memprediksi kelas dengan lebih akurat dan menemukan semua data dalam kelas itu.

$$F1 - Score = \frac{2TP}{2TP + FP + FN} \quad (4)$$

Keterangan:

TP: True Positive, adalah kondisi ketika model berhasil memprediksi suatu data sebagai kelas tertentu dan prediksi tersebut memang benar sesuai dengan label aslinya.

TN: True Negative, merupakan keadaan ketika model memprediksi suatu data bukan bagian dari kelas tertentu, dan hasil tersebut juga benar.

FP: False Positive, ketika kesalahan muncul saat model memberikan label sebuah data ke kelas tertentu, meskipun data tersebut tidak seharusnya berada pada kelas itu.

FN: False Negative, kondisi ketika model memprediksi data bukan termasuk dalam suatu kelas tertentu, padahal data tersebut sebenarnya adalah bagian dari kelas tersebut
Keempat nilai di atas menjadi dasar dalam penyusunan *confusion matrix* untuk menilai kinerja model klasifikasi secara lebih detail.

3. HASIL DAN PEMBAHASAN

Sebelum melakukan tahapan *pre-processing* data, *splitting* data, dan evaluasi model BERT, langkah pertama yang harus dilakukan adalah proses pengumpulan data sebagai tahap pertama penelitian. Data yang digunakan dalam analisis berupa komentar pengguna Instagram di Sumatera Selatan yang membahas Program Makan Siang Gratis. Komentar yang telah dikumpulkan kemudian melewati tahapan *pre-processing* yang mencakup pembersihan teks, penyeragaman huruf, pemecahan kata, normalisasi, penghapusan stopword, dan stemming. Setelah itu, data dibagi ke dalam data pelatihan, validasi, serta pengujian untuk digunakan dalam pemodelan BERT. Lalu, model dievaluasi dengan beberapa metrik utama seperti presisi, *recall*, skor f1, akurasi dan *confusion matrix* guna memberikan informasi menyeluruh mengenai ketepatan klasifikasi sentimen yang dilakukan model.

3.1 Pengumpulan Data

Data yang berhasil terkumpul dengan *scarping* komentar berdasarkan hastag *mbg* adalah sebanyak 14.742 komentar. Jumlah postingannya sebanyak 93 postingan dari tanggal 15 Februari 2025 sampai 12 Oktober 2025. Berikut adalah contoh datasetnya.

Tabel 1. Dataset komentar MBG

No.	Text
1.	MBG diganti lapangan pekerjaan bagi ortu murid yg kurang mampu, lebih banyak manfaatnya.
2.	Penyiksaan....🤔🤔🤔🤔smoga ga knapa2
3.	Innalillahi wainna ilahi roji'un...aku.lebih setuju kebijakan bapa Gubernur Jabar bapa aing ...lebih baik uangnya ..jd biar ibunya yg memasak di rumah ..dengan nial uang segitu bisa di masakin untuk anak lebih higienis 🤔
4.	Paling higienis makanan dari Rumah masing masing
5.	Bener kata kang @dedimulyadi71 kasihkan ke ortunya aja uang makan itu nanti sampai disekolah mereka bisa tukaran lauknya dan terjamin kebersihannya...

Tabel 1 diatas merupakan contoh dataset berdasarkan komentar pengguna instagram terhadap postingan makan siang gratis. Dataset tersebut masih mengandung kalimat tidak baku, karakter khusus, dan elemen lain yang tidak relevan. Oleh karena itu, untuk memperoleh hasil analisis sentimen yang lebih akurat, dataset perlu melalui tahap *pre-processing* agar kualitas data menjadi lebih baik sebelum dilakukan analisis sentimen.

3.2 Pre-processing Data

Selanjutnya, data yang telah dikumpulkan akan dilakukan *pre-processing* untuk memastikan bahwa dataset berkualitas baik sebelum digunakan ke dalam proses pelatihan model. Proses ini sangat penting untuk mengurangi kesalahan model dalam memahami makna yang sama tetapi ditulis dalam gaya berbeda, terutama ketika terdapat variasi yang tinggi dalam gaya bahasa di media sosial. Setelah semua langkah *pre-processing* selesai, data yang telah bersih dan terstruktur bisa digunakan dalam proses pelatihan model BERT.

- Cleaning* data: tahapan ini akan menghilangkan elemen yang tidak sesuai seperti URL, hastag, karakter spesial dan elemen lainnya yang tidak memiliki makna khusus

	text	cleaning
0	MBG diganti lapangan pekerjaan bagi ortu murid yg kurang mampu, lebih banyak manfaatnya.	MBG diganti lapangan pekerjaan bagi ortu murid yg kurang mampu lebih banyak manfaatnya
1	Penyiksaan....🤔🤔🤔🤔smoga ga knapa2	Penyiksaan smoga ga knapa
2	Innalillahi wainna ilahi roji'un...aku.lebih setuju kebijakan bapa Gubernur Jabar bapa aing ...lebih baik uangnya ..jd biar ibunya yg memasak di rumah ..dengan nial uang segitu bisa di masakin untuk anak lebih higienis 🤔	Innalillahi wainna ilahi roji un aku lebih setuju kebijakan bapa Gubernur Jabar bapa aing lebih baik uangnya jd biar ibunya yg memasak di rumah dengan nial uang segitu bisa di masakin untuk anak lebih higienis
3	Paling higienis makanan dari Rumah masing masing	Paling higienis makanan dari Rumah masing masing
4	Bener kata kang @dedimulyadi71 kasihkan ke ortunya aja uang makan itu nanti sampai disekolah mereka bisa tukaran lauknya dan terjamin kebersihannya...	Bener kata kang kasihkan ke ortunya aja uang makan itu nanti sampai disekolah mereka bisa tukaran lauknya dan terjamin kebersihannya...

