

DETEKSI BERITA HOAKS BAHASA INDONESIA MENGGUNAKAN MODEL HYBRID INDOBERT-BILSTM DENGAN TEKNIK HYPERPARAMETER TUNING

I Wayan Bagus Satria Putra Pratama, Made Adi Paramartha Putra, A.A. Istri Ita Paramitha

Teknik Informatika, Universitas Primakara, Bali
Jln. Tukad Badung No 135 Denpasar, Bali, Indonesia
2201020056@primakara.ac.id

ABSTRAK

Penyebaran berita hoaks, khususnya di Indonesia, menjadi tantangan serius bagi integritas informasi di Indonesia. Penelitian ini bertujuan mengembangkan model deteksi hoaks berbasis arsitektur hybrid IndoBERT dan BiLSTM untuk meningkatkan akurasi klasifikasi teks. Metode penelitian menggunakan kerangka CRISP-DM yang mencakup tahap persiapan data, pembangunan model, dan evaluasi. Model memanfaatkan kemampuan representasi kontekstual IndoBERT serta pemodelan dependensi sekuensial BiLSTM, dengan optimasi hyperparameter menggunakan Optuna. Hasil pengujian menunjukkan model mencapai F1-Score sebesar 98,59% dan performa lebih baik dibandingkan konfigurasi baseline. Selain itu, model diimplementasikan dalam prototipe berbasis Gradio untuk pengujian real-time. Temuan ini menunjukkan bahwa integrasi arsitektur Transformer dan RNN yang dioptimasi dapat meningkatkan kinerja deteksi hoaks secara efektif.

Kata kunci : Deteksi Hoaks, IndoBERT, BiLSTM, Hyperparameter Tuning, CRISP-DM, Optuna.

1. PENDAHULUAN

Transformasi digital di Indonesia telah mengubah cara masyarakat mengakses informasi secara instan melalui platform digital [1]. Data menunjukkan urgensi masalah yang signifikan dengan keberadaan sekitar 800.000 situs penyebar disinformasi yang aktif, di mana konten tersebut menyebar secara cepat melalui media sosial seperti WhatsApp dan Facebook. Fenomena ini tidak hanya mengancam stabilitas sosial tetapi juga melemahkan kepercayaan publik, sehingga metode verifikasi manual yang lambat dan tidak skalabel dinilai tidak lagi mencukupi untuk membendung volume informasi di era digital [2]. Hal ini mendasari kebutuhan mendesak akan sistem deteksi otomatis yang mampu bekerja secara presisi dan konsisten.

Upaya otomatisasi awal yang mengandalkan algoritma *machine learning* konvensional seperti SVM, KNN, dan Naive Bayes sering kali menemui jalan buntu saat berhadapan dengan kompleksitas linguistik Bahasa Indonesia [3]. Model-model tersebut cenderung kesulitan menangkap aspek kontekstual, sarkasme, dan muatan emosi yang kerap disisipkan dalam narasi hoaks [4]. Sebagai solusinya, pemanfaatan model berbasis Transformer seperti IndoBERT mulai mendominasi karena kemampuannya dalam merepresentasikan makna kata secara mendalam berdasarkan konteksnya [5]. Kendati demikian, IndoBERT memiliki keterbatasan dalam memodelkan hubungan ketergantungan jarak jauh pada teks berita yang panjang, sehingga diperlukan arsitektur tambahan untuk menutupi celah tersebut [6].

Integrasi arsitektur hybrid antara IndoBERT dan *Bidirectional Long Short-Term Memory* (BiLSTM) muncul sebagai pendekatan yang sangat prospektif dalam meningkatkan performa deteksi [7]. Dalam skema ini, IndoBERT berfungsi sebagai pengekstraksi fitur kontekstual yang kaya, sementara BiLSTM

berperan menangkap urutan narasi dari dua arah untuk memahami pola manipulatif yang tersirat dalam alur cerita. Penelitian terdahulu membuktikan bahwa kombinasi ini mampu meningkatkan akurasi hingga melampaui 95% dibandingkan penggunaan model tunggal [8]. Untuk mencapai performa maksimal, penelitian ini juga menerapkan teknik optimasi melalui *hyperparameter tuning* menggunakan Optuna, yang memungkinkan pencarian parameter kritis secara otomatis dan efisien guna menjamin generalisasi model yang lebih baik serta meminimalkan risiko overfitting [9].

Melalui latar belakang tersebut, penelitian ini difokuskan pada perancangan dan pengembangan model deteksi hoaks hybrid IndoBERT-BiLSTM yang dioptimalkan secara sistematis. Secara teoretis, studi ini diharapkan memperkaya literatur mengenai implementasi NLP pada bahasa lokal, sedangkan secara praktis diharapkan mampu menghasilkan prototipe sistem yang aplikatif untuk membantu masyarakat memitigasi penyebaran informasi palsu. Ruang lingkup penelitian ini dibatasi pada analisis dataset berita Bahasa Indonesia dari sumber terpercaya periode 2020-2025 dengan evaluasi menggunakan metrik performa standar (akurasi, presisi, *recall*, dan *f1-score*), yang diimplementasikan melalui antarmuka prototipe berbasis web menggunakan Gradio.

2. TINJAUAN PUSTAKA

2.1. Artificial Intelligence (AI)

Artificial Intelligence merupakan bidang ilmu komputer yang berfokus pada pengembangan sistem yang mampu melakukan tugas kognitif seperti pengambilan keputusan, pembelajaran pola, dan pemrosesan bahasa secara otomatis [10]. Dalam penelitian ini digunakan konsep *Narrow AI*, yaitu sistem yang dirancang untuk menyelesaikan satu tugas spesifik secara optimal melalui pemanfaatan *machine*

learning dan *deep learning*. Pendekatan ini relevan untuk klasifikasi teks karena memungkinkan model mempelajari pola linguistik secara adaptif dari data pelatihan tanpa perlu aturan eksplisit [11].

