

ANALISIS SENTIMEN SIKAP PUBLIK TERHADAP KONTEN DEEFAKE DI TIKTOK MELALUI MODEL DEEP LEARNING INDOBERT

Firlana Eza Putra¹, Yayak Kartika Sari², Taufiq Agung Cahyono³

^{1, 2, 3}Universitas Bhinneka PGRI, Jl. Mayor Sujadi No.7, Tulungagung, Jawa Timur, Indonesia
Email: firlana.student@ubhi.ac.id

Article History

Received: 13-05-2026

Revision: 09-06-2026

Accepted: 12-06-2026

Published: 01-07-2026

Abstract. The rapid advancement of artificial intelligence has contributed to the increasing spread of deepfake content on social media, raising concerns regarding misinformation, opinion manipulation, and digital ethics. This study aims to analyze public sentiment toward deepfake content on TikTok using the IndoBERT model. This research employs a quantitative approach with deep learning-based sentiment analysis. The dataset consists of 4,000 Indonesian-language comments collected from TikTok through a web scraping technique, which involves the automated extraction of public comments using the Apify TikTok Scraper API based on deepfake-related keywords. Data analysis was conducted through several stages, including text preprocessing, sentiment labeling using a lexicon-based approach, dataset splitting into training and testing sets, fine-tuning the IndoBERT model, and performance evaluation using accuracy, precision, recall, and F1-score metrics. The results indicate that negative sentiment dominates with 48.8% of the total comments, followed by neutral sentiment (34.2%) and positive sentiment (17.0%). The IndoBERT model achieved an accuracy score of 78.16% in classifying sentiment. These findings suggest that most TikTok users are concerned about the potential misuse of deepfake technology and demonstrate the effectiveness of IndoBERT for sentiment analysis of Indonesian social media data.

Keywords: Sentiment Analysis, Deepfake, TikTok, IndoBERT, Deep Learning

Abstrak. Kemajuan teknologi kecerdasan buatan telah mendorong meningkatnya penyebaran konten *deepfake* di media sosial yang berpotensi menimbulkan misinformasi, manipulasi opini, dan permasalahan etika digital. Penelitian ini bertujuan menganalisis sentimen publik terhadap konten *deepfake* di TikTok menggunakan model *IndoBERT*. Penelitian ini merupakan penelitian kuantitatif dengan pendekatan analisis sentimen berbasis *deep learning*. Data penelitian berupa 4.000 komentar berbahasa Indonesia yang dikumpulkan dari TikTok melalui teknik *web scraping*, yaitu proses pengambilan data komentar publik secara otomatis menggunakan *Apify TikTok Scraper API* berdasarkan kata kunci terkait *deepfake*. Analisis data dilakukan melalui tahapan *preprocessing* teks, pelabelan sentimen menggunakan pendekatan *lexicon-based*, pembagian data menjadi data latih dan data uji, *fine-tuning* model *IndoBERT*, serta evaluasi menggunakan metrik *accuracy*, *precision*, *recall*, dan *F1-score*. Hasil penelitian menunjukkan bahwa sentimen negatif mendominasi sebesar 48,8%, diikuti sentimen netral sebesar 34,2% dan sentimen positif sebesar 17,0%. Model *IndoBERT* memperoleh tingkat akurasi sebesar 78,16% dalam mengklasifikasikan sentimen komentar. Temuan ini menunjukkan bahwa mayoritas pengguna TikTok memiliki kekhawatiran terhadap potensi penyalahgunaan teknologi *deepfake* serta membuktikan bahwa *IndoBERT* efektif digunakan untuk analisis sentimen pada data media sosial berbahasa Indonesia.

