

ANALISIS SENTIMEN PUBLIK TERHADAP KABINET MERAH PUTIH PADA APLIKASI X DENGAN METODE BIDIRECTIONAL ENCODER REPRESENTATION FROM TRANSFORMER

Darwis Robinson Manalu^{1*}, Jusup Sopian Sihotang², Edward Rajagukguk³

^{1,2,3} Fakultas Ilmu Komputer, Universitas Methodist Indonesia

^{1*}manaludarwis@gmail.com, ²sihotangjusup@gmail.com, ³edw4rd.raja@gmail.com

DOI: <https://doi.org/10.46880/mtk.v12i2.5580>

ABSTRACT

The X application has become a platform for Indonesians to express their opinions, including those related to the Kabinet Merah Putih. This study aims to analyze public sentiment on the X application regarding the Kabinet Merah Putih using the IndoBERT method. A total of 1944 tweets were collected from the X platform using the keyword “Kabinet Merah Putih”, then processed through preprocessing stages including cleaning, normalization, case folding, and stopword removal, automatic labeling with TextBlob, and data division for model training and testing. The finetuned IndoBERT model was used to classify sentiment into positive, negative, and neutral. The results of the study show that public sentiment is dominated by positive and neutral sentiments, while IndoBERT's performance shows good accuracy in sentiment classification, with a model accuracy value of 85.96%, precision of 75%, recall of 74%, and F1-Score of 74%. This study provides an objective picture of public perception of the Kabinet Merah Putih on the X application and is expected to serve as evaluation material for the government.

Keywords: BERT, Natural Language Processing, Platform X, Sentiment Analysis

I. PENDAHULUAN

Perkembangan teknologi internet yang semakin pesat telah memberikan dampak yang sangat besar pada masyarakat dunia [1]. Sebelum internet ditemukan, masyarakat menyampaikan kritik dan pandangan mereka melalui media cetak. Namun, seiring dengan perkembangan teknologi, media sosial muncul dan menjadi sarana bagi masyarakat untuk menyampaikan aspirasinya. Media sosial telah mengubah dinamika komunikasi dan penyebaran informasi secara signifikan dalam beberapa dekade terakhir [2]. X menjadi salah satu dari begitu banyaknya media sosial yang berkembang dan menjadi aplikasi populer bagi masyarakat untuk memberikan opini, kritik, serta tanggapan mereka terhadap berbagai aspek politik di negeri ini. Pada tahun 2024, ada sebanyak 24,85 juta warga Indonesia yang menggunakan aplikasi ini [3]. Pengguna yang begitu banyak ini menyebabkan semakin banyak dan kompleks data yang dihasilkan, sehingga dibutuhkan metode untuk melakukan analisis sentimen sehingga hasil yang didapatkan optimal.

Kabinet Merah Putih menjadi salah satu topik perbincangan di aplikasi X. Para menteri yang dilantik dalam Kabinet Merah Putih, dengan pengalaman dan tanggung jawab mereka, tentu akan menghadapi berbagai opini publik terkait kinerja mereka [4]. Berbagai pendapat, mulai dari mengenai kebijakan serta kinerja para menteri dapat dilihat di berbagai media sosial, termasuk aplikasi X. Berbagai kebijakan dan keputusan yang diambil oleh Kabinet Merah Putih menjadi kontroversi dan menimbulkan kegaduhan ditengah masyarakat. Hal ini terbukti dengan viralnya tagar #IndonesiaGelap pada aplikasi X Februari 2025 yang lalu. Menurut data dari [5], tagar #IndonesiaGelap telah digunakan sebanyak 366.000 unggahan di aplikasi X. Bukan hanya protes melalui media sosial, berbagai aksi demo bertajuk “Indonesia Gelap” terjadi di beberapa kota di Indonesia. Hal ini menunjukkan pentingnya pemetaan opini publik secara objektif dan

akurat. Dengan melakukan analisis sentimen masyarakat terhadap keputusan dan kebijakan yang dikeluarkan, potensi kegaduhan dapat diminimalisir.

Untuk melakukan analisis terhadap sentimen dan persepsi publik, dibutuhkan metode yang dapat memahami makna dan konteks kata secara mendalam. Algoritma BERT (*Bidirectional Encoder Representations from Transformers*) dipilih karena keunggulannya dalam memahami konteks secara bidirectional, yang memungkinkan deteksi makna tersirat dalam kalimat. Metode BERT merupakan metode dari deep learning yang digunakan dalam pemrosesan bahasa alami (NLP) yang diperkenalkan oleh Google pada tahun 2018. Metode ini sangat efektif dalam menganalisis sentimen karena kemampuannya untuk mengerti makna kata secara mendalam dengan memperhatikan kata lain yang berada pada kalimat yang sama sehingga BERT mampu memahami makna dari sebuah kata tersebut sesuai konteks kalimatnya. BERT yang dilatih melalui pre training dan fine-tuning dapat disesuaikan dengan dataset khusus sehingga meningkatkan akurasi dalam konteks spesifik [6].

