

IMPLEMENTASI DAN ANALISIS SENTIMEN PADA ULASAN APLIKASI GOOGLE PLAY STORE MENGGUNAKAN ALGORITMA RoBERTa

BRIAN NATANAEL NAINGGOLAN¹,
KOKO HANDOKO²

¹ Mahasiswa Program Studi Teknik Informatika, Universitas Putera Batam

² Dosen Program Studi Teknik Informatika, Universitas Putera Batam

email: pb210210086@upbatam.ac.id

ABSTRACT

The rapid growth of Instagram user reviews on the Google Play Store poses challenges in understanding sentiment quickly and accurately. This research aims to develop an automatic sentiment analysis dashboard based on the Indonesian-RoBERTa model. Review data was collected using the google_play_scraper library and analyzed using a fine-tuned model. Fine-tuning was performed on the w11wo/indonesian-roberta-base-sentiment-classifier model using Indonesian tweet datasets during the PPKM period with 23,645 labeled data points (positive, neutral, negative). The preprocessing process included text cleaning, tokenization, and class weighting. Model evaluation used precision, recall, F1-score, and confusion matrix metrics. Test results showed good performance on positive and neutral classes, but performance on the negative class still needs improvement. The dashboard successfully performed scraping, sentiment prediction, and data visualization automatically. This research demonstrates the potential application of transformer-based models in Indonesian language and supports data-driven decision making.

Keywords: Fine-tuning; Google Play Store; Indonesian-RoBERTa; Instagram, Sentiment analysis.

PENDAHULUAN

Pertumbuhan teknologi digital telah mendorong peningkatan signifikan dalam aktivitas daring, termasuk penyampaian opini melalui platform seperti *Google Play Store*. Ulasan pengguna terhadap aplikasi mencerminkan sentimen yang dapat bernilai positif, netral, atau negatif, dan menjadi indikator penting bagi pengembang dalam mengevaluasi kualitas aplikasi.

Bagi perusahaan seperti Meta, pemilik aplikasi Instagram, memahami sentimen pengguna merupakan aspek krusial untuk

menjaga kepuasan dan loyalitas pengguna, namun volume ulasan yang sangat besar menyulitkan proses analisis secara manual. Oleh karena itu, pemanfaatan *Natural Language Processing* (NLP) dengan pendekatan *deep learning* menjadi solusi efektif untuk mengotomatisasi klasifikasi sentimen.

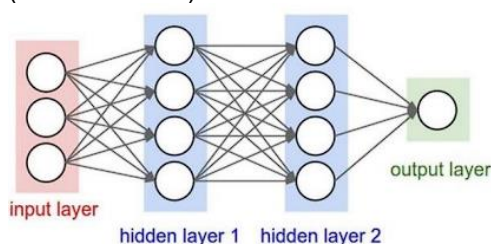
RoBERTa, sebagai pengembangan dari model BERT, menawarkan kinerja unggul dalam tugas-tugas NLP. Model *Indonesian-RoBERTa* variasi RoBERTa yang dilatih khusus pada data berbahasa Indonesia memungkinkan analisis

konteks kalimat secara lebih akurat. Dalam penelitian ini, *Indonesian-RoBERTa* digunakan untuk membangun sistem *dashboard* analisis sentimen yang mengklasifikasikan ulasan pengguna Instagram di *Google Play Store* secara otomatis dan efisien. Sistem ini diharapkan dapat mendukung pengambilan keputusan berbasis data secara lebih tepat.

KAJIAN TEORI

2.1. Deep learning

Deep learning merupakan cabang dari *machine learning* yang menggunakan jaringan saraf tiruan (*artificial neural networks*) untuk mempelajari pola dan representasi data secara otomatis. Berbeda dengan *machine learning* tradisional yang membutuhkan rekayasa fitur manual, *deep learning* mampu mengekstraksi fitur kompleks melalui arsitektur berlapis (*multilayered*), sehingga lebih efektif dalam menangani data tidak terstruktur seperti teks (Yudistira 2021).



