

Implementasi Model BERT untuk Klasifikasi Konten Hoaks pada Berita Politik

DOI: <http://dx.doi.org/10.35889/jutisi.v15i1.3418>

Creative Commons License 4.0 (CC BY – NC)

Galvan Adam Maliki^{1*}, Christian Sri Kusuma Aditya²

Informatika, Universitas Muhammadiyah Malang, Indonesia

*e-mail Corresponding Author. galpanpanprg@webmail.umm.ac.id

Abstract

The rapid growth of digital media accelerates the spread of information, including political hoaxes that can influence public opinion and disrupt social stability. Distinguishing factual news from hoaxes is challenging due to linguistic complexity and the high volume of textual data that cannot be handled manually. To address this issue, the present study employs an automated approach based on the BERT (Bidirectional Encoder Representations from Transformers) model, which is designed to capture linguistic context in a more comprehensive manner. The dataset consists of 3,807 Indonesian political news articles verified as factual or hoaxes. The results show that BERT achieves 98% accuracy, performing strongly on factual news and reasonably well on hoax detection. These findings confirm that BERT is effective for political hoax classification and serves as a relevant foundation for developing future NLP-based misinformation detection systems.

Keywords: BERT; Fake news; Political news; Text classification

Abstrak

Perkembangan media digital mempercepat penyebaran informasi, termasuk berita hoaks yang berpotensi memengaruhi opini publik dan stabilitas sosial. Tantangan utama dalam memilah fakta dan hoaks terletak pada kompleksitas bahasa serta tingginya volume teks yang sulit ditangani secara manual. Untuk mengatasi permasalahan tersebut, penelitian ini mengadopsi metode otomatis berbasis model BERT (*Bidirectional Encoder Representations from Transformers*) yang dirancang untuk memahami konteks bahasa secara lebih mendalam. Karena keunggulannya dalam memproses hubungan antarkata secara *bidirectional* sehingga mampu menangkap pola linguistik yang tidak dapat ditangani metode tradisional. Dataset yang digunakan mencakup 3.807 berita politik berbahasa Indonesia yang telah diverifikasi sebagai fakta atau hoaks. Hasil penelitian menunjukkan bahwa BERT mencapai akurasi 98% dengan performa sangat baik pada kelas fakta dan cukup baik pada kelas hoaks. Temuan ini menegaskan bahwa BERT efektif digunakan untuk klasifikasi berita hoaks politik dan relevan sebagai dasar pengembangan sistem deteksi misinformasi berbasis NLP di masa mendatang.

Kata kunci: BERT; Berita hoaks; Berita politik; Klasifikasi teks

1. Pendahuluan

Penyebaran informasi melalui media sosial dan situs berita yang semakin cepat turut memicu meningkatnya jumlah berita hoaks, yang dapat membingungkan dan menyesatkan masyarakat. Perkembangan pesat platform media sosial menyebabkan berita hoaks menyebar dengan lebih cepat dan lebih sering menarik perhatian dibandingkan berita yang valid [1]. Kondisi ini memberikan dampak signifikan pada opini publik, membentuk persepsi yang keliru, dan bahkan berpotensi menimbulkan konflik sosial. Dalam konteks politik, berita hoaks dapat dimanfaatkan untuk memanipulasi pemilih dan merusak kepercayaan terhadap institusi [2]. Tantangan utama adalah bahwa berita hoaks sering kali menyerupai berita valid, sehingga sulit untuk mengklasifikasikannya dengan metode tradisional seperti pemeriksaan fakta manual [3]. Untuk mengatasi masalah ini, diperlukan pendekatan yang lebih canggih, salah satunya dengan memanfaatkan teknologi pemrosesan bahasa alami seperti BERT (*Bidirectional Encoder*

Representations from Transformers) untuk mengklasifikasikan berita hoaks secara otomatis dan berbasis teks.

