

Analisis Sentimen Topik #Kaburajadulu di Platform X Berbasis Model IndoBERTweet

Rivaldo Nugraha^{1*}, Shafira Febriani²

^{1,2} Program Studi Teknik Informatika, STMIK AMIK Bandung

*Corresponding Author: rivaldo@stmik-amikbandung.ac.id

Article History

Received: 25-6-2025

Revised: 29-7-2025

Published: 13-8-2025

Key Words:

Classification,
IndoBERTweet, Natural
Language processing,
Sentiment Analysis, X

Abstract: Social media has emerged as a prominent medium for individuals to express opinions and engage in discourse on various trending topics. These expressions often reflect diverse public sentiments, typically categorized as supportive (pro), opposing (contra), or neutral. Analysing these sentiment variations is essential for understanding the dominant public stance toward a particular issue. This study aims to perform sentiment analysis on a currently trending topic in Indonesia using IndoBERTweet, a transformer-based language model pre-trained on Indonesian Twitter data. The model is specifically optimized for processing informal language structures commonly found on social media platforms. The selected case study focuses on a widely discussed issue circulating under the hashtag #kaburajadulu on platform X (formerly Twitter). The primary objective is to determine the prevailing sentiment—whether the majority of users exhibit a pro, contra, or neutral stance on the topic. Data collection involved scraping tweets associated with the hashtag, followed by preprocessing and sentiment labelling. The IndoBERTweet model was then employed to classify the sentiments and assess performance using four standard evaluation metrics: accuracy, precision, recall, and F1-score. The experimental results indicate strong model performance, with all evaluation metrics exceeding 94%, suggesting high reliability in sentiment classification for Indonesian-language tweets. These findings underscore the effectiveness of transformer-based models in sentiment analysis tasks, particularly within the context of social media data. Moreover, the study contributes to the growing body of research that leverages natural language processing for real-time public opinion monitoring in the Indonesian digital landscape.

Kata Kunci:

Analisis Sentimen,
IndoBERTweet,
Klasifikasi, Natural
Language Processing, X.

Abstrak: Media sosial telah menjadi sarana utama bagi individu untuk menyampaikan pendapat dan terlibat dalam diskusi mengenai berbagai topik yang sedang tren. Ekspresi tersebut sering kali mencerminkan beragam sentimen publik, yang umumnya dikategorikan sebagai mendukung (pro), menentang (kontra), atau netral. Analisis terhadap variasi sentimen ini sangat penting untuk memahami sikap dominan masyarakat terhadap suatu isu tertentu. Penelitian ini bertujuan untuk melakukan analisis sentimen terhadap sebuah topik yang sedang tren di Indonesia dengan menggunakan IndoBERTweet, sebuah model bahasa berbasis transformer yang telah dilatih sebelumnya menggunakan data Twitter berbahasa Indonesia. Model ini secara khusus dioptimalkan untuk memproses struktur bahasa informal yang umum digunakan di platform media sosial. Studi kasus yang dipilih berfokus pada isu yang banyak diperbincangkan dengan tagar #kaburajadulu di platform X (sebelumnya Twitter). Tujuan utama dari penelitian ini adalah untuk menentukan sentimen yang paling dominan—apakah mayoritas pengguna menunjukkan sikap pro, kontra, atau netral terhadap topik tersebut. Pengumpulan data dilakukan dengan melakukan scraping terhadap tweet yang terkait dengan tagar tersebut, kemudian dilanjutkan dengan proses pra-pemrosesan dan pelabelan sentimen. Model IndoBERTweet digunakan untuk mengklasifikasikan sentimen, dan kinerjanya dievaluasi menggunakan empat metrik standar: akurasi, presisi, recall, dan F1-score. Hasil eksperimen menunjukkan performa model yang sangat baik, dengan semua metrik evaluasi melebihi 94%, yang menunjukkan tingkat keandalan tinggi dalam klasifikasi sentimen untuk tweet berbahasa Indonesia.



Pendahuluan

Kemunculan teknologi digital telah mengubah cara masyarakat dalam menyampaikan opini mengenai berbagai topik seperti politik, isu sosial, dan regulasi pemerintahan. Di media sosial, berita hangat—yang juga dikenal sebagai topik viral atau trending—disebarkan secara luas. Platform X, yang sebelumnya dikenal sebagai Twitter, merupakan salah satu sumber informasi internet paling populer untuk mengetahui berita dan topik yang sedang tren.

Saat ini, sedang ramai diperbincangkan di Indonesia mengenai fenomena warga negara Indonesia yang memilih untuk pindah atau meninggalkan tanah air. Tagar #kaburajadulu mencerminkan isu ini, sekaligus mengisyaratkan bahwa masyarakat Indonesia, khususnya yang berpendidikan tinggi atau memiliki keahlian khusus, merasa tidak mendapatkan penghargaan yang layak di dalam negeri. Banyak di antara mereka merasa kurang dihargai, tidak memperoleh layanan publik yang memadai, atau kesulitan mendapatkan pekerjaan yang layak. Di platform X, tagar ini memicu berbagai diskusi dengan pandangan yang beragam. Perbedaan sudut pandang tersebut tentu merefleksikan sikap pengguna platform X, khususnya dari kalangan masyarakat Indonesia. Salah satu pendekatan yang dapat digunakan untuk menganalisis sikap publik terhadap isu ini adalah analisis sentimen. Analisis sentimen bertujuan untuk menentukan apakah suatu opini bernada positif, negatif, atau netral. Oleh karena itu, penting untuk mengamati proporsi atau sebaran statistik sentimen masyarakat di platform X terkait isu ini.

