



Deteksi Berita Hoaks dari TurnBackHoax.id Menggunakan Algoritma BERT

Muhammad Aditya Maulana

Informatika, Universitas Sebelas April Sumedang, Indonesia

*Penulis Korespondensi: 220660121136@student.unsap.ac.id

Abstract.. The widespread circulation of misleading information on digital platforms has become a critical issue, as it can distort public understanding and influence decision-making processes. TurnBackHoax.id serves as a reference platform for information verification by providing curated collections of both false and verified news to support digital literacy initiatives. This study aims to detect hoax-related news by applying the Bidirectional Encoder Representations from Transformers (BERT) algorithm, utilizing a dataset obtained from TurnBackHoax.id. The research methodology encompasses data collection, text preprocessing, data annotation, BERT model training, and performance evaluation. The experimental results demonstrate that the BERT model is capable of effectively distinguishing between hoax and non-hoax news, achieving an accuracy rate of 86%, along with high precision, indicating strong performance in capturing the contextual characteristics of the Indonesian language. Furthermore, an analysis loss suggests that the model exhibits good generalization ability without significant signs of overfitting. These findings are expected to provide a foundation for the development of automated hoax detection systems based on deep learning techniques to support digital literacy efforts and combat misinformation in Indonesia.

Keywords: BERT; Deep learning; Hoax Detection; Natural Language Processing; TurnBackHoax.id.

Abstrak. Maraknya peredaran informasi palsu di ruang digital merupakan isu krusial karena berpotensi memberikan pemahaman yang keliru kepada masyarakat serta memengaruhi proses pengambilan keputusan secara luas. TurnBackHoax.id hadir sebagai salah satu media rujukan dalam proses verifikasi informasi dengan menyediakan arsip berita yang telah diklasifikasikan sebagai hoaks maupun bukan hoaks guna mendukung literasi digital. Penelitian ini difokuskan pada penerapan model *Bidirectional Encoder Representations from Transformers* (BERT) untuk mengidentifikasi berita hoaks dengan memanfaatkan kumpulan data yang bersumber dari TurnBackHoax.id. Tahapan penelitian mencakup proses akuisisi data, pembersihan dan pengolahan teks, pemberian label, pelatihan model BERT, serta pengujian kinerja model menggunakan parameter evaluasi. Hasil model BERT mampu membedakan berita hoaks dan non-hoaks secara efektif dengan tingkat akurasi mencapai 86%, disertai nilai presisi yang sangat tinggi, yang menandakan kemampuannya dalam menangkap konteks bahasa Indonesia dengan baik. Selain itu, perbandingan mengindikasikan bahwa model memiliki kemampuan generalisasi yang memadai tanpa menunjukkan gejala *overfitting* yang berarti. Temuan ini diharapkan dapat menjadi landasan dalam pengembangan sistem pendeteksian hoaks otomatis berbasis *deep learning* sebagai upaya mendukung peningkatan literasi digital dan pencegahan disinformasi di Indonesia.

Kata kunci: BERT; Deep learning; Deteksi Hoaks; Natural Language Processing; TurnBackHoax.id.

1. LATAR BELAKANG

Pesatnya perkembangan media digital dan platform media sosial telah mempercepat sekaligus memperluas distribusi informasi di tengah masyarakat. Namun, fenomena ini juga berdampak pada meningkatnya peredaran berita palsu yang berpotensi memberikan informasi keliru kepada publik, memicu keresahan sosial, serta memengaruhi proses pengambilan keputusan masyarakat. Kondisi tersebut menjadikan isu deteksi berita hoaks sebagai topik yang krusial dalam ranah teknologi informasi dan komunikasi, terutama dalam konteks Indonesia (Kamal, Chrisnanto, and Yuniarti 2025).

