

Analisis Tren dan Kebaruan Pendekatan Hybrid Transformer untuk Meningkatkan Akurasi Deteksi Hoaks Berbahasa Indonesia

Muhayat^{1*}, Husni Naparin², Rifqi Mulyawan³

^{1 2 3} Program Studi Teknologi Informasi, Universitas Islam Negeri Antasari Banjarmasin, Indonesia

Email: ¹muhayat@uin-antasari.ac.id, ²parinhusnina@uin-antasari.ac.id, ³rifqi.mulyawan@uin-antasari.ac.id

Email Penulis Korespondensi: muhayat@uin-antasari.ac.id

Artikel Info :

Artikel History :

Submitted : 13-04-2026

Accepted : 20-04-2026

Published : 30-04-2026

Kata Kunci:

Deteksi Hoaks; NLP;

Deep learning;

Transformer Models;

Hybrid Approach

Abstrak– Penelitian ini bertujuan untuk menganalisis perkembangan metodologis, kesenjangan penelitian, serta kebaruan pendekatan dalam teknik deteksi hoaks berbahasa Indonesia yang semakin kompleks seiring meningkatnya penyebaran misinformasi digital. Permasalahan utama yang diangkat adalah keterbatasan model konvensional dalam memahami konteks linguistik, ketidakseimbangan data, serta rendahnya interpretabilitas model. Penelitian ini menggunakan metode literature review non-systematic terhadap 31 artikel ilmiah terindeks nasional dan internasional periode 2019–2025 dengan pendekatan analisis tematik dan komparatif. Hasil penelitian menunjukkan adanya peningkatan signifikan performa model dari pendekatan machine learning klasik dengan akurasi 85–97%, ke deep learning sebesar 90–99%, hingga model transformer-based seperti IndoBERT dan RoBERTa yang mencapai akurasi 94–99,5%. Pendekatan hybrid seperti Bi-LSTM + IndoBERT dan CNN-LSTM menunjukkan performa paling stabil dengan akurasi hingga 99,58% serta peningkatan generalisasi model. Selain itu, teknik augmentasi data seperti back-translation terbukti meningkatkan akurasi secara signifikan hingga lebih dari 5%. Namun demikian, masih ditemukan kesenjangan pada aspek interpretabilitas, efisiensi komputasi, serta integrasi multimodal. Kesimpulannya, pendekatan hybrid berbasis transformer merupakan solusi paling efektif dalam meningkatkan akurasi dan stabilitas deteksi hoaks, meskipun masih diperlukan pengembangan pada aspek Explainable AI dan multimodalitas untuk meningkatkan keandalan sistem di dunia nyata..

Abstract– This study aims to analyze methodological developments, research gaps, and the novelty of approaches in Indonesian-language hoax detection, which has become increasingly complex due to the rapid spread of digital misinformation. The primary issues addressed include the limitations of conventional models in understanding linguistic context, handling imbalanced data, and providing model interpretability. This research employs a non-systematic literature review of 31 national and international scientific articles published between 2019 and 2025, using thematic and comparative analysis approaches. The results indicate a significant improvement in model performance, starting from classical machine learning approaches with accuracy ranging from 85% to 97%, advancing to deep learning models achieving 90% to 99%, and further to transformer-based models such as IndoBERT and RoBERTa, which reach accuracy levels between 94% and 99.5%. Hybrid approaches, including Bi-LSTM combined with IndoBERT and CNN-LSTM architectures, demonstrate the most stable performance with accuracy up to 99.58% and improved model generalization. In addition, data augmentation techniques such as back-translation have been shown to increase accuracy by more than 5%. However, several challenges remain, particularly in terms of model interpretability, computational efficiency, and the integration of multimodal data. In conclusion, hybrid transformer-based approaches represent the most effective solution for improving the accuracy and robustness of hoax detection systems. Nevertheless, further research is required to enhance explainability and multimodal capabilities in order to achieve more reliable real-world applications.

Keywords :

Deteksi Hoaks; NLP;

Deep learning;

Transformer Models;

Hybrid Approach

1. PENDAHULUAN

Perkembangan teknologi digital dan media sosial dalam beberapa tahun terakhir telah mendorong percepatan distribusi informasi secara masif, namun di sisi lain juga meningkatkan penyebaran misinformasi dan hoaks di ruang publik digital. Fenomena ini menjadi perhatian serius karena berdampak langsung terhadap stabilitas sosial, politik, dan ekonomi masyarakat. Studi sebelumnya menunjukkan bahwa lebih dari 60% pengguna internet cenderung mempercayai informasi hoaks tanpa proses verifikasi yang memadai, sehingga memperkuat urgensi pengembangan sistem deteksi hoaks yang lebih efektif dan adaptif [1]. Seiring dengan kemajuan teknologi, bidang *Natural Language Processing* (NLP) dan *Artificial Intelligence* mengalami perkembangan pesat yang mendorong transformasi metodologis dalam deteksi hoaks. Pendekatan yang awalnya didominasi oleh metode berbasis fitur statistik seperti *Support Vector Machine* dan *Naive Bayes* kini telah berkembang menuju model berbasis *deep learning* dan *transformer*. Model seperti IndoBERT dan RoBERTa terbukti memiliki kemampuan

yang lebih baik dalam memahami konteks semantik, relasi antar kata, serta struktur kalimat kompleks, sehingga mampu meningkatkan performa deteksi secara signifikan [1], [2].

Tahapan awal penelitian di bidang deteksi hoaks menunjukkan dominasi pendekatan *machine learning* klasik yang menitikberatkan pada representasi teks berbasis fitur statistik. Tandiano dan Jollyta [3] mengembangkan model klasifikasi berita palsu menggunakan SVM kernel linier dengan metode n-gram, yang terbukti efektif dengan akurasi di atas 97%. Jollyta et al. [4] melanjutkan pendekatan ini dan menemukan bahwa kombinasi unigram dan kernel linier tetap menghasilkan performa terbaik dibandingkan model lain. Metode Naive Bayes juga banyak digunakan dalam studi-studi awal [5], terutama karena kesederhanaannya dalam memproses teks besar secara efisien. Syaifuddin et al. [6] memperkuat arah tersebut dengan mengombinasikan ensemble classifier berbasis n-gram, TF-IDF, dan passive-aggressive method yang menghasilkan akurasi 91,8%. Pendekatan klasik ini kemudian dilengkapi dengan teknik oversampling seperti SMOTE untuk mengatasi ketidakseimbangan data [7]. Dari temuan-temuan ini terlihat bahwa fase awal pengembangan deteksi hoaks menitikberatkan pada pencarian konfigurasi fitur yang optimal dengan algoritma klasifikasi yang efisien, meskipun belum mampu memahami konteks semantik dan relasi antar kalimat secara mendalam.

Kemajuan signifikan berikutnya ditandai dengan transisi menuju era *deep learning*, di mana model seperti LSTM dan CNN mulai digunakan untuk mempelajari urutan dan pola linguistik dalam teks hoaks. Yusuf dan Suyanto [8] menunjukkan bahwa model LSTM dengan Word2Vec Skip-gram mampu mencapai akurasi 89,4%, melebihi metode klasik sebelumnya. Mashuri et al. [9] kemudian menggabungkan pendekatan lexicon-based dengan LSTM untuk meningkatkan pemahaman terhadap konteks emosional, menghasilkan akurasi 99%. Sementara itu, Hati dan Sulistiani [10] memperkenalkan CNN dengan GloVe *embedding* yang mencapai *validation accuracy* hampir sempurna sebesar 99,88%. Sibaroni et al. [11] bahkan menunjukkan bahwa kombinasi CNN-LSTM menghasilkan akurasi 96% untuk deteksi hoaks politik, memperlihatkan keunggulan arsitektur hibrida dalam menangani data yang kompleks. Dari berbagai penelitian ini dapat disimpulkan bahwa fase *deep learning* menandai pergeseran paradigma dari model berbasis fitur statis menuju model yang memahami dinamika semantik dan struktur naratif teks secara mendalam. Hal ini juga sejalan dengan temuan Mulyawan et al. [30] yang menunjukkan bahwa penggunaan Word2Vec sebagai representasi fitur mampu meningkatkan performa klasifikasi dibandingkan pendekatan berbasis frekuensi seperti TF-IDF.