Gambar 2. Cleaning Data

- Case folding*: data yang ada dilakukan perubahan menjadi huruf kecil semua tanpa ada huruf kapital

	cleaning	case_folding
0	MBG diganti lapangan pekerjaan bagi ortu murid yg kurang mampu lebih banyak manfaatnya	mbg diganti lapangan pekerjaan bagi ortu murid yg kurang mampu lebih banyak manfaatnya
1	Penyiksaan smoga ga knpa	penyiksaan smoga ga knpa
2	Innalillahi wainna ilahi roji un aku lebih setuju kebijakan bapa Gubernur Jabar bapa aing lebih baik uangnya jd biar ibunya yg memasak di rumah dengan nial uang segitu bisa di masakin untuk anak lebih higienis	innalillahi wainna ilahi roji un aku lebih setuju kebijakan bapa gubernur jabar bapa aing lebih baik uangnya jd biar ibunya yg memasak di rumah dengan nial uang segitu bisa di masakin untuk anak lebih higienis
3	Paling higienis makanan dari Rumah masing masing	paling higienis makanan dari rumah masing masing
4	Bener kata kang kasihkan ke ortunya aja uang makan itu nanti sampai disekolah mereka bisa tukaran lauknya dan terjamin kebersihannya	bener kata kang kasihkan ke ortunya aja uang makan itu nanti sampai disekolah mereka bisa tukaran lauknya dan terjamin kebersihannya

Gambar 3. Case Folding

Gambar 3 merupakan tahapan *case folding*, dimana semua huruf kapital diubah menjadi huruf kecil, seperti kata “MBG” yang semuanya diubah menjadi “mbg”. Awal kata pada setiap kalimat juga diubah menjadi huruf kecil semua.

c. *Tokenization*: tahapan memecah teks menjadi unit-unit kecil, agar model memahami makna tiap kata lebih baik lagi

	case_folding	tokenisasi
0	mbg diganti lapangan pekerjaan bagi ortu murid yg kurang mampu lebih banyak manfaatnya	[mbg, diganti, lapangan, pekerjaan, bagi, ortu, murid, yg, kurang, mampu, lebih, banyak, manfaatnya]
1	penyiksaan smoga ga knpa	[penyiksaan, smoga, ga, knpa]
2	innalillahi wainna ilahi roji un aku lebih setuju kebijakan bapa gubernur jabar bapa aing lebih baik uangnya jd biar ibunya yg memasak di rumah dengan nial uang segitu bisa di masakin untuk anak lebih higienis	[innalillahi, wainna, ilahi, roji, un, aku, lebih, setuju, kebijakan, bapa, gubernur, jabar, bapa, aing, lebih, baik, uangnya, jd, biar, ibunya, yg, memasak, di, rumah, dengan, nial, uang, segitu, bisa, di, masakin, untuk, anak, lebih, higienis]
3	paling higienis makanan dari rumah masing masing	[paling, higienis, makanan, dari, rumah, masing, masing]
4	bener kata kang kasihkan ke ortunya aja uang makan itu nanti sampai disekolah mereka bisa tukaran lauknya dan terjamin kebersihannya	[bener, kata, kang, kasihkan, ke, ortunya, aja, uang, makan, itu, nanti, sampai, disekolah, mereka, bisa, tukaran, lauknya, dan, terjamin, kebersihannya]

Gambar 4. Tokenization

Gambar 4 diatas merupakan tahapan *tokenization*, dimana teks dipecah menjadi unit terkceil, seperti data nomor tiga “penyiksaan smga ga knpa” dipecah menjadi “penyiksaan, smga, ga, knpa”.

d. *Normalization*: kata tidak baku akan diubah menjadi baku, sesuai dengan pedoman Bahasa Indonesia

	tokenisasi	normalisasi
0	[mbg, diganti, lapangan, pekerjaan, bagi, ortu, murid, yg, kurang, mampu, lebih, banyak, manfaatnya]	[mbg, diganti, lapangan, pekerjaan, bagi, orang tua, murid, yang, kurang, mampu, lebih, banyak, manfaatnya]
1	[penyiksaan, smoga, ga, knpa]	[penyiksaan, semoga, tidak, kenapa]
2	[innalillahi, wainna, ilahi, roji, un, aku, lebih, setuju, kebijakan, bapa, gubernur, jabar, bapa, aing, lebih, baik, uangnya, jd, biar, ibunya, yg, memasak, di, rumah, dengan, nial, uang, segitu, bisa, di, masakin, untuk, anak, lebih, higienis]	[innalillahi, wainna, ilahi, roji, un, aku, lebih, setuju, kebijakan, bapak, gubernur, jawa barat, bapak, saya, lebih, baik, uangnya, jadi, agar, ibunya, yang, memasak, di, rumah, dengan, nilai, uang, segitu, bisa, di, masakin, untuk, anak, lebih, higienis]
3	[paling, higienis, makanan, dari, rumah, masing, masing]	[paling, higienis, makanan, dari, rumah, masing, masing]
4	[bener, kata, kang, kasihkan, ke, ortunya, aja, uang, makan, itu, nanti, sampai, disekolah, mereka, bisa, tukaran, lauknya, dan, terjamin, kebersihannya]	[bener, kata, kang, kasihkan, ke, orang tuanya, saja, uang, makan, itu, nanti, sampai, disekolah, mereka, bisa, tukaran, lauknya, dan, terjamin, kebersihannya]