Kata Kunci: Sentiment Analysis, Deepfake, Tiktok, IndoBERT, Deep Learning

How to Cite: Putra, F. E., Sari, Y. K., & Cahyono, T. A. (2026). Analisis Sentimen Sikap Publik terhadap Konten Deepfake di Tiktok melalui Model Deep Learning Indobert. *HORIZON: Indonesian Journal of Multidisciplinary*, 4 (3), 2682-2692. <http://doi.org/10.54373/hijm.v4i3.5700>

PENDAHULUAN

Perkembangan pesat teknologi kecerdasan buatan (*Artificial Intelligence/AI*) telah mendorong munculnya berbagai inovasi dalam produksi konten digital, salah satunya adalah teknologi *deepfake*. *Deepfake* merupakan teknik manipulasi media sintetis yang memanfaatkan model *deep learning*, khususnya *Generative Adversarial Networks* (GANs), untuk menghasilkan video, gambar, atau audio yang tampak sangat realistis dan sulit dibedakan dari konten asli (Franco et al., 2025). Di satu sisi, teknologi ini dimanfaatkan untuk hiburan dan kreativitas digital, namun di sisi lain berpotensi disalahgunakan untuk penipuan, disinformasi, pencemaran nama baik, serta manipulasi opini publik (Franco et al., 2025).

Di Indonesia, penyebaran konten *deepfake* mengalami peningkatan signifikan seiring tingginya adopsi media sosial. Laporan menunjukkan adanya lonjakan kasus penipuan berbasis *deepfake* hingga 1.550% dalam periode 2022–2023 (News, 2024). Beberapa kasus menonjol melibatkan manipulasi video tokoh publik yang disebarakan melalui platform TikTok untuk tujuan penipuan dan propaganda, sehingga menimbulkan kerugian sosial dan ekonomi (DetikNews, 2025). TikTok menjadi platform yang relevan untuk dikaji karena memiliki lebih dari 100 juta pengguna aktif dewasa di Indonesia dan berperan sebagai sarana utama penyebaran konten berbasis video pendek (DataReportal, 2025).

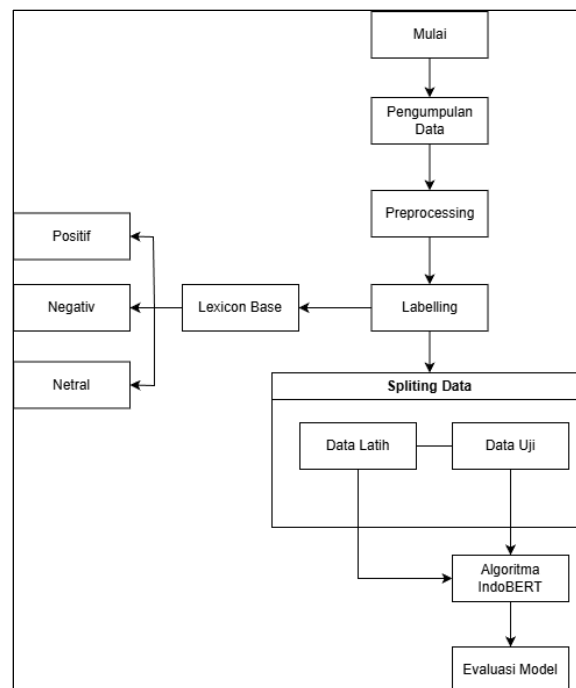
Kolom komentar TikTok merepresentasikan ruang ekspresi opini publik yang bersifat spontan dan emosional. Namun, rendahnya tingkat literasi digital menyebabkan sebagian masyarakat kesulitan dalam membedakan antara konten asli dan konten sintetis, sehingga berpotensi memunculkan dominasi sentimen negatif serta meningkatkan risiko terjadinya perpecahan di masyarakat (Iskandar et al., 2022). Oleh karena itu, pemahaman terhadap sikap publik terhadap fenomena *deepfake* menjadi penting sebagai dasar dalam mengidentifikasi dan mengantisipasi dampak negatif yang ditimbulkan oleh teknologi tersebut.

Analisis sentimen merupakan pendekatan dalam *Natural Language Processing* yang digunakan untuk mengidentifikasi kecenderungan opini publik berdasarkan teks (Ren, 2023). Model berbasis *Transformer*, khususnya *IndoBERT*, telah terbukti unggul dalam menangkap konteks bahasa Indonesia yang kompleks dan informal pada data media sosial dibandingkan metode klasik seperti *Naïve Bayes* dan *Support Vector Machine* (Khofifah et al., 2025). Meskipun penelitian analisis sentimen telah banyak dilakukan, kajian mengenai sentimen publik terhadap konten *deepfake* di TikTok masih terbatas. Sebagian besar penelitian sebelumnya berfokus pada *platform* seperti YouTube atau X (Twitter), sehingga belum memberikan gambaran spesifik mengenai respons pengguna TikTok terhadap fenomena *deepfake* (Alexander et al., 2023).

