

ANALISIS SENTIMEN PADA ULASAN APLIKASI ASTRO – GROCERIES IN MINUTES DI GOOGLE PLAY STORE MENGGUNAKAN METODE BIDIRECTIONAL ENCODER REPRESENTATION FROM TRANSFORMERS (BERT)

Aghi Kalam Ibrahim^{1*}, Metty Mustikasari², Irwan Bastian³

^{1,2,3}Sistem Informasi, Universitas Gunadarma; Jl. Margonda Raya 100, Kota Depok

Received: 2 Januari 2025

Accepted: 14 Januari 2025

Published: 20 Januari 2025

Keywords:

Sentiment Analysis, Astro – Groceries in Minutes, BERT

Correspondent Email:

aghikalam2@gmail.com

Analisis sentimen merupakan proses penting dalam memahami persepsi pengguna terhadap suatu aplikasi, terutama dalam konteks aplikasi yang tengah berkembang pesat namun masih memiliki tantangan dalam mengenalkan layanannya kepada khalayak luas. Penelitian ini bertujuan untuk menganalisis sentimen pada ulasan pengguna aplikasi "Astro – Groceries in Minutes" yang terdapat di Google Play Store dengan menggunakan metode Bidirectional Encoder Representation from Transformers (BERT). Aplikasi Astro telah diunduh lebih dari satu juta kali, namun masih ada tantangan dalam meningkatkan kesadaran dan penggunaannya di kalangan masyarakat yang lebih luas. BERT adalah model pembelajaran mesin berbasis transformer yang mampu menangkap konteks kata secara bidirectional, sehingga dapat menghasilkan representasi teks yang lebih akurat. Dalam penelitian ini, data ulasan dikumpulkan, dipreproses, dan kemudian diklasifikasikan ke dalam kategori sentimen positif, negatif, atau netral. Hasil analisis menunjukkan bahwa metode BERT dapat secara efektif mengklasifikasikan sentimen dalam ulasan aplikasi dengan tingkat akurasi yang tinggi. Temuan ini diharapkan dapat memberikan wawasan bagi pengembang aplikasi dalam memahami umpan balik pengguna dan meningkatkan kualitas layanan aplikasi Astro.

Sentiment analysis is a crucial process in understanding user perceptions of an application, especially for apps that are rapidly growing but still face challenges in raising awareness among a broader audience. This study aims to analyze the sentiment in user reviews of the "Astro – Groceries in Minutes" app available on the Google Play Store using the Bidirectional Encoder Representations from Transformers (BERT) method. Although the Astro app has been downloaded over one million times, it still faces challenges in expanding its user base and awareness. BERT is a transformer-based machine learning model that captures word context bidirectionally, allowing for more accurate text representation. In this study, review data is collected, preprocessed, and then classified into positive, negative, or neutral sentiment categories. The analysis results show that the BERT method effectively classifies sentiment in app reviews with a high level of accuracy. These findings are expected to provide insights for app developers in understanding user feedback and improving the quality of the Astro app's services.

1. PENDAHULUAN

Di era digital saat ini, banyak aspek kehidupan, termasuk sektor belanja, telah

terpengaruh oleh kemajuan teknologi. Aplikasi Astro – Groceries in Minutes adalah salah satu layanan *e-commerce* yang menyediakan berbagai kebutuhan rumah tangga dengan akses mudah melalui perangkat seluler. Berdasarkan data dari Google Play Store per April 2024, aplikasi ini telah diunduh lebih dari 1 juta kali dengan *rating* 4.9/5.0 dari 20.500 ulasan pengguna [1].

Meskipun aplikasi Astro menawarkan layanan yang praktis dan mendapatkan *rating* yang tinggi, masih terdapat tantangan terkait kurangnya kesadaran masyarakat tentang keberadaan aplikasi ini. Banyak potensi pengguna yang belum mengetahui atau memanfaatkan aplikasi ini, sehingga perlu dilakukan upaya untuk memahami persepsi dan sentimen pengguna yang ada.

Untuk memahami persepsi tersebut, analisis sentimen menjadi salah satu pendekatan yang relevan. Analisis sentimen merupakan bidang penelitian yang berkembang pesat dalam pemrosesan bahasa alami, penambahan data, dan penggalian informasi dari web. Bidang ini juga bisa disebut ekstraksi opini, menggunakan teknik pemrosesan bahasa alami dan metode ekstraksi informasi untuk mengenali dan menganalisis teks[2]. Proses analisis sentimen ini juga memanfaatkan teknologi untuk mengategorikan pendapat dalam teks menjadi positif, negatif, atau netral. Hasil dari analisis sentimen sering kali menjadi dasar pengembangan sistem rekomendasi dan visualisasi data, membantu pengambilan keputusan berbasis informasi yang lebih mendalam.