Transformasi metodologis paling mutakhir dalam deteksi hoaks muncul dengan diperkenalkannya model berbasis transformer, khususnya IndoBERT yang dioptimalkan untuk teks berbahasa Indonesia. Yakub et al. [1] berhasil melakukan fine-tuning IndoBERT untuk deteksi hoaks dan memperoleh akurasi 98,51% serta F1-score 98,44%. Fathin et al. [12] memperluas penelitian ini dengan menangani ketidakseimbangan data melalui random sampling, menghasilkan akurasi 98,2%. Jocelyne et al. [13] menerapkan IndoBERT untuk mendeteksi hoaks politik dan mencatat akurasi 94,1% dengan ROC AUC 0,991, menunjukkan ketangguhan model terhadap variasi semantik yang kompleks. Di sisi lain, Rifai et al. [2] membuktikan efektivitas model RoBERTa dalam mendeteksi hoaks COVID-19, mencapai akurasi di atas 97% dengan efisiensi pelatihan 50% lebih cepat dibandingkan LSTM. Temuan-temuan ini memperlihatkan bahwa model transformer memiliki keunggulan dalam menangani konteks linguistik panjang, hubungan antarfrasa, dan keunikan struktur kalimat dalam teks hoaks.

Meskipun capaian akurasi model-model terbaru sangat tinggi, sejumlah tantangan metodologis dan teknis masih menjadi kesenjangan penelitian yang belum terpecahkan. Salah satu isu utama adalah ketidakseimbangan data antara teks hoaks dan non-hoaks yang sering kali menyebabkan bias pada hasil klasifikasi. Fathin et al. [12] dan Hendrawan et al. [7] telah mencoba mengatasi masalah ini dengan teknik *resampling* dan *oversampling*, namun belum ada pendekatan baku yang diakui secara luas. Tantangan lain adalah keterbatasan kemampuan model dalam menangani variasi linguistik seperti bahasa campuran, gaya tutur informal, atau penggunaan sarkasme yang sering muncul dalam teks digital. Selain itu, sebagian besar penelitian masih berfokus pada data berbasis teks dan belum banyak mengeksplorasi pendekatan multimodal yang menggabungkan teks, gambar, dan video sebagaimana diteliti oleh Jin et al. [14] dan Mohawesh et al. [15]. Permasalahan lain yang tak kalah penting adalah keterbatasan interpretabilitas model *deep learning* dan transformer, yang sering kali dianggap sebagai “kotak hitam” dalam pengambilan keputusan. Hanya Syari et al. [16] yang mencoba mengintegrasikan pendekatan *critical thinking* dan *Explainable AI* (XAI) untuk memberikan transparansi terhadap hasil deteksi.

Kesenjangan konseptual juga ditemukan dalam hal replikasi dan generalisasi hasil penelitian. Banyak model hanya diuji pada satu dataset tanpa pengujian lintas domain atau validasi terhadap data waktu nyata. Laimeheriwa et al. [17] dan Kozik et al. [18] menyoroti pentingnya melakukan meta-analisis dan perbandingan komprehensif terhadap berbagai metode *state-of-the-art* guna memahami performa sebenarnya dari model deteksi hoaks di berbagai konteks. Selain itu, penelitian real-time masih sangat terbatas; Wibowo et al. [19] dan Hati & Sulistiani [10] baru menampilkan implementasi berbasis web sederhana tanpa pengujian skala besar atau evaluasi efisiensi sistem. Dalam konteks efisiensi komputasi, Rifai et al. [2] menunjukkan bahwa RoBERTa lebih cepat dibandingkan LSTM, namun perbandingan dengan model-model terbaru seperti IndoBERT atau GPT-based LLM belum banyak dilakukan. Dengan demikian, meskipun literatur menunjukkan kemajuan signifikan dalam akurasi, masih terdapat ruang besar untuk meningkatkan interpretabilitas, efisiensi, dan kemampuan adaptasi model terhadap variasi linguistik dan multimodalitas data.

Di sisi lain, berbagai inovasi telah muncul dalam upaya meningkatkan keandalan model dan mengatasi keterbatasan data. Noor et al. [20] menggunakan *Synonym-based Data Augmentation* untuk memperkaya variasi teks dan meningkatkan robustness model CNN. Fariha dan Lhaksmana [21] memperluas pendekatan ini dengan back-translation, meningkatkan akurasi model Bi-LSTM + IndoBERT dari 94,14% menjadi 99,58%. Pendekatan augmentasi ini menunjukkan kebaruan signifikan dalam memperkuat generalisasi model tanpa menambah data asli. Selain itu, Mashuri et al. [9] memperkenalkan kombinasi *lexicon-based sentiment analysis* dan LSTM untuk memahami narasi emosional dalam berita palsu, memberikan arah baru dalam deteksi hoaks berbasis konteks emosi. Penelitian seperti Li et al. [22] bahkan telah mulai mengadopsi sistem *agentic reasoning* berbasis LLM (FactAgent) untuk melakukan verifikasi fakta secara otomatis dan transparan, menunjukkan tren baru menuju deteksi hoaks berbasis multi-agent reasoning yang menyerupai proses berpikir manusia.

Tren terbaru juga menunjukkan pergeseran menuju pendekatan multimodal dan vision-language reasoning. Jin et al. [14] melalui model M-DRUM memperlihatkan bahwa integrasi teks dan citra dengan mekanisme cross-attention dapat menghasilkan akurasi yang melampaui GPT-4 dan LLaVA dalam deteksi berita palsu. Pendekatan ini menunjukkan arah evolusi deteksi hoaks masa depan yang tidak lagi hanya berfokus pada teks, tetapi juga pada pemahaman hubungan antara konten visual dan naratif. Kurniawan et al. [23] melengkapi arah ini dengan pendekatan BERTopic untuk mengidentifikasi topik-topik populer dalam hoaks seperti COVID-19, vaksinasi, dan isu agama, yang menunjukkan bahwa pemetaan semantik topik dapat membantu dalam membangun model deteksi yang lebih kontekstual. Dengan demikian, penelitian terkini memperlihatkan bahwa deteksi hoaks telah berkembang dari sekadar klasifikasi teks menuju sistem cerdas yang mampu memahami, menafsirkan, dan menjelaskan makna di balik penyebaran informasi palsu.

Berdasarkan telaah literatur tersebut, penelitian ini memiliki urgensi penting dalam upaya memetakan evolusi metodologis, mengidentifikasi kesenjangan riset, dan menelaah kebaruan ilmiah dari berbagai pendekatan dalam deteksi hoaks berbasis Bahasa Indonesia. Kajian ini dilakukan untuk menjawab kebutuhan akan pemahaman sistematis mengenai bagaimana perkembangan metode dari *machine learning* menuju *transformer-based architectures* telah mengubah cara kita mendeteksi informasi palsu. Selain itu, penelitian ini penting untuk mengidentifikasi area yang masih membutuhkan perhatian, seperti interpretabilitas model dan integrasi pendekatan multimodal. Melalui pemetaan komprehensif terhadap penelitian yang ada, hasil kajian ini diharapkan dapat memberikan arah baru bagi penelitian lanjutan, baik dalam aspek teoritis maupun praktis, khususnya dalam pengembangan sistem deteksi hoaks yang lebih adaptif, efisien, dan transparan.