Gambar 5. Normalization

Gambar 5 diatas merupakan tahapan *normalization*, dimana kata tidak baku akan diubah menjadi kata yang baku, seperti kata “ga” diubah menjadi “tidak”, “higenis” diubah menjadi “higienis”.

e. *Stopword removal*: tahapan untuk membuang kata yang tidak berkontribusi penting pada makna kalimat, misalnya “yang”, “ke”, dan “dan”.

	normalisasi	stopword_removal
0	[mbg, diganti, lapangan, pekerjaan, bagi, orang tua, murid, yang, kurang, mampu, lebih, banyak, manfaatnya]	[mbg, diganti, lapangan, pekerjaan, orang tua, murid, kurang, mampu, lebih, banyak, manfaatnya]
1	[penyiksaan, semoga, tidak, kenapa]	[penyiksaan, semoga]
2	[innalillahi, wainna, ilahi, roji, un, aku, lebih, setuju, kebijakan, bapak, gubernur, jawa barat, bapak, saya, lebih, baik, uangnya, jadi, agar, ibunya, yang, memasak, di, rumah, dengan, nilai, uang, segitu, bisa, di, masakin, untuk, anak, lebih, higienis]	[innalillahi, wainna, ilahi, roji, un, aku, lebih, setuju, kebijakan, bapak, gubernur, jawa barat, bapak, lebih, baik, uangnya, jadi, ibunya, memasak, rumah, nilai, uang, segitu, masakin, anak, lebih, higienis]
3	[paling, higienis, makanan, dari, rumah, masing, masing]	[paling, higienis, makanan, rumah, masing, masing]
4	[bener, kata, kang, kasihkan, ke, orang tuanya, saja, uang, makan, itu, nanti, sampai, disekolah, mereka, bisa, tukaran, lauknya, dan, terjamin, kebersihannya]	[bener, kata, kang, kasihkan, orang tuanya, uang, makan, disekolah, tukaran, lauknya, terjamin, kebersihannya]

Gambar 6. Stopword Removal

Gambar 5 diatas merupakan tahapan *stopword removal*, dimana kata yang tidak memiliki makna penting seperti kata “dari”, “ke”, “itu” dihilangkan.

- f. *Stemming*: tahapan menghapus imbuhan yang terletak di awal, akhir, dan disisipan. Jadi hanya ada kata dasarnya saja

	stopword_removal	stemming
0	[mbg, diganti, lapangan, pekerjaan, orang tua, murid, kurang, mampu, lebih, banyak, manfaatnya]	[mbg, ganti, lapang, kerja, orang tua, murid, kurang, mampu, lebih, banyak, manfaat]
1	[penyiksaan, semoga]	[siksa, moga]
2	[innalillahi, wainna, ilahi, roji, un, aku, lebih, setuju, kebijakan, bapak, gubernur, jawa barat, bapak, lebih, baik, uangnya, jadi, ibunya, memasak, rumah, nilai, uang, segitu, masakin, anak, lebih, higienis]	[innalillahi, wainna, ilahi, roji, un, aku, lebih, tuju, bijak, bapak, gubernur, jawa barat, bapak, lebih, baik, uang, jadi, ibu, masak, rumah, nilai, uang, segitu, masakin, anak, lebih, higienis]
3	[paling, higienis, makanan, rumah, masing, masing]	[paling, higienis, makan, rumah, masing, masing]
4	[bener, kata, kang, kasihkan, orang tuanya, uang, makan, disekolah, tukaran, lauknya, terjamin, kebersihannya]	[bener, kata, kang, kasih, orang tua, uang, makan, seko, tukar, lauk, jamin, bersih]

Gambar 7. Stemming

Gambar 7 diatas merupakan tahapan *stemming*, dimana kata seperti “memasak” menjadi “masak”, “makanan”, menjadi “makan”. Setelah data di *pre-processing*, data yang awalnya berjumlah 14.742 berkurang menjadi 14.051 data.

3.3 Splitting Data

Setelah seluruh tahapan pre-processing selesai dilaksanakan, dataset kemudian dipisahkan ke dalam tiga subset utama. Pemisahan ini diperlukan agar model tidak hanya belajar dari data *training* saja, namun juga kemampuannya diuji dalam data yang tidak pernah dilihat sebelumnya untuk melakukan generalisasi. Pada penelitian ini, komposisi pembagian data dilakukan secara proporsional, yaitu 70% sebagai data pelatihan sebanyak 9.835 komentar, 20% digunakan sebagai data validasi sejumlah 2.810 komentar untuk memantau kinerja model selama proses pelatihan, dan 10% sisanya diperuntukkan sebagai data pengujian sebanyak 1.406 komentar untuk menilai performa akhir model. Setiap subset memiliki peran yang saling mendukung dalam memastikan bahwa model yang dibangun memiliki kualitas prediksi yang baik serta mampu mengenali pola sentimen secara akurat.

Data pelatihan atau training data digunakan pada tahap awal untuk mengajarkan model BERT memahami pola hubungan antara isi teks dan label sentimennya. Pada proses ini, model akan belajar mengenai struktur bahasa, konteks yang muncul dalam kalimat, serta bagaimana sentimen positif, negatif, dan netral diekspresikan dalam komentar pengguna di Instagram. Semakin baik proses pembelajaran ini, semakin tinggi pula kemampuan model dalam mengenali sentimen secara tepat.