Terdapat beberapa metode yang dapat digunakan dalam analisis sentimen. Salah satu metode yang dapat digunakan adalah metode BERT. BERT merupakan teknologi open source yang didasarkan pada jaringan saraf, dirancang khusus untuk melakukan pre-training dalam pemrosesan bahasa alami (NLP). Teknologi ini memungkinkan sistem komputer untuk memahami bahasa dengan cara yang lebih alami, mirip dengan pemahaman manusia. Salah satu keunggulannya adalah kemampuan dalam mengenali konteks kata, yang pada gilirannya menghasilkan pencarian dengan hasil yang lebih relevan [3].

Metode BERT telah banyak digunakan dalam berbagai penelitian yang mendukung pengembangan analisis sentimen. Sebagai

contoh, Zulwida Nurul Huda (2023) melakukan analisis sentimen pada ulasan aplikasi Alfagift dan berhasil mencapai akurasi 65% menggunakan model IndoBERT-base [4]. Penelitian lain oleh Zilvi Azus Sriyanti, Dhian Satria Yudha Kartika, dan Abdul Rezha Efrat Najaf (2024) dengan menggunakan metode BERT untuk menganalisis sentimen pengguna Twitter terhadap aksi boikot produk israel, mencapai akurasi tertinggi sebesar 85%[5]. Selain itu, Raden Mas Rizqi Wahyu Panca Kusuma Atmaja dan Wiyli Yustanti (2021) dalam analisis sentimen ulasan aplikasi Ruang Guru juga menemukan bahwa metode BERT sangat efektif dengan akurasi hingga 99% [6].

Keberhasilan dari penelitian tersebut tidak terlepas dari kemampuan unggul BERT dalam memahami konteks bahasa secara mendalam. Bidirectional Encoder Representations from Transformers (BERT) diakui karena kemampuannya dalam memahami konteks kalimat secara mendalam, menjadikannya sangat efektif dalam berbagai tugas, termasuk analisis sentimen. BERT menggunakan arsitektur transformer yang sepenuhnya bidirectional, yang memungkinkan model ini untuk mempertimbangkan konteks dari kedua sisi kalimat, sehingga memberikan pemahaman yang lebih komprehensif terhadap makna kalimat secara keseluruhan.

Diperkenalkan oleh Google pada tahun 2018, BERT telah merevolusi bidang Natural Language Processing (NLP) melalui pendekatan transfer learning. Model ini dilatih menggunakan korpus teks besar seperti Wikipedia dan BookCorpus, yang memfasilitasi proses fine-tuning dengan sedikit penyesuaian pada lapisan output untuk menangani berbagai tugas NLP, seperti klasifikasi teks, menjawab pertanyaan, dan pengenalan entitas.[7]

Salah satu fitur utama BERT adalah dua tugas pre-training yang diterapkan dalam proses pelatihannya, yaitu Masked Language Modeling (MLM) dan Next Sentence Prediction (NSP). Dalam MLM, beberapa kata dalam kalimat diacak, dan model dilatih untuk memprediksi kata-kata yang hilang berdasarkan konteks dua arah. Sementara itu, NSP mengharuskan model untuk memprediksi apakah sebuah kalimat secara logis mengikuti kalimat lainnya. Kombinasi dari kedua tugas ini memungkinkan BERT untuk mempelajari

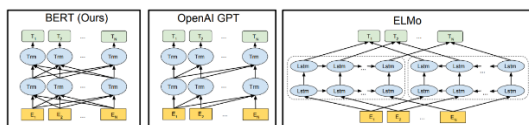
hubungan antar kalimat, menjadikannya sangat sesuai untuk tugas-tugas seperti kesamaan kalimat dan rangkuman.[7]

Berdasarkan uraian di atas, penelitian ini bertujuan untuk menganalisis sentimen pengguna terhadap aplikasi Astro melalui ulasan yang terdapat di Google Play Store dengan menggunakan metode BERT.

2. TINJAUAN PUSTAKA

2.1 Bidirectional Encoder Representation from Transformers (BERT)

Bidirectional Encoder Representation from Transformers (BERT) merupakan model pemrosesan bahasa alami (NLP) yang diperkenalkan oleh Google pada tahun 2018. BERT dirancang dengan memanfaatkan teknik-teknik pembelajaran mendalam (deep learning) serta berbagai metode, termasuk pembelajaran semi-terawasi. Model ini terinspirasi oleh beberapa pendekatan sebelumnya, seperti ELMo, ULMFiT, OpenAI Transformers, dan Transformers.[8]



Gambar 1. Perbedaan Antara Arsitektur BERT dengan OpenAI dan ELMo

Pada gambar 1, terdapat perbedaan dalam arsitektur model pra-pelatihan. BERT memanfaatkan *Transformer* dengan pendekatan dua arah, sementara OpenAI GPT mengandalkan *Transformer* yang beroperasi dari kiri ke kanan. ELMo, di sisi lain, menggabungkan LSTM yang dilatih secara independen dari kiri ke kanan dan dari kanan ke kiri untuk menghasilkan fitur yang dapat digunakan dalam tugas-tugas lanjutan. Di antara ketiga model tersebut, hanya representasi BERT yang secara simultan mempertimbangkan konteks dari sisi kiri dan kanan di semua lapisannya. Selain perbedaan dalam arsitektur, baik BERT maupun OpenAI GPT menerapkan pendekatan *fine-tuning*, sedangkan ELMo menerapkan pendekatan berbasis fitur.[8]

