

IMPLEMENTASI INDOROBERTA UNTUK KLASIFIKASI SENTIMEN PADA MEDIA SOSIAL X TERHADAP PROGRAM MAKAN BERGIZI GRATIS

Salma Nur Fithriyah¹⁾, Alam Rahmatulloh²⁾

^{1,2}Prodi Informatika, Informatika, Fakultas Teknik, Universitas Siliwangi

Correspondence author: A.Rahmatulloh, alam@unsil.ac.id, Kota Tasikmalaya, Indonesia

Abstract

The MBG (Free Nutritious Meals) program remains a hot topic on social media, sparking a variety of comments, both positive and negative, regarding the policy. The diverse responses and comments on social media serve as a relevant data source and can be used as research objects. This study aims to classify sentiment in netizen posts and comments regarding the MBG program implemented in Indonesia. The classification model uses the IndoRoBERTa method, implemented within a sentiment analysis scheme, to classify sentiment in text. The process includes collecting social media comment text data, *preprocessing*, training the IndoRoBERTa model, and evaluating its performance. The results show that the developed sentiment classification model achieved an *accuracy* of 85.3% and an *F1-score* of 82.6%. Sentiment classification tends to produce negative sentiments from the overall text data.

Keywords: *classification, sentiment, social media x, free nutritious meals, IndoRoBERTa*

Abstrak

Program MBG (Makan Bergizi Gratis) masih menjadi salah satu topik yang masif di media sosial dan memicu berbagai komentar pro dan kontra terhadap kebijakan tersebut. Beragam respon dan komentar yang muncul di media sosial menjadi sebuah sumber data yang relevan dan bisa dijadikan sebagai sebuah objek penelitian. Penelitian ini, bertujuan untuk melakukan klasifikasi terhadap sentimen pada postingan dan komentar warganet terhadap program MBG yang diterapkan di Indonesia. Model Klasifikasi menggunakan metode IndoRoBERTa yang diimplementasikan dalam skema analisis sentimen untuk mengklasifikasi sentimen pada teks. Proses yang dilakukan mencakup pengumpulan data teks komentar di media sosial, tahap *preprocessing*, lalu implementasi model IndoRoBERTa dan mengukur hasil dari evaluasi kinerja model yang dibangun. Hasil penelitian menunjukkan nilai evaluasi model klasifikasi sentimen yang dikembangkan mencapai tingkat akurasi sebesar 85,3% dan nilai *F1-score* sebesar 82,6%. Klasifikasi sentimen cenderung menghasilkan sentimen berlabel negatif dari keseluruhan data teks.

Kata Kunci: *klasifikasi, sentimen, media sosial x, makan bergizi gratis, indoroberta*

A. PENDAHULUAN

Pola komunikasi masyarakat telah berubah secara fundamental yang disebabkan oleh revolusi digital yang terjadi dan menjadikan media sosial sebagai platform utama bagi semua individu untuk mengekspresikan opini, sikap dan emosi terhadap berbagai isu publik dan kebijakan pemerintah (Fitriani et al., 2025). Tingginya penetrasi internet dan pengguna aktif media sosial di Indonesia, seperti pada platform X menjadi data tekstual yang kaya dan *real-time* untuk mengetahui dan memantau reaksi yang diberikan masyarakat (Arifin & Al-Idrus, 2024). Analisis terhadap reaksi dan tanggapan dari masyarakat terutama dalam konteks penerapan program makan bergizi gratis (MBG) yang diterapkan di Indonesia dilakukan untuk menilai keberhasilan dari program MBG, mengidentifikasi kekurangan, dan merespons sentimen publik secara cepat dan akurat (Aprianti et al., 2025). Analisis sentimen umumnya melakukan klasifikasi teks ke dalam polaritas (positif, negatif, tunggal), namun ekspresi dan ungkapan emosi seringkali bersifat kompleks dan tidak eksklusif. Teks ulasan dalam media sosial bisa berupa kombinasi dari beberapa emosi dasar seperti senang dan terkejut (Aripin et al., 2023). Pada penelitian ini, berfokus pada *Multi-Label Classification* untuk mengatasi keterbatasan dalam klasifikasi tradisional. Penelitian ini memungkinkan sebuah teks dapat diberi label polaritas negatif, positif atau netral berdasarkan dengan penerapan *lexicon* dan model IndoRoBERTa.