Selanjutnya, data validasi berperan sebagai pengawas selama proses pelatihan. Data ini digunakan untuk mengevaluasi apakah model *overfitting*, yaitu kondisi ketika model terlalu terpaku pada data latih sehingga fungsinya berkurang untuk mengenali pola baru pada data yang tidak dikenal. Melalui proses validasi, peneliti dapat melihat apakah hasil pelatihan terlalu bias terhadap data latih atau sudah secara optimal bekerja pada data baru. Jika performa data validasi menurun jauh dibandingkan data pelatihan, maka diperlukan tindakan seperti penyesuaian hyperparameter, penambahan data, atau peningkatan teknik *pre-processing* agar model menjadi lebih seimbang.

Tahap terakhir dalam pembagian data adalah penggunaan data uji. Data ini digunakan sebagai pengujian final untuk mengetahui kemampuan sesungguhnya dari model BERT dalam memprediksi sentimen. Karena data uji sama sekali tidak diberikan pada proses pelatihan maupun validasi, performa model pada tahap ini menjadi indikator yang mencerminkan efektivitas model di kondisi nyata. Hasil evaluasi dari data uji menunjukkan seberapa baik model mampu memahami komentar baru dan memberikan prediksi sentimen yang sesuai dengan label aslinya. Pemisahan dataset menjadi data latih, validasi, dan uji merupakan bagian penting dalam penelitian berbasis *machine learning*. Strategi ini memastikan bahwa model tidak hanya unggul dalam proses pembelajaran, tetapi juga mampu berfungsi secara akurat ketika diterapkan pada analisis sentimen di dunia nyata.

3.4 Fine-tuning BERT

Pada tahap ini dilakukan proses *fine-tuning* terhadap model BERT menggunakan data komentar Instagram yang telah melalui proses *pre-processing* dan dibagi menjadi data latih, data validasi, dan data uji. Tahap ini bertujuan untuk menyesuaikan model dengan karakteristik bahasa pada data penelitian, sehingga model mampu mengenali pola penulisan, gaya bahasa informal, singkatan, serta ragam ekspresi yang muncul pada komentar terkait Program Makan Siang Gratis. Sebelum data dimasukkan ke dalam model, setiap komentar diproses menggunakan tokenizer BERT untuk diubah ke dalam bentuk token dan representasi numerik agar dapat dibaca oleh model. Selain itu, label sentimen yang semula berbentuk teks (positif, netral, negatif) dikonversi menjadi nilai numerik agar dapat digunakan dalam proses pelatihan.

Data yang telah ditokenisasi kemudian dibentuk menjadi dataset khusus untuk proses pelatihan, yang terdiri dari pasangan input token dan label sentimen. Dataset ini digunakan sebagai masukan bagi model selama proses *fine-tuning*. Pada tahap ini, model dilatih menggunakan data latih dengan tujuan agar model mampu mempelajari hubungan antara pola teks pada komentar dan kategori sentimen yang sesuai. Proses pelatihan dilakukan secara bertahap melalui sejumlah

iterasi, di mana parameter model diperbarui secara terus-menerus berdasarkan kesalahan prediksi yang dihasilkan pada setiap langkah pembelajaran.

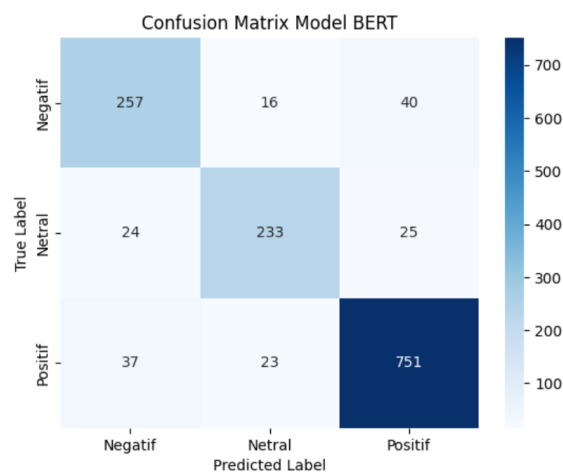
Selama proses *fine-tuning*, model tidak hanya mengandalkan pengetahuan umum yang dimilikinya pada tahap *pre-training*, tetapi disesuaikan kembali dengan konteks data penelitian yang lebih spesifik. Dengan demikian, model lebih menyesuaikan diri terhadap karakteristik bahasa pada media sosial, termasuk penggunaan kata tidak baku, variasi gaya penulisan, serta kecenderungan opini yang muncul pada komentar pengguna Instagram. Proses penyesuaian ini diharapkan dapat meningkatkan kemampuan model dalam membedakan kecenderungan sentimen positif, netral, dan negatif secara lebih akurat.