Untuk mengisi kesenjangan tersebut, penelitian ini menerapkan model *IndoBERT* untuk menganalisis sentimen publik terhadap konten *deepfake* di TikTok. Model ini dipilih karena mampu memahami konteks bahasa Indonesia, termasuk bahasa informal yang umum digunakan di media sosial. Melalui pendekatan ini, penelitian diharapkan dapat memberikan gambaran yang lebih komprehensif mengenai persepsi masyarakat terhadap *deepfake* serta mendukung upaya peningkatan literasi digital. Berdasarkan latar belakang dan kesenjangan penelitian tersebut, penelitian ini bertujuan untuk menganalisis sentimen sikap publik terhadap konten *deepfake* di TikTok menggunakan model *deep learning IndoBERT*. Hasil penelitian diharapkan dapat menjadi dasar bagi penguatan literasi digital dan pengawasan konten berbasis kecerdasan buatan di Indonesia.

METODE

Metode Penelitian ini merupakan penelitian kuantitatif dengan pendekatan analisis sentimen menggunakan model *deep learning* berbasis *IndoBERT*. Metode ini dipilih karena terbukti unggul dalam menangkap konteks bahasa Indonesia yang kompleks dan cenderung informal, khususnya pada data media social (Muftie & Haris, 2023). Tahapan penelitian meliputi pengumpulan data, *preprocessing*, pelabelan menggunakan pendekatan *lexicon-based*, pembagian data (*data splitting*), proses *fine-tuning*, serta evaluasi model. Seluruh tahapan tersebut disajikan secara rinci pada Gambar 1.



Gambar 1. Alur model penelitian

Tahap pengumpulan data dilakukan menggunakan Google Colaboratory (Colab) dengan memanfaatkan Apify TikTok Scraper API untuk mengambil komentar publik pada video bertema deepfake. Seluruh proses pengambilan, pembersihan, dan pengolahan data dilakukan secara komputasional di Google Colab, kemudian hasilnya disimpan dalam format CSV untuk keperluan analisis lebih lanjut. Penelitian ini dilaksanakan selama tiga bulan, yaitu Oktober hingga Desember 2025. Selanjutnya, dilakukan tahap *preprocessing* untuk membersihkan dan menormalkan teks (Wilie et al., 2020). Data yang telah diproses kemudian diberi label sentimen ke dalam tiga kategori, yaitu positif, netral, dan negatif (Fathoni et al., 2024). Setelah proses pelabelan, data dibagi menjadi data latih (*training set*) dan data uji (*testing set*) (Putri, 2025). Dalam penelitian ini, model *IndoBERT* yang telah dilatih sebelumnya (*pre-trained*) *di-fine-tune* secara langsung untuk tugas klasifikasi sentimen. Proses *fine-tuning* dilakukan dengan menyesuaikan parameter model menggunakan data latih sehingga model mampu memahami karakteristik data secara lebih spesifik. Tahap akhir adalah evaluasi kinerja model menggunakan metrik *accuracy*, *precision*, *recall*, *F1-score*, serta confusion matrix untuk mengukur performa model secara menyeluruh.