Penelitian terkini banyak berfokus pada pengembangan dan penerapan model bahasa untuk teks media sosial berbahasa Indonesia, khususnya dari Twitter. Salah satu model yang dikembangkan adalah IndoBERTweet, sebuah model yang dilatih sebelumnya (*pretrained*) secara khusus untuk data Twitter berbahasa Indonesia dengan memperluas model BERT monolingual Bahasa Indonesia menggunakan kosakata yang disesuaikan untuk domain media sosial (Koto, 2021). IndoBERTweet juga terbukti efektif dalam mengidentifikasi bahasa pada tweet campuran Indonesia-Jawa-Inggris, mengungguli metode lain seperti BLSTM dan CRF (Hidayatullah, 2023). Keberhasilan IndoBERTweet dalam berbagai aplikasi ini disebabkan oleh kemampuannya memahami konteks kata dalam urutan teks serta inisialisasi kosakata domain-spesifik yang efisien, yang membuat proses pretraining lima kali lebih cepat dan efektif (Koto, 2021).

Beberapa penelitian sebelumnya yang menggunakan IndoBERTweet dalam analisis sentimen antara lain oleh Setiawan (2023) melakukan analisis sentimen terhadap ulasan pengguna TikTok, dan hasilnya menunjukkan bahwa model IndoBERTweet memiliki performa lebih baik dibandingkan model LSTM berdasarkan empat metrik utama klasifikasi. Selanjutnya, Fadhel dan Maharani (2024) melakukan klasifikasi gejala depresi dari data platform X menggunakan IndoBERTweet, dengan akurasi terbaik mencapai 82%. Kusuma dan Andry (2023) meneliti deteksi ujaran kebencian di media sosial menggunakan pendekatan deep learning. Untuk membedakan ujaran kebencian dari kebebasan berpendapat, mereka menggabungkan model IndoBERTweet dengan lapisan BiLSTM, dan berhasil mencapai akurasi sebesar 93,7%, meningkat dari penelitian sebelumnya dalam klasifikasi ujaran kebencian.

Samosir dan Riyaldi (2024) mengembangkan model berbasis IndoBERT untuk klasifikasi sentimen komentar TikTok tentang pemilihan presiden Indonesia. Proses penelitian meliputi praproses data, pelabelan, pelatihan, validasi, dan pengujian. Indriani

(2024) meneliti klasifikasi multilabel terhadap umpan balik mahasiswa dalam bahasa informal untuk mengidentifikasi aspek-aspek yang dikritik. Dengan membandingkan tiga model BERT—IndoBERT, IndoBERTtweet, dan mBERT—studi ini menunjukkan bahwa IndoBERTtweet memiliki performa terbaik (macro F1: 0,8462) pada panjang input 64 token dengan metode pemotongan di akhir, sehingga menjadikannya model paling sesuai untuk analisis otomatis komentar mahasiswa berbahasa Indonesia.

Meskipun penelitian sebelumnya telah membuktikan keunggulan IndoBERTtweet pada berbagai domain, terdapat celah penelitian yang jelas. Penerapannya untuk menganalisis sentimen publik terkait fenomena sosial spesifik seperti tren "brain drain" atau keinginan warga negara untuk pindah, yang direpresentasikan oleh tagar #KaburAjaDulu, merupakan area yang belum dieksplorasi. Studi-studi terdahulu belum secara khusus mengkaji bagaimana model ini dapat menangkap nuansa sentimen pro, kontra, dan netral pada diskursus mengenai kepuasan warga negara dan aspirasi untuk mencari kehidupan yang lebih baik di luar negeri.

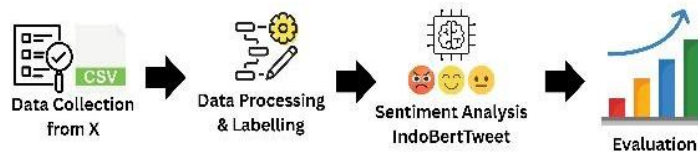
Oleh karena itu, penelitian ini tidak hanya bertujuan menguji efektivitas teknis model IndoBERTtweet pada korpus data yang unik, tetapi juga memiliki implikasi praktis yang signifikan. Bagi pemangku kepentingan seperti pemerintah, hasil analisis ini dapat menjadi masukan berharga untuk memahami akar ketidakpuasan publik dan merumuskan kebijakan yang lebih responsif terkait ketenagakerjaan, layanan publik, dan apresiasi talenta lokal. Bagi peneliti sosial, studi ini menawarkan metode kuantitatif untuk memetakan dinamika opini publik secara real-time, memberikan gambaran terukur mengenai sikap masyarakat terhadap isu krusial yang diangkat oleh tagar tersebut. Evaluasi performa model akan dilakukan menggunakan metrik akurasi, presisi, recall, dan F1-score.