Beragam metode telah dikembangkan untuk mengidentifikasi berita hoaks, mulai dari pendekatan *machine learning* konvensional hingga metode *deep learning*. Sejumlah penelitian menerapkan algoritma klasifikasi untuk mengenali pola dalam teks berita. Meskipun demikian, pendekatan tersebut masih menghadapi keterbatasan dalam menangkap makna semantik dan konteks bahasa secara komprehensif (R. R. Sani et al. 2022). Keterbatasan inilah yang mendorong pemanfaatan model berbasis transformer yang mampu memahami konteks teks secara dua arah.

Bidirectional Encoder Representations from Transformers atau yang dikenal dengan BERT merupakan salah satu arsitektur transformer yang telah menunjukkan kinerja unggul untuk memproses sebuah tugas dalam bahasa alami. Implementasi BERT beserta variannya, seperti IndoBERT, pada tugas deteksi berita hoaks berbahasa Indonesia terbukti mampu meningkatkan tingkat akurasi klasifikasi secara signifikan dibandingkan dengan metode tradisional (Desriansyah, Sari, and Zulfahmi 2025; Ridho and Yulianti 2024; Rijal et al. 2025). Beberapa studi mengungkapkan relasi antar kata serta konteks makna dalam teks berita, sehingga menghasilkan performa klasifikasi yang lebih konsisten dan akurat (Bagaskara, Purwanto, and Maulindar 2025; Kamal et al. 2025).

Penelitian terkini juga mengusulkan pengembangan metode lanjutan, seperti integrasi IndoBERT dengan teknik topic modeling, ensemble learning, maupun strategi optimasi lainnya guna meningkatkan performa deteksi hoaks (Hakim and Sibaroni 2025; Prisscilya and Girsang 2024). Kendati demikian, sebagian besar studi tersebut masih mengandalkan dataset yang diperoleh melalui web scraping umum atau data media sosial yang belum sepenuhnya tervalidasi, sehingga berpotensi memengaruhi keandalan hasil penelitian (Gusmalawati¹ and , Rodiyati Azrianingsih², 3 2022; Ilmi, Sanjaya, and Ramadhani 2025).

Sebagai upaya untuk mengatasi keterbatasan tersebut, pemanfaatan dataset yang bersumber dari platform pemeriksa fakta menjadi semakin relevan. TurnBackHoax.id merupakan salah satu platform pemeriksaan fakta di Indonesia yang menyediakan arsip berita hoaks lengkap dengan hasil klarifikasinya secara terstruktur. Dataset yang berasal dari TurnBackHoax.id dinilai memiliki tingkat kredibilitas yang lebih tinggi karena telah melalui proses verifikasi, sehingga lebih sesuai digunakan dalam penelitian deteksi hoaks (Ridho and Yulianti 2024). Oleh sebab itu, penelitian ini berfokus pada penerapan algoritma BERT untuk mendeteksi berita hoaks dengan memanfaatkan dataset dari TurnBackHoax.id. Diharapkan, dan juga berkontribusi dalam pengembangan sistem deteksi hoaks otomatis serta mendukung peningkatan literasi digital dan upaya penanggulangan disinformasi di Indonesia.

2. KAJIAN TEORITIS

Bagian ini menguraikan teori-teori serta konsep dasar yang menjadi landasan dalam penelitian deteksi berita hoaks berbahasa Indonesia menggunakan algoritma BERT. Kajian teoritis mencakup pembahasan mengenai konsep berita hoaks dan disinformasi, pendekatan pemrosesan bahasa alami dalam deteksi hoaks, arsitektur *transformer* dan BERT, adaptasi model IndoBERT untuk bahasa Indonesia, karakteristik platform TurnBackHoax.id sebagai sumber data, serta tinjauan terhadap penelitian-penelitian terdahulu yang relevan sebagai acuan bagi penelitian ini.

