



ANALISA SENTIMEN TERHADAP ULASAN PENGGUNA PADA APLIKASI POLRI PRESISI MENGGUNAKAN METODE BERT-BIDIRECTIONAL ENCODER REPRESENTATION FROM TRANSFORMERS

Sentiment Analysis of User Reviews on the Polri Presisi Application Using the BERT Method.

Linda Wahyu Widianti^{1*}, Mohamad Saefudin²

Program Studi Sistem Informasi,^{1,2}
STMIK Jakarta STI&K^{1,2}

*Correspondent Author: lindawewe100@gmail.com

Email : lindawewe100@gmail.com¹, saefudin@gmail.com²

Received: October 18,2025. **Revised:** November 20,2025. **Accepted:** December 03, 2025. **Issue Period:** Vol.9 No.2 (2025), Pp. 277-290

Abstrak: Aplikasi Penelitian ini bertujuan untuk menganalisis sentimen ulasan pengguna pada aplikasi Polri Presisi menggunakan metode BERT (Bidirectional Encoder Representations from Transformers). Sebanyak 10.000 ulasan dari Google Play Store dikumpulkan melalui proses scraping dan diproses melalui tahapan pre-processing seperti case folding, pembersihan data, tokenisasi, penghapusan stopwords, normalisasi, dan stemming. Model BERT dilatih dan diuji menggunakan 3.000 data uji, menghasilkan akurasi 86,6% dan Matthews Correlation Coefficient (MCC) 0,720. Model memiliki kinerja tinggi pada sentimen negatif (precision 0,89; recall 0,92) dan positif (precision 0,90; recall 0,86), namun masih rendah pada sentimen netral (precision 0,20; recall 0,18). Hasil ini menunjukkan efektivitas BERT dalam analisis sentimen berbahasa Indonesia, sekaligus mengindikasikan perlunya peningkatan performa pada kelas netral.

Kata kunci: Analisis, Sentimen, Polri, Presisi, Ulasan, Pengguna

Abstract: This study aims to analyze user review sentiments on the Polri Presisi application using the BERT (Bidirectional Encoder Representations from Transformers) method. A total of 10,000 reviews were collected from the Google Play Store through a scraping process and pre-processed using steps such as case folding, data cleaning, tokenization, stopwords removal, normalization, and stemming. The BERT model was trained and tested using 3,000 test data, achieving an accuracy of 86.6% and a Matthews Correlation Coefficient (MCC) of 0.720. The model demonstrated strong performance in classifying negative sentiment (precision 0.89; recall 0.92) and positive sentiment (precision 0.90; recall 0.86), but relatively



DOI: 10.52362/jisicom.v9i2.2140

Ciptaan disebarluaskan di bawah [Lisensi Creative Commons Atribusi 4.0 Internasional](https://creativecommons.org/licenses/by/4.0/).



low performance in neutral sentiment (precision 0.20; recall 0.18). These results indicate the effectiveness of BERT in Indonesian-language sentiment analysis while highlighting the need for improvements in detecting neutral sentiment.

Keywords: *Sentiment, Analysis, Polri, Presisi Application, User Reviews.*

I. PENDAHULUAN

Perkembangan teknologi informasi telah mendorong instansi pemerintah untuk terus melakukan inovasi dalam memberikan pelayanan publik yang efektif dan efisien. Kepolisian Negara Republik Indonesia (Polri) turut mengikuti perkembangan ini dengan menghadirkan aplikasi digital bernama Polri Presisi. Aplikasi ini bertujuan untuk mempermudah masyarakat dalam mengakses berbagai layanan kepolisian, seperti pengaduan, pengecekan status laporan, hingga pelaporan tindak kejahatan secara daring.

Sebagai aplikasi layanan publik, Polri Presisi telah diunduh dan digunakan oleh ribuan masyarakat Indonesia. Pengguna secara aktif memberikan ulasan mereka melalui platform Google Play Store, baik berupa pujian, kritik, maupun saran. Ulasan-ulasan ini mengandung informasi penting mengenai persepsi dan pengalaman pengguna terhadap kualitas dan kinerja aplikasi. Namun, banyaknya jumlah ulasan dan sifatnya yang tidak terstruktur menyulitkan proses analisis secara manual.

Untuk mengatasi permasalahan tersebut, dibutuhkan metode yang mampu mengolah data teks dalam jumlah besar dan mengklasifikasikan opini pengguna secara otomatis. Analisis sentimen merupakan salah satu solusi yang umum digunakan dalam mengelompokkan opini ke dalam kategori positif, negatif, atau netral. Salah satu pendekatan terbaru dan paling efektif dalam bidang pemrosesan bahasa alami (Natural Language Processing/NLP) adalah penggunaan model BERT (Bidirectional Encoder Representations from Transformers).