2.2. Natural Language Processing (NLP)

Natural Language Processing adalah cabang AI yang bertujuan memungkinkan komputer memahami, menafsirkan, dan menghasilkan bahasa manusia secara bermakna [12]. Pendekatan NLP tradisional berbasis statistik memiliki keterbatasan dalam memahami konteks panjang dan relasi antar kata, terutama pada teks kompleks seperti berita. Untuk mengatasi keterbatasan tersebut, arsitektur Transformer diperkenalkan dengan mekanisme *self-attention* yang mampu memodelkan dependensi global secara paralel sehingga representasi teks menjadi lebih kontekstual [13]. Salah satu implementasi utama arsitektur ini adalah BERT yang menggunakan *pre-training bidirectional* untuk memahami hubungan kata berdasarkan konteks sebelum dan sesudahnya dalam kalimat [14].

2.3. IndoBERT

IndoBERT merupakan adaptasi arsitektur BERT yang dioptimalkan khusus untuk bahasa Indonesia melalui proses *pre-training* pada korpus besar teks lokal, sehingga mampu menghasilkan representasi semantik yang lebih relevan dibanding model multilingual umum [5]. Penelitian yang dilakukan oleh Juarto dan Yulianto [15] bawasannya model IndoBERT mampu menghasilkan skor akurasi terbaik sebesar 94,5% membuktikan model IndoBERT mampu mengidentifikasi berita dengan akurasi yang tinggi. Dalam penelitian ini, IndoBERT berfungsi sebagai pengekstraksi fitur utama sebelum diproses oleh lapisan klasifikasi lanjutan.

2.4. Bidirectional Long Short-Term Memory (BiLSTM)

BiLSTM merupakan pengembangan dari *Long Short-Term Memory* (LSTM) yang mampu mempelajari pola data dari dua arah secara simultan maju dan mundur [16]. Pendekatan dua arah ini memungkinkan model memahami konteks kata tidak hanya berdasarkan urutan sebelumnya tetapi juga informasi setelahnya, sehingga representasi urutan teks menjadi lebih komprehensif [7]. Integrasi BiLSTM memperkuat pemahaman model terhadap urutan kata secara eksplisit, yang sangat krusial untuk mengenali pola naratif dalam berita [7].

2.5. Model Hybrid IndoBERT-BiLSTM

Model hybrid menggabungkan kekuatan representasi kontekstual IndoBERT dengan kemampuan sekuensial BiLSTM untuk memahami hubungan antar kata secara lebih mendalam [7]. Penelitian sebelumnya menunjukkan bahwa kombinasi IndoBERT dan LSTM sudah mampu menghasilkan performa deteksi hoaks yang kuat. Studi

oleh Yefferson *et al.* [17] melaporkan akurasi 93,2 persen dengan F1-Score 90,8 persen, membuktikan bahwa integrasi representasi kontekstual IndoBERT dan pemodelan sekuensial LSTM efektif dalam memahami karakteristik teks hoaks berbahasa Indonesia. Berdasarkan hasil tersebut, penggunaan IndoBERT yang dipadukan dengan BiLSTM menjadi pendekatan yang lebih unggul secara konseptual karena BiLSTM memproses konteks dua arah sehingga mampu menangkap relasi kata secara lebih lengkap, yang berpotensi meningkatkan akurasi identifikasi pola manipulatif dalam berita.

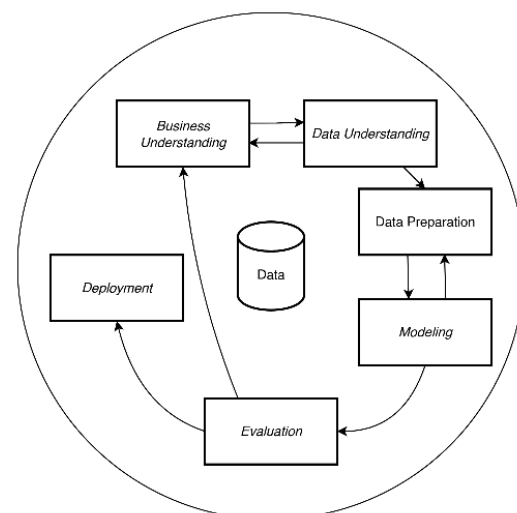
2.6. Hyperparameter Tuning

Hyperparameter tuning merupakan proses optimasi parameter eksternal model yang tidak dipelajari secara langsung selama pelatihan, seperti *learning rate*, *batch size*, dan *dropout*. Pemilihan konfigurasi yang tepat sangat memengaruhi performa model karena menentukan stabilitas proses pelatihan dan kemampuan generalisasi terhadap data baru. Penelitian ini menggunakan pendekatan *Bayesian Optimization* melalui Optuna karena metode ini mampu mengeksplorasi ruang parameter secara adaptif dan efisien dibandingkan teknik konvensional seperti *grid search* atau *random search* [18].

2.7. Metrik Evaluasi Kinerja Model

Evaluasi performa model bertujuan menilai kemampuan sistem dalam mengklasifikasikan teks secara akurat dan konsisten [7]. Pengukuran dilakukan menggunakan *Confusion Matrix* yang memuat nilai *True Positive* (TP), *True Negative* (TN), *False Positive* (FP), dan *False Negative* (FN) sebagai dasar perhitungan metrik evaluasi. Metrik yang digunakan meliputi akurasi, presisi, *recall*, *f1-score*, dan ROC-AUC karena kombinasi metrik tersebut memberikan gambaran performa yang komprehensif, khususnya dalam menilai kemampuan model meminimalkan kesalahan klasifikasi hoaks sebagai berita valid [4].

2.8. Metode CRISP-DM



Gambar 1. Metode CRISP-DM

CRISP-DM merupakan metodologi *data mining* dengan enam fase iteratif yaitu *Business Understanding*, *Data Understanding*, *Data Preparation*, *Modeling*, *Evaluation*, dan *Deployment* [19]. Kerangka ini menyediakan alur sistematis untuk memastikan proses pengembangan model berjalan terstruktur mulai dari identifikasi permasalahan hingga implementasi sistem, sehingga hasil penelitian tidak hanya akurat secara teknis tetapi juga relevan terhadap kebutuhan pengguna.

