

PERBANDINGAN MODEL *TRANSFORMER*, *DEEP LEARNING*, DAN *MACHINE LEARNING* UNTUK DETEKSI BERITA PALSU: STUDI KASUS PADA TEKS BERBAHASA INDONESIA

Arief Rahman Hakim¹, Alva Hendi Muhammad²

^{1,2} Magister Informatika, Universitas Amikom Yogyakarta

Jl. Ring Road Utara, Ngringin, Condongcatur, Kec. Depok, Kabupaten Sleman, Daerah Istimewa Yogyakarta 55281

¹ ariefrh@students.amikom.ac.id, ² alva@amikom.ac.id

Abstract

Detecting fake news in Indonesian remains a significant challenge in natural language processing (NLP). This study compares six methods—RoBERTa, BERT, IndoBERT, SVM, LSTM, and CNN—for fake news identification. The dataset used was preprocessed through cleaning and tokenization before being applied to each model. Results show that Transformer-based models outperform traditional approaches such as SVM, CNN, and LSTM. Additionally, IndoBERT, which is specifically trained for Indonesian, achieves better performance than the standard BERT model. RoBERTa demonstrates the highest accuracy at 99.5%, followed by IndoBERT (98.6%), BERT (98.2%), SVM (95.9%), CNN (93.9%), and LSTM (92.3%). These findings can serve as a reference for developing more accurate and efficient fake news detection systems in Indonesian.

Keywords : Fake News, RoBERTa, BERT, IndoBERT, SVM, LSTM, CNN, NLP, Transformer.

Abstrak

Deteksi berita palsu dalam bahasa Indonesia masih menjadi tantangan dalam pemrosesan bahasa alami (NLP). Penelitian ini membandingkan enam metode: RoBERTa, BERT, IndoBERT, SVM, LSTM, dan CNN dalam mengidentifikasi berita palsu. Dataset yang digunakan telah melalui proses pembersihan dan tokenisasi sebelum diterapkan pada masing-masing model. Penelitian ini memberikan analisis komprehensif terhadap keunggulan model *Transformer* dibandingkan dengan metode klasik seperti SVM, CNN, dan LSTM. Selain itu, penelitian ini juga menegaskan bahwa model yang dilatih khusus untuk bahasa Indonesia, seperti IndoBERT, memiliki performa lebih baik dibandingkan BERT standar. Hasil evaluasi menunjukkan bahwa model berbasis *Transformer* memiliki performa terbaik, dengan RoBERTa sebagai model paling akurat. Temuan ini dapat menjadi referensi bagi pengembangan sistem deteksi berita palsu yang lebih akurat dan efisien dalam bahasa Indonesia. Akurasi yang diperoleh dari masing-masing model adalah sebagai berikut: RoBERTa (99,5%), IndoBERT (98,6%), BERT (98,2%), SVM (95,9%), CNN (93,9%), dan LSTM (92,3%).

Kata kunci : Berita Palsu, RoBERTa, BERT, IndoBERT, SVM, LSTM, CNN, NLP, Transformer.

1. PENDAHULUAN

Berita palsu telah menjadi tantangan besar di era digital [1], terutama dengan meningkatnya penyebaran informasi melalui media sosial dan platform online. Penyebaran berita palsu dapat menimbulkan dampak yang signifikan, termasuk mempengaruhi opini publik, memicu

disinformasi, dan bahkan berkontribusi pada ketidakstabilan sosial dan politik. Oleh karena itu, deteksi berita palsu telah menjadi bidang penelitian yang sangat penting dalam pemrosesan bahasa alami (*Natural Language Processing*, NLP).

Sejumlah penelitian sebelumnya telah berkontribusi dalam upaya mengembangkan



metode deteksi berita palsu. [2] mengkaji pendekatan berbasis *deep learning* seperti *Long Short-Term Memory* (LSTM) dan *Convolutional Neural Network* (CNN) dalam mendeteksi berita palsu, menunjukkan bahwa model berbasis fitur teks masih memiliki keterbatasan dalam menangkap konteks yang lebih kompleks [3] mengeksplorasi metode berbasis *machine learning* klasik seperti *Support Vector Machine* (SVM) dan *Random Forest*, yang menunjukkan performa cukup baik, tetapi masih kalah dibandingkan dengan metode berbasis *deep learning* [4].

Dalam perkembangan lebih lanjut, [5] memperkenalkan *Bidirectional Encoder Representations from Transformers* (BERT) sebagai model berbasis *Transformer* yang mampu menangkap konteks lebih baik dalam berbagai tugas NLP, termasuk klasifikasi teks. [6], [7], [8] memperkenalkan RoBERTa, yang merupakan penyempurnaan dari BERT dengan proses pre-training yang lebih ekstensif, dan terbukti meningkatkan akurasi dalam berbagai tugas NLP.

Dalam konteks bahasa Indonesia, penelitian mengenai deteksi berita palsu masih tergolong terbatas. [9] mengembangkan IndoBERT, sebuah model *Transformer* yang dioptimalkan untuk bahasa Indonesia, yang menunjukkan performa lebih unggul dibandingkan dengan BERT standar dalam berbagai tugas NLP. Studi lain [10], [11], [12] menguji efektivitas model berbasis *Transformer* seperti BERT dan RoBERTa dalam mendeteksi berita palsu berbahasa Indonesia, dengan hasil yang lebih akurat dibandingkan dengan pendekatan berbasis *machine learning* klasik.