BERT adalah model berbasis transformer yang dikembangkan oleh Google dan memiliki kemampuan untuk memahami konteks kata secara dua arah (kiri dan kanan). Berbeda dengan model NLP tradisional, BERT mampu mempertimbangkan seluruh kalimat secara utuh dalam memahami makna suatu kata. Hal ini memungkinkan BERT untuk menangani ambiguitas dalam bahasa alami dengan lebih akurat.

Dalam penelitian ini, BERT akan digunakan untuk menganalisis sentimen terhadap ulasan pengguna pada aplikasi Polri Presisi yang diperoleh dari Google Play Store. Penelitian ini tidak hanya bertujuan untuk mengetahui persepsi pengguna, tetapi juga mengevaluasi efektivitas BERT dalam klasifikasi sentimen berbahasa Indonesia.

Berdasarkan dari latar belakang yang sudah diuraikan sebelumnya maka ruang lingkup dari penelitian adalah data yang diambil berdasarkan ulasan yang digunakan berjumlah 10000 ulasan dari Google Play Store. Data yang digunakan merupakan data dengan Bahasa Indonesia. Bahasa pemrograman yang digunakan adalah Python. Analisis sentimen diolah menggunakan metode Bidirectional Encoder Representation From Transformers (BERT). Dataset yang diolah menggunakan metode BERT, diklasifikasikan menjadi tiga kategori yaitu positif, negative dan netral. Data yang diklasifikasikan merupakan data berupa ulasan yang diambil dari Google Play Store pada periode Mei 2022 sampai dengan Agustus 2025

II. METODE DAN MATERI

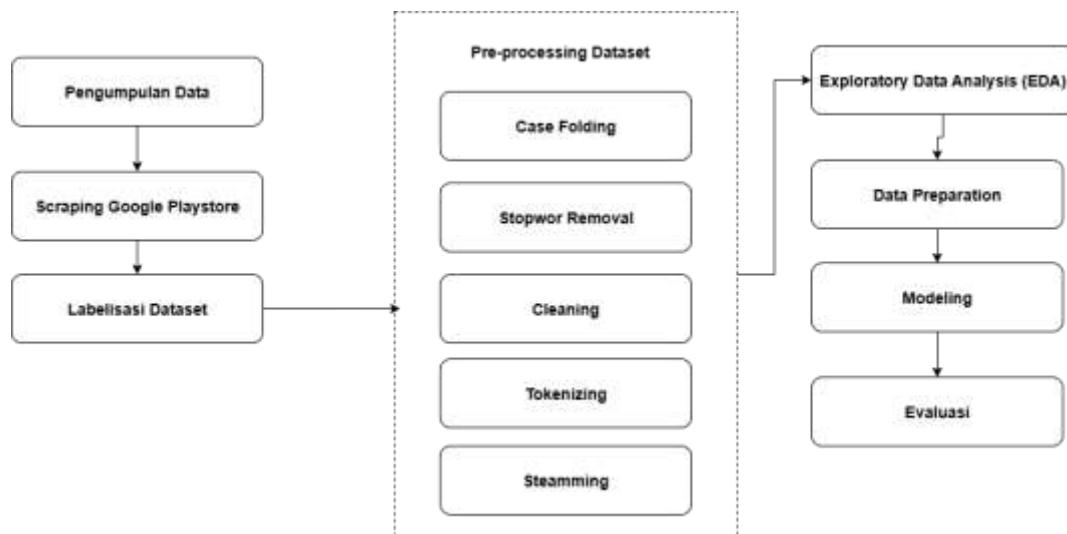
2.1 Tahapan Penelitian

Penelitian ini berfokus pada analisis sentimen terhadap aplikasi Polri Presisi, yaitu layanan publik digital yang dikembangkan oleh Kepolisian Negara Republik Indonesia (Polri) guna mempermudah masyarakat dalam mengakses berbagai layanan kepolisian secara terpadu. Tujuannya adalah untuk mewujudkan pelayanan yang cepat, transparan, dan akuntabel.

Melalui ulasan pengguna di Google Play Store, penelitian ini menganalisis persepsi publik terhadap kualitas layanan aplikasi menggunakan model Bidirectional Encoder Representations from Transformers (BERT), khususnya IndoBERT. Tahapan penelitian meliputi:



1. Pengumpulan data
2. Web scraping
3. Labelisasi dataset
4. Pra-pemrosesan (pre-processing)
5. Pembagian data
6. Pelatihan model
7. Evaluasi hasil



Gambar 1. Tahapan Analisis Sentimen BERT

(Diagram alur: Pengumpulan Data → Pre-processing → EDA → Data Preparation → Modeling → Evaluation)

2.2 Pengumpulan Data

Data penelitian diperoleh dari ulasan pengguna aplikasi Polri Presisi di Google Play Store tahun 2025 sebanyak 10.000 entri. Pengumpulan data dilakukan dengan web scraping menggunakan Google Colaboratory dan pustaka google-play-scraper pada bahasa pemrograman Python. Hanya ulasan berbahasa Indonesia yang digunakan, dan data disimpan dalam format CSV untuk memudahkan pengolahan lebih lanjut.