Untuk menghadapi masalah ini, pendekatan menggunakan BERT (*Bidirectional Encoder Representations from Transformers*) diusulkan karena kemampuannya dalam memahami konteks kata secara lebih komprehensif melalui teknik pemrosesan bahasa alami yang mutakhir. Berbeda dengan metode tradisional seperti RNN, BERT menggunakan arsitektur transformer yang memproses teks dalam dua arah (*bidirectional*), sehingga mampu menangkap hubungan antar kata dengan lebih akurat. Pada klasifikasi berita hoaks, BERT mengubah kata-kata menjadi vektor numerik yang mempertimbangkan makna dalam keseluruhan kalimat, bukan hanya berdasarkan urutan, sehingga dapat meningkatkan akurasi dalam klasifikasi berita hoaks. Penelitian lainnya, seperti yang dilakukan oleh Minaee et al. (2020), mengkaji penerapan BERT dalam deteksi berita hoaks dengan memanfaatkan representasi kontekstual dua arah serta *deep learning* untuk memperoleh akurasi klasifikasi yang lebih tinggi. Penelitian ini menggarisbawahi potensi BERT dalam mengatasi tantangan dalam deteksi berita hoaks berbasis teks dan memperkuat penerapan teknik pemrosesan bahasa alami dalam memahami konteks kalimat secara lebih menyeluruh [4]. Implementasi ini diharapkan mampu memberikan solusi yang lebih efektif dalam klasifikasi berita hoaks berbasis teks. Selain BERT, pendekatan lain yang menggunakan CNN-LSTM juga telah terbukti efektif dalam mendeteksi berita hoaks. Umer et al. (2020) mengusulkan model CNN-LSTM untuk mendeteksi *stance* (sikap) terhadap berita hoaks, yang memungkinkan model untuk memahami hubungan antara konten teks dan konteks sosial di media sosial, sehingga memperbaiki akurasi klasifikasi berita hoaks [5].

Kaliyar et al. dalam penelitian mereka meneliti sejumlah pendekatan klasifikasi berita hoaks di media sosial dengan memasukkan model BERT sebagai salah satu metode yang digunakan. Studi ini mengindikasikan bahwa pemanfaatan model BERT memberikan peningkatan performa dibandingkan metode konvensional dalam mendeteksi berita hoaks, khususnya yang berbasis fitur statistik. Penelitian ini mengungkapkan bahwa pola linguistik berita hoaks cenderung berbeda dari berita valid, yang dapat dimanfaatkan untuk meningkatkan akurasi klasifikasi dengan memanfaatkan model BERT yang memproses konteks secara *bidirectional*. Temuan ini menunjukkan adanya kebutuhan akan penggunaan teknik yang lebih canggih, seperti BERT, yang dapat memanfaatkan pemahaman kontekstual dua arah untuk meningkatkan kemampuan klasifikasi berita hoaks [6].

Berdasarkan temuan sebelumnya, penelitian ini mengusulkan pemanfaatan BERT untuk menangani teks secara lebih optimal dan memahami konteks dengan lebih baik melalui pendekatan pemrosesan *bidirectional*. Dengan arsitektur transformer yang memungkinkan pemodelan informasi dari keseluruhan kalimat, BERT dapat mengidentifikasi makna dan pola pada berita hoaks dengan lebih akurat. Peningkatan kinerja model diharapkan dapat tercapai dengan menggunakan kumpulan data yang lebih berkualitas dan beragam untuk memperkaya representasi teks. Dalam studi ini, BERT akan diterapkan untuk klasifikasi berita hoaks guna meningkatkan efektivitas dan akurasi klasifikasi dibandingkan dengan metode konvensional.

Pemanfaatan BERT dalam penelitian ini menjadi penting karena meningkatnya kompleksitas bahasa pada konten hoaks politik yang tidak dapat ditangani secara optimal oleh metode tradisional [7]. Penyebaran hoaks yang semakin masif menuntut adanya model yang mampu memahami konteks linguistik secara mendalam dan akurat. BERT, dengan kemampuan pemrosesan *bidirectional* dan representasi kontekstual yang lebih unggul, menawarkan solusi yang lebih tepat dalam mengidentifikasi pola kebahasaan yang sulit dikenali pendekatan sebelumnya. Oleh karena itu, penelitian ini tidak hanya diarahkan untuk meningkatkan akurasi klasifikasi berita hoaks, tetapi juga untuk memberikan dasar yang kuat bagi pengembangan sistem deteksi misinformasi berbasis NLP yang lebih adaptif dan relevan di era digital [4].

2. Tinjauan Pustaka

Sejumlah studi sebelumnya dalam klasifikasi konten hoaks memanfaatkan teknik pemrosesan bahasa alami (NLP) yang dikombinasikan dengan algoritma pembelajaran mesin. Pada tahap awal, sebagian besar penelitian menggunakan metode pembelajaran mesin klasik berbasis fitur statistik dan leksikal. Wibowo et al. [8] melakukan studi komparatif terhadap algoritma *Naïve Bayes*, *Support Vector Machine* (SVM), dan *Logistic Regression* dalam klasifikasi berita hoaks. Hasil penelitian menunjukkan bahwa SVM memberikan performa yang lebih baik dibandingkan metode lainnya, namun masih memiliki keterbatasan dalam memahami konteks semantik kata secara mendalam karena sangat bergantung pada representasi fitur manual.