Dalam beberapa tahun terakhir, berbagai metode telah dikembangkan untuk meningkatkan akurasi dalam deteksi berita palsu. [13], [14] mengevaluasi kombinasi metode *deep learning* dan fitur tambahan seperti metadata berita, termasuk sumber, tanggal publikasi, dan pola distribusi berita di media sosial. Hasil penelitian tersebut menunjukkan bahwa dengan mempertimbangkan informasi kontekstual selain isi teks, akurasi model dapat meningkat secara signifikan. Selain itu, [15] membandingkan metode berbasis *Transformer* dengan pendekatan berbasis graf, seperti *Graph Neural Networks* (GNN), untuk menangkap hubungan antara berbagai artikel berita dalam mendeteksi pola penyebaran berita palsu. Studi tersebut menemukan bahwa dengan mempertimbangkan

hubungan antarberita dan pola interaksi pengguna, model dapat lebih efektif dalam membedakan berita palsu dari berita asli.

Selain pendekatan berbasis *deep learning*, penelitian lain juga mengeksplorasi metode hybrid yang menggabungkan teknik *machine learning* dan NLP. [16], [17], [18] meneliti efektivitas pendekatan ensemble yang menggabungkan BERT dengan SVM untuk meningkatkan kinerja klasifikasi berita palsu. Pendekatan ini menggabungkan kekuatan representasi semantik dari *Transformer* dengan kemampuan klasifikasi berbasis margin dari SVM, menghasilkan peningkatan akurasi dibandingkan dengan metode tunggal. Sementara itu, [19], [20] mengembangkan model berbasis *Attention Mechanism* yang dikombinasikan dengan LSTM untuk menangkap informasi penting dalam teks berita. Hasil penelitian mereka menunjukkan bahwa dengan memanfaatkan mekanisme perhatian, model dapat lebih fokus pada kata-kata kunci yang lebih relevan dalam menentukan apakah suatu berita bersifat palsu atau tidak.

Meskipun berbagai metode telah dikembangkan, masih terdapat tantangan dalam mendeteksi berita palsu secara akurat, terutama untuk bahasa Indonesia. Oleh karena itu, penelitian ini bertujuan untuk membandingkan berbagai metode deteksi berita palsu, termasuk RoBERTa, BERT, IndoBERT, SVM, LSTM, dan CNN, guna menentukan model yang paling efektif untuk bahasa Indonesia. Dengan menguji berbagai pendekatan, penelitian ini diharapkan dapat memberikan wawasan yang lebih mendalam mengenai efektivitas masing-masing metode serta memberikan rekomendasi dalam pemilihan model yang optimal untuk sistem deteksi berita palsu dalam bahasa Indonesia. Keunikan dari penelitian ini terletak pada perbandingan langsung antara model berbasis *Transformer* yang lebih baru, seperti RoBERTa dan IndoBERT, dengan metode klasik seperti SVM, LSTM, dan CNN dalam konteks bahasa Indonesia, yang masih terbatas dalam literatur yang ada. Penelitian ini juga memberikan evaluasi komprehensif mengenai keunggulan dan keterbatasan setiap pendekatan dalam mendeteksi berita palsu berbahasa Indonesia.

Selain itu, penelitian ini juga mengeksplorasi keunggulan dan keterbatasan dari setiap metode. Model berbasis *Transformer* seperti RoBERTa, BERT, dan IndoBERT diharapkan mampu menangkap hubungan semantik yang lebih baik



dibandingkan dengan metode berbasis machine learning klasik seperti SVM. Sementara itu, LSTM dan CNN, yang merupakan bagian dari deep learning, juga memiliki keunggulan dalam menangani urutan teks tetapi mungkin tidak seefisien Transformer dalam memahami konteks yang kompleks. Oleh karena itu, dengan melakukan evaluasi menyeluruh terhadap keenam metode ini, penelitian ini berupaya memberikan kontribusi bagi pengembangan sistem deteksi berita palsu yang lebih akurat dan andal.

Hasil penelitian ini diharapkan dapat menjadi referensi dalam pengembangan sistem berbasis kecerdasan buatan untuk menangani berita palsu di Indonesia, sekaligus membantu regulator, jurnalis, dan masyarakat dalam mengidentifikasi serta menangkal penyebaran informasi yang tidak valid di berbagai platform digital.

2. TINJAUAN PUSTAKA

Penelitian sebelumnya mengenai penerapan metode *Long Short-Term Memory* (LSTM) dalam sistem pendeteksi berita palsu berbahasa Indonesia menunjukkan bahwa perkembangan teknologi informasi telah memungkinkan masyarakat untuk mengakses berita secara luas, namun juga meningkatkan risiko penyebaran berita palsu (hoaks). Hoaks sering kali disajikan dengan judul provokatif dan bahasa yang meyakinkan, sehingga mudah menyebar akibat kurangnya kesadaran masyarakat dalam menyaring informasi. Untuk mengatasi masalah ini, penelitian tersebut mengembangkan model pendeteksi berita palsu menggunakan LSTM dengan IndoBERT, yang dilatih pada dataset berisi 3.309 berita hoaks dan non-hoaks. Hasil evaluasi menunjukkan bahwa model ini mampu mencapai tingkat akurasi 94,2%, dengan *precision* 0,944, *recall* 0,940, dan *F1-score* 0,941, yang lebih unggul dibandingkan metode pembandingan. Temuan ini mengindikasikan bahwa penggunaan LSTM dan IndoBERT dapat menjadi solusi efektif dalam meningkatkan akurasi sistem pendeteksi berita palsu berbahasa Indonesia [21].