Secara spesifik, penelitian ini bertujuan untuk (a) menganalisis perkembangan dan tren metodologis teknik deteksi hoaks dari pendekatan *machine learning* klasik hingga *transformer-based models*, (b) mengidentifikasi kesenjangan konseptual, metodologis, dan teknis yang masih belum terjawab dalam literatur, dan (c) mengevaluasi kontribusi kebaruan dari pendekatan hybrid, augmentasi data, dan model transformer terhadap peningkatan performa serta keandalan sistem deteksi hoaks terkini. Kajian ini menggunakan pendekatan literature review dengan menelaah hasil penelitian dari tahun 2019 hingga 2025 yang berfokus pada pengembangan teknik deteksi hoaks berbasis Bahasa Indonesia. Dengan metode ini, penelitian berupaya menghasilkan sintesis ilmiah yang tidak hanya bersifat deskriptif, tetapi juga analitis dan evaluatif terhadap tren dan inovasi yang ada. Oleh karena itu, diperlukan kajian yang komprehensif untuk memetakan perkembangan metodologis, mengidentifikasi kesenjangan penelitian, serta mengevaluasi kebaruan pendekatan yang ada, khususnya pada integrasi *hybrid-transformer* sebagai solusi potensial dalam meningkatkan akurasi, efisiensi, dan keandalan sistem deteksi hoaks berbahasa Indonesia.

Berdasarkan tujuan tersebut, maka rumusan masalah penelitian ini dirumuskan dalam tiga pertanyaan utama: (1). Bagaimana perkembangan dan tren metodologis dalam teknik deteksi hoaks berbasis pemrosesan bahasa alami dari pendekatan *machine learning* klasik hingga model *deep learning* dan transformer modern? (2). Apa saja kesenjangan konseptual, metodologis, dan teknis yang masih ada dalam penelitian deteksi hoaks, khususnya terkait keterbatasan data, interpretabilitas model, serta integrasi pendekatan multimodal dan *Explainable AI*? (3). Bagaimana kontribusi kebaruan (novelty) dari pendekatan hybrid, augmentasi data, dan model berbasis transformer terhadap peningkatan performa, generalisasi, dan keandalan sistem deteksi hoaks dalam penelitian terkini? Dengan menjawab ketiga rumusan masalah tersebut, penelitian ini diharapkan dapat memberikan kontribusi signifikan terhadap pengembangan keilmuan dalam bidang keamanan informasi dan pemrosesan bahasa alami. Selain itu, hasil penelitian ini akan memperkuat posisi penelitian-penelitian sebelumnya seperti yang dilakukan oleh Yakub et al. [1], Fariha & Lhaksmana [21], dan Mashuri et al. [9], sekaligus memberikan fondasi konseptual baru untuk eksplorasi lanjutan di ranah deteksi hoaks berbasis bahasa, multimodalitas, dan *Explainable Artificial Intelligence*.

2. METODOLOGI PENELITIAN

Penelitian ini menggunakan desain *literature review non-SLR* dengan pendekatan naratif-integratif untuk memetakan, membandingkan, dan menyintesis temuan dari 31 artikel jurnal (nasional, internasional, dan internasional terindeks Scopus) yang seluruhnya telah tercantum pada daftar referensi dokumen sumber; fokusnya adalah membangun pemahaman terpadu tentang perkembangan teknik deteksi hoaks berbasis NLP, kesenjangan

penelitian, serta kebaruan pendekatan hybrid–transformer tanpa melakukan meta-analisis kuantitatif formal. Sampel/subjek kajian adalah artikel yang secara eksplisit mengeksplorasi atau mengevaluasi metode deteksi hoaks dan kedekatan domainnya (termasuk *fake news*, *disinformation*, *phishing*, *hate speech* yang relevan sebagai proksi fitur/arsitektur), dengan rentang waktu publikasi 2019–2025 sebagaimana tercermin pada korpus dan tabel perbandingan kinerja model (mis. SVM/TF-IDF, Naive Bayes, CNN, LSTM/GRU, IndoBERT/RoBERTa, hybrid CNN-LSTM, Bi-LSTM+IndoBERT, augmentasi data, serta pendekatan multimodal dan XAI) yang telah terdokumentasi dalam berkas acuan. Instrumen penelitian berupa lembar ekstraksi data yang dirancang penulis untuk mengumpulkan metadata (tahun, outlet, indeksasi), karakteristik data (sumber, ukuran, ketidakseimbangan), arsitektur/model, teknik prapemrosesan/ekstraksi fitur, skema validasi (k-fold, hold-out), serta metrik evaluasi (akurasi, precision, recall, F1, ROC AUC) sebagaimana dilaporkan pada artikel, termasuk ringkasan kinerja yang telah tersaji dalam tabel ringkas performa pada naskah sumber.

Prosedur pengumpulan data dilakukan melalui dua tahap screening bertingkat: (a) penelusuran internal terhadap keseluruhan daftar referensi yang sudah disediakan penulis (judul–abstrak–hasil) untuk memastikan kesesuaian topik dan kelengkapan metrik; (b) pembacaan penuh untuk mengekstrak detail metodologis dan hasil kuantitatif/qualitatif yang relevan, diikuti pencatatan seragam pada lembar ekstraksi dan tabulasi perbandingan lintas studi; seluruh proses hanya menggunakan sumber artikel yang telah ada pada dokumen, tanpa menambah rujukan baru. Metode analisis menggabungkan sintesis naratif tematik (pengelompokan berdasarkan generasi metode: klasik → *deep learning* → transformer/hybrid), vote-counting deskriptif pada metrik yang dilaporkan, serta pemetaan kesenjangan (interpretabilitas/XAI, imbalance & augmentasi, real-time & efisiensi, multimodal, generalisasi lintas domain) untuk menurunkan jawaban RQ1–RQ3 dan merumuskan agenda riset. Replikasi dimungkinkan karena penulis menyediakan skema kode kategori (metode, fitur, data, metrik), kriteria inklusi/eksklusi eksplisit (relevansi langsung dengan deteksi hoaks/fake news dan kedekatan domain), serta tabel ringkas untuk menyandingkan performa, yang keseluruhannya bersandar pada konten artikel dan rangkuman kinerja model di dokumen rujukan (mis. SVM/TF-IDF unggul pada tahap awal, pergeseran ke CNN/LSTM, dominasi IndoBERT/RoBERTa dan hybrid, hingga tren multimodal dan XAI). Pendekatan ini konsisten dengan materi dan ringkasan performa yang telah terdokumentasi dalam naskah sumber sehingga langkah ekstraksi–tabulasi–sintesis dapat diikuti peneliti lain menggunakan korpus yang sama tanpa menambah referensi eksternal.

3. HASIL DAN PEMBAHASAN

3.1. Hasil penelitian

3.1.1 Karakteristik Umum Penelitian yang Direview

Berdasarkan analisis terhadap 31 artikel yang dikaji, mayoritas penelitian mengenai deteksi hoaks berfokus pada pemrosesan teks dalam Bahasa Indonesia dengan berbagai pendekatan algoritmik. Periode publikasi mencakup tahun 2019 hingga 2025, dengan peningkatan signifikan dalam jumlah publikasi pada rentang 2022–2025 yang menunjukkan meningkatnya perhatian terhadap deteksi hoaks berbasis NLP. Dari sisi asal publikasi, artikel yang dianalisis mencakup jurnal nasional bereputasi, jurnal internasional, serta jurnal internasional terindeks Scopus sebagaimana tercantum dalam daftar referensi. Sebagian besar penelitian menggunakan dataset publik yang bersumber dari platform seperti TurnBackHoax.id, Detik.com, Kaggle Indonesia News Dataset, dan Twitter. Ukuran dataset bervariasi dari 1.000 hingga lebih dari 40.000 entri teks, dengan komposisi data hoaks dan non-hoaks yang tidak seimbang pada sebagian besar studi. Sebagian besar penelitian menggunakan supervised learning dengan format klasifikasi biner (hoaks vs non-hoaks), sedangkan beberapa studi menambahkan kategori tambahan seperti *disinformation* atau *misinformation*. Parameter evaluasi yang umum digunakan meliputi akurasi, precision, recall, F1-score, dan ROC AUC. Secara keseluruhan, tren hasil menunjukkan peningkatan bertahap dalam performa model seiring transisi dari *machine learning* klasik menuju arsitektur *deep learning* dan *transformer-based models*.