Sesuai dengan namanya, BERT memanfaatkan mekanisme *Transformer*. *Transformer* adalah suatu metode yang

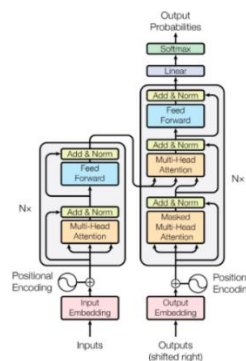
mempelajari hubungan kontekstual antara kata-kata dalam teks [9]. Mekanisme ini mampu memahami dan mengonversi pemahaman yang diperoleh melalui suatu proses yang dikenal sebagai *self-attention mechanism*. *Self-attention mechanism* merupakan cara yang digunakan oleh *Transformer* untuk mengubah "pemahaman" mengenai kata-kata terkait menjadi kata-kata yang akan diproses lebih lanjut. Dalam struktur *Transformer* terdapat dua mekanisme utama, yaitu:

2.1.1 Encoder

Encoder berfungsi untuk membaca seluruh input teks secara bersamaan. Struktur *encoder* terdiri dari tumpukan (*stack*) $N = 6$ lapisan yang identik. Setiap lapisan memiliki dua sub-lapisan yang memungkinkan node untuk tidak hanya fokus pada kata yang sedang dianalisis, tetapi juga memperoleh konteks semantik dari kata tersebut. Setiap posisi dalam *encoder* dapat mengakses semua posisi di lapisan sebelumnya.

2.1.2 Decoder

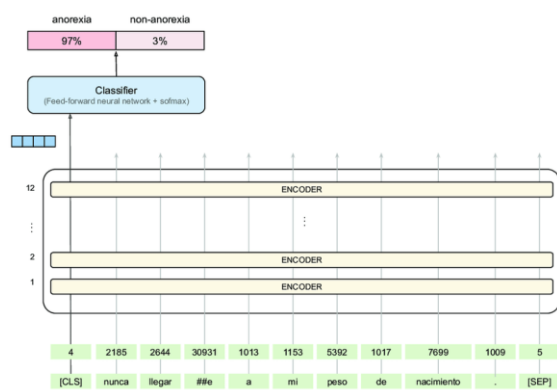
Decoder bertugas untuk menghasilkan urutan output yang berupa prediksi. Seperti halnya *encoder*, *decoder* juga terdiri dari tumpukan $N = 6$ lapisan yang identik. Setiap lapisan dalam *decoder* terdiri dari dua sub-lapisan yang mirip dengan yang ada pada *encoder*, dengan tambahan lapisan perhatian di antara kedua lapisan tersebut untuk membantu node saat ini dalam mendapatkan konten kunci yang membutuhkan perhatian melalui penerapan *multi-head attention* pada output dari *encoder*. [9] Sama seperti pada *encoder*, lapisan *self-attention* di *decoder* memungkinkan setiap posisi untuk menangani semua posisi sebelumnya serta posisi saat ini.



Gambar 2. Encoder (kiri) Decoder (kanan)

Pada gambar 2, arsitektur BERT terdiri dari encoder *Transformer* yang bersifat multi-layer

dan bidirectional. Transformer merupakan suatu mekanisme perhatian (attention) yang mempelajari hubungan kontekstual antara kata atau sub-kata dalam teks. Mekanisme self-attention dalam Transformer berfungsi untuk memahami representasi dari input dan output. Transformer memiliki dua mekanisme terpisah, yaitu encoder yang bertugas membaca teks input dan decoder yang berfungsi untuk menghasilkan prediksi. Setiap layer encoder terdiri dari dua sub-layer, yang pertama adalah mekanisme multi-head self-attention, dan yang kedua adalah jaringan feed-forward yang sepenuhnya terhubung. Dalam encoder, setiap input akan melewati layer self-attention, yang membantu encoder untuk mempertimbangkan kata-kata lain yang terdapat dalam kalimat. Output dari layer attention kemudian digunakan sebagai input untuk jaringan neural feed-forward.

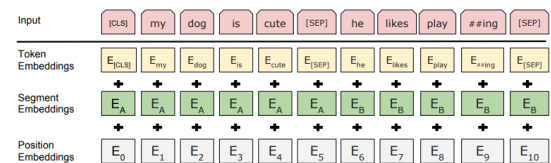


Gambar 3. Arsitektur BERT

Pada gambar 3 menggambarkan penerapan model BERT dalam klasifikasi teks, khususnya untuk menentukan apakah suatu teks mengindikasikan adanya anoreksia. Proses ini dimulai dengan pengolahan input teks yang di-tokenisasi menjadi token-token yang relevan. Hal ini ditunjukkan di bagian bawah gambar, di mana token [CLS] menandakan awal teks dan token [SEP] menandakan akhir teks.