Selain itu, penelitian mengenai implementasi algoritma *Naïve Bayes Classifier* dalam mendeteksi berita palsu di media sosial menunjukkan bahwa penyebaran hoaks di internet telah menjadi masalah global yang dapat mengganggu stabilitas sosial, politik, dan ekonomi. Berdasarkan survei Masyarakat Telematika Indonesia, 44,3%

responden menerima berita palsu setiap hari, sementara laporan dari Kominfo hingga 11 Agustus 2021 mencatat 1.848 hoaks terkait pandemi Covid-19 dan 290 hoaks terkait vaksin Covid-19. Penelitian ini menggunakan *Naïve Bayes Classifier*, metode klasifikasi berbasis Teorema Bayes, dalam mendeteksi berita palsu menggunakan model CRISP-DM (Cross-Industry Standard Process for Data Mining). Data pelatihan yang digunakan berasal dari situs Kumparan, dengan 300 sampel data. Proses analisis menggunakan pustaka NLP Python, yaitu "satro", serta evaluasi model menggunakan *confusion matrix*. Pada tahap penerapan, model diunggah ke Heroku, memungkinkan pengguna untuk melakukan prediksi berita secara langsung melalui antarmuka pengguna. Hasil penelitian ini menunjukkan bahwa metode *Naïve Bayes Classifier* dapat menjadi alat yang efektif dalam mendeteksi berita palsu secara otomatis dan real-time [22].

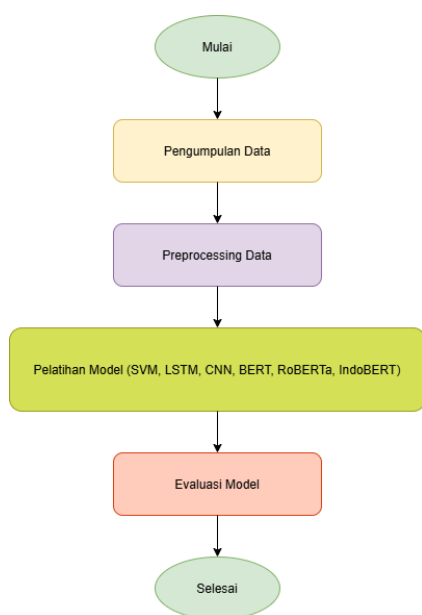
Penelitian terkait deteksi berita hoaks menunjukkan bahwa berita hoaks merupakan informasi palsu yang dapat memicu provokasi dan kebencian di masyarakat. Dengan meningkatnya akses internet, penyebaran hoaks menjadi semakin masif, sehingga diperlukan metode deteksi yang lebih akurat. Salah satu pendekatan yang digunakan adalah *deep learning* dengan mengintegrasikan *text mining* untuk menemukan pola berita hoaks. Dataset yang digunakan berasal dari Kaggle dengan 2.700 data yang telah melalui proses *text Preprocessing* agar lebih terstruktur. Kemudian, dilakukan *feature engineering* menggunakan BERT, sehingga data dapat diproses menggunakan *machine learning* dengan tiga metode klasifikasi, yaitu BERT, SVM, dan Random Forest. Hasil penelitian menunjukkan bahwa BERT menghasilkan kinerja terbaik dengan akurasi 99% dan ROC-AUC 0.99, dibandingkan dengan model *machine learning* tradisional seperti SVM dan Random Forest [23].

3. METODOLOGI PENELITIAN

3.1. Tahapan Penelitian

Penelitian ini mengadaptasi pendekatan multi-model, di mana beberapa teknik seperti *Support Vector Machine* (SVM), *Long Short-Term Memory* (LSTM), *Convolutional Neural Network* (CNN), BERT, RoBERTa, dan IndoBERT digunakan untuk mendeteksi berita palsu. Setiap model diuji

dan dibandingkan untuk mengidentifikasi pendekatan yang paling optimal. Proses penelitian ini meliputi beberapa tahapan utama, seperti pengumpulan dataset, *Preprocessing* data, ekstraksi fitur, pelatihan model, serta evaluasi menggunakan metrik seperti akurasi, *precision*, *recall*, dan *F1-score*. Alur penelitian secara keseluruhan ditunjukkan pada Gambar 1, yang mengilustrasikan setiap langkah utama dalam penelitian ini.



Gambar 1. Tahap Alur Penelitian

Alur metode penelitian ini dimulai dengan pengumpulan dataset berita palsu dan berita asli yang akan digunakan untuk pelatihan dan pengujian model. Setelah itu, dilakukan *Preprocessing* data, termasuk pembersihan teks, tokenisasi, dan konversi ke bentuk yang dapat diproses oleh model. Selanjutnya, data yang telah diproses dikonversi ke representasi numerik menggunakan teknik seperti TF-IDF untuk SVM atau *embedding* untuk model berbasis *deep learning*. Setelah representasi fitur dibuat, dilakukan pembagian dataset menjadi data latih dan data uji dengan proporsi tertentu. Tahap berikutnya adalah pelatihan model, di mana beberapa model seperti RoBERTa, BERT, IndoBERT, SVM, LSTM, dan CNN digunakan untuk mendeteksi berita palsu berdasarkan teks. Setelah

model selesai dilatih, dilakukan evaluasi kinerja menggunakan metrik seperti akurasi, *precision*, *recall*, dan *F1-score* untuk menilai efektivitas masing-masing metode. Akhirnya, hasil evaluasi dibandingkan untuk mengetahui model terbaik dalam mendeteksi berita palsu, yang kemudian digunakan untuk menarik kesimpulan dan memberikan rekomendasi untuk penelitian selanjutnya.