Pada implementasinya, proses *training* dijalankan menggunakan beberapa pengaturan *hyperparameter* yang telah ditentukan. Penelitian ini menggunakan ukuran *batch* sebesar 16, jumlah *epoch* sebanyak tiga kali pelatihan, nilai *learning rate* sebesar $2e-5$, serta *weight decay* sebesar 0,01 untuk menjaga keseimbangan proses pembelajaran. Pemilihan pengaturan tersebut bertujuan agar proses penyesuaian bobot model berlangsung secara bertahap, tidak terlalu cepat, dan tetap mampu menjaga kemampuan model dalam memprediksi data baru. Selama pelatihan berlangsung, data validasi digunakan untuk memantau perkembangan kinerja model pada setiap *epoch*. Selain itu, pada proses *fine-tuning* juga diterapkan fungsi evaluasi otomatis yang menghitung metrik akurasi, presisi, *recall*, dan *F1-score* berbasis *weighted average* pada data validasi. Penggunaan pendekatan *weighted* dilakukan karena distribusi jumlah data pada setiap kelas sentimen tidak seimbang, sehingga penilaian kinerja model tidak hanya didominasi oleh kelas dengan jumlah data terbesar. Melalui mekanisme ini, stabilitas performa model dapat diamati selama proses pelatihan berlangsung dan sekaligus membantu memastikan bahwa model tidak mengalami *overfitting* terhadap data latih.

Setelah proses *fine-tuning* dan evaluasi pada data validasi selesai, model yang telah terlatih kemudian digunakan untuk melakukan prediksi pada data uji. Data uji diproses melalui tokenizer dengan tahapan yang sama seperti pada data latih, sehingga representasi numeriknya konsisten. Model memetakan setiap komentar ke dalam salah satu kategori sentimen, yaitu positif, netral, atau negatif. Hasil prediksi ini kemudian digunakan sebagai dasar untuk mengevaluasi kinerja model dalam mengklasifikasikan sentimen komentar, yang akan dibahas lebih lanjut pada sub-bab evaluasi model.

3.5 Evaluasi Model

Setelah proses *fine-tuning* selesai, model yang telah terlatih digunakan untuk melakukan prediksi pada data uji. Tahap evaluasi model bertujuan untuk menilai sejauh mana model BERT mampu memprediksi sentimen komentar secara akurat berdasarkan label asli pada data uji, yang sebelumnya tidak pernah dilihat oleh model. Setelah model menghasilkan prediksi terhadap seluruh data uji, kualitas hasil klasifikasi kemudian dievaluasi menggunakan metrik presisi, *recall*, skor *f1*, akurasi dan *confusion matrix*. Masing-masing metrik tersebut memberikan perspektif berbeda dalam menilai bagaimana model membedakan setiap kelas sentimen, sehingga menghasilkan evaluasi yang lebih menyeluruh terhadap performa model. Sebelum menampilkan metrik evaluasi lainnya, terlebih dahulu ditampilkan *confusion matrix* sebagai alat untuk memahami distribusi hasil prediksi model, baik yang tepat maupun yang salah, pada setiap kategori terhadap setiap kelas sentimen. *Confusion matrix* memberikan gambaran detail mengenai seberapa sering model melakukan kesalahan klasifikasi, misalnya komentar negatif yang diprediksi sebagai netral atau sebaliknya. Informasi visual ini sangat penting dalam proses analisis karena membantu mengidentifikasi kelas mana yang masih sulit dikenali model. Hasil *confusion matrix* ditampilkan pada bagian berikut ini untuk menggambarkan kinerja model secara lebih komprehensif.



Gambar 8. Confusion Matrix

Pada Gambar 8 *confusion matrix* di atas memberikan rincian terperinci tentang klasifikasi model BERT terhadap tiga kelas sentimen yaitu *positive*, *neutral*, dan *negative*. Matriks ini menggambarkan distribusi prediksi yang benar dan salah pada data uji, untuk membantu pemahaman yang lebih mendalam tentang pola kesalahan model. Analisis ini tidak hanya mengevaluasi kinerja keseluruhan model berdasarkan akurasi atau skor *F1*, tetapi juga mengidentifikasi kelas mana yang paling sering salah diklasifikasikan dan hubungan antar kelas mana yang menyebabkan kesalahan tersebut.

Jumlah prediksi akurat tertinggi tercatat pada kelas positif, dengan total 751 data, ini membuktikan model mengidentifikasi sentimen positif dengan efektivitas tinggi. Komentar yang dipenuhi dengan dukungan, pujian, dan apresiasi dalam hal struktur kalimat lebih mudah dikenali oleh model BERT, sehingga kesalahan mengkelaskan relatif rendah. Selain itu, akurasi tinggi dalam kelas positif menunjukkan bahwa model secara efektif menangkap karakteristik umum komentar positif. Komentar-komentar ini mengekspresikan ulasan dengan cara yang jelas dan langsung. Meskipun demikian, masih ditemukan kesalahan klasifikasi pada model, di mana sebanyak 37 komentar yang sebenarnya bernada positif terdeteksi sebagai komentar negatif, dan 23 komentar lainnya salah dikenali sebagai komentar netral. Klasifikasi yang salah sebagai negatif dapat terjadi ketika komentar positif menggunakan kalimat perbandingan dengan nada kritis atau sarkastis. Klasifikasi yang salah antara positif dan netral biasanya terjadi pada komentar yang sopan, sederhana, atau komentar yang tidak secara langsung mengungkapkan perasaan positif yang kuat.