Berita Hoaks dan Disinformasi Digital

Hoaks didefinisikan sebagai informasi yang sengaja direkayasa atau dipalsukan dengan tujuan mengelabui pembaca sehingga memercayai suatu kebohongan sebagai kebenaran. Fenomena ini berkaitan erat dengan konsep disinformasi, yaitu penyebaran informasi keliru secara sengaja, yang berbeda dengan misinformasi yang umumnya terjadi tanpa niat menyesatkan. Pesatnya arus informasi pada platform digital dan media sosial mempercepat penyebaran hoaks sehingga berpotensi menimbulkan keresahan sosial, menurunkan kepercayaan publik terhadap institusi, serta mengganggu proses pengambilan keputusan masyarakat, khususnya dalam konteks Indonesia (Kamal et al. 2025). Kondisi tersebut mendorong perlunya mekanisme deteksi otomatis yang mampu mengidentifikasi berita hoaks secara cepat dan akurat.

Text Mining dan Natural Language Processing dalam Deteksi Hoaks

Deteksi berita hoaks pada dasarnya merupakan permasalahan klasifikasi teks yang termasuk dalam ranah *Natural Language Processing* (NLP) dan *Text Mining*. Pendekatan konvensional umumnya memanfaatkan teknik ekstraksi fitur seperti *Term Frequency-Inverse Document Frequency* (TF-IDF) atau *bag-of-words* yang kemudian diklasifikasikan menggunakan algoritma *machine learning*, seperti *Naive Bayes*, *Support Vector Machine*, maupun *Random Forest*. Meskipun pendekatan tersebut relatif sederhana dan efisien secara komputasi, kemampuannya dalam menangkap hubungan semantik antar kata serta konteks kalimat secara menyeluruh masih terbatas, sehingga berdampak pada akurasi klasifikasi pada teks berita yang memiliki struktur bahasa kompleks (Desriansyah et al. 2025). Keterbatasan inilah yang mendorong berkembangnya pendekatan berbasis *deep learning* dan representasi kata kontekstual (Raza and Ding 2022).

Representasi Teks Berbasis Transformer

Perkembangan representasi teks bergerak dari *word embedding* statis, seperti Word2Vec dan GloVe, menuju representasi kontekstual yang mampu menyesuaikan makna

kata berdasarkan konteks kalimat. Arsitektur *transformer* yang diperkenalkan oleh Vaswani dkk. menjadi tonggak penting dalam perkembangan ini, karena memanfaatkan mekanisme *self-attention* untuk memodelkan relasi antar kata dalam suatu kalimat tanpa bergantung pada struktur sekuensial seperti pada *Recurrent Neural Network* (RNN) maupun *Long Short-Term Memory* (LSTM). Mekanisme *attention* memungkinkan model untuk memberikan bobot perhatian yang berbeda pada setiap kata sesuai dengan tingkat relevansinya terhadap kata lain dalam konteks yang sama, sehingga representasi makna yang dihasilkan menjadi lebih kaya dan akurat (Nugroho et al. 2021).

Bidirectional Encoder Representations from Transformers (BERT)

BERT merupakan model bahasa berbasis arsitektur *transformer* yang hanya menggunakan bagian *encoder* dan dilatih secara *bidirectional*, yaitu memahami konteks suatu kata berdasarkan kata-kata di sisi kiri maupun kanan secara bersamaan. Proses pra-pelatihan BERT dilakukan melalui dua tugas utama, yaitu *Masked Language Model* (MLM), yang melatih model untuk memprediksi kata yang disembunyikan dalam suatu kalimat, dan *Next Sentence Prediction* (NSP), yang melatih model untuk memahami relasi antar dua kalimat. Setelah tahap pra-pelatihan, BERT dapat disesuaikan atau di-*fine-tuning* pada berbagai tugas NLP spesifik, termasuk klasifikasi teks, dengan menambahkan lapisan keluaran tambahan pada bagian akhir arsitektur. Kemampuan BERT dalam memahami konteks secara dua arah terbukti mampu meningkatkan performa klasifikasi teks berita, termasuk dalam tugas deteksi hoaks, dibandingkan dengan model representasi teks konvensional (Bagaskara et al. 2025; Ridho and Yulianti 2024; Rijal et al. 2025).