3.2. Pengumpulan Data

Dalam penelitian ini, dataset dikumpulkan dari dua sumber utama, yaitu TurnbackHoax dan Detik News [24], [25]. TurnbackHoax adalah platform yang dikelola oleh Masyarakat Anti Fitnah Indonesia (MAFINDO) yang menyediakan data berita hoaks yang telah diverifikasi. Sementara itu, Detik News adalah portal berita daring yang menyajikan berita aktual dan kredibel. Dengan menggabungkan kedua sumber ini, dataset yang diperoleh mencakup berita palsu dan berita faktual, sehingga dapat digunakan untuk melatih model deteksi berita palsu secara lebih akurat.

Setelah data dikumpulkan, proses *Preprocessing* dilakukan untuk membersihkan teks dari karakter khusus, angka, tanda baca, dan kata-kata yang tidak relevan. Selain itu, teks berita diubah menjadi bentuk standar dengan melakukan lowercasing serta menghapus stopwords menggunakan daftar stopwords bahasa Indonesia. Setiap berita kemudian diberi label, di mana berita dari TurnbackHoax dilabeli sebagai hoaks (1) dan berita dari Detik News dilabeli sebagai fakta (0).

Selanjutnya, dilakukan filtering untuk memastikan hanya berita dengan panjang teks yang memadai (lebih dari 50 kata) yang digunakan dalam penelitian. Hal ini bertujuan untuk menghindari teks yang terlalu pendek, yang dapat mengganggu kualitas klasifikasi. Setelah proses filtering selesai, dataset diklasifikasikan ke dalam dua kategori utama, yaitu berita palsu dan berita faktual, yang kemudian digunakan dalam proses pelatihan model.

Berikut adalah ringkasan jumlah data yang digunakan dalam penelitian ini di jelaskan pada Tabel 1:



Tabel 1. Data Yang Di Kumpulkan

Sumber Data	Kategori	Jumlah Data
TurnbackHoax	Berita Palsu	4304
Detik News	Berita Faktual	2160
Total		6464

3.3. Preprocessing Data

Preprocessing data merupakan tahap penting dalam pemrosesan bahasa alami (*Natural Language Processing, NLP*) untuk meningkatkan kualitas data sebelum digunakan dalam pelatihan model. Pada penelitian ini, *Preprocessing* dilakukan melalui beberapa tahapan utama, yaitu pembersihan teks (*cleaning*), normalisasi, tokenisasi, stopword removal, dan transformasi teks ke dalam representasi numerik menggunakan TF-IDF dan word embeddings.

1. Pembersihan Teks

Tahapan ini bertujuan untuk menghapus karakter yang tidak relevan, seperti angka, tanda baca, dan karakter khusus, serta mengonversi semua teks menjadi huruf kecil (*lowercasing*). Hal ini dilakukan untuk memastikan bahwa perbedaan huruf besar dan kecil tidak mempengaruhi hasil klasifikasi.

2. Normalisasi dan Tokenisasi

Setelah pembersihan teks, dilakukan normalisasi kata dengan mengubah istilah tidak baku menjadi bentuk standar (misalnya "gak" menjadi "tidak"). Kemudian, teks dibagi menjadi token menggunakan pendekatan tokenisasi berbasis kata (*word tokenization*).

3. Stopword Removal

Stopwords adalah kata-kata yang sering muncul tetapi tidak memiliki makna penting dalam analisis teks, seperti "yang", "dan", "di", "ke", dan sebagainya. Penghapusan stopwords dilakukan menggunakan daftar stopwords Bahasa Indonesia untuk mengurangi dimensi data dan meningkatkan relevansi fitur.

4. Representasi Teks (TF-IDF dan *Word Embeddings*)

Setelah teks diproses, dilakukan transformasi teks menjadi bentuk numerik agar dapat digunakan oleh model *machine learning* dan *deep learning*. Dua pendekatan utama yang digunakan dalam penelitian ini adalah TF-IDF (Term Frequency-Inverse Document Frequency) dan word embeddings (Word2Vec atau BERT embeddings).

- TF-IDF dihitung menggunakan rumus berikut:

$$TF - IDF(t, d) = TF(t, d) \times IDF(t) \quad (1)$$

$$IDF(t) = \log \frac{N}{df(t)}$$

Di mana:

- $TF(t, d)$ adalah frekuensi kemunculan kata t dalam dokumen d .
- $IDF(t)$ adalah inverse document frequency, yang dihitung berdasarkan jumlah dokumen N dibagi dengan jumlah dokumen yang mengandung kata t (document frequency $df(t)$).
- Word Embeddings, seperti yang digunakan dalam BERT atau IndoBERT, mengubah kata-kata menjadi vektor berdimensi tinggi yang menangkap makna semantik dari teks.