Kinerja klasifikasi dalam ranah *natural language processing* (NLP), telah ditingkatkan secara signifikan dengan model *deep learning* berbasis *transformer* seperti BERT dan RoBERTa (Lazuardi et al., 2025). Pada bahasa Indonesia, secara khusus terdapat model yang telah dilatih yaitu IndoBERT dan IndoBERTa dengan memiliki keunggulan komparatif dalam memahami konteks dan nuansa bahasa

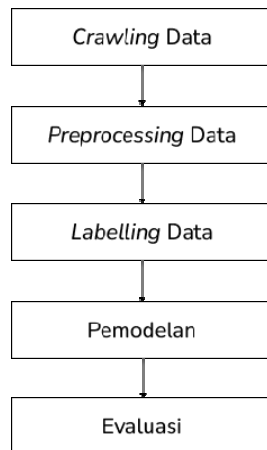
lokal (Lazuardi et al., 2025). Model IndoRoBERTa menjadi versi yang lebih optimal, dengan teknik *pre-training* yang disempurnakan terbukti memberikan hasil yang lebih tinggi dan lebih unggul dari IndoBERT dalam beberapa skenario tertentu (Apriansyah et al., 2025).

Bidang NLP telah mengalami perkembangan seperti penggunaan model berbasis *transformer* seperti BERT dan RoBERTa menjadi sebuah kebaruan dan standar baru dalam pemrosesan bahasa alami. Pada beberapa penelitian terkait penerapan IndoRoBERTa, seperti yang dilakukan oleh penelitian (Apriansyah et al., 2025), menggunakan penerapan *Fine-tuning* IndoRoBERTa untuk klasifikasi emosi pada teks mencapai akurasi sebesar 70% dan menghadapi kesulitan yaitu model sulit mengenali kelas yang memiliki ambiguitas yang tinggi seperti kelas netral dan sedih. Adapun penelitian yang dilakukan oleh (Apriansyah et al., 2025), melakukan perbandingan IndoBERT dan IndoRoBERTa untuk analisis sentimen komentar publik film dokumenter *Dirty Vote* menunjukkan IndoBERT memiliki akurasi yang lebih unggul yaitu 99% jika dibandingkan dengan IndoRoBERTa hanya menghasilkan akurasi sebesar 94%. Selain itu hasil dari penelitian yang dilakukan (Apriansyah et al., 2025), juga disebutkan bahwa mode IndoRoBERTa cenderung mengkategorikan komentar yang memiliki ambiguitas atau kritik tersirat sebagai sentimen netral, sementara IndoBERT menunjukkan prediksi yang lebih tegas.

Berdasarkan penelitian yang telah dilakukan oleh beberapa penelitian sebelumnya, pada penelitian ini mengusulkan untuk melakukan klasifikasi sentimen dengan menggunakan model dengan penerapan IndoRoBERTa dan objek yang diteliti yaitu terkait respon dan komentar yang diberikan *warganet* terhadap program makan bergizi gratis.

B. METODE PENELITIAN

Penelitian ini menggunakan *dataset* dari postingan dan komentar mengenai program makanan bergizi gratis di aplikasi X.



Gambar 1. Tahapan Penelitian

Adapun tahapan penelitian yang dilakukan meliputi, *crawling* data, *preprocessing* data, *labelling* data, pemodelan dan evaluasi.

1. *Crawling* Data

Tahap *crawling* data dilakukan menggunakan fungsi pada python untuk mengumpulkan postingan dan komentar dari aplikasi X mengenai program makan bergizi gratis (MBG) dengan menggunakan kata kunci mbg dengan rentang waktu dari Januari 2025 sampai Oktober 2025.

2. *Preprocessing* Data

Tahap *preprocessing* dilakukan untuk menyiapkan data mentah untuk dilakukan analisis yang akan digunakan dalam pemodelan. Tahap awal dari *preprocessing* yaitu melakukan *case folding* dengan mengubah penggunaan huruf kapital menjadi huruf kecil (Prasetyo et al., 2023). *Preprocessing* lainnya yang dilakukan diantaranya seperti melakukan normalisasi teks untuk menghapus penyebutan username pada komentar, spasi diawal/diakhir kalimat, menghapus tautan/ URL, menghapus karakter khusus, emotikon, angka dan melakukan normalisasi kata yang tidak baku menggunakan kamus normalisasi

kata. Tahap *preprocessing* setelah normalisasi, dilakukan tokenisasi untuk memecah teks menjadi kata yang terpisah. Tahap selanjutnya, yaitu melakukan penghapusan *stopword* yang terdapat pada teks seperti kata - kata seperti "dan", "atau", "tapi", dan lainnya (Muttakin et al., 2025), tujuannya untuk menghilangkan kata-kata yang kurang bermakna dan tidak mengandung sentimen tetapi muncul sering dalam teks (Hoiriyah et al., 2023).