BERT memanfaatkan WordPiece embeddings yang terdiri dari 30.000 token dalam vocabularynya. WordPiece Embedding dikembangkan untuk mengatasi tantangan segmentasi dalam bahasa Jepang dan Korea, yang merupakan bagian dari sistem pengenalan suara Google.[10] Masalah segmentasi sering kali menghasilkan banyak kata yang tidak terdaftar dalam vocabulary. WordPiece dilatih untuk secara otomatis mempelajari unit-unit kata dari kumpulan data yang besar, sehingga

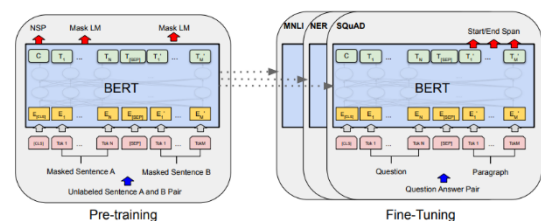
meminimalkan kemungkinan terjadinya out of vocabulary (OOV). Kata-kata yang sering muncul akan tetap dipertahankan sebagai satu kesatuan, sementara kata-kata yang jarang digunakan akan dipecah menjadi sub-kata atau bahkan karakter.



Gambar 4. Representasi Input BERT

Pada gambar 4 mengilustrasikan proses pemrosesan input teks oleh BERT dengan memanfaatkan tiga jenis embedding, yaitu token embedding, segment embedding, dan position embedding.

Framework BERT terdiri dari dua tahap utama, yaitu pre-training dan fine-tuning. Pada tahap pre-training, model mempelajari bahasa dan konteksnya melalui dua metode, yaitu *Mask Language Modeling* (MLM) dan *Next Sentence Prediction* (NSP) secara bersamaan. MLM memberikan kemampuan untuk mengintegrasikan konteks dari sisi kiri dan kanan, yang mendukung pelatihan transformator dua arah yang mendalam, sedangkan NSP berfungsi untuk menggabungkan representasi dari pasangan teks yang telah dilatih sebelumnya.



Gambar 5. Ilustrasi Pre-training dan Fine-tuning

Pada gambar 5 menggambarkan proses pre-training dan fine-tuning pada model BERT. Selama tahap pre-training, model dilatih untuk memahami bahasa dengan menggunakan teknik Masked Language Model (MLM) dan Next Sentence Prediction (NSP). Setelah itu, model akan menjalani proses fine-tuning untuk disesuaikan dengan tugas-tugas spesifik, seperti klasifikasi teks (MNLI/NER) dan pemahaman teks (SQuAD).

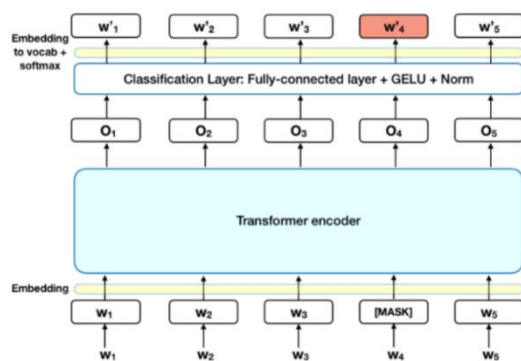
2.1.3 Pre-training BERT

BERT merupakan salah satu model bahasa yang inovatif dalam bidang pemrosesan bahasa

alami. Proses pre-training BERT mencakup langkah-langkah sebagai berikut

2.1.3.1 Masked LM

Pada tahap ini, teks yang tidak diawasi digunakan untuk melatih model. BERT mengambil sebuah kalimat dan secara acak menggantikan beberapa token dengan token khusus "[MASK]". Model kemudian memprediksi ID *vocabulary* dari kata yang sebenarnya berdasarkan konteks kalimat tersebut. Untuk melakukan prediksi, diperlukan penambahan lapisan klasifikasi di atas *output encoder*. Selanjutnya, dilakukan perkalian vektor output dengan matriks embedding, yang kemudian diubah menjadi dimensi *vocabulary*, serta menghitung probabilitas dari setiap kata dalam *vocabulary* menggunakan fungsi softmax.[11] Gambar 6 memperlihatkan proses *Masked LM*.



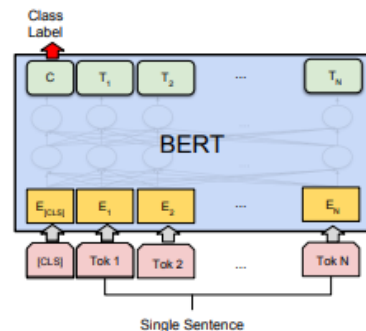
Gambar 6. Proses *Masked Language Modeling*

Pada gambar 6 menggambarkan proses Masked Language Modelling (MLM) yang diterapkan pada model BERT. Pada tahap ini, sejumlah token dalam kalimat digantikan dengan token "[MASK]". Model kemudian melakukan prediksi terhadap kata yang hilang dengan mempertimbangkan konteks yang tersedia. Proses ini melibatkan penambahan lapisan klasifikasi di atas output dari encoder, dilanjutkan dengan perkalian menggunakan matriks embedding dan penerapan fungsi softmax untuk menghitung probabilitas kata dalam kosakata.

