

Klasifikasi Berita Hoaks Berbahasa Indonesia Menggunakan Algoritma *Machine Learning* dan Model Transformer

Lyvia Valentina^{*1}, Msy Anggita Shafira², Ken Ditha Tania³, Ahmad Rifai⁴

^{1,2,3,4}Jurusan Sistem Informasi, Universitas Sriwijaya, Indralaya, Indonesia
Email: ¹valentinalyvia@gmail.com, ²msyanggitashafira1@gmail.com, ³kenya.tania@gmail.com, ⁴rifai.bae@gmail.com

Abstrak

Peningkatan penyebaran informasi digital di Indonesia yang turut memperbesar risiko beredarnya berita hoaks dan disinformasi yang dapat memicu keresahan serta polarisasi masyarakat. Penelitian ini bertujuan untuk mengembangkan sistem klasifikasi berita hoaks otomatis serta membandingkan performa algoritma *Machine Learning* klasik dengan model berbasis transformer. Tahapan penelitian meliputi pengumpulan dataset dari TurnBackHoax.id sebagai sumber berita hoaks serta AntaraNews, Kompas, dan Detik.com sebagai sumber berita faktual, dilanjutkan dengan pra-pemrosesan teks, ekstraksi fitur menggunakan *Term Frequency-Inverse Document Frequency* (TF-IDF), serta *fine-tuning* model IndoBERT. Model yang diuji meliputi *Support Vector Machine* (SVM), *Random Forest*, dan IndoBERT, dengan evaluasi menggunakan metrik *Accuracy*, *Precision*, *Recall*, dan *Macro F1-Score*. Hasil penelitian menunjukkan bahwa IndoBERT memberikan performa terbaik dengan akurasi sebesar 99,38% dan *Macro F1-Score* sebesar 0,99, diikuti oleh SVM dengan akurasi 98,04% dan *Random Forest* sebesar 97,09%. Keunggulan IndoBERT terletak pada kemampuannya dalam menangkap konteks semantik bahasa secara lebih mendalam dibandingkan pendekatan berbasis frekuensi. Meskipun demikian, SVM tetap menunjukkan performa yang kompetitif dengan kompleksitas komputasi yang lebih rendah. Dengan demikian, dapat disimpulkan bahwa model transformer IndoBERT merupakan pendekatan paling efektif untuk klasifikasi berita hoaks berbahasa Indonesia, sementara algoritma *Machine Learning* klasik tetap relevan sebagai alternatif pada lingkungan dengan keterbatasan sumber daya komputasi.

Kata kunci: Deteksi Hoaks, IndoBERT, Klasifikasi Teks, Machine Learning, Transformer.

Classification of Indonesian Hoax News Using Machine Learning Algorithms and Transformer Models

Abstract

The increased spread of digital information in Indonesia which also enlarges the risk of circulating hoax news and disinformation that can trigger unrest and societal polarization. This study aims to develop an automatic hoax news classification system and compare the performance of classical Machine Learning algorithms with transformer-based models. The research stages include dataset collection from TurnBackHoax.id as a source of hoax news and AntaraNews, Kompas, and Detik.com as sources of factual news, followed by text pre-processing, feature extraction using *Term Frequency-Inverse Document Frequency* (TF-IDF), and *fine-tuning* of the IndoBERT model. The evaluated models include *Support Vector Machine* (SVM), *Random Forest*, and IndoBERT, using *Accuracy*, *Precision*, *Recall*, and *Macro F1-Score* as evaluation metrics. The results show that IndoBERT achieves the best performance with an accuracy of 99.38% and a *Macro F1-Score* of 0.99, followed by SVM with 98.04% and *Random Forest* with 97.09%. The superiority of IndoBERT lies in its ability to capture deeper semantic context compared to frequency-based approaches. However, SVM remains a competitive alternative due to its lower computational complexity. In conclusion, the IndoBERT transformer model is the most effective approach for Indonesian hoax news classification, while classical Machine Learning algorithms remain relevant in resource-constrained environments.

Keywords: Hoax Detection, IndoBERT, Machine Learning, Text Classification, Transformer.