3. *Labelling* Data

Tahap *labelling* data dilakukan dengan menggunakan *lexicon based*. Pelabelan menggunakan *lexicon based* merupakan suatu proses pemilihan kata penting pada dokumen berdasarkan suatu kamus/leksikon yang sudah ada (Sumitro et al., 2021). Penggunaan *lexicon InSet* telah terbukti efektif dalam menganalisis sentimen pada data berbahasa Indonesia (Muttakin et al., 2025). *Lexicon InSet* memiliki 2 kamus kata yaitu negatif dan positif. Kamus *lexicon Inset* positif terdiri dari 3609 kata dan kamus negatif terdiri dari 6.609 kata dengan setiap kata memiliki rentang bobot nilai atau polaritas daari -5 sampai +5 (Muttakin et al., 2025).

4. Pemodelan

Tahap pemodelan diawali dengan memuat data yang sudah melalui tahap *preprocessing* dan tahap *labelling*. Penerapan IndoRoBERTa dilakukan penambahan lapisan klasifikasi di atas lapisan *transformer*, yang disesuaikan dengan jumlah kelas target penelitian yaitu kelas positif, negatif dan netral. IndoRoBERTa merupakan pemrosesan bahasa alami yang dikembangkan berdasarkan arsitektur pada RoBERTa dan merupakan hasil dari optimalisasi BERT (Apriansyah et al., 2025). Model IndoRoBERTa dibangun dengan arsitektur *transformer* yang efektif digunakan dalam tugass pemrosesan bahasa alami seperti klasifikasi teks (Lazuardi et al., 2025). Proses pemodelan dengan IndoRoBERTa menggunakan 7 *epoch* sebagai parameter

utama. Pada penelitian ini menggunakan model IndoRoBERTa cahya/roberta-base-indonesian-522M dengan data yang telah melalui tahap *preprocessing* dan pelabelan.

5. Evaluasi

Evaluasi model dilakukan untuk mengukur kinerja model dalam melakukan klasifikasi sentimen. Evaluasi ini dilakukan menggunakan *confusion matrix*, yang memberikan representasi kinerja model melalui jumlah prediksi yang benar (*true netral*, *true positive* dan *true negative*) maupun prediksi yang keliru (*false netral*, *false positive* dan *false negative*) untuk masing-masing kelas sentimen.

Penilaian dilakukan dengan mengukur *accuracy*, *precision*, *recall*, dan *F1-score* berdasarkan *confusion matrix* dari model yang dihasilkan (Khairani & Zufria, 2025). Metrik - metrik ini berperan penting dalam menilai sejauh mana model mampu mengenali dan mengklasifikasikan pola sentimen secara akurat pada data uji.

C. HASIL DAN PEMBAHASAN

Berdasarkan hasil klasifikasi sentimen dari model IndoRoBERTa yang dilakukan dalam beberapa tahapan, berikut hasil dari tahapan metode yang telah dilakukan.

Crawling Data

Tahap *crawling* data dilakukan dengan menggunakan library yang tersedia di python dari aplikasi X menghasilkan data sebanyak 7891 data. Proses *crawling* data mengambil beberapa atribut seperti “*created_at*” untuk waktu tanggal dibuat komentar, “*full_text*” berisi komentar yang diposting, “*reply_count*” jumlah balasan yang didapatkan pada komentar yang diposting dan “*user_id_str*” menampilkan id akun yang digunakan user, seperti yang ditampilkan pada tabel 1.