Metode Penelitian

Analisis sentimen telah menjadi bidang penelitian penting dalam pemrosesan bahasa alami (Natural Language Processing/NLP) dan penambangan teks, dengan berbagai metode dan aplikasi yang telah dikembangkan selama bertahun-tahun. Pendekatan berbasis leksikal dan pembelajaran mesin terawasi (supervised machine learning) merupakan teknik umum yang digunakan untuk mengukur sentimen (Ribeiro, 2015). Proses analisis sentimen umumnya mencakup tahap praproses data, ekstraksi fitur, dan klasifikasi untuk menentukan polaritas opini (Wankhede, 2018). Kemajuan terbaru difokuskan pada peningkatan akurasi serta penanganan tantangan komputasional dalam NLP (Rastogi, 2021). Dalam hal ini, leksikon sentimen memainkan peran penting, dengan berbagai jenis yang berbeda dalam hal cakupan, metode pengembangan, dan tingkat kedetailan (granularitas) (Lahase, 2022). Alat-alat ini sangat bermanfaat bagi dunia bisnis, pemerintah, maupun individu dalam menganalisis opini terkait produk, layanan, dan isu sosial yang tersebar di berbagai platform seperti media sosial, situs ulasan, dan artikel berita (Ribeiro, 2015. Rastogi, 2021).

Bagian ini secara sistematis menjelaskan tahapan-tahapan yang dilakukan dalam penelitian analisis sentimen terhadap tagar #KaburAjaDulu di platform X menggunakan model IndoBERTtweet. Dengan memanfaatkan model ini, penelitian dapat mengklasifikasikan apakah opini yang diungkapkan oleh warga negara Indonesia atau pengguna platform X secara umum mengandung sentimen positif, negatif, atau netral

(Sainger, 2021). Selain itu, opini tersebut juga dapat diinterpretasikan dalam kategori pro, kontra, atau netral terhadap isu yang sedang diperbincangkan.



Gambar 1. Alur Sistem

Berdasarkan Gambar 1, penelitian ini dimulai dengan proses pengumpulan data dari platform X yang berisi opini publik, khususnya dari warga Indonesia, terkait topik #kaburajadulu. Data yang diperoleh disimpan dalam format CSV. Setelah data terkumpul, langkah selanjutnya adalah pemrosesan data (*data preprocessing*), seperti mengubah seluruh teks menjadi huruf kecil (*lowercasing*), menghapus karakter atau angka yang tidak diperlukan, menghapus kata-kata atau singkatan yang tidak baku, serta melakukan penyeimbangan kelas agar distribusi data seimbang. Setelah data bersih, dilakukan penyesuaian lanjutan (*finetuning*) model IndoBERTweet untuk mengarahkan model agar dapat mengklasifikasikan sentimen publik terhadap topik #kaburajadulu ke dalam tiga kategori: pro, kontra, atau netral. Hasil klasifikasi ini kemudian dievaluasi menggunakan metrik evaluasi seperti akurasi, precision, recall, dan F1-score.

Pengumpulan Data

Pengumpulan data dilakukan dengan menghimpun opini publik yang disimpan dalam file CSV. Opini-opini ini telah diberi label berdasarkan kategorisasi sentimen, yaitu pro, kontra, atau netral. Format data saat pengumpulan dari platform X ditampilkan pada Gambar 2.

	id	username	text	date
0	b69f9edc-d821-497c-9ccc-44ed94078659	lhateUall	satu satunya copss yang dapat gua percaya adal...	2025-03-23
1	315c414a-1766-493c-a0c6-e6a1a8338551	reey	pengen bgt #KaburAjaDulu	2025-03-23
2	262e1430-bf31-4625-86ec-ba101b9fc0bd	nanayo	Ayo kita usahakan #KaburAjaDulu itu	2025-03-23
3	94bbf54f-96af-4e72-92e6-69cc3a5176d6	weirdozz	kapan yh #KaburAjaDulu	2025-03-23
4	81f203a8-e5e5-4a2c-87b2-f657ea243826	sei	@yuuyufever ayo realisasikan #KaburAjaDulu htt...	2025-03-23
...	...	...	...	...
10290	be55abfc-bff4-40dc-99c0-0cb92c76136f	Whoareyou	@asumsico Tetaplah bodoh jangan pintar #KaburA...	2025-02-20
10291	c838c3dd-d937-486a-b9d2-0c0a496451ed	mimin	Buka apk defcal, kalo bb udah lebih 50. Artiny...	2025-02-20
10292	5e0f0595-3380-47e8-bcdb-5ff83cbb1651	nanæ.	@zaarahyt HAHAAHAHJSJSJS benerrr, sumpek bang...	2025-02-20
10293	505fb272-c376-4760-a340-9ab985120d17	Cokusi Kameumeut	ada yang mau ikutan quiz Puzzle bareng Minsi? ...	2025-02-20
10294	03ead38a-7f5e-4f2e-9607-cc0abce386f9	.	Ga ngerti subtansinya berarti. Hastag #KaburAj...	2025-02-20

10295 rows x 4 columns

Gambar 2. Format Dataset

Berdasarkan Gambar 2, format dataset tersebut berisi opini publik dari tagar yang dikumpulkan di platform X. Dataset ini terdiri dari beberapa kolom, yaitu “id”, “username”, “text”, dan “date”. Data ini selanjutnya akan diproses agar dapat digunakan dalam model analisis sentimen.