2.9. Berita Hoaks

Berita hoaks merupakan bentuk informasi palsu yang dapat dikategorikan sebagai disinformasi, misinformasi, atau malinformasi yang sengaja disusun untuk menyesatkan publik [20]. Di Indonesia, penyebaran hoaks meningkat seiring tingginya penetrasi media digital dan rendahnya tingkat literasi informasi Masyarakat [21]. Kondisi ini menuntut adanya solusi otomatis berbasis AI yang mampu mengidentifikasi konten palsu secara cepat dan akurat, sehingga proses mitigasi penyebaran informasi menyesatkan dapat dilakukan secara lebih efisien dan sistematis [7].

3. METODE PENELITIAN

Penelitian ini menggunakan pendekatan *development research* dengan kerangka kerja CRISP-DM karena sifatnya yang sistematis, adaptif, dan iteratif dalam menangani eksplorasi data kompleks. Proses penelitian mencakup tahapan mulai dari identifikasi fenomena hoaks politik di Indonesia hingga publikasi hasil pengembangan model deteksi hoaks berbasis arsitektur hybrid IndoBERT-BiLSTM.

3.1. Business Understanding

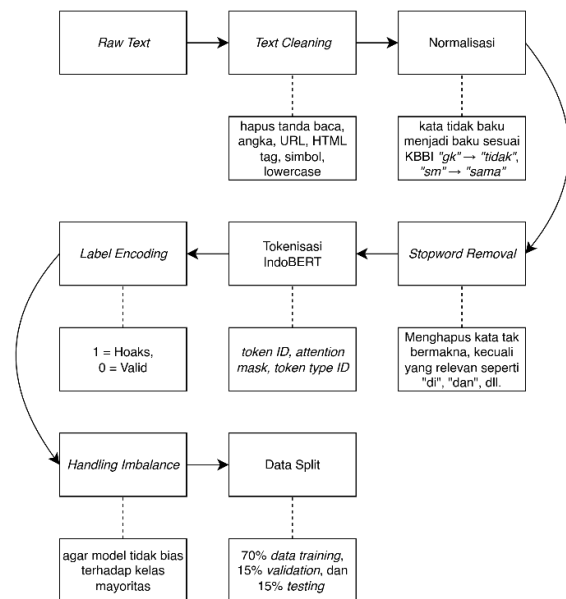
Tahap ini menyelaraskan tujuan teknis dengan kebutuhan sosial masyarakat digital. Mengingat hoaks politik dapat mengancam stabilitas sosial, penelitian ini diorientasikan untuk menghasilkan solusi praktis yang dapat digunakan oleh institusi atau media dalam menyaring informasi palsu secara konsisten.

3.2. Data Understanding

Eksplorasi dilakukan terhadap *dataset* "Deteksi Berita Hoaks Indo - Dataset" dari Kaggle yang berisi 24.658 sampel berita. Pada tahap ini, peneliti memverifikasi kualitas data dan distribusi label, di mana ditemukan keseimbangan antara berita hoaks sebanyak 12.012 data (48,7%) dan berita valid sebanyak 12.646 data (51,3%). Jika ditemukan masalah kualitas pada tahap ini, proses dapat dikembalikan ke tahap *Business Understanding* untuk merevisi ruang lingkup penelitian.

3.3. Data Preparation

Data teks diproses melalui rangkaian pipeline agar layak menjadi masukan model:



Gambar 2. Pipeline data preparation

Berdasarkan gambar 2 secara garis besar *Data Preparation* melalui *preprocessing data* meliputi beberapa langkah utama, yaitu *text cleaning* untuk menghapus tanda baca, angka, URL, dan tag HTML, normalisasi kata dengan mengonversi bentuk tidak baku menjadi bentuk baku sesuai KBBI, *stopword removal* secara selektif untuk menghilangkan kata umum yang tidak berkontribusi signifikan terhadap makna, serta tokenisasi menggunakan *tokenizer* IndoBERT guna mengubah teks menjadi representasi numerik berupa token ID dan *attention mask* [22]. Setelah seluruh proses *preprocessing* selesai, *dataset* kemudian dibagi secara stratifikasi menjadi data pelatihan (70%), validasi (15%), dan pengujian (15%) untuk memastikan distribusi kelas tetap seimbang pada setiap subset.

3.4. Modeling

Tahap inti di mana arsitektur hybrid dibangun. IndoBERT bertugas mengekstraksi fitur kontekstual dua arah, yang hasilnya diteruskan ke lapisan BiLSTM untuk menangkap informasi sekuensial naratif berita. Optimasi dilakukan menggunakan Optuna untuk mencari *learning rate*, unit LSTM, dan *dropout rate* terbaik.

3.5. Evaluation

Model diuji menggunakan metrik akurasi, presisi, *recall*, *f1-score*, dan ROC-AUC.

Tabel 1. Target metrik evaluasi kinerja model

Metrik	Target Minimum
<i>F1-score</i> (Hoaks)	≥ 80%
ROC-AUC	≥ 0.85
Presisi (Hoaks)	≥ 75%
<i>Recall</i> (Hoaks)	≥ 75%

Jika performa belum mencapai target yang telah ditentukan oleh peneliti seperti pada tabel 1, proses dikembalikan ke tahap *Data Preparation* untuk perbaikan kualitas data.

3.6. Deployment

Model diimplementasikan ke dalam prototipe antarmuka web menggunakan Gradio. Sistem ini memungkinkan pengguna memasukkan teks berita secara manual atau memilih kartu berita yang disediakan untuk dianalisis secara otomatis, dengan hasil akhir berupa label klasifikasi disertai nilai probabilitas keyakinan model.

4. HASIL DAN PEMBAHASAN

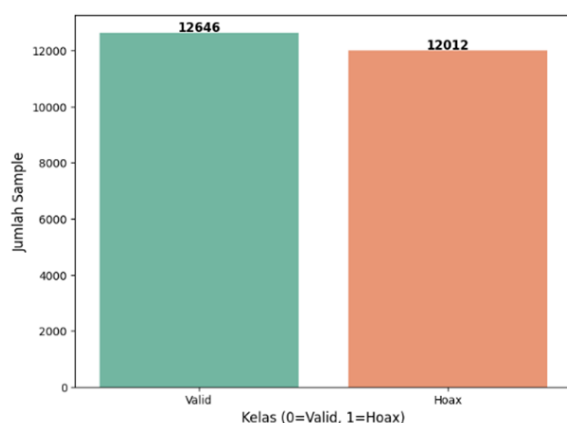
Memuat hasil dan pembahasan dari penelitian yang telah dilakukan. Hasil penelitian mencakup implementasi dari seluruh tahapan metodologi CRISP-DM, mulai dari *business understanding* hingga *deployment* sistem.