Penelitian mengenai Kabinet Merah Putih ini masih termasuk baru. Penelitian yang dilakukan oleh [4] menunjukkan bahwa sentimen negatif terhadap Kabinet Merah Putih mendominasi ruang media sosial pada masa awal pelantikannya. Studi tersebut berhasil menggambarkan persepsi awal publik, namun belum menyentuh dinamika evaluatif pasca-pelantikan, yaitu bagaimana masyarakat menilai kebijakan dan tindakan nyata yang telah diambil oleh kabinet.

Berdasarkan pemaparan di atas, penelitian ini menjadi penting dilakukan untuk mengetahui bagaimana sentimen publik setelah Kabinet Merah Putih mulai bekerja dan mengeluarkan berbagai kebijakan konkret yang berdampak langsung terhadap masyarakat.



II. METODE

Text Mining

Text Mining adalah ekstraksi informasi dari data sumber yang belum terstruktur yang mengacu pada teknik penambangan data untuk menganalisis dan memproses data [7]. Tujuan utama dari text mining adalah mengubah teks yang tidak terstruktur menjadi informasi yang terstruktur dan bermakna sehingga dapat dianalisis lebih lanjut. Penggunaan *text mining* dapat menyelesaikan masalah seperti analisis sentimen, *document clustering*, *document classification*, *information extraction*, *information retrieval*, dan *web mining* [1]. Oleh karena itu, metode text mining digunakan untuk membersihkan, mengolah, dan mengubah teks menjadi data yang dapat dianalisis lebih lanjut dengan algoritma seperti BERT.

Natural Language Processing

Salah satu bidang *Artificial Intelligence* yang mempelajari dan mengembangkan bagaimana komputer dapat memahami, dan memproses bahasa alami dalam bentuk teks atau ucapan dikenal sebagai *Natural Language Processing* [1]. Natural Language Processing adalah komponen kecerdasan buatan, atau AI dalam program komputer, dan memiliki kemampuan untuk memahami bahasa manusia seperti yang diucapkan [8]. Tujuan utama dari NLP adalah agar komputer dapat memahami, menafsirkan dan menghasilkan teks atau ucapan dalam bahasa manusia.

Analisis Sentimen

Analisis sentimen merupakan teknik yang digunakan untuk mengekstrak dan menilai opini serta perasaan yang terkandung dalam teks [9]. Analisis sentimen bertujuan untuk mengekstraksi atribut pada dokumen atau teks yang berisi komentar untuk mengetahui respon didalamnya agar dapat digolongkan menjadi respon positif, negatif dan netral [10]. Analisis sentimen ini dapat menentukan pendapat atau opini dari suatu teks atau dokumen kedalam kelas positif, netral, maupun negatif. [11]

Metode Bidirectional Encoder Representation From Transformer

Metode Bidirectional Encoder Representation from Transformers merupakan salah satu metode dari deep learning yang sering digunakan untuk pemrosesan bahasa alami (NLP). Model ini dibangun menggunakan 12 layer dan memberikan output berupa vektor dengan ukuran hidden size sebesar 768 [12]. Ada dua tahapan utama dalam BERT, yaitu tahapan *pre-training* dan tahap *fine tuning*. Proses *pre-training* BERT dilakukan dengan dua tugas utama dari BERT, yaitu :

1. Mask Language Modelling

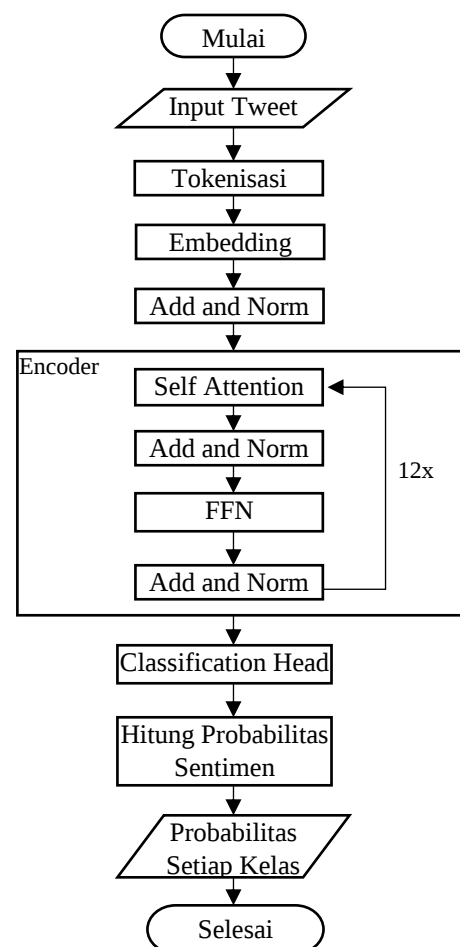
Pada tahap ini, model BERT akan diberikan input berupa kalimat, kemudian salah satu kata dari kalimat tersebut akan disembunyikan (*masking*) dan BERT harus memprediksi kata apa yang disembunyikan tersebut. Tujuan dari tahap ini adalah supaya model dapat memahami kata berdasarkan konteks dari kedua arah (sebelum dan sesudah kata), bukan hanya berdasarkan satu arah seperti pada model tradisional.

