

**Sentimen Analisis terhadap Penggemar Artis karena Endorsement
Produk Minuman yang masuk Daftar Hitam menggunakan Metode BERT**

LAPORAN TUGAS AKHIR

Laporan ini disusun guna memenuhi salah satu syarat untuk menyelesaikan program studi Teknik Informatika S-1 pada Fakultas Teknologi Industri Universitas Islam Sultan Agung



Disusun Oleh :

Nama : Furbaya Mutiara Ashar

NIM : 32602000027

**PROGRAM STUDI TEKNIK INFORMATIKA
FAKULTAS TEKNOLOGI INDUSTRI
UNIVERSITAS ISLAM SULTAN AGUNG
SEMARANG
2025**

FINAL PROJECT

**SENTIMENT ANALYSIS OF ARTIST FANS DUE TO BLACKLISTED BEVERAGE
PRODUCT ENDORSEMENTS USING THE BERT METHOD**

This report was prepared to fulfill one of the requirements to complete the S-1 Informatics Engineering study program at the Faculty of Industrial Technology, Sultan Agung Islamic University



Arranged By :

Name : Furbaya Mutiara Ashar

NIM : 32602000027

**MAJORING OF INFORMATICS ENGINEERING
INDUSTRIAL TECHNOLOGY FACULTY
SULTAN AGUNG ISLAMIC UNIVERSITY
SEMARANG**

2025

LEMBAR PENGESAHAN

TUGAS AKHIR

**SENTIMEN ANALISIS TERHADAP PENGGEAR ARTIS KARENA
ENDORSEMENT PRODUK MINUMAN YANG MASUK DAFTAR HITAM
MENGUNAKAN METODE BERT**

**FIRBAYA MUTIARA ASHAR
32602000027**

Telah dipertahankan di depan tim penguji ujian sarjana tugas akhir
Program Studi Teknik Informatika
Universitas Islam Sultan Agung
Pada tanggal : 27 Februari 2025

TIM PENGUJI UJIAN SARJANA :

Mustafa, S.T., MM., M.Kom
NIDN. 0623117703
(Ketua Penguji)

Arief Marwanto, S.T., M.Eng., Ph.D
NIDN. 0628097501
(Anggota Penguji)

Imam Much Ibnu Subroto, S.T., M.Sc., Ph.D
NIDN. 0613037301
(Pembimbing)

Semarang, 7 Maret 2025

Mengetahui,

Ketua Program Studi Teknik Informatika
Universitas Islam Sultan Agung


Moch Taufik, S.T., MIT
NIDN. 0622037502

SURAT PERNYATAAN KEASLIAN TUGAS AKHIR

Yang bertanda tangan dibawah ini :

Nama : Furbaya Mutiara Ashar

NIM : 32602000027

Judul Tugas Akhir : SENTIMEN ANALISIS TERHADAP PENGGEMAR ARTIS
KARENA ENDORSEMENT PRODUK MINUMAN YANG MASUK DAFTAR HITAM
MENGUNAKAN METODE BERT

Dengan bahwa ini saya menyatakan bahwa judul dan isi Tugas Akhir yang saya buat dalam rangka menyelesaikan Pendidikan Strata Satu (S1) Teknik Informatika tersebut adalah asli dan belum pernah diangkat, ditulis ataupun dipublikasikan oleh siapapun baik keseluruhan maupun sebagian, kecuali yang secara tertulis diacu dalam naskah ini dan disebutkan dalam daftar pustaka, dan apabila di kemudian hari ternyata terbukti bahwa judul Tugas Akhir tersebut pernah diangkat, ditulis ataupun dipublikasikan, maka saya bersedia dikenakan sanksi akademis. Demikian surat pernyataan ini saya buat dengan sadar dan penuh tanggung jawab.

Semarang, 28 Januari 2025

Yang Menyatakan,



Furbaya Mutiara Ashar

PERNYATAAN PERSETUJUAN PUBLIKASI KARYA ILMIAH

Saya yang bertanda tangan dibawah ini :

Nama : Furbaya Mutiara Ashar

NIM : 32602000027

Program Studi : Teknik Informatika

Fakultas : Teknologi Industri

Dengan ini menyatakan Karya Ilmiah berupa Tugas akhir dengan Judul : SENTIMEN ANALISIS TERHADAP PENGGEAR ARTIS KARENA ENDORSEMENT PRODUK MINUMAN YANG MASUK DAFTAR HITAM MENGGUNAKAN METODE BERT Menyetujui menjadi hak milik Universitas Islam Sultan Agung serta memberikan Hak bebas Royalti Non-Eksklusif untuk disimpan, dialihmediakan, dikelola dan pangkalan data dan dipublikasikan diinternet dan media lain untuk kepentingan akademis selama tetap menyantumkan nama penulis sebagai pemilik hak cipta. Pernyataan ini saya buat dengan sungguh-sungguh. Apabila dikemudian hari terbukti ada pelanggaran Hak Cipta/Plagiarisme dalam karya ilmiah ini, maka segala bentuk tuntutan hukum yang timbul akan saya tanggung secara pribadi tanpa melibatkan Universitas Islam Sultan agung.

Semarang, 28 Januari 2025

Yang menyatakan,



Furbaya Mutiara Ashar

KATA PENGANTAR

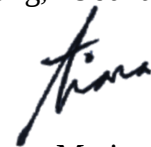
Puji syukur penulis panjatkan kepada Allah SWT, yang telah memberikan rahmat, taufik serta hidayah-Nya, sehingga penulis dapat menyelesaikan laporan Tugas Akhir dengan judul “Sentimen Analisis Terhadap Penggemar Artis Karena Endorsement Produk Minuman Yang Masuk Daftar Hitam Menggunakan Metode BERT” untuk memenuhi salah satu syarat menyelesaikan studi dan memperoleh gelar sarjana (S-1) DI Program Studi Teknik Informatika Fakultas Teknologi Industri Universitas Islam Sultan Agung.

Tak lupa penulis mengucapkan terima kasih kepada :

1. Rektor UNISSULA Bapak Prof. Dr. H. Gunarto, SH., M.Hum;
2. Dekan Fakultas Teknologi Industri Ibu Dr. Ir. Novi Marlyana, ST., MT., IPU., ASEAN Eng;
3. Kaprodi Teknik Informatika UNISSULA Bapak Moch Taufik, S.T., M.T;
4. Dosen Pembimbing Bapak Imam Much Ibnu Subroto, ST., M.Sc., Ph. D yang telah membimbing penulis sehingga penulis dapat menyusun laporan tugas akhir ini dengan benar.
5. Orang tua penulis yang telah mendukung dan mendoakan penulis untuk mengerjakan laporan tugas akhir ini.
6. Teman-teman penulis yang selalu menemani dan membantu setiap proses pengerjaan tugas akhir ini.
7. Dan kepada semua pihak yang tidak dapat saya sebutkan satu persatu.

Dengan kerendahan hati, penulis menyadari bahwa dalam penyusunan proposal tugas akhir ini masih terdapat banyak kekurangan, untuk itu penulis mengharap kritik dan saran untuk sempurnanya proposal ini. Semoga dengan ditulisnya proposal ini dapat menjadi sumber ilmu bagi setiap pembaca.

Semarang, 28 Januari 2025



Firda Mutiara Ashar

DAFTAR ISI

HALAMAN JUDUL	i
LEMBAR PENGESAHAN	iii
SURAT PERNYATAAN KEASLIAN TUGAS AKHIR.....	iv
PERNYATAAN PERSETUJUAN PUBLIKASI KARYA ILMIAH	v
KATA PENGANTAR.....	v
DAFTAR ISI.....	vi
DAFTAR GAMBAR.....	viii
DAFTAR TABEL	ix
ABSTRAK	1
BAB I PENDAHULUAN.....	2
1.1 Latar Belakang	2
1.2 Perumusan Masalah	3
1.3 Batasan Masalah	4
1.4 Tujuan Penelitian	4
1.5 Manfaat Penelitian	4
1.6 Sistematika Penulisan	4
BAB II TINJAUAN PUSTAKA DAN DASAR TEORI	6
2.1 Tinjauan Pustaka	6
2.2 Dasar Teori.....	7
2.2.1 <i>NLP (Natural Language Processing)</i>	7
2.2.2 Sentimen Analisis	7
2.2.3 <i>Text Mining</i>	8
2.2.4 Instagram.....	8
2.2.5 <i>Transformers</i>	10
2.2.6 <i>Bidirectional Encoder Representations from Transformers (BERT)</i>	11
2.2.7 IndoBERT	12
2.2.8 <i>Confusion Metrics</i>	12
BAB III METODE PENELITIAN	15
3. 1 Alur Penelitian	15
3. 2 Pengumpulan Data	16
3. 3 Pelabelan kelas sentimen	16