Gambar 1 Struktur *Neural Network Deep learning*

(Sumber: Data Penelitian, 2025)

Struktur dasar jaringan saraf terdiri dari lapisan input, lapisan tersembunyi (*hidden layers*), dan lapisan output. Setiap neuron dalam jaringan ini terhubung melalui parameter *weight* dan *bias* yang disesuaikan selama pelatihan.

Proses pelatihan melibatkan dua tahap utama: *forward propagation* untuk menghasilkan prediksi, dan *backpropagation* untuk memperbaiki kesalahan prediksi melalui optimasi parameter (Drewiek-Ossowicka, Pietrolaj, and Rumiński 2021).

Fungsi aktivasi seperti ReLU dan *sigmoid* digunakan untuk menangani hubungan *non-linear* dalam data, sedangkan fungsi kehilangan (*loss function*) seperti *categorical cross entropy* digunakan untuk mengukur kesalahan prediksi (Sekhar and Meghana 2020). Kombinasi komponen ini memungkinkan *deep learning* menghasilkan model prediktif yang akurat, menjadikannya pendekatan unggul dalam tugas pemrosesan bahasa alami seperti analisis sentimen (Terven et al. 2023).

2.2 Natural Language Processing

Natural Language Processing (NLP) merupakan cabang kecerdasan buatan yang memungkinkan komputer untuk memahami, menafsirkan, dan menghasilkan bahasa manusia. Teknologi ini menjadi dasar dari berbagai aplikasi modern seperti *chatbot*, asisten virtual, terjemahan mesin, sistem pencarian informasi, serta analisis sentimen (Yudistira 2021).

Dalam konteks analisis sentimen, NLP bekerja melalui serangkaian tahap, mulai dari *preprocessing* hingga representasi semantik. Tahap awal mencakup proses *tokenisasi* (pemecahan teks menjadi unit kata), *stop word removal* (menghapus kata-kata umum yang tidak memiliki makna penting), serta *stemming* dan *lemmatization* untuk menyederhanakan kata ke bentuk dasarnya. Analisis struktur kalimat dilakukan melalui teknik seperti *Part-of-Speech (POS) tagging*

dan *dependency parsing*, yang membantu dalam memahami fungsi gramatikal dan hubungan antar kata. Representasi teks juga menjadi aspek penting dalam NLP. Pendekatan tradisional seperti *TF-IDF* mengukur pentingnya kata berdasarkan frekuensi kemunculan dalam dokumen dan korpus, sementara pendekatan modern seperti *word embeddings* (*Word2Vec*, *GloVe*, *FastText*) merepresentasikan kata sebagai vektor numerik dalam ruang semantik, memungkinkan model memahami makna kontekstual kata. Teknik seperti *n-grams* juga digunakan untuk menangkap hubungan antar kata secara berurutan. Dengan kemampuannya untuk memahami konteks linguistik dan semantik dalam teks, NLP menjadi fondasi utama dalam pengembangan sistem analisis sentimen berbasis *deep learning*, seperti yang digunakan dalam penelitian ini

2.2 Indonesian-RoBERTa

Indonesian-RoBERTa merupakan varian dari RoBERTa, model bahasa berbasis transformer yang dikembangkan oleh Facebook AI sebagai penyempurnaan dari BERT (*Bidirectional Encoder Representations from Transformers*). RoBERTa menghilangkan komponen *Next Sentence Prediction* (NSP), menggunakan data pelatihan yang lebih besar, serta menerapkan *dynamic masking*, sehingga menghasilkan performa yang lebih baik dalam berbagai tugas NLP (Naseer et al. 2022).