2. Next Sentence Prediction (NSP)

Pada tahap ini, model BERT akan diberikan input berupa kalimat, kemudian salah satu kata dari kalimat tersebut akan disembunyikan (*masking*) dan BERT harus

memprediksi kata apa yang disembunyikan tersebut. Tujuan dari tahap ini adalah supaya model dapat memahami kata berdasarkan konteks dari kedua arah (sebelum dan sesudah kata), bukan hanya berdasarkan satu arah seperti pada model tradisional.

Setelah melakukan *pre-training*, akan dilakukan proses penyesuaian ulang (*fine-tuning*). *Fine-tuning* adalah proses pelatihan lanjutan pada model *pre-trained* (yang telah dilatih sebelumnya) dengan data spesifik pada tugas tertentu, seperti klasifikasi sentimen. Dalam konteks ini, model yang telah dilatih secara umum akan disesuaikan ulang menggunakan dataset spesifik (dalam hal ini kumpulan tweet berbahasa Indonesia untuk tugas analisis sentimen). *Flowchart* dari metode BERT dapat dilihat pada gambar 1 berikut:



Gambar 1. Tahapan Metode BERT

Pada tahap awal, data berupa teks akan dilakukan tokenisasi, yaitu mengubah kalimat menjadi token (kata atau sub-kata). Ada beberapa token khusus yang ditambahkan kedalam BERT untuk menandakan hubungan antara satu kalimat dengan kalimat lain, diantaranya token [SEP] untuk menandakan akhir dari sebuah kalimat, token [CLS] digunakan untuk memberitahu model bahwa kita sedang melakukan klasifikasi dan penanda dari awal kalimat, token [PAD] untuk padding, token [UNK] untuk data yang tidak diketahui [13]. Setelah melewati tahap tokenisasi, setiap token akan dilakukan embedding. *Embedding* merupakan tahapan yang bertujuan mengubah token pada sebuah kalimat menjadi vektor numerik sebagai representasi dari token tersebut. Setiap vektor akan memiliki sebanyak 768 komponen. Sebelum masuk ke lapisan encoder, nilai dari

embedding akan dinormalisasi terlebih dahulu dengan formula sebagai berikut:

$$\hat{u}_i = \gamma_i \cdot \frac{u_i - \mu}{\sqrt{e + \sigma^2}} + \beta_i \quad (1)$$

dimana γ_i adalah bobot, u_i merupakan komponen ke- i dari vektor u , μ merupakan rata-rata dari seluruh komponen u untuk satu token, dan σ^2 merupakan varians dari seluruh komponen u untuk satu token. Selanjutnya, hasil dari normalisasi ini akan masuk ke dalam *encoder*. Ada dua komponen utama dari *encoder* ini, yaitu *self-attention* dan juga *feed forward network*. *Self-attention* adalah mekanisme yang memungkinkan setiap token dalam kalimat memperhatikan token lainnya, termasuk dirinya sendiri, untuk memahami konteks secara lebih luas. Mekanisme ini memungkinkan model untuk menangkap hubungan antara komponen-komponen yang berjauhan dalam sebuah kalimat dan menggunakannya sebagai pertimbangan dalam menghasilkan output akhir. Berikut rumus untuk *attention mechanism* ini:

$$\text{Attention}(Q, K, V) = \text{softmax}\left(\frac{\sum Q \cdot K^T}{\sqrt{d_k}}\right) V \quad (2)$$

dimana Q adalah vektor Query dari token yang sedang dianalisa, K merupakan Key, yaitu vektor dari semua token pada sebuah kalimat, dan V merupakan Value, yaitu vektor yang digunakan untuk menghasilkan output, sementara K^T adalah transpose dari Key, digunakan untuk mencocokkan (*dot product*) dengan Query. Nilai dari Q , K , dan V ini didapatkan dengan menggunakan formula sebagai berikut:

$$Q, K, V = E_i \cdot W_{QKV} + b_{QKV} \quad (3)$$

dimana E_i adalah vektor embedding sepanjang 768, W_{QKV} adalah matrik bobot untuk masing-masing komponen berukuran 768x768, serta b_{QKV} adalah bias untuk masing-masing komponen sepanjang 768. Setelah itu, hasil dari mekanisme ini akan dijumlahkan dengan input sebelumnya (*add/residual*) dengan rumus:

$$u = x + \text{sublayer}(x) \quad (4)$$

dimana x input awal pada *encoder* BERT, dan $\text{sublayer}(x)$ adalah *output* dari *self-attention*. Hasil penjumlahan ini akan dinormalisasi menggunakan formula yang sama dengan normalisasi yang dilakukan setelah tahap *embedding*.