2.1.4 Fine-Tuning BERT

Fine-tuning pada model BERT adalah proses melatih kembali model BERT yang telah dilatih secara umum pada tugas spesifik yang lebih terbatas menggunakan dataset berlabel. Untuk melatih sebuah model bahasa, classifier

perlu dilatih dengan sedikit perubahan pada model BERT selama fase pelatihan yang disebut *fine-tuning*. *Fine-tuning* sangat mudah dilakukan karena mekanisme *self-attention* pada transformer membuat BERT bisa membuat model untuk berbagi tugas, baik pada kalimat tunggal atau kalimat berpasangan, dengan menukar masukan dan keluaran yang sesuai.[11]



Gambar 7. Ilustrasi Fine Tuning pada Tugas *Single Sentence*

Pada gambar 7 menggambarkan proses fine-tuning pada tugas single sentence menggunakan model BERT. Pada tahap ini, model yang telah dilatih sebelumnya diadaptasi untuk tugas tertentu dengan menggunakan dataset berlabel. Proses fine-tuning dilakukan dengan menyesuaikan model BERT untuk memprediksi label kelas dari kalimat tunggal, memanfaatkan mekanisme *self-attention* pada transformer untuk mengelola masukan dan keluaran yang relevan.

3. METODE PENELITIAN

Tahapan penelitian yang digunakan dalam analisis sentimen pengguna aplikasi Astro – Groceries in Minutes menggunakan metode Bidirectional Encoder Representation from Transformers (BERT) terdiri dari beberapa langkah seperti berikut:

3.1 Scraping Data

Data yang digunakan dalam penelitian ini didapat dari reviews aplikasi Astro – Groceries in Minutes di Google Play Store menggunakan teknik pengambilan data menggunakan teknik web scraping. web scraping adalah praktik mengumpulkan data dengan cara selain program yang berinteraksi dengan API (atau, tentunya, melalui manusia yang menggunakan browser web).[12] Python sendiri memiliki beberapa library untuk melakukan scraping data

salah satunya yaitu google play scraper. Untuk penelitian ini penulis menggunakan google play scraper untuk memperoleh data mentah.

3.2 Labelisasi Dataset

Dalam melakukan analisis sentimen dengan metode *supervised learning*, diperlukan *dataset* dengan label atau anotasi karena metode ini membutuhkan contoh. *Supervised Learning* adalah suatu mekanisme yang dapat menghasilkan generalisasi sehingga *output* dari model dapat melihat, memahami dan mengerti bagaimana ulasan yang memiliki kategori *positive*, *neutral*, dan *negative*. Dalam proses labelisasi *dataset* jika ulasan memiliki nilai 4-5 akan dimasukkan kedalam kategori *positive*, jika nilai ulasan bernilai 3 maka akan dimasukkan kedalam kategori *neutral*, dan jika nilai ulasan bernilai 1-2 maka akan dimasukkan kedalam kategori *negative*.

3.3 Pre-processing Data

Pada tahap ini pre-processing dilakukan untuk mengubah dataset yang semula tidak terstruktur menjadi terstruktur sehingga mempermudah data untuk diproses dengan beberapa tahapan seperti case folding, data cleaning, tokenisasi dan normalisasi. Teknik ini melibatkan validasi dan imputasi data. Validasi memiliki tujuan untuk memberi nilai Tingkat kelengkapan dan akurasi dari data yang sudah tersaring, dan imputasi data memiliki tujuan untuk memperbaiki data dari kesalahan dan memasukkan nilai yang hilang baik secara manual maupun otomatis.

3.4 Dataset Splitting

Dataset splitting adalah suatu Teknik yang digunakan untuk melihat kinerja model dengan melakukan pembagian terhadap data yang akan diolah menjadi beberapa bagian, dalam hal ini training, validation dan testing. Dataset training digunakan untuk melatih model, dataset validasi digunakan untuk meminimalisir overfitting yang sering terjadi pada jaringan syaraf tiruan, sedangkan dataset testing sendiri digunakan sebagai tes akhir untuk melihat keakuratan jaringan yang sudah dilatih dengan dataset training. Proporsi split dataset pada penelitian ini adalah 70% train set. 20% validation set, 10% test set.

3.5 Implementasi BERT

Pada penelitian ini menggunakan sebuah model *indobert-base-p1* dari *indoBERT*. Library yang digunakan adalah library *Transformers* yang disediakan oleh *Hugging Face*. Library tersebut berisikan sejumlah besar model yang telah terlatih dan dapat digunakan untuk berbagai macam tugas seperti klasifikasi, translasi, ekstrak informasi dan lain-lain dalam 100 bahasa. *Transformers* menggunakan library populer seperti *PyTorch* dan *TensorFlow* untuk membuat model mesin pembelajaran.