3. 4 <i>Preprocessing</i>	17
3. 5 Split Data	17
3. 6 Klasifikasi IndoBERT	17
3. 7 Evaluasi.....	18
BAB IV HASIL DAN ANALISIS PENELITIAN	19
4.1 Hasil Penelitian	19
4.1.1 Deskripsi Dataset	19
4.1.2 <i>Data Preprocessing</i>	19
4.1.3 Eksperimen Model	21
4.2 Hasil Perancangan Model	22
4.3 Hasil Implementasi <i>Streamlit</i>	27
BAB V KESIMPULAN DAN SARAN	31
5.1 Kesimpulan	31
5.2 Saran	31
DAFTAR PUSTAKA	



DAFTAR GAMBAR

Gambar 2. 1 Model Arsitektur Transformers (Vaswani, 2017).....	10
Gambar 2. 2 Pre-training dan Fine-Tuning.....	11
Gambar 3. 1 Flowchart Alur Penelitian	15
Gambar 3. 2 Sampel Scraping Data	16
Gambar 4. 1 Informasi dataset	19
Gambar 4. 2 Hasil Labeling dataset	20
Gambar 4. 3 Hasil data cleaning	21
Gambar 4. 4 Menunjukkan hasil eksperimen dataset inbalance dengan iterasi epoch 5	22
Gambar 4. 5 Performa Model	23
Gambar 4. 6 Diagram Confusion Matrix	23
Gambar 4. 7 Jumlah data hasil balancing	25
Gambar 4. 8 menunjukkan hasil eksperimen dataset balancing dengan iterasi epoch 5	25
Gambar 4. 9 Performa Model	26
Gambar 4. 10 Diagram Confusion Matrix	27
Gambar 4. 11 Tampilan Awal Streamlit	28
Gambar 4. 12 Inputan Text Komentar	28
Gambar 4. 13 Tampilan Tombol klik dan Hasil Analisis Sentimen	29
Gambar 4. 14 Hasil Sentimen Positif	29
Gambar 4. 15 Hasil Sentimen Netral	30

DAFTAR TABEL

Tabel 4. 1 Perbandingan sebelum dan sesudah data cleaning	21
Tabel 4. 2 Hyperparameter	22



ABSTRAK

Media sosial, khususnya Instagram, telah menjadi wadah utama bagi masyarakat untuk mengekspresikan opini, termasuk dalam merespons fenomena budaya populer seperti *Korean Wave* dan Kpop. Baru-baru ini, keputusan grup Kpop NCT untuk berkolaborasi dengan Starbucks memicu beragam reaksi dari penggemar, terutama di tengah gerakan boikot terhadap brand tersebut. Penelitian ini bertujuan untuk menganalisis sentimen komentar penggemar NCT terkait kolaborasi ini menggunakan model *Indonesian Bidirectional Encoder Representations from Transformers (IndoBERT)*. Eksperimen dilakukan pada data yang tidak seimbang dan data yang telah diseimbangkan menggunakan metode *Synthetic Minority Over-sampling Technique (SMOTE)*. Hasil penelitian menunjukkan bahwa data yang tidak seimbang menyebabkan model mengalami *overfitting* dengan akurasi sebesar 83%, *precision* 84%, *recall* 83%, dan *f1-score* 83%. Setelah dilakukan *balancing* data menggunakan SMOTE, performa model meningkat dengan *akurasi*, *precision*, *recall*, dan *f1-score* sebesar 86%. Hal ini membuktikan bahwa *balancing* data berperan penting dalam meningkatkan performa model dalam analisis sentimen.

Kata Kunci : analisis sentimen, IndoBERT, SMOTE, Kpop, Boikot, Instagram

ABSTRACT

Social media, especially Instagram, has become a primary platform for people to express their opinions, including in responding to popular culture phenomena such as the Korean Wave and Kpop. Recently, Kpop group NCT's decision to collaborate with Starbucks sparked various reactions from fans, especially amidst the boycott movement against the brand. This study aims to analyze the sentiment of NCT fans' comments regarding this collaboration using the *Indonesian Bidirectional Encoder Representations from Transformers (IndoBERT)* model. Experiments were conducted on unbalanced data and balanced data using the *Synthetic Minority Over-sampling Technique (SMOTE)* method. The results showed that unbalanced data caused the model to overfit with an accuracy of 83%, precision of 84%, recall of 83%, and *f1-score* of 83%. After data balancing using SMOTE, the model's performance increased with an accuracy, precision, recall, and *f1-score* of 86%. This proves that data balancing plays an important role in improving model performance in sentiment analysis.

Keywords : sentiment analysis, IndoBERT, SMOTE, Kpop, Boycott, Instagram

BAB I

PENDAHULUAN

1.1 Latar Belakang

Berdasarkan penelitian yang telah dilakukan dibidang sentimen analisis pada e-marketplace yaitu oleh (Kosasih dan Alberto, 2021) yang melakukan sentimen analisis pada ulasan produk mainan di shopee. Menggunakan data sebanyak 1000, yang dibagi menjadi 700 data latih dan 300 data uji dengan kategori positif, negatif dan netral. Kemudian memperoleh akurasi sebesar 79,3333%, dimana akurasi tersebut diperoleh dengan cara pembobotan data menggunakan TF-IDF dan mengklasifikasi data uji menggunakan metode *K-Nearest Neighbor*(K-NN).

Media sosial saat ini telah menjadi bagian fundamental dalam kehidupan sehari-hari, tidak hanya untuk keperluan pribadi tetapi juga bisa digunakan untuk berbisnis, serta berbagai aktivitas lainnya. Media sosial ini menghasilkan banyak data termasuk gambar, komentar teks atau emoticon, video dan lain-lain, sehingga memungkinkan masyarakat untuk bebas beropini. Dengan adanya analisis sentimen terhadap opini yang berkembang di media sosial, dapat menghasilkan data dan informasi yang bermanfaat (Permatasari, Linawati dan Jasa, 2021).

Instagram adalah salah satu platform media sosial yang memungkinkan penggunaanya untuk mengirim dan menerima informasi dalam bentuk gambar, video dan cerita dalam berbagai format, termasuk unggaham feed, Instagram stories, reels dan siaran langsung (Instagram live). Pada awal peluncuran, Instagram hanya memiliki pengguna yang berjumlah 100 ribu orang. Dalam waktu sekitar 2,5 bulan, jumlah pengguna Instagram meningkat pesat menjadi 1 juta pengguna. Instagram telah menjadi salah satu platform media sosial terbesar di dunia dengan lebih dari 2 miliar pengguna aktif bulanan. Cepatnya penyebaran informasi melalui Instagram dapat memfasilitasi pertukaran budaya, seperti budaya korea yang dikenal sebagai Korean Wave atau Hallyu, yang mencakup drama Korea dan musik Kpop. Budaya Korea ini mulai dikenal

di Indonesia sejak tahun 2002 dan terus berkembang hingga sekarang, menarik minat remaja maupun orang dewasa. Kpop sebagai salah satu aspek budaya Korea yang paling digemari, telah menjadi kunci kesuksesan dari Korean Wave (Rizkina dan Hasan, 2023).

Namun belakangan ini salah satu *agency* dari boygrup Kpop menjadi sorotan netizen di Indonesia dikarenakan kerjasamanya dengan brand coffee shop yang telah diboikot oleh beberapa negara. Berawal dari ketegangan yang terjadi di Timur Tengah, Gerakan *Free Palestine* mewakili perjuangan yang kompleks untuk kemerdekaan negara Palestina dan hak asasi manusia. Seruan boikot terhadap merek dagang starbucks telah ramai diserukan sejak tahun lalu. Dampak dari boikot ini sudah dirasakan secara global oleh starbucks, yang mengumumkan penurunan saham sebesar 17% di awal tahun 2024. Ditengah gerakan boikot ini, keputusan SM dan NCT untuk berkolaborasi dengan Starbucks menimbulkan beragam tanggapan dari para penggemar, yang menyebabkan penurunan jumlah pengikut media sosial NCT sebanyak 500 ribu (Melati, 2024).

Penelitian ini akan menganalisis bagaimana *BERT* mengklasifikasikan sentimen positif, negatif dan netral pada respon penggemar dari topik NCT dan Starbucks berdasarkan komentar dari media sosial Instagram. Bert merupakan model pembelajaran mendalam yang telah menunjukkan hasil yang luar biasa dalam berbagai tugas NLP. Karena memiliki enam lapisan transformer berlapis di atas setiap *encoder* dan *decoder*, file proses pelatihan sangat kompleks, konfigurasinya sangat sulit, waktu pelatihan sangat lama dan biayanya sangat tinggi. Oleh karena itu, dari hasil penelitian ini diharapkan dapat mengetahui tingkat akurasi pada analisis sentimen menggunakan metode *BERT*.

1.2 Perumusan Masalah

Berdasarkan latar belakang tersebut, penulis mengidentifikasi permasalahan yang akan dibahas yaitu :

1. Mengetahui apa pola sentimen yang muncul di kalangan penggemar NCT terhadap *Endorsement* dengan Starbucks.

2. Mengetahui seberapa akurat metode klasifikasi yang ada saat mengklasifikasi komentar penggemar pada topik diatas.

1.3 Batasan Masalah

Batasan masalah dari penulisan proposal ini adalah sebagai berikut :

1. Data yang digunakan hanya diambil dari media sosial Instagram.
2. Data yang dikumpulkan hanya pada periode setelah pengumuman kerjasama antara boygrup dengan brand minuman tersebut.
3. Penelitian ini menggunakan metode *IndoBERT* sebagai alat utama untuk analisis sentimen.
4. Data hanya diambil dari penggemar boygrup NCT

1.4 Tujuan Penelitian

Tujuan dari penelitian ini adalah untuk menganalisis sentimen terhadap komentar penggemar NCT terhadap kerjasamanya dengan Starbucks dengan menggunakan metode *BERT* berdasarkan data dari media sosial instagram.