Tabel 1. Data Hasil *Crawling*

id	created_at	full_text	user_id
0	Thu Jan 30 16:24:50 +0000 2025	Dokter tan pernah bilang kalau stunting itu dicegah bukan dikurangi . Ini berarti MBG sama sekali ga relevan klo dikaitkan sama isu stunting alias MBG itu bukti nyata penggunaan anggaran yg sgt tidak efisien.	18850012344 42260000
1	Thu Jan 30 16:26:48 +0000 2025	Kalo anggaran ya bagus sdh pasti diambil ternak ya mulyono Bak Bunuh Diri Apabila Memaksakan Diri jadi Mitra Program MBG https://t.co/FigO9BSz9R #inilahcom @inilahdotcom	18850017296 88880000
...	...	...	...
7889	Wed Oct 29 12:16:31 +0000 2025	@shiyiershier mana aku anak sekolahan yg tiap hari dapet mbg ☐	19835082419 05750000
7890	Wed Oct 29 12:03:21 +0000 2025	@ARSIPAJA Solusi jangka pendek: Hentikan modelan MBG sentralisasi lewat SPPG/BGN serahkan MBG ke kantin sekolah. Solusi jangka menengah dan panjang: Sederhanakan rantai distribusi bahan pangan. Jadi pejabat males mikir heran.	19835049302 04710000

Preprocessing Data

Tahap *preprocessing* data dilakukan dengan beberapa tahap yang dilakukan salah satunya menghapus data duplikat yang terdapat pada *dataset* berdasarkan kolom *full_text*. Data yang telah melalui tahapan penghapusan duplikat menjadi 6144 data.

Adapun tahap *preprocessing* yang dilakukan yaitu :

1. Normalisasi Kata

Normalisasi kata dilakukan untuk mengubah semua huruf pada kolom *full_text* menjadi huruf kecil (*Case folding*), menghapus spasi diawal dan diakhir kalimat, menghapus tautan/ URL, karakter khusus, emotikon, angka dan merubah kata tida baku/ tidak standar menjadi bentuk baku dengan menggunakan kamus normalisasi kata yang diambil dari repositori kaggle serta kata yang disesuaikan dari *dataset*. Hasil dari normalisasi kata ditampilkan pada kolom *normalisasi_text* pada tabel 2

Tabel 2. Data Hasil Normalisasi

id	full_text	normalisasi_text
0	Dokter tan pernah bilang kalau stunting itu dicegah bukan dikurangi . Ini berarti MBG sama sekali ga relevan klo dikaitkan sama isu stunting alias MBG itu bukti nyata penggunaan anggaran yg sgt tidak efisien.	dokter tan pernah bilang kalau stunting itu dicegah bukan dikurangi ini berarti mbg sama sekali ga relevan klo dikaitkan sama isu stunting alias mbg itu bukti nyata penggunaan anggaran yg sangat tidak efisien
1	Kalo anggaran ya bagus sdh pasti diambil ternak ya mulyono Bak Bunuh Diri Apabila Memaksakan Diri jadi Mitra Program MBG https://t.co/FigO9BSz9R #inilahcom @inilahdotcom	kalo anggaran ya bagus sudah pasti diambil ternak ya mulyono bak bunuh diri apabila memaksakan diri jadi mitra program mbg
...	...	...
7889	@shiyershier mana aku anak sekolahan yg tiap hari dapet mbg ☐	mana aku anak sekolahan yang tiap hari dapat mbg
7890	@ARSIPAJA Solusi jangka pendek: Hentikan modelan MBG sentralisasi lewat SPPG/BGN serahkan MBG ke kantin sekolah. Solusi jangka menengah dan panjang: sederhanakan rantai distribusi bahan pangan. Jadi pejabat males mikir heran.	solusi jangka pendek hentikan modelan mbg sentralisasi lewat sppgbgn serahkan mbg ke kantin sekolah solusi jangka menengah dan panjang sederhanakan rantai distribusi bahan pangan jadi pejabat males mikir heran

2. Tokenisasi

Tokenisasi dilakukan untuk memecah teks komentar menjadi kata yang terpisah yang berfungsi untuk mempermudah analisis pada saat *modelling* dilakukan. Pada tabel 3, menampilkan perbandingan kolom *full_text* yang berisi teks mentah, kolom normalisasi untuk teks yang sudah dilakukan *preprocessing*, dan kolom tokenisasi yang dilakukan berdasarkan pada kolom normalisasi.