4.1. Business Understanding

Penelitian ini bertujuan mengembangkan sistem klasifikasi berita hoaks berbahasa Indonesia berbasis arsitektur hybrid IndoBERT-BiLSTM yang akurat dan siap diimplementasikan secara praktis. Keberhasilannya diukur melalui tiga indikator utama, yaitu performa tinggi pada data uji (akurasi dan *f1-score*), peningkatan kinerja signifikan melalui optimasi *hyperparameter* menggunakan Optuna dibanding parameter *default*, serta keberhasilan integrasi model ke dalam prototipe aplikasi web berbasis Gradio yang mampu melakukan deteksi secara *real-time*.

4.2. Data Understanding

Dataset penelitian ini merupakan hasil integrasi berbagai sumber kredibel untuk merepresentasikan berita di Indonesia, dengan data hoaks (label 1) diperoleh dari TurnBackHoax dan data valid (label 0) dari *scraping* artikel Kompas, Detik, dan CNN Indonesia. Setelah proses penggabungan dan pembersihan awal, total *dataset* berjumlah 24.658 artikel.



Gambar 3. Distribusi kelas *dataset*

Distribusi kelas *dataset* tergolong seimbang (*balanced*), sebagaimana ditunjukkan pada gambar 3, di mana label Valid berjumlah 12.646 artikel (51,29%) dan Hoaks 12.012 artikel (48,71%), dengan selisih kurang dari 3% yang menandakan tidak adanya masalah *data imbalance* signifikan. Kondisi ini memberikan dasar evaluasi yang kuat karena metrik akurasi dapat digunakan secara andal tanpa bias kelas, meskipun penelitian tetap menggunakan *precision*, *recall*, dan *f1-score* untuk memastikan konsistensi performa model IndoBERT-BiLSTM dalam mendeteksi kedua kelas secara setara.

4.3. Data Preparation

Tahap persiapan data dilakukan untuk memastikan kualitas teks dan kompatibilitas format data dengan arsitektur model IndoBERT-BiLSTM. Proses ini mencakup dua tahapan utama data *preprocessing* dan pembagian dataset.

Tabel 2. Perbandingan data asli dan bersih

Teks Asli (Sebelum)	Teks Bersih (Sesudah)
Mahfud: Setiap Hari Kita Dengar Berita Korupsi Triliunan bahwa saat ini berita terkaitkorupsiselalu hadir setiap harinya . disampaikan Mahfud dalam dialog publik yang mengangkat tema "Enam Bulan Pemerintahan Prabowo", digelar di Universitas Paramadina, Jakarta, pada Kamis . ujar Mahfud dikutip dari kanal Youtube	mahfud setiap hari kita dengar berita korupsi triliunan bahwa saat ini berita terkaitkorupsiselalu hadir setiap harinya . disampaikan mahfud dalam dialog publik yang mengangkat tema enam bulan pemerintahan prabowo digelar di universitas paramadina jakarta pada kamis ujar mahfud dikutip dari kanal youtube
[PENIPUAN] Tautan Lowongan Kerja PT Ajinomoto Indonesia pemeriksaan Fakta Tim Pemeriksa Fakta Mafindo (TurnBackHoax) mengakses tautan yang tersemat di bio akun pengunggah . untuk melamar lowongan pekerjaan untuk bidang section manager, sales, regulation staff, production, hingga operational untuk wilayah Jakarta dan Surabaya .	penipuan tautan lowongan kerja pt ajinomoto indonesia pemeriksaan fakta tim pemeriksa fakta mafindo turnbackhoax mengakses tautan yang tersemat di bio akun pengunggah untuk melamar lowongan pekerjaan untuk bidang section manager sales regulation staff production hingga operational untuk wilayah jakarta dan surabaya

Text preprocessing dilakukan untuk mengurangi *noise* pada data mentah melalui serangkaian *text cleaning* yang meliputi *case folding* (huruf kecil), penghapusan URL, pembersihan tanda baca dan karakter non-alfabet menggunakan regex, serta normalisasi spasi. Tahap ini penting agar model berfokus pada fitur semantik tanpa terdistraksi variasi penulisan, dengan hasil pembersihan yang membuat teks lebih rapi dan terstruktur tanpa menghilangkan

makna kontekstual, sebagaimana ditunjukkan pada tabel 2.

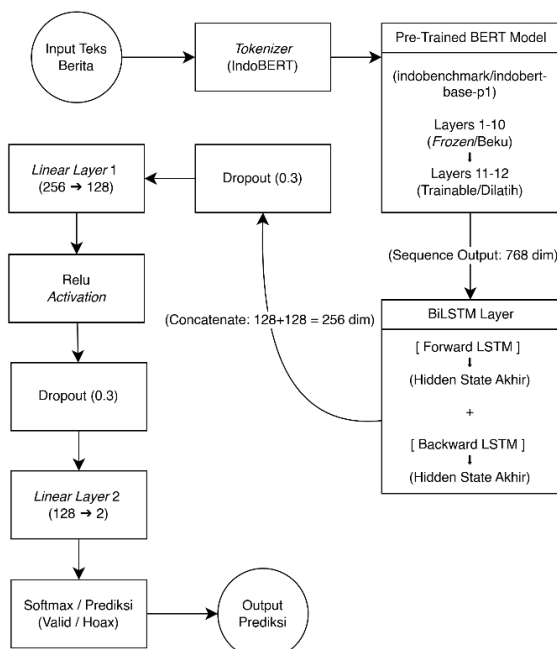
Tabel 3. Distribusi pembagian dataset

Dataset	Jumlah Sampel	Persentase
Data Latih (Training)	17.260	70%
Data Validasi (Validation)	3.699	15%
Data Uji (Testing)	3.699	15%
Total	24.658	100%

Dataset yang telah dibersihkan dibagi menjadi tiga subset dengan rasio 70% data latih, 15% validasi, dan 15% uji, di mana data validasi digunakan selama *hyperparameter tuning* Optuna untuk mencegah *overfitting*, sedangkan data uji digunakan khusus untuk evaluasi akhir seperti pada tabel 3. Setelah itu, dilakukan tokenisasi menggunakan *tokenizer* IndoBERT untuk mengubah teks menjadi input IDs dan *attention mask* sesuai format *input* model sehingga siap diproses oleh lapisan *embedding*.