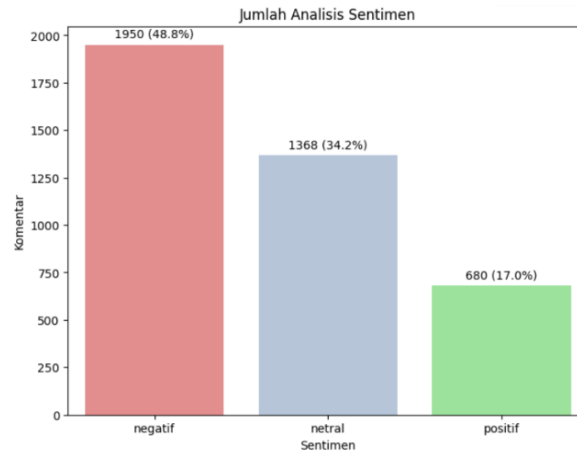
HASIL DAN DISKUSI

Pengumpulan Data

Pengumpulan data dalam penelitian ini dilakukan menggunakan teknik *web scraping* pada platform TikTok dengan menggunakan bahasa pemrograman Python dan pustaka *TikTokApi*. Data yang diambil berupa komentar publik dari video yang mengandung konten *deepfake*, dengan kata kunci seperti *deepfake*, *AI-generated*, dan *#deepfake*. Proses pengambilan data difokuskan pada komentar berbahasa Indonesia untuk memastikan kesesuaian dengan model *IndoBERT*. Data yang diperoleh kemudian disimpan dalam format CSV dengan atribut utama berupa teks komentar, waktu unggahan, dan jumlah interaksi. Setelah melalui tahap seleksi dan pembersihan data, seperti penghapusan duplikasi dan komentar yang tidak relevan, diperoleh sebanyak 4.000 komentar.

Preprocessing Data

Dalam tahap ini, data mengalami serangkaian proses *preprocessing*, yang meliputi *cleaning*, *casefolding*, normalisasi, *stopword removal*, dan *stemming*. Tahap ini bertujuan untuk menstandarisasi struktur teks sebelum dilakukan klasifikasi. Rincian hasil *preprocessing* disajikan dalam Tabel 1.



Gambar 3. Distribusi label sentimen

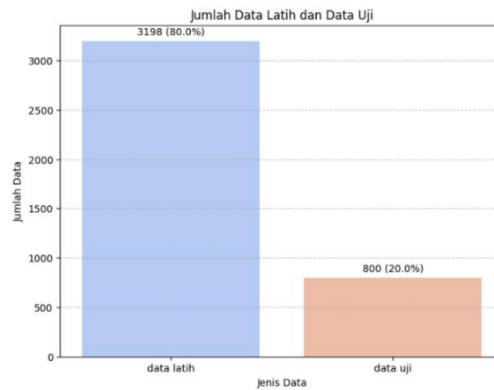
Temuan ini menunjukkan bahwa mayoritas opini publik terhadap konten *deepfake* di TikTok cenderung bersifat negatif, yang mengindikasikan adanya kekhawatiran terhadap dampak dan potensi penyalahgunaan teknologi tersebut. Meskipun demikian, keberadaan sentimen netral dan positif menunjukkan bahwa sebagian pengguna masih memandang teknologi ini sebagai hal yang informatif maupun menarik. Ketidakseimbangan distribusi kelas sentimen ini berpotensi memengaruhi kinerja model klasifikasi, terutama pada kelas dengan jumlah data yang lebih sedikit. Contoh komentar untuk masing-masing kategori sentimen disajikan pada tabel 2.

Tabel 2. Text berdasarkan hasil *labelling*

Text	Sentiment Score	Sentiment Label
udah sering liat utama artis yg promosi judl gw dah tau ulah oknum yg pake ai jaman sekarang teknologi makin canggih malah salah guna korbn biasa orang awam	0	Negatif
udah sering liat utama artis yg promosi judl gw dah tau ulah oknum yg pake ai jaman sekarang teknologi makin canggih malah salah guna korbn biasa orang awam	1	Positif
aku pernah vc baim wong alhamdulillah gk kena tipu aku bilang gin eehh tipu langsung mati hp nya	2	Netral

Split Data

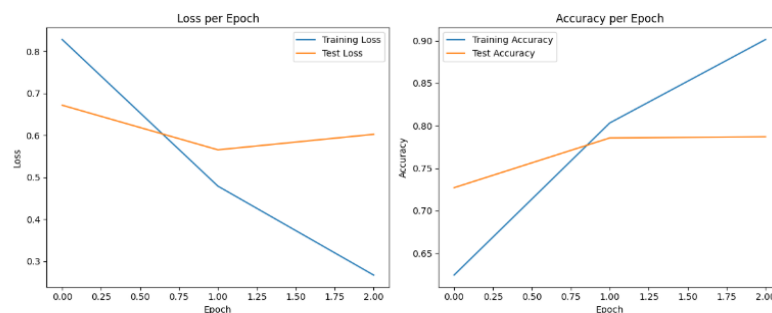
Tahap splitting data dilakukan untuk membagi dataset menjadi data latih dan data uji guna mengukur kemampuan model dalam melakukan generalisasi. Pada penelitian ini digunakan rasio 80:20, yaitu 80% data latih dan 20% data uji. Dari total 4.000 data, diperoleh sebanyak 3.198 data latih dan 800 data uji. Data latih digunakan untuk proses fine-tuning model IndoBERT, sedangkan data uji digunakan untuk evaluasi kinerja model. Gambar 3 menggambarkan distribusi pembagian data.