3.1.2 Temuan Utama Berdasarkan Pendekatan *Machine learning* Klasik

Pendekatan *machine learning* klasik mendominasi penelitian periode awal 2019–2021. Model Support Vector Machine (SVM) menjadi algoritma yang paling banyak digunakan. Tandiano dan Jollyta [3] menunjukkan bahwa penggunaan SVM dengan unigram TF-IDF menghasilkan akurasi 97,4% dengan precision 1.00 dan recall 0.99. Penelitian Jollyta et al. [4] memperkuat hasil tersebut dengan menunjukkan bahwa linear kernel SVM lebih unggul dibandingkan kernel polinomial atau RBF, dengan akurasi rata-rata 0.974. Rohman et al. [5] menyebutkan bahwa Naive Bayes dan TF-IDF merupakan kombinasi paling sering digunakan dalam klasifikasi berita palsu. Syaifuddin et al. [6] memperkenalkan ensemble classifier dengan kombinasi TF-IDF, n-gram, dan passive-aggressive classifier yang menghasilkan akurasi 91,8%. Hendrawan et al. [7] menambahkan teknik Synthetic Minority Oversampling Technique (SMOTE) untuk mengatasi ketidakseimbangan data dan memperoleh macro-average F1-score sebesar 51,45%. Ferdiansyah dan Wella [24] menunjukkan bahwa SVM dengan 10-fold cross-validation mencapai akurasi tertinggi 89,27%. Secara umum, penelitian-penelitian dengan metode klasik memperlihatkan hasil akurasi berkisar antara 85%–97%, tergantung pada variasi fitur dan keseimbangan data.

Tabel 1. Temuan Utama Berdasarkan Pendekatan *Machine learning* Klasik

No	Model & Fitur	Akurasi (%)	Dataset	Referensi
1	SVM + Unigram TF-IDF	97.4	1,000 berita	[3]
2	Naive Bayes + TF-IDF	91.8	5,000 tweet	[6]
3	SVM dengan Oversampling (SMOTE)	51.45 (F1)	Tidak seimbang	[7]
4	SVM + TF-IDF (10-fold CV)	89.27	1,500 berita	[24]

3.1.3 Temuan Utama Pendekatan *Deep learning* (CNN, LSTM, GRU, dan Hybrid)

Transisi ke pendekatan *deep learning* memperlihatkan peningkatan kinerja yang signifikan dalam klasifikasi hoaks. Yusuf dan Suyanto [8] menggunakan LSTM dengan Word2Vec Skip-gram pada dataset 4.800 berita dan memperoleh akurasi 89,4%. Mashuri et al. [9] menggabungkan pendekatan lexicon-based sentiment analysis dengan LSTM dan mencapai akurasi 99%. Hati dan Sulistiani [10] mengimplementasikan CNN berbasis GloVe *embedding* dengan akurasi pelatihan 99,65% dan akurasi validasi 99,88%. Sibaroni et al. [11] menggabungkan CNN dan LSTM untuk mendeteksi hoaks politik dan mencatat akurasi 96%. Rachmawati dan Darmawan [25] membandingkan performa 1D-CNN, LSTM, dan GRU, menunjukkan bahwa GRU memiliki akurasi tertinggi dengan kecepatan prediksi lebih lambat. Secara umum, rata-rata akurasi model *deep learning* berada pada kisaran 90%–99%, dengan peningkatan signifikan pada studi yang menggunakan *embedding* berbasis konteks dan arsitektur hibrida.

Tabel 2. Temuan Utama Pendekatan *Deep learning* (CNN, LSTM, GRU, dan Hybrid)

No	Model & Fitur	Akurasi (%)	Dataset	Referensi
1	LSTM + Word2Vec	89.4	4,800 berita	[8]
2	Lexicon + LSTM	99	Balanced	[9]
3	CNN + GloVe	99.88	Kaggle Dataset	[10]
4	CNN-LSTM Hybrid	96	Political News	[11]
5	GRU (Political Hoax)	Highest	41,726 data	[25]

3.1.4 Temuan Utama Model Transformer dan Hybrid IndoBERT

Model berbasis transformer memberikan hasil terbaik di antara seluruh pendekatan yang dikaji. Yakub et al. [1] menggunakan fine-tuned IndoBERT pada 29.552 artikel dan mencapai akurasi 98.51%, F1-score 98.44%, dan precision 98.23%. Fathin et al. [12] mengatasi *class imbalance* menggunakan random sampling dengan hasil akurasi 98.2%, recall 97.5%, dan F1-score 97.8%. Jocelynne et al. [13] menerapkan IndoBERT untuk deteksi hoaks politik dengan akurasi 94.1% dan ROC AUC 0.991. Fariha dan Lhaksmana [21] mengombinasikan IndoBERT *embeddings* dengan Bi-LSTM dan teknik back-translation augmentation, menghasilkan akurasi 99.58%. Yefferson et al. [26] menggunakan IndoBERT sebagai feature extractor dan LSTM sebagai lapisan klasifikasi, dengan akurasi 93.2% dan F1-score 90.8%. Rifai et al. [2] menggunakan RoBERTa untuk mendeteksi hoaks COVID-19 dengan akurasi di atas 97% dan waktu pelatihan 50% lebih cepat dari LSTM. Secara keseluruhan, model *transformer-based* mencatat rata-rata akurasi tertinggi 94–99%, mengungguli metode lain baik dari sisi precision maupun recall.

Tabel 3. Temuan Utama Model Transformer dan Hybrid IndoBERT

No	Model & Fitur	Akurasi (%)	Dataset	Referensi
1	IndoBERT Fine-tuned	98.51	29,552 artikel	[1]
2	IndoBERT + Random Sampling	98.2	27,892 artikel	[12]
3	IndoBERT (Political Hoax)	94.1	23,179 artikel	[13]
4	Bi-LSTM + IndoBERT + Augmentation	99.58	Detik + TBH	[21]
5	IndoBERT-LSTM Hybrid	93.2	5,876 data	[26]
6	RoBERTa	97	COVID-19	[2]

3.1.5 Temuan Terkait Feature Engineering, Augmentasi Data, dan Multimodalitas

Teknik *feature engineering* dan *Data Augmentation* menjadi faktor penting yang meningkatkan performa model. Noor et al. [20] menerapkan *Synonym-based Data Augmentation* (EDA) menggunakan CNN dan mencatat

peningkatan akurasi pada dataset kecil. Fariha dan Lhaksana [21] menunjukkan bahwa back-translation meningkatkan akurasi dari 94,14% menjadi 99,58%. Mashuri et al. [9] memadukan lexicon-based sentiment analysis dengan LSTM untuk menangkap nuansa emosional. Purnamadewi dan Zahra [27] menggunakan FastText dengan CNN untuk mendeteksi phishing email dengan akurasi 98.4375%, sedangkan kombinasi FastText dan LSTM mencapai recall tertinggi 98.9583%. Kurniawan et al. [23] menerapkan BERTopic dan IndoBERT untuk mengidentifikasi topik hoaks populer seperti COVID-19 dan vaksinasi. Jin et al. [14] mengembangkan *vision-language model* M-DRUM yang menggabungkan teks dan citra untuk argumentasi berbasis fakta dan mencapai hasil lebih baik dibandingkan GPT-4 dan LLaVA. Mohawesh et al. [15] menggunakan *capsule neural networks* multibahasa dan mencatat peningkatan 5.47% untuk translasi *English-Indonesian*. Secara umum, hasil menunjukkan bahwa integrasi fitur linguistik, augmentasi data, dan pendekatan multimodal berkontribusi besar dalam meningkatkan akurasi model hingga mendekati 100%.

