

BAB 1

PENDAHULUAN

1.1 Latar Belakang

Program Makan Bergizi Gratis (MBG) merupakan salah satu kebijakan unggulan yang diumumkan oleh presiden Prabowo Subianto dalam masa kampanye Pemilu 2024. Program ini mulai diimplementasikan secara bertahap sejak awal masa pemerintahan pada Oktober 2024 dengan tujuan utama menyediakan makanan bergizi gratis bagi siswa sekolah dasar hingga menengah, termasuk peserta didik di lingkungan pesantren. Secara umum, program ini dirancang untuk mengatasi permasalahan gizi serta meningkatkan kualitas sumber daya manusia di Indonesia, khususnya pada anak-anak dan ibu hamil. Dasar hukum utama pelaksanaan program ini diatur dalam Peraturan Presiden Nomor 83 Tahun 2024 tentang Badan Gizi Nasional serta Peraturan Presiden Nomor 115 Tahun 2025 yang mengatur tata kelola teknis program dalam rangka peningkatan gizi masyarakat [1].

Sejak mulai diterapkan, Program Makan Bergizi Gratis (MBG) menimbulkan berbagai tanggapan dari masyarakat. Sebagian masyarakat mendukung program tersebut karena dinilai mampu membantu pemenuhan gizi anak sekolah dan menekan angka stunting. Namun, sebagian lainnya memberikan kritik terkait kesiapan infrastruktur, pemerataan distribusi makanan, kualitas pelaksanaan program, serta efektivitas penggunaan anggaran negara. Perbedaan persepsi tersebut memunculkan diskusi yang luas di media sosial, khususnya pada platform X, yang menjadi salah satu media komunikasi digital paling aktif dalam menyampaikan opini publik terhadap kebijakan pemerintah. Karakteristik media

sosial *X* yang bersifat *real-time*, terbuka, dan cepat dalam menyebarkan informasi menjadikan platform ini sebagai sumber data yang relevan untuk mengamati respons masyarakat terhadap Program MBG [2].

Besarnya jumlah data opini pada media sosial menyebabkan proses analisis secara manual menjadi sulit dilakukan. Selain itu, data teks pada media sosial umumnya bersifat tidak terstruktur, menggunakan bahasa informal, singkatan, simbol, dan variasi bahasa yang kompleks sehingga memerlukan metode otomatis yang mampu memahami konteks bahasa secara lebih baik. Salah satu pendekatan yang dapat digunakan adalah analisis sentimen, yaitu teknik dalam *Natural Language Processing* (NLP) dan *text mining* yang bertujuan untuk mengidentifikasi serta mengklasifikasikan opini publik ke dalam kategori positif, negatif, dan netral. Dalam perkembangannya, metode berbasis *deep learning* menunjukkan performa yang lebih baik dibandingkan metode *machine learning* tradisional karena mampu memahami hubungan kontekstual antar kata secara lebih mendalam [3].

Salah satu model *deep learning* yang banyak digunakan dalam analisis sentimen adalah *RoBERTa* (*Robustly Optimized BERT Approach*). Model ini merupakan pengembangan dari *BERT* dengan optimasi pada proses *pre-training* sehingga mampu menghasilkan representasi bahasa yang lebih akurat. *RoBERTa* terbukti memiliki performa tinggi dalam berbagai tugas klasifikasi teks, termasuk analisis sentimen pada data media sosial yang bersifat dinamis dan tidak terstruktur. Kemampuan tersebut menjadikan *RoBERTa* relevan digunakan dalam menganalisis opini masyarakat terhadap kebijakan publik yang berkembang di media sosial [4].

Beberapa penelitian terdahulu menunjukkan bahwa analisis sentimen dapat digunakan untuk memahami persepsi masyarakat terhadap kebijakan publik. Penelitian oleh Saputro mengkaji klasifikasi emosi pada teks pidato politik menggunakan model *RoBERTa* dengan dataset sebanyak 2.952 paragraf. Hasil penelitian menunjukkan bahwa model *RoBERTa* mampu mencapai akurasi sebesar 90% dengan nilai *precision* 91%, *recall* 90%, dan *F1-score* sebesar 0,91. Temuan ini menunjukkan bahwa model berbasis *transformer* efektif dalam memahami konteks bahasa formal serta mampu mengklasifikasikan emosi secara akurat pada teks pidato politik [5]. Selain itu, penelitian oleh Alwi Jaya juga menerapkan model *RoBERTa* dalam analisis sentimen terhadap opini masyarakat mengenai profesi Pegawai Negeri Sipil (PNS) pada media sosial *Twitter* dengan jumlah data sebanyak 394 *tweet*. Hasil penelitian menunjukkan bahwa mayoritas sentimen yang dihasilkan bersifat netral sebesar 86,4%, dengan akurasi model mencapai sekitar 90%. Hal ini menunjukkan bahwa model *RoBERTa* memiliki kemampuan yang baik dalam mengolah data teks yang bersifat informal dan dinamis serta mampu menangkap makna dalam konteks media sosial [6].

Berdasarkan uraian diatas, penelitian ini dilakukan untuk menganalisis sentimen publik Indonesia terhadap Program Makan Bergizi Gratis (MBG) berdasarkan data media sosial *X* dengan menggunakan model *RoBERTa*. Penelitian ini diharapkan dapat memberikan gambaran yang lebih objektif mengenai persepsi masyarakat terhadap program tersebut, serta dapat menjadi bahan evaluasi bagi pemerintah dalam meningkatkan kualitas kebijakan publik di masa mendatang.

1.2 Rumusan Masalah

Berdasarkan latar belakang yang dikemukakan, maka dapat dirumuskan suatu masalah dalam penelitian ini yaitu, “Bagaimana penerapan model *RoBERTa* dalam menganalisis sentimen publik terhadap Program Makan Bergizi Gratis (MBG) berdasarkan data yang diperoleh dari media sosial X”.

1.3 Tujuan Penelitian

Tujuan dari penelitian ini adalah untuk menerapkan model *RoBERTa* dalam menganalisis sentimen publik terhadap Program Makan Bergizi Gratis (MBG) berdasarkan data yang diperoleh dari media sosial X.

1.4 Batasan Masalah

Batasan masalah dalam penelitian ini adalah sebagai berikut:

1. Data yang digunakan dalam penelitian ini berasal dari media sosial X sebanyak 5.964 *tweet* yang berkaitan dengan Program Makan Bergizi Gratis (MBG) dalam rentang waktu Januari sampai Juni 2026.
2. Analisis yang dilakukan terbatas pada klasifikasi sentimen ke dalam tiga kategori, yaitu positif, negatif, dan netral.
3. Metode yang digunakan dalam penelitian ini adalah model *deep learning RoBERTa* untuk proses klasifikasi sentimen.
4. Tahapan penelitian meliputi pengumpulan data, pelabelan data, pra-pemrosesan data, pelatihan model, serta evaluasi kinerja model.
5. Evaluasi model dilakukan menggunakan metrik seperti *accuracy*, *precision*, *recall*, dan *F1-score*.