1. PENDAHULUAN

Perkembangan teknologi informasi dan komunikasi serta penggunaan media sosial yang semakin luas telah mengubah cara masyarakat memperoleh dan menyebarkan informasi. Informasi kini dapat tersebar dengan sangat cepat melalui berbagai platform digital sehingga menjangkau audiens yang sangat luas dalam waktu singkat. Meskipun memberikan kemudahan dalam akses informasi, kondisi ini juga memunculkan permasalahan serius berupa penyebaran informasi palsu atau hoaks. Penyebaran hoaks dapat menyebabkan misinformasi, memicu keresahan masyarakat, serta berpotensi menimbulkan konflik sosial apabila tidak segera ditangani secara efektif [1], [2].

Fenomena berita hoaks menjadi salah satu tantangan utama dalam ekosistem informasi digital. Konten hoaks umumnya disajikan dengan gaya bahasa yang provokatif dan sensasional untuk menarik perhatian pembaca. Selain itu, penggunaan huruf kapital berlebihan, tanda seru yang berulang, serta struktur kalimat yang tidak mengikuti kaidah jurnalistik sering ditemukan dalam berita hoaks. Karakteristik linguistik tersebut berbeda dengan berita faktual yang cenderung menggunakan bahasa yang lebih formal dan informatif [3], [4].

Untuk mengatasi permasalahan tersebut, berbagai penelitian telah dilakukan dengan memanfaatkan teknik *Machine Learning* dan *Natural Language Processing* (NLP) untuk mendeteksi berita hoaks secara otomatis. Salah satu tahap penting dalam klasifikasi teks adalah proses representasi fitur yang mengubah data teks menjadi bentuk numerik yang dapat diproses oleh algoritma klasifikasi. Metode yang umum digunakan dalam tahap ini adalah *Term Frequency-Inverse Document Frequency* (TF-IDF), yang menghitung bobot kata berdasarkan frekuensinya kemunculannya dalam dokumen dan tingkat kepentingannya terhadap keseluruhan korpus dokumen [5], [6].

Setelah proses representasi fitur dilakukan, berbagai algoritma klasifikasi dapat digunakan untuk membangun model prediksi. Dua algoritma yang sering digunakan dalam klasifikasi teks adalah *Support Vector Machine* (SVM) dan *Random Forest*. SVM dikenal memiliki kemampuan yang baik dalam menangani data berdimensi tinggi serta mampu menghasilkan *hyperplane* optimal untuk memisahkan kelas data. Sementara itu, *Random Forest* merupakan metode ensemble berbasis *decision tree* yang mampu meningkatkan akurasi klasifikasi dengan menggabungkan hasil dari beberapa pohon keputusan [7], [8]. Beberapa penelitian menunjukkan bahwa kombinasi TF-IDF dengan algoritma SVM maupun *Random Forest* dapat memberikan performa yang baik pada berbagai tugas klasifikasi teks seperti analisis sentimen dan deteksi informasi palsu [9], [10].

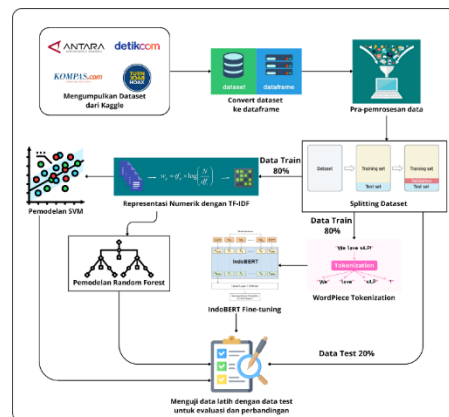
Seiring dengan perkembangan teknologi NLP, pendekatan berbasis deep learning khususnya model transformer juga semakin banyak digunakan dalam klasifikasi teks. Salah satu model yang dikembangkan khusus untuk bahasa Indonesia adalah IndoBERT, yaitu model bahasa berbasis transformer yang telah dilatih menggunakan korpus besar bahasa Indonesia sehingga mampu memahami konteks kata dalam kalimat secara lebih baik. Beberapa penelitian menunjukkan bahwa IndoBERT mampu menghasilkan performa tinggi dalam klasifikasi teks berbahasa Indonesia, dengan peningkatan akurasi dibandingkan metode *Machine Learning* tradisional [1], [2], [11].