Pendekatan pembelajaran mendalam awal mulai diterapkan untuk meningkatkan performa klasifikasi hoaks. Tan et al. [9] mengusulkan model CNN-LSTM yang mengombinasikan *Convolutional Neural Network* (CNN) dan *Long Short-Term Memory* (LSTM) untuk mendeteksi stance terhadap berita hoaks. CNN digunakan untuk mengekstraksi fitur lokal dari teks, sedangkan LSTM berperan dalam memodelkan dependensi sekuensial. Meskipun pendekatan ini mampu meningkatkan akurasi dibandingkan metode statistik murni, model tersebut masih memiliki keterbatasan dalam menangkap konteks kata secara dua arah serta bergantung pada *embedding* statis.

Penelitian lain oleh Sunan et al. [10] mengembangkan metode klasifikasi berita hoaks berbasis LSTM dengan memanfaatkan word *embedding* statis. Hasil penelitian menunjukkan bahwa model sekuensial lebih efektif dalam mempelajari urutan kata dibandingkan pendekatan bag-of-words. Namun demikian, penggunaan *embedding* statis menyebabkan satu kata memiliki representasi yang sama pada berbagai konteks, sehingga model kurang adaptif terhadap ambiguitas makna yang sering muncul pada teks berita politik.

Secara umum, keterbatasan utama dari metode klasik dan model pembelajaran mendalam awal terletak pada ketidakmampuan dalam memahami konteks linguistik secara komprehensif, terutama pada teks panjang dan kompleks. Sebagian besar pendekatan tersebut memproses informasi secara satu arah atau hanya menangkap fitur lokal dan sekuensial, sehingga kurang optimal dalam memodelkan hubungan semantik global antar kata dalam sebuah kalimat.

Sejumlah penelitian dalam bidang klasifikasi teks secara umum menunjukkan bahwa pemahaman konteks kata secara *bidirectional* berperan penting dalam meningkatkan performa model. Deping et al. [11] membuktikan bahwa model BERT mampu meningkatkan akurasi klasifikasi teks melalui pemahaman konteks kata secara dua arah. Sabiri et al. [12] melakukan analisis komparatif dan menemukan bahwa BERT secara konsisten mengungguli metode tradisional seperti SVM dan Naïve Bayes pada berbagai tugas klasifikasi teks. Sementara itu, Guo [13] mengembangkan pendekatan BERT-Capsules yang mengombinasikan BERT dengan capsule network untuk memperkaya representasi fitur teks, sehingga mampu menangkap informasi semantik yang lebih kompleks.

Meskipun menunjukkan performa yang unggul dalam tugas klasifikasi teks secara umum, sebagian besar penelitian tersebut masih berfokus pada data berbahasa Inggris atau domain teks umum, serta belum secara khusus mengkaji penerapan BERT pada klasifikasi konten hoaks berita politik berbahasa Indonesia yang memiliki karakteristik linguistik dan konteks sosial yang berbeda.

Dari berbagai penelitian terdahulu dapat disimpulkan bahwa metode klasifikasi tradisional maupun model deep learning awal masih kurang optimal dalam merepresentasikan konteks linguistik secara mendalam, terutama pada konten berita politik yang kontekstual. Oleh karena itu, penelitian ini mengusulkan penggunaan model BERT untuk klasifikasi konten hoaks berita politik berbahasa Indonesia menggunakan dataset terverifikasi dari Indonesian Fact and Hoax Dataset (IFAHG). Penelitian ini menekankan evaluasi performa model secara komprehensif melalui analisis *confusion matrix* serta metrik *precision*, *recall*, dan *F1-score* guna menilai keseimbangan performa antar kelas fakta dan hoaks.