1.5 Manfaat Penelitian

Penelitian ini diharapkan memberikan manfaat akademis dan praktis, yaitu sebagai referensi dalam pengembangan kajian analisis sentimen berbasis *Natural*

Language Processing (NLP) khususnya menggunakan model *RoBERTa*, serta memperkaya literatur ilmiah terkait analisis opini publik pada media sosial *X*. Secara praktis, penelitian ini dapat memberikan gambaran mengenai kecenderungan sentimen masyarakat terhadap Program Makan Bergizi Gratis (MBG), menjadi bahan pertimbangan bagi pemerintah atau pihak terkait dalam evaluasi kebijakan, serta mendukung pengembangan metode analisis sentimen yang lebih akurat dan efektif untuk penelitian selanjutnya.

1.6 Sistematika Penulisan

Sistematika penulisan dari skripsi ini terdiri dari pokok-pokok permasalahan yang akan dibahas pada masing-masing yang diuraikan menjadi beberapa bagian sebagai berikut :

BAB 1 PENDAHULUAN

Bab ini berisi latar belakang, rumusan masalah, tujuan penelitian, batasan masalah, manfaat penelitian, dan sistematika penulisan.

BAB 2 LANDASAN TEORI

Bab ini berisi teori-teori yang mendasari penelitian, meliputi konsep analisis sentimen, *text mining*, *Natural Language Processing* (NLP), media sosial *X*, serta model *deep learning RoBERTa*. Selain itu, dijelaskan juga konsep *preprocessing* teks, pembobotan kata, serta penelitian terdahulu yang relevan sebagai dasar dalam penelitian.

BAB 3 METODOLOGI PENELITIAN

Bab ini menjelaskan tahapan penelitian, mulai dari pengumpulan data dari media sosial *X*, proses *preprocessing* (*case folding*, *cleaning*, *tokenizing*), pelabelan data, pembagian data, *fine-tuning* model *RoBERTa*, serta

evaluasi model menggunakan metrik seperti *accuracy*, *precision*, *recall*, dan *F1-score*.

BAB 4 ANALISIS DAN PERANCANGAN SISTEM

Bab ini berisi analisis kebutuhan sistem dalam melakukan analisis sentimen, serta perancangan sistem yang meliputi alur proses pengolahan data, arsitektur model *RoBERTa*, dan perancangan alur klasifikasi sentimen dari data media sosial.

BAB 5 IMPLEMENTASI DAN PENGUJIAN

Bab ini menjelaskan implementasi sistem analisis sentimen menggunakan model *RoBERTa*, hasil pengujian terhadap dataset yang digunakan, serta analisis hasil klasifikasi sentimen berdasarkan metrik evaluasi yang telah ditentukan.

BAB 6 PENUTUP

Bab ini berisi kesimpulan dari penelitian yang telah dilakukan serta saran untuk pengembangan penelitian selanjutnya.

BAB 2

LANDASAN TEORI

2.1 *Artificial Intelligence (AI)*

Artificial Intelligence (AI) merupakan cabang ilmu komputer yang berfokus pada pengembangan sistem atau mesin yang mampu meniru kemampuan kognitif manusia, seperti belajar, bernalar, dan mengambil keputusan secara mandiri. Secara konseptual, *AI* tidak memiliki definisi tunggal yang baku, namun secara umum dipahami sebagai teknologi yang mengintegrasikan algoritma, data besar, serta kemampuan komputasi untuk menghasilkan sistem cerdas yang adaptif terhadap lingkungan. Karakteristik utama *AI* meliputi kemampuan pembelajaran (*learning*), penalaran (*reasoning*), dan pengambilan keputusan (*decision-making*), yang membedakannya dari sistem komputasi konvensional [7].

Perkembangan *AI* dalam beberapa tahun terakhir menunjukkan kemajuan signifikan, terutama dengan munculnya *generative artificial intelligence* yang mampu menghasilkan konten baru seperti teks, gambar, dan kode program secara otomatis. Teknologi ini memanfaatkan model pembelajaran mendalam (*deep learning*) dan data dalam jumlah besar untuk menciptakan *output* yang menyerupai hasil kreativitas manusia. Transformasi ini menandai pergeseran dari *AI* yang bersifat analitis menjadi *AI* yang juga bersifat kreatif, sehingga memperluas penerapan di berbagai bidang seperti pendidikan, industri, dan

layanan digital [8].

Selain menawarkan berbagai manfaat, implementasi *Artificial Intelligence (AI)* juga menghadirkan tantangan yang kompleks, terutama terkait transparansi, etika, dan akuntabilitas sistem. Banyak model *AI* modern, khususnya berbasis *deep learning*, bekerja sebagai “*black box*” yang sulit dipahami oleh manusia karena proses internalnya tidak dapat dijelaskan secara langsung. Kondisi ini berpotensi menimbulkan risiko, terutama ketika *AI* digunakan dalam sektor kritis seperti kesehatan, keuangan, dan sistem otonom, di mana kesalahan keputusan dapat berdampak signifikan terhadap kehidupan manusia [9].

2.2 Natural Language Processing (NLP)

Natural Language Processing (NLP) merupakan cabang dari kecerdasan buatan yang berfokus pada interaksi antara komputer dan bahasa manusia, baik dalam bentuk teks maupun suara. NLP bertujuan untuk memungkinkan sistem komputer memahami, menginterpretasikan, dan menghasilkan bahasa manusia secara otomatis. Dalam perkembangannya, NLP telah banyak digunakan dalam berbagai aplikasi seperti analisis sentimen, klasifikasi teks, *machine translation*, serta *chatbot*. Dengan meningkatnya volume data teks dari media digital, NLP menjadi sangat penting dalam proses ekstraksi informasi dan pengambilan keputusan berbasis data [10].

Dalam implementasinya, NLP melibatkan beberapa tahapan penting seperti *preprocessing* teks, ekstraksi fitur, serta pemodelan menggunakan algoritma *machine learning* atau *deep learning*. Tahapan *preprocessing* meliputi proses seperti *case folding*, *tokenizing*, *stopword removal*, dan *stemming* yang

bertujuan untuk membersihkan dan menormalisasi data teks. Setelah itu, teks direpresentasikan dalam bentuk numerik menggunakan metode seperti *TF-IDF* atau *word embedding*, sehingga dapat diproses oleh model pembelajaran mesin. Pendekatan ini memungkinkan sistem untuk mengenali pola bahasa dan mengklasifikasikan teks berdasarkan makna yang terkandung di dalamnya [11].

Seiring perkembangan teknologi, NLP mengalami kemajuan pesat dengan hadirnya model berbasis *deep learning* seperti *BERT* dan *transformer* lainnya yang mampu memahami konteks bahasa secara lebih mendalam. Model-model ini terbukti lebih unggul dibandingkan metode tradisional dalam berbagai tugas NLP, termasuk analisis sentimen dan klasifikasi teks. Selain itu, NLP juga telah diterapkan dalam berbagai sektor seperti pemerintahan, kesehatan, dan pendidikan untuk mendukung analisis data berbasis teks secara otomatis dan efisien [12].