Tabel 3. Data Hasil Tokenisasi

id	normalisasi_text	tokenisasi
0	dokter tan pernah bilang kalau stunting itu dicegah bukan dikurangi ini berarti mbg sama sekali ga relevan klo dikaitkan sama isu stunting alias mbg itu bukti nyata penggunaan anggaran yg sangat tidak efisien	['dokter', 'tan', 'pernah', 'bilang', 'kalau', 'stunting', 'itu', 'dicegah', 'bukan', 'dikurangi', 'ini', 'berarti', 'mbg', 'sama', 'sekali', 'ga', 'relevan', 'kalo', 'dikaitkan', 'sama', 'isu', 'stunting', 'alias', 'mbg', 'itu', 'bukti', 'nyata', 'penggunaan', 'anggaran', 'yang', 'sangat', 'tidak', 'efisien']
1	kalo anggaran ya bagus sudah pasti diambil ternak ya mulyono bak bunuh diri apabila memaksakan diri jadi mitra program mbg	['kalau', 'anggaran', 'iya', 'bagus', 'sdh', 'pasti', 'diambil', 'ternak', 'iya', 'mulyono', 'bak', 'bunuh', 'diri', 'apabila', 'memaksakan', 'diri', 'jadi', 'mitra', 'program', 'mbg']
...	...	...
7889	mana aku anak sekolahan yang tiap hari dapat mbg	['mana', 'aku', 'anak', 'sekolahan', 'yang', 'tiap', 'hari', 'dapat', 'mbg']
7890	solusi jangka pendek hentikan modelan mbg sentralisasi lewat sppgbgn serahkan mbg ke kantin sekolah solusi jangka menengah dan panjang sederhanakan rantai distribusi bahan pangan jadi pejabat males mikir heran	['solusi', 'jangka', 'pendek', 'hentikan', 'modelan', 'mbg', 'sentralisasi', 'lewat', 'sppgbgn', 'serahkan', 'mbg', 'ke', 'kantin', 'sekolah', 'solusi', 'jangka', 'menengah', 'panjang', 'sederhanakan', 'rantai', 'distribusi', 'bahan', 'pangan', 'jadi', 'pejabat', 'males', 'mikir', 'heran']

3. Penghapusan Stopword

Tahap penghapusan *stopword* dilakukan untuk menghapus kata - kata tidak penting dan tidak berpengaruh dalam proses

sentimen menggunakan kamus daftar *stopword*, data hasil *stopword* ditampilkan pada tabel 4 hasil tahapn *stopword_removal*.

Tabel 4. Data Hasil Penghapusan *Stopword*

id	full_text	tokenisasi	stopword_removal
0	Dokter tan pernah bilang kalau stunting itu dicegah bukan dikurangi . Ini berarti MBG sama sekali ga relevan klo dikaitkan sama isu stunting alias MBG itu bukti nyata penggunaan anggaran yg sgt tidak efisien.	['dokter', 'tan', 'pernah', 'bilang', 'kalau', 'stunting', 'itu', 'dicegah', 'bukan', 'dikurangi', 'ini', 'berarti', 'mbg', 'sama', 'sekali', 'ga', 'relevan', 'kalau', 'dikaitkan', 'sama', 'isu', 'stunting', 'alias', 'mbg', 'itu', 'bukti', 'nyata', 'penggunaan', 'anggaran', 'yang', 'sangat', 'tidak', 'efisien']	['dokter', 'tan', 'pernah', 'bilang', 'kalau', 'stunting', 'itu', 'dicegah', 'bukan', 'dikurangi', 'ini', 'berarti', 'mbg', 'sama', 'sekali', 'relevan', 'kalau', 'dikaitkan', 'sama', 'isu', 'stunting', 'alias', 'mbg', 'itu', 'bukti', 'nyata', 'yang', 'sangat', 'tidak', 'efisien']
1	Kalo anggaran ya bagus sdh pasti diambil ternak ya mulyono Bak Bunuh Diri Apabila Memaksakan Diri jadi Mitra Program MBG https://t.co/FigO9BSz9R #inilahcom @inilahdotcom	['kalau', 'anggaran', 'iya', 'bagus', 'sdh', 'pasti', 'diambil', 'ternak', 'iya', 'mulyono', 'bak', 'bunuh', 'diri', 'apabila', 'memaksakan', 'diri', 'jadi', 'mitra', 'program', 'mbg']	['kalau', 'anggaran', 'iya', 'bagus', 'sdh', 'pasti', 'diambil', 'ternak', 'iya', 'mulyono', 'bak', 'bunuh', 'diri', 'apabila', 'memaksakan', 'diri', 'jadi', 'mitra', 'program', 'mbg']
...	...	...	...
6142	@shiyiershi er mana aku anak sekolahan yg tiap hari dapat mbg □	['mana', 'aku', 'anak', 'sekolahan', 'yang', 'setiap', 'hari', 'dapat', 'mbg']	['mana', 'aku', 'anak', 'sekolahan', 'yang', 'setiap', 'hari', 'dapat', 'mbg']
6143	@ARSIPAJ A Solusi jangka pendek: Hentikan modelan	['solusi', 'jangka', 'pendek', 'hentikan', 'modelan', 'mbg']	['solusi', 'jangka', 'pendek', 'hentikan', 'modelan', 'lewat']