1.5 Manfaat Penelitian

Manfaat dari penelitian ini yaitu untuk membantu mengidentifikasi perasaan dan pendapat penggemar terhadap kerjasama boygrup tersebut dengan produk minuman yang masuk daftar hitam.

1.6 Sistematika Penulisan

Dalam penyusunan laporan tugas akhir ini menggunakan sistematika penulisan sebagai berikut :

BAB 1 : PENDAHULUAN

Bab ini penulis menjelaskan latar belakang pemilihan judul, perumusan masalah, batasan masalah, tujuan penelitian, manfaat penelitian, dan sistematika penulisan yang digunakan dalam penyusunan laporan tugas akhir.

BAB 2 : TINJAUAN PUSTAKA DAN DASAR TEORI

Pada bab ini mencakup penelitian-penelitian sebelumnya dan dasar teori yang berguna untuk membantu memahami teori yang berhubungan dengan sentimen analisis dan metode Bert.

BAB 3 : METODE PENELITIAN

Bab ini penulis memaparkan proses tahapan-tahapan penelitian yang diawali dari mendapat data hingga klasifikasi data.

BAB 4 : HASIL DAN ANALISIS PENELITIAN

Selanjutnya pada bab 4 ini mengungkapkan hasil dari penelitian yaitu hasil identifikasi sentimen positif, negatif atau netral dari teks komentar menggunakan metode Bert.

BAB 5 : KESIMPULAN DAN SARAN

Pada bab ini berisikan kesimpulan dan saran dari proses penelitian mulai dari penelitian dari awal hingga akhir.



BAB II

TINJAUAN PUSTAKA DAN DASAR TEORI

2.1 Tinjauan Pustaka

Sentimen analisis adalah ungkapan perasaan yang terkandung dalam kalimat melalui suatu proses dengan melibatkan pengolahan data teks, seperti ulasan. Tujuan dari proses tersebut yaitu untuk menentukan apakah ulasan yang terkandung dalam sentimen bersifat positif, negatif atau netral (Amardita dkk., 2022). Selanjutnya, dengan menganalisa teks tersebut, dapat ditentukan apakah ulasan dari subjek yang dibahas memiliki nilai kecenderungan positif, negatif atau netral.

Salah satu penelitian yang telah dilakukan dibidang sentimen analisis pada e- marketplace yaitu oleh (Kosasih dan Alberto, 2021) yang melakukan sentimen analisis pada ulasan produk mainan di shopee. Menggunakan data sebanyak 1000, yang dibagi menjadi 700 data latih dan 300 data uji dengan kategori positif, negatif dan netral. Kemudian memperoleh akurasi sebesar 79,3333%, dimana akurasi tersebut diperoleh dengan cara pembobotan data menggunakan TF-IDF dan mengklasifikasi data uji menggunakan metode *K-Nearest Neighbor*(K-NN).

Penelitian selanjutnya yaitu tentang penerapan analisis sentimen pada pengguna twitter yang dilakukan oleh (Deviyanto dan Wahyudi, 2018) dengan jumlah data set sebanyak 2000 data berbahasa Indonesia yang terbagi menjadi dua kategori positif dan negatif. Tahapan selanjutnya yaitu melakukan klasifikasi data dengan menggunakan metode *K-Nearest Neighbor*(K-NN) dan memperoleh nilai akurasi terbesar yaitu 67,2%.

Berdasarkan penelitian sebelumnya yaitu menggunakan penerapan sentimen analisis dan metode klasifikasi *K-Nearest Neighbor*(K-NN), sehingga pada penelitian ini akan dibuat penelitian yang menggunakan teknik sentimen analisis dengan kategori positif, negatif dan netral dengan menerapkan model IndoBERT.

2.2 Dasar Teori

2.2.1 NLP (*Natural Language Processing*)

NLP (*Natural Language Processing*) adalah salah satu cabang ilmu Kecerdasan buatan (AI) yang berfokus pada pengolahan bahasa natural. Bahasa natural adalah Bahasa yang biasa digunakan manusia untuk berbicara satu sama lain (Rohman, Utami dan Raharjo, 2019).

Dalam bidang *Natural Language Processing*, terdapat berbagai topik yang dibahas, salah satunya adalah *Tokenization*. *Tokenization* yaitu memecahkan kata-kata menjadi unit lebih kecil sebelum diproses. *Natural Language Processing* dalam pemrosesan teks dan analisis sentimen dapat membantu seseorang untuk menuliskan kata-kata yang tepat dan bahkan menilai sentimen yang muncul dari kalimat tersebut.

2.2.2 Sentimen Analisis

Analisis sentimen adalah teknik untuk menganalisis pikiran, perasaan, dan penilaian pengguna. Menurut beberapa sudut pandang, itu digunakan untuk menentukan apa yang diyakini pengguna berdasarkan informasi seperti opini tertulis. Karena perkembangan era digital yang tidak dapat dihindari, orang lebih sering mengungkapkan dan mengunggah ide-ide mereka di media sosial (Chinnasamy dkk., 2022). Proses sentimen analisis terdiri dari dua bagian yaitu *Sentimen Extraction* yang telah dievaluasi dan dianalisis, dan *Sentimen Classification* adalah aspek yang bermakna positif, negatif atau netral. Selanjutnya, proses penentuan perspektif mengenai masing-masing aspek tersebut (Mubaraq dan Maharani, 2022).

Berdasarkan penjelasan diatas, analisis sentimen adalah metode untuk mengklasifikasikan komentar atau pendapat yang diungkapkan dalam teks di media digital. Analisis sentimen banyak digunakan untuk mengumpulkan data tentang pendapat atau respon dari konsumen. Sejumlah literature penting lainnya membahas analisis perasaan dengan penekanan pada penggunaan tertentu untuk mengkategorikan pendapat positif, negatif dan netral (Syah dan Witanti, 2022) (Isnain, Marga dan Alita, 2021).

2.2.3 Text Mining

Text mining menambang data dalam bentuk teks, dengan sumber data biasanya berasal dari data, *text mining* bertujuan untuk menemukan kata-kata yang dapat menunjukkan isi data, sehingga dapat dilakukan analisis hubungan antar data. (Analisis Pada Twitter CommuterLine dkk., 2022)

Text mining menggunakan kumpulan teks yang berformat tidak beraturan, terstruktur atau minimal semi-terstruktur untuk mendapatkan informasi berguna dari sekumpulan data. *Text mining* melakukan dua tugas khusus yaitu mengkategorikan teks (*text categorization*) dan pengelompokan teks (*text clustering*). Proses *text mining* menggunakan konsep dan teknik data mining untuk menemukan pola dalam teks, khususnya dalam proses menganalisis teks, untuk menemukan informasi yang berguna untuk tujuan tertentu. Proses ini memerlukan beberapa langkah awal untuk mempersiapkan dasar agar teks dapat dianalisis (Apri Wenando, 2023). *Text mining* dan *data mining* berbeda karena preprocessingnya. Data mining berkonsentrasi pada penomoran (*indexing*) dan normalisasi data, sedangkan *text mining* berkonsentrasi pada identifikasi dan ekstraksi fitur (Isnain dkk., 2021).

2.2.4 Instagram

Instagram merupakan salah satu media sosial yang paling cepat berkembang saat ini karena memungkinkan pengguna untuk berinteraksi dengan orang lain kapan saja dan dimana saja dari komputer maupun perangkat selular. Instagram cukup efektif sebagai media promosi dalam memberikan rangsangan melalui konten-konten yang diunggah untuk meningkatkan pengetahuan konsumen terkait produk yang ditawarkan (Ramadan dan Fatchiya, 2021).