Setelah dinormalisasi, hasil dari normalisasi tersebut akan masuk ke dalam tahapan *feed forward network*, dimana tahapan ini berfungsi untuk mengubah representasi token secara non-linear, meningkatkan kemampuan jaringan dalam memodelkan hubungan kompleks. Berikut adalah formula yang digunakan dalam FFN:

$$\text{FFN}(x) = (\text{GELU}(xW_1 + b_1))W_2 + b_2 \quad (5)$$

dimana x adalah vektor input dari token, W_1 merupakan matriks bobot layer pertama berukuran 768x3072, b_1 adalah bias layer pertama, $\text{GELU}()$ adalah *Gaussian Error Linear Unit*, merupakan fungsi

aktivasi yang digunakan pada FFN BERT, W_2 merupakan matriks bobot untuk layer kedua, dan b_2 merupakan bias layer kedua. Output dari tahap ini akan dijumlahkan dengan input sebelumnya, dan akan dinormalisasi dengan formula yang sama dengan tahap sebelumnya.

Selanjutnya akan ditambahkan *classification-head* berupa *fully connected layer (dense layer)* di atas model BERT. Lapisan dense ini dapat digunakan untuk mengubah representasi teks menjadi prediksi kelas yang spesifik [12]. Berikut adalah rumus yang digunakan dalam *fully connected layer* ini :

$$y = [\text{CLS}] \cdot w + b \quad (6)$$

dimana $[\text{CLS}]$ adalah vektor representasi dari hasil token $[\text{CLS}]$, w merupakan bobot pada *fully connected layer* serta b adalah bias pada *fully connected layer*. Setelah hasil y didapatkan, akan dilanjutkan dengan fungsi aktivasi softmax untuk menghitung nilai probabilitas dari setiap kelas sebagai berikut :

$$P(y_i) = \frac{e^{y_i}}{\sum_{j=1}^K e^{y_j}} \quad (7)$$

dimana y_i adalah hasil dari *fully connected layer* C merupakan jumlah total kelas. (dalam penelitian ini ada 3 kelas)

Penelitian ini akan menggunakan *Cross Entropy Loss* yang umum digunakan untuk menghitung nilai *loss* dalam klasifikasi multi-kelas. Adapun rumus dari *Cross Entropy Loss* ini adalah sebagai berikut:

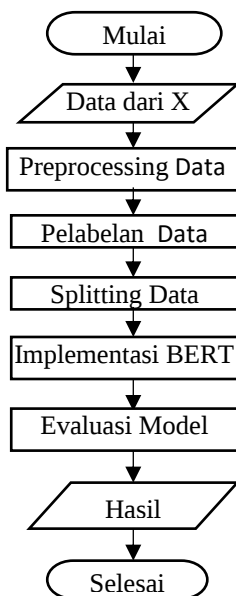
$$Y = - \sum_{i=1}^C y_i \log P(\hat{y}_i) \quad (8)$$

Dimana C merupakan jumlah kelas (3 kelas), y_i merupakan nilai sebenarnya (*ground truth*), $\hat{y}_i$ adalah probabilitas prediksi dari model.

Model dilatih dengan *optimizer* AdamW. *Learning rate* yang digunakan adalah 1e-5 dan *batch-size* 16 yang umumnya digunakan dalam *fine-tuning* BERT untuk analisis sentimen. Nilai ini telah terbukti memberikan hasil yang baik dalam menjaga keseimbangan antara konvergensi model dan mencegahnya dari *overshooting* atau *divergensi*. [9]

Framework Penelitian

Framework penelitian adalah struktur konseptual yang digunakan untuk mengarahkan peneliti dalam melaksanakan dan mengembangkan penelitian. Framework ini membantu memastikan bahwa semua aspek penelitian terstruktur dengan baik dan terhubung satu sama lain secara logis. Framework penelitian ini dapat dilihat pada gambar 2 berikut:

**Gambar 2.** Framework Penelitian

Penelitian ini akan menggunakan algoritma *indoBERT*, yang merupakan varian dari algoritma BERT. Alasan penggunaan *indoBERT* adalah karena *indoBERT* merupakan model yang dilatih khusus menggunakan Bahasa Indonesia sehingga *indoBERT* sangat cocok digunakan untuk melakukan tugas *Natural Language Processing* dalam Bahasa Indonesia.[14]

III. HASIL DAN PEMBAHASAN

Pengumpulan Data

Proses pengumpulan data dilakukan dengan menggunakan bahasa pemrograman python. Dari proses pengumpulan data ini, diperoleh data mentah sebanyak 1944, dimana ada sebanyak 921 tweet positif, 914 netral, serta 109 tweet negatif yang kemudian akan dilakukan preprocessing untuk memastikan kualitas data sebelum diproses. Contoh hasil proses pengumpulan data ini dapat dilihat pada tabel 1 berikut ini:

Tabel 1. Hasil Pengumpulan Data

No	Tweet
1	Ini maksudnya apa? Ko bisa cakar2an dalam kabinet? Bagaimana mau bikin indonesia cerah kalo caranya begini.
2	Dengan Kabinet Merah Putih yang semakin solid semangat optimisme untuk Indonesia Maju kini semakin nyata di depan mata.
...	...
1944	@dinopattidjalal @prabowo @Menlu_RI Sependapat focus benahi kabinet merah putih ganti menteri2 termul sy yakin dukungan rakyat makin besar dan tdk kecewa !