2.3 Scrapping Data

Proses web scraping dilakukan untuk mengambil teks ulasan, nama pengguna, serta nilai rating dari aplikasi Polri Presisi. Data tersebut menjadi sumber utama analisis sentimen. Hasil scraping dibersihkan dan diklasifikasikan ke dalam tiga kategori sentimen, yaitu positif, netral, dan negatif.



Gambar 2. Tahapan Scrapping Data

2.4 Labelisasi Dataset



DOI: 10.52362/jisicom.v9i2.2140

Ciptaan disebarluaskan di bawah [Lisensi Creative Commons Atribusi 4.0 Internasional](https://creativecommons.org/licenses/by/4.0/).

Pelabelan dilakukan secara manual rule-based dengan mengacu pada nilai rating (1–5).

Rating 1–2: Sentimen Negatif

Rating 3: Sentimen Netral

Rating 4–5: Sentimen Positif

Tabel 1. Pemetaan Labelisasi

Skor Rating	Label Numerik	Kategori Sentimen
1	0	Negatif
2	0	Negatif
3	2	Netral
4	1	Positif
5	1	Positif

Metode ini menjamin konsistensi pelabelan sesuai karakteristik data.

2.5 Pre-Processing Dataset

Tahap pra-pemrosesan bertujuan mengubah data tidak terstruktur menjadi data yang siap dianalisis. Proses ini mencakup:

1. Case Folding-Mengubah huruf menjadi *lowercase* untuk menghindari perbedaan kata akibat kapitalisasi.
2. Stopword Removal-Menghapus kata umum yang tidak memiliki makna penting (seperti “dan”, “yang”, “di”).
3. Cleaning Dataset-Menghilangkan angka, tanda baca, emoji, serta karakter berulang.
4. Tokenizing-Memecah teks menjadi unit kata (token) agar dapat dianalisis secara individual.
5. Stemming-Mengubah menjadi kata dasar (contoh: berlari, pelari, menjadi lari) menurunkan variasi kata.



Gambar 3. Tahapan Pre-Processing Dataset

(Diagram alur: Case Folding → Stopword Removal → Cleaning → Tokenizing → Stemming)

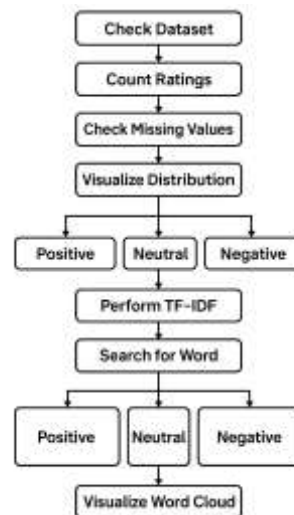
Tahapan ini menghasilkan data teks bersih yang siap digunakan dalam pemodelan.

2.6 Exploratory Data Analysis (EDA)

EDA dilakukan untuk memahami karakteristik data sebelum pemodelan. Langkah-langkahnya meliputi:

1. Pemeriksaan ukuran dataset dan distribusi rating.
2. Pengecekan missing values.
3. Visualisasi distribusi ulasan dengan count plot dan pie chart.
4. Analisis TF-IDF untuk menentukan kata paling berpengaruh pada masing-masing sentimen.
5. Pembuatan Word Cloud untuk menampilkan kata dominan pada sentimen positif, netral, dan negatif.

Hasil EDA memberikan gambaran pola bahasa dan persebaran data yang digunakan dalam proses modeling.

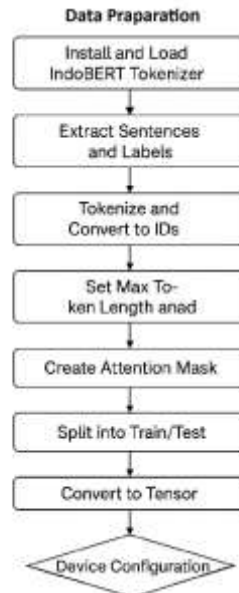


Gambar 4. Tahapan EDA

2.7 Data Preparation

Data preparation dilakukan untuk menyiapkan data agar dapat diproses oleh model IndoBERT. Proses meliputi:

1. Tokenisasi menggunakan BertTokenizer dari model `indolem/indobert-base-uncased`.
2. Penentuan panjang maksimum (65 token) dan penerapan padding untuk menyeragamkan input.
3. Pembuatan attention mask untuk menandai token aktif.
4. Pembagian dataset menjadi data latih (70%) dan data uji (30%) dengan proporsi label seimbang.
5. Konversi data ke tensor menggunakan PyTorch agar dapat diproses pada GPU.
6. Pembuatan DataLoader untuk pembagian data dalam batch saat pelatihan dan pengujian.