2.3 *Transformer*

Transformer merupakan arsitektur *deep learning* yang diperkenalkan untuk meningkatkan kemampuan pemrosesan data berurutan pada bidang *Natural Language Processing* (NLP). Berbeda dengan model sebelumnya seperti *Recurrent Neural Network* (RNN) dan *Long Short-Term Memory* (LSTM) yang memproses data secara berurutan, *transformer* menggunakan mekanisme *self-attention* sehingga mampu memproses seluruh kata dalam kalimat secara paralel. Pendekatan ini memungkinkan model memahami hubungan antar kata dalam konteks kalimat secara lebih efektif dan efisien. Kehadiran *transformer* menjadi tonggak penting dalam perkembangan NLP modern karena mampu meningkatkan performa berbagai tugas pemrosesan bahasa seperti *machine translation*, *text*

classification, dan *sentiment analysis* [13].

Arsitektur *transformer* terdiri dari *encoder* dan *decoder* yang memanfaatkan mekanisme *multi-head self-attention* serta *positional encoding* untuk memahami urutan dan hubungan antar kata dalam teks. Mekanisme *self-attention* memungkinkan model memberikan fokus pada kata-kata tertentu yang memiliki keterkaitan makna dalam suatu kalimat, sehingga model dapat memahami konteks bahasa secara lebih mendalam. Selain itu, *transformer* juga mendukung proses komputasi paralel yang membuat pelatihan model menjadi lebih cepat dibandingkan metode berbasis *recurrent network*. Kemampuan tersebut menyebabkan *transformer* banyak digunakan sebagai dasar pengembangan model bahasa modern seperti *BERT*, *GPT*, dan *RoBERTa* [14].

Dalam perkembangannya, *transformer* telah menjadi arsitektur utama dalam berbagai penelitian *deep learning* berbasis bahasa karena mampu menghasilkan representasi teks yang lebih akurat dan kontekstual. Model berbasis *transformer* terbukti memiliki performa tinggi dalam memahami struktur bahasa alami, terutama pada data media sosial yang memiliki karakteristik tidak terstruktur dan dinamis [15].

2.4 RoBERTa

RoBERTa (*Robustly Optimized BERT Approach*) merupakan model pengembangan dari *BERT* yang dirancang untuk meningkatkan performa dalam tugas *Natural Language Processing* (NLP). Model ini melakukan optimasi pada proses *pretraining*, seperti penggunaan data yang lebih besar, penghapusan *Next Sentence Prediction*, serta peningkatan jumlah *batch* dan iterasi pelatihan.

RoBERTa terbukti mampu memberikan hasil yang lebih baik dibandingkan *BERT* pada berbagai tugas klasifikasi teks karena kemampuannya dalam memahami konteks bahasa secara lebih mendalam [16].

RoBERTa memanfaatkan arsitektur *transformer* dengan mekanisme *self-attention* yang memungkinkan model untuk menangkap keterkaitan antar kata dalam satu kalimat secara menyeluruh. Pendekatan ini membuat *RoBERTa* mampu mengolah data teks yang kompleks dan tidak terstruktur dengan lebih efektif. Dalam pemrosesan bahasa alami, kemampuan ini sangat penting karena bahasa manusia memiliki variasi struktur dan makna yang tinggi, terutama pada data dari media sosial yang bersifat dinamis [17].

Dalam penerapannya, *RoBERTa* banyak digunakan dalam berbagai tugas pemrosesan teks seperti klasifikasi, analisis sentimen, dan pemahaman bahasa alami. Model ini dikenal memiliki performa yang tinggi karena mampu memahami konteks secara lebih baik dibandingkan metode tradisional maupun model berbasis *machine learning*. Dengan kemampuan tersebut, *RoBERTa* menjadi salah satu metode yang relevan untuk digunakan dalam analisis sentimen pada data media sosial, termasuk *Twitter*[18].

2.5 Preprocessing Text

Preprocessing text merupakan tahapan awal dalam proses *Natural Language Processing* (NLP) yang bertujuan untuk membersihkan dan menyiapkan data teks mentah agar dapat diproses lebih lanjut oleh algoritma *machine learning*. Data teks yang diperoleh dari berbagai sumber umumnya bersifat tidak terstruktur, mengandung noise seperti kesalahan penulisan, simbol,

maupun informasi yang tidak relevan. Oleh karena itu, *preprocessing* dilakukan untuk mentransformasikan data menjadi format yang lebih terstruktur sehingga dapat meningkatkan kualitas input dan kinerja model analisis teks secara keseluruhan [19].

Secara umum, *preprocessing text* mencakup berbagai teknik seperti *case folding*, *tokenization*, penghapusan *stopwords*, dan *stemming*. Tahapan ini bertujuan untuk menyederhanakan teks tanpa menghilangkan makna utama, sehingga proses ekstraksi fitur menjadi lebih efektif. Penelitian menunjukkan bahwa pemilihan teknik *preprocessing* yang tepat dapat meningkatkan akurasi model *NLP*, seperti pada analisis sentimen dan klasifikasi teks, karena fitur yang dihasilkan menjadi lebih representatif terhadap data yang dianalisis [20].

Lebih lanjut, efektivitas *preprocessing text* sangat bergantung pada jenis data dan metode yang digunakan dalam pemodelan. Studi terbaru menunjukkan bahwa *preprocessing* tetap memiliki peran penting, bahkan dalam model modern seperti *transformer*, karena mampu memengaruhi performa secara signifikan jika diterapkan dengan tepat. Namun demikian, teknik yang digunakan perlu disesuaikan dengan konteks dataset dan tujuan analisis agar tidak menghilangkan informasi penting dalam teks [21].

2.5.1 Case Folding

Case folding merupakan salah satu tahapan awal dalam *text preprocessing* yang bertujuan untuk menyeragamkan bentuk huruf dalam teks dengan mengubah seluruh karakter menjadi huruf kecil (*lowercase*) atau bentuk tertentu yang konsisten. Proses ini dilakukan karena perbedaan penggunaan huruf besar dan

kecil tidak memberikan makna yang signifikan dalam banyak tugas *Natural Language Processing* (NLP), sehingga perlu dinormalisasi agar tidak dianggap sebagai kata yang berbeda. Dengan demikian, kata seperti “Data” dan “data” akan diperlakukan sebagai entitas yang sama, sehingga dapat mengurangi kompleksitas dan variasi kosakata dalam dataset [22].

Secara matematis, proses *case folding* dapat dinyatakan sebagai berikut:

$$T_{lower} = lower(T).....(1)$$

Keterangan:

T = Teks asli

T_{lower} = Teks setelah diubah menjadi huruf kecil

$lower()$ = Fungsi yang mengubah seluruh karakter menjadi huruf kecil

Selain itu, penerapan *case folding* juga berfungsi untuk meningkatkan efisiensi proses analisis teks karena mampu menyederhanakan struktur data serta mempermudah proses pencocokan kata (*matching*). Dalam praktiknya, tahap ini sering dikombinasikan dengan penghapusan karakter non-alfabet seperti angka, tanda baca, atau simbol yang tidak relevan [23].

2.5.2 Pembersihan data

Pembersihan data merupakan tahapan penting dalam *text preprocessing* yang bertujuan untuk membersihkan data teks dari berbagai elemen yang tidak relevan atau mengandung noise. Pada data yang berasal dari media sosial, seperti *X (Twitter)*, teks sering kali mengandung *URL*, simbol, angka, tanda baca, serta karakter khusus yang tidak memberikan kontribusi terhadap analisis sentimen

[18].