Instagram telah menjadi salah satu platform media sosial terbesar di dunia, dengan lebih dari 2 miliar pengguna aktif per bulan. Karena fiturnya yang luas dan pengguna yang besar, serta kemampuan untuk berbagi pendapat tentang isu-isu terhangat dan kontroversial, instagram telah menjadi pilihan yang populer dibandingkan dengan platform sosial media lainnya

(Pilar dkk., 2023). Salah satu fitur utama Instagram adalah kemampuan untuk mengupload foto atau video di feed atau cerita serta mengirimkan pesan kepada pengguna lain. Berikut adalah fitur media sosial Instagram lainnya :

a. Feed Instagram

Fitur ini merupakan fitur utama pengguna dapat melihat unggahan foto dan video dari akun yang mereka ikuti. Algoritma Instagram menentukan urutan tampilan postingan berdasarkan interaksi pengguna (suka, komentar, bagikan, waktu yang dihabiskan untuk melihat postingan, dan lain-lain).

b. Story Instagram

Fitur ini hanya tersedia selama 24 jam setelah mengunggah dan kemudian hilang secara otomatis. Dilengkapi dengan fitur tambahan seperti filter, GIF, stiker, polling, Tanya jawab, musik dan banyak lagi untuk meningkatkan keterlibatan dengan pengikut anda.

c. Reels Instagram

Fitur video pendek yang mirip tiktok, berdurasi antara 15 sampai 90 detik, dilengkapi berbagai alat pengeditan, efek dan musik untuk membuat konten lebih menarik.

d. Live Instagram

Fitur siaran langsung yang memungkinkan pengguna berkomunikasi dengan pengikutnya secara real time. Dapat digunakan untuk berbagai tujuan, termasuk pertanyaan dan jawaban, forum diskusi, promosi produk, dan banyak lagi.

e. Direct Message (DM)

Fitur pesan pribadi yang memungkinkan pengguna mengirim teks, gambar, video, dan melakukan panggilan audio dan video.

f. Explore Page

Halaman jelajah yang menampilkan konten yang mungkin menarik bagi pengguna lain berdasarkan aktivitas mereka. Instagram menggunakan AI dan algoritma untuk menyesuaikan konten yang Anda lihat di halaman ini agar sesuai dengan minat Anda.

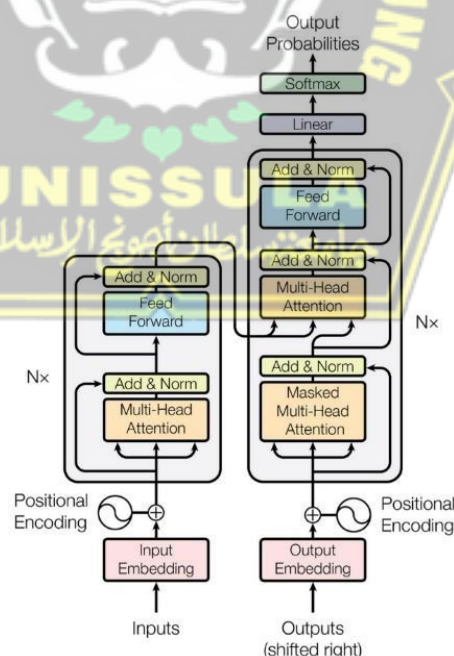
g. Instagram Shop

Fitur yang memungkinkan akun bisnis menjual produk langsung di Instagram melalui postingan, cerita, dan tab khusus. Terintegrasi dengan Facebook Shop untuk pengalaman berbelanja yang lebih baik.

2.2.5 Transformers

Transformers adalah model deep learning yang menggunakan mekanisme *self-attention*, yang mengidentifikasi hubungan antar kata dalam konteks. *Transformers* digunakan terutama dalam pemrosesan Bahasa alami dan *computer vision*. Model ini dimaksudkan untuk memproses tipe data sekuens seperti menerjemahkan bahasa dan merangkum teks. *Transformer* terdiri dari dua *stacks*.(Vaswani, 2017)

Tranformator menggunakan *self-attention Mechanism* untuk mempelajari hubungan antar kata dalam kalimat. Mekanisme ini memungkinkan model untuk memahami bagaimana kata-kata saling terkait satu sama lain, serta bagaimana maknanya berubah-ubah sesuai dengan konteks kalimat. Untuk arsitektur transformer dapat dilihat pada gambar 2.1 ;

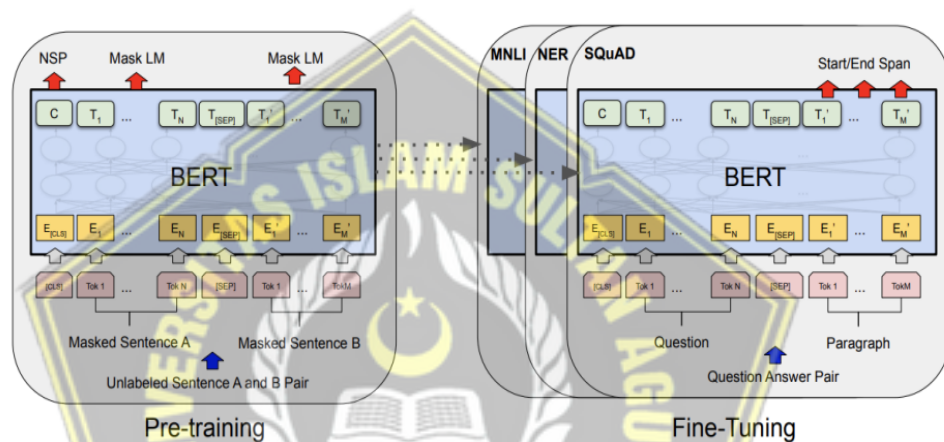


Gambar 2. 1 Model Arsitektur *Transformers* (Vaswani, 2017)

Pada gambar 2.1 menunjukkan arsitektur *transformers* yang terdiri dari dua bagian utama yaitu *encoder* dan *decoder*.

2.2.6 Bidirectional Encoder Representations from Transformers (BERT)

BERT (*Bidirectional Encoder Representations from Transformers*) adalah sebuah model bahasa yang revolusioner dalam bidang pemrosesan Bahasa alami atau biasa disebut NLP (*Natural Language Processing*). Bert memiliki transformer dengan 6 lapisan yang ditumpangkan di atas encoder dan decoder masing-masing, ini menyebabkan file proses pelatihan yang sangat kompleks, konfigurasi tinggi, dan waktu pelatihan yang sangat lama dan biaya yang mahal (Mas dkk., 2021).



Gambar 2.2 Pre-training dan Fine-Tuning

Gambar 2.2 menunjukkan proses *pre-training* dan *fine-tuning* BERT secara keseluruhan menggunakan arsitektur yang sama, terlepas dari lapisan *output* parameter *pre-training* yang sama digunakan untuk menginisialisasi model untuk berbagai model untuk berbagai tugas down-stream. Semua parameter disetel dengan baik selama proses *fine-tuning* (Mubaraq dan Maharani, 2022).

Pada tahap *pre-training* BERT, model dilatih untuk memprediksi kata yang diacak dalam konteks kalimat melalui korpus teks yang sangat besar. Setelah tahap *training*, model BERT yang sudah dilatih dapat di *fine-tune* untuk tugas tertentu. Selama tahap *fine-tuning*, beberapa lapisan di atas model BERT dapat disesuaikan untuk tugas-tugas tertentu yang digunakan dalam penelitian ini untuk menjawab pertanyaan. Ini memungkinkan BERT untuk

menyesuaikan diri dengan tugas tertentu dan menghasilkan hasil yang lebih baik (Devlin dkk., 2019).

2.2.7 IndoBERT

IndoBERT (*Indonesian Bidirectional Encoder Representations from Transformers*) merupakan model pemrosesan bahasa alami (NLP) yang dikembangkan khusus untuk bahasa Indonesia. Secara total, *transformers* IndoBERT dilatih dari 220 juta kata dalam Bahasa Indonesia. Dimana data tersebut berasal dari tiga sumber utama yaitu Wikipedia Indonesia dengan 74 juta kata, artikel berita dari platform seperti Kompas, Tempo, dan Liputan 6 dengan 55 juta kata, serta Web Corpus Indonesia dengan 90 juta kata (Wijaya dkk., 2023). IndoBERT dirancang untuk meningkatkan kinerja berbagai tugas NLP dalam bahasa Indonesia, seperti analisis sentimen, ekstraksi informasi, klasifikasi teks, dan pemrosesan tanya jawab.

IndoBERT adalah ide yang dikembangkan oleh tim IndoNLU. IndoBERT adalah versi modifikasi dari model BERT dasar yang digunakan untuk menganalisis bahasa Indonesia. Model ini juga merupakan model yang belakangan ini terkenal dikarenakan memiliki 4 miliar korpus kata yang dilatih (Mubaraq & Maharani, 2022). IndoBERT memiliki arsitektur yang sama dengan BERT. Satu-satunya perbedaan adalah dataset yang digunakan untuk pelatihan. Dataset IndoBERT dikenal sebagai Indo4B. Dataset ini berisi kalimat formal dan sehari-hari dalam bahasa Indonesia yang dikumpulkan dari 15 sumber data. Dua dari sumber data yang berfokus pada bahasa sehari-hari bahasa Indonesia, delapan sumber lainnya yang menekan pada bahasa formal, dan sisanya mengandung gaya campuran bahasa sehari-hari dan formal (Dharmawan dkk., 2023).

2.2.8 Confusion Metrics

Confusion metrics merupakan termasuk pada tahap evaluasi pada *performance vector* dari *machine learning* berupa tingkat akurasi, presisi, dan *recall* pada kata-kata yang tidak bernilai atau mengurangi tingkat performansi pada *performance vector*.

Rumus *confusion metrics* yang digunakan untuk menghitung *accuracy*, *precision*, *recall* dan *f1-score* sebagai berikut.

a. *Accuracy*

Merupakan rasio prediksi benar (“positif” dan “negatif”) dengan keseluruhan data. Akurasi akan menjawab pertanyaan berapa persen positif dan negatif yang disampaikan melalui instagram terhadap akun Instagram tersebut.