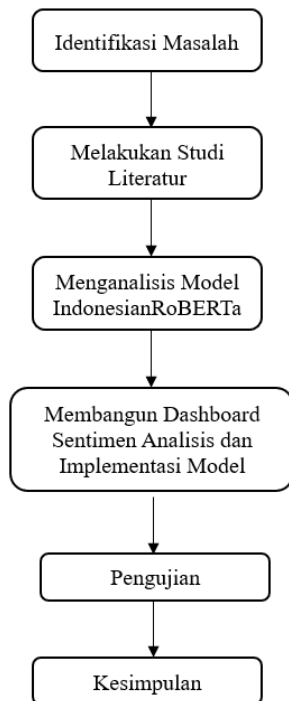
Indonesian-RoBERTa mempertahankan arsitektur dasar RoBERTa dan diadaptasi khusus untuk bahasa Indonesia. Model ini dilatih menggunakan korpus berbahasa Indonesia, termasuk dataset SmSA yang berisi komentar dan ulasan yang telah

dilabeli berdasarkan sentimen. Tokenisasi dilakukan menggunakan teknik *Byte-Pair Encoding* (BPE) yang mampu menangani keragaman morfologi dalam bahasa Indonesia secara efisien. *Embedding* yang digunakan mencakup *token embeddings* dan *positional embeddings*, tanpa *segment embeddings*, sesuai dengan karakteristik RoBERTa yang tidak menggunakan NSP (Sinha et al. 2021). Dalam arsitekturnya, *Indonesian-RoBERTa* menggunakan tumpukan *encoder* yang terdiri dari mekanisme *self-attention* dan *feed-forward neural networks* untuk memahami hubungan antar token dalam dua arah. Model ini terbukti efektif dalam berbagai tugas NLP dalam bahasa Indonesia, seperti klasifikasi teks, *named entity recognition* (NER), dan analisis sentimen. Kinerja superior *Indonesian-RoBERTa* dibandingkan model konvensional menjadikannya pilihan tepat untuk tugas klasifikasi sentimen terhadap ulasan pengguna aplikasi di *Google Play Store*.

METODE PENELITIAN

3.1 Desain Penelitian

Desain penelitian ini disusun sebagai kerangka kerja sistematis yang mengarahkan proses penelitian dari tahap identifikasi masalah hingga penarikan kesimpulan. Gambar 3.1 menggambarkan alur penelitian yang terdiri dari enam tahap utama.



Gambar 2. Alur Penelitian
(Sumber: Data Penelitian, 2025)

1. Identifikasi Masalah

Penelitian ini dimulai dengan mengidentifikasi permasalahan terkait perlunya pemahaman sentimen pengguna terhadap aplikasi Instagram di *Google Play Store*. Fokus penelitian adalah mengembangkan sistem yang dapat mengumpulkan, menganalisis, dan menyajikan data sentimen secara otomatis dan efektif.

2. Studi Literatur

Studi pustaka dilakukan untuk mengkaji teori-teori terkait NLP, algoritma berbasis transformer seperti BERT dan *RoBERTa*, serta penelitian terdahulu dalam bidang analisis sentimen. Kajian ini menjadi dasar teoritis dalam pemilihan metode dan pengembangan sistem.

3. Analisis Model *Indonesian-RoBERTa*
Pada tahap ini, dilakukan telaah mendalam terhadap arsitektur dan performa *Indonesian-RoBERTa* sebagai model utama dalam klasifikasi sentimen. Tujuannya adalah memastikan kesesuaian model dengan konteks bahasa Indonesia dan kebutuhan penelitian.

4. Pembangunan Dashboard dan Implementasi Model

Penulis mengembangkan dashboard menggunakan framework Laravel yang terintegrasi dengan modul *scraping* untuk mengambil data ulasan aplikasi dari *Google Play Store*. Model *Indonesian-RoBERTa* diimplementasikan secara langsung dalam proses *scraping* agar data dapat dianalisis secara *real-time* dan ditampilkan melalui antarmuka dashboard.