3.6 Training Data (Fine Tuning)

Tahap ini juga disebut tahap training. Training dilakukan dengan menggunakan model yang telah dilatih sebelumnya lalu belajar sedikit lagi untuk mencapai titik optimal. Pada tahap ini dataset akan dibagi menjadi tiga kategori yaitu dataset untuk training, validation, dan testing. Hal ini bertujuan untuk membuat machine learning mengetahui tujuan dari setiap dataset.

3.7 Evaluasi Model dan Visualisasi Hasil

Tahap evaluasi dilakukan untuk melihat hasil analisis sentiment terhadap kalimat yang ada pada dataset. Nilai akurasi tertinggi yang didapat pada proses sebelumnya akan menjadi nilai dari akurasi model.

4. HASIL DAN PEMBAHASAN

Metode yang digunakan untuk scraping menggunakan google play scraper yang ada pada library python. Scraping dilakukan agar mendapatkan komentar pengguna aplikasi *Astro – Groceries in Minutes*. Ulasan yang telah didapat dari Google Play Store berdasarkan penilaian sebanyak 5000 data. Data yang digunakan dari penilaian dan komentar yang diberikan oleh pengguna. Pada tabel 1 adalah total dari hasil scraper

Tabel 1. Hasil Scraping

Sumber Data	Total Data
Google Play Store	5000

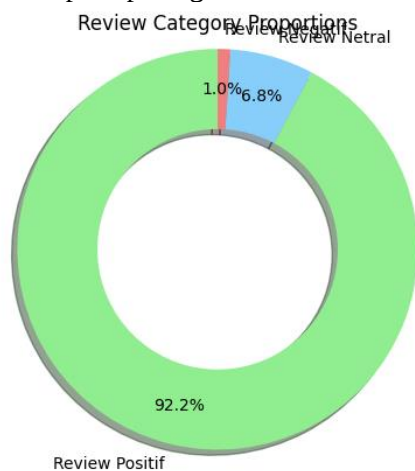
Detail pada kolom data yang tersedia tidak dapat di visualisasikan dalam artikel ini karena jumlah kolom yang terlalu banyak untuk

dimuat. Selanjutnya melabelisasi dataset. Pada tabel 2. Adalah hasil proses labelisasi data hasil dari scraping aplikasi Astro – Groceries in Minutes.

Tabel 2. Hasil Labelisasi

Ulasan	Rating
sy salah satu pelanggan yg setiap mggu pasti order di astro, pelayanan cepat, produk smua bagus kecuali baru kmrn beli ikan tp ud bau, ga sprti biasanya...jdnya ikannya smua dibuang....tolong diperhatikan lg produk ² yg dijual ya trutama ikan ² segarnya....tq	Positif
bener" cepat pengirimannya. jdi mudah untuk belanja	Positif
Kalo buat ngewarung kayaknya kurang cocok, harga beli sama harga jual di umumnya lebih tinggi di harga beli	Netral

Pada proses labelisasi akan diberikan 3 klasifikasi yaitu positif, netral, ataupun negatif. Untuk rating 1-2 akan diberikan label negatif, rating 3 diberikan label netral, dan rating 4-5 akan diberikan label positif. Selanjutnya dibuat jumlah persentasi dari setiap klasifikasi sentimen seperti pada gambar 8.

**Gambar 8.** Persentase Klasifikasi Sentimen

Pada gambar 8 menunjukkan jumlah persentasi dari setiap klasifikasi sentimen

terlihat terdapat 92.2% review positif, 6.8% review netral, dan 1% review negatif.

Setelah dataset dilabelisasi, selanjutnya data akan di lakukan pre-processing yaitu data akan dilakukan beberapa proses seperti *case folding*, *data cleaning*, tokenisasi dan normalisasi. Tabel 3 adalah contoh hasil dari pre-processing

Tabel 3. Hasil Pre-processing

Ulasan	Rating
harga nya sini lebih mahal bukan harga pasar contoh cleo ml harga pasar kak kak sini kak belcube pasar kak kak sini kak masih banyak lagi beda harga nya belum lagi daging tentu kayak ikan tongkol kemarin saya pesan sudah enggak fresh momogi hampir e sudah enggak bisa makan	Netral
guna baru dapat diskon harga nya ribu total pesan percuma ada minimal harga cowok lag banget juga aplikasi	Negatif
jelek banget sistem pesan batal cara pihak padahal saya sudah konfirmasi dengan driver untuk temu di lobby tapi tiba batal sama pusat tanpa alas jelas kasihan juga drivernya sudah jauh nganter	Negatif

Selanjutnya setelah dilakukan pre-procesing, data akan dilakukan proses splitting. Dataset dibagi menjadi tiga bagian yaitu *dataset training*, *dataset validation*, dan *dataset test*. proporsi pembagian dataset adalah 70:20:10 seperti pada tabel 4. Adalah hasil split antara data *training*, *validation*, dan *test*

Tabel 4. Hasil Split Antara Data Latih dan Data Uji

Type Data	Persentase	Jumlah
Training	70%	3496
Validation	20%	1004
Test	10%	495

Pada tabel 4 dapat dilihat hasil dari split training mendapatkan jumlah data sebesar 3496 data, validation sebesar 1004 data, dan test sebesar 495 data.