Gambar 4. Pembagian data latih dan data uji

Model IndoBERT

Model *IndoBERT* digunakan untuk melakukan klasifikasi sentimen terhadap komentar TikTok ke dalam tiga kategori, yaitu positif, negatif, dan netral. Sebelum proses pelatihan, data teks ditokenisasi dan diubah menjadi representasi numerik agar dapat diproses oleh model. Selanjutnya, dilakukan proses *fine-tuning* menggunakan data latih untuk menyesuaikan parameter model terhadap tugas klasifikasi sentimen.



Gambar 5. Grafik *loss* dan akurasi model *IndoBERT*

Berdasarkan Gambar, nilai training loss menunjukkan penurunan yang signifikan dari sekitar 0,80 pada epoch awal menjadi sekitar 0,26 pada epoch terakhir, yang menandakan bahwa model mampu mempelajari pola data dengan baik. Sementara itu, test loss mengalami penurunan di awal pelatihan dan cenderung stabil pada epoch berikutnya. Dari sisi akurasi, *training accuracy* meningkat secara konsisten hingga mencapai sekitar 0,91, sedangkan *test accuracy* meningkat hingga sekitar 0,78 dan kemudian cenderung stabil. Perbedaan antara nilai pada data latih dan data uji menunjukkan adanya indikasi *overfitting* ringan, di mana model memiliki performa yang sangat baik pada data latih namun sedikit menurun pada data uji. Secara keseluruhan, hasil ini menunjukkan bahwa model *IndoBERT* mampu mempelajari representasi sentimen dengan baik dan memiliki kemampuan prediksi yang cukup optimal, meskipun masih terdapat peluang untuk peningkatan performa melalui optimasi model lebih lanjut.

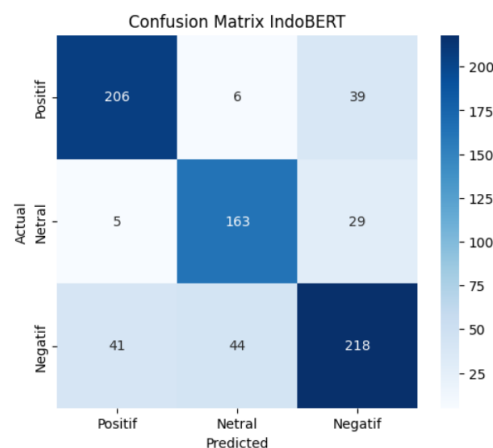
Evaluasi Model

Evaluasi kinerja model *IndoBERT* dilakukan menggunakan beberapa metrik, yaitu *accuracy*, *precision*, *recall*, dan *F1-score*, serta analisis *confusion matrix*. Evaluasi ini bertujuan untuk mengukur kemampuan model dalam mengklasifikasikan sentimen komentar ke dalam tiga kategori, yaitu positif, netral, dan negatif.

	precision	recall	f1-score	support
positive	0.82	0.82	0.82	251
neutral	0.77	0.83	0.80	197
negative	0.76	0.72	0.74	303
accuracy			0.78	751
macro avg	0.78	0.79	0.78	751
weighted avg	0.78	0.78	0.78	751

Gambar 6. Hasil evaluasi model

Berdasarkan hasil pengujian pada data uji, model *IndoBERT* memperoleh nilai akurasi sebesar 0,7816, yang menunjukkan bahwa model mampu mengklasifikasikan sekitar 78,16% data dengan benar. Nilai *precision*, *recall*, dan *F1-score* rata-rata tertimbang masing-masing sebesar 0,78, yang mengindikasikan keseimbangan kinerja model dalam meminimalkan kesalahan klasifikasi baik *false positive* maupun *false negative*. Analisis lebih lanjut berdasarkan kategori sentimen menunjukkan bahwa model memiliki performa terbaik pada kelas positif dengan nilai *F1-score* sebesar 0,82, diikuti oleh kelas netral sebesar 0,80, dan kelas negatif sebesar 0,74. Hal ini menunjukkan bahwa model lebih mampu mengenali pola sentimen positif dibandingkan dengan sentimen negatif.