Meskipun model berbasis transformer seperti IndoBERT menunjukkan performa yang sangat baik, pendekatan berbasis fitur seperti TF-IDF yang dikombinasikan dengan algoritma SVM dan *Random Forest* masih banyak digunakan karena memiliki kompleksitas komputasi yang lebih rendah dan tetap mampu memberikan performa yang kompetitif pada dataset dengan ukuran terbatas [7], [8], [9]. Selain itu, metode *Machine Learning* juga telah banyak digunakan dalam berbagai penelitian analisis data seperti analisis sentimen aplikasi digital, *knowledge discovery*, dan prediksi berbasis data menggunakan teknik pembelajaran mesin [12], [13], [14], [15].

Berdasarkan uraian tersebut, penelitian ini bertujuan untuk mengembangkan model klasifikasi berita hoaks berbahasa Indonesia dengan membandingkan performa pendekatan TF-IDF yang dikombinasikan dengan algoritma *Support Vector Machine* (SVM) dan *Random Forest* dengan pendekatan transformer menggunakan IndoBERT. Evaluasi performa model dilakukan menggunakan metrik seperti *Macro-Precision*, *Macro-Recall*, dan *Macro-F1*, serta analisis *confusion matrix* untuk memahami pola kesalahan klasifikasi pada setiap kelas. Penelitian ini diharapkan dapat memberikan kontribusi dalam pengembangan sistem deteksi hoaks yang lebih efektif serta mendukung upaya mitigasi penyebaran informasi palsu di ekosistem digital Indonesia.

2. METODE PENELITIAN

Metodologi yang diusulkan dalam penelitian ini mengintegrasikan pendekatan statistik berbasis frekuensi dan pemodelan bahasa berbasis konteks untuk mengklasifikasikan berita hoaks. Alur penelitian dimulai dari akuisisi data melalui korpus publik, dilanjutkan dengan rangkaian pra-pemrosesan teks, ekstraksi fitur menggunakan TF-IDF untuk model *Machine Learning* (SVM dan *Random Forest*), serta proses *fine-tuning* pada model transformer IndoBERT.



Gambar 1. Tahapan Penelitian

2.1. Dataset

Penelitian ini menggunakan dataset Deteksi Berita Hoaks Indo yang diakses melalui platform Kaggle. Dataset ini merupakan korpus komprehensif yang dikumpulkan melalui teknik scraping dari berbagai portal berita kredibel dan situs verifikasi fakta di Indonesia. Karakteristik data yang digunakan mencakup dua label utama: hoaks dan faktual. Sumber berita faktual berasal dari situs AntaraNews, Kompas, dan Detik.com, sedangkan data hoaks bersumber dari TurnBackHoax.id. Dataset berita hoaks dilabelkan dengan nomor 0 sementara berita faktual dilabelkan dengan nomor 1. Total entri data dalam dataset ini berjumlah 23.942 dokumen dengan rincian sebagai berikut:

Tabel 1. Rincian Dataset

Sumber Data	Kategori	Jumlah Dokumen
AntaraNews	Faktual	4.200
Kompas	Faktual	3.500
Detik.com	Faktual	2.948
TurnBackHoax.id	Hoaks	12.744
Total Keseluruhan		23.942

2.2. Pra-pemrosesan Data

Tahap pra-pemrosesan data dilakukan secara naratif untuk mengubah data mentah menjadi bentuk yang siap diolah oleh model klasifikasi. Proses ini diawali dengan tahapan *case folding* untuk menyeragamkan seluruh karakter menjadi huruf kecil guna menghindari redundansi kata akibat perbedaan kapitalisasi. Selanjutnya, dilakukan data cleaning menggunakan ekspresi reguler untuk menghapus elemen non-tekstual seperti URL, simbol khusus, angka, dan tanda baca yang tidak memberikan kontribusi semantik dalam deteksi hoaks. Setelah teks bersih, dilakukan proses normalisasi untuk memperbaiki kata-kata tidak baku atau singkatan menjadi bentuk standar sesuai Pedoman Umum Ejaan Bahasa Indonesia (PUEBI) menggunakan kamus leksikon.