5. Pengujian Sistem

Evaluasi dilakukan terhadap akurasi klasifikasi model, keandalan sistem *scraping*, serta aspek *usability dashboard*. Tahap ini memastikan sistem berjalan sesuai dengan tujuan dan memberikan hasil yang dapat diandalkan.

6. Penarikan Kesimpulan

Berdasarkan hasil pengujian, disusun kesimpulan terkait temuan utama, kontribusi penelitian, serta rekomendasi untuk pengembangan lebih lanjut. Kesimpulan ini menjadi dasar dalam penyusunan laporan ilmiah dan diseminasi hasil penelitian.

3.2 Metode Pengumpulan Data

Pengumpulan data merupakan tahap penting dalam penelitian ini karena kualitas data secara langsung mempengaruhi kinerja model yang digunakan. Penelitian ini menggunakan pendekatan berbasis *pre-trained model Indonesian-RoBERTa* yang telah dilatih sebelumnya menggunakan dataset khusus untuk bahasa Indonesia. Selain itu,

dilakukan juga *fine-tuning* dengan dataset tambahan agar model lebih adaptif terhadap konteks data yang relevan dengan kasus analisis sentimen.

1. Dataset *Finetuning*

Dataset yang digunakan untuk *fine-tuning* diperoleh dari platform publik Kaggle. Dataset ini terdiri dari 23.645 tweet berbahasa Indonesia yang dikumpulkan selama masa PPKM. Data bersifat sekunder dan telah memiliki label sentimen (positif, netral, atau negatif), namun masih dalam bentuk mentah dan memerlukan proses pembersihan sebelum digunakan untuk pelatihan model.

2. Dataset Pelatihan Model *Indonesian-RoBERTa*

Model *Indonesian-RoBERTa* merupakan model *pre-trained* dengan arsitektur *RoBERTa* dan memiliki 124 juta parameter. Model ini dilatih menggunakan dataset SmSA (*Sentiment and Subjectivity Analysis*), yang terdiri dari opini-opini berbahasa Indonesia yang telah dilabeli sentimen. Dataset ini secara luas digunakan dalam riset NLP Indonesia karena mewakili ekspresi sentimen dalam bahasa sehari-hari dengan cukup baik.

Dataset SmSA menjadi sumber utama dalam pelatihan model dasar, sedangkan dataset tweet dari Kaggle digunakan untuk *fine-tuning* guna menyesuaikan model dengan domain data yang lebih spesifik.

3.3 Metode Perancangan

Penelitian ini terdiri dari beberapa tahapan, mulai dari preprocessing data,

pembagian dataset, tokenisasi, pelatihan dan evaluasi model *Indonesian-RoBERTa*. Model yang finetuning juga akan diimplementasikan kedalam *dashboard*, dengan perancangan sistem website yang dijelaskan melalui diagram UML untuk menggambarkan alur proses dan struktur komponen sistem.

1. *Preprocessing Data*

Data mentah berupa *tweet* PPKM sebanyak 23.671 data dibersihkan melalui proses seperti penghapusan URL, mention, simbol, tanda baca, dan spasi berlebih. Teks juga dikonversi ke huruf kecil. Nama kolom disesuaikan, duplikasi dihapus, dan baris kosong dihilangkan untuk memastikan data bersih dan konsisten.

2. Perhitungan Bobot Kelas

Karena distribusi kelas tidak seimbang, bobot kelas dihitung menggunakan rumus berbasis proporsi jumlah data tiap kelas. Bobot ini digunakan saat pelatihan model agar kesalahan pada kelas minoritas mendapatkan penalti lebih besar, sehingga model tidak bias terhadap kelas mayoritas.

3. Pembagian Dataset

Dataset dibagi menjadi data latih (81%), validasi (9%), dan uji (10%). Data latih digunakan untuk pelatihan model, data validasi untuk pemantauan performa selama *training*, dan data uji untuk evaluasi akhir terhadap kemampuan generalisasi model.