4.4. Modeling

Tahap ini merupakan inti penelitian yang mencakup implementasi arsitektur model hybrid IndoBERT-BiLSTM serta proses optimasi *hyperparameter*. Arsitektur dirancang dengan menggabungkan kemampuan representasi kontekstual dari model pralatih IndoBERT (*indobenchmark/indobert-base-p1*) dan kemampuan pemodelan dependensi sekuensial dari *Bidirectional Long Short-Term Memory* (BiLSTM).



Gambar 4. Diagram arsitektur model

Pada tahap awal, teks ditokenisasi dengan *tokenizer* IndoBERT menjadi representasi numerik lalu diproses oleh IndoBERT dengan sepuluh lapisan

awal dibekukan dan lapisan akhir *di-fine-tune* agar sesuai karakteristik data hoaks. *Last hidden state* diteruskan ke BiLSTM untuk menangkap dependensi dua arah, kemudian diproses oleh *fully connected layer* dengan aktivasi ReLU dan *dropout* guna meningkatkan generalisasi serta menekan *overfitting*, sebelum diklasifikasikan menggunakan *sigmoid* atau *softmax* seperti pada gambar 4. Penelitian ini menguji dua skenario, yaitu model baseline dan model hasil *hyperparameter tuning*, untuk membandingkan performa sebelum dan sesudah optimasi parameter.

Pada skenario pertama, model dilatih dengan *hyperparameter* yang ditentukan manual berdasarkan praktik umum penelitian sebelumnya meliputi.

Tabel 4. Parameter baseline

Parameter	Nilai
<i>lstm_hidden_size</i>	128
<i>lstm_num_layers</i>	1
<i>lstm_dropout</i>	0.3
<i>classifier_dropout</i>	0.3
<i>learning_rate</i>	3e-5
<i>batch_size</i>	16
<i>l2_reg</i>	0.01
<i>freeze_bert_layers</i>	10
<i>max_length</i>	128

Pada skenario kedua, model menggunakan *hyperparameter* optimal hasil tuning Optuna sebanyak 15 *trial*, yang dipilih untuk menyeimbangkan performa dan efisiensi komputasi GPU. Optimasi difokuskan pada parameter tambahan karena IndoBERT telah memiliki representasi bahasa kuat dari *pre-training*, dan jumlah *trial* dinilai cukup karena *f1-score* validasi telah konvergen tanpa peningkatan signifikan pada percobaan tambahan.

Tabel 5. Search space hyperparameter tuning

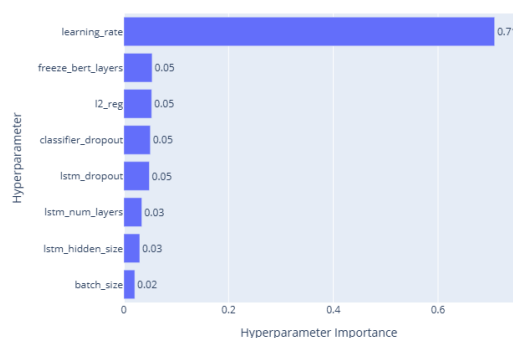
Parameter	Tipe	Rentang Nilai (Search Space)
<i>lstm_hidden_size</i>	Categorical	[64, 128, 256]
<i>lstm_num_layers</i>	Integer	1 - 2
<i>lstm_dropout</i>	Float	0.1 - 0.4
<i>classifier_dropout</i>	Float	0.1 - 0.4
<i>learning_rate</i>	Float (LogUniform)	1e-6 - 5e-5
<i>batch_size</i>	Categorical	[8, 16, 32]
<i>l2_reg</i>	Float (LogUniform)	1e-5 - 1e-3
<i>freeze_bert_layers</i>	Integer	8 - 11

Ruang pencarian (*search space*) parameter mencakup variasi arsitektur LSTM, regularisasi, ukuran batch, jumlah lapisan yang dibekukan pada BERT, serta laju pembelajaran. Optuna menggunakan pendekatan *define-by-run* yang memungkinkan eksplorasi ruang parameter secara adaptif berdasarkan hasil *trial* sebelumnya. Strategi ini dimulai dengan tahap eksplorasi untuk memetakan performa awal, dilanjutkan dengan eksploitasi pada area parameter yang menjanjikan guna menghindari *local optima*.

Tabel 6. Hasil *hyperparameter tuning* dengan Optuna

Trial	Value (Score)	LSTM Hidden Size	LSTM Layers	LSTM Dropout	Classifier Dropout	Learning Rate	Batch Size	L2 Reg	Freeze BERT Layers
0	0.9640	64	1	0.2214	0.2133	5.29e-06	16	0.000116	10
1	0.9849	256	2	0.3133	0.1823	4.31e-05	32	5.10e-05	10
2	0.9738	64	2	0.1762	0.3434	5.01e-06	8	0.000748	9
3	0.9735	128	2	0.3178	0.1670	4.37e-06	16	0.000252	8
4	0.9686	256	1	0.3424	0.1220	8.38e-06	16	0.000170	11
5	0.9727	256	2	0.3422	0.1689	5.64e-06	32	1.04e-05	8
6	0.9808	128	1	0.2515	0.3582	1.47e-05	16	1.01e-05	10
7	0.9822	128	2	0.3544	0.1651	3.22e-05	32	0.000506	9
8	0.9757	128	2	0.2566	0.2220	1.55e-05	16	2.85e-05	11
9	0.9313	64	1	0.3236	0.3230	1.40e-06	16	1.88e-05	8
10	0.9859	256	2	0.1275	0.2828	4.36e-05	32	5.06e-05	10
11	0.9849	256	2	0.1086	0.2870	4.74e-05	32	5.17e-05	10
12	0.9819	256	2	0.1203	0.2924	2.59e-05	32	5.54e-05	10
13	0.9832	256	2	0.1022	0.2783	5.00e-05	32	5.45e-05	11
14	0.9854	256	2	0.1538	0.2728	2.01e-05	8	2.73e-05	9

Hasil optimasi menunjukkan bahwa performa model mengalami fluktuasi selama proses *tuning*, dengan nilai terendah sebesar 0,9313 pada *trial* ke-9 dan nilai tertinggi sebesar 0,9859 pada *trial* ke-10. Secara keseluruhan, sebagian besar *trial* menghasilkan skor di atas 0,96, yang menunjukkan bahwa ruang pencarian yang dirancang telah efektif mengarahkan model menuju konfigurasi optimal. Pola ini mengindikasikan bahwa strategi eksplorasi parameter berjalan stabil dan efisien.