Secara konseptual, pembersihan data dapat dinyatakan sebagai proses transformasi dari teks mentah menjadi teks yang lebih bersih melalui fungsi tertentu, yang dirumuskan sebagai:

$$T_{clean} = F(T_{raw}) \dots \dots \dots (2)$$

T_{raw} = teks awal yang masih mengandung noises

T_{clean} = hasil teks yang telah dibersihkan melalui fungsi f

Transformasi ini bertujuan untuk mempertahankan informasi penting sekaligus mengurangi redundansi data, sehingga memudahkan model dalam mengekstraksi fitur yang relevan. Pendekatan ini umum digunakan dalam berbagai penelitian NLP modern untuk meningkatkan efisiensi dan akurasi model [24].

2.5.3 *Tokenizing*

Tokenizing merupakan proses memecah teks menjadi unit-unit kecil yang disebut token, seperti kata atau frasa, sehingga data dapat diolah oleh sistem *Natural Language Processing* (NLP). Tahapan ini penting karena teks mentah perlu diubah menjadi bentuk terstruktur agar dapat digunakan sebagai fitur dalam proses analisis seperti klasifikasi atau analisis sentimen [25].

Secara umum proses tokenisasi dapat dinyatakan sebagai:

$$T_{tokens} = T_{tokens}(T_{clean}) \dots \dots \dots (3)$$

Keterangan:

T_{tokens} = kumpulan token hasil pemecahan teks

T_{clean} = teks yang telah melalui proses *preprocessing*

$T_{okenize}()$ = fungsi pemecahan teks menjadi token

Dalam implementasinya, tokenizing dilakukan dengan memisahkan teks berdasarkan spasi, tanda baca, atau aturan tertentu. Ketepatan proses ini sangat memengaruhi kualitas hasil analisis, karena tokenisasi yang tidak tepat dapat menyebabkan hilangnya informasi penting dan menurunkan performa model [26].

2.6 *Confusion Matrix*

Confusion matrix merupakan salah satu metode evaluasi kinerja model klasifikasi dalam bidang *machine learning* yang digunakan untuk membandingkan hasil prediksi model dengan data aktual secara sistematis dalam bentuk tabel. Matriks ini umumnya berbentuk persegi yang terdiri dari empat komponen utama pada klasifikasi biner, yaitu *true positive* (TP), *true negative* (TN), *false positive* (FP), dan *false negative* (FN), yang masing-masing merepresentasikan jumlah prediksi benar dan salah dari model terhadap kelas tertentu. Melalui struktur tersebut, *confusion matrix* mampu memberikan gambaran rinci mengenai jenis kesalahan klasifikasi yang terjadi, sehingga tidak hanya menilai akurasi secara umum tetapi juga memberikan wawasan terhadap distribusi kesalahan antar kelas. Pendekatan ini banyak digunakan dalam berbagai bidang seperti data mining, sistem deteksi intrusi, hingga analisis medis karena kemampuannya dalam mengevaluasi performa model secara komprehensif [27].

Adapun metrik evaluasi yang sering digunakan dalam klasifikasi adalah sebagai berikut:

1. *Accuracy*

Accuracy digunakan untuk mengukur proporsi jumlah prediksi yang benar

terhadap seluruh data yang diuji.

$$Accuracy = \frac{TP+TN}{TP+TN+FP+FN} \dots\dots\dots(4)$$

2. Precision

Precision menunjukkan tingkat ketepatan model dalam memprediksi kelas positif.

$$Precision = \frac{TP}{TP+FP} \dots\dots\dots(5)$$

3. Recall

Recall atau sensitivitas mengukur kemampuan model dalam mengidentifikasi seluruh data yang termasuk dalam kelas positif.

$$Recall = \frac{TP}{TP+FN} \dots\dots\dots(6)$$

4. F1-Score

F1-Score merupakan rata-rata harmonis antara *precision* dan *recall* yang digunakan untuk menyeimbangkan kedua metrik tersebut.

$$F1 = \frac{2 \times (Precision \times Recall)}{Precision + Recall} \dots\dots\dots(7)$$

Metrik-metrik tersebut memberikan gambaran yang lebih lengkap mengenai performa model dibandingkan hanya menggunakan *accuracy* saja. Hal ini menjadi sangat penting terutama pada dataset yang tidak seimbang (*imbalanced data*), di mana distribusi kelas tidak merata. Dalam kondisi tersebut, metrik seperti *precision*, *recall*, dan *F1-score* lebih relevan digunakan untuk mengevaluasi kinerja model secara objektif [28].

2.7 Python

Python merupakan bahasa pemrograman tingkat tinggi yang bersifat *general-purpose* dan banyak digunakan dalam berbagai bidang seperti pengembangan perangkat lunak, analisis data, serta kecerdasan buatan. Bahasa ini dirancang dengan sintaks yang sederhana dan mudah dibaca sehingga mendukung produktivitas pengembang serta mempermudah proses pembelajaran. *Python* juga mendukung berbagai paradigma pemrograman, termasuk pemrograman berorientasi objek dan prosedural, serta memiliki pustaka standar yang luas untuk berbagai kebutuhan komputasi. Karakteristik ini menjadikan *Python* sebagai salah satu bahasa yang paling populer dalam pengembangan sistem modern [29].

Python juga banyak digunakan dalam bidang *data science* dan *machine learning* karena kemampuannya dalam mengelola data dan integrasinya dengan berbagai *library* seperti *NumPy*, *Pandas*, dan *TensorFlow*. Penelitian terbaru menunjukkan bahwa *Python* menjadi pilihan utama dalam pengembangan aplikasi berbasis kecerdasan buatan karena kemudahan sintaks, fleksibilitas, serta dukungan komunitas yang luas. Hal ini membuat *Python* efektif dalam membangun model analisis teks, termasuk analisis sentimen, karena mampu menangani data dalam skala besar secara efisien dan terstruktur [30].

2.8 Google Colaboratory

Google Colaboratory merupakan platform berbasis *cloud* yang disediakan oleh *Google* untuk mendukung pengembangan dan eksekusi kode pemrograman, khususnya *Python*, melalui lingkungan *Jupyter Notebook*. Platform ini memungkinkan pengguna untuk menjalankan komputasi tanpa perlu melakukan instalasi perangkat lunak secara lokal, sehingga sangat mendukung kegiatan

penelitian di bidang *data science* dan *machine learning*. Selain itu, *Google Colab* menyediakan akses ke GPU dan TPU secara gratis dalam batas tertentu, sehingga mempercepat proses pelatihan model berbasis *deep learning* [18].

Google Collaboratory banyak dimanfaatkan karena kemudahan akses, fleksibilitas, serta kemampuannya dalam mendukung kolaborasi daring. *Notebook* dapat disimpan dan dibagikan melalui *Google Drive* sehingga memudahkan replikasi eksperimen dan kerja tim [31].