4. Tokenisasi

Teks yang telah dibersihkan diubah menjadi token menggunakan *tokenizer* BPE milik *RoBERTa*. Tokenisasi ini bertujuan memecah teks menjadi unit-unit kecil yang dapat dikonversi ke representasi numerik sebelum diproses oleh model.

5. Proses Training (Fine-tuning)

Model *w11wo/Indonesian-RoBERTa-base-sentiment-classifier* digunakan sebagai model dasar dan dilatih ulang menggunakan dataset yang telah disiapkan. Lapisan klasifikasi disesuaikan untuk 3 kelas, dan *loss function* dimodifikasi dengan bobot kelas melalui implementasi *WeightedTrainer*.

6. Evaluasi Model

Model yang telah dilatih dievaluasi menggunakan data uji dengan metrik *precision*, *recall*, dan *F1-score*, baik per kelas maupun secara agregat (*macro*, *micro*, dan *weighted*). Evaluasi ini mengukur performa model dalam mengklasifikasikan sentimen secara menyeluruh dan adil.

7. Analisis Prediksi Model

Prediksi model dianalisis melalui proses transformasi dari logits ke probabilitas menggunakan fungsi *softmax*, lalu ditentukan kelas akhir berdasarkan nilai probabilitas tertinggi (MAP). Analisis ini memberikan pemahaman tentang cara model membuat keputusan klasifikasi.

8. Confusion Matrix

Confusion matrix digunakan untuk mengidentifikasi distribusi kesalahan prediksi model. Matriks ini menunjukkan jumlah data dari tiap kelas aktual yang diprediksi ke kelas tertentu, sehingga dapat dianalisis bagian mana yang perlu perbaikan.

9. Perancangan UML

Perancangan sistem dalam penelitian ini divisualisasikan menggunakan *Unified Modeling Language* (UML) untuk menggambarkan struktur dan alur kerja sistem. Empat diagram UML digunakan, yaitu *Use Case Diagram* untuk memodelkan interaksi pengguna dan

sistem, *Activity Diagram* untuk menunjukkan alur aktivitas pengguna, *Sequence Diagram* untuk menggambarkan komunikasi antar objek, serta *Class Diagram* untuk merepresentasikan struktur dan relasi antar komponen sistem. Visualisasi ini bertujuan memberikan gambaran sistem yang terstruktur dan sistematis.

HASIL DAN PEMBAHASAN

4.1 Hasil Penelitian

Penelitian ini bertujuan mengembangkan sistem *dashboard* untuk analisis sentimen ulasan pengguna aplikasi Instagram di *Google Play Store* dengan memanfaatkan model transformer *Indonesian-RoBERTa* yang telah melalui proses fine-tuning. Tahapan eksperimen dilakukan dengan fine-tuning model *w11wo/Indonesian-RoBERTa-base-sentiment-classifier* menggunakan dataset berbahasa Indonesia yang terdiri atas 23.645 *tweet* dari masa PPKM, yang telah dilabeli ke dalam tiga kelas sentimen: positif, netral, dan negatif.

1. Hasil Finetuning Model

Selama proses pelatihan, model dilatih menggunakan pendekatan *class weighting* untuk mengatasi ketidakseimbangan distribusi label. Evaluasi dilakukan menggunakan metrik *accuracy*, *precision*, *recall*, dan *F1-score*, serta divisualisasikan dalam *confusion matrix*. Tabel berikut menunjukkan hasil pelatihan selama 4 *epoch*:

Tabel 1. Hasil Training Model

Epoch.	Training Loss	Validation Loss	Accuracy	F1 Macro
1	0.4412	0.4509	85.37%	0.7707
2	0.3069	0.5550	88.86%	0.8091
3	0.1748	0.7882	89.48%	0.8195
4	0.1370	0.9267	89.01%	0.8109

(Sumber: Data Penelitian, 2025)

Berdasarkan hasil di atas, terlihat bahwa nilai *training loss* mengalami penurunan konsisten di setiap *epoch*, menunjukkan bahwa model terus belajar dari data latih. Nilai *accuracy* dan *F1-score* secara umum menunjukkan peningkatan performa, di mana nilai *F1-score* tertinggi dicapai pada *epoch* ke-3 sebesar 0.8195.