Berikut Tabel 2 adalah contoh hasil *Preprocessing* data:

Tabel 2. Contoh Hasil *Preprocessing* Data

Teks Awal	Teks Setelah <i>Preprocessing</i>
"Pemerintah mengumumkan kebijakan baru untuk ekonomi!"	"pemerintah umum kebijakan baru ekonomi"
"4erit aini benar atau hanya hoaks?"	"berita benar hoaks"
"Presiden akan meakukan konferensi pers hari ini"	"presiden konferensi pers hari"

Tahap *Preprocessing* ini memastikan bahwa model yang digunakan dapat belajar dari data dengan lebih efektif, mengurangi noise, dan meningkatkan performa dalam mendeteksi berita palsu.

3.4. Pelatihan Model (SVM, LSTM, CNN, BERT, RoBERTa, IndoBERT)

Pelatihan model merupakan tahapan utama dalam penelitian ini, di mana enam metode klasifikasi—Support Vector Machine (SVM), Long Short-Term Memory (LSTM), Convolutional Neural Network (CNN), Bidirectional Encoder



Representations from *Transformers* (BERT), Robustly Optimized BERT Pretraining Approach (RoBERTa), dan IndoBERT—digunakan untuk mendeteksi berita palsu dalam bahasa Indonesia.

1. Support Vector Machine (SVM)

SVM adalah algoritma pembelajaran mesin berbasis supervised learning yang bekerja dengan menemukan hyperplane terbaik untuk memisahkan dua kelas data. Dalam penelitian ini, teks berita diubah menjadi representasi TF-IDF sebelum digunakan sebagai input untuk model SVM dengan kernel linear [26]. Pelatihan dilakukan selama 5 epoch, dan 20% dari data digunakan sebagai data uji. Persamaan dasar SVM dapat dituliskan sebagai:

$$f(x) = w^T x + b \quad (2)$$

Di mana w adalah bobot, x adalah input, dan b adalah bias.

2. Long Short-Term Memory (LSTM)

LSTM merupakan jenis jaringan saraf tiruan berbasis Recurrent Neural Network (RNN) yang mampu menangkap hubungan jangka panjang dalam teks. Model ini digunakan untuk memproses urutan kata dalam berita dan mempertahankan konteks penting. Pelatihan dilakukan selama 5 epoch dengan test size 0.2, menggunakan optimizer Adam dan loss function categorical crossentropy.

3. Convolutional Neural Network (CNN)

CNN sering digunakan dalam pemrosesan gambar, tetapi juga dapat diterapkan dalam klasifikasi teks. Model ini menggunakan lapisan konvolusi untuk mengekstrak fitur penting dari teks sebelum dikirim ke lapisan fully connected untuk klasifikasi. Filter konvolusi digunakan untuk menangkap pola dalam urutan kata, yang dihitung sebagai:

$$h_i = f(W \cdot x_{i:i+h-1} + b) \quad (3)$$

Dimana W adalah bobot filter, x adalah input teks yang telah diubah menjadi vektor, dan f adalah fungsi aktivasi.

Model CNN dilatih dengan 5 epoch dan menggunakan test size 0.2. Lapisan dropout

juga diterapkan untuk menghindari overfitting.

4. Bidirectional Encoder Representations from Transformers (BERT)

BERT adalah model berbasis *Transformer* yang mampu menangkap konteks teks secara dua arah (bidirectional). Model ini menggunakan teknik Masked Language Model (MLM) untuk memahami hubungan kata dalam kalimat. Model BERT dilatih selama 5 epoch, menggunakan batch size 8. Representasi teks dalam BERT dihitung menggunakan self-attention mechanism, dengan skor perhatian dihitung sebagai:

$$Attention(Q, K, V) = \text{softmax}\left(\frac{QK^T}{\sqrt{d_k}}\right)V \quad (4)$$

Dimana Q, K , dan V adalah matriks kueri, kunci, dan nilai, serta d_k adalah dimensi vektor [27].

5. Robustly Optimized BERT Pretraining Approach (RoBERTa)

RoBERTa adalah pengembangan dari BERT yang menggunakan lebih banyak data dalam pretraining dan tidak menggunakan Next Sentence Prediction (NSP). Model ini juga dilatih selama 5 epoch dengan test size 0.2, mengikuti skema pretraining yang lebih luas.

6. IndoBERT

IndoBERT adalah varian BERT yang dilatih secara khusus untuk bahasa Indonesia. Model ini menggunakan dataset teks dalam bahasa Indonesia, sehingga lebih optimal dibandingkan BERT standar dalam tugas NLP berbasis bahasa Indonesia. Model dilatih selama 5 epoch dengan test size 0.2, menggunakan hyperparameter yang sama dengan BERT.

4. HASIL DAN PEMBAHASAN

Pada tahap ini, dilakukan evaluasi terhadap keenam model yang telah dilatih, yaitu SVM, LSTM, CNN, BERT, RoBERTa, dan IndoBERT. Evaluasi dilakukan menggunakan metrik *accuracy*, *precision*, *recall*, dan *F1-score*, serta analisis menggunakan *confusion matrix* untuk melihat jumlah prediksi yang benar dan salah.