Gambar 5. Alur Data Preparation

(Diagram: Tokenisasi → Padding → Attention Mask → Split Data → Tensor Conversion)

2.8 Modeling

Pemodelan dilakukan dengan menggunakan IndoBERT, model bahasa berbasis BERT yang telah disesuaikan untuk bahasa Indonesia. Model BertForSequenceClassification digunakan untuk klasifikasi multi-kelas dengan tiga label (positif, netral, negatif).

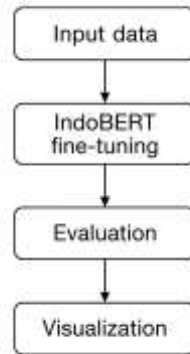
Parameter pelatihan meliputi:

- Tahap Optimisasi (AdamW): Algoritma AdamW digunakan untuk menyesuaikan bobot model secara efisien dengan mengurangi risiko overfitting melalui regulasi bobot (weight decay).
- Learning Rate ($2e-5$): Nilai ini menunjukkan kecepatan pembaruan bobot model di setiap iterasi pelatihan. Nilai kecil (2×10^{-5}) membantu proses konvergensi lebih stabil.
- Epoch (10): Model dilatih sebanyak 10 kali pengulangan penuh terhadap seluruh data latih untuk mencapai hasil optimal.
- Scheduler (get_linear_schedule_with_warmup): Mengatur learning rate agar meningkat perlahan di awal (warm-up) lalu menurun secara bertahap, sehingga pelatihan lebih stabil.
- Random Seed (42): Nilai acuan acak agar hasil pelatihan dapat direplikasi dengan hasil yang sama pada setiap percobaan.
- Model dijalankan menggunakan GPU untuk mempercepat pelatihan. Nilai loss setiap epoch dicatat, dan visualisasi training loss curve dibuat untuk memantau proses konvergensi model.



DOI: 10.52362/jisicom.v9i2.2140

Ciptaan disebarluaskan di bawah [Lisensi Creative Commons Atribusi 4.0 Internasional](https://creativecommons.org/licenses/by/4.0/).



Gambar 6. Tahapan Modeling

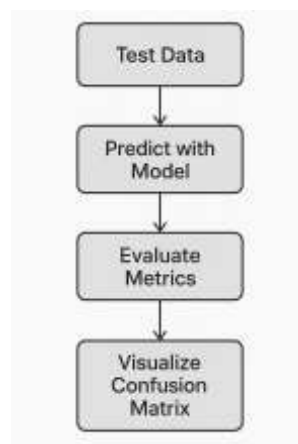
2.9 Predict & Evaluate

Tahap evaluasi bertujuan mengukur kemampuan generalisasi model pada data uji. Model dijalankan dalam evaluation mode, dan prediksi dilakukan menggunakan `torch.no_grad()` untuk efisiensi memori.

Evaluasi kinerja dilakukan menggunakan beberapa metrik:

1. Confusion Matrix – Menampilkan perbandingan prediksi dan label aktual.
2. Classification Report – Menyajikan nilai precision, recall, dan f1-score tiap kelas.
3. Matthews Correlation Coefficient (MCC) – Mengukur keseimbangan prediksi antar kelas.
4. Accuracy (ACC) – Menunjukkan tingkat akurasi keseluruhan model.

Pendekatan ini memberikan penilaian menyeluruh terhadap kinerja model IndoBERT dalam menganalisis sentimen publik terhadap aplikasi Polri Presisi.



Gambar 7. Tahapan Evaluasi Model

III. PEMBAHASA DAN HASIL

3.1 Hasil Pengumpulan Data

Data penelitian dikumpulkan dari ulasan pengguna aplikasi **Polri Presisi** di **Google Play Store** menggunakan metode *web scraping*. Total diperoleh **10.000 entri** berbahasa Indonesia dalam rentang waktu 5



DOI: 10.52362/jisicom.v9i2.2140

Ciptaan disebarluaskan di bawah [Lisensi Creative Commons Atribusi 4.0 Internasional](https://creativecommons.org/licenses/by/4.0/).

September 2022-2 Agustus 2025. Data mencakup *username*, tanggal, isi ulasan, dan skor rating, kemudian disimpan dalam format CSV untuk analisis lebih lanjut.