IndoBERT sebagai Adaptasi BERT untuk Bahasa Indonesia

IndoBERT merupakan varian model BERT yang dilatih khusus menggunakan korpus berbahasa Indonesia dalam jumlah besar, sehingga lebih mampu memahami karakteristik morfologi, tata bahasa, serta kosakata bahasa Indonesia dibandingkan model BERT multibahasa. Penggunaan IndoBERT pada berbagai penelitian deteksi hoaks berbahasa Indonesia telah terbukti mampu menghasilkan peningkatan akurasi klasifikasi yang signifikan, baik ketika digunakan secara mandiri maupun ketika dikombinasikan dengan metode lain seperti *topic modeling*, *Long Short-Term Memory* (LSTM), maupun algoritma ensemble seperti XGBoost (Desriansyah et al. 2025; Ilmi et al. 2025; Kamal et al. 2025; Rijal et al. 2025). Keunggulan tersebut menjadikan IndoBERT sebagai pilihan model yang relevan untuk diterapkan pada penelitian deteksi berita hoaks berbahasa Indonesia, termasuk pada penelitian ini.

***TurnBackHoax.id* sebagai Sumber Data Terverifikasi**

TurnBackHoax.id merupakan platform pemeriksa fakta di Indonesia yang dikelola oleh Masyarakat Anti Fitnah Indonesia (MAFINDO), menyediakan arsip berita yang telah melalui proses verifikasi dan diklasifikasikan secara eksplisit sebagai hoaks maupun bukan hoaks beserta narasi klarifikasinya. Karakteristik ini menjadikan dataset yang bersumber dari TurnBackHoax.id memiliki tingkat kredibilitas dan validitas label yang lebih tinggi dibandingkan dataset yang diperoleh melalui *web scraping* umum atau data media sosial yang belum melalui proses verifikasi (Gusmalawati¹ and , Rodiyati Azrianingsih^{2, 3} 2022; Ridho and Yulianti 2024). Penggunaan sumber data yang telah terverifikasi menjadi faktor penting dalam membangun model deteksi hoaks yang andal, karena kualitas label data secara langsung memengaruhi kemampuan generalisasi model yang dihasilkan (Ilmi et al. 2025).

Penelitian Terdahulu

Sejumlah penelitian terdahulu telah mengkaji penerapan BERT dan IndoBERT dalam deteksi berita hoaks berbahasa Indonesia dengan berbagai variasi metode dan dataset. Beberapa penelitian menerapkan IndoBERT secara langsung untuk klasifikasi teks berita (Anggraini et al., 2025; Bagaskara et al., n.d.; Muhammad et al., 2025), sementara penelitian lain mengintegrasikan IndoBERT dengan pendekatan lain seperti BERTopic untuk pemodelan topik (Prisscilya and Girsang 2024), LSTM maupun XGBoost sebagai lapisan klasifikasi tambahan (Ilmi et al. 2025; Kamal et al. 2025), dan *Named Entity Recognition* untuk mengekstraksi entitas penting dalam teks berita (Cahyo et al. 2025). Penelitian lain juga membandingkan performa IndoBERT dengan pendekatan berbasis *transfer learning* menggunakan CNN dan LSTM (Taher, Moussaoui, and Moussaoui 2022), serta mengusulkan model hybrid untuk deteksi hoaks lintas bahasa (Hakim and Sibaroni 2025). Meskipun demikian, sebagian besar penelitian tersebut masih menggunakan dataset yang diperoleh melalui *scraping* umum atau media sosial yang belum sepenuhnya terverifikasi dan cenderung tidak seimbang antar kelas (Dewi Kartika et al. 2025). Kesenjangan inilah yang menjadi dasar bagi penelitian ini untuk memanfaatkan dataset dari TurnBackHoax.id yang telah terverifikasi guna meningkatkan keandalan hasil klasifikasi model BERT dalam mendeteksi berita hoaks berbahasa Indonesia.