Pengolahan Data

Pemrosesan data merupakan tahapan penting karena model membutuhkan input data yang bersih agar performa hasil klasifikasinya tidak terganggu. Proses ini mencakup beberapa langkah berikut:

a. *Lower Casing*

Pada tahap awal, semua karakter dalam teks diubah menjadi huruf kecil. Proses ini bertujuan untuk mengurangi redundansi data akibat perbedaan kapitalisasi, sehingga kata seperti “Bagus” dan “bagus” dapat dikenali sebagai entitas yang sama (Musfiroh, 2021. Rahmi dan Dari, 2024).

b. Penggantian Mention Username

Ketika terdapat mention terhadap nama pengguna seperti *@username*, bagian tersebut diganti dengan token generik *@USER*. Tujuannya adalah untuk menghapus elemen spesifik yang mengacu pada identitas individu, sehingga fokus model tetap pada konteks kalimat, bukan pada nama pengguna.

c. Penggantian URL

Seluruh tautan (link) yang terdapat dalam teks akan digantikan dengan token seperti *HTTPURL*. Hal ini dilakukan agar isi spesifik dari URL tidak memengaruhi pemodelan konteks teks (Jianqiang, 2017).

d. Konversi Emoji ke Teks

Dalam tweet, pengguna sering menggunakan emoji untuk mengekspresikan pendapat. Oleh karena itu, emoji dalam teks perlu dikonversi menjadi representasi teks atau deskripsi maknanya. Tujuannya adalah agar model dapat memahami ekspresi emosional dalam emoji secara efektif, sehingga meningkatkan pemahaman konteks dalam data teks media sosial (Wankhede, 2018).

e. Normalisasi Spasi (*Whitespace Normalization*)

Langkah ini menghapus spasi berlebih dan menggabungkan spasi yang tidak perlu, sehingga teks menjadi lebih rapi dan mudah diproses (Garg dan Sharma, 2022).

f. Penyeimbangan Kelas (*Class Balancing*)

Ketidakeimbangan kelas merupakan permasalahan umum dalam klasifikasi dengan menggunakan machine learning atau deep learning. Dalam hal ini, label sentimen (pro, netral, dan kontra) sering kali tidak terdistribusi secara merata. Algoritma pembelajaran mesin pada umumnya mengasumsikan distribusi data yang seimbang antar kelas. Ketidakeimbangan dapat menyebabkan model lebih condong untuk memprediksi kelas yang dominan. Oleh karena itu, proses penyeimbangan kelas sangat penting untuk memastikan model mampu mengenali seluruh kelas secara adil dan akurat. Beberapa strategi telah diajukan dalam literatur untuk menangani masalah ini (Jadhav, 2022. Rogic, 2021).

Analisis Sentimen

Model IndoBERTweet merupakan adaptasi dari model dasar IndoBERT. IndoBERT sendiri adalah model bahasa berbasis arsitektur transformer yang dikembangkan khusus untuk memahami nuansa dan struktur bahasa Indonesia. Kemampuan pemrosesan teks secara bidirectional memungkinkan model ini memahami konteks kalimat secara menyeluruh, baik dari sisi kiri maupun kanan kata yang dianalisis.

IndoBERT dilatih menggunakan korpus bahasa Indonesia yang besar, sehingga mampu menangkap makna dan hubungan semantik dalam teks secara mendalam. Oleh karena itu, IndoBERT sangat efektif dalam tugas-tugas seperti klasifikasi sentimen karena kemampuannya memahami konteks yang kaya dan relevan (Fachry, 2025).

Model IndoBERTweet merupakan model pralatih yang dikembangkan dengan memperluas arsitektur BERT untuk secara khusus menangani karakteristik unik bahasa Indonesia di media sosial seperti Twitter. Karena dilatih menggunakan data tweet berbahasa Indonesia, model ini dapat memahami pola bahasa informal, akronim, serta struktur kalimat yang tidak baku yang sering muncul di platform digital. Oleh karena itu, IndoBERTweet sangat unggul dalam tugas-tugas pemrosesan bahasa alami pada data media sosial, seperti deteksi opini publik dan analisis sentimen (Khairani, 2024. Subarkah, 2022).

Parameter Evaluasi

Untuk menilai performa model, digunakan empat metrik klasifikasi standar, yaitu: akurasi, precision, recall, dan F1-score (Vujivic, 2021. Rainio, 2024. Muller, 2022). Keempat metrik ini memberikan gambaran menyeluruh mengenai seberapa baik model dalam mengenali masing-masing kelas sentimen, terutama dalam konteks data yang tidak seimbang (imbalanced data).

a. Akurasi

Akurasi mengukur proporsi prediksi yang benar terhadap jumlah keseluruhan data.

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

Diketahui bahwa:

- TP (*True Positive*): Jumlah data yang benar-benar termasuk dalam kelas positif dan diprediksi benar oleh model.
- TN (*True Negative*): Jumlah data yang benar-benar termasuk dalam kelas negatif dan juga diprediksi benar oleh model.
- FP (*False Positive*): Jumlah data yang sebenarnya bukan kelas positif, tetapi salah diprediksi sebagai positif oleh model.
- FN (*False Negative*): Jumlah data yang sebenarnya termasuk kelas positif, tetapi salah diprediksi sebagai negatif oleh model.

b. Presisi

Presisi mengukur proporsi prediksi positif yang benar terhadap semua prediksi positif yang dilakukan oleh model.