Preprocessing Data

Ada beberapa jenis *preprocessing* yang dilakukan untuk penelitian ini seperti *cleaning* untuk menghapus elemen elemen yang tidak terlalu relevan untuk analisis sentimen, *case-folding* untuk mengubah huruf kapital menjadi huruf kecil, normalisasi untuk mengubah kata gaul/informal menjadi kata baku, serta penghapusan *stopwords*, yaitu kata-kata yang tidak terlalu memiliki pengaruh untuk klasifikasi sentimen, seperti yang, dan, serta dll.

1. Cleaning

Proses ini bertujuan membersihkan data dari elemen elemen yang tidak relevan, seperti mention, URL, emoji, dan karakter khusus. Hasil *cleaning* akan ditunjukkan pada tabel 2 berikut:

Tabel 2. Hasil Cleaning

No	Tweet	Hasil Cleaning
1	Ini maksudnya apa? Ko bisa cakar2an dalam kabinet? Bagaimana mau bikin indonesia cerah kalo caranya begini	Ini maksudnya apa Ko bisa cakaran dalam kabinet Bagaimana mau bikin indonesia cerah kalo caranya begini
2	Dengan Kabinet Merah Putih yang semakin solid semangat optimisme untuk Indonesia Maju kini semakin nyata di depan mata.	Dengan Kabinet Merah Putih yang semakin solid semangat optimisme untuk Indonesia Maju kini semakin nyata di depan mata
...	...	...
1944	@dinopattidjalal @prabowo @Menlu_RI Sependapat fokus benahi kabinet merah putih ganti menteri2 termul sy yakin dukungan rakyat makin besar dan tdk kecewa !	Sependapat fokus benahi kabinet merah putih g anti menteri termul sy yakin dukungan rakyat makin besar dan tdk kecewa

2. Case-folding

Proses *case folding* mengonversi seluruh teks dalam dokumen menjadi huruf kecil untuk mengatasi ketidakkonsistenan dalam penggunaan huruf kapital [6]. Hal ini penting dilakukan karena sebuah kata yang terdiri dari huruf yang sama akan dibaca oleh sistem sebagai kata yang berbeda jika huruf tersebut memiliki ukuran huruf yang berbeda [15]. Berikut adalah contoh hasil *case-folding* pada tabel 3 berikut:

Tabel 3. Hasil *case-folding*

N	Tweet	Hasil <i>case-folding</i>
0		
1	Ini maksudnya apa? Ko bisa cakaran dalam kabinet Bagaimana mau bikin indonesia cerah kalo caranya begini	ini maksudnya apa ko bisa cakaran dalam kabinet bagaimana mau bikin indonesia cerah kalo caranya begini
2	Dengan Kabinet Merah Putih yang semakin solid semangat optimisme untuk Indonesia Maju kini semakin nyata di depan mata	Dengan kabinet merah putih yang semakin solid semangat optimisme untuk indonesia maju kini semakin nyata di depan mata
...	...	...
19	Sependapat fokus	sependapat fokus
44	benahi kabinet merah putih ganti menteri termul sy yakin dukungan rakyat makin besar dan tdk kecewa	benahi kabinet merah putih ganti menteri termul sy yakin dukungan rakyat makin besar dan tdk kecewa

3. Normalisasi

Proses ini bertujuan untuk mengubah kata informal menjadi bentuk kata baku. Hal ini dilakukan karena pada umumnya banyak sekali yang menggunakan kata kata gaul atau pun singkatan seperti: yg, okeh, bnr, pngn, km, mantep, manteppppp, dan lain-lain [16]. Berikut akan ditunjukkan hasil normalisasi pada tabel 4 berikut:

Tabel 4. Hasil Normalisasi

No	Tweet	Hasil Normalisasi
1	ini maksudnya apa ko bisa cakaran dalam kabinet bagaimana mau bikin indonesia cerah kalo caranya begini	ini maksudnya apa kenapa bisa cakaran dalam kabinet bagaimana mau membuat indonesia cerah kalau caranya seperti ini
2	dengan kabinet merah putih yang semakin solid semangat optimisme untuk indonesia maju kini semakin nyata di depan mata	dengan kabinet merah putih yang semakin solid semangat optimisme untuk indonesia maju kini semakin nyata di depan mata
...	...	...
1944	sependapat focus benahi kabinet merah putih ganti menteri termul saya yakin dukungan rakyat makin besar dan tdk kecewa	sependapat fokus benahi kabinet merah putih ganti menteri termul saya yakin dukungan rakyat semakin besar dan tidak kecewa