4.1. Evaluasi Performa Model

Berdasarkan hasil pengujian, performa masing-masing model dapat dirangkum pada Tabel 3 sebagai berikut:

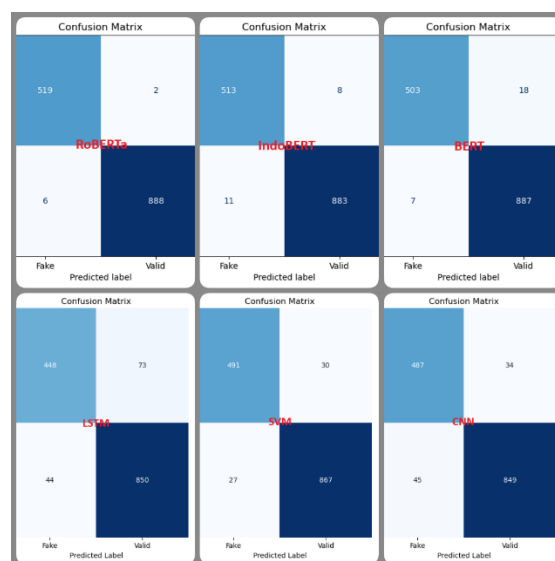
Tabel 3. Hasil Model Yang Di Uji

Model	Accuracy (%)	Precision	Recall	F1-score
RoBERTa	99.5	0.99	0.99	0.99
IndoBERT	98.6	0.99	0.98	0.98
BERT	98.2	0.98	0.99	0.98
SVM	95.9	0.94	0.96	0.96
CNN	93.9	0.92	0.95	0.93
LSTM	92.3	0.96	0.94	0.95

Dari hasil tersebut, dapat dilihat bahwa model RoBERTa mencatatkan performa terbaik dengan akurasi tertinggi sebesar 99.5%, diikuti oleh IndoBERT (98.6%) dan BERT (98.2%). Model berbasis *Transformer* ini menunjukkan keunggulan dalam menangkap konteks teks yang lebih kompleks dibandingkan model berbasis *machine learning* klasik maupun *deep learning* konvensional. Model SVM (95.9%), meskipun masih memiliki akurasi tinggi, menunjukkan performa yang lebih rendah dibandingkan model *Transformer* karena keterbatasannya dalam memahami konteks kalimat yang lebih kompleks. CNN (93.9%) dan LSTM (92.3%) memiliki akurasi yang lebih rendah, meskipun tetap mampu melakukan klasifikasi dengan cukup baik. Dari sisi *precision*, *recall*, dan *F1-score*, model RoBERTa dan IndoBERT tetap unggul dengan skor hampir mendekati 1.00, menunjukkan bahwa model ini memiliki tingkat kesalahan yang sangat rendah dalam mendeteksi berita palsu dan valid.

4.2. Analisis Confusion Matrix

Untuk memahami lebih dalam kesalahan prediksi yang terjadi, dilakukan analisis *Confusion Matrix* pada model terbaik dapat dilihat pada Gambar 2.



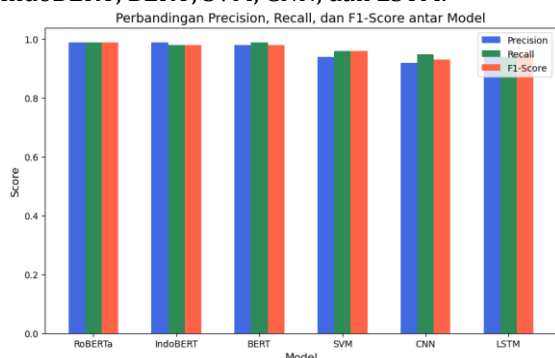
Gambar 2. Confusion Matrix Semua Model

Analisis *Confusion Matrix* dilakukan untuk mengevaluasi performa enam model yang digunakan, yaitu CNN, SVM, LSTM, BERT, IndoBERT, dan RoBERTa. Berdasarkan hasil *Confusion Matrix* yang diperoleh, model BERT dan IndoBERT menunjukkan performa terbaik dengan tingkat akurasi dan presisi yang tinggi, ditandai dengan jumlah True Positives (TP) yang lebih besar dan False Negatives (FN) yang lebih kecil dibandingkan model lainnya. Sementara itu, CNN dan LSTM memiliki kinerja yang cukup baik, tetapi masih terdapat jumlah False Positives (FP) yang lebih tinggi, yang menunjukkan adanya kesalahan dalam mengklasifikasikan data. Model SVM memiliki performa terendah dibandingkan yang lain, dengan jumlah kesalahan klasifikasi yang lebih banyak, yang mengindikasikan bahwa pendekatan berbasis Support Vector Machine kurang efektif dalam menangani dataset yang digunakan. Sedangkan RoBERTa menunjukkan hasil yang kompetitif dengan BERT dan IndoBERT, namun masih terdapat sedikit ketidakseimbangan dalam klasifikasinya. Secara keseluruhan, hasil *Confusion Matrix* mengindikasikan bahwa model berbasis *Transformer* seperti BERT, IndoBERT, dan RoBERTa lebih unggul dibandingkan model berbasis CNN, LSTM, dan SVM dalam menangani tugas klasifikasi pada dataset ini.

4.3. Grafik Perbandingan Precision, Recall, dan F1-score Antar Model

Grafik yang ditampilkan Pada Gambar 3 menggambarkan perbandingan tiga metrik

evaluasi utama, yaitu *Precision*, *Recall*, dan *F1-score*, untuk enam model yang diuji: RoBERTa, IndoBERT, BERT, SVM, CNN, dan LSTM.