id	full_text	tokenisasi	stopword_removal
	MBG sentralisasi lewat SPPG/BGN serahkan MBG ke kantin sekolah. Solusi jangka menengah dan panjang: Sederhanakan an rantai distribusi bahan pangan. Jadi pejabat males mikir heran.	['sentralisasi', 'lewat', 'sppgbgn', 'serahkan', 'mbg', 'ke', 'kantin', 'sekolah', 'solusi', 'jangka', 'menengah', 'dan', 'panjang', 'sederhanakan', 'rantai', 'distribusi', 'bahan', 'pangan', 'jadi', 'pejabat', 'males', 'mikir', 'heran']	['sppgbgn', 'serahkan', 'mbg', 'ke', 'kantin', 'sekolah', 'solusi', 'jangka', 'menengah', 'dan', 'panjang', 'sederhanakan', 'rantai', 'distribusi', 'bahan', 'pangan', 'jadi', 'pejabat', 'males', 'mikir', 'heran']

Labelling Data

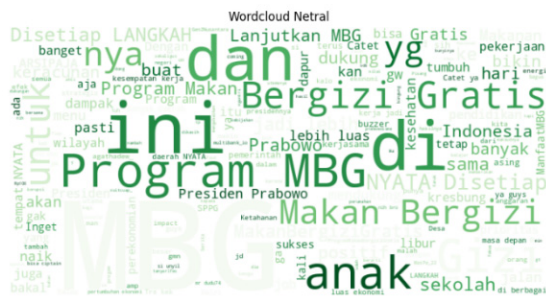
Tahap *labelling* dilakukan setelah tahap *preprocessing* data telah dilakukan. Pada penelitian ini menggunakan kamus *lexicon* yang berisi kata serta bobot untuk setiap kata tersebut. Kamus *lexicon* digunakan sebagai acuan untuk memberikan label pada setiap teks dalam data yang telah melalui proses *preprocessing*. Pelabelan data menghasilkan 209 label netral, 900 label positif dan 5035 label negatif seperti pada tabel 5.

Tabel 5. Data Hasil *Labelling*

id	normalisasi_text	label
0	dokter tan pernah bilang kalau stunting itu dicegah bukan dikurangi ini berarti mbg sama sekali ga relevan klo dikaitkan sama isu stunting alias mbg itu bukti nyata penggunaan anggaran yg sangat tidak efisien	negatif
1	kalo anggaran ya bagus sudah pasti diambil ternak ya mulyono bak bunuh diri apabila memaksakan diri jadi mitra program mbg	negatif
...	..	
6142	mana aku anak sekolahan yang tiap hari dapat mbg	positif
6143	solusi jangka pendek hentikan modelan mbg sentralisasi lewat sppgbgn serahkan mbg ke kantin sekolah solusi jangka menengah dan panjang sederhanakan rantai	negatif

id	normalisasi_text	label
	distribusi bahan pangan jadi pejabat males mikir heran	

Berdasarkan hasil dari label sentimen yang sudah diberikan kepada data yang dilabeli, berikut merupakan visualisasi *wordcloud* untuk setiap sentimen.



Gambar 2. *Wordcloud* sentimen netral



Gambar 3. *Wordcloud* sentimen positif



Gambar 4. *Wordcloud* sentimen negatif

Berdasarkan gambar dari hasil *wordcloud* pada gambar 2, 3 dan 4 kata yang kemunculannya lebih sering ditandai dengan warna dan ketebalannya lebih kontras, seperti kalimat mbg, yang, program dan kata anak.