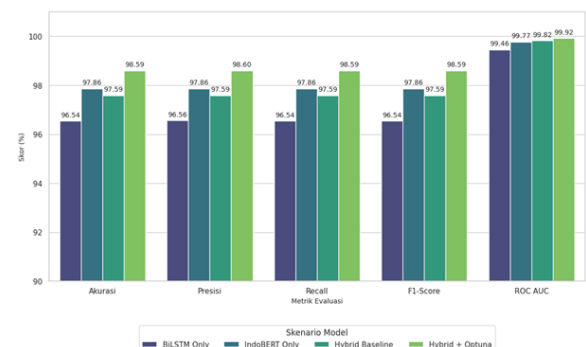
Gambar 5. *Hyperparameter importance*

Analisis *hyperparameter importance* menunjukkan bahwa learning rate menjadi faktor paling dominan dengan kontribusi sekitar 71% terhadap performa model, menegaskan peran krusialnya dalam proses *fine-tuning* arsitektur berbasis transformer. Parameter lain seperti jumlah lapisan BERT yang dibekukan, regulasi L2, dan dropout memberikan pengaruh lebih kecil dan relatif merata, sedangkan parameter arsitektur LSTM termasuk jumlah lapisan, ukuran *hidden state*, dan *batch size* memiliki dampak paling rendah.

4.5. Evaluation

Pada tahap ini, seluruh model BiLSTM *only*, IndoBERT *only*, hybrid *baseline*, dan hybrid + Optuna dievaluasi menggunakan *test set* yang tidak terlibat dalam proses pelatihan. Evaluasi dilakukan menggunakan metrik akurasi, presisi, *recall*, *f1-score*,

dan ROC AUC untuk memberikan gambaran performa klasifikasi secara menyeluruh.



Gambar 6. Perbandingan metrik evaluasi model

Berdasarkan gambar 6, model BiLSTM *only* memperoleh performa terendah dengan akurasi dan *f1-score* sebesar 96,54%, menunjukkan keterbatasan arsitektur RNN dalam menangkap konteks semantik kompleks. Model IndoBERT *only* meningkat menjadi 97,86%, menegaskan keunggulan representasi bahasa dari model pralatih. Sementara itu, hybrid *baseline* sedikit menurun menjadi 97,59%, mengindikasikan bahwa penambahan kompleksitas arsitektur tanpa optimasi parameter belum tentu meningkatkan performa. Setelah dilakukan *tuning* menggunakan Optuna, model hybrid + Optuna mencapai hasil terbaik dengan akurasi 98,59% dan ROC AUC 0,9992.

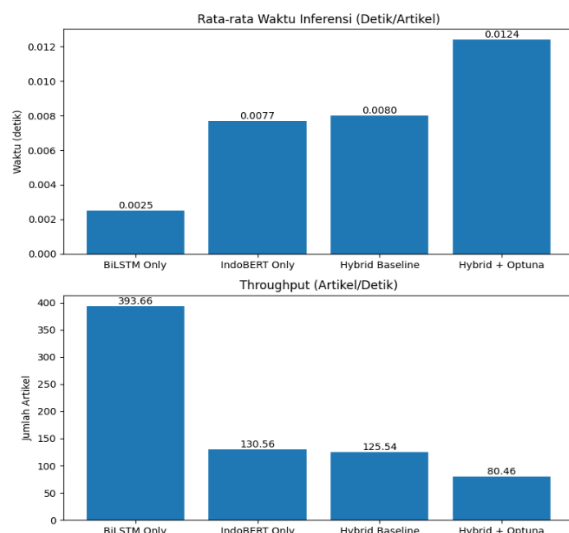
Tabel 7. *Confusion matrix* model terbaik

	Prediksi: Valid	Prediksi: Hoaks
Aktual: Valid	1861 (TN)	36 (FP)
Aktual: Hoaks	16 (FN)	1786 (TP)

Analisis kesalahan melalui tabel 7 menunjukkan bahwa model terbaik mampu mengklasifikasikan sebagian besar data secara benar dengan kesalahan minimal. Kesalahan *false negative* umumnya disebabkan oleh berita hoaks yang menggunakan gaya bahasa formal sehingga menyerupai teks valid,

sedangkan *false positive* banyak dipicu oleh *noise* teks seperti repetisi kata akibat kesalahan ekstraksi.

Setelah itu didapatkan juga nilai ROC-AUC sebesar 0,9992 yang menandakan kemampuan diskriminasi kelas yang sangat tinggi. Selain akurasi, efisiensi komputasi juga diuji menggunakan metrik waktu inferensi dan *throughput* pada 100 sampel dengan panjang maksimum 128 token untuk menjaga konsistensi beban proses.

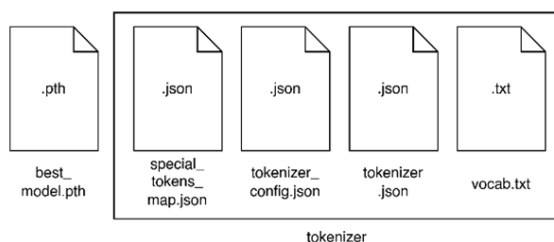


Gambar 7. Inference dan troughput time

Hasil pada gambar 7 menunjukkan bahwa model BiLSTM saja memiliki waktu inferensi tercepat (0,0025 detik/artikel) dan *throughput* tertinggi (393,66 artikel/detik) karena arsitekturnya sederhana, sedangkan model berbasis Transformer membutuhkan waktu lebih lama IndoBERT only 0,0077 detik, hybrid *baseline* 0,0080 detik, dan hybrid + Optuna 0,0124 detik dengan *throughput* 80,46 artikel/detik. Secara keseluruhan, terdapat *trade-off* antara akurasi dan efisiensi komputasi; meskipun hybrid + Optuna paling lambat, model ini menghasilkan performa prediksi terbaik sehingga dinilai sebagai konfigurasi paling optimal untuk implementasi sistem deteksi hoaks.