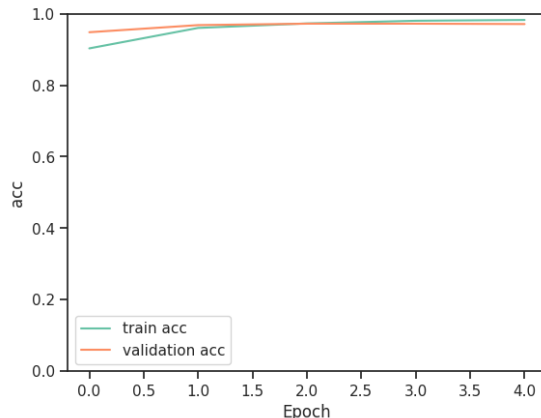
Setelah proses splitting data, dilakukanlah implementasi BERT dengan menggunakan model IndoBERT yang diterapkan dalam proses pre-training. Selanjutnya data dilakukan training data menggunakan teknik fine tuning, proses ini dilakukan melalui beberapa langkah, yaitu mendefinisikan optimizer Adam dengan learning rate yang rendah. Seperti pada tabel 5 adalah parameter model BERT

Tabel 5. Parameter BERT model

Batch Size	32
Learning Rate	3e-6
epochs	5

Pada tahap ini menggunakan batch size sebesar 32. Untuk epoch adalah 10 epoch dikarenakan agar tidak terjadi model kehilangan kemampuan untuk generalisasi ke data baru dan menghindari waktu pelatihan yang lama.

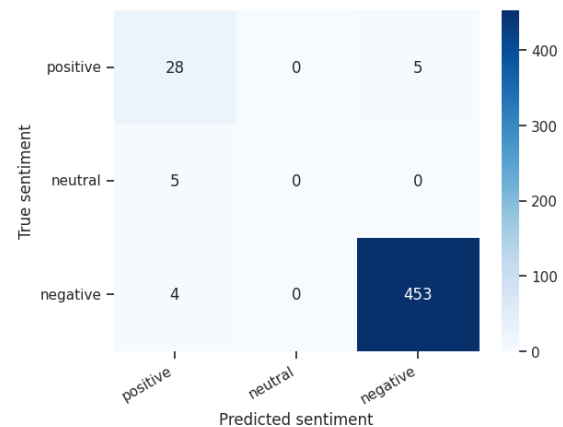
Training history



Gambar 9. Kurva Hasil Performa Training dan validasi

Pada gambar 9. Dari hasil pengamatan terhadap percobaan tersebut, didapat bahwa akurasi saat *training* lebih baik dari pada saat validasi. Kurva menunjukkan adanya peningkatan hasil akurasi yang diperoleh saat *training*. Meskipun hasil validasi juga meningkat, peningkatannya tidak sebesar peningkatan selama *training*.

Tahap terakhir yaitu tahap evaluasi. Pada tahap ini melakukan evaluasi menggunakan confusion matrix yang berfungsi untuk mengukur akurasi prediksi sentimen dalam data uji.



Gambar 10. Diagram confusion matrix

Dari gambar 10 diagram confusion matrix, dapat disimpulkan bahwa sebanyak 28 data termasuk kategori True Positive (TP), 0 data termasuk dalam kategori True Neutral (TNt), 453 data termasuk dalam kategori True Negative (TN), 5 data termasuk dalam kategori False Positive (FP), 5 data termasuk dalam kategori False Neutral (FNt), dan 4 data termasuk dalam kategori False Negative (FN).

Setelah mendapatkan hasil dari *confusion matrix*, dilakukan perhitungan *accuracy*, *precision*, *recall*, dan *fi-score* menggunakan fungsi *classification_report()* dari *sklearn*. Dapat dilihat pada Tabel 6.

Tabel 6. Klasifikasi Report.

	Precision	Recall	Fi-score	Support
Positive	0.76	0.85	0.80	33
Neutral	0.00	0.00	0.00	5
Negative	0.99	0.99	0.99	457
Accuracy			0.97	495
Macro avg	0.58	0.61	0.60	495
Weighted avg	0.96	0.97	0.97	495

Dari hasil classification report tersebut, diketahui bahwa tingkat akurasi dari model dalam memprediksi data testing yaitu 97%. Model mampu mengklasifikasikan dengan benar sebanyak 481 data sentimen dari total 495 data testing. Selain itu, model mendapatkan tingkat precision untuk kelas negatif yang sangat tinggi dibanding kelas netral dan positif, yaitu 99%. Untuk recall, kelas negatif juga memiliki nilai yang relatif lebih tinggi dibanding kelas positif dan netral, yaitu 99%.