Gambar 8. Aplikasi Polri Presisi

3.2 Hasil Scraping Data

Proses *scraping* dilakukan dengan pustaka **google-play-scraper** di **Google Colab**. Data yang dihasilkan bersifat tidak terstruktur dan ambigu, sehingga diperlukan tahap pelabelan dan pembersihan sebelum dapat digunakan untuk analisis sentimen.



Gambar 9. Ulasan Aplikasi Polri Presisi



Gambar 10. Hasil Scraping Data

3.3 Hasil Labelisasi Dataset

Data diberi label sentimen berdasarkan nilai rating:

- 1–2: Negatif
- 3: Netral
- 4–5: Positif

Proses dilakukan secara otomatis menggunakan Python dan fungsi *mapping*, dengan verifikasi manual untuk menjaga akurasi. Hasilnya terbentuk dua kolom baru, yaitu *label* (angka) dan *category* (teks sentimen).



```
# Labeling
sns_learn['label'] = sns_learn['score'].map([1:0, 2:0, 3:1, 4:1, 5:1])
sns_learn['category'] = sns_learn['score'].map([1:'negative', 2:'negative', 3:'neutral', 4:'positive', 5:'positive'])
sns_learn.head()
```

	username	score	at	content	label	category
0	Pengguna Google	3	7/26/2025 11:09	Bintang tiga dulu deh, pertama kali bikin skck...	2	neutral
1	Pengguna Google	1	7/27/2025 11:21	masi perpanjang SIM online harapan mudah, malah...	0	negative
2	Pengguna Google	3	7/31/2025 6:34	saya pikir dengan adanya SKCK online membantu...	2	neutral
3	Pengguna Google	1	7/22/2025 4:29	persuna online tapi tetap di ribetkan dengan m...	0	negative
4	Pengguna Google	4	7/17/2025 8:34	tolong di bagian syarat untuk pembuatan SKCK d...	1	positive

Gambar 11. Source Code dan Hasil Lebelisasi Dataset

3.4 Hasil Pre-Processing Dataset

Tahap ini bertujuan membersihkan dan menyeragamkan data agar siap dianalisis. Beberapa langkah yang dilakukan:

1. **Case Folding:** Mengubah semua huruf menjadi huruf kecil dan menghapus karakter non-alfabet.
2. **Stopword Removal:** Menghapus kata umum yang tidak memiliki makna penting, kecuali kata “tidak”.
3. **Cleaning Dataset:** Menghapus mention, URL, tag HTML, emoji, angka, dan simbol.
4. **Tokenizing:** Memecah teks menjadi token (kata) untuk mempermudah pemrosesan.
5. **Stemming:** Mengembalikan setiap kata ke bentuk dasar menggunakan pustaka **Sastrawi**, agar variasi kata dengan makna sama dianggap serupa.

Tahapan ini menghasilkan data teks bersih, terstruktur, dan siap untuk analisis lebih lanjut.

3.5 Hasil Exploratory Data Analysis (EDA)

EDA dilakukan untuk memahami pola data sebelum pemodelan.

- **Distribusi Rating:** Dari 10.000 data, rating 1 mendominasi (5.665 ulasan), menunjukkan kecenderungan negatif terhadap aplikasi.

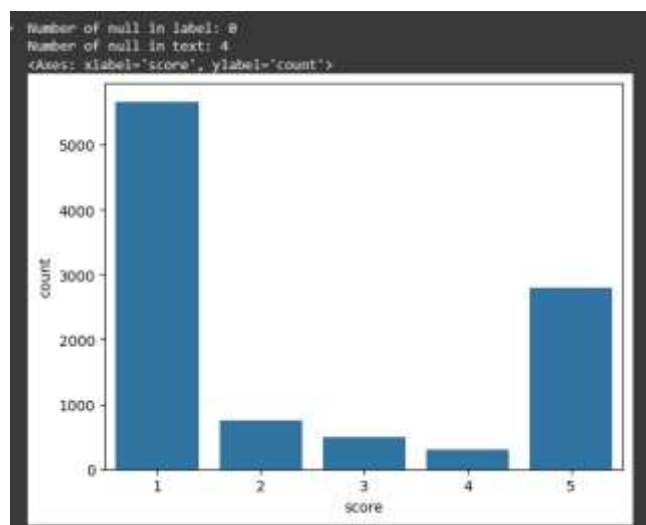
```

Input data has 10000 rows and 9 columns
rating 1 = 5665 rows
rating 2 = 757 rows
rating 3 = 487 rows
rating 4 = 302 rows
rating 5 = 2789 rows

```

Gambar 12. Hasil Hasil Distribusi Data Berdasarkan Rating

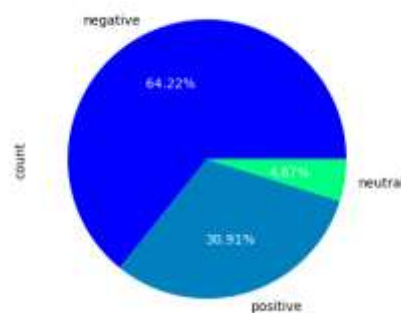
- **Missing Values:** Kolom *score* tidak memiliki nilai kosong, sedangkan *stemmed_review* memiliki 4 nilai kosong.