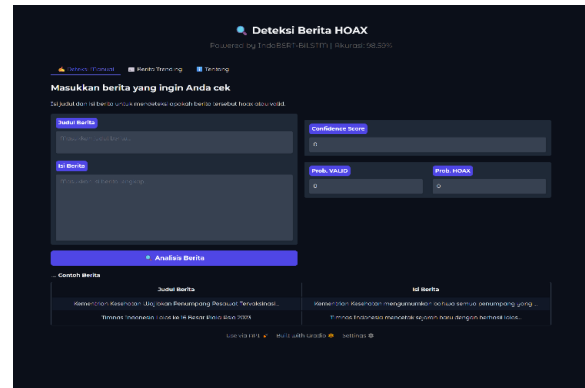
4.6. Deployment

Tahap deployment diwujudkan dalam bentuk prototipe aplikasi web interaktif menggunakan pustaka Gradio.

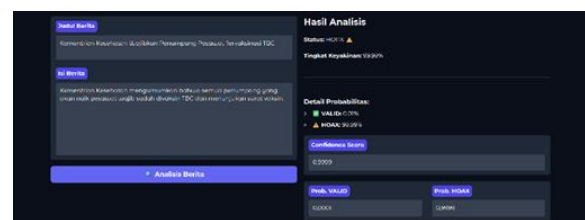


Gambar 7. Komponen file model IndoBERT-BiLSTM

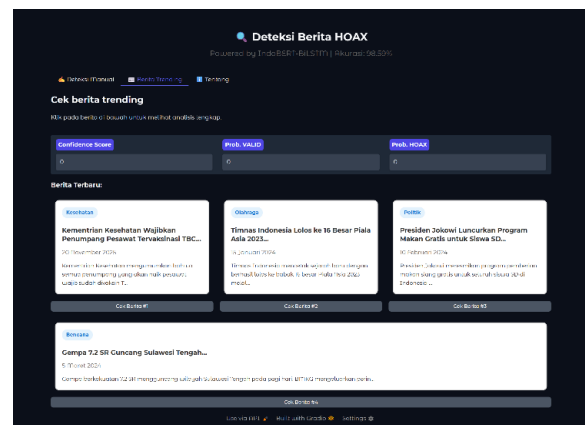
Model terbaik hasil proses *hyperparameter tuning* dengan Optuna beserta *tokenizer* IndoBERT disimpan dan kemudian dimuat kembali ke dalam skrip aplikasi untuk keperluan inferensi. Komponen file yang diperlukan ditunjukkan pada Gambar X, yang meliputi model terlatih berformat `.pth` sebagai penyimpanan bobot jaringan, file *tokenizer* (konfigurasi dan kosakata), serta berkas konfigurasi *hyperparameter*.



Gambar 8. Tampilan halaman deteksi berita



Gambar 9. Tampilan halaman deteksi berita saat test



Gambar 10. Tampilan halaman berita trending

Antarmuka aplikasi memungkinkan pengguna memasukkan teks berita secara langsung melalui kotak input. Setelah tombol analisis ditekan, sistem menjalankan rangkaian proses otomatis yang meliputi *text preprocessing*, tokenisasi, inferensi model, dan penampilan hasil klasifikasi. Tampilan antarmuka prototipe dapat dilihat pada gambar 8, sedangkan contoh hasil prediksi ditunjukkan pada gambar 9. Fitur tambahan berupa halaman berita *trending* juga disediakan untuk menguji model pada teks aktual sebagaimana ditampilkan pada gambar 10.

5. KESIMPULAN DAN SARAN

Penelitian ini mengembangkan sistem deteksi hoaks otomatis berbasis arsitektur hybrid IndoBERT-BiLSTM untuk mengatasi keterbatasan verifikasi manual dan metode konvensional, menggunakan metodologi CRISP-DM serta *dataset* seimbang 24.658 artikel berita. Hasil menunjukkan model hybrid tanpa optimasi belum optimal (akurasi 97,59%), sehingga tuning Optuna dengan *learning rate* sebagai faktor dominan (71%) menjadi kunci peningkatan performa. Model terbaik mencapai akurasi dan F1-score 98,59% serta ROC-AUC 0,9992 dengan waktu inferensi 0,0124 detik per artikel, dan berhasil diimplementasikan dalam prototipe web.