2. Hasil analisis prediksi

Evaluasi prediksi dilakukan untuk mengamati proses klasifikasi oleh model

Indonesian-RoBERTa terhadap teks bahasa Indonesia. Contoh yang digunakan adalah kalimat: "Aplikasi ini sangat berguna dan saya suka sekali!". Proses prediksi dianalisis mulai dari keluaran *logits*, transformasi *softmax*, hingga pemilihan kelas akhir menggunakan metode *Maximum A Posteriori (MAP)*.

Tabel 2. Ringkasan Hasil Prediksi Model

Tahap.	Input	Output	Interpretasi
<i>Logits</i>	Teks: "Aplikasi ini sangat berguna dan saya suka sekali!"	[6.3, -3.0, -3.1]	Model memberikan skor tertinggi pada kelas positif (6.3), rendah pada lainnya
<i>Softmax</i>	<i>Logits</i> : [6.3, -3.0, -3.1]	[0.9998, 0.000091, 0.000082]	Probabilitas tertinggi pada kelas positif (99.98%)
<i>MAP Decision</i>	Probabilitas hasil <i>softmax</i>	Kelas 0 (Positif) – 99.98%	Model memutuskan kelas positif sebagai hasil akhir

(Sumber: Penelitian, 2025)

Dari hasil tersebut, model menunjukkan tingkat keyakinan yang sangat tinggi dalam mengklasifikasikan teks sebagai positif, dengan probabilitas 99.98%. Perbedaan skor *logits* yang mencolok antara kelas positif dan dua kelas lainnya

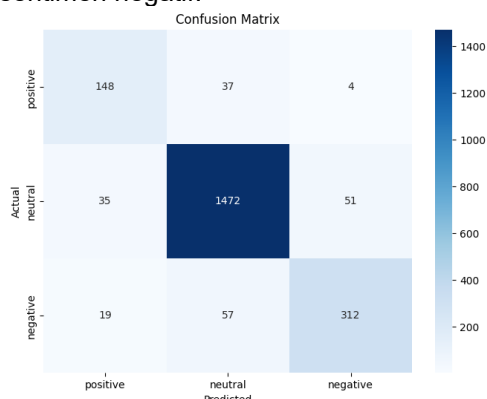
menunjukkan bahwa model memiliki *decision boundary* yang tajam dan mampu menangkap konteks positif dari kalimat input. Nilai *confidence* yang tinggi ini mengindikasikan bahwa model telah menginternalisasi pola bahasa dan

ekspresi sentimen positif dalam data pelatihan dengan baik.

Secara keseluruhan, hasil ini memperkuat temuan bahwa proses *fine-tuning* pada model Indonesian-*RoBERTa* berhasil meningkatkan sensitivitas model terhadap ekspresi positif dalam bahasa Indonesia, serta menunjukkan kapabilitas model dalam memberikan prediksi yang presisi dan terukur dalam konteks klasifikasi sentimen.

3. Confusion Matrix

Confusion matrix menunjukkan bahwa model memiliki tingkat prediksi yang tinggi untuk kelas positif dan netral, tetapi relatif lebih rendah dalam mengklasifikasikan sentimen negatif.