Untuk kelas netral, model ini berhasil melakukan prediksi akurat pada 233 data. Meskipun angka ini lebih rendah dibandingkan dengan kelas positif, hasil ini menunjukkan bahwa model dapat mengidentifikasi sentimen netral yang kurang jelas pengekspresiannya. Komentar netral umumnya mencakup informasi, pertanyaan, fakta, atau ide umum tanpa emosional yang kuat. Namun, ada 24 komentar netral yang diklasifikasikan sebagai negatif, yang menunjukkan bahwa model tersebut menganggap beberapa komentar netral sebagai tanda kritik atau opini yang cukup emosional. Selain itu, 25 komentar netral diklasifikasikan secara tidak tepat sebagai positif, meskipun jumlah ini setara dengan jumlah kesalahan dalam komentar negatif. Hal ini menunjukkan bahwa kelas netral sebenarnya sangat mirip dengan dua kelas lainnya, karena komentar netral dapat berisi kata-kata yang mengungkapkan persetujuan atau ketidaksetujuan, tetapi tidak terlalu kuat. Kelas netral sebenarnya adalah yang paling kompleks, karena batasannya sering kali tidak jelas.

Pada kelas negatif, model berhasil mengidentifikasi 257 data dengan tepat. Hal ini menandakan model mengenali teks yang bernada negatif, meskipun performanya tidak sekuat pada kelas positif. Sentimen negatif umumnya diwujudkan dalam bentuk kritik, keluhan, atau penolakan. Akan tetapi, tidak semua komentar yang bersifat negatif menuliskan emosi tersebut secara eksplisit. Masih terdapat 16 komentar negatif yang salah diprediksi sebagai netral, sehingga menunjukkan bahwa model terkadang menilai komentar yang sebenarnya mengandung ketidaksetujuan sebagai tidak memihak. Hal ini kemungkinan terjadi karena penggunaan bahasa yang kurang jelas menggambarkan emosi negatif sehingga kemiripannya lebih dekat dengan komentar netral. Selain itu, ditemukan pula 40 komentar negatif yang keliru diklasifikasikan sebagai positif. Kesalahan tersebut dapat muncul ketika komentar negatif disampaikan melalui sarkasme, humor bernada kritik, atau pujian yang bersifat ironis sehingga secara tekstual tampak positif, namun makna sesungguhnya adalah sebaliknya. Berdasarkan hasil *confusion matrix*, diperoleh nilai metrik presisi, *recall*, skor f1, dan akurasi yang ditampilkan pada tabel berikut.

Tabel 2. Hasil Evaluasi Model BERT

Kelas	Precision	Recall	F1-Score	Support
Negatif	0.81	0.82	0.81	313
Netral	0.86	0.83	0.84	282
Positif	0.92	0.93	0.92	811
Accuracy	-	-	0.88	1406
Macro Avg	0.86	0.86	0.86	1406
Weighted Avg	0.88	0.88	0.88	1406

Tabel 2 menyajikan hasil performa model dalam mengelompokkan sentimen komentar Instagram ke dalam, yaitu positive, netral, dan negative. Setiap kategori dievaluasi menggunakan metrik presisi, *recall*, dan skor f1 yang merupakan metrik utama untuk menilai kualitas model klasifikasi, serta dilengkapi informasi jumlah data uji pada masing-masing kelas. Selain itu, ditampilkan nilai *macro average* dan *weighted average* yang memberikan gambaran mengenai efektivitas model secara keseluruhan, baik tanpa mempertimbangkan proporsi data maupun berdasarkan distribusi data di setiap kelas.

Berdasarkan hasil yang diperoleh, kelas positif menunjukkan performa paling unggul dibanding dua kelas lainnya. Model menghasilkan nilai presisi sebesar 0,92, *recall* 0,93, serta skor f1 0,92 dari 811 data uji. Nilai *recall* yang tinggi memperlihatkan bahwa model mampu mengenali sebagian besar komentar positif secara tepat. Hal ini dapat terjadi karena jumlah data pelatihan pada kelas positif lebih besar sehingga pola bahasanya lebih mudah dipelajari dan diidentifikasi model. Tingginya nilai *presisi* juga menandakan bahwa kesalahan prediksi untuk kelas ini relatif rendah.

Pada kelas netral, model meraih presisi sebesar 0,86, *recall* 0,83, dan skor f1 0,84 dari 282 data. Hasil tersebut menunjukkan bahwa mayoritas komentar yang diprediksi sebagai netral memang sesuai dengan label aslinya, sehingga akurasi dalam mengenali sentimen netral tergolong baik. Meski demikian, nilai *recall* yang sedikit lebih rendah, yaitu 0,83, menandakan bahwa masih terdapat sejumlah komentar dengan sentimen netral yang belum berhasil diidentifikasi secara akurat oleh model. Hal ini dapat terjadi karena model menganggap komentar yang sebenarnya berisi pendapat netral sebagai negatif atau positif karena kata-kata yang ambigu atau sedikit emosional yang berada di antara netral dan sentimen lain. Namun, skor F1 sebesar 0,84 menunjukkan kinerja yang secara umum baik dan relatif stabil.

Selanjutnya, pada kelas negatif, model memperoleh nilai presisi sebesar 0,81, *recall* sebesar 0,82, serta skor f1 sebesar 0,81 dari 313 data uji. Nilai presisi tersebut memperlihatkan mayoritas komentar yang diprediksi sebagai negatif terbukti benar mengandung sentimen negatif. *Recall* sebesar 0,82 menunjukkan bahwa model mampu menemukan 82%

komentar negatif dari semua data. Skor F1 sebesar 0.81 yang sama dengan kedua metrik tadi menunjukkan performa yang selaras.