Sedangkan untuk f1-score, kelas negatif lebih tinggi daripada kelas lainnya, yaitu 99%.

5. KESIMPULAN

Berdasarkan hasil pengujian analisis data ulasan aplikasi *Astro-Groceries in Minutes* pada Google Play Store dengan menggunakan metode *Bidirectional Encoder Representations from Transformers* (BERT). Pengujian ini menggunakan *dataset* dari ulasan aplikasi *Astro-Groceries in Minutes* sebanyak 5000 ulasan dengan data latih sebanyak 3496, data validasi sebanyak 1004, dan data uji sebanyak 495. Analisis sentimen ini dilakukan menggunakan *pre-trained* model IndoBERT-base-p1 dengan teknik *fine-tuning*. Didapatkan hasil presisi klasifikasi ulasan aplikasi *Astro-Groceries in Minutes* yaitu sebesar 76% ulasan positif, 0% ulasan netral dan 99% ulasan negatif. Model menghasilkan akurasi sebesar 97% dengan pemilihan *hyperparameter* yaitu batch size 32, *learning rate* $3e-6$, dan *epoch* 5.

UCAPAN TERIMA KASIH

Penulis mengucapkan terima kasih kepada pihak-pihak terkait yang telah memberi dukungan terhadap penelitian ini.

DAFTAR PUSTAKA

- [1] M. Aziz Muslim *et al.*, *Data Mining Algoritma C4.5 Disertai Contoh Kasus dan Penerapannya Dengan Program Computer*, vol. 1. 2019.
- [2] Reza Zulfiqui, Betha Nurina Sari, and Tesa Nur Padilah, "Analisis Sentimen Ulasan Pengguna Aplikasi Media Sosial Instagram Pada Situs Google Play Store Menggunakan Naive Bayes Classifier," *Jurnal Informatika dan Teknik Elektro Terapan*, vol. 12, no. 3, Aug. 2024, doi: 10.23960/jitet.v12i3.4995.
- [3] Tri Buwono Bagus Wicaksono and Rama Dian Syah, "Implementasi Metode Bidirectional Encoder Representations From Transformers Untuk Analisis Sentimen Terhadap Ulasan Aplikasi Access," *Jurnal Ilmiah Informatika Komputer*, vol. 29, no. 3, pp. 254–265, Dec. 2024, doi: 10.35760/ik.2024.v29i3.12514.
- [4] Z. Nurul Huda, "Analisis Sentimen Pada Ulasan Aplikasi Alfagift di Google Play Store Menggunakan Metode Bidirectional Encoder Representation from Transformers (BERT)," *Gunadarma Library*, 2023.
- [5] Zilvi Azus Sriyanti, Dhian Satria Yudha Kartika, and Abdul Rezha Efrat Najaf, "Implementasi Model BERT Pada Analisis Sentimen Pengguna Twitter Terhadap Aksi Boikot Produk Israel," *Jurnal Informatika dan Teknik Elektro Terapan*, vol. 12, no. 3, Aug. 2024, doi: 10.23960/jitet.v12i3.4743.
- [6] R. M. Rizqi Wahyu Panca Kusuma Atmaja and W. Yustanti, "Analisis Sentimen Customer Review Aplikasi Ruang Guru dengan Metode BERT (Bidirectional Encoder Representations from Transformers)," *JEISBI*, vol. 02, 2021.
- [7] Amandeep, *Implement NLP Use-cases Using BERT Explore the Implementation of NLP Tasks Using the Deep Learning Framework and Python Amandeep*, 1st ed. New Delhi: Manish Jain for BPB Publications, 2021. [Online]. Available: www.bpbonline.com
- [8] J. Devlin, M.-W. Chang, K. Lee, K. T. Google, and A. I. Language, "BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding," pp. 4171–4186, [Online]. Available: <https://github.com/tensorflow/tensor2tensor>
- [9] A. Vaswani *et al.*, "Attention is All you Need," in *Advances in Neural Information Processing Systems*, I. Guyon, U. Von Luxburg, S. Bengio, H. Wallach, R. Fergus, S. Vishwanathan, and R. Garnett, Eds., Curran Associates, Inc., 2017. [Online]. Available: https://proceedings.neurips.cc/paper_files/paper/2017/file/3f5ee243547dee91fbd053c1c4a845aa-Paper.pdf
- [10] Y. Wu *et al.*, "Google's Neural Machine Translation System: Bridging the Gap between Human and Machine Translation," Sep. 2016, [Online]. Available: <http://arxiv.org/abs/1609.08144>
- [11] Bella Rahmatullah, "Sentimen analisis pelaksanaan work from home di indonesia pada masa pandemi covid 19 menggunakan IndoBERT," *Institut Teknologi Sepuluh Nopember*, 2021.
- [12] R. Mitchell, *Web Scraping with Python COLLECTING MORE DATA FROM THE MODERN WEB*, 2nd ed. O'Reilly Media, 2018. [Online]. Available: www.allitebooks.com