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

Presisi yang tinggi menunjukkan bahwa model jarang salah dalam memberikan label positif pada data yang tidak sesuai.

c. Recall

Recall mengukur proporsi data positif yang berhasil dikenali dengan benar oleh model.

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

Recall yang tinggi menunjukkan bahwa sebagian besar data yang seharusnya masuk dalam kelas positif berhasil ditangkap oleh model.

d. *F1 Score*

F1-score adalah rata-rata harmonik dari *precision* dan *recall*. Metrik ini berguna ketika kedua metrik tersebut tidak seimbang, karena *F1-score* memberikan nilai agregat yang memperhitungkan keduanya secara proporsional (Riyanto, 2023).

$$F1\ score = \frac{Precision \times Recall}{Precision + Recall} \quad (4)$$

F1-score memberikan nilai yang seimbang ketika *precision* tinggi tetapi *recall* rendah, atau sebaliknya.

Hasil dan Pembahasan

Hasil Pengolahan Data

Bagian ini menampilkan hasil dari proses pra-pemrosesan data mentah yang berasal dari platform X (berisi opini publik) dan disimpan dalam file berformat CSV, hingga data tersebut siap digunakan dalam pelatihan model. Gambar berikut menunjukkan hasil dari proses pra-pemrosesan data.

	text	sentiment_label
0	satu satunya copss yang dapat gua percaya adal...	LABEL_1
1	pengen bgt #KaburAjaDulu	LABEL_2
2	Ayo kita usahakan #KaburAjaDulu itu	LABEL_2
3	kapan yh #KaburAjaDulu	LABEL_2
4	@yuuyufever ayo realisasikan #KaburAjaDulu htt...	LABEL_1
...	...	...
10290	@asumsico Tetaplah bodoh jangan pintar #KaburA...	LABEL_2
10291	Buka apk defcal, kalo bb udah lebih 50. Artiny...	LABEL_1
10292	@zaarahyt HAHAAHAJSJSJS benerrr, sumpek bang...	LABEL_2
10293	ada yang mau ikutan quiz Puzzle bareng Minsi? ...	LABEL_1
10294	Ga ngerti subtansinya berarti. Hastag #KaburAj...	LABEL_2

10295 rows x 2 columns

Gambar 3. Format Dataframe setelah Pemrosesan Data

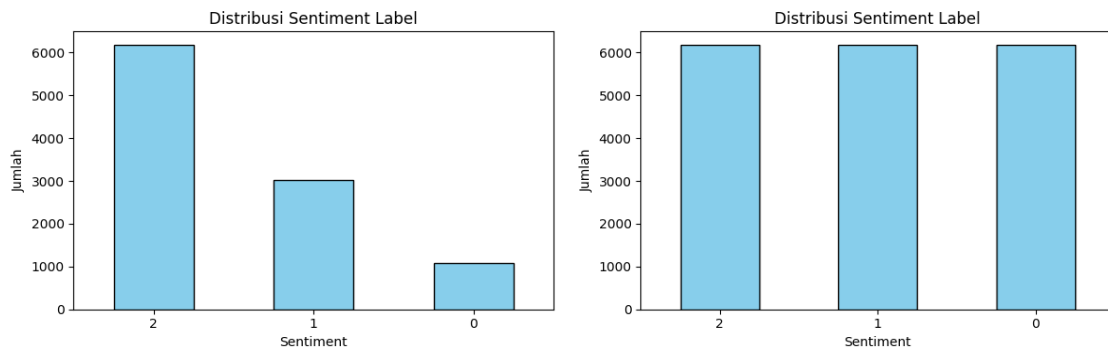
Seperti yang ditunjukkan pada Gambar 3, data yang sebelumnya terdiri atas empat kolom, yakni id, username, text, dan date, telah mengalami penyederhanaan. Setelah pemrosesan, hanya kolom opini (text) dan label sentimen yang dipertahankan karena merupakan elemen yang paling penting. Label sentimen dikategorikan sebagai berikut:

LABEL_0: Netral

LABEL_1: Positif (pro)

LABEL_2: Negatif (kontra)

Semua tahapan pemrosesan dari bagian 2.2 telah diterapkan pada kolom opini. Tujuan dari pembentukan dataset pasca pemrosesan ini adalah untuk meningkatkan kualitas dan konsistensi input, sehingga model dapat berlatih dan melakukan prediksi dengan lebih baik.



Gambar 4. Distribusi Kelas: (a) Sebelum (b) Setelah Proses Penyeimbangan

Selain pada kolom opini, analisis juga dilakukan terhadap distribusi label atau kelas pada keseluruhan data. Ketidakseimbangan kelas (imbalanced class) dalam klasifikasi dapat menyebabkan model cenderung memprediksi kelas mayoritas karena secara statistik memberikan nilai akurasi tinggi, meskipun tidak benar-benar memahami pola dari kelas minoritas. Akibatnya, performa model dalam mengenali kelas minoritas menjadi buruk.

Untuk mengatasi masalah ini, dilakukan teknik oversampling pada kelas minoritas, seperti yang terlihat pada Gambar 4. Proses ini memastikan distribusi data menjadi lebih seimbang sehingga model dapat belajar secara adil terhadap semua kelas.