2.3. Dataset Splitting

Untuk mengukur efektivitas model, dataset dibagi menjadi dua subset utama dengan proporsi 80% sebagai data latih dan 20% sebagai data uji menggunakan metode stratified sampling. Evaluasi kinerja dilakukan secara menyeluruh menggunakan metrik dari *Confusion matrix*, yang mencakup *Accuracy*, *Precision*, *Recall*, dan *F1-Score*. *F1-Score* menjadi metrik utama yang diperhatikan karena mampu memberikan gambaran performa yang seimbang antara presisi dan daya tangkap model terhadap berita hoaks. Hasil akhir dari ketiga pendekatan tersebut dibandingkan untuk menganalisis sejauh mana model transformer mengungguli algoritma *Machine Learning* konvensional dalam mendeteksi disinformasi pada kumpulan data teks berita bahasa Indonesia yang kompleks.

2.4. Representasi Numerik dengan TF-IDF

Teks yang telah melalui tahap pra-pemrosesan kemudian ditransformasikan ke dalam representasi numerik menggunakan metode *Term Frequency-Inverse Document Frequency* (TF-IDF) untuk kebutuhan algoritma SVM dan *Random Forest*. TF-IDF bekerja dengan menghitung bobot setiap kata berdasarkan frekuensi

kemunculannya dalam dokumen tertentu relatif terhadap kelangkaannya di seluruh korpus. Formulasi pembobotan fitur didefinisikan secara matematis sebagai berikut:

$$W_{i,j} = tf_{i,j} \times \log\left(\frac{N}{df_i}\right) \quad (1)$$

Dalam implementasinya, ekstraksi fitur dilakukan dengan menggabungkan fitur unigram dan bigram untuk menangkap hubungan antar kata yang berdekatan yang sering muncul dalam narasi provokatif pada berita hoaks.

2.5. Pemodelan *Machine Learning*

1) *Support Vector Machine* (SVM)

Algoritma *Support Vector Machine* (SVM) diimplementasikan menggunakan kelas SVC dari pustaka *Scikit-Learn* dengan inisialisasi parameter `class_weight='balanced'`. Penambahan parameter ini sangat krusial untuk menangani potensi ketidakseimbangan kelas dalam dataset dengan memberikan bobot penalti yang lebih besar pada kelas minoritas secara otomatis. Model SVM dilatih menggunakan matriks fitur TF-IDF untuk mencari *hyperplane* optimal yang memisahkan berita hoaks dan faktual dengan margin maksimal. Strategi ini dipilih karena SVM memiliki kemampuan generalisasi yang kuat pada data teks berdimensi tinggi.

2) *Random Forest* (RF)

Algoritma *Random Forest* diterapkan sebagai metode *ensemble learning* yang berbasis pada pembangunan sekumpulan pohon keputusan (*decision trees*). Model *Random Forest* membangun sekumpulan pohon keputusan secara independen melalui teknik bagging untuk meningkatkan stabilitas prediksi. Keputusan akhir klasifikasi ditentukan melalui mekanisme *majority voting* dari seluruh pohon yang terbentuk, yang secara efektif mampu menangkap interaksi non-linear antar fitur kata sekaligus mencegah terjadinya *overfitting* yang umum ditemukan pada pohon keputusan tunggal.

3) Model Transformer IndoBERT

Pendekatan transformer diimplementasikan menggunakan model IndoBERT untuk menangkap konteks semantik bahasa Indonesia secara bidireksional. Berbeda dengan pendekatan statistik, data teks pada jalur ini diproses menggunakan metode *WordPiece Tokenization* yang memecah kalimat menjadi unit sub-kata untuk menangani kosakata di luar kamus secara lebih efektif. Proses *fine-tuning* dilakukan dengan menambahkan satu lapisan klasifikasi linear di atas *output* token. Konfigurasi hyperparameter ditetapkan dengan *learning rate* sebesar 2×10^{-5} , *batch size* 32, dan jumlah epoch sebanyak 5 untuk menjamin konvergensi model yang stabil. Strategi *early stopping* juga diterapkan untuk menjaga kemampuan generalisasi model terhadap data baru.

3. HASIL DAN PEMBAHASAN

Proses evaluasi pada penelitian ini dilakukan dengan membandingkan performa algoritma *Machine Learning* konvensional berbasis statistik, yaitu *Support Vector Machine* (SVM) dan *Random Forest*, terhadap model deep learning berbasis transformer yaitu IndoBERT. Evaluasi model dilakukan menggunakan metrik *Accuracy*, *Precision*, *Recall*, dan *F1-Score*, dengan fokus utama pada *Macro F1-Score* untuk menilai performa model secara seimbang pada kedua kelas. Berdasarkan hasil eksperimen yang dijalankan pada dataset berita hoaks berbahasa Indonesia dengan total 23.942 dokumen, ditemukan bahwa ketiga model menunjukkan efektivitas yang sangat tinggi dalam membedakan berita fakta dan hoaks. Namun, secara konsisten model IndoBERT memberikan performa yang paling superior di seluruh metrik evaluasi dibandingkan dengan model lainnya.