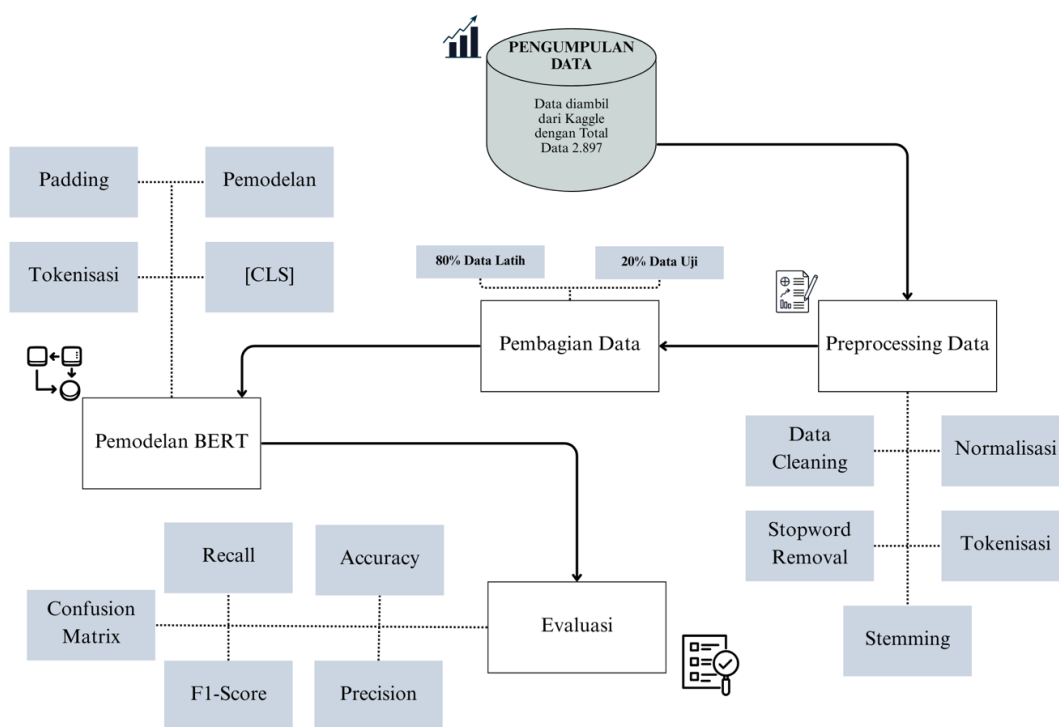
3. Metodologi

Bagian ini menjelaskan metode yang digunakan dalam penelitian untuk mencapai tujuan klasifikasi berita hoaks menggunakan model BERT. Metodologi disusun secara sistematis mulai dari pengumpulan data, proses pra-proses teks, pembagian data, penerapan model BERT, hingga evaluasi performa model. Setiap tahapan dirancang untuk memastikan bahwa proses analisis berlangsung secara terstruktur dan menghasilkan keluaran yang valid. Dengan penjelasan metodologi yang komprehensif, penelitian ini memberikan gambaran yang jelas mengenai alur kerja serta pendekatan teknis yang diterapkan dalam pengembangan sistem klasifikasi berita hoaks.

3.1 Alur Penelitian

Penelitian ini menerapkan alur kerja yang sistematis dalam melakukan klasifikasi berita hoaks berbasis BERT, yang mencakup tahapan pengumpulan data, pra-proses data, pembangunan model, serta evaluasi hasil. Tahapan-tahapan ini dijelaskan secara rinci untuk

menggambarkan alur kerja penelitian secara menyeluruh sebagaimana ditampilkan pada Gambar 1.



Gambar 1. Alur Penelitian

Alur penelitian yang digunakan dalam studi ini dijelaskan secara sistematis pada **Gambar 1**. Penelitian ini bertujuan untuk melakukan klasifikasi berita hoaks (*fake news classification*) menggunakan model bert. Alur tersebut terdiri dari beberapa tahapan utama, yaitu: pengumpulan data, preprocessing data, pemodelan menggunakan BERT, dan evaluasi performa model. Berikut adalah penjelasan tiap tahapan secara rinci :

3.1.1 Pengumpulan Data

Data yang digunakan dalam penelitian ini merupakan subset dari dataset *Indonesian Fact and Hoax Dataset* (IFAHF) yang berfokus pada kategori berita politik. Dataset tersebut disusun oleh Masyarakat Anti Fitnah Indonesia (MAFINDO) dengan memanfaatkan sumber berita daring yang telah diverifikasi seperti CNN Indonesia dan Kompas.id. Dataset ini terdiri atas dua label utama, yaitu fakta dan hoaks, yang masing-masing telah melalui proses verifikasi oleh tim pengecek fakta (*fact-checkers*) MAFINDO.

Secara keseluruhan, dataset ini berjumlah 3.807 data berita, dengan rincian 1.500 berita fakta dan 2.307 berita hoaks, yang disajikan dalam format CSV. Setiap entri berita berisi beberapa kolom, antara lain: title (judul berita), content (isi berita), label (kategori fakta atau hoaks), serta source (sumber berita). Dataset ini dipilih karena memiliki struktur yang jelas, bersifat terbuka, dan relevan dengan tujuan penelitian yaitu melakukan klasifikasi berita politik antara fakta dan hoaks menggunakan model BERT.