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

b. *Precision*

Merupakan rasio prediksi benar positif dibandingkan dengan keseluruhan hasil yang diprediksi positif. *Precision* akan menjawab pertanyaan “berapa persen komentar “positif” yang disampaikan masyarakat terhadap instagram.

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

c. *Recall (Sensitivity)*

Merupakan rasio prediksi benar positif dibandingkan dengan keseluruhan data yang benar positif. *Recall* akan menjawab pertanyaan “berapa persen komentar “positif” yang disampaikan masyarakat terhadap Instagram.

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

d. *F1-Score*

Merupakan ukuran rata-rata dari recall dan precision, memberikan gambaran seimbang antara kedua metric tersebut.

$$F1 - score = 2 \frac{(recall \times precision)}{(recall + precision)} \quad (4)$$

Keterangan :

TP (*True Positive*) = jumlah dokumen dari kelas 0 yang benar diklasifikasikan sebagai kelas 0.

TN (*True Negative*) = jumlah dokumen dari kelas 1 yang benar diklasifikasikan sebagai kelas 1.

FP (*False Positive*) = jumlah dokumen dari kelas 1 yang benar diklasifikasikan sebagai kelas 0.

FN (*False Negative*) = jumlah dokumen dari kelas 0 yang benar diklasifikasikan sebagai kelas 1.

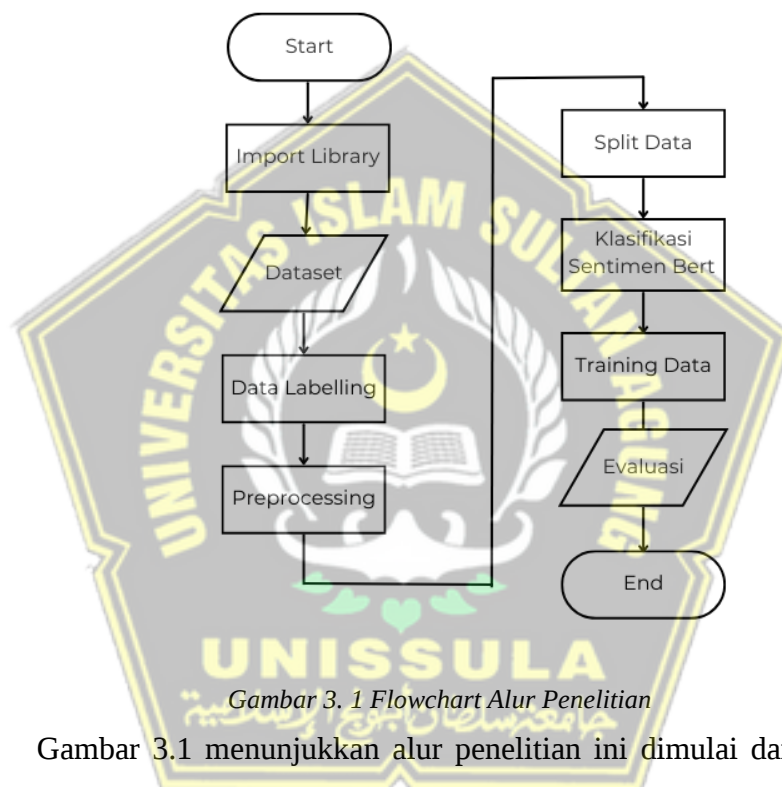


BAB III

METODE PENELITIAN

3. 1 Alur Penelitian

Dalam penelitian ini sistem yang digunakan yaitu sistem sentimen analisis terhadap penggemar NCT mengenai endorsement dengan Starbucks yang masuk daftar hitam menggunakan model BERT.



Gambar 3. 1 Flowchart Alur Penelitian

Gambar 3.1 menunjukkan alur penelitian ini dimulai dari mengimport *library* dan memuat dataset, penulis akan memuat dataset yang didapat dari proses scraping data pada komentar postingan instagram. Setelah itu, dataset diberi pelabelan positif, negatif dan netral yang dilambangkan dengan [0, 1, 2]. Setelah itu, dataset masuk ke tahap pre-processing. Pada tahap ini menyiapkan data yang sebelumnya tidak terstruktur menjadi data yang lebih terstruktur melalui proses *data cleaning* dan *tokenizing*. Lalu, data yang dihasilkan dari tahap ini dibagi menjadi data latih dan data uji. Dataset kemudian dilatih dengan menggunakan model BERT setelah melewati tahapan tersebut. Selanjutnya, data diklasifikasikan menjadi tiga sentimen

untuk menentukan ke dalam kategori positif, negatif, atau netral. Setelah melewati tahap pelatihan dan klasifikasi, data kemudian dievaluasi untuk memastikan keakuratannya.

3. 2 Pengumpulan Data

Data yang digunakan dalam penelitian ini didapatkan melalui *web scraping* pada kolom komentar postingan kerjasama antara NCT dan starbucks di instagram. Data tersebut berisi username, komentar dan like Penggemar NCT mengenai *endorsement* nya dengan starbucks yang masuk daftar hitam. Tujuan dari *web scraping* ialah untuk mendapatkan data kemudian melakukan ekstraksi informasi yang dimiliki data tersebut.

	A	B	C
1	username	comment	likes
2	2813ceminnn	udh pernah nyoba ga enak amis rasa darah	7 suka
3	ilysm.vr	Percayalah enakkan nutrisari jeruk peras 🤔🤔	26 suka
4	nunuyy79	BOIKOT MEREKA	182 suka
5	kezzaa12	Kenapa harus nct, kenapa harus nct sih...?	14 suka
6	ffiasha	ini udah selesai kan plis collabnya??	6 suka
7	adelliarisma5	Hihhh gara gara collab collab nihh gue harus ngejauh in idol guee 🤔	552 suka
8	sfyamdadaa	jgn di beli ya guys	212 suka
9	estt_asya	Siapa yang mau beli anji...	1 suka
10	stifanyaml	enakan jg kopikap	64 suka
11	na4yale	oh ini kah yang jualan tapi bayar nya pake nyawa orang?	634 suka
12	thv4u1l	GA NCT GA LAKU? 🤔🤔	2 suka
13	minnnjae.03	ikan sepat ikan tongkol, minuman lu ga enak	268 suka
14	abelaa_x	I bagus beli pop ice lagi terus murah ga kekk strabak udah mahal bauk amis pulak lagi! iuuu	1 suka
15	mo0nyss	ga bakal ada yang beli udahh	13 suka

Gambar 3. 2 Sampel *Scraping Data*

Hasil dari *web scraping* yang diambil dari kolom *Username*, *Comment*, *likes*, sesuai dengan yang tertera pada gambar 3.2.

3. 3 Pelabelan kelas sentimen

Pelabelan kelas sentimen data komentar penggemar dikelompokkan menjadi tiga kelas yaitu positif, negatif dan netral. Berdasarkan pada kelas negatif bermakna ejekan, keluhan dan penolakan terhadap keputusan yang diambil oleh pihak NCT dan *agency*, sedangkan kelas netral teks yang bermakna tidak menyatakan perasaan atau penilaian emosional positif dan negatif, lalu yang terakhir kelas positif bermakna pujian, dukungan dan menyatakan setuju.

3. 4 Preprocessing

Tahapan selanjutnya yaitu perlu melakukan pembersihan data dengan tujuan supaya data dapat digunakan pada tahap selanjutnya. Adapun tahapan yang dilakukan untuk penelitian yaitu sebagai berikut :

a. *Data Cleaning*

Pada tahap ini, data cleaning digunakan untuk menghilangkan tanda baca, huruf besar, angka, symbol, URL, Username, hashtag, spasi berlebih dan pengulangan karakter pada dataset.

b. *Tokenizing*

Pada tahap ini merupakan proses untuk memecah kalimat menjadi per kata sehingga lebih mudah untuk diolah secara terpisah.

3. 5 Split Data

Setelah tahap tokenisasi, dataset dibagi menjadi 30% data uji dan 70% data latih. Training/latihan membantu model mengenali pola dalam data. Sedangkan testing/pengujian memastikan bahwa model yang telah dilatih mampu memprediksi label yang belum dipelajari dengan baik.

3. 6 Klasifikasi IndoBERT

Sistem pada laporan tugas akhir ini menerapkan metode sentimen analisis *Indonesian Bidirectional Encoder Representations from Transformers* (IndoBERT). Sistem ini akan mengidentifikasi sentimen positif, negatif atau netral dari komentar penggemar NCT pada instagram dengan menggunakan model IndoBERT-base-p1. Model tersebut merupakan model dari IndoBERT dengan arsitekturnya yaitu Base yang memiliki 12 layer encoder, 768 hidden nodes dan 124 juta parameter. *Library* yang digunakan pada tugas akhir ini yaitu *Transformer* yang disediakan oleh HuggingFace. Sistem pada tugas akhir ini nantinya akan berbasis *Streamlit*.