Dataset yang digunakan telah melalui proses pra-pemrosesan dan penghapusan duplikasi untuk menghindari potensi data *leakage*. Selanjutnya dataset dibagi menggunakan metode *stratified split* dengan proporsi 80% data latih dan 20% data uji sehingga distribusi kelas tetap seimbang pada kedua subset.

Evaluasi model dilakukan menggunakan metrik *Accuracy*, *Precision*, *Recall*, dan *F1-Score*, dengan fokus utama pada *Macro F1-Score* untuk menilai performa model secara seimbang pada kedua kelas.

Jumlah data yang digunakan pada eksperimen adalah:

- Data latih: 16.904 dokumen
- Data uji: 4.226 dokumen

3.1. Distribusi Dataset

Sebelum proses pelatihan model, dataset dianalisis untuk mengetahui distribusi kelas pada data uji. Distribusi ini penting untuk memastikan bahwa proses evaluasi dilakukan pada dataset yang relatif seimbang.

Tabel 2. Distribusi Dataset		
Label	Kelas	Jumlah Data
0	Berita Valid	11.200
1	Berita Hoaks	12.744
	Total	23.944

Distribusi kelas menunjukkan bahwa jumlah berita valid dan hoaks relatif seimbang. Hal ini penting untuk menghindari bias model terhadap salah satu kelas.

3.2. Hasil Klasifikasi Menggunakan *Support Vector Machine* (SVM)

Tabel 3. Hasil Klasifikasi <i>Support Vector Machine</i> (SVM)				
Kelas	<i>Precision</i>	<i>Recall</i>	<i>F1-Score</i>	<i>Support</i>
Valid (0)	0.99	1.00	0.99	2175
Hoaks (1)	1.00	0.99	0.99	2051
Accuracy			99.38%	
Macro Avg	0.99	0.99	0.99	4226
Weighted Avg	0.99	0.99	0.99	4226

Model IndoBERT memberikan performa terbaik dengan nilai akurasi sebesar 99.38% dan Macro *F1-Score* sebesar 0.99. Nilai *Precision* dan *Recall* yang sangat tinggi pada kedua kelas menunjukkan bahwa model mampu mengenali pola linguistik pada berita hoaks dan berita valid secara sangat akurat.

Keunggulan IndoBERT berasal dari kemampuannya dalam memahami konteks bahasa secara mendalam melalui mekanisme *self-attention* pada arsitektur transformer. Berbeda dengan TF-IDF yang hanya merepresentasikan frekuensi kata, IndoBERT mampu menangkap hubungan semantik antar kata dalam kalimat sehingga menghasilkan representasi teks yang lebih informatif.

3.3. Hasil Klasifikasi Menggunakan *Random Forest*

Tabel 4. Hasil Klasifikasi <i>Random Forest</i>				
Kelas	<i>Precision</i>	<i>Recall</i>	<i>F1-Score</i>	<i>Support</i>
Valid (0)	0.97	0.97	0.97	2175
Hoaks (1)	0.97	0.97	0.97	2051
Accuracy			97.09%	
Macro Avg	0.97	0.97	0.97	4226
Weighted Avg	0.97	0.97	0.97	4226

Model *Random Forest* menghasilkan akurasi sebesar 97.09% dengan nilai Macro *F1-Score* sebesar 0.97. Meskipun performanya masih tergolong tinggi, hasil yang diperoleh sedikit lebih rendah dibandingkan SVM.

Hal ini disebabkan oleh karakteristik algoritma *Random Forest* yang berbasis pohon keputusan. Model ini cenderung kurang optimal dalam menangani data teks yang memiliki jumlah fitur sangat besar dan bersifat *sparse* setelah proses representasi TF-IDF. Meskipun demikian, *Random Forest* tetap mampu memberikan performa klasifikasi yang baik dengan akurasi di atas 97%.