Gambar 3. Hasil Grafik Perbandingan

- 1) Grafik Pada Gambar 3 menunjukkan perbandingan metrik evaluasi *Precision*, *Recall*, dan *F1-score* dari berbagai model NLP, yaitu RoBERTa, IndoBERT, BERT, SVM, CNN, dan LSTM. RoBERTa memiliki performa terbaik dengan *Precision*, *Recall*, dan *F1-score* masing-masing sebesar 0.99, menunjukkan keseimbangan antara prediksi positif yang akurat dan cakupan deteksi yang tinggi. IndoBERT dan BERT juga menunjukkan kinerja yang kuat dengan skor yang hampir mendekati RoBERTa. Sementara itu, SVM, CNN, dan LSTM memiliki performa yang lebih rendah, dengan CNN mencatat *Precision* terendah (0.92) dan LSTM memiliki *Recall* lebih rendah (0.94), meskipun tetap memiliki keseimbangan yang cukup baik dalam *F1-score*. Secara keseluruhan, model berbasis *Transformer* (RoBERTa, IndoBERT, dan BERT) unggul dibandingkan metode klasik seperti SVM dan model berbasis jaringan saraf konvolusi (CNN) atau rekuren (LSTM).

5. KESIMPULAN DAN SARAN

Berdasarkan analisis *confusion matrix* dan metrik evaluasi, RoBERTa menunjukkan performa terbaik dengan akurasi 99.5% serta *Precision*, *Recall*, dan *F1-score* yang seimbang di angka 0.99. Model berbasis *Transformer* seperti IndoBERT dan BERT juga memberikan hasil yang kompetitif, sedangkan metode klasik seperti SVM dan model berbasis jaringan saraf seperti CNN dan LSTM menunjukkan performa yang lebih rendah. Dari

hasil ini, disarankan untuk menggunakan RoBERTa atau IndoBERT dalam aplikasi yang membutuhkan prediksi dengan akurasi tinggi. Sebagai rekomendasi, penelitian lanjutan dapat mengeksplorasi teknik *fine-tuning* lebih lanjut pada model *Transformer* untuk meningkatkan kinerja pada dataset yang lebih kompleks. Selain itu, *future work* dapat mencakup penggabungan teknik *ensemble learning* atau *hybrid models* untuk mengoptimalkan keakuratan model serta menerapkan metode *explainable AI* (XAI) untuk meningkatkan interpretabilitas prediksi, terutama dalam aplikasi yang berkaitan dengan pengambilan keputusan kritis.

DAFTAR PUSTAKA:

- [1] M. Guderlei and M. Aßenmacher, "Evaluating Unsupervised Representation Learning for Detecting Stances of Fake News," *Proceedings of the 28th International Conference on Computational Linguistics*, pp. 6339–6349, 2020, [Online]. Available: <https://github.com/magud/fake-news-detection>
- [2] P. Bahad, P. Saxena, and R. Kamal, "Fake News Detection using Bi-directional LSTM-Recurrent Neural Network," in *Procedia Computer Science*, Elsevier B.V., 2019, pp. 74–82. doi: 10.1016/j.procs.2020.01.072.
- [3] R. Jehad and S. Yousif, "Classification of fake news using multi-layer perceptron," in *AIP Conference Proceedings*, American Institute of Physics Inc., Mar. 2021. doi: 10.1063/5.0042264.
- [4] P. Dedeepya, M. Yarrarapu, P. P. Kumar, S. K. Kaushik, P. N. Raghavendra, and P. Chandu, "Fake News Detection on Social Media Through a Hybrid SVM-KNN Approach Leveraging Social Capital Variables," in *2024 3rd International Conference on Applied Artificial Intelligence and Computing (ICAAIC)*, 2024, pp. 1168–1175. doi: 10.1109/ICAAIC60222.2024.10575681.
- [5] Zhang Tong et al., "BDANN: BERT-Based Domain Adaptation Neural Network for Multi-Modal Fake News Detection," *2020 International Joint Conference on Neural Networks (IJCNN)*, 2020.