3. 7 Evaluasi

Setelah proses train model selesai, dataset masuk ke tahap evaluasi. Proses evaluasi yang sangat penting dilakukan untuk menentukan kualitas program berdasarkan *accuracy*, *precision*, *recall* dan *f1-score*. yang dihasilkan dengan menggunakan *confusion metrics*. Kemudian simpan model untuk digunakan nantinya.



HASIL DAN ANALISIS PENELITIAN

4.1 Hasil Penelitian

4.1.1 Deskripsi Dataset

Pada penelitian ini, data diambil dari komentar pada aplikasi Instagram menggunakan *Web scraping*, teknik ini digunakan untuk mengambil informasi dari berbagai situs web. File didownload dalam bentuk *excel*, selanjutnya diberi label secara manual dengan membaca kalimat atau komentar satu persatu, data yang diambil 1084 dikategorikan positif 70, negatif 808 dan netral 206.

	Comment	label
0	udh pernah nyoba ga enak amis rasa darah	negatif
1	Percayalah enakkan nutrisari jeruk peras 🙄🙄	negatif
2	BOIKOT MEREKA	negatif
3	Kenapa harus nct,kenapa harus nct sihhh...	netral
4	ini udah selesai kan plis collabnya??	netral
...	...	...
1079	ga sabarrrr 🤔💖	positif
1080	ini bakalan ada diindo ga kak	netral
1081	Kecewa bgt:)))	netral
1082	Wkwkwkwkwkwkwkwkwkwkw 🤔🤔	netral
1083	Delete our comments? 🙄🙄🙄🙄	netral

[1084 rows x 2 columns]

Gambar 4. 1 Informasi *dataset*

Gambar 4.1 memperlihatkan sampel *dataset* untuk memberikan wawasan mengenai struktur dan fitur *dataset* yang digunakan dalam pelatihan model pada penelitian ini. Dataset yang diambil terdiri dari kolom Comment dan label dengan data yang berjumlah 1084 sesuai dengan kode diatas.

4.1.2 Data Preprocessing

4.1.2.1 Labeling

Pada proses *labeling*, dilakukan secara manual dengan membaca kalimat atau komentar satu persatu. Kategori pelabelan meliputi kalimat positif, negatif dan netral. Setelah mengganti nilai, mengkonversi tipe data

pada kolom “label” menjadi integer. Sehingga menjadi positif (0), negatif (1) dan netral (2)

	Comment	label
0	udh pernah nyoba ga enak amis rasa darah	1
1	Percayalah enakkan nutrisari jeruk peras 🙄🙄	1
2	BOIKOT MEREKA	1
3	Kenapa harus nct,kenapa harus nct sih..	2
4	ini udah selesai kan plis collabnya??	2
...	...	...
1079	ga sabarrrr 🙄🙄	0
1080	ini bakalan ada diindo ga kak	2
1081	Kecewa bgt:)))	2
1082	Wkwkwkwkwkwkwkwkw 🙄🙄	2
1083	Delete our comments? 🙄🙄🙄🙄	2

[1084 rows x 2 columns]

Gambar 4. 2 Hasil *Labeling dataset*

Gambar 4.2 merupakan hasil dari perubahan pada kolom “label” dari tipe data string menjadi tipe data integer.

4.1.2.2 Data Cleaning

Proses *data cleaning* yang dilakukan adalah membersihkan data yang akan dijadikan input model yaitu data pada kolom komentar dengan menghapus karakter selain huruf, angka serta spasi. Selain itu, diterapkan juga *lowercase* untuk mengkonversi semua huruf dalam setiap baris di kolom komentar menjadi huruf kecil. Pembersihan data ini berfungsi untuk memberikan data yang konsisten untuk pelatihan model, sehingga model dapat memahami pola teks dengan baik serta dapat meningkatkan akurasi dan kinerjanya. Pada proses ini, memberikan hasil pada *dataset* menjadi seperti gambar

Tabel 4. 2 Hyperparameter

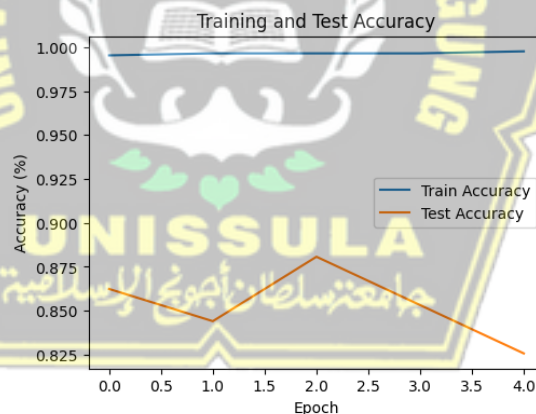
No	Hyperparameter	Ukuran
1.	Batch Size	32
2.	Learning Rate	5e-5
3.	Epoch	5
4.	Max_length	256

Berdasarkan tabel 4.2 *hyperparameter* yang digunakan menghasilkan perbandingan antara dataset yang mengalami inbalance dan data yang sudah mengalami proses balancing.

4.2 Hasil Perancangan Model

4.2.1 Data inbalance

Hasil eksperimen pada dataset inbalance yang telah dilakukan dengan menggunakan *hyperparameter batch size 32, learning rate 5e-5, epoch 5, serta max length 256* dapat dilihat pada gambar 4.4.



Gambar 4. 4 Menunjukkan hasil eksperimen dataset inbalance dengan iterasi epoch 5

Gambar 4.4 menunjukkan bahwa selama proses pelatihan, akurasi training tetap sangat tinggi, mendekati 100%, tanpa perubahan yang signifikan. Hal ini mengindikasikan bahwa model sangat baik dalam menyesuaikan diri dengan data training. Namun, akurasi testing justru mengalami overfitting dan menurun setelah awal pelatihan. Pada grafik epoch 0.0, akurasi testing berada di sekitar 86%, kemudian sempat meningkat pada epoch 2.0, tetapi menurun kembali hingga sekitar 82% pada epoch 4.0.

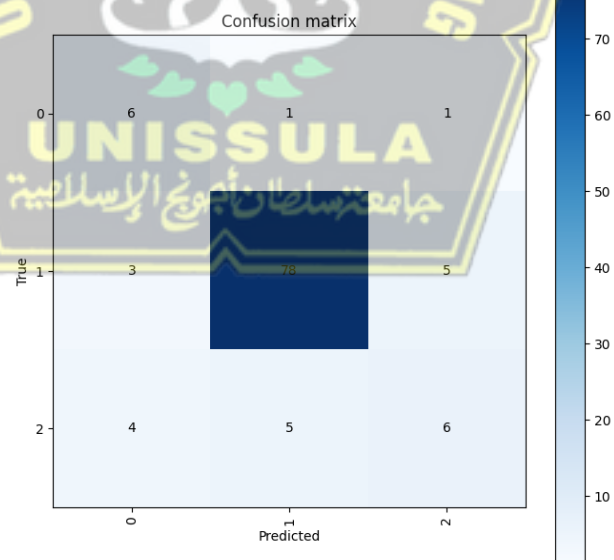
Setelah proses pelatihan dilakukan, tahap selanjutnya yaitu proses menguji performa model. Hasil dari implementasi model terhadap data training dan data testing menunjukkan hasil *accuracy*, *precision*, *recall* dan *f1-score* dari dataset inbalance yang diperoleh ditunjukkan pada gambar 4.5.

Classification Report:

	precision	recall	f1-score	support
0	0.46	0.75	0.57	8
1	0.93	0.91	0.92	86
2	0.50	0.40	0.44	15
accuracy			0.83	109
macro avg	0.63	0.69	0.64	109
weighted avg	0.84	0.83	0.83	109

Gambar 4. 5 Performa Model

Pada gambar 4.5 menunjukkan performa model yang dihasilkan berdasarkan data yang tidak seimbang, sehingga memperoleh *accuracy* sebesar 83%, *precision* 84%, *recall* 83%, *f1-score* 83%.



Gambar 4. 6 Diagram Confusion Matrix

Pada gambar 4.6 menunjukkan performa model klasifikasi dengan tiga kelas yang diberi label 0, 1 dan 2. Nilai diagonal utama (6, 78, 6) menunjukkan jumlah prediksi yang benar untuk masing-masing kelas, dimana model berhasil

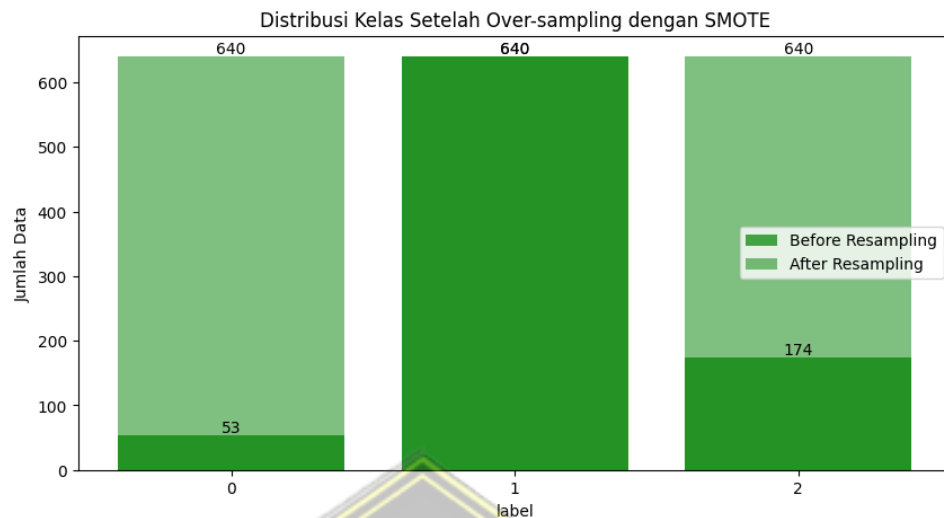
mengklasifikasikan *true positif* 6 data, *true negative* 78 data, dan *true netral* 6 data.