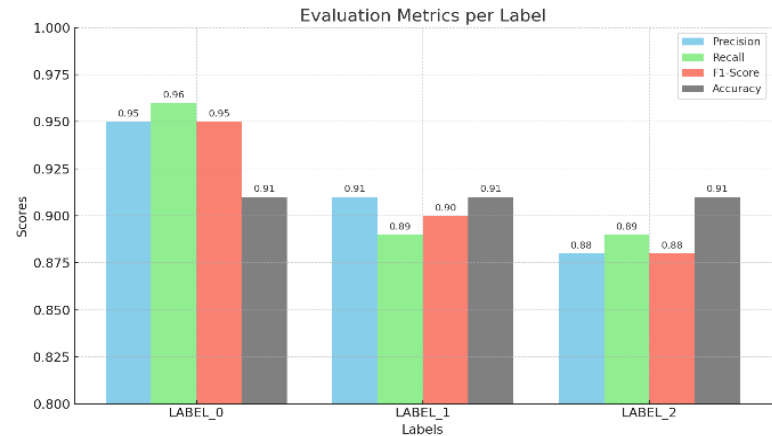
Performa Model

Bagian ini menampilkan hasil evaluasi model dalam mengklasifikasikan opini ke dalam masing-masing kelas sentimen (pro, netral, dan kontra) dengan menggunakan sejumlah metrik evaluasi penting seperti akurasi, presisi, recall, dan F1-score. Tujuan dari analisis ini adalah untuk memberikan gambaran menyeluruh mengenai seberapa baik model memahami dan memprediksi opini pengguna di media sosial.



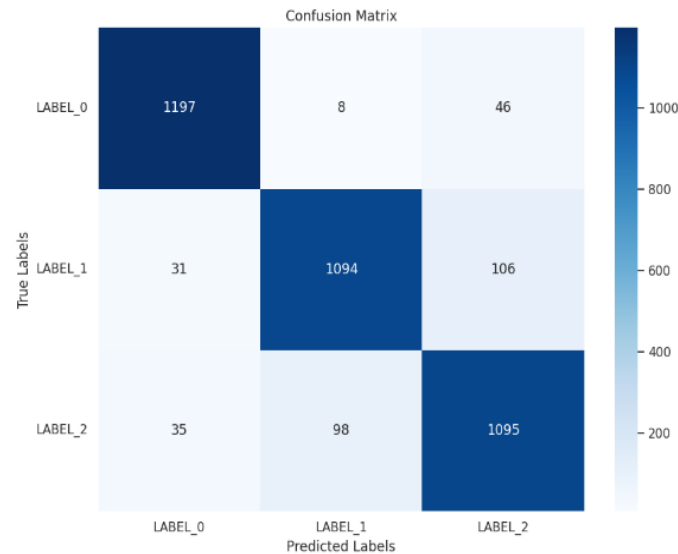
Gambar 5. Pemantauan Performa pada Setiap Iterasi Pelatihan

Seiring bertambahnya jumlah langkah pelatihan (training steps), keempat metrik evaluasi menunjukkan peningkatan yang konsisten, menandakan bahwa model semakin baik dalam mengenali pola dari data. Performa model mulai stabil setelah sekitar 500 langkah pelatihan, ditunjukkan dengan nilai metrik yang cenderung konvergen. Namun, terdapat penurunan sementara di sekitar langkah ke-1100, yang kemungkinan disebabkan oleh noise atau data validasi yang tidak representatif. Meski demikian, model dengan cepat kembali ke performa sebelumnya.



Gambar 6. Klasifikasi

Selain pemantauan selama pelatihan, performa model juga dievaluasi berdasarkan laporan klasifikasi yang ditampilkan dalam bentuk diagram batang pada Gambar 6. Hasil ini menunjukkan bahwa meskipun proses resampling berhasil menghasilkan distribusi label yang lebih seimbang, model masih mengalami kesulitan dalam mengenali kelas dengan fitur yang tidak terlalu menonjol, terutama pada LABEL_1 dan LABEL_2.



Gambar 7. Confusion Matrix

Confusion matrix menunjukkan bahwa model memiliki tingkat akurasi yang tinggi dalam memprediksi setiap label. Namun, terdapat beberapa kasus misklasifikasi

yang perlu ditinjau lebih lanjut, khususnya pada LABEL_2 yang seringkali diprediksi salah sebagai LABEL_1. Pola ini mengindikasikan bahwa kedua kelas tersebut memiliki kemiripan karakteristik, dan bisa menjadi fokus pengembangan lanjutan untuk meningkatkan ketepatan klasifikasi. Hasil akurasi 94.2% mendukung hipotesis bahwa IndoBERTweet efektif untuk analisis sentimen bahasa informal

Kesimpulan

Penelitian ini membuktikan kemampuan luar biasa dari model IndoBERTweet dalam melakukan analisis sentimen secara mendalam (*fine-grained*) terhadap wacana media sosial berbahasa Indonesia. Dengan berfokus pada sentimen publik terhadap tagar #KaburAjaDulu, model ini berhasil mengklasifikasikan opini ke dalam tiga kategori yang telah ditentukan: pro, netral, dan kontra. Hasil empiris yang diperoleh sangat meyakinkan: model mencapai nilai akurasi, presisi, *recall*, dan F1-score di atas 94%, yang menunjukkan kinerja klasifikasi yang sangat baik.