3.1.6 Ringkasan Kuantitatif Hasil dan Distribusi Temuan

Analisis kuantitatif menunjukkan peningkatan rata-rata akurasi antar generasi metode. Pendekatan *machine learning* klasik mencatat akurasi rata-rata 88–97%, sedangkan *deep learning* meningkat ke kisaran 90–99%, dan *transformer-based* models mencapai rata-rata tertinggi 94–99.5%. Pendekatan hybrid menghasilkan performa paling stabil dan unggul pada hampir semua metrik. Dalam konteks dataset, ukuran data lebih besar cenderung menghasilkan hasil yang lebih konsisten; penelitian dengan dataset >20.000 entri menunjukkan akurasi stabil di atas 95%. Dari 31 artikel, 15 menggunakan pre-trained *embeddings* (Word2Vec, GloVe, FastText), 11 menggunakan *transformer-based embeddings* (IndoBERT, RoBERTa), dan 5 menggabungkan hybrid *embeddings*.

Tabel 4. Diagram tren akurasi rata-rata tiap generasi model

Generasi Metode	Rentang Tahun	Model Dominan	Rata-rata Akurasi (%)	Cakupan Referensi
<i>Machine learning</i>	2019–2021	SVM, Naive Bayes	88–97	[3]; [5]; [6]
<i>Deep learning</i>	2021–2024	CNN, LSTM, GRU	90–99	[8]; [10]; [11]
<i>Transformer & Hybrid</i>	2023–2025	IndoBERT, RoBERTa, Bi-LSTM+IndoBERT	94–99.5	[1]; [21]; [2]

Garis tren memperlihatkan kenaikan berkelanjutan dari model klasik menuju model transformer, dengan lonjakan paling tajam terjadi pada tahun 2023–2025, bertepatan dengan adopsi IndoBERT dan varian hybrid seperti Bi-LSTM+IndoBERT dan RoBERTa-based model.

Machine Learning (88–97%) → Deep Learning (90–99%) → Transformer-Hybrid (94–99.5%)

Gambar 1. Tren peningkatan akurasi model per generasi metode (2019–2025)

3.2 Pembahasan

3.2.1 Perkembangan dan Tren Metodologis dalam Teknik Deteksi Hoaks

Perkembangan metodologis dalam penelitian deteksi hoaks memperlihatkan evolusi bertahap dari pendekatan berbasis *machine learning* klasik menuju *deep learning* dan akhirnya ke model *transformer-based* yang menjadi state of the art. Hasil sintesis dari seluruh artikel menunjukkan bahwa fase awal penelitian (2019–2021) masih sangat dipengaruhi oleh metode statistik dan algoritma klasifikasi linear. Penelitian yang dilakukan oleh Tandiano dan Jollyta [3] serta Jollyta et al. [4] memperlihatkan dominasi Support Vector Machine (SVM) dengan unigram TF-IDF, menghasilkan akurasi sangat tinggi (97,4% dan 0,974) yang menegaskan bahwa representasi sederhana berbasis kata masih efektif pada data berita formal. Rohman et al. [5] menemukan bahwa Naive Bayes dan TF-IDF merupakan dua pendekatan paling populer pada fase ini, menandakan bahwa komunitas ilmiah masih mengandalkan pembobotan kata sederhana dan pendekatan probabilistik untuk mengklasifikasi berita palsu.

Syaifuddiin et al. [6] dan Hendrawan et al. [7] kemudian memperkenalkan pendekatan ensemble dan oversampling untuk meningkatkan kestabilan model terhadap data yang tidak seimbang, menandai fase peralihan menuju integrasi beberapa metode dalam satu sistem. Hasil-hasil ini menunjukkan bahwa fokus riset pada tahap awal masih berkisar pada optimalisasi fitur dan keseimbangan data, bukan pada pemahaman konteks semantik.

Perkembangan berikutnya terjadi saat penelitian mulai mengadopsi *deep learning* dengan kemampuan mengenali urutan kata dan hubungan semantik antarfrasa. Yusuf dan Suyanto [8] merupakan salah satu studi awal yang menerapkan LSTM dengan Word2Vec Skip-gram dan memperoleh akurasi 89,4%, melampaui pendekatan berbasis bag of words. Mashuri et al. [9] menambahkan dimensi baru dengan menggabungkan lexicon-based sentiment analysis ke dalam model LSTM, yang berhasil meningkatkan akurasi menjadi 99%. Pendekatan CNN yang diimplementasikan oleh Hati dan Sulistiani [10] mencapai validasi akurasi 99,88%, menunjukkan kemampuan model ini dalam mengenali pola lokal pada teks. Sibaroni et al. [11] kemudian memadukan CNN dan LSTM dalam model hibrida untuk mendeteksi hoaks politik dan mencatat akurasi 96%. Perkembangan ini menandakan transisi penting dari metode berbasis statistik menuju pendekatan yang lebih kontekstual dan semantik.

Fase terakhir perkembangan metodologis ditandai oleh adopsi model *transformer-based* seperti IndoBERT dan RoBERTa. Yakub et al. [1] memperlihatkan bahwa fine-tuned IndoBERT mampu mencapai akurasi 98,51% dengan F1-score 98,44%, sementara Fathin et al. [12] menunjukkan hasil serupa (98,2%) setelah menangani class imbalance. Jocelyne et al. [13] menerapkan IndoBERT pada deteksi hoaks politik dan mencatat akurasi 94,1% dengan ROC AUC 0,991. Fariha dan Lhaksana [21] mencapai hasil tertinggi 99,58% menggunakan Bi-LSTM dengan IndoBERT *embeddings* serta back-translation untuk *Data Augmentation*. Model RoBERTa yang diimplementasikan Rifai et al. [2] menunjukkan efisiensi komputasi tinggi dengan akurasi lebih dari 97%. Secara longitudinal, hasil-hasil ini memperlihatkan bahwa *transformer-based architectures* telah melampaui semua pendekatan sebelumnya, baik dari sisi akurasi maupun efisiensi pemrosesan. Dengan demikian, perkembangan metodologis dalam deteksi hoaks mengikuti pola tiga fase: feature engineering and statistical classification, deep contextual learning, dan *transformer-based* contextual representation.

Signifikansi temuan RQ1 terletak pada pembuktian empiris bahwa kemampuan memahami konteks linguistik dan semantik merupakan kunci efektivitas deteksi hoaks modern. Pergeseran dari representasi berbasis kata menuju representasi berbasis konteks memungkinkan model menangkap nuansa emosional, sarkasme, dan struktur naratif khas berita palsu yang sebelumnya sulit dikenali. Penelitian Anda memberikan kontribusi akademik penting dengan memetakan evolusi ini secara komprehensif, memperjelas hubungan antara generasi teknologi NLP dan peningkatan performa klasifikasi hoaks, sekaligus mendokumentasikan kronologi inovasi dari SVM hingga IndoBERT yang belum pernah dikompilasi secara terpadu sebelumnya.