3.1.2 Praproses Data

Tahap praproses data berfungsi sebagai fondasi dalam pipeline pemodelan teks karena kualitas data masukan berpengaruh langsung terhadap kinerja model. Pada penelitian ini, praproses dirancang untuk menyesuaikan format teks dengan kebutuhan arsitektur BERT. Proses yang dilakukan meliputi pembersihan data dari simbol, angka, URL, tanda baca berlebihan, dan karakter non-alfabet, dilanjutkan dengan case folding untuk menyeragamkan

huruf menjadi huruf kecil. Selanjutnya dilakukan penghapusan *stopword* guna mengurangi kata-kata yang tidak memiliki nilai informatif dalam klasifikasi berita hoaks.

Tabel 1. Tahapan Praproses Teks sebelum Pelatihan Model BERT

Tahap Praproses	Deskripsi
Data Cleaning	Menghapus simbol, angka, URL, karakter khusus
Case Folding	Mengubah seluruh huruf menjadi huruf kecil
Stopword Removal	Menghapus kata-kata umum yang tidak signifikan
Tokenisasi	Memecah teks menjadi token (kata/sub-kata) sesuai format BERT
Stemming	Mengembalikan kata ke bentuk dasarnya

Usai dilakukan penghapusan *stopword*, teks kemudian diproses melalui tahap tokenisasi untuk membagi kalimat menjadi token-token individual. BERT menggunakan *token* berbasis sub-kata, sehingga kata “berita” bisa dipecah menjadi “ber” dan “##ita” sesuai kamus internal *tokenizer*. Setelah seluruh tahapan sebelumnya, dilakukan stemming sebagai tahap akhir untuk menyederhanakan kata-kata turunan menjadi bentuk dasar, misalnya “berlari”, “pelari”, dan “berlari-lari” yang dipetakan ke kata “lari”. Seluruh tahapan ini bertujuan untuk mengurangi kompleksitas teks, meningkatkan konsistensi representasi semantik, dan meminimalkan noise selama pelatihan model. Proses praproses ini merujuk pada praktik umum yang telah banyak dibahas dalam studi terdahulu [14], [3].

3.1.3 Pembagian Data

Setelah proses praproses selesai, data kemudian dibagi menjadi dua subset, yaitu data latih dan data uji, dengan komposisi 80% untuk pelatihan dan 20% untuk pengujian.

Melalui pembagian data tersebut, sebagian besar data digunakan untuk melatih model BERT, sedangkan sisanya dimanfaatkan sebagai data uji untuk mengukur kinerja model pada data yang belum pernah dilihat. Dengan demikian, kemampuan generalisasi model dalam mengklasifikasikan berita hoaks dapat dievaluasi secara lebih objektif.

3.1.4 Pemodelan dengan BERT

Setelah data melalui tahapan praproses, data kemudian dikonversi ke dalam format input yang sesuai dengan model BERT. Proses ini diawali dengan tokenisasi menggunakan *tokenizer* BERT untuk menghasilkan *token-token* yang sesuai dengan kamus BERT. Selanjutnya, ditambahkan *token* khusus seperti [CLS] yang berfungsi sebagai representasi keseluruhan kalimat dan [SEP] yang digunakan untuk memisahkan segmen teks. Untuk memastikan panjang input seragam, dilakukan proses padding dan truncation. *Token* [PAD] ditambahkan pada teks yang lebih pendek dari panjang maksimum input, sedangkan teks yang melebihi batas dipotong (truncated). Secara matematis, sebuah teks masukan direpresentasikan sebagai urutan *token*

$$X = \{x_1, x_2, \dots, x_n\} \quad (1)$$

yang kemudian dipetakan menjadi vektor *embedding* dengan mengombinasikan *token embedding*, *segment embedding*, dan *position embedding*. Proses ini dapat dinyatakan sebagai:

$$E_i = E_{token}(x_i) + E_{segment}(x_i) + E_{position}(i) \quad (2)$$

Embedding tersebut kemudian diproses oleh lapisan Transformer pada BERT yang menggunakan mekanisme *self-attention*. Mekanisme *self-attention* dihitung menggunakan persamaan:

$$Attention(Q, K, V) = softmax\left(\frac{QK^T}{\sqrt{d_k}}\right)V \quad (3)$$

di mana Q , K , dan V masing-masing merepresentasikan *query*, *key*, dan *value*, serta d_k adalah dimensi vektor