2.9 Analisis Sentimen

Analisis sentimen merupakan teknik dalam *text mining* yang digunakan untuk mengidentifikasi dan mengklasifikasikan opini atau emosi yang terkandung dalam suatu teks. Tujuan utama dari analisis sentimen adalah untuk menentukan apakah suatu opini bersifat positif, negatif, atau netral terhadap suatu topik tertentu. Teknik ini banyak digunakan dalam analisis media sosial, ulasan produk, maupun opini publik terhadap suatu kebijakan atau layanan [32].

Dalam implementasinya, analisis sentimen biasanya memanfaatkan metode *machine learning* untuk melakukan klasifikasi terhadap teks. Algoritma seperti *Naive Bayes*, *Support Vector Machine*, dan *Logistic Regression* sering digunakan karena mampu menghasilkan performa yang cukup baik dalam mengklasifikasikan data teks. Model klasifikasi ini dilatih menggunakan dataset yang telah diberi label sehingga dapat mempelajari pola bahasa yang mencerminkan suatu sentimen tertentu [33].

Analisis sentimen memiliki peranan penting dalam memahami persepsi

masyarakat secara cepat dan efisien. Dengan memanfaatkan teknik ini, organisasi dapat memperoleh gambaran umum mengenai opini publik terhadap suatu isu tanpa harus melakukan analisis manual terhadap setiap data yang tersedia. Hal ini menjadikan analisis sentimen sebagai salah satu metode yang populer dalam penelitian berbasis data teks dan media sosial [34].

2.10 Program Makan Gratis

Program makan bergizi gratis merupakan salah satu kebijakan yang dirancang untuk meningkatkan kualitas gizi masyarakat, khususnya pada anak usia sekolah dan ibu hamil, melalui pemberian makanan dan susu secara cuma-cuma. Program ini diusulkan sebagai bagian dari agenda pembangunan nasional oleh presiden terpilih Prabowo Subianto bersama wakil presiden Gibran Rakabuming Raka untuk periode pemerintahan 2024–2029, dengan tujuan utama mengatasi permasalahan stunting, malnutrisi, serta meningkatkan kualitas sumber daya manusia Indonesia sejak usia dini [35].

Dalam perspektif akademik, program makan bergizi gratis tidak hanya berdampak pada aspek kesehatan, tetapi juga berkontribusi terhadap peningkatan kualitas pendidikan. Pemenuhan gizi yang optimal terbukti berpengaruh terhadap konsentrasi belajar, tingkat kehadiran siswa, serta kemampuan kognitif. Oleh karena itu, program ini dipandang sebagai intervensi strategis yang mengintegrasikan sektor kesehatan dan pendidikan guna meningkatkan hasil belajar peserta didik secara menyeluruh [36].

Meskipun demikian, program ini juga menjadi perhatian dalam berbagai kajian karena menimbulkan pro dan kontra di kalangan masyarakat dan

akademisi. Beberapa penelitian menunjukkan bahwa tantangan utama terletak pada aspek pembiayaan, kesiapan infrastruktur, serta mekanisme distribusi yang merata di seluruh wilayah. Dengan demikian, diperlukan perencanaan yang matang, pengawasan yang berkelanjutan, serta evaluasi berbasis data agar implementasi program dapat berjalan efektif dan memberikan dampak yang optimal [37].

2.11 Media Sosial X

Media sosial X (sebelumnya dikenal sebagai *Twitter*) merupakan salah satu platform komunikasi digital yang memungkinkan pengguna untuk berbagi informasi, opini, dan interaksi secara *real-time* melalui pesan singkat yang disebut *tweet*. Platform ini mengalami transformasi signifikan sejak diakuisisi dan berganti nama menjadi X pada tahun 2022, yang turut memengaruhi dinamika penggunaan dan persepsi pengguna secara global. X menjadi salah satu media sosial yang banyak digunakan untuk menyampaikan opini publik karena karakteristiknya yang terbuka, cepat, dan berbasis teks singkat [38].

Dalam konteks komunikasi digital, media sosial X memiliki peran penting dalam membentuk opini publik dan menyebarkan informasi, khususnya dalam isu sosial, politik, dan ekonomi. Berbagai penelitian menunjukkan bahwa X sering dimanfaatkan sebagai sarana diskusi publik, kampanye politik, serta penyebaran informasi yang dapat memengaruhi persepsi masyarakat. Karakteristiknya yang memungkinkan interaksi langsung dan penggunaan fitur seperti *hashtag* menjadikan platform ini efektif dalam membangun gerakan sosial serta memperkuat partisipasi publik di ruang digital [2].

Selain itu, *X* juga banyak digunakan sebagai sumber data dalam penelitian, khususnya dalam analisis sentimen dan kajian perilaku pengguna di media sosial. Data yang dihasilkan dari aktivitas pengguna, seperti *tweet*, *retweet*, dan komentar, dapat diolah menggunakan teknik *data mining* dan *natural language processing* untuk memahami opini masyarakat terhadap suatu isu tertentu. Meskipun demikian, penggunaan *X* juga memiliki tantangan, seperti potensi penyebaran informasi yang tidak akurat, polarisasi opini, serta munculnya misinformasi yang dapat memengaruhi kualitas diskursus publik [39].

2.12 Penelitian Terkait

Berikut adalah penelitian sebelumnya yang dijadikan acuan dan referensi dalam menyelesaikan masalah ini.

Tabel 2. 1 Penelitian Sebelumnya

NO	Nama	Tahun	Judul	Hasil
1	Dewa Gede Mahesta Parawangsa, dkk.	2026	Implementasi Metode <i>XLM-RoBERTa</i> Untuk Analisis Sentimen Ulasan Pengunjung Objek Wisata Kerta Gosa Klungkung	Penelitian melakukan analisis sentimen pada ulasan pengunjung objek wisata Kerta Gosa Klungkung menggunakan model <i>XLM-RoBERTa</i> . Setelah melalui tahap <i>preprocessing</i> , pelabelan data, dan pelatihan model multibahasa, hasil klasifikasi menunjukkan akurasi sebesar 91,37% dan <i>F1-score</i> sebesar 87,26%, sehingga metode ini cukup efektif dalam memahami opini publik pada data teks multibahasa meskipun masih terdapat

				keterbatasan pada kelas minoritas [40].
2	Akbar Rusmanto et al.	2026	Sentiment Analysis <i>Tweet Covid-19</i> Menggunakan <i>RoBERTa</i>	Penelitian melakukan analisis sentimen terhadap <i>tweet</i> terkait COVID-19 menggunakan model <i>RoBERTa</i> . Setelah <i>preprocessing</i> dan <i>fine-tuning</i> model, diperoleh akurasi sebesar 96,30% dan <i>F1-score</i> 0,96, menunjukkan bahwa metode ini sangat efektif dalam menganalisis teks media sosial berbahasa informal [41].
3	Cindy Setyowati et al.	2026	Analisis Sentimen dan Karakteristik Linguistik Komentar Publik terhadap Kebijakan Militer Menggunakan <i>RoBERTa</i>	Penelitian melakukan analisis sentimen dan karakteristik linguistik terhadap komentar publik terkait kebijakan militer. Setelah <i>preprocessing</i> dan pelatihan model <i>RoBERTa</i> , diperoleh akurasi sebesar 95% dengan <i>F1-score</i> rata-rata sekitar 0,94–0,95. Analisis menunjukkan komentar negatif cenderung emosional, netral bersifat informatif, dan positif menunjukkan dukungan. Model efektif dalam memahami nuansa bahasa dan opini publik [42].
4	Wahyu Adinda Nur Ashifa dkk	2025	Analisis Sentimen Tanggapan Masyarakat terhadap Layanan Kesehatan di Kota Surabaya dengan <i>RoBERTa</i>	Penelitian melakukan analisis sentimen terhadap tanggapan masyarakat mengenai layanan kesehatan di Kota Surabaya. Setelah <i>preprocessing</i> dan klasifikasi menggunakan <i>RoBERTa</i> , diperoleh akurasi sebesar 87,2%. Sentimen masyarakat