3.4. Hasil Klasifikasi Menggunakan IndoBERT

Tabel 5. Hasil Klasifikasi IndoBERT				
Kelas	<i>Precision</i>	<i>Recall</i>	<i>F1-Score</i>	<i>Support</i>
Valid (0)	0.99	1.00	0.99	2175
Hoaks (1)	1.00	0.99	0.99	2051
Accuracy			99.38%	
Macro Avg	0.99	0.99	0.99	4226
Weighted Avg	0.99	0.99	0.99	4226

Model IndoBERT memberikan performa terbaik dengan nilai akurasi sebesar 99.38% dan Macro *F1-Score* sebesar 0.99. Nilai *Precision* dan *Recall* yang sangat tinggi pada kedua kelas menunjukkan bahwa model mampu mengenali pola linguistik pada berita hoaks dan berita valid secara sangat akurat.

Keunggulan IndoBERT berasal dari kemampuannya dalam memahami konteks bahasa secara mendalam melalui mekanisme *self-attention* pada arsitektur transformer. Berbeda dengan TF-IDF yang hanya merepresentasikan frekuensi kata, IndoBERT mampu menangkap hubungan semantik antar kata dalam kalimat sehingga menghasilkan representasi teks yang lebih informatif.

3.5. Perbandingan Performa Model

Tabel 6. Hasil Klasifikasi IndoBERT

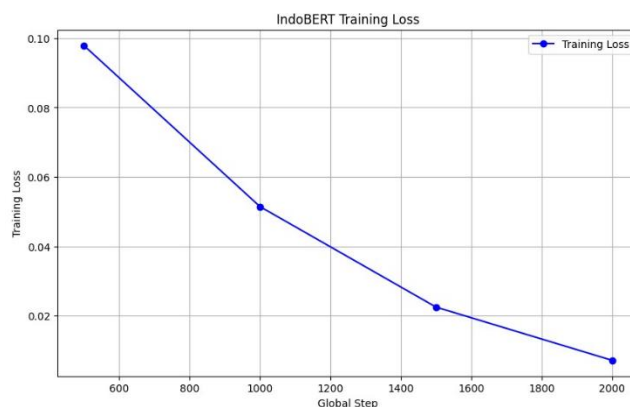
Model	Accuracy	Macro F1
<i>Random Forest</i>	97.09%	0.9708
<i>Support Vector Machine</i> (SVM)	98.04%	0.9803
IndoBERT	99.38%	0.9938

Berdasarkan hasil perbandingan, model IndoBERT menunjukkan performa terbaik dibandingkan dua model lainnya. Selisih performa antara IndoBERT dan SVM mencapai sekitar 1.34%, sedangkan perbedaan dengan *Random Forest* mencapai lebih dari 2%.

Hasil ini menunjukkan bahwa model berbasis transformer memiliki kemampuan yang lebih baik dalam memahami konteks bahasa dibandingkan metode *Machine Learning* klasik yang bergantung pada representasi frekuensi kata.

Namun demikian, model SVM masih memberikan performa yang sangat kompetitif dengan akurasi diatas 98%, sehingga dapat menjadi alternatif yang efisien ketika sumber daya komputasi terbatas. Sementara itu, *Random Forest* tetap mampu memberikan hasil yang baik meskipun performanya sedikit lebih rendah dibandingkan kedua model lainnya.

Secara keseluruhan, hasil penelitian ini menunjukkan bahwa penggunaan model berbasis transformer seperti IndoBERT memberikan peningkatan performa yang signifikan dalam tugas klasifikasi berita hoaks berbahasa Indonesia.



Gambar 2. Training loss

Grafik *training loss* pada model IndoBERT menunjukkan performa proses pembelajaran yang sangat optimal dan konvergen. Secara visual, kurva *loss* dimulai dari titik tertinggi di sekitar 0,10 pada langkah awal (*step* 500) dan mengalami penurunan yang tajam serta stabil hingga mencapai titik di bawah 0,02 pada langkah akhir (*step* 2000). Penurunan nilai *loss* yang konsisten tanpa adanya fluktuasi ekstrem ini menandakan bahwa model mampu meminimalkan kesalahan prediksi secara efektif selama proses *fine-tuning*. Hal ini membuktikan bahwa arsitektur IndoBERT berhasil mengidentifikasi pola-pola bahasa yang kompleks dalam dataset berita hoaks Indonesia, di mana parameter internal model terus menyesuaikan diri untuk mengenali perbedaan halus antara narasi berita bohong dan berita faktual.