Namun, terdapat kesalahan klasifikasi yang terlihat pada elemen non-diagonal. Pada kelas 0 salah diklasifikasikan sebagai kelas 1 sebanyak 1 kali dan sebagai kelas 2 sebanyak 1 kali. Selain itu, kelas 1 salah diprediksi sebagai kelas 0 sebanyak 3 kali dan sebagai kelas 2 sebanyak 5 kali. Kelas 2 juga memiliki beberapa kesalahan, dimana 4 sampel kelas 2 diklasifikasikan sebagai kelas 0 dan 5 sampel kelas 2 diklasifikasikan sebagai kelas 1.

Hal ini dapat terjadi karena model terlalu kompleks untuk dataset yang digunakan, atau karena jumlah data training yang tidak mencukupi untuk menangkap pola yang lebih umum.

4.2.2 Data Balance

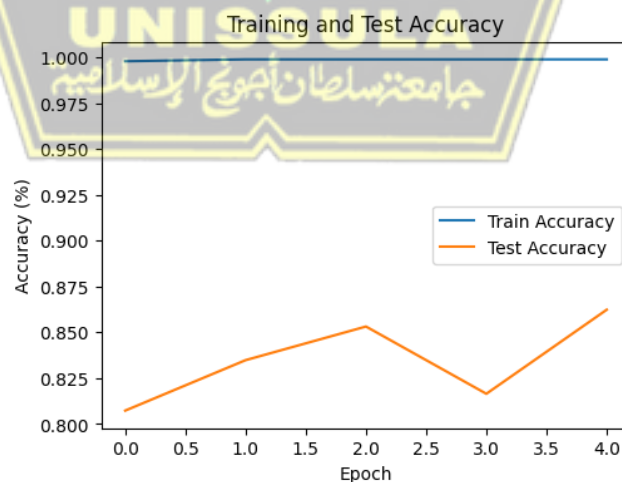
Dikarenakan terdapat ketidakseimbangan distribusi data untuk setiap kelas pada dataset, dapat berpengaruh pada kinerja model. Model yang dilatih dengan data yang tidak seimbang akan cenderung mengklasifikasikan kelas mayoritas secara berlebihan dan mengabaikan kelas minoritas, sehingga akan mengakibatkan banyak kesalahan klasifikasi. Selain itu, distribusi data yang tidak seimbang juga dapat mengakibatkan *overfitting*. *Overfitting* merupakan kondisi dimana model memiliki kinerja yang baik terhadap data training, namun tidak pada data testing. Oleh karena itu, pada penelitian ini dilakukan penyeimbangan data training menggunakan SMOTE sebelum data digunakan untuk pelatihan model. SMOTE melakukan *oversampling* untuk kelas minoritas agar jumlah datanya menjadi seimbang dengan kelas mayoritas. Hasil dari proses *balancing* ini ditunjukkan pada gambar 4.7.



Gambar 4. 7 Jumlah data hasil *balancing*

Gambar 4.7 memperlihatkan distribusi data *training* setelah dilakukan *resampling* menggunakan SMOTE. Warna hijau tua menunjukkan jumlah data asli sebelum di seimbangkan, dan warna hijau muda menunjukkan jumlah data setelah dilakukan *resampling* dengan SMOTE. Jumlah data kelas minoritas diseimbangkan dengan kelas mayoritas sehingga jumlah datanya menjadi 640 untuk setiap kelas.

Hasil eksperimen pada dataset yang sudah melalui proses balancing dapat dilihat pada gambar 4.8.



Gambar 4. 8 menunjukkan hasil eksperimen dataset *balancing* dengan iterasi epoch 5

Gambar 4.8 menunjukkan bahwa akurasi training tetap sangat tinggi dan hampir mendekati 100% di sepanjang epoch. Namun, akurasi test mengalami

overfitting, dengan train awal meningkat hingga epoch 2.0, kemudian menurun pada epoch 3.0, dan kembali meningkat pada epoch 4.0. Selain itu, perbedaan besar antara akurasi training dan akurasi test menandakan bahwa model lebih banyak menghafal data training daripada memahami pola yang dapat diterapkan ke data baru.

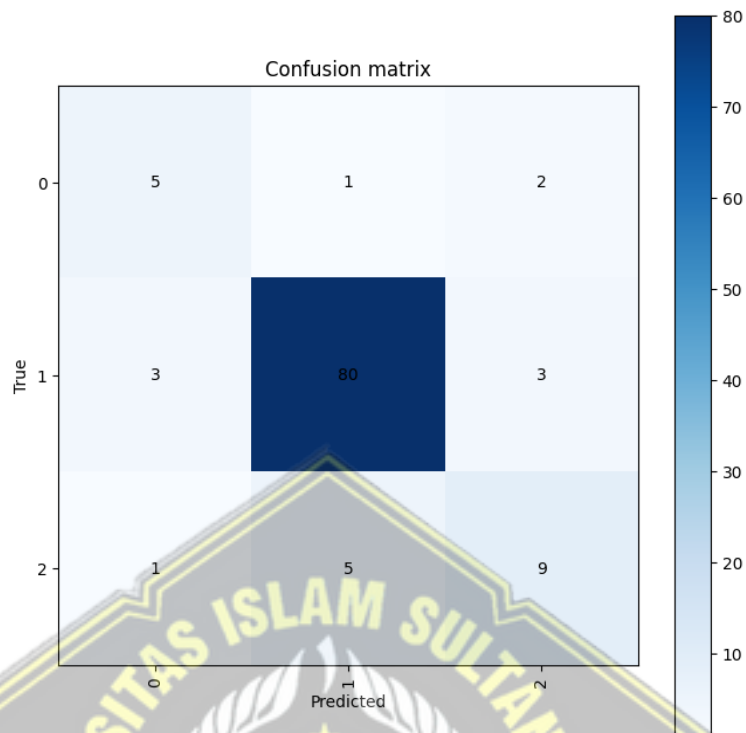
Setelah proses pelatihan dilakukan, tahap selanjutnya yaitu proses menguji performa model. Hasil dari implementasi model terhadap data training dan data testing menunjukkan hasil *accuracy*, *precision*, *recall* dan *f1-score* dari dataset inbalance yang diperoleh ditunjukkan pada gambar 4.9.

Classification Report:

	precision	recall	f1-score	support
0	0.56	0.62	0.59	8
1	0.93	0.93	0.93	86
2	0.64	0.60	0.62	15
accuracy			0.86	109
macro avg	0.71	0.72	0.71	109
weighted avg	0.86	0.86	0.86	109

Gambar 4. 9 Performa Model

Pada gambar 4.9 menunjukkan performa model yang dihasilkan berdasarkan data yang sudah melalui proses *resampling*, sehingga memperoleh *accuracy* sebesar 86%, *precision* 86%, *recall* 86%, dan *f1-score* 86%.



Gambar 4. 10 Diagram Confusion Matrix

Pada gambar 4.10 menunjukkan performa model klasifikasi dengan tiga kelas yang diberi label 0, 1 dan 2. Nilai diagonal utama (5, 80, 9) menunjukkan jumlah prediksi yang benar untuk masing-masing kelas, dimana model berhasil mengklasifikasikan *true positif* menghasilkan 5 data, *true negatif* 80 data, dan *true netral* 9 data.

4.3 Hasil Implementasi Streamlit

Hasil dari pemodelan yang sebelumnya sudah dilakukan dengan menggunakan model IndoBERT, selanjutnya akan dibuat sistem yang menggunakan *framework Streamlit*. Sehingga nantinya model ini memiliki sistem yang dapat melakukan inputan yang dilakukan oleh fans NCT. Hal tersebut ditunjukkan pada gambar 4.11.

Analisis Sentimen Teks dengan IndoBERT

Masukkan teks yang ingin dianalisis sentimennya



The screenshot shows a web application interface. At the top, there is a title 'Analisis Sentimen Teks dengan IndoBERT'. Below the title is a prompt 'Masukkan teks yang ingin dianalisis sentimennya'. Underneath this is a large, empty text input area with a placeholder text 'Masukkan teks di sini...'. At the bottom of the input area is a button labeled 'Analisis Sentimen'.