				didominasi netral pada masalah antrean (0,53), sentimen positif tertinggi pada fasilitas kesehatan (0,45), dan sentimen negatif tertinggi pada tenaga kesehatan (0,33). Studi ini memberikan insight penting untuk peningkatan layanan kesehatan [43].
5	Affa Fahmi Zain, dkk.	2025	Analisis Sentimen Ulasan Pengguna <i>iPhone</i> dengan Pendekatan Hibrida <i>RoBERTa</i> dan <i>XGBoost</i>	Penelitian melakukan analisis sentimen pada ulasan pengguna <i>iPhone</i> menggunakan pendekatan hibrida <i>RoBERTa</i> dan <i>XGBoost</i> . Setelah melalui tahap <i>preprocessing</i> , ekstraksi fitur, serta penyeimbangan data menggunakan SMOTE, hasil klasifikasi menunjukkan akurasi sebesar 99,74% dan <i>F1-score</i> mencapai 1,00, sehingga metode ini sangat efektif dalam meningkatkan performa klasifikasi sentimen terutama pada data yang tidak seimbang [44].
6	Syukriyansyah, Annisa Damayant	2025	Integrasi Model <i>RoBERTa</i> untuk Analisis Sentimen dan Deteksi Ujaran Kebencian pada Komentar YouTube: Pendekatan NLP dan Netiket	Penelitian melakukan analisis sentimen dan deteksi ujaran kebencian pada komentar YouTube terkait isu pertambangan di Papua menggunakan model Indonesian <i>RoBERTa</i> . Setelah melalui tahap <i>preprocessing</i> yang meliputi pembersihan teks, normalisasi, <i>stemming</i> , dan <i>stopword removal</i> . Hasil klasifikasi menunjukkan bahwa 57,7% komentar memiliki sentimen negatif,

				dan 52,4% dari komentar negatif tersebut mengandung ujaran kebencian. Temuan ini menunjukkan bahwa model RoBERTa efektif dalam mengidentifikasi sentimen dan ujaran kebencian serta mampu mengungkap degradasi kualitas diskusi publik di media social [45].
7	Elvanro Marthen Pusung, Ika Novita Dewi	2024	Optimasi RoBERTa dengan Hyperparameter Tuning untuk Deteksi Emosi Berbasis Teks	Penelitian melakukan deteksi emosi berbasis teks menggunakan model RoBERTa dengan optimasi hyperparameter tuning. Setelah melalui tahap preprocessing, fine-tuning, dan optimasi menggunakan Bayesian Optimization, hasil klasifikasi menunjukkan akurasi sebesar 83,64% dengan nilai precision, recall, dan F1-score yang seimbang, sehingga optimasi hyperparameter terbukti meningkatkan performa model secara signifikan [46].
8	Yogie Oktavianus Sihombing, Nita Vibriyanti Situmorang	2024	Prediksi Sentimen Pada Teks Media Sosial Corporate University Menggunakan RoBERTa	Penelitian melakukan analisis sentimen pada teks media sosial corporate university menggunakan model RoBERTa. Setelah melalui tahap preprocessing dan fine-tuning model, hasil klasifikasi menunjukkan akurasi sebesar 96,2% dan F1-score 95,2%, sehingga metode ini sangat efektif dalam memahami opini publik terhadap layanan pendidikan dan pelatihan

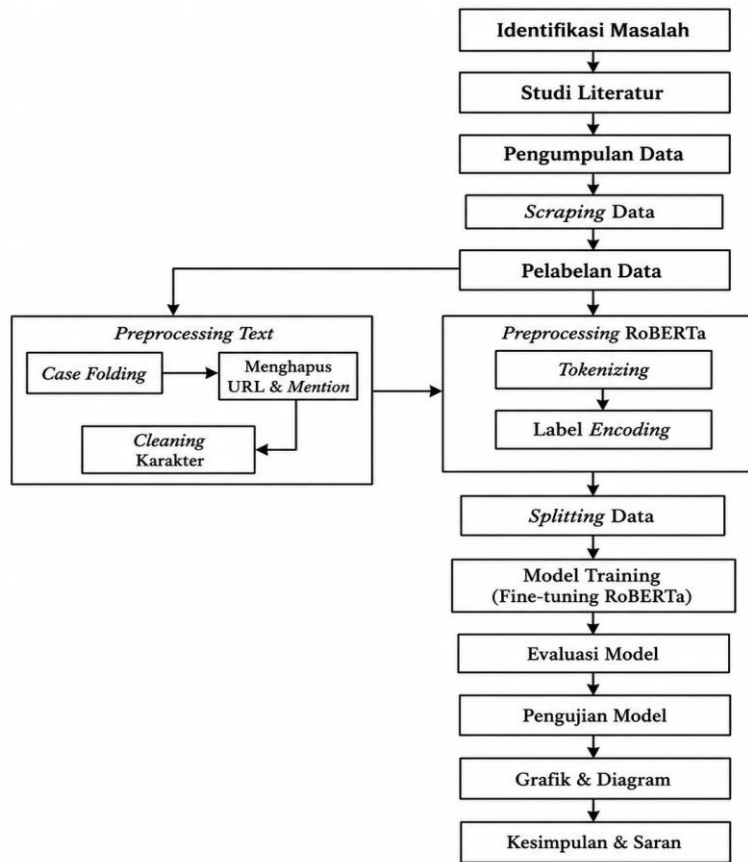
				[47].
9.	Bagas Setiadi et al.	2024	Optimisasi Klasifikasi Sentimen Pada <i>Review Hotel Bahasa Inggris Dengan Model RoBERTa Twitter</i>	Penelitian berhasil mengembangkan model analisis sentimen multi-aspek menggunakan <i>RoBERTa Twitter</i> dengan akurasi keseluruhan 88%. Akurasi per aspek yaitu fasilitas 75%, pelayanan 78%, kamar 81%, serta lokasi dan harga 84%. Model terbukti efektif dalam klasifikasi sentimen netral [48].
10	Alwi Jaya	2023	Analisis Sentimen Pandangan Publik terhadap Profesi PNS dari <i>Twitter</i> Menggunakan Indonesian <i>RoBERTa Base</i>	Penelitian ini berhasil mengklasifikasikan sentimen masyarakat terhadap profesi PNS berdasarkan 394 <i>tweet</i> menjadi kategori positif, negatif, dan netral. Hasilnya menunjukkan adanya variasi persepsi publik terhadap profesi PNS yang dapat diukur dalam bentuk persentase sentimen [6].