Rendahnya nilai *loss* di akhir tahapan pelatihan memiliki korelasi langsung dengan tingginya tingkat akurasi yang dicapai, yakni sebesar 99,38%. Keadaan ini menunjukkan bahwa model tidak hanya sekedar menghafal data (*overfitting*), tetapi benar-benar mempelajari fitur semantik dari teks, mengingat penurunan *loss* terjadi secara gradual dan halus. Dalam konteks klasifikasi teks, nilai *loss* yang sangat kecil ini mengindikasikan bahwa distribusi probabilitas yang dihasilkan oleh model sangat mendekati label sebenarnya. Dengan demikian,

stabilitas pada grafik ini memberikan jaminan bahwa model IndoBERT memiliki tingkat kepercayaan yang tinggi dan ketangguhan yang sangat baik dalam melakukan inferensi terhadap data berita baru yang belum pernah dilihat sebelumnya.

4. KESIMPULAN

Penelitian ini berhasil menyimpulkan bahwa implementasi model transformer IndoBERT memberikan efektivitas tertinggi dalam klasifikasi berita hoaks berbahasa Indonesia dibandingkan algoritma *Machine Learning* konvensional, dengan capaian akurasi sebesar 99,38% dan *F1-Score* 0,99. Hasil eksperimen terhadap 23.942 dokumen menunjukkan bahwa IndoBERT secara konsisten mengungguli *Support Vector Machine* (98,04%) dan *Random Forest* (97,09%) dalam membedakan narasi berita fakta dan hoaks. Keunggulan ini didorong oleh kemampuan arsitektur transformer dalam memahami konteks semantik secara bidireksional dan penggunaan *WordPiece Tokenization* yang adaptif terhadap keragaman kosa kata, didukung oleh proses pelatihan yang stabil dengan penurunan *loss* dari 0,10 hingga di bawah 0,02. Dengan demikian, penggunaan model monolingual yang di-*fine-tune* secara spesifik terbukti menjadi solusi paling akurat dan tangguh untuk memperkuat sistem mitigasi disinformasi di ruang digital Indonesia.