4. Penghapusan Stopwords

Stopwords adalah kata-kata yang umumnya dianggap tidak memberikan informasi penting dalam pemrosesan teks karena kemunculannya yang sering, seperti 'yang', 'di', dll [17]. Proses ini membantu meningkatkan kualitas analisis teks dengan fokus pada kata-kata kunci atau kata-kata yang lebih informatif [18]. Hasil penghapusan stopwords dapat dilihat pada tabel 5 berikut:

Tabel 5. Hasil Stopwords Removal

No	Tweet	Hasil Stopword removal
1	ini maksudnya apa kenapa bisa cakaran dalam kabinet bagaimana mau membuat indonesia cerah kalau caranya begini	maksudnya cakaran kabinet mau membuat Indonesia cerah caranya begini
2	dengan kabinet merah putih yang semakin solid semangat optimisme untuk indonesia maju kini semakin nyata di depan mata	kabinet merah putih semakin solid semangat optimisme indonesia maju semakin nyata depan mata
...	...	...
1944	sependapat fokus benahi kabinet merah putih ganti menteri termul saya yakin dukungan rakyat makin besar dan tidak kecewa	sependapat fokus benahi kabinet merah putih ganti menteri termul saya yakin dukungan rakyat makin besar tidak kecewa

Pelabelan Data

Pelabelan data ini dilakukan secara otomatis dengan bantuan library python, yaitu textblob. Library ini tergolong ringan dan cukup mudah digunakan. Ada 3 label yang akan digunakan pada penelitian ini, yaitu positif, negatif dan netral. Hasil pelabelan dapat dilihat pada tabel 6 berikut:

Tabel 6. Hasil Pelabelan

No	Tweet	Label
1	maksudnya cakaran kabinet mau membuat indonesia cerah begini	netral
2	kabinet merah putih semakin solid semangat optimisme indonesia maju semakin nyata depan mata	positif
...	...	...
1944	sependapat fokus benahi kabinet merah putih ganti menteri termul saya yakin dukungan rakyat makin besar tidak kecewa	netral

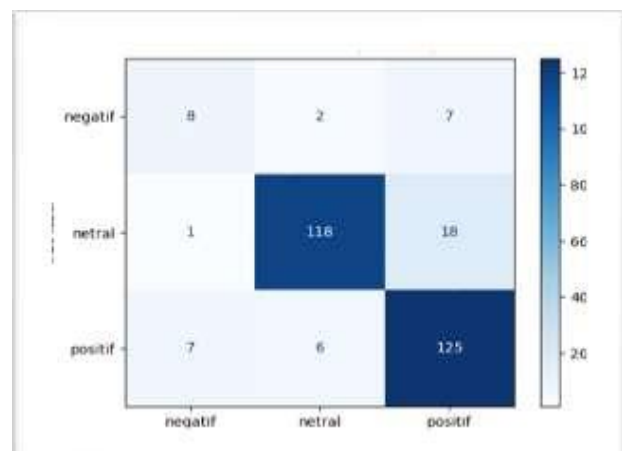
Splitting Data

Dataset yang terdiri dari 1944 data dibagi menjadi 3 bagian, yaitu 70% data digunakan untuk training model, 15% data digunakan untuk validasi dan 15% data digunakan sebagai data uji.

Implementasi BERT

Ada beberapa tahapan yang dilakukan dalam implementasi BERT, diantaranya sebagai berikut:

1. **Tokenisasi BERT**
Proses ini dilakukan untuk mengubah kalimat menjadi bagian yang lebih kecil (token).
2. **Embedding**
Embedding merupakan proses mengubah data yang awalnya berupa teks menjadi vektor numerik, supaya data tersebut dapat diproses oleh komputer.
3. **Self-attention**



Gambar 3. Confusion Matrix

Self-attention adalah mekanisme yang memungkinkan setiap token dalam kalimat memperhatikan token lainnya, termasuk dirinya sendiri, untuk memahami konteks secara lebih luas. Mekanisme ini memungkinkan

model untuk menangkap hubungan antara komponen-komponen yang berjauhan dalam sebuah kalimat dan menggunakannya sebagai pertimbangan dalam menghasilkan output akhir.

4. Feed Forward Network (FFN)

Feed Forward Network dalam BERT adalah jaringan saraf dua layer (*dense/fully connected layer*) yang berfungsi mengubah representasi token secara non-linear, meningkatkan kemampuan jaringan dalam memodelkan hubungan kompleks.