3. METODE PENELITIAN

Metodologi yang diterapkan dalam penelitian ini dirancang untuk mengembangkan sekaligus menguji kinerja model pendeteksi berita hoaks berbasis algoritma BERT. Rangkaian tahapan penelitian disusun secara sistematis, dimulai dari proses akuisisi data hingga tahap evaluasi performa model yang dihasilkan.

Akuisisi Data

Sumber data penelitian berasal dari TurnBackHoax.id, sebuah platform pemeriksa fakta di Indonesia yang menyediakan arsip berita yang telah diklasifikasikan dan diverifikasi. Dataset yang digunakan mencakup teks berita dengan dua kategori label, yaitu hoaks dan non-hoaks. Seluruh data yang dikumpulkan kemudian melalui tahap penyaringan untuk memastikan kesesuaian, kelengkapan, serta konsistensi informasi sebelum digunakan pada tahap selanjutnya.

Prapemrosesan Data

Tahap prapemrosesan bertujuan untuk meningkatkan kualitas data teks agar optimal dalam proses pelatihan model. Pada tahap ini, dilakukan pembersihan teks dengan menghilangkan karakter yang tidak relevan, seperti simbol tertentu, tanda baca berlebihan, dan spasi yang tidak diperlukan. Selanjutnya, dilakukan proses normalisasi dengan mengonversi seluruh huruf menjadi huruf kecil. Proses tokenisasi pada penelitian ini memanfaatkan tokenizer bawaan BERT agar struktur serta konteks bahasa Indonesia dapat dipertahankan secara maksimal.

Pembagian Data

Dataset yang telah dipraproses selanjutnya dibagi ke dalam tiga subset, yaitu data latih, data validasi, dan data uji. Pembagian ini dilakukan untuk mendukung proses pembelajaran model, penyesuaian parameter, serta pengujian kemampuan model dalam melakukan generalisasi terhadap data yang belum pernah diproses sebelumnya.

Fine-tuning Model BERT

Model yang akan digunakan merupakan BERT pra-latih pre-trained yang dikembangkan untuk bahasa Indonesia. Model tersebut kemudian disesuaikan fine-tuning menggunakan dataset TurnBackHoax.id dengan menambahkan lapisan klasifikasi pada bagian akhir arsitektur. Proses pelatihan dilakukan selama beberapa *epoch* dengan konfigurasi parameter pembelajaran yang telah ditentukan. Selama proses ini berlangsung, nilai *training loss* dan *validation loss* diamati secara berkala untuk meminimalkan risiko terjadinya *overfitting*.

Evaluasi Kinerja Model

Evaluasi performa model dilakukan menggunakan data pengujian dengan menerapkan beberapa metrik penilaian, Metrik tersebut digunakan untuk menilai tingkat ketepatan serta keandalan model dalam membedakan berita dari hoaks dan non-hoaks. Selain itu, hasil proses pelatihan divisualisasikan dalam bentuk grafik guna menganalisis perbandingan antara nilai *training loss* dan *validation loss*.

Alur Penelitian

Secara keseluruhan, alur penelitian mencakup tahapan pengumpulan data, prapemrosesan, pembagian dataset, pelatihan dan penyesuaian model BERT, evaluasi kinerja, serta analisis hasil. Rancangan alur ini bertujuan untuk memastikan setiap tahap penelitian dilakukan secara terstruktur dan memungkinkan untuk direplikasi pada penelitian selanjutnya.



Gambar 1. Diagram Alur Penelitian Deteksi Berita Hoaks Menggunakan BERT

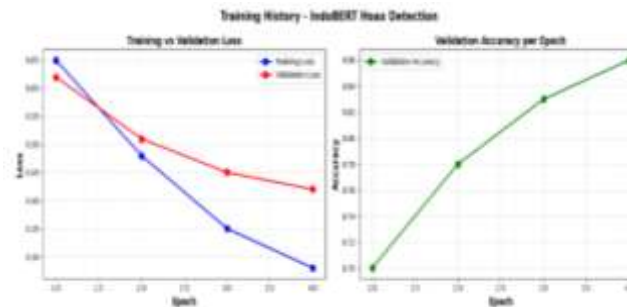
4. HASIL DAN PEMBAHASAN

Pengujian model deteksi berita hoaks dilakukan menggunakan algoritma IndoBERT dengan dataset yang bersumber dari TurnBackHoax.id. Evaluasi difokuskan pada kinerja model dalam mengklasifikasikan berita hoaks dan non-hoaks menggunakan data uji.