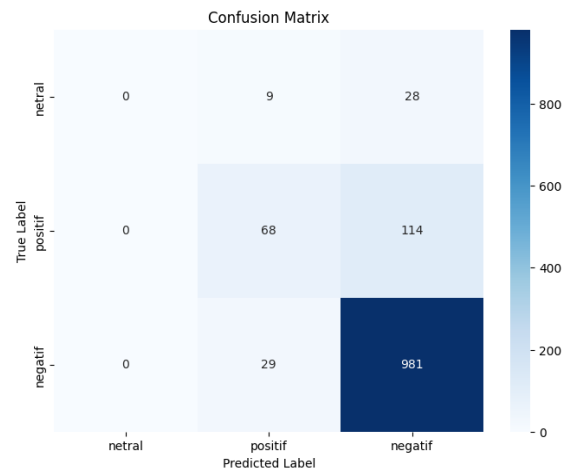
Pemodelan dengan IndoRoBERTa

Tahap pemodelan dengan menggunakan IndoRoBERTa, dilakukan menggunakan

data yang telah melewati tahap *preprocessing* dan tahap pelabelan. Pada pelatihan model, menggunakan data yang memiliki label dan dilakukan inisialisasi menggunakan nilai seperti, 0 untuk sentimen netral, 1 untuk sentimen positif dan 2 untuk sentimen negatif. Data yang diproses dilakukan *split* menjadi data *training* dan *data test* dengan rasio 80% untuk *training* dan 20% untuk *testing*. Pembagian data dilakukan untuk memastikan bahwa model dievaluasi menggunakan data yang benar benar belum pernah dilatih sebelumnya, sehingga penilaian kinerjanya dapat dilakukan secara objektif (Yarkhamsetiawan et al., 2025).

Evaluasi

Model IndoRoBERTa yang dibangun untuk melakukan klasifikasi sentimen pada postingan program makan bergizi gratis menghasilkan akurasi sebesar 85,3% dengan *F1-score* sebesar 82,6%.



Gambar 5. *Confusion matrix* Model IndoRoBERTa

Berdasarkan *confusion matrix* dari gambar 5, kinerja model IndoRoBERTa dievaluasi berdasarkan kemampuannya mengklasifikasikan tiga kategori sentimen yaitu netral, positif dan negatif dari 1229 data *testing*. Hasil pengujian menunjukkan pola kesalahan klasifikasi yang spesifik, terutama di sekitar label kelas minoritas.

Model menghasilkan kinerja yang baik pada kelas negatif, dengan mengenali 926 data dari 970 data *testing* sentimen negatif. Pada sentimen positif, model mengklasifikasikan 155 data dengan benar dari 208 data *testing* sentimen positif, serta pada kelas netral, model mengklasifikasikan 32 dengan benar dari keseluruhan data *testing* yaitu 51 data netral.

Secara keseluruhan, model cenderung menghasilkan kesalahan klasifikasi pada wilayah peralihan kelas seperti terjadi pada 53 kelas positif keliru dikenali sebagai kelas negatif dan pada kelas sentiment netral yang keliru diklasifikasikan pada sentiment negatif serta kasus kelas netral yang keliru diklasifikasikan sebagai kelas sentiment positif.

D. PENUTUP

Penelitian ini bertujuan untuk menghasilkan model yang mampu mengklasifikasi data postingan dan komentar menggunakan mode IndoRoBERTa. Melalui penerapan IndoRoBERTa mampu membangun model klasifikasi sentimen yang mampu mengidentifikasi opini publik secara efektif.

Hasil evaluasi dari model IndoRoBERTa untuk klasifikasi sentimen menunjukkan bahwa model yang dibangun mampu mencapai akurasi sebesar 85,3% dengan *F1-score* sebesar 82,6%, dengan nilai presisi, *recall*, dan *F1-score* rata-rata masing-masing sebesar 0,81. Kinerja model lebih optimal dalam mengklasifikasikan sentimen negatif dengan F2-score sebesar 0,85 dibandingkan sentimen positif yang memiliki *F1-score* sebesar 0,76. Berdasarkan hasil evaluasi dan label sentimen yang dihasilkan, perbedaan hasil dari setiap sentimen yang cukup jauh menimbulkan peluang untuk melakukan penambahan data teks yang lebih bervariasi agar hasil yang didapatkan untuk setiap sentimen tidak terlalu berbeda.