5. Clasification Head

Lapisan ini merupakan *dense layer (fully connected layer)* yang akan ditambahkan di atas lapisan BERT. Lapisan *dense* ini dapat digunakan untuk mengubah representasi teks menjadi prediksi kelas yang spesifik.[12]

Evaluasi

Evaluasi dilakukan dengan menggunakan *confusion matrix* untuk melihat seberapa baik performa dari model. Evaluasi ini penting untuk melihat efektivitas model dalam menangani masalah klasifikasi sentimen. Performa model pada tahap uji dapat dilihat pada tabel 7 berikut ini:

Tabel 7. Hasil Evaluasi

	Prediksi Negatif	Prediksi Netral	Prediksi Positif	Class Recall
True Negatif	8	2	7	47%
True Netral	1	118	18	86%
True Positif	7	6	125	91%
Class Precision	50%	94%	83%	

Berdasarkan tabel di atas, dapat dilihat bahwa:

1. Model menunjukkan performa yang kurang baik, yaitu precision 50% dan recall 47% pada kelas sentimen negatif yang diakibatkan oleh jumlah data yang terlalu sedikit pada kelas sentimen tersebut.
2. Performa model dalam mengklasifikasikan kelas sentimen netral cukup baik dengan nilai precision 94% dan recall 86%.
3. Performa model dalam mengklasifikasikan kelas sentimen positif cukup baik dengan nilai precision 83% dan recall 91%.

Visualisasi dari confusion matrix dapat dilihat pada gambar 3 berikut ini:

Kontribusi Penelitian

Penelitian ini berkontribusi penting dalam bidang analisis sentimen berbasis *deep learning*, dimana hasil dari penelitian ini menunjukkan bahwa performa model *deep learning*, yaitu BERT untuk analisis sentimen khususnya mengenai Kabinet Merah Putih dan data yang diambil dari platform X cukup baik. Selain itu, penelitian ini juga memberikan gambaran objektif mengenai sentimen masyarakat terhadap Kabinet Merah Putih, khususnya di platform X.

IV. KESIMPULAN

Berdasarkan temuan dari penelitian ini, diperoleh beberapa Kesimpulan, diantaranya: Dari 1944 data yang didapatkan, ada sebanyak 921 data positif, 914 netral dan 109 data negatif. Data dengan sentiment negative yang hanya 109 (6%) menunjukkan minimnya kritik atau ketidakpuasan masyarakat khususnya di platform X terhadap Kabinet Merah Putih. Berdasarkan distribusi data ini, citra Kabinet Merah Putih di media sosial X tampak relatif positif, tidak ada dominasi besar dari sentimen negatif yang menunjukkan krisis kepercayaan atau penolakan publik secara signifikan.

Berdasarkan hasil evaluasi terhadap data uji, model indobert-base-p1 yang digunakan pada penelitian ini memiliki performa yang cukup baik dengan nilai akurasi 85.96%. Nilai *precision*, *recall*, dan F1-Score masing-masing mencapai 0.75, 0.74 dan 0.74 bahwa model cukup baik dalam mengklasifikasikan sentimen publik terhadap Kabinet Merah Putih. Namun performa model pada kelas negatif masih relatif rendah dengan nilai F1-score 0.48, yang menunjukkan adanya keterbatasan model dalam mengenali sentimen negatif. Hal ini disebabkan ketidakseimbangan distribusi data, dimana jumlah tweet negatif jauh lebih sedikit dibandingkan tweet positif dan netral.

V. REFERENSI

- [1] B. Kurniawan, A. Ari Aldino, and A. Rahman Isnain, "Sentimen Analisis terhadap Kebijakan Penyelenggara Sistem Elektronik (PSE) Menggunakan Algoritma Bidirectional Encoder Representations from Transformers (Bert)," *J. Teknol. dan Sist. Inf.*, vol. 3, no. 4, pp. 98–106, 2022, [Online]. Available: <http://jim.teknokrat.ac.id/index.php/JTSI>
- [2] Diana Dwi Rahayu, Muhammad Fatchan, and Alfonsus Ligouri, "Analisis Sentimen Twitter Terpilihnya Prabowo - Gibran Menggunakan Metode Neural Network," *Tematik*, vol. 11, no. 1, pp. 85–91, 2024, doi: 10.38204/tematik.v11i1.1943.
- [3] M. A. Rizaty, "Daftar Negara Pengguna X Terbanyak di Dunia per April 2024, Indonesia Keempat," *dataindonesia.id*. Accessed: May 05, 2025. [Online]. Available: <https://dataindonesia.id/internet/detail/daftar-negara-pengguna-x-terbanyak-di-dunia-per-april-2024-indonesia-keempat>
- [4] A. Krisna, Haycal, and Z. Fathanul, "Media 2024 ANALISIS SENTIMEN MASYARAKAT TERHADAP PELANTIKAN KABINET MERAH PUTIH PADA MEDIA SOSIAL MENGGUNAKAN ALGORITMA NAIVE BAYES," pp. 1080–1089, 2024.
- [5] M. Iqbal and R. E. A. Sartika, "Tagar #IndonesiaGelap Jadi Trending, Bentuk Sentimen Negatif Terhadap Pemerintah," *Kompas.com*. Accessed: May 23, 2025. [Online]. Available: https://www.kompas.com/tren/read/2025/02/20/095352665/tagar-indonesiagelap-jadi-trending-bentuk-sentimen-negatif-terhadap?lgm_method=google&google_btn=onetap
- [6] R. R. S. Elvika Alya Junita, "Analisis Sentimen Hate Speech Mengenai Calon Wakil Presiden Indonesia Menggunakan Algoritma Bert," *JUPI (Jurnal Ilm. Penelit. dan Pembelajaran Inform.*, vol. 8, no. 3, pp. 1–13, 2024.
- [7] D. Duei Putri, G. F. Nama, and W. E. Sulistiono, "Analisis Sentimen Kinerja Dewan Perwakilan Rakyat (DPR) Pada Twitter Menggunakan Metode Naive Bayes Classifier," *J. Inform. dan Tek. Elektro Terap.*, vol. 10, no. 1, pp. 34–40, 2022, doi: 10.23960/jitet.v10i1.2262.