Gambar 13. Visualisasi Diagram Batang berdasarkan Rating

- **Distribusi Sentimen:** Sentimen negatif 64,22%, positif 30,91%, dan netral 4,87%.

Reviews Percentages Vs Category Sentiment



Gambar 13. Visualisasi Diagram Pie Berdasarkan Kategori

- **Jumlah Token:** Positif 25.305 token, netral 8.447, negatif 109.854 — ulasan negatif cenderung lebih panjang.

Total Tokens in Positive Category: 25305
Total Tokens in Neutral Category: 8447
Total Tokens in Negative Category: 109854

Gambar 14. Jumlah Token

- **Analisis TF-IDF:**
 - Positif: kata dominan “mudah”, “bantu”, “mantap”.
 - Negatif: “aplikasi”, “gagal”, “bayar”, “data”.
 - Netral: tidak menunjukkan kata dominan karena datanya sedikit.

	Word	TF-IDF	Percentage
1459	mudah	354.587931	4.875359
715	bantu	880.466008	4.111174
123	aplikasi	214.188398	2.045610
409	cepat	173.816488	2.300486
1188	layan	153.993889	2.185488
...	...	...	...
1637	omernya	0.002701	0.001275
414	erita	0.002701	0.001275
578	doff	0.002701	0.001275
317	big	0.002701	0.001275
727	glosy	0.002701	0.001275
[2563 rows x 3 columns]			

	Word	TF-IDF	Percentage
789	aplikasi	391.000461	1.796317
4052	nya	258.842627	1.152168
5411	stock	240.195516	1.103250
531	bayar	219.695064	1.005094
1165	data	200.470436	0.957537
...	...	...	...
1167	Injutkan	0.120887	0.000555
4456	peranjg	0.120887	0.000555
5433	slama	0.120887	0.000555
541	lckgrnd	0.120887	0.000555
4697	prpnjg	0.120887	0.000555
[6447 rows x 3 columns]			

	Word	TF-IDF	Percentage
78	aplikasi	26.971992	1.610100
134	bayar	24.944066	1.469389
1243	stock	26.842194	1.245121
940	nya	19.470263	1.163640
281	data	16.996400	1.134858
...	...	...	...
331	di	0.143749	0.000588
667	kg	0.138195	0.000256
1867	plus	0.138195	0.000256
391	fix	0.138195	0.000256
1134	ribu	0.138195	0.000256
[1000 rows x 3 columns]			

Gambar 15. Hasil Analisis TF-IDF Kategori Positif, Negatif dan Netral

- **Word Cloud:**
 - Positif menampilkan kata terkait kemudahan dan pelayanan.
 - Negatif menunjukkan keluhan seperti “gagal”, “ribet”, “verifikasi”.

EDA menunjukkan mayoritas ulasan bersifat negatif, dengan fokus keluhan pada aspek teknis aplikasi.

3.6 Hasil Data Preparation

Data disiapkan agar kompatibel dengan model **IndoBERT** melalui tiga tahap:

1. Tokenizer dan Encoding:

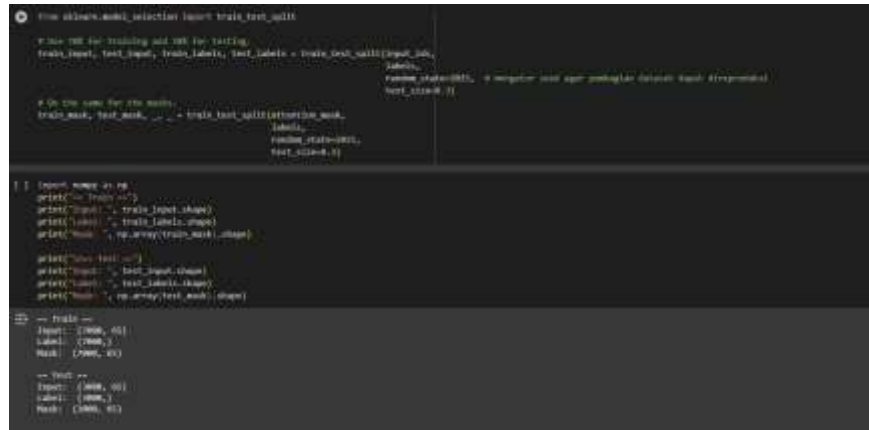
- Tokenisasi teks dengan *BertTokenizer* (indolem/indobert-base-uncased).
- Penambahan token khusus [CLS] dan [SEP].
- Penerapan *padding* (panjang 65 token) dan *attention mask* untuk menandai token aktif.