BAB 3

METODOLOGI PENELITIAN

3.1 Tahapan Penelitian

Penelitian ini dilakukan secara bertahap dengan pendekatan yang terstruktur dan sistematis untuk menganalisis sentimen publik terhadap Program Makan Bergizi Gratis (MBG) berdasarkan data dari media sosial *X* menggunakan model *RoBERTa*. Setiap tahapan saling berkaitan guna menghasilkan model klasifikasi yang optimal dan akurat. Alur metode penelitian ini disajikan dalam bentuk skema yang menggambarkan proses penelitian secara jelas dan teratur. Adapun tahapan lengkap dari penelitian ini dapat dilihat pada Gambar 3.1 berikut.



Gambar 3. 1 Tahapan Penelitian

3.1.1 Identifikasi Masalah

Dalam analisis data media sosial, salah satu tantangan utama adalah bagaimana mengolah data teks yang tidak terstruktur menjadi informasi yang bermakna. Proses ini dikenal dengan analisis sentimen, yang bertujuan untuk mengklasifikasikan opini publik ke dalam kategori tertentu seperti positif, negatif, dan netral. Data pada media sosial *X* yang memiliki karakteristik singkat, beragam, serta mengandung bahasa tidak baku sering kali sulit diproses menggunakan metode konvensional berbasis *machine learning*. Oleh karena itu, diperlukan pendekatan yang mampu memahami konteks bahasa secara lebih mendalam untuk menghasilkan klasifikasi yang lebih akurat. Dalam hal ini, model

RoBERTa menjadi salah satu metode yang potensial untuk diterapkan karena memiliki kemampuan dalam memahami konteks linguistik secara lebih baik melalui pendekatan *deep learning* berbasis *transformer*. Berdasarkan hal tersebut, masalah yang diidentifikasi dalam penelitian ini adalah bagaimana mengimplementasikan model *RoBERTa* untuk analisis sentimen publik terhadap Program Makan Bergizi Gratis (MBG) pada media sosial X secara efektif dan akurat.

3.1.2 Studi Literatur

Pada tahap ini, dilakukan pengumpulan dan kajian terhadap berbagai referensi yang relevan dengan topik penelitian. Referensi yang digunakan meliputi jurnal ilmiah, buku akademik, serta sumber terpercaya lainnya yang berkaitan dengan analisis sentimen, *Natural Language Processing* (NLP), dan *model deep learning* khususnya *RoBERTa*. Kegiatan ini bertujuan untuk memperoleh landasan teori yang kuat serta memahami perkembangan penelitian terkini dalam bidang yang diteliti. Melalui studi literatur, peneliti dapat mengidentifikasi metode-metode yang telah digunakan dalam penelitian sebelumnya, serta membandingkan kelebihan dan kekurangan dari masing-masing metode tersebut. Dalam penelitian ini, model *RoBERTa* dipilih karena memiliki kemampuan yang lebih baik dalam memahami konteks bahasa dibandingkan metode *machine learning konvensional*.

Hasil dari studi literatur digunakan sebagai dasar dalam penyusunan metodologi penelitian, mulai dari tahap *preprocessing* hingga evaluasi model. Dengan adanya studi literatur yang komprehensif, penelitian ini diharapkan dapat

memberikan kontribusi yang relevan serta memiliki dasar ilmiah yang kuat dalam analisis sentimen terhadap Program Makan Bergizi Gratis (MBG).

3.1.3 Pengumpulan Data

Pengumpulan data merupakan tahap awal yang penting dalam penelitian ini, karena berfungsi untuk memperoleh data utama yang akan dianalisis. Data yang digunakan bersumber dari media sosial X, berupa opini, komentar, dan tanggapan masyarakat terkait Program Makan Bergizi Gratis (MBG). Pengambilan data dilakukan berdasarkan kata kunci tertentu, seperti “MBG” dan “Makan Bergizi Gratis”, sehingga data yang diperoleh relevan dengan topik penelitian. Data yang dikumpulkan sebanyak 5.964 *tweet* berupa teks tweet atau cuitan pengguna yang mencerminkan opini publik terhadap program tersebut.

Selanjutnya, data yang diperoleh akan melalui tahap seleksi awal untuk menghapus data duplikat, spam, serta konten yang tidak relevan. Proses ini bertujuan untuk meningkatkan kualitas dataset agar hasil analisis menjadi lebih akurat dan representatif. Selain itu, data yang digunakan diupayakan mencakup berbagai jenis sentimen, yaitu positif, negatif, dan netral, sehingga mampu memberikan gambaran yang lebih objektif mengenai persepsi masyarakat terhadap Program Makan Bergizi Gratis (MBG).

3.1.4 Scraping Data

Pada tahapan ini dilakukan proses *scraping* data untuk mengambil data secara otomatis dari media sosial X guna memperoleh dataset yang relevan dengan tujuan penelitian. Teknik *scraping* dilakukan dengan bantuan *tools* atau *library* pemrograman yang mendukung proses ekstraksi data dari platform X.

Pengambilan data dilakukan berdasarkan kata kunci yang telah ditentukan sebelumnya, yaitu “MBG” dan “Makan Bergizi Gratis”, sehingga data yang diperoleh sesuai dengan topik penelitian. Selain itu, proses *scraping* dapat dibatasi berdasarkan rentang waktu tertentu agar data yang dikumpulkan mencerminkan kondisi terkini. Data yang diperoleh dari proses *scraping* meliputi teks *tweet* sebagai data utama analisis serta atribut pendukung seperti tanggal publikasi. Seluruh data kemudian disimpan dalam format terstruktur, seperti *CSV* atau *JSON*, untuk memudahkan proses pengolahan pada tahap selanjutnya.

3.1.5 Pelabelan Data

Tahap ini merupakan proses pemberian kategori terhadap data teks yang digunakan dalam penelitian. Dalam analisis sentimen, pelabelan bertujuan untuk mengelompokkan opini masyarakat berdasarkan polaritas sentimen yang terkandung dalam teks. Pada data diklasifikasikan ke dalam tiga kategori utama, yaitu:

1. Sentimen positif, yaitu opini yang mengandung dukungan atau tanggapan baik terhadap Program Makan Bergizi Gratis.
2. Sentimen negatif, yaitu opini yang berisi kritik atau tanggapan kurang baik terhadap program tersebut.
3. Sentimen netral, yaitu opini yang bersifat informatif, tidak memihak, atau tidak menunjukkan kecenderungan sentimen tertentu terhadap program MBG.

3.1.6 Pre-processing Data

Data teks mentah yang dikumpulkan biasanya belum siap untuk diproses

lebih lanjut. Tahap *preprocessing* bertujuan untuk membersihkan dan mempersiapkan data agar sesuai dengan kebutuhan algoritma klasifikasi yang digunakan, yaitu *RoBERTa*. Proses *preprocessing* dilakukan guna mengolah data mentah menjadi dataset yang lebih terstruktur dan siap digunakan dalam proses analisis sentimen. Selain itu, tahap ini juga bertujuan untuk mengurangi noise, menyederhanakan teks, serta memilih data yang relevan untuk diproses lebih lanjut dalam dokumen penelitian.