Gambar 7. *Confusion matrix* model *IndoBERT*

Berdasarkan Gambar menampilkan *confusion matrix* yang memberikan gambaran distribusi hasil prediksi model. Pada kelas positif, sebanyak 206 data berhasil diklasifikasikan dengan benar, sementara sebagian kecil salah diklasifikasikan ke dalam kelas netral (6 data) dan negatif (39 data). Pada kelas netral, terdapat 163 data yang terklasifikasi dengan benar,

dengan kesalahan klasifikasi ke kelas positif (5 data) dan negatif (29 data). Sementara itu, pada kelas negatif, sebanyak 218 data berhasil diklasifikasikan dengan benar, namun masih terdapat kesalahan prediksi ke kelas positif (41 data) dan netral (44 data).

Hasil tersebut menunjukkan bahwa model masih mengalami kesulitan dalam membedakan sentimen negatif dengan kategori lainnya. Hal ini kemungkinan disebabkan oleh karakteristik komentar media sosial yang sering menggunakan bahasa tidak baku, singkatan, atau makna implisit. Meskipun demikian, model *IndoBERT* tetap mampu mencapai akurasi sebesar 78,16%, yang menunjukkan kemampuan yang cukup baik dalam memahami konteks bahasa Indonesia pada data media sosial. Selain itu, dominasi sentimen negatif sebesar 48,8% menunjukkan bahwa sebagian besar pengguna TikTok memiliki kekhawatiran terhadap potensi penyalahgunaan teknologi *deepfake*. Temuan ini sejalan dengan penelitian Franco et al., (2025) yang menyatakan bahwa *deepfake* berpotensi menimbulkan disinformasi dan menurunkan kepercayaan publik terhadap informasi digital. Hasil penelitian ini juga mendukung penelitian Khofifah et al., (2025) yang menunjukkan bahwa model *IndoBERT* memiliki performa yang baik dalam tugas klasifikasi sentimen berbahasa Indonesia dibandingkan metode konvensional. Temuan penelitian menunjukkan bahwa model *IndoBERT* efektif digunakan untuk analisis sentimen pada data media sosial berbahasa Indonesia dan mampu mengidentifikasi kecenderungan persepsi masyarakat terhadap fenomena *deepfake* di TikTok.

KESIMPULAN

Berdasarkan hasil penelitian yang telah dilakukan, analisis sentimen terhadap komentar TikTok terkait konten *deepfake* menggunakan model *IndoBERT* berhasil dilakukan pada 4.000 komentar berbahasa Indonesia. Hasil pelabelan menunjukkan bahwa sentimen negatif mendominasi sebesar 48,8%, diikuti sentimen netral sebesar 34,2% dan sentimen positif sebesar 17,0%. Temuan ini mengindikasikan bahwa sebagian besar pengguna memiliki kekhawatiran terhadap potensi penyalahgunaan teknologi *deepfake* di media sosial. Berdasarkan hasil evaluasi, model *IndoBERT* memperoleh akurasi sebesar 78,16% dengan nilai *precision*, *recall*, dan *F1-score* yang relatif seimbang. Hasil tersebut menunjukkan bahwa *IndoBERT* cukup efektif dalam memahami konteks bahasa Indonesia dan dapat digunakan untuk analisis sentimen pada data media sosial. Namun, model masih mengalami keterbatasan dalam mengidentifikasi sentimen negatif secara optimal karena karakteristik komentar yang cenderung tidak terstruktur dan mengandung makna implisit.

Keterbatasan penelitian ini terletak pada penggunaan pelabelan sentimen berbasis *lexicon-based*, penggunaan data yang hanya berasal dari TikTok, serta ketidakseimbangan jumlah data pada setiap kelas sentimen yang berpotensi memengaruhi performa model. Oleh karena itu, penelitian selanjutnya dapat menggunakan metode pelabelan yang lebih akurat, sumber data yang lebih beragam, serta teknik penyeimbangan data untuk meningkatkan hasil klasifikasi.