Gambar 16. Hasil Tokenizer dan Data Encoding

2. Data Splitting:

- Pembagian data menjadi **70% latih** dan **30% uji** menggunakan `train_test_split`.



```
from sklearn.model_selection import train_test_split

# Use 80% for training and 20% for testing
train_input, test_input, train_label, test_label = train_test_split(
    input_data,
    label_data,
    random_state=42, # mengontrol nilai agar pembagian dataset tidak divergen
    test_size=0.2)

# On the same for the mask
train_mask, test_mask, = train_test_split(attrition_mask,
    label_data,
    random_state=42,
    test_size=0.2)

[ ] In[ ]: import numpy as np
print("Train input shape")
print("label")
print("train_label shape")
print("test")
print("np.array(train_mask) shape")

print("train test")
print("train input shape")
print("label")
print("test_label shape")
print("test")
print("np.array(test_mask) shape")

-- train --
Input: (1000, 40)
label: (1000, 1)
Mask: (1000, 40)

-- Test --
Input: (200, 40)
label: (200, 1)
Mask: (200, 40)
```

Gambar 17. Hasil Data Splitting

3. Creating Data Loaders:

- Data latih menggunakan *RandomSampler* (acak), data uji dengan *SequentialSampler* (berurutan).
- *Batch size* = 16 untuk efisiensi memori dan waktu pelatihan.

3.7 Hasil Modeling

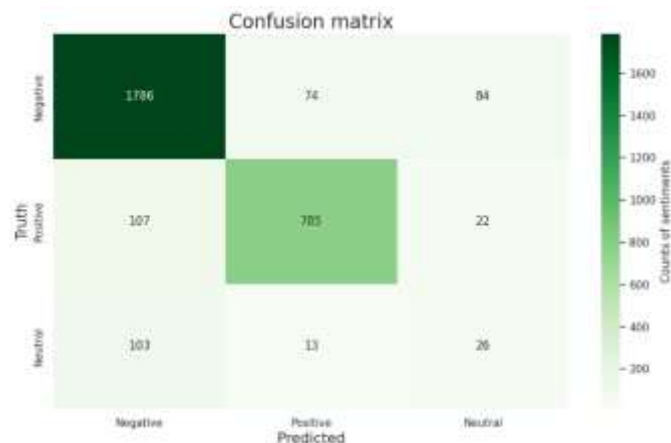
Model yang digunakan adalah **IndoBERT-base-uncased** untuk klasifikasi teks menjadi tiga kelas sentimen.

- **Optimisasi:** Menggunakan algoritma **AdamW** dengan *learning rate scheduler* untuk menjaga stabilitas pelatihan.
- **Pelatihan:** Dilakukan selama **10 epoch**, dengan *training loss* menurun dari 0,52 menjadi 0,13.
- **Akurasi Model:** Mencapai **98%** pada data uji, menunjukkan efektivitas model mengenali pola sentimen.
- **Waktu Pelatihan:** ±1 menit 27 detik per epoch.

3.8 Hasil Predict & Evaluate

Evaluasi dilakukan terhadap **3.000 data uji** menggunakan berbagai metrik:

- **Confusion Matrix:** Model sangat akurat pada sentimen positif dan negatif, tetapi masih kurang pada kategori netral.
- **Metrik Evaluasi:**
 - Negatif → *Precision* 0,89, *Recall* 0,92
 - Positif → *Precision* 0,90, *Recall* 0,86
 - Netral → *Precision* 0,20, *Recall* 0,18 (masih lemah)



Gambar 18. confusion matrix berdasarkan data uji

- **Akurasi Keseluruhan:** 86,6%

	precision	recall	f1-score	support
0	0.89	0.92	0.91	1944
1	0.90	0.86	0.88	914
2	0.20	0.18	0.19	142
accuracy			0.87	3000
macro avg	0.66	0.65	0.66	3000
weighted avg	0.86	0.87	0.86	3000

Gambar Hasil Evaluasi Model Terhadap Data Uji

- **Matthews Correlation Coefficient (MCC):** 0,72 (kategori korelasi kuat)

Hasil ini menunjukkan bahwa **IndoBERT** sangat efektif dalam menganalisis ulasan positif dan negatif, meskipun performa terhadap ulasan netral masih perlu ditingkatkan.