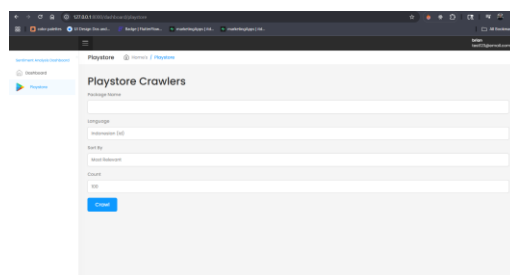
Gambar 3. Confusion Matrix
(Sumber: Data Penelitian, 2025)

Berdasarkan *confusion matrix*, model menunjukkan performa terbaik pada klasifikasi kelas netral dengan 1.472 prediksi benar. Kelas negatif memiliki akurasi yang baik dengan 312 prediksi benar namun terdapat 76 kesalahan klasifikasi. Kelas positif menunjukkan performa terendah dengan 148 prediksi benar dan kecenderungan salah diklasifikasikan sebagai netral (37 kasus). Model menunjukkan bias terhadap kelas

netral, mengindikasikan kesulitan dalam memisahkan fitur kelas positif dan negatif.

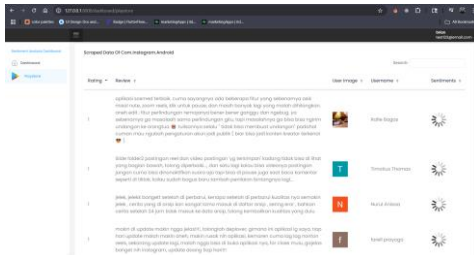
4. Implementasi model

Setelah model selesai di *finetuning* dan di evaluasi, model kemudian di implementasikan kedalam sebuah web *dashboard*. *Dashboard* ini dirancang untuk mengambil data ulasan dengan teknik *web scraping*, mengklasifikasi dan menampilkan hasil klasifikasi sentimen secara informatif, sehingga memudahkan pengguna dalam memahami opini publik secara cepat dan efisien.



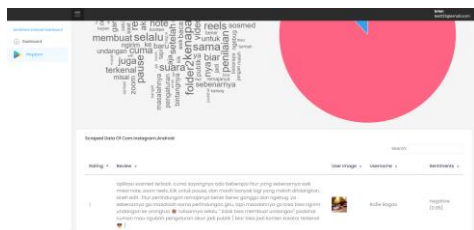
Gambar 4. Halaman Klasifikasi Dashboard
(Sumber: Data Penelitian, 2025)

Halaman *Playstore Crawlers* merupakan antarmuka input untuk melakukan proses *crawling* ulasan aplikasi dari *Google Play Store*. Pengguna dapat memasukkan nama paket aplikasi, memilih bahasa ulasan, metode penyortiran, serta jumlah ulasan yang diambil. Setelah parameter diisi, sistem akan mengeksekusi proses pengambilan data menggunakan pustaka *google_play_scraper*, yang selanjutnya digunakan dalam analisis sentimen.



Gambar 4. 5. Hasil data ulasan yang berhasil di ambil
(Sumber: Data Penelitian, 2025)

Halaman ini menampilkan data ulasan yang di-scrape dari aplikasi Instagram melalui *Google Play Store*. Setiap entri dalam tabel mencakup rating yang diberikan oleh pengguna, isi ulasan, gambar profil, nama pengguna, serta kolom *sentiments* yang akan menampilkan hasil analisis sentimen. Tampilan ini juga dilengkapi dengan fitur pencarian di bagian kanan atas untuk memudahkan pengguna dalam menyaring ulasan tertentu berdasarkan kata kunci.



Gambar 5. Hasil Klasifikasi Ulasan dan Visualisasi Data
(Sumber: Data Penelitian, 2025)

Halaman ini menampilkan hasil klasifikasi sentimen dari data ulasan yang telah dikumpulkan, lengkap dengan visualisasi dalam bentuk *word cloud* dan diagram lingkaran. *Word cloud* menampilkan frekuensi kata yang sering muncul dalam ulasan, sedangkan *pie chart* menunjukkan distribusi label sentimen (positif, netral,

negatif). Di bagian bawah, terdapat tabel hasil klasifikasi yang mencantumkan isi ulasan, nama pengguna, gambar profil, label sentimen, serta skor probabilitas hasil prediksi dari model *IndoRoBERTa*. Fitur pencarian juga disediakan untuk memudahkan pengguna dalam menelusuri ulasan tertentu.