- [6] A. Pritzkau, "Multi-class fake news detection of news articles and domain identification with RoBERTa-a baseline model," *CEUR Workshop Proc*, Sep. 2021, [Online]. Available: <https://www.fkie.fraunhofer.de/>
- [7] N. Lin, S. Fu, and S. Jiang, "Fake News Detection in the Urdu Language using CharCNN-RoBERTa," *CEUR Workshop Proc*, vol. 2826, pp. 447–451, 2020, [Online]. Available: <https://huggingface.co/urduhack/urdu-roberta>
- [8] E. Raja, B. Soni, and S. K. Borgohain, "Fake News Detection in Malayalam using Optimized XLM-RoBERTa Model," in *DravidianLangTech 2023 - 3rd Workshop on Speech and Language Technologies for Dravidian Languages, associated with 14th International Conference on Recent Advances in Natural Language Processing, RANLP 2023 - Proceedings*, Incoma Ltd, 2023, pp. 186–191. doi: 10.26615/978-954-452-085-4_026.
- [9] S. M. Isa, G. Nico, and M. Permana, "IndoBERT for Indonesian Fake News Detection," *ICIC Express Letters*, vol. 16, no. 3, pp. 289–297, Mar. 2022, doi: 10.24507/icicel.16.03.289.
- [10] L. B. Angizeh and M. R. Keyvanpour, "Detecting Fake News using Advanced Language Models: BERT and RoBERTa," in *2024 10th International Conference on Web Research (ICWR)*, 2024, pp. 46–52. doi: 10.1109/ICWR61162.2024.10533352.
- [11] S. F. N. Azizah, H. D. Cahyono, S. W. Sihwi, and W. Widiarto, "Performance Analysis of Transformer Based Models (BERT, ALBERT, and RoBERTa) in Fake News Detection," in *2023 6th International Conference on Information and Communications Technology (ICOIAC)*, 2023, pp. 425–430. doi: 10.1109/ICOIAC59844.2023.10455849.
- [12] V. Kolev, G. Weiss, and G. Spanakis, "FOREAL: RoBERTa Model for Fake News Detection based on Emotions," in *International Conference on Agents and Artificial Intelligence*, Science and Technology Publications, Lda, 2022, pp. 429–440. doi: 10.5220/0010873900003116.
- [13] M. Umer, Z. Imtiaz, S. Ullah, A. Mehmood, G. S. Choi, and B. W. On, "Fake news stance detection using deep learning architecture (CNN-LSTM)," *IEEE Access*, vol. 8, pp. 156695–156706, 2020, doi: 10.1109/ACCESS.2020.3019735.
- [14] R. Jehad and S. A.Yousif, "Fake News Classification Using Random Forest and Decision Tree (J48)," *Al-Nahrain Journal of Science*, vol. 23, no. 4, pp. 49–55, Dec. 2020, doi: 10.22401/ANJS.23.4.09.
- [15] H. T. Phan, N. T. Nguyen, and D. Hwang, "Fake news detection: A survey of graph neural network methods," May 01, 2023, *Elsevier Ltd*. doi: 10.1016/j.asoc.2023.110235.
- [16] J. Alghamdi, Y. Lin, and S. Luo, "A Comparative Study of Machine learning and Deep learning Techniques for Fake News Detection," *Information (Switzerland)*, vol. 13, no. 12, Dec. 2022, doi: 10.3390/info13120576.
- [17] D. G. Dev and V. Bhatnagar, "Hybrid RFSVM: Hybridization of SVM and Random Forest Models for Detection of Fake News," *Algorithms*, vol. 17, no. 10, Oct. 2024, doi: 10.3390/a17100459.
- [18] P. Jain, S. Sharma, Monica, and P. K. Aggarwal, "Classifying Fake News Detection Using SVM, Naive Bayes and LSTM," in *2022 12th International Conference on Cloud Computing, Data Science & Engineering (Confluence)*, 2022, pp. 460–464. doi: 10.1109/Confluence52989.2022.9734129.
- [19] R. Mohawesh, H. Bany Salameh, Y. Jararweh, M. Alkhalaileh, and S. Maqsood, "Fake review detection using Transformer-based enhanced LSTM and RoBERTa," *International Journal of Cognitive Computing in Engineering*, vol. 5, pp. 250–258, Jan. 2024, doi: 10.1016/j.ijcce.2024.06.001.
- [20] S. Ghazi Arkaan, A. Rialdy Atmadja, and M. D. Firdaus, "Fake News Detection in the 2024 Indonesian General Election Using Bidirectional Long Short-Term Memory (BI-LSTM) Algorithm," *Jurnal Ilmiah Ilmu Komputer dan Matematika*, vol. 21, no. 2, pp. 693–7554, 2024, doi: 10.33751/komputasi.v21i2.5260.



- [21] M. Rifky Maulana *et al.*, "PENERAPAN METODE LSTM DALAM PEMBUATAN SISTEM PENDETEKSI BERITA PALSU BERBAHASA INDONESIA," 2023. [Online]. Available: <https://drive.google.com/drive/folders/1eYKx7agpzgvXGCpyE1E57HxDhonk>
- [22] N. Agustina, A. Adrian, and M. Hermawati, "Implementasi Algoritma Naïve Bayes Classifier untuk Mendeteksi Berita Palsu pada Sosial Media," *Faktor Exacta*, vol. 14, no. 4, p. 206, Jan. 2022, doi: 10.30998/faktorexacta.v14i4.11259.
- [23] A. Ripa'i, F. Santoso, and F. Lazim, "Deteksi Berita Hoax dengan Perbandingan Website Menggunakan Pendekatan *Deep learning* Algoritma BERT," *G-Tech: Jurnal Teknologi Terapan*, vol. 8, no. 3, pp. 1749–1758, Jul. 2024, doi: 10.33379/gtech.v8i3.4541.
- [24] A. A. Kurniawan, M. Mustikasari, and P. Korespondensi, "EVALUASI KINERJA MLLIB APACHE SPARK PADA KLASIFIKASI BERITA PALSU DALAM BAHASA INDONESIA," *Jurnal Teknologi Informasi dan Ilmu Komputer (JTIK)*, vol. 9, no. 3, Jun. 2022, doi: 10.25126/jtiik.202293538.
- [25] R. N. Rahayu and & Sensusiyati, "ANALISIS BERITA HOAX COVID-19 DI MEDIA SOSIAL DI INDONESIA," *INTELEKTIVA : JURNALEKONOMI, SOSIAL & HUMANIORA*, vol. 01, Apr. 2020, [Online]. Available: <https://www.suara.com/news>,
- [26] C. F. Navarro and C. A. Perez, "Color-texture pattern classification using global-local feature extraction, an SVM classifier, with bagging ensemble post-processing," *Applied Sciences (Switzerland)*, vol. 9, no. 15, Aug. 2019, doi: 10.3390/app9153130.
- [27] M. Szczepański, M. Pawlicki, R. Kozik, and M. Choraś, "New explainability method for BERT-based model in fake news detection," *Sci Rep*, vol. 11, no. 1, Dec. 2021, doi: 10.1038/s41598-021-03100-6.