Gambar 4. 11 Tampilan Awal *Streamlit*

Gambar 4.11 merupakan tampilan awal dari *Streamlit*, sebelumnya agar tampilan awal bisa muncul harus menjalankan perintah `streamlit run namafile.py` di command prompt terlebih dahulu. Setelah tampilan awal muncul, tahap selanjutnya *User* dapat memasukkan pendapat mereka terhadap *Endorsement* NCT dengan Starbucks. Seperti yang terlihat pada gambar 4.12.

Masukkan teks yang ingin dianalisis sentimennya



The screenshot shows the same web application interface as in Gambar 4.11, but now the text input area is filled with the following text: 'PLEASE BUKA MATA BUKA HATI!!! KEHILANGAN ENDORSEMENT GA AKAN MEMBUAT KALIAN KEHILANGAN PENGEMAR!!!! FANS HANYA BUTUH IDOL YANG MANUSIAWII!'. The text is in all caps and appears to be a comment or endorsement. The 'Analisis Sentimen' button is still visible at the bottom.

Gambar 4. 12 *Inputan* Text Komentar

Gambar 4.12 menunjukkan inputan yang telah diisi oleh *User*. Tahap selanjutnya *User* dapat mengklik tombol hasil Analisis Sentimen maka hasil analisisnya nanti akan keluar.

Masukkan teks yang ingin dianalisis sentimennya

PLEASE BUKA MATA BUKA HATI!!! KEHILANGAN ENDORSEMENT GA AKAN MEMBUAT KALIAN
KEHILANGAN PENGEMAR!!!! FANS HANYA BUTUH IDOL YANG MANUSIAW!!!

Analisis Sentimen

Hasil Sentimen: negatif

Gambar 4. 13 Tampilan Tombol klik dan Hasil Analisis Sentimen

Pada gambar 4.13 setelah *User* memasukkan text kemudian *User* dapat meng-klik tombol Analisis Sentimen setelah itu hasil analisisnya akan keluar dan menghasilkan sentimen negatif seperti yang terlihat pada gambar 4.13

Analisis Sentimen Teks dengan IndoBERT

Masukkan teks yang ingin dianalisis sentimennya

Cantik sekali, aku ingin membelinya juga haha.

Analisis Sentimen

Hasil Sentimen: positif

Gambar 4. 14 Hasil Sentimen Positif

Percobaan kedua dapat dilihat pada gambar 4.14 menunjukkan hasil analisis sentimen kategori positif.

Analisis Sentimen Teks dengan IndoBERT

Masukkan teks yang ingin dianalisis sentimennya

ini orang Starbucks liat komen nya ga sii, kena mentall kata gue mah 🙏🙏

Analisis Sentimen

Hasil Sentimen: netral

Gambar 4. 15 Hasil Sentimen Netral

Percobaan ketiga dapat dilihat pada gambar 4.15 menunjukkan hasil analisis sentimen netral.



BAB V

KESIMPULAN DAN SARAN

5.1 Kesimpulan

Berdasarkan hasil penelitian yang telah dilakukan, dapat disimpulkan bahwa, komentar para fans NCT terhadap *endorsement* dengan Starbucks memiliki data sentimen yang tidak seimbang yaitu lebih banyak pola sentimen negatif, hal itu dapat dilihat dari hasil eksperimen yang telah dilakukan. Bahwa dari eksperimen data yang tidak seimbang dengan data yang sudah diseimbangkan, akurasi yang dihasilkan lebih baik data yang sudah diseimbangkan daripada data yang tidak seimbang. Sehingga pada eksperimen awal data yang tidak seimbang mengalami *overfitting* dan performa model yang tidak terlalu baik. Sedangkan pada data yang sudah melalui proses *balancing* mendapatkan akurasi data dan performa model yang cukup baik, yaitu *accuracy* sebesar 86%, *precision* 86%, *recall* 86%, dan *f1-score* 86%.

5.2 Saran

Saran yang diberikan untuk penelitian yang akan datang, yaitu:

1. Dari hasil penelitian ini sistem memiliki akurasi terbesar 86%, *precision* 86%, *recall* 86%, dan *f1-score* 86%, namun penelitian ini hanya dapat menganalisis sentimen pada ulasan atau komentar fans NCT saja. Diharapkan untuk penelitian selanjutnya sistem dapat dikembangkan menjadi sistem yang lebih lengkap.
2. Menggunakan dataset yang lebih banyak dan kalimat yang lebih panjang
3. Diharapkan nantinya peneliti dapat mengatasi *overfitting* yang terjadi.
4. Model pada tugas akhir ini dilatih dengan maksimal iterasi 5 epoch karena spesifikasi laptop yang kurang, sehingga diharapkan penelitian selanjutnya dapat menggunakan laptop yang memiliki spesifikasi tinggi.

DAFTAR PUSTAKA

- Analisis Pada Twitter CommuterLine, S. dkk. (2022) "Text Mining untuk Sentimen Analisis dengan Metode Naïve Bayes, SMOTE, N-Gram dan AdaBoost Pada Twitter CommuterLine," *Jurnal Sains Komputer & Informatika (J-SAKTI)*, 6(2), hal. 961–973.
- Apri Wenando, F. (2023) "Analisis Opini Publik Terhadap Undang-Undang KUHP Tahun 2022 Menggunakan Algoritma Naïve Bayes Classifier," *Jurnal Fasikom*, 13(02), hal. 334–339. doi: 10.37859/jf.v13i02.5670.
- Chinnasamy, P. dkk. (2022) "COVID-19 vaccine sentiment analysis using public opinions on Twitter," *Materials Today: Proceedings*, 64, hal. 448–451. doi: 10.1016/j.matpr.2022.04.809.
- Devlin, J. dkk. (2019) "BERT: Pre-training of deep bidirectional transformers for language understanding," *NAACL HLT 2019 - 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies - Proceedings of the Conference*, 1(Mlm), hal. 4171–4186.
- Isnain, A. R. dkk. (2021) "Sentimen Analisis Publik Terhadap Kebijakan Lockdown Pemerintah Jakarta Menggunakan Algoritma Svm," *Jurnal Data Mining dan Sistem Informasi*, 2(1), hal. 31. doi: 10.33365/jdmsi.v2i1.1021.
- Isnain, A. R., Marga, N. S. dan Alita, D. (2021) "Sentiment Analysis Of Government Policy On Corona Case Using Naive Bayes Algorithm," *IJCCS (Indonesian Journal of Computing and Cybernetics Systems)*, 15(1), hal. 55. doi: 10.22146/ijccs.60718.
- Mas, R. dkk. (2021) "Analisis Sentimen Customer Review Aplikasi Ruang Guru dengan Metode BERT (Bidirectional Encoder Representations from Transformers)," *Jeisbi*, 02(03), hal. 2021.
- Melati, D. (2024) "No Title," *GEMA KETIDAKSETUJUAN: GERAKAN BOIKOT SM ENTERTAINMENT DAN NCT TERKAIT KERJASAMA STARBUCKS*, hal. 3.
- Mubaraq, M. F. dan Maharani, W. (2022) "Sentiment Analysis on Twitter Social Media towards Climate Change on Indonesia Using IndoBERT Model," *Jurnal Media Informatika Budidarma*, 6(4), hal. 2426. doi: 10.30865/mib.v6i4.4570.
- Permatasari, P. A., Linawati, L. dan Jasa, L. (2021) "Survei Tentang Analisis

Sentimen Pada Media Sosial,” *Majalah Ilmiah Teknologi Elektro*, 20(2), hal. 177. doi: 10.24843/mite.2021.v20i02.p01.

Pilar, G. D. dkk. (2023) “A novel flexible feature extraction algorithm for Spanish tweet sentiment analysis based on the context of words,” *Expert Systems with Applications*, 212(September 2022). doi: 10.1016/j.eswa.2022.118817.

Ramadan, A. dan Fatchiya, A. (2021) “Efektivitas Instagram sebagai Media Promosi Produk ‘ Rendang Uninam ’ The Effectiveness of Instagram as a Promotional Media of Products ‘ Rendang Uninam ,” 05(01), hal. 64–82.

Rizkina, N. Q. dan Hasan, F. N. (2023) “Analisis Sentimen Komentar Netizen Terhadap Pembubaran Konser NCT 127 Menggunakan Metode Naive Bayes,” *Journal of Information System Research (JOSH)*, 4(4), hal. 1136–1144. doi: 10.47065/josh.v4i4.3803.

Rohman, A. N., Utami, E. dan Raharjo, S. (2019) “Deteksi Kondisi Emosi pada Media Sosial Menggunakan Pendekatan Leksikon dan Natural Language Processing,” *Eksplora Informatika*, 9(1), hal. 70–76. doi: 10.30864/eksplora.v9i1.277.

Syah, H. dan Witanti, A. (2022) “Analisis Sentimen Masyarakat Terhadap Vaksinasi Covid-19 Pada Media Sosial Twitter Menggunakan Algoritma Support Vector Machine (Svm),” *Jurnal Sistem Informasi dan Informatika (Simika)*, 5(1), hal. 59–67. doi: 10.47080/simika.v5i1.1411.

Vaswani, A. (2017) “Attention is All you Need,” *Advances in Neural Information Processing Systems*, hal. I. Tersedia pada: <https://cir.nii.ac.jp/crid/1370849946232757637>.