DAFTAR PUSTAKA

- [1] L. Alexander, R. Gian, A. Ananta, J. Harefa, L. Alexander, and K. Jingga, "Evaluating IndoBERT for Indonesian Hoax News Detection: A Evaluating IndoBERT for Ensemble Indonesian Hoax News Detection: Comparative Study with and Models Comparative Study with Ensemble and CNN-LSTM Models," *Procedia Comput. Sci.*, vol. 269, pp. 1625–1633, 2025, doi: 10.1016/j.procs.2025.09.105.
- [2] C. J. L. Tobing, I. G. N. L. Wijayakusuma, and L. P. I. Harini, "Detection of Political Hoax News Using Fine-Tuning IndoBERT," *J. Appl. Informatics Comput.*, vol. 9, no. 2, pp. 354–360, 2025, [Online]. Available: <https://doi.org/10.30871/jaic.v9i2.8989>
- [3] A. Kunaefi, Z. Abidin, and R. Kusumawati, "Klasifikasi Berita Hoaks Bahasa Indonesia Menggunakan IndoBERT," *J. Ilm. Penelit. dan Pembelajaran Inform.*, vol. 10, no. 2, pp. 1706–1714, 2025, [Online]. Available: <https://doi.org/10.29100/jipi.v10i2.7811>
- [4] C. J. L. Tobing, I. G. N. L. Wijayakusuma, and L. P. Ida, "Perbandingan Kinerja IndoBERT dan MBERT untuk Deteksi Berita Hoaks Politik dalam Bahasa Indonesia," *J. Sains dan Teknol.*, vol. 14, no. 1, pp. 114–123, 2025, [Online]. Available: <https://doi.org/10.23887/jstundiksha.v14i1.92126%0A>
- [5] A. A. Nabil, F. R. Wildantama, D. Satrianto, M. G. Bakara, I. Budiawan, and D. Mulyat, "Klasifikasi Hoax Menggunakan Metode TF-IDF + SVM," *J. Pengabd. Masy. dan Ris. Pendidik.*, vol. 4, no. 3, pp. 17653–17660, 2026, [Online]. Available: <https://doi.org/10.31004/jerkin.v4i3.5078>
- [6] I. Fahrur, I. Maulidia, M. Hani, R. Arianto, D. R. Anantaf, and Y. A. Yuli, "Comparison of feature extraction in *Support Vector Machine* (SVM) based sentiment analysis system," *J. Ilm. Kursor*, vol. 13, no. 1, pp. 1–12, 2025, [Online]. Available: <https://doi.org/10.21107/kursor.v13i1.417>
- [7] A. U. Hamdani, S. Setiawati, Z. D. Mentari, and M. H. Purnomo, "COMPARISON OF K-NN , SVM , AND RANDOM FOREST ALGORITHM FOR DETECTING HOAX ON INDONESIAN ELECTION 2024," *J. Nas. Pendidik. Tek. Inform.*, vol. 13, no. 1, pp. 166–179, 2024, [Online]. Available: <https://doi.org/10.23887/janapati.v13i1.76079>
- [8] N. Rochmawati, A. K. Zyen, and N. A. Widiastuti, "Comparison of *Support Vector Machine* (SVM) and *Random Forest* Algorithms in the Analysis of Social Media X User Sentiment Towards the TNI Bill," *J. Appl. Informatics Comput.*, vol. 9, no. 5, pp. 2854–2860, 2025, [Online]. Available: <https://doi.org/10.30871/jaic.v9i5.10883>
- [9] R. P. Fernandes and R. T. Shita, "Penerapan Metode SVM dan *Random Forest* untuk Mendeteksi Berita Hoaks pada PT . Global Arrow," *J. Ticom Technol. Inf. Commun.*, vol. 12, no. 3, pp. 102–107, 2024, [Online]. Available: <https://doi.org/10.70309/ticom.v12i3.129>
- [10] N. S. Pramudya and Y. P. Santosa, "COMPARISON BETWEEN CNN AND RANDOM FOREST PERFORMANCE IN DETECTING HOAX INDONESIAN NEWS ARTICLES," *Proxies J.*, vol. 7, no. 1, pp. 1–14, 2023, [Online]. Available: <https://doi.org/10.24167/proxies.v7i1.12463>
- [11] M. Al Fachrozi and K. D. Tania, "COMPARISON OF NAÏVE BAYES, SVM, K-NN, *DECISION TREE*, AND *RANDOM FOREST* IN SENTIMENT ANALYSIS BASED ON SEABANK APPLICATION ASPECTS," *J. Teknol. dan Sist. Inf.*, vol. XII, no. 1, pp. 69–80, 2025, [Online]. Available: <http://dx.doi.org/10.33330/jurteks.v12i1.4189>

- [12] A. D. S. Davinka, K. D. Tania, and P. E. Sevtiyuni, "Penerapan Metode *Machine Learning* Dan Teknik SMOTE untuk Prediksi Diabetes," J. Sist. Komput. Dan Inform., vol. 7, no. 2, pp. 436–447, 2025, doi: 10.30865/json.v7i2.9032.
- [13] A. Fatihaturrahmah and K. D. Tania, "Review : A Hybrid Approach of Aspect-Based Sentiment Analysis and Knowledge Extraction for Evaluating Security Perceptions in Digital Payment Applications," Sci. J. Informatics, vol. 12, no. 4, pp. 635–650, 2025, doi: 10.15294/sji.v12i4.31557.
- [14] D. A. Ardhani and K. D. Tania, "*Knowledge discovery* on E-Commerce Customer Churn Using Interpretable *Machine Learning*: A Comparative Study of SHAP-Based Classifiers," J. Appl. Informatics Comput., vol. 9, no. 5, pp. 2695–2702, 2025, [Online]. Available: <https://doi.org/10.30871/jaic.v9i5.10811>
- [15] P. Ayuningtiyas, K. D. Tania, and W. K. Sari, "Sentiment-Based *Knowledge discovery* pada Aplikasi iPusnas Menggunakan Metode *Machine Learning* dan Deep Learning," J. Appl. Informatics Comput., vol. 9, no. 5, pp. 2486–2497, 2025, [Online]. Available: <https://doi.org/10.30871/jaic.v9i5.10258>