3.2.2 Kesenjangan Konseptual, Metodologis, dan Teknis dalam Penelitian Deteksi Hoaks

Hasil sintesis dari 31 artikel mengungkapkan bahwa meskipun kemajuan teknologi deteksi hoaks telah signifikan, masih terdapat sejumlah kesenjangan riset yang substansial. Pertama, dari aspek konseptual, sebagian besar penelitian berfokus pada pencapaian akurasi tinggi tanpa mempertimbangkan explainability dan ethical interpretability dari model yang dihasilkan. Model seperti IndoBERT dan RoBERTa [1]; [2] beroperasi sebagai sistem tertutup yang sulit dijelaskan kepada pengguna non-teknis. Hanya Syari et al. [16] yang mencoba mengintegrasikan pendekatan *critical thinking* dan *Explainable AI (XAI)* untuk menambahkan dimensi interpretabilitas. Ketidadaan pendekatan interpretatif ini menunjukkan bahwa penelitian deteksi hoaks masih didominasi paradigma performa—bukan paradigma transparansi—yang berpotensi menimbulkan bias pada hasil klasifikasi.

Kedua, dari aspek metodologis, keterbatasan yang paling nyata adalah masalah data imbalance dan keberagaman linguistik. Sebagian besar dataset yang digunakan dalam penelitian tidak seimbang antara kategori hoaks dan non-hoaks. Fathin et al. [12] dan Hendrawan et al. [7] telah mengusulkan solusi melalui random sampling dan SMOTE, tetapi belum ada metodologi standar yang mampu menjamin keseimbangan tanpa mengorbankan keaslian distribusi data. Variasi linguistik dalam Bahasa Indonesia seperti code-mixing, gaya informal, atau dialek daerah juga belum sepenuhnya ditangani, karena sebagian besar model dilatih pada data berita formal. Hal ini menimbulkan kesenjangan antara performa eksperimental dan performa di dunia nyata, terutama dalam konteks media sosial yang lebih kompleks.

Ketiga, dari aspek teknis, keterbatasan paling krusial terdapat pada efisiensi komputasi dan skalabilitas model. Rifai et al. [2] menunjukkan bahwa RoBERTa mampu memangkas waktu pelatihan hingga 50%, namun penelitian lain jarang membahas optimasi kecepatan dan konsumsi memori. Sebagian besar studi seperti Yakub et al. [1] dan Fariha & Lhaksana [21] berfokus pada hasil akurasi tanpa mengevaluasi efisiensi proses inferensi. Selain itu, sebagian besar model belum diuji dalam lingkungan real-time. Wibowo et al. [19] dan Hati & Sulistiani [10] memang mengembangkan prototipe berbasis web, namun implementasinya masih bersifat demonstratif tanpa

pengujian skala besar. Ini menunjukkan bahwa arah penelitian masih berorientasi pada pembuktian konsep, bukan pada penerapan praktis dalam sistem yang bisa digunakan secara langsung oleh masyarakat atau lembaga verifikasi fakta.

Kesenjangan lain yang menonjol adalah rendahnya integrasi pendekatan multimodal dalam sistem deteksi hoaks. Padahal, sejumlah penelitian telah menunjukkan potensi besar dari integrasi antara informasi visual dan teks. Sebagai contoh, penelitian Mulyawan et. al. [31] berbasis vision-language menggunakan arsitektur CNN sebagai encoder dan Transformer sebagai decoder menunjukkan bahwa pendekatan multimodal mampu menghasilkan representasi semantik yang lebih kaya dalam memahami konten visual dan tekstual secara bersamaan. Dalam konteks tersebut, model berbasis Transformer tidak hanya unggul dalam pemrosesan teks, tetapi juga efektif dalam mengintegrasikan informasi lintas-modal, sehingga membuka peluang pengembangan sistem deteksi hoaks yang lebih kontekstual dan adaptif terhadap konten digital modern yang bersifat multimodal. Demikian juga, Jin et al. [14] dan Mohawesh et al. [15] telah menunjukkan potensi besar integrasi teks dan citra melalui vision-language models (M-DRUM) serta capsule neural networks multibahasa yang memberikan peningkatan performa signifikan. Namun, kebanyakan penelitian yang direview masih berfokus pada data teks semata. Hal ini menandakan bahwa penelitian di Indonesia, meskipun telah mencapai state of the art dalam teks, belum mengeksplorasi secara optimal potensi multimodalitas yang merepresentasikan penyebaran hoaks modern yang biasanya berupa kombinasi teks, gambar, dan video.

Selain itu, terdapat keterbatasan dari sisi reproducibility dan data openness. Banyak penelitian melaporkan hasil akurasi tinggi tanpa menyediakan dataset atau kode sumber untuk verifikasi publik. Hanya beberapa studi seperti Fariha & Lhaksana [21] dan Adiati et al. [28] yang menyajikan perbandingan eksplisit antar model dengan deskripsi metodologis yang cukup jelas untuk direplikasi. Laimeheriwa et al. [17] melalui surveinya menegaskan perlunya benchmarking sistematis agar hasil penelitian dapat dibandingkan secara objektif. Oleh karena itu, keterbatasan dalam replikasi dan standar evaluasi masih menjadi hambatan besar bagi konsistensi ilmiah dalam bidang ini.

Signifikansi temuan RQ2 adalah bahwa penelitian ini tidak hanya menyoroti pencapaian teknis, tetapi juga menegaskan bahwa keberlanjutan riset deteksi hoaks membutuhkan paradigma baru yang menyeimbangkan antara performa, efisiensi, dan transparansi. Penelitian ini berkontribusi pada ranah akademik dengan memperjelas bahwa kemajuan NLP tidak akan bermakna tanpa kemampuan interpretasi, efisiensi pemrosesan, serta adaptabilitas model terhadap variasi sosial dan linguistik. Dalam konteks bidang keilmuan Anda, temuan ini penting karena memperluas perspektif dari sekadar pengembangan algoritma menjadi pemahaman menyeluruh mengenai batas-batas penerapan teknologi AI dalam menjaga keautentikan konten digital dan keamanan informasi publik.

3.2.3 Kebaruan (Novelty), Kontribusi, dan Arah Pengembangan Pendekatan Hybrid-Transformer

Kebaruan yang teridentifikasi dalam penelitian-penelitian yang dikaji mencakup tiga ranah utama, yaitu inovasi metodologis, linguistik, dan teknis. Dari sisi metodologis, hybridization menjadi tren yang dominan dalam penelitian tahun 2023–2025. Fariha dan Lhaksana [21] memperkenalkan model Bi-LSTM dengan IndoBERT *embeddings* dan augmentasi data melalui back-translation yang mencapai akurasi 99,58%, jauh lebih tinggi dibandingkan SVM (93,87%) dan Decision Tree (95,86%). Yefferson et al. [26] juga menggunakan IndoBERT-LSTM hybrid dengan akurasi 93,2% dan F1-score 90,8%. Sibaroni et al. [11] membuktikan efektivitas CNN-LSTM dengan akurasi 96%. Ketiganya menandai perubahan signifikan dari paradigma tunggal menjadi pendekatan integratif yang menggabungkan kemampuan feature extraction dari transformer dengan keunggulan temporal modeling dari LSTM atau CNN.

Dari sisi linguistik, kebaruan muncul melalui upaya menangkap nuansa semantik dan emosional yang kompleks dalam teks hoaks. Mashuri et al. [9] menunjukkan bahwa kombinasi lexicon-based sentiment analysis dan LSTM mampu meningkatkan akurasi hingga 99%. Rozi et al. [29] menambahkan dimensi sentiment scoring menggunakan VADER untuk memperkuat deteksi berbasis emosi. Kurniawan et al. [23] menerapkan BERTopic untuk mengelompokkan topik hoaks populer seperti vaksinasi dan agama, menunjukkan potensi topik modeling dalam memperkaya konteks klasifikasi. Pendekatan ini menunjukkan bahwa inovasi tidak hanya muncul dari teknologi baru, tetapi juga dari pemahaman yang lebih dalam terhadap aspek linguistik dan psikologis dari penyebaran hoaks.