IV. KESIMPULAN

Berdasarkan hasil penelitian, model BERT yang dilatih pada data ulasan aplikasi Polri Presisi berhasil mencapai akurasi 86,6% dan nilai Matthews Correlation Coefficient (MCC) sebesar 0,720, yang menunjukkan adanya korelasi kuat antara prediksi dan label sebenarnya. Model memiliki kinerja sangat baik pada klasifikasi sentimen Negatif (precision 0,89; recall 0,92) dan Positif (precision 0,90; recall 0,86). Namun, performa pada sentimen Netral masih rendah (precision 0,20; recall 0,18) akibat tingginya tingkat misclassification dengan kelas positif dan negatif. Hal ini menunjukkan bahwa model sudah cukup andal dalam analisis sentimen berbahasa Indonesia, namun masih memerlukan peningkatan khususnya pada pengenalan sentimen netral.

REFERENASI

- [1.] H. M. Kim, S. P. Lee, and J. W. Choi, 2022, "A Comprehensive Study on Sentiment Analysis Techniques and Their Applications," Journal of Computer Science, vol. 18, no. 5, pp. 315-328.
- [2.] A. B. Prabowo, 2021, "Analisis Sentimen pada Ulasan Produk E-Commerce Menggunakan Metode Deep Learning," Jurnal Informatika, vol. 15, no. 1, pp. 45-56, 2021.



DOI: 10.52362/jisicom.v9i2.2140

Ciptaan disebarluaskan di bawah [Lisensi Creative Commons Atribusi 4.0 Internasional](https://creativecommons.org/licenses/by/4.0/).



e-ISSN : 2597-3673 (Online) , p-ISSN : 2579-5201 (Printed)

Vol.9 No.2 (December 2025)

Journal of Information System, Informatics and Computing

Website/URL: <http://journal.stmikjayakarta.ac.id/index.php/jisicom>

Email: jisicom@stmikjayakarta.ac.id , jisicom2017@gmail.com

- [3.] C. J. Chen and L. D. Wang, 2023, "Social Media Sentiment Analysis for Public Policy Decision-Making," *Big Data Analytics*, vol. 7, no. 3, pp. 112-125.
- [4.] J. A. Lee, 2020, "Foundations of Neural Networks and Their Role in Pattern Recognition," *International Journal of Artificial Intelligence*, vol. 12, no. 4, pp. 210-225.
- [5.] B. A. M. Nur, 2022, "Pendekatan Linguistik dan Komputasional dalam Natural Language Processing (NLP)," *Jurnal Bahasa dan Teknologi*, vol. 10, no. 2, pp. 78-91.
- [6.] M. F. Rasyid, 2023, "Perbandingan Algoritma Machine Learning dalam Klasifikasi Teks Berbahasa Indonesia," *Prosiding Konferensi Nasional Teknologi Informasi dan Komunikasi*.
- [7.] P. A. Johnson, 2021, "Understanding the Core Principles of Machine Learning Algorithms," *Journal of Data Science*, vol. 9, no. 1, pp. 55-68.
- [8.] B. T. Han, 2024, "A Survey of Supervised, Unsupervised, and Reinforcement Learning in Modern AI," *Expert Systems*, vol. 40, no. 6, pp. 412-428.
- [9.] K. L. Wu, A. R. Singh, and S. M. Patel, 2024, "Adaptive Learning Systems: Principles and Applications," *IEEE Transactions on Systems, Man, and Cybernetics*, vol. 54, no. 2, pp. 301-315.
- [10.] S. A. Ginting, 2023, "Ekstraksi Informasi Berbasis Kunci dalam Text Mining untuk Analisis Big Data," *Jurnal Sains Komputer*, vol. 22, no. 3, pp. 199-211.
- [11.] M. A. Fitri, 2021, "Penerapan Model BERT untuk Klasifikasi Teks dengan Data Berbahasa Inggris," *Seminar Nasional Informatika*.
- [12.] J. S. Li and R. M. Davis, 2022, "Contextual Embeddings with Bidirectional Transformers: A Deep Dive into BERT," *Computational Linguistics*, vol. 48, no. 1, pp. 50-65.
- [13.] N. P. Sari, 2023, "Fine-Tuning BERT for Sentiment Analysis: A Comparative Study," *International Journal of AI and Society*, vol. 16, no. 4, pp. 789-801.
- [14.] I. B. Setiawan, 2021, "Pengembangan IndoBERT: Model Bahasa Transformasi untuk Bahasa Indonesia," *Konferensi Internasional Bahasa dan Teknologi*.
- [15.] R. M. Wibowo, 2022, "Analisis Kinerja IndoBERT pada Tugas Natural Language Understanding (NLU) Bahasa Indonesia," *Jurnal Teknologi Informasi*, vol. 14, no. 2, pp. 110-125.



DOI: 10.52362/jisicom.v9i2.2140

Ciptaan disebarluaskan di bawah [Lisensi Creative Commons Atribusi 4.0 Internasional](https://creativecommons.org/licenses/by/4.0/).