Secara metodologis, pelatihan model menunjukkan konvergensi yang stabil. Meskipun terdapat fluktuasi sementara pada nilai validation loss, hal tersebut diasumsikan sebagai akibat dari noise data spesifik per batch, bukan disebabkan oleh ketidakstabilan sistemik pada model. Temuan analisis utama menunjukkan bahwa pada tahap awal, model mengalami kesulitan dalam membedakan antara kelas netral (LABEL_1) dan pro (LABEL_2)—sebuah tantangan umum dalam analisis sentimen, terutama ketika ekspresi sentimen bersifat halus atau tidak eksplisit. Namun, implementasi strategis dari teknik resampling berhasil mengatasi kemiripan antar kelas tersebut, sehingga menghasilkan kemampuan klasifikasi yang seimbang dan kuat, sebagaimana tercermin dalam laporan klasifikasi akhir.

Kontribusi utama dari penelitian ini adalah demonstrasi terhadap alat otomatis yang andal untuk mengukur opini publik dalam konteks linguistik yang spesifik dan berskala besar. Model yang telah divalidasi ini memiliki implikasi praktis yang signifikan, menawarkan metode yang dapat diskalakan untuk memantau wacana publik mengenai isu sosial-politik maupun komersial secara nyaris real-time.

Rekomendasi

Untuk penelitian selanjutnya, model ini dapat diuji daya generalisasinya terhadap berbagai topik lain serta lintas platform media sosial. Penelitian lebih lanjut juga dapat difokuskan pada rekayasa fitur lanjutan atau modifikasi arsitektur model guna meningkatkan kemampuan model dalam menangani ambiguitas antara kelas sentimen yang berdekatan.

Selain rekomendasi teknis tersebut, penelitian ini juga memberikan implikasi penting bagi pemangku kebijakan. Model analisis sentimen ini dapat diadopsi sebagai alat pemantauan opini publik secara real-time untuk mendeteksi secara dini sentimen negatif terhadap isu-isu krusial seperti "*brain drain*". Dengan demikian, pemerintah dapat menggunakan temuan ini sebagai landasan berbasis bukti (*evidence-based*) untuk mengevaluasi kebijakan yang ada dan merancang intervensi yang lebih tepat sasaran. Misalnya, jika sentimen negatif dominan terkait kurangnya penghargaan atau layanan publik, kebijakan dapat difokuskan pada perbaikan ekosistem kerja dan peningkatan

kualitas birokrasi. Ini mengubah analisis sentimen dari sekadar laporan akademis menjadi instrumen strategis untuk pengambilan keputusan yang responsif terhadap aspirasi publik.

Referensi

- A. F. Hidayatullah, R. A. Apong, D. T. C. Lai and A. Qazi, "Corpus creation and language identification for code-mixed Indonesian-Japanese-English Tweets," *PeerJ Comput. Sci.*, vol. 9, e1312, 2023, doi: 10.7717/peerj-cs.1312.
- A. Fachry, A. Farizi, and Y. Sibaroni, "Implementation of Bilstm and Indobert For," vol. 10, no. 1, pp. 96–106, 2025.
- A. Jadhav, S. M. Mostafa, H. Elmannai, and F. K. Karim, "An Empirical Assessment of Performance of Data Balancing Techniques in Classification Task," *Appl. Sci.*, vol. 12, no. 8, 2022, doi: 10.3390/app12083928.
- A. R. Lahase, M. Shelke, R. Jagdale and S. Deshmukh, "A Survey on Sentiment Lexicon Creation and Analysis," in *IOT with Smart Systems*, T. Senjyu, P. Mahalle, T. Perumal and A. Joshi, Eds. Singapore: Springer, 2022, pp. 647–658, doi: 10.1007/978-981-16-3945-6_57.
- A. Rastogi, R. Singh and D. Ather, "Sentiment Analysis Methods and Applications—A Review," in *Proc. 2021 10th Int. Conf. Syst. Modeling & Adv. Res. Trends (SMART)*, Moradabad, India, 2021, pp. 391–395, doi: 10.1109/SMART52563.2021.9676260.
- D. G. Sainger, "Sentiment Analysis – An Assessment of Online Public Opinion: A Conceptual Review," *Turkish J. Comput. Math. Educ. (TURCOMAT)*, vol. 12, no. 5, pp. 1881–1887, Apr. 2021.
- D. Müller, I. Soto-Rey, and F. Kramer, "Towards a guideline for evaluation metrics in medical image segmentation," *BMC Res. Notes*, vol. 15, no. 1, pp. 1–8, 2022, doi: 10.1186/s13104-022-06096-y.
- D. Musfiroh, U. Khaira, P. E. P. Utomo, and T. Suratno, "Analisis Sentimen terhadap Perkuliahan Daring di Indonesia dari Twitter Dataset Menggunakan InSet Lexicon," *MALCOM Indones. J. Mach. Learn. Comput. Sci.*, vol. 1, no. 1, pp. 24–33, 2021, doi: 10.57152/malcom.v1i1.20.
- F. Indriani, R. A. Nugroho, M. R. Faisal, and D. Kartini, "Comparative Evaluation of IndoBERT, IndoBERTweet, and mBERT for Multilabel Student Feedback Classification," *J. RESTI*, vol. 8, no. 6, pp. 748–757, 2024, doi: 10.29207/resti.v8i6.6100.
- F. Koto, J. H. Lau and T. Baldwin, "IndoBERTweet: A Pretrained Language Model for Indonesian Twitter with Effective Domain-Specific Vocabulary Initialization," in *Proc. 2021 Conf. Empirical Methods in Natural Lang. Process.*, Punta Cana, Dominican Republic, 2021, pp. 10660–10668.
- F. N. Ribeiro, M. Araújo, P. Gonçalves, M. A. Gonçalves and F. Benevenuto, "A Benchmark Comparison of State-of-the-Practice Sentiment Analysis Methods," *arXiv preprint arXiv:1512.01818*, 2015.
- F. V. P. Samosir and S. Riyaldi, "Sentiment Analysis of TikTok Comments on Indonesian Presidential Elections Using IndoBERT," in *Proc. 2024 3rd Int. Conf. Creative Commun. Innovative Technol. (ICCIT)*, Tangerang, Indonesia, 2024, pp. 1–7, doi: 10.1109/ICCIT62134.2024.10701256.
- Garg, N., & Sharma, K. (2022). Text pre-processing of multilingual for sentiment analysis based on social network data. *International Journal of Electrical & Computer Engineering* (2088-8708), 12(1).