SIMPULAN

Penelitian ini berhasil mengimplementasikan model *Indonesian RoBERTa* ke dalam sistem *dashboard* analisis sentimen berbasis web, yang mengotomatisasi proses pengambilan data, klasifikasi sentimen, dan visualisasi hasil. Data dikumpulkan secara otomatis dari *Google Play Store* menggunakan *google_play_scraper*, dan model *diffinetuning* dengan 23.645 *tweet* berbahasa Indonesia dari masa PPKM. Hasil evaluasi terhadap model menunjukkan performa klasifikasi yang baik, dengan nilai *micro-averaged precision*, *recall*, dan *F1-score* sebesar 0,9049. Berdasarkan *confusion matrix*, model mencapai total *True Positive* (TP) sebanyak 1932, serta *False Positive* (FP) dan *False Negative* (FN) masing-masing sebesar 203. Model menunjukkan akurasi tertinggi dalam mengklasifikasikan sentimen netral (*precision*: 0,95 dan *recall*: 0,94), namun masih terdapat kesalahan klasifikasi pada kelas positif dan negatif yang cenderung tertukar dengan kelas netral. Visualisasi hasil pada dashboard memungkinkan identifikasi pola sentimen yang lebih sistematis dan terukur terhadap ulasan pengguna aplikasi Instagram.

DAFTAR PUSTAKA

- Drewek-Ossowicka, Anna, Mariusz Pietrolaj, and Jacek Rumiński. 2021. "A Survey of Neural Networks Usage for Intrusion Detection Systems." *Journal of Ambient Intelligence and Humanized Computing* 12(1):497–514. doi: 10.1007/s12652-020-02014-x.
- Naseer, Muchammad, Jauzak Hussaini Windiatmaja, Muhamad Asvial, and Riri Fitri Sari. 2022. "RoBERTaEns: Deep Bidirectional Encoder Ensemble Model for Fact Verification." *Big Data and Cognitive Computing* 6(2):33. doi: 10.3390/BDCC6020033.
- Sekhar, Ch, and P. Sai Meghana. 2020. "A Study on Backpropagation in Artificial Neural Networks." *Asia-Pacific Journal of Neural Networks and Its Applications* 4(1):21–28. doi: 10.21742/ajnnia.2020.4.1.03.
- Sinha, Koustuv, Robin Jia, Dieuwke Hupkes, Joelle Pineau, Adina Williams, and Douwe Kiela. 2021. "Masked Language Modeling and the Distributional Hypothesis: Order Word Matters Pre-Training for Little." *EMNLP 2021 - 2021 Conference on Empirical Methods in Natural Language Processing, Proceedings* 2888–2913. doi: 10.18653/v1/2021.emnlp-main.230.
- Terven, Juan, Diana M. Cordova-Esparza, Alfonso Ramirez-Pedraza, Edgar A. Chavez-Urbiola, and Julio A. Romero-Gonzalez. 2023. "Loss Functions and Metrics in Deep Learning."
- Yudistira, Novanto. 2021. "Peran Big Data Dan Deep Learning Untuk Menyelesaikan Permasalahan Secara Komprehensif." *EXPERT: Jurnal Manajemen Sistem Informasi Dan Teknologi* 11(2):78. doi: 10.36448/expert.v11i2.2063.

	Penulis Pertama, Brian Natanael Nainggolan merupakan mahasiswa Prodi Teknik Informatika Universitas Putera Batam Mahasiswa yang aktif dalam bidang informatika
	Penulis kedua, Koko Handoko, S. Kom., M. Kom, yang merupakan Dosen Pembimbing Prodi Teknik Informatika Universitas Putera Batam. Penulis Aktif sebagai tenaga kerja dan peneliti