- [8] N. Putu, V. D. Saraswati, N. Yudistira, and P. P. Adikara, "Analisis Sentimen terhadap Perundungan Siber pada Twitter menggunakan Algoritma Bidirectional Encoder Representations from Transformer (BERT)," *J. Pengemb. Teknol. Inf. dan Ilmu Komput.*, vol. 7, no. 2, pp. 909–916, 2023, [Online]. Available: <https://j-ptiik.ub.ac.id/index.php/j-ptiik/article/view/12345>
- [9] D. F. Sjoraida, B. W. K. Guna, and D. Yudhakusuma, "Analisis Sentimen Film Dirty Vote Menggunakan BERT (Bidirectional Encoder Representations from Transformers)," *J. JTik (Jurnal Teknol. Inf. dan Komunikasi)*, vol. 8, no. 2, pp. 393–404, 2024, doi: 10.35870/jtik.v8i2.1580.
- [10] H. Syah and A. Witanti, "Analisis Sentimen Masyarakat Terhadap Vaksinasi Covid-19 Pada Media Sosial Twitter Menggunakan Algoritma Support Vector Machine (Svm)," *J. Sist. Inf. dan Inform.*, vol. 5, no. 1, pp. 59–67, 2022, doi: 10.47080/simika.v5i1.1411.
- [11] D. R. Manalu, M. C. L. Tobing, and M. Yohanna, "ANALISIS SENTIMEN TWITTER TERHADAP WACANA PENUNDAAN PEMILU DENGAN METODE SUPPORT VECTOR MACHINE," vol. 6, no. 2, pp. 149–156, 2024.
- [12] N. Karimah and A. Baita, "Multi-Aspect Sentiment Analysis Pada Review Film Menggunakan Metode Bidirectional Encoder Representations From Transformers (BERT) Multi-Aspect Sentiment Analysis of Film Review Using Bidirectional Encoder Representations from Transformers (BERT)," *J. Sist. Komput.*, vol. 13, no. 1, p. 2020, 2024, doi: 10.34010/komputika.v13i1.11098.
- [13] R. Kusnadi, Y. Yusuf, A. Andriantony, R. Ardian Yaputra, and M. Caintan, "Analisis Sentimen Terhadap Game Genshin Impact Menggunakan Bert," *Rabit J. Teknol. dan Sist. Inf. Univrab*, vol. 6, no. 2, pp. 122–129, 2022, doi: 10.36341/rabit.v6i2.1765.
- [14] S. Dharmawan, V. C. Mawardi, and N. J. Perdana, "Klasifikasi Ujaran Kebencian Menggunakan Metode FeedForward Neural Network (IndoBERT)," *J. Ilmu Komput. dan Sist. Inf.*, vol. 11, no. 1, pp. 1–6, 2023, doi: 10.24912/jiksi.v11i1.24066.
- [15] M. mahrus Zain, "Analisis Sentimen Pendapat Masyarakat Mengenai Vaksin Covid-19 Pada Media Sosial Twitter dengan Robustly Optimized BERT Pretraining Approach," *J. Komput. Terap.*, vol. 7, no. 2, pp. 280–289, 2021, doi: 10.35143/jkt.v7i2.4782.
- [16] Nelly Sofi, Tri Sulistyorini, and Muhammad Nazaruddin, "Analisis Sentimen Masyarakat Pengguna Media Sosial Twitter Terhadap Motogp Mandalika Lombok Menggunakan Metode Bidirectional Encoder Representation From Transformers (BERT)," *J. Informasi, Sains dan Teknol.*, vol. 6, no. 1, pp. 120–130, 2023, doi: 10.55606/isaintek.v6i1.103.
- [17] R. Merdiansah and A. Ali Ridha, "Sentiment Analysis of Indonesian X Users Regarding Electric Vehicles Using IndoBERT," *J. Ilmu Komput. dan Sist. Inf. (JIKOMSI)*, vol. 7, no. 1, pp. 221–228, 2024.
- [18] A. Prabowo and F. Indra Sanjaya, "Penerapan Metode Transfer Learning Pada Indobert Untuk Analisis Sentimen Teks Bahasa Jawa Ngoko Lugu," *Simkom*, vol. 9, no. 2, pp. 205–217, 2024, doi: 10.51717/simkom.v9i2.478.