Tahapan *preprocessing* yang dilakukan meliputi beberapa langkah sebagai berikut:

1. *Case Folding*

Tahap ini bertujuan untuk mengubah seluruh huruf dalam teks menjadi huruf kecil (*lowercase*). Hal ini dilakukan untuk menghindari perbedaan penulisan kata yang sebenarnya memiliki arti yang sama. Sebagai contoh, kata “Gratis”, “gratis”, dan “GRATIS” akan dianggap sebagai satu kata yang sama.

2. Penghapusan *URL* dan Simbol

Tahap ini dilakukan dengan menghilangkan tautan (*URL*), tanda baca, serta karakter khusus yang tidak berkaitan langsung dengan isi teks, sehingga tidak mengganggu proses analisis.

3. Ekstraksi Teks dari *Hashtag*

Kata-kata yang terdapat dalam *hashtag* dipisahkan dan diolah menjadi bentuk teks biasa agar dapat dipahami sebagai bagian dari kalimat secara utuh.

4. Normalisasi Kata

Pada tahap ini, kata-kata tidak baku, singkatan, atau bahasa informal disesuaikan ke dalam bentuk yang lebih umum digunakan agar maknanya lebih jelas dan mudah dipahami.

5. *Tokenizing*

Teks diuraikan menjadi unit-unit kata (token) sehingga setiap bagian dalam kalimat dapat diidentifikasi dan dianalisis secara lebih terstruktur.

6. *Label Encoding*

Kategori sentimen diubah ke dalam bentuk numerik untuk mempermudah proses pengolahan data, yaitu sentimen positif diberi nilai 2, netral 1, dan negatif 0.

3.1.7 ***Splitting Data***

Tahap ini merupakan proses pembagian dataset yang telah melalui tahap *preprocessing* dan pelabelan ke dalam tiga bagian, yaitu data latih (*training data*), data validasi (*validation data*), dan data uji (*testing data*). Data latih digunakan untuk melatih model agar mampu memahami pola dan karakteristik data berdasarkan kategori sentimen yang telah ditentukan, yaitu positif, negatif, dan netral. Selanjutnya, data validasi digunakan selama proses pelatihan untuk memantau kinerja model serta membantu dalam penyesuaian parameter agar model dapat bekerja secara optimal. Sementara itu, data uji digunakan untuk mengukur kinerja akhir model dalam mengklasifikasikan sentimen pada data yang belum pernah dipelajari sebelumnya.

Pembagian dataset dalam penelitian ini menggunakan perbandingan 80%

untuk data *training*, 10% untuk data *validation*, dan 10% untuk data *testing*. Proporsi ini dipilih agar model memiliki jumlah data yang cukup untuk proses pembelajaran, sekaligus menyediakan data yang memadai untuk evaluasi performa secara objektif dalam menganalisis sentimen masyarakat terhadap Program Makan Bergizi Gratis (MBG).

3.1.8 Model Training

Pada tahap ini, model *RoBERTa* digunakan untuk mempelajari pola dari data yang telah dipersiapkan melalui proses *preprocessing* dan pelabelan. Data tersebut menjadi dasar bagi model dalam mengenali kecenderungan sentimen pada setiap pernyataan terkait Program Makan Bergizi Gratis (MBG). Melalui proses *fine-tuning*, model yang telah dilatih sebelumnya akan disesuaikan dengan karakteristik data penelitian sehingga mampu memahami makna kata serta hubungan antar kata *training* yang telah ditentukan sebelumnya. Selama proses ini, model mengoptimalkan parameternya agar mampu membedakan sentimen positif, netral, dan negatif secara lebih akurat. Dengan demikian, tahap ini diharapkan menghasilkan model yang mampu mengenali pola sentimen publik terhadap Program Makan Bergizi Gratis (MBG) secara lebih baik dan konsisten.

3.1.9 Evaluasi Model

Tahap ini dilakukan untuk mengetahui performa algoritma dalam mengklasifikasikan data. Evaluasi model menggunakan *confusion matrix* untuk membandingkan hasil prediksi dengan label sebenarnya.

Berdasarkan *confusion matrix*, digunakan beberapa metrik evaluasi, yaitu:

- A. *Accuracy*, untuk mengukur tingkat ketepatan model.

- B. *Precision*, untuk mengetahui ketepatan prediksi positif.
- C. *Recall*, untuk mengukur kemampuan mendeteksi data positif.
- D. *F1-Score*, sebagai keseimbangan antara *precision* dan *recall*.

3.1.10 Pengujian Model

Setelah melalui tahap evaluasi, model selanjutnya diuji menggunakan data pengujian (*testing data*) untuk menilai kemampuan akhir dalam mengklasifikasikan sentimen pada data baru yang belum pernah digunakan sebelumnya. Tahap ini bertujuan untuk memperoleh gambaran yang objektif mengenai performa model dalam kondisi yang mendekati penggunaan nyata.

Pengujian dilakukan untuk mengetahui sejauh mana model berbasis *RoBERTa* mampu mengidentifikasi sentimen positif, netral, dan negatif terhadap Program Makan Bergizi Gratis (MBG) berdasarkan data dari media sosial X. Hasil pengujian ini menjadi indikator penting dalam menilai kualitas dan keandalan model, sehingga dapat ditentukan apakah model sudah layak digunakan atau masih memerlukan penyempurnaan lebih lanjut.

3.1.11 Grafik dan Diagram

Hasil evaluasi dan pengujian model kemudian disajikan dalam bentuk visual untuk memudahkan pemahaman terhadap performa model. Penyajian ini bertujuan memberikan gambaran yang lebih jelas mengenai kemampuan model dalam mengklasifikasikan sentimen publik terhadap Program Makan Bergizi Gratis dalam kalimat secara kontekstual. Proses pelatihan dilakukan secara bertahap dengan memanfaatkan data (MBG) berdasarkan data dari media sosial X. Dengan adanya visualisasi, hasil analisis tidak hanya ditampilkan dalam bentuk

numerik, tetapi juga dapat dipahami secara lebih sederhana dan informatif.

Bentuk visualisasi yang digunakan dalam penelitian ini meliputi *confusion matrix* dan grafik performa model. *Confusion matrix* berfungsi untuk memperlihatkan perbandingan antara prediksi yang benar dan kesalahan klasifikasi pada setiap kategori sentimen, sedangkan grafik performa digunakan untuk menunjukkan tingkat akurasi model secara keseluruhan. Melalui penyajian ini, kinerja model berbasis *RoBERTa* dapat dijelaskan secara lebih sistematis, sehingga memudahkan dalam proses analisis dan penarikan kesimpulan.

3.1.12 Kesimpulan dan Saran

Tahap ini merupakan tahap akhir dalam penelitian yang bertujuan untuk merangkum dan menyimpulkan seluruh hasil yang telah diperoleh. Selain itu, dilakukan analisis terhadap kelebihan dan kekurangan metode yang digunakan dalam penelitian. Pada tahap ini juga diberikan saran sebagai bahan pertimbangan untuk pengembangan.