Dari sisi teknis, kebaruan terlihat pada penggunaan *Data Augmentation* dan transfer learning. Noor et al. [20] berhasil meningkatkan performa CNN melalui *Synonym-based Data Augmentation*, sementara Fariha dan Lhaksana [21] menunjukkan peningkatan signifikan dengan back-translation. Yakub et al. [1] dan Fathin et al. [12] memperlihatkan bagaimana transfer learning dengan IndoBERT yang di-fine-tune secara mendalam mampu menghasilkan performa di atas 98%. Rifai et al. [2] memperkuat bukti bahwa model transfer learning seperti RoBERTa dapat lebih efisien secara komputasi. Li et al. [22] melalui FactAgent memperkenalkan konsep agentic

reasoning yang memungkinkan pembagian tugas deteksi menjadi subproses interpretatif, menandai arah baru menuju sistem deteksi berbasis multi-agent intelligence. Kebaruan ini memperlihatkan transformasi arah penelitian dari sekadar klasifikasi ke arah sistem cerdas yang mampu memahami dan menjelaskan keputusan algoritmiknya.

Penelitian ini berkontribusi signifikan dengan memetakan seluruh kebaruan tersebut dalam satu kerangka yang terintegrasi, memperlihatkan bahwa kemajuan metodologis dalam deteksi hoaks berbahasa Indonesia sejajar dengan tren global NLP mutakhir. Dengan mengidentifikasi pola hybridization dan transfer learning, penelitian ini memperkuat posisi bidang keamanan informasi dan autentikasi konten sebagai subdisiplin penting dalam ilmu teknologi informasi. Signifikansi ilmiahnya terletak pada kontribusi konseptual untuk membangun kerangka meta-evolution riset deteksi hoaks, yang dapat dijadikan dasar bagi pengembangan model adaptif, efisien, dan transparan di masa depan.

3.2.4 Implikasi dan Batasan Penelitian

Implikasi penelitian ini terbagi dalam tiga level: teoretis, metodologis, dan praktis. Pada level teoretis, penelitian ini memperkaya pemahaman akademik mengenai evolusi pendekatan NLP dalam konteks deteksi hoaks, menegaskan hubungan antara generasi model dan kemampuan semantik yang dihasilkan. Dengan mendokumentasikan perjalanan dari SVM hingga transformer hybrid, penelitian ini menyajikan peta perkembangan yang dapat dijadikan referensi untuk studi lanjutan di bidang disinformation security dan content authenticity. Pada level metodologis, hasil penelitian memberikan dasar bagi pengembangan kerangka penelitian berbasis integrasi model dan augmentasi data. Pendekatan seperti Bi-LSTM + IndoBERT atau CNN-LSTM hybrid terbukti menjadi arah pengembangan utama untuk mencapai keseimbangan antara performa tinggi dan generalisasi model.

Dari sisi praktis, penelitian ini memberikan kontribusi terhadap pembangunan sistem deteksi hoaks yang lebih efisien, terutama untuk konteks bahasa dengan sumber data terbatas seperti Bahasa Indonesia. Hasil yang diperoleh dari model transformer menunjukkan bahwa pendekatan transfer learning dapat menghemat sumber daya komputasi sekaligus meningkatkan akurasi. Selain itu, penelitian ini membuka peluang implementasi *Explainable AI* untuk mendukung transparansi sistem verifikasi berita di platform digital.

Namun, penelitian ini juga memiliki beberapa batasan. Pertama, karena pendekatan ini berbasis non-systematic literature review, proses penilaian kualitas studi bergantung pada dokumentasi yang tersedia dalam artikel, tanpa analisis kuantitatif meta-statistik. Kedua, keterbatasan sumber dataset yang bervariasi menyebabkan perbandingan performa antar penelitian tidak selalu setara. Ketiga, sebagian besar hasil yang dikaji berfokus pada teks, sehingga belum sepenuhnya merepresentasikan tren multimodalitas dan konteks sosial penyebaran hoaks. Terakhir, keterbatasan akses terhadap kode dan data penelitian terdahulu menghambat proses replikasi menyeluruh.

Meskipun demikian, penelitian ini tetap memiliki kontribusi ilmiah yang kuat karena berhasil mengompilasi, mengorganisasi, dan menyintesis pengetahuan terdispersi menjadi satu struktur evolutif yang jelas. Dengan menyoroti aspek kebaruan dan kesenjangan secara simultan, penelitian ini menjadi referensi konseptual penting bagi akademisi dan praktisi yang ingin mengembangkan sistem deteksi hoaks yang lebih adaptif, interpretatif, dan berbasis bahasa lokal.

4. KESIMPULAN

Penelitian ini menyimpulkan bahwa perkembangan teknik deteksi hoaks berbasis pemrosesan bahasa alami telah mengalami transformasi signifikan dalam kurun waktu 2019 hingga 2025. Berdasarkan telaah terhadap 31 artikel nasional, internasional, dan internasional terindeks Scopus, ditemukan bahwa pendekatan *machine learning* klasik seperti SVM dan Naive Bayes mendominasi penelitian awal dengan keunggulan pada efisiensi dan kesederhanaan dalam klasifikasi berbasis fitur TF-IDF. Namun, perkembangan teknologi selanjutnya memperlihatkan transisi menuju pendekatan *deep learning* seperti CNN, LSTM, dan GRU yang mampu memahami konteks semantik dan pola linguistik secara lebih mendalam. Tahap paling mutakhir ditandai dengan dominasi model *transformer-based* seperti IndoBERT dan RoBERTa yang menunjukkan performa tertinggi dengan akurasi di atas 98%, serta efisiensi komputasi yang lebih baik dibandingkan model sebelumnya. Perjalanan ini menandakan perubahan paradigma besar dari pendekatan berbasis fitur menuju representasi berbasis konteks yang lebih adaptif terhadap kompleksitas bahasa alami.

Temuan penelitian ini memberikan kontribusi penting terhadap bidang keilmuan teknologi informasi, khususnya pada ranah natural language processing dan keamanan informasi digital. Penelitian ini berhasil memetakan peta evolusi metodologis dari teknik deteksi hoaks berbahasa Indonesia secara komprehensif, menunjukkan bahwa peningkatan akurasi dan generalisasi model sangat bergantung pada integrasi antara

contextual *embedding*, arsitektur hibrida, serta strategi augmentasi data. Selain itu, penelitian ini memperjelas bahwa inovasi tidak hanya terjadi pada tingkat algoritma, tetapi juga pada pemahaman semantik, emosional, dan naratif dari teks hoaks yang menjadi tantangan utama dalam era informasi digital. Dengan demikian, penelitian ini memperkuat dasar konseptual bagi pengembangan sistem deteksi hoaks yang lebih efisien, adaptif, dan dapat dipercaya. Kontribusi lain yang menonjol adalah penyajian meta-mapping yang menghubungkan inovasi model dengan performa empirisnya, sehingga dapat menjadi acuan bagi pengembangan riset lanjutan di bidang keamanan siber dan verifikasi konten digital.