Pada gambar juga menunjukkan akurasi keseluruhan model sebesar 0,88. Hal ini berarti bahwa 88% dari data uji diprediksi dengan benar. Selain itu, nilai *macro average* sebesar 0,86 menunjukkan bahwa model mampu bekerja secara cukup seimbang pada semua kategori sentimen, karena perhitungannya tidak bergantung pada jumlah data tiap kelas. Di sisi lain, *weighted average* memperhitungkan distribusi data pada setiap kelas, sehingga kelas yang memiliki jumlah komentar lebih banyak berdampak besar terhadap hasil akhir performa model. Pada hasil evaluasi ini, *weighted average* mencapai 0,88, yang menandakan bahwa model mampu bekerja sangat baik berdasarkan distribusi data aktual di mana kelas positif merupakan kelas dengan data terbanyak sehingga lebih mendominasi perhitungan nilai rata-rata. Setelah analisis evaluasi ini, ditampilkan *word cloud* untuk menunjukkan kata-kata yang sering dibahas dalam komentar.



Gambar 9. Word Cloud Komentar MBG

Gambar 9 merupakan *word cloud* yang sering muncul pada komentar mengenai Program Makan Siang Gratis, yaitu kata “makan”, “gratis”, “anak”, “sekolah”, “orang”, dan “program” terlihat paling besar, menandakan bahwa topik pembahasan utama berfokus pada inti kebijakan tersebut. Kemunculan kata “gizi” dan “baik” menunjukkan adanya komentar yang menyoroti manfaat dan harapan publik terhadap peningkatan kualitas makanan bagi siswa. Di sisi lain, keberadaan kata seperti “racun”, “kurang”, atau “korban” menandakan sebagian komentar bernada negatif dan berisi kritik atau keraguan terhadap efektivitas pelaksanaannya. Selain itu, kata “presiden”, “pemerintah”, dan “prabowo” menunjukkan bahwa pembahasan program ini juga kuat dikaitkan dengan kebijakan dan tokoh politik. Secara keseluruhan, *word cloud* memberikan gambaran bahwa opini masyarakat beragam, dari dukungan positif hingga kritik yang menyoroti kendala implementasi program.

4. KESIMPULAN

Penelitian ini bertujuan untuk menganalisis sentimen pengguna Instagram di Sumatera Selatan mengenai program MBG dan untuk mengetahui seberapa banyak komentar positif, netral, dan negatif dengan menggunakan model BERT. Berdasarkan hasil dan pembahasan pada penelitian ini, model mampu memetakan sentimen ke dalam kelas negatif, netral, dan positif dengan tingkat akurasi yang cukup tinggi, yaitu 88%. Distribusi sentimen pada data uji memperlihatkan bahwa mayoritas komentar termasuk dalam kategori positif, sementara sisanya tersebar pada kategori netral dan negatif. Dari total 1.406 data uji, sebanyak 811 komentar diklasifikasikan sebagai positif, 282 komentar sebagai netral, dan 313 komentar sebagai negatif. Selain itu, analisis melalui presisi, *recall*, dan skor f1 menunjukkan bahwa performa model paling kuat terdapat pada kelas positif, sedangkan kelas negatif dan netral masih menunjukkan beberapa kesalahan prediksi, terutama karena kemiripan pola bahasa pada komentar yang bernada halus atau ambigu. Meskipun demikian, hasil penelitian ini secara keseluruhan berhasil menjawab tujuan utama, yaitu mengidentifikasi dan menggambarkan kecenderungan sentimen masyarakat terhadap program tersebut. Namun, penelitian ini memiliki keterbatasan, terutama pada distribusi data yang tidak seimbang serta adanya kesalahan klasifikasi antara kelas negatif dan netral. Keterbatasan ini membuka peluang bagi penelitian selanjutnya untuk menggunakan dataset yang lebih seimbang, menambah teknik preprocessing yang lebih kaya konteks, atau menerapkan model lanjutan agar performa klasifikasi dapat lebih optimal.

UCAPAN TERIMA KASIH

Penulis menyampaikan terima kasih kepada Jurusan Sistem Informasi, Fakultas Ilmu Komputer, dan Universitas Sriwijaya atas dukungan dan fasilitas yang telah diberikan dalam pelaksanaan penelitian ini.