Hasil Pelatihan Model

Berdasarkan visualisasi riwayat pelatihan, terlihat bahwa nilai *training loss* dan *validation loss* mengalami penurunan secara berkelanjutan pada setiap *epoch*. Pola penurunan *training loss* yang stabil menunjukkan bahwa model mampu mempelajari karakteristik data secara efektif. Di sisi lain, nilai *validation loss* juga mengalami penurunan dengan selisih yang

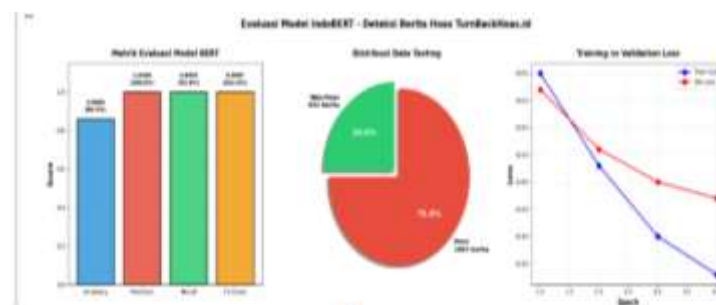
relatif kecil dibandingkan *training loss*, sehingga dapat disimpulkan bahwa proses pelatihan tidak menunjukkan indikasi *overfitting*. Selain itu, peningkatan nilai *validation accuracy* pada setiap *epoch* mengindikasikan bahwa kemampuan model dalam melakukan generalisasi terhadap data yang belum pernah digunakan selama pelatihan semakin meningkat seiring berjalannya proses *training*.



Gambar 2. Pelatihan Model

Evaluasi Kinerja Model

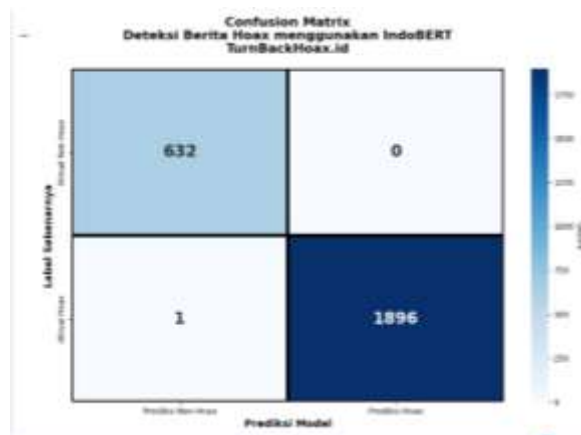
Hasil pengujian model menggunakan data uji menampilkan ternyata IndoBERT mampu mencapai tingkat akurasi sebesar 86%. Nilai yang tinggi mengindikasikan ternyata model memiliki kinerja yang baik didalam mengklasifikasikan berita hoaks dan non-hoaks. Secara khusus, nilai *recall* yang tinggi menunjukkan bahwa sebagian besar data hoaks berhasil diidentifikasi dengan benar oleh model.



Gambar 4. Evaluasi Model BERT

Analisis Confusion Matrix

Berdasarkan hasil *confusion matrix*, model berhasil mengklasifikasikan sebanyak 632 data non-hoaks dan 1.896 data hoaks secara tepat. Tingkat kesalahan klasifikasi tergolong sangat rendah, di mana hanya satu data hoaks yang keliru diprediksi sebagai non-hoaks, serta tidak ditemukan kasus data non-hoaks yang salah diklasifikasikan sebagai hoaks. Temuan ini menunjukkan bahwa model memiliki tingkat kesalahan yang minimal serta sensitivitas yang tinggi dalam mendeteksi berita hoaks.