DAFTAR PUSTAKA

- [1] I Nyoman Purnama and Ni Nengah Widya Utami, "IMPLEMENTASI PERINGKAS DOKUMEN BERBAHASA INDONESIA MENGGUNAKAN METODE TEXT TO TEXT TRANSFER TRANSFORMER (T5)," *Jurnal Teknologi Informasi dan Komputer*, vol. 9, no. 4, pp. 381–391, Aug. 2023, doi: 10.36002/jutik.v9i4.2531.
- [2] M. Laia, Ayuliana, Wasiran, M. L. Hakim, and D. Suryadi, "Analisis Big Data untuk Deteksi Hoaks dan Disinformasi di Platform Berita Online," *Jurnal JTik (Jurnal Teknologi Informasi dan Komunikasi)*, vol. 9, no. 2, pp. 776–782, Feb. 2025, doi: 10.35870/jtik.v9i2.3859.
- [3] R. DickiPrabowo, I. Widaningrum, and J. Karaman, "SISTEM DETEKSI BERITA HOAX PEMILU 2024 INDONESIA MENGGUNAKAN ALGORITMA K-NEAREST NEIGHBOR (KNN) DAN SUPPORT VECTOR MACHINE (SVM)," *JIKO (Jurnal Informatika dan Komputer)*, vol. 9, no. 1, p. 93, Feb. 2025, doi: 10.26798/jiko.v9i1.1424.
- [4] D. F. S. Arianti, U. S. Aesyi, and A. Himawan, "Klasifikasi Berita Hoaks Pasca Pemilihan Umum Presiden Dan Wakil Presiden Republik Indonesia Tahun 2024 Menggunakan Algoritma K-Nearest Neighbour," *Angkasa: Jurnal Ilmiah Bidang Teknologi*, vol. 17, no. 1, p. 64, May 2025, doi: 10.28989/angkasa.v17i1.2596.
- [5] B. Wilie *et al.*, "IndoNLU: Benchmark and Resources for Evaluating Indonesian Natural Language Understanding," in *Proceedings of the 1st Conference of the Asia-Pacific Chapter of the Association for Computational Linguistics and the 10th International Joint Conference on Natural Language Processing*, Stroudsburg, PA, USA: Association for Computational Linguistics, 2020, pp. 843–857. doi: 10.18653/v1/2020.aacl-main.85.
- [6] S. Cahyawijaya *et al.*, "IndoNLG: Benchmark and Resources for Evaluating Indonesian Natural Language Generation," in *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing*, Stroudsburg, PA, USA: Association for Computational Linguistics, 2021, pp. 8875–8898. doi: 10.18653/v1/2021.emnlp-main.699.
- [7] L. B. Wijaya, Y. N. Wicaksana, S. S. Widhiasari, and A. Saptawijaya, "Pengembangan Model Deteksi Hoaks Berbahasa Indonesia Menggunakan Kombinasi IndoBERT dan BiLSTM.," *Buletin Gemastik*, vol. 2, no. 1, pp. 12–16, Apr. 2024.
- [8] M. D. Maulana, C. Sri, and K. Aditya, "Perbandingan IndoBERT dan Bi-LSTM Dalam Mendeteksi Pelanggaran Undang-Undang ITE," *Science and Information Technology (SINTECH)*, vol. 8, no. 1, pp. 52–59, Apr. 2025, doi: <https://doi.org/10.31598>.
- [9] Anugerah Simanjuntak *et al.*, "Research and Analysis of IndoBERT Hyperparameter Tuning in Fake News Detection," *Jurnal Nasional Teknik Elektro dan Teknologi Informasi*, vol. 13, no. 1, pp. 60–67, Feb. 2024, doi: 10.22146/jnteti.v13i1.8532.
- [10] S. Russell and P. Norvig, *Artificial Intelligence: A Modern Approach*, 4th (Fourth). Pearson, 2021.
- [11] A. A. I. I. Paramitha and I Nyoman Mahayasa Adiputra, "DETEKSI KENDARAAN PADA LALU LINTAS MENGGUNAKAN ARTIFICIAL INTELLIGENCE UNTUK MENDUKUNG DENPASAR SMART CITY," *Jurnal Informatika Teknologi dan Sains*, vol. 4, no. 4, pp. 353–358, Nov. 2022, doi: 10.51401/jinteks.v4i4.2074.
- [12] R. C. Tarumingkeng, "Natural Language Processing (NLP)," rudict.com.
- [13] T. B. Brown *et al.*, "Language Models are Few-Shot Learners," vol. 2020-December, Jul. 2020, [Online]. Available: <http://arxiv.org/abs/2005.14165>
- [14] C. L. Goh, "Applying Masked Language Models to Search for Suitable Verbs Used in Academic Writing," in *Proceedings of the 35th Pacific Asia Conference on Language, Information and Computation*, K. Hu, J.-B. Kim, C. Zong, and E. Chersoni, Eds., Shanghai, China: Association for Computational Linguistics, Jan. 2021, pp. 180–188.
- [15] B. Juarto, "Indonesian News Classification Using IndoBert," *Original Research Paper International Journal of Intelligent Systems and Applications in Engineering IJISAE*, no. 2, pp. 454–460, Feb. 2023.
- [16] M. A. P. Putra, D.-S. Kim, and J.-M. Lee, "DB-BiLSTM: Euclidean Distance-Based Sensor Data Prediction for IoT Applications," in *2021 International Conference on Information and Communication Technology Convergence (ICTC)*, IEEE, Oct. 2021, pp. 814–817. doi: 10.1109/ICTC52510.2021.9620877.

-
- [17] D. Y. Yefferson, V. Lawijaya, and A. S. Girsang, "Hybrid model: IndoBERT and long short-term memory for detecting Indonesian hoax news," *IAES International Journal of Artificial Intelligence (IJ-AI)*, vol. 13, no. 2, p. 1913, Jun. 2024, doi: 10.11591/ijai.v13.i2.pp1913-1924.
- [18] M. A. P. Putra, K. R. P. Utama, N. W. Utami, and I. G. J. E. Putra, "Enhancing Federated Learning Performance through Adaptive Client Optimization with Hyperparameter Tuning," *Journal of Applied Data Sciences*, vol. 5, no. 2, pp. 747–755, May 2024, doi: 10.47738/jads.v5i2.251.
- [19] A. Rianti, N. W. A. Majid, and A. Fauzi, "CRISP-DM: Metodologi Proyek Data Science," in *Prosiding Seminar Nasional Teknologi Informasi dan Bisnis (SENATIB)*, Jul. 2023, pp. 107–114.
- [20] T. L. Khasanah *et al.*, "Analisis Literasi Digital Masyarakat Indonesia Terkait Hoaks Libur Sekolah Saat Ramadhan Tahun 2025," *Jurnal Ilmiah Multidisiplin*, vol. 2, no. 6, pp. 406–412, 2025, doi: 10.62017/merdeka.
- [21] Fauzi and Marhamah, "Pengaruh Literasi Digital Terhadap Pencegahan Informasi Hoaks pada Remaja di SMANegeri 7 Kota Lhokseumawe," *Jurnal Pekommas Vol. 6 No.*, vol. 2, pp. 77–84, 2021, doi: 10.30818/jpkm.2021.2060210.
- [22] I Made Artana and N. W. Utami, "PENERAPAN DATA MINING UNTUK MENENTUKAN STRATEGI PROMOSI PRODUK INDUSTRI KREATIF UMKM KOTA DENPASAR PASCA PANDEMI COVID 19," *Jurnal Informatika Teknologi dan Sains*, vol. 4, no. 2, pp. 124–133, May 2022, doi: 10.51401/jinteks.v4i2.2032.