REFERENCES

- [1] Y. Zaen Vebrian, "A Sentiment Analysis Of Free Meal Plans On Social Media Using Naïve Bayes Algorithms Analisis Sentimen Terhadap Rencana Makan Gratis Di Sosial Media X Menggunakan Algoritma Naive Bayes," *JURNAL INOVTEK POLBENG - SERI INFORMATIKA*, vol. 10, no. 1, 2025.
- [2] B. Rahmatullah, S. A. Saputra, P. Budiono, and D. P. Wigandi, "Sentimen Analisis Makan Bergizi Gratis Menggunakan Algoritma Naive Bayes," *JIFOTECH (JOURNAL OF INFORMATION TECHNOLOGY)*, vol. 05, no. 01, 2025.
- [3] A. Merlinda and Y. Yusuf, "Analisis Program Makan Gratis Prabowo Subianto Terhadap Strategi Peningkatan Motivasi Belajar Siswa di Sekolah Tinjauan dari Perspektif Sosiologi Pendidikan," vol. 7, 2025.
- [4] S. Fatimah, A. Rasyid, Anirwan, Qamal, and H. O. Arwakon, "Kebijakan Makan Bergizi Gratis di Indonesia Timur: Tantangan, Implementasi, dan Solusi untuk Ketahanan Pangan," vol. 4, no. 1, pp. 14–21, 2024.
- [5] R. A. Yahya, "Dana MBG dari Mana? Simak Penjelasan dan Sumbernya," *Tirto.id*.
- [6] A. Amelia and R. Yusuf, "Analisis Sentimen Masyarakat Indonesia Pada Platform X Terhadap Isu Fufufafa Menggunakan Bidirectional Encoder Representations From Transformers," 2025.
- [7] J. Ouyang, Z. Yang, S. Liang, B. Wang, Y. Wang, and X. Li, "Aspect-Based Sentiment Analysis with Explicit Sentiment Augmentations," 2024. [Online]. Available: www.aaai.org
- [8] A. Moreno-Ortiz, "Making Sense of Large Social Media Corpora Keywords, Topics, Sentiment, and Hashtags in the Coronavirus Twitter Corpus," 2024. doi: <https://doi.org/10.1007/978-3-031-52719-7>.
- [9] A. K. Ibrahim, M. Mustikasari, and I. Bastian, "Analisis Sentimen Pada Ulasan Aplikasi Astro - Groceries in Minutes di Google Play Store Menggunakan Metode Bidirectional Encoder Representation From Transformers (BERT)," *Jurnal Informatika dan Teknik Elektro Terapan*, vol. 13, no. 1, Jan. 2025, doi: 10.23960/jitet.v13i1.6033.
- [10] V. Chandradev, I. M. A. D. Suarjaya, and I. P. A. Bayupati, "Analisis Sentimen Review Hotel Menggunakan Metode Deep Learning BERT," Bali, 2023.
- [11] Z. Liu and Y. Lin Maosong Sun Editors, "Representation Learning for Natural Language Processing Second Edition," 2023. doi: <https://doi.org/10.1007/978-981-99-1600-9>.
- [12] M. Faruqziddan, E. Daniati, and M. N. Muzaki, "Pendekatan BERT Dalam Analisis Sentimen Terhadap Kominfo Di Media Sosial X," *IJCSR: The Indonesian Journal of Computer Science Research*, vol. 4, Jul. 2025.
- [13] A. Farhan, Istiadi, and A. Y. Rahman, "Analisis Sentimen Ulasan Aplikasi Identitas Kependudukan Digital di Google Play Store Dengan BERT," Jun. 2025.
- [14] R. M. R. W. P. K. Atmaja and W. Yustanti, "Analisis Sentimen Customer Review Aplikasi Ruang Guru dengan Metode BERT (Bidirectional Encoder Representations from Transformers)," *JEISBI*, vol. 02, p. 2021, 2021.
- [15] A. Aljabar and B. M. Karomah, "Mengungkap Opini Publik: Pendekatan BERT-based-caused untuk Analisis Sentimen pada Komentar Film," 2024.
- [16] M. F. Faturrian, J. H. Jaman, and I. Maulana, "Analisis Sentimen Terhadap Game Palworld di Steam Menggunakan Algoritma Bidirectional Encoder Representation From Transformers," 2025.
- [17] Y. P. Papani, Komaruddin, and M. A. Setiawan, "Penggunaan Media Sosial Sebagai Media Komunikasi Publik Pemerintah Provinsi Sumatera Selatan, *Jurnal Ilmiah Teknologi Informasi dan Komunikasi (JTIK)*, vol. 16, no. 1, pp. 90–99, 2025, [Online]. Available: <http://ejurnal.provisi.ac.id/index.php/JTIKP> page90
- [18] We Are Social dan Meltwater, "Digital 2025 : Indonesia," 2025.
- [19] K. Sikumbang et al., "Peranan Media Sosial Instagram terhadap Interaksi Sosial dan Etika pada Generasi Z," *Journal on Education*, vol. 06, no. 02, 2024.
- [20] J. B. Budiman and I. Deu, "Studi Penetrasi Aplikasi Media Sosial Instagram Sebagai Media Pemasaran Digital Pada UMKM," *Jurnal Riset Sistem Informasi Dan Teknik Informatika (JURASIK)*, vol. 10, pp. 274–281, 2025, [Online]. Available: <https://tunasbangsa.ac.id/ejurnal/index.php/jurasik>
- [21] M. Riswan, A. Primajaya, and A. S. Y. Irawan, "Analisis Sentimen Terhadap Pemberitaan Hasil Rekapitulasi Pemilu Presiden 2024 pada Media Sosial Instagram Menggunakan Naive Bayes Classifier," *Jurnal Informatika dan Teknik Elektro Terapan*, vol. 13, no. 1, Jan. 2025, doi: 10.23960/jitet.v13i1.5559.
- [22] M. Kozharinova and L. Manovich, "Instagram as a narrative platform," 2024.
- [23] Wikipedia, "Instagram."
- [24] J. U. S. Lazuardi and A. Juarna, "Analisis Sentimen Ulasan Pengguna Aplikasi Joox Pada Android Menggunakan Metode Bidirectional Encoder Representation From Transformer (BERT)," *Jurnal Ilmiah Informatika Komputer*, vol. 28, no. 3, pp. 251–260, 2023, doi: 10.35760/ik.2023.v28i3.10090.
- [25] E. Helmud, E. Helmud, Fitriyani, and P. Romadiana, "Classification Comparison Performance of Supervised Machine Learning Random Forest and Decision Tree Algorithms Using Confusion Matrix," *Jurnal Sisfokom (Sistem Informasi dan Komputer)*, vol. 13, no. 1, pp. 92–97, Feb. 2024, doi: 10.32736/sisfokom.v13i1.1985.