REKOMENDASI

Berdasarkan hasil penelitian, penelitian selanjutnya disarankan untuk menggunakan jumlah dataset yang lebih besar dan lebih beragam agar model dapat mempelajari pola bahasa yang lebih kompleks. Selain itu, penyeimbangan jumlah data pada setiap kelas sentimen perlu dilakukan untuk mengurangi bias terhadap kelas mayoritas. Penggunaan metode pelabelan yang lebih akurat, seperti pelabelan manual oleh beberapa annotator atau kombinasi metode *lexicon-based* dan *machine learning*, juga dapat meningkatkan kualitas dataset. Penelitian mendatang dapat membandingkan performa *IndoBERT* dengan model lain seperti LSTM, CNN, maupun model *Transformer* terbaru untuk memperoleh hasil yang lebih optimal. Selain itu, pengembangan sistem analisis sentimen berbasis aplikasi yang mampu melakukan pemantauan opini publik secara *real-time* terhadap isu *deepfake* menjadi peluang penelitian yang menarik untuk dikembangkan.

REFERENSI

- Alexander, J., Pratama, R., & Yusuf, D. (2023). Analisis Sentimen Komentar YouTube Indonesia terhadap Deepfake Tokoh Publik. *Jurnal Teknologi dan Informasi*, 12(2), 101–115. <https://doi.org/10.1234/jti.v12i2.4567>
- DataReportal. (2025). *Digital 2025: Indonesia*. <https://datareportal.com/reports/digital-2025-indonesia>
- DetikNews. (2025). *Babak Baru Perkara Deepfake Catut Prabowo*. <https://news.detik.com/berita/d-7884867/babak-baru-perkara-deepfake-catut-prabowo>
- Fathoni, M. F. N., Puspaningrum, E. Y., & Sihananto, A. (2024). Lexicon-Based Sentiment Labeling in Indonesian Social Media Analysis. *Jurnal Teknologi Informasi Indonesia*, 13(1), 21–33. <https://doi.org/10.54059/jtii.v13i1.214>
- Franco, M., Rossi, L., & Nguyen, P. (2025). Deepfakes: The New Evil of the New Age of (MIS) Information and Post-Truth. *Journal of Digital Ethics and AI Studies*, 18(1), 44–61. <https://doi.org/10.1016/j.jdeas.2025.03.002>
- Iskandar, J., Prasetya, A., Sari, Y. K., & Cahyono, A. (2022). Analisis Penerimaan Sistem Informasi Akademik Universitas Bhinneka PGRI Menggunakan Integrasi Model TPB dan TAM.
- Khofifah, N., Santoso, B., & Rahayu, E. (2025). Perbandingan Model IndoBERT, SVM, Naïve Bayes, dan KNN pada Analisis Sentimen Ulasan Aplikasi M-Paspor. *Jurnal Sains Komputer dan Informatika (J-SAKTI)*, 9(1), 75–86. <https://doi.org/10.21009/jsakti.2025.09.1.07>

- Muftie, F., & Haris, M. (2023). *IndoBERT Based Data Augmentation for Indonesian Text Classification*. <https://doi.org/10.1109/icitri59340.2023.10250061>
- News, A. (2024). *VIDA Catat Penipuan “Deepfake” di Indonesia Melonjak 1.550 Persen*. <https://www.antaranews.com/berita/4437365/vida-catat-penipuan-deepfake-di-indonesia-melonjak-1550-persen>
- Putri, R. A. (2025). Adapting Language Models to Indonesian Local Languages. In *arXiv preprint arXiv:2507.01645*.
- Ren, Z.-G. (2023). *Sentiment Analysis*. https://doi.org/10.1007/978-1-4842-9496-3_6
- Wilie, B., Vincentio, K., Winata, G. I., Cahyawijaya, S., Li, X., Lim, Z. Y., Soleman, S., Mahendra, R., & Baldwin, T. (2020). IndoNLU: Benchmark and Resources for Evaluating Indonesian NLP. *ArXiv Preprint*. <https://arxiv.org/abs/2009.05387>