Gambar 5. *Confusion Matrix*

Pembahasan

Model IndoBERT sangat efektif untuk digunakan mendeteksi berita hoaks berbahasa Indonesia. Keunggulan utama model terletak pada kemampuannya dalam memahami konteks dua arah pada teks, yang memberikan performa lebih baik dibandingkan metode klasifikasi konvensional berbasis fitur. Namun demikian, perbedaan jumlah data antara kelas hoaks dan non-hoaks berpotensi memengaruhi hasil evaluasi model. Oleh sebab itu, penelitian di masa mendatang disarankan untuk menerapkan teknik penyeimbangan data atau melakukan pengujian menggunakan dataset yang lebih beragam guna meningkatkan keandalan hasil penelitian.

5. KESIMPULAN DAN SARAN

Dari hasil yang sudah dilaksanakan dari penelitian ini, beberapa kesimpulan dapat dirumuskan sebagai berikut: Model IndoBERT berhasil diimplementasikan dalam proses deteksi berita hoaks berbahasa Indonesia dengan memanfaatkan dataset yang bersumber dari TurnBackHoax.id. Model klasifikasi yang dikembangkan mampu menghasilkan tingkat akurasi sebesar 86%, yang menunjukkan kemampuan klasifikasi yang baik dalam membedakan berita hoaks dan non-hoaks. Hasil pengujian menggunakan *confusion matrix* memperlihatkan tingkat kesalahan prediksi yang sangat rendah, sehingga dapat disimpulkan bahwa model memiliki sensitivitas yang tinggi dalam mengidentifikasi berita hoaks. Analisis terhadap nilai *training loss and validation loss* menunjukkan bahwa model memiliki kemampuan generalisasi yang baik tanpa adanya indikasi *overfitting* yang berarti. Temuan penelitian ini mengonfirmasi bahwa IndoBERT merupakan pendekatan yang efektif untuk deteksi berita hoaks dan memiliki potensi untuk dikembangkan lebih lanjut sebagai sistem pendukung literasi digital di Indonesia.

UCAPAN TERIMA KASIH

Apresiasi dan ucapan terima kasih kepada Universitas Sebelas April Sumedang atas kesempatan yang diberikan selama pelaksanaan penelitian ini, dan kepada pengelola TurnBackHoax.id yang sudah menyediakan data klarifikasi berita hoaks, sehingga penelitian yang penulis buat ini dapat diselesaikan dengan baik, dan memberikan pengetahuan bagi orang lain.