- J. C. Setiawan, K. M. Lhaksana, and B. Bunyamin, "Sentiment Analysis of Indonesian TikTok Review Using LSTM and IndoBERTweet Algorithm," *JPI (Jurnal Ilm. Penelit. dan Pembelajaran Inform.*, vol. 8, no. 3, pp. 774–780, 2023, doi: 10.29100/jipi.v8i3.3911.
- J. Forry Kusuma and A. Chowanda, "International Journal On Informatics Visualization Journal Homepage: Www.Joiv.Org/Index.Php/Joiv International Journal On Informatics Visualization Indonesian Hate Speech Detection Using IndoBERTweet and BiLSTM on Twitter," vol. 7, no. September, pp. 773–780, 2023.
- M. Fadhel and W. Maharani, "Depression Detection of Users in Social Media X using IndoBERTweet," *Sinkron*, vol. 8, no. 2, pp. 885–891, 2024, doi: 10.33395/sinkron.v9i2.13354.
- M. Z. Subarkah, M. Hilda, and E. Zukhronah, "Analisis Sentimen Review Tempat Wisata Pada Data Online Travel Agency Di Yogyakarta Menggunakan Model Neural Network IndoBERTweet Fine Tuning," *Semin. Nas. Off. Stat.*, vol. 2022, no. 1, pp. 543–552, 2022, doi: 10.34123/semnasoffstat.v2022i1.1246.
- N. Alinda Rahmi and R. Wulan Dari, "Implementation of Natural Language Processing (Nlp) in Consumer Sentiment Analysis of Product Comments on the Marketplace," *J. Tek. Inform.*, vol. 5, no. 3, pp. 693–701, 2024, [Online]. Available: <https://doi.org/10.52436/1.jutif.2024.5.3.1666>
- O. Rainio, J. Teuho, and R. Klén, "Evaluation metrics and statistical tests for machine learning," *Sci. Rep.*, vol. 14, no. 1, pp. 1–14, 2024, doi: 10.1038/s41598-024-56706-x.
- S. Rogić and L. Kaščelan, "Class balancing in customer segments classification using support vector machine rule extraction and ensemble learning," *Comput. Sci. Inf. Syst.*, vol. 18, no. 3, pp. 893–925, 2021, doi: 10.2298/CSIS200530052R.
- S. Wankhede, R. Patil, S. Sonawane and P. A. Save, "Data Preprocessing for Efficient Sentimental Analysis," in *Proc. 2018 2nd Int. Conf. Inventive Commun. Comput. Technol. (ICICCT)*, Coimbatore, India, 2018, pp. 723–726, doi: 10.1109/ICICCT.2018.8473277.
- Slamet Riyanto, Imas Sukaesih Sitanggang, Taufik Djatna and Tika Dewi Atikah, "Comparative Analysis using Various Performance Metrics in Imbalanced Data for Multi-class Text Classification" *International Journal of Advanced Computer Science and Applications(IJACSA)*, 14(6), 2023.
- U. Khairani, V. Mutiawani, and H. Ahmadian, "Pengaruh Tahapan Preprocessing Terhadap Model Indobert Dan IndoBERTweet Untuk Mendeteksi Emosi Pada Komentar Akun Berita Instagram," *J. Teknol. Inf. dan Ilmu Komput.*, vol. 11, no. 4, pp. 887–894, 2024, doi: 10.25126/jtiik.1148315.
- Z. Jianqiang and G. Xiaolin, "Comparison Research on Text Pre-processing Methods on Twitter Sentiment Analysis," *IEEE Access*, vol. 5, pp. 2870–2879, 2017, doi: 10.1109/ACCESS.2017.2672677.
- Ž. Vujović, "Classification Model Evaluation Metrics," *Int. J. Adv. Comput. Sci. Appl.*, vol. 12, no. 6, pp. 599–606, 2021, doi: 10.14569/IJACSA.2021.0120670.