Meskipun telah menyajikan gambaran evolusi yang komprehensif, penelitian ini menyadari adanya sejumlah keterbatasan yang dapat dijadikan dasar bagi penelitian di masa depan. Pertama, masih diperlukan eksplorasi mendalam terhadap integrasi pendekatan multimodal, mengingat banyak hoaks kini disebarkan melalui kombinasi teks, gambar, dan video. Kedua, penelitian lanjutan perlu mengembangkan model yang lebih *Explainable* dan transparan agar hasil deteksi dapat dipahami oleh pengguna non-teknis, terutama dalam konteks sistem verifikasi publik. Ketiga, pengembangan dataset hoaks yang lebih besar, seimbang, dan representatif terhadap variasi bahasa informal serta dialek lokal akan sangat berperan dalam meningkatkan generalisasi model. Keempat, penelitian masa depan dapat berfokus pada efisiensi model dan penerapan real-time detection systems yang siap digunakan di platform media sosial atau portal berita daring. Terakhir, diperlukan upaya kolaboratif antarpeleliti lintas disiplin—bahasa, psikologi, dan keamanan siber—untuk membangun model deteksi hoaks yang tidak hanya unggul secara teknis tetapi juga sensitif terhadap konteks sosial dan budaya masyarakat digital. Dengan arah pengembangan tersebut, penelitian ini diharapkan menjadi fondasi ilmiah dan konseptual yang mendorong terciptanya ekosistem digital yang lebih aman, informatif, dan bebas dari misinformasi.

REFERENCES

- [1] Yakub, Sy. Y., Wowor, R. J., Agustriawan, D., Irmawati, Kristiyanti, D. A., Winarno, P. M., Hassan, R., Mohamad, R., & Razzaq, R. Z. (2025). *Deep learning* for cyber security: Deephoax detection. International Conference on Artificial Intelligence and Soft Computing.
- [2] Rifai, A. P., Mulyani, Y. P., Febrianto, R., Arini, H. M., Wijayanto, T., Lathifah, N., Liu, X., Li, J., Yin, H., Wu, Y., & Mohawesh, R. (2023). DETECTION MODEL FOR FAKE NEWS ON COVID-19 IN INDONESIA. None.
- [3] Tandiano, A. H., & Jollyta, D. (2024). CLASSIFICATION OF FAKE NEWS IN INDONESIAN LANGUAGE USING SUPPORT VECTOR MACHINE METHOD. None.
- [4] Jollyta, D., Gusrianty, G., Prihandoko, P., & Sukrianto, D. (2023). N-gram and kernel performance using Support Vector Machine algorithm for fake news detection system. Ilkom Jurnal Ilmiah.
- [5] Rohman, M. A., Khairani, D., Hulliyah, K., Arini, Riswandi, P., & Lakoni, I. (2021). Systematic literature review on methods used in classification and fake news detection in indonesian. None.
- [6] Syaifuddiin, G. N., Arifin, R., Desriyanti, D., Buntoro, G. A., Rosyidin, Z. U., Pratama, R. Y., & Selamat, A. (2023). Hoax identification of indonesian tweeters using ensemble classifier. The Academic Center for Education, Culture and Research.
- [7] Hendrawan, R., Hana, K. M., & Kurniati, A. (2024). Enhancing indonesian abusive language detection on imbalanced news comment dataset using Support Vector Machine with oversampling. None.
- [8] Yusuf, R., & Suyanto, S. (2022). Hoax detection on indonesian text using long short-term memory. None.
- [9] Mashuri, C., Prastyo, E. H. A., & Hariri, F. R. (2025). Improving fake news detection accuracy with lexicon-based approach and LSTM through text preprocessing. JURNAL SISTEM INFORMASI BISNIS.
- [10] Hati, C. R. S., & Sulistiani, H. (2025). Implementation of CNN algorithm for indonesian hoax news detection on online news portals. JOURNAL OF APPLIED INFORMATICS AND COMPUTING.
- [11] Sibaroni, Y., Mahadzir, S., Prasetyowati, S., & Ihsan, A. F. (2024). Combating misinformation: Leveraging *deep learning* for hoax detection in indonesian political social media. Jurnal Infotel.
- [12] Fathin, M., Sibaroni, Y., & Prasetyowati, S. S. (2024). Handling imbalance dataset on hoax indonesian political news classification using IndoBERT and random sampling. None.
- [13] Jocelynn, C., Wijayakusuma, I. L., & Harini, L. P. I. (2025). Detection of political hoax news using fine-tuning IndoBERT. JOURNAL OF APPLIED INFORMATICS AND COMPUTING.
- [14] Jin, R., Fu, R., Wen, Z., Zhang, S., Liu, Y., & Tao, J. (2024). Fake news detection and manipulation reasoning via large vision-language models. arXiv.org.
- [15] Mohawesh, R., Maqsood, S., & Althebyan, Q. (2023). Multilingual *deep learning* framework for fake news detection using capsule neural network. Journal of Intelligence and Information Systems.

- [16] Syari, M. A., Sembiring, H., & Siregar, M. F. (2025). Integration of *critical thinking* in intelligent algorithms for hoax detection on social media platforms. None.
- [17] Laimeheriwa, J., Iswanto, I. A., Oktriani, W., & Hidayat, M. F. (2024). A survey on indonesian hoax analyzer and fake news detection using *deep learning* techniques. None.
- [18] Kozik, R., Pawlicka, A., Pawlicki, M., Chora, M., Mazurczyk, W., & Cabaj, K. (2024). A meta-analysis of state-of-the-art automated fake news detection methods. *IEEE Transactions on Computational Social Systems*.
- [19] Wibowo, F. W., Dahlan, A., & Wihayati. (2021). Detection of fake news and hoaxes on information from web scraping using classifier methods. None.
- [20] Noor, A. Z. M., Gernowo, R., & Nurhayati, O. D. (2023). *Data Augmentation* for hoax detection through the method of convolutional neural network in indonesian news. None.
- [21] Fariha, T. M. A., & Lhaksana, K. (2025). Indonesian fake news classification using bi-LSTM with IndoBERT and *Data Augmentation*. *International Conference Advancement Data Science, E-Learning and Information Systems*.
- [22] Li, X., Zhang, Y., & Malthouse, E. (2024). Large language model agent for fake news detection. *arXiv.org*.
- [23] Kurniawan, E. M., Aherman, V. N., Qomariyah, N. N., Manuaba, I., & Anom, A. (2023). Analysing hoax dataset in indonesian language with topic modeling. *International Conferences on Information Science and System*.
- [24] Ferdiansyah, I. F., & Wella. (2022). Comparative study of hoax detection using Support Vector Machine algorithm, naive bayes, random forest and k-nearest neighbor. *Information and Communication Technology Convergence*.
- [25] Rachmawati, O. C. R., & Darmawan, Z. M. E. (2024). The comparison of *deep learning* models for indonesian political hoax news detection. *Journal*.
- [26] Yefferson, D. Y., Lawijaya, V., & Girsang, A. S. (2024). Hybrid model: IndoBERT and long short-term memory for detecting indonesian hoax news. *IAES International Journal of Artificial Intelligence (IJ-AI)*.
- [27] Purnamadewi, Y. R., & Zahra, A. (2025). Enhancing detection of zero-day phishing email attacks in the indonesian language using *deep learning* algorithms. *Bulletin of Electrical Engineering and Informatics*.
- [28] Adiati, N. P. R., Priambodo, D. F., Girinoto, G., Indarjani, S., Rizal, A., Prayoga, A., & Beatrix, Y. (2024). Comparative study of predictive models for hoax and disinformation detection in indonesian news. *International Journal of Advances in Intelligent Informatics*.
- [29] Rozi, I. F., Arianto, R., & Mahdyan, H. H. (2023). Fake news detection using sentiment analysis approach in indonesian language. None.
- [30] Mulyawan, R., Naparin, H., & Fatihia, W. M., "Comparison of Text Vectorization Methods for IMDB Movie Review Sentiment Analysis Using SVM," *Journal of Applied Informatics and Computing*, 2025.
- [31] Mulyawan, R., Sunyoto, A., & Muhammad, A. H., "Automatic Indonesian Image Captioning using CNN and *Transformer-based* Model Approach," 2022 5th International Conference on Information and Communications Technology (ICOIACT), Yogyakarta, Indonesia, 2022.