E. DAFTAR PUSTAKA

- Apriansyah, F. M., Ramadhan, T. I., Hidayat, C. R., & Wijaya, A. K. (2025). Perbandingan IndoBERT dan IndoRoBERTa Untuk Analisis Sentimen Pada Film Dokumenter Dirty Vote. *INFORMATIKA : Jurnal Pengembangan IT*, 10(3), 593–605. <https://doi.org/10.30591/jpit.v10i3.8607>
- Aprianti, N. N., Desmayani, N. M. M. R., Libraeni, L. G. B., Indrawan, I. G. A., & Radhitya, M. L. (2025). Public Sentiment Analysis of the Free Nutritious Meals Program (MBG) on Social Media X Using the Naive Bayes Method. *Journal of Applied Informatics and Computing*, 9(6), 3929–3936. <https://doi.org/10.30871/jaic.v9i6.11420>
- Arifin, K., & Al-Idrus, S. I. (2024). Klasifikasi Emosi Pengguna Twitter Terhadap Bakal Calon Presiden Pada Pemilu 2024 Menggunakan Algoritma Naive Bayes. *SAINTIKOM: Jurnal Sains Manajemen Informatika Dan Komputer*, 23(1), 37–45. <https://doi.org/10.53513/jis.v23i1.9558>
- Aripin, Santoso, S. A., & Haryanto, H. (2023). Optimizing the Accuracy of the Semantic-Based Compound Emotion Classifications using the XLM-RoBERTa. *JNTETI: Jurnal Nasional Teknik Elektro Dan Teknologi Informasi*, 12(1), 29–36. <https://doi.org/10.22146/jnteti.v12i1.6084>
- Fitriani, K. E., Faisal, M. R., Mazdadi, M. I., Indriani, F., Nugrahadi, D. T., & Prastya, S. E. (2025). Enhancing Natural Disaster Monitoring: A Deep Learning Approach to Social Media Analysis Using Indonesian BERT Variants. *IJEEEMI: Indonesian Journal of Electronics, Electromedical Engineering, and Medical Informatics*, 7(1), 77–89. <https://doi.org/10.35882/t158qq37>
- Hoiriyah, Qomariya, N., Darmawan, A. K.,

- Walid, M., & Efenie, Y. (2023). Sentiment Analysis on LGBT Issues in Indonesia With Lexicon-Based and Support Vector Machine Algorithms. *Pilar Nusa Mandiri: Journal of Computing and Information System*, 19(1), 27–36. <https://doi.org/10.33480/pilar.v19i1.4183>
- Khairani, M., & Zufria, I. (2025). Klasifikasi Tingkatan Perokok dengan Analisis Data Survei Masyarakat menggunakan Algoritma K-Means dan XGBoost. *JEPIN: Jurnal Edukasi & Penelitian Informatika*, 11(2), 230–241. <https://doi.org/10.26418/jp.v11i2.96678>
- Lazuardi, M. D. B., Fatyanosa, T. N., & Marji. (2025). Klasifikasi Emosi Multikelas Berbasis Teks Bahasa Indonesia Menggunakan IndoRoBERTa. *JPTIHK: Jurnal Pengembangan Teknologi Informasi Dan Ilmu Komputer*, 9(9), 1–9. <https://j-ptiik.ub.ac.id/index.php/j-ptiik/article/view/15276>
- Muttakin, F., Andrika, N., & Salsabila. (2025). Sentiment Analysis of Shoe Product Reviews on Indonesian E-Commerce Platform Using Lexicon Based and Support Vector Machine. *JUTIF: Jurnal Teknik Informatika*, 6(2), 839–854. <https://doi.org/10.52436/1.jutif.2025.6.2.3800>
- Prasetyo, S. D., Hilabi, S. S., & Nurapriani, F. (2023). Analisis Sentimen Relokasi Ibukota Nusantara Menggunakan Algoritma Naïve Bayes dan KNN. *Jurnal KomtekInfo*, 10(1), 1–7. <https://doi.org/10.35134/komtekinfo.v10i1.330>
- Sumitro, P. A., Rasiban, Mulyana, D. I., & Saputro, W. (2021). Analisis Sentimen Terhadap Vaksin Covid-19 di Indonesia pada Twitter Menggunakan Metode Lexicon Based. *J-ICOM: Jurnal Informatika Dan Teknologi Komputer*, 2(2), 50–56. <https://doi.org/10.55377/j-icom.v2i2.4009>
- Yarkhamsetiawan, Y., Akbar, M., & Satrianansyah. (2025). Analisis Sentimen Pengguna Aplikasi Qur'an Kemenag Menggunakan Metode Support Vector Machine (SVM). *JAMIKA: Jurnal Manajemen Informatika*, 15(2), 181–193. <https://doi.org/10.34010/jamika.v15i2.16843>