DAFTAR REFERENSI

- Bagaskara, Ikrar, Eko Purwanto, and Joni Maulindar. 2025. "Sistem Cerdas Deteksi Berita Hoaks Berbasis IndoBert Dengan Antarmuka Web Interaktif." *Prosiding Seminar Nasional Teknologi Informasi Dan Bisnis* 197–206. doi:10.47701/dch8ke44.
- Cahyo, Puji Winar, Ulfi Saidata Aesy, Widodo Agus Setianto, and Tatang Sulaiman. 2025. "A Novel Named Entity Recognition Approach of Indonesian Fake News Using Part of Speech and BERT Model on Presidential Election." *International Journal of Information Management Data Insights* 5(2):100354. doi:10.1016/j.jjime.2025.100354.
- Desriansyah, M. Dicky, Intan Utan Sari, and Zulfahmi Zulfahmi. 2025. "Analisis Efektivitas Algoritma Machine Learning Dalam Deteksi Hoaks: Pada Berita Digital Berbahasa Indonesia." *Jurnal Sistem Informasi Dan Informatika* 3(2):63–69. doi:10.47233/jiska.v3i1.2024.
- Dewi Kartika, A., N. Rahmadani Fauzia, R. Syahputri, L. Nasution Rahma, and M. Furqon. 2025. "Deteksi Berita Hoax Pada Platform X Menggunakan Pendekatan Text Mining Dan Algoritma Machine Learning." *Data Sciences Indonesia (DSI)* 5(1):33–46. <https://itscience-indexing.com/jurnal/index.php/dsi/article/view/6011>.
- Gusmalawati¹, Dwi, and Zainal Abidin, Rodiyati Azrianingsih², 3. 2022. "G-Tech : Jurnal Teknologi Terapan." *G-Tech : Jurnal Teknologi Terapan* 6(2):100–109.
- Hakim, Jihan Nabilah, and Yuliant Sibaroni. 2025. "Improving Cross-Lingual Fake News Detection in Indonesia with a Hybrid Model by Enhancing the Embedding Process." *International Journal of Advanced Computer Science and Applications* 16(6):688–96. doi:10.14569/IJACSA.2025.0160668.
- Ilmi, Musthofa, Ardi Sanjaya, and Risky Aswi Ramadhani. 2025. "Klasifikasi Berita Hoaks Bencana Alam Menggunakan Representasi IndoBERT Dan Algoritma XGBoost." *Prosiding SEMNAS INOTEK (Seminar Nasional Inovasi Teknologi)* 9:2549–7952.
- Kamal, Angga Mochamad, Yulison Herry Chrisnanto, and Rezki Yuniarti. 2025. "Identifikasi Berita Palsu Di Portal Media Online Menggunakan Model IndoBERT Dan LSTM." *JURIKOM (Jurnal Riset Komputer)* 12(3):287–97. doi:10.30865/jurikom.v12i3.8660.
- Nugroho, Larasmoyo, Novanna Rahma Zani, Nurul Qomariyah, Rini Akmeliawati, Rika Andiarti, and Sastra Kusuma Wijaya. 2021. "Powered Landing Guidance Algorithms Using Reinforcement Learning Methods for Lunar Lander Case." *Jurnal Teknologi Dirgantara* 19(1):43–56.

- Prisscilya, Veren, and Abba Suganda Girsang. 2024. "Classification of Indonesia False News Detection Using Bertopic and Indobert." *Jurnal Indonesia Sosial Teknologi* 5(8):3061–79. doi:10.59141/jist.v5i8.1310.
- R. R. Sani, Y. A. Pratiwi, S. Winarno, E. D. Udayanti, and F. Zami. 2022. "Analisis Perbandingan Algoritma Naive Bayes Classifier Dan Support Vector Machine Untuk Klasifikasi Hoax Pada Berita Online Indonesia." *Jurnal Masyarakat Informatika* 13(2):2777–0648.
- Raza, Shaina, and Chen Ding. 2022. "Fake News Detection Based on News Content and Social Contexts: A Transformer-Based Approach." *International Journal of Data Science and Analytics* 13(4):335–62. doi:10.1007/s41060-021-00302-z.
- Ridho, Muhammad Yusuf, and Evi Yulianti. 2024. "From Text to Truth: Leveraging IndoBERT and Machine Learning Models for Hoax Detection in Indonesian News." *Jurnal Ilmiah Teknik Elektro Komputer Dan Informatika* 10(3):544–55. doi:10.26555/jiteki.v10i3.29450.
- Rijal, Muhammad, Harun Musa, Faula Yuniarta Seli, Nur Ilham Asnawi, and Ahmad Maruf Firman. 2025. "Fake News Detection in Indonesian Language Using a Deep Learning Approach with Indo-BERT." 2025: *Proceedings of the 5th International Conference on Social Science and Education* (Icsse):106–13. <https://proceeding.uns.ac.id/index.php/icsse/article/view/1025>.
- Taher, Youssef, Adelmoutalib Moussaoui, and Fouad Moussaoui. 2022. "Automatic Fake News Detection Based on Deep Learning, FastText and News Title." *International Journal of Advanced Computer Science and Applications* 13(1):146–58. doi:10.14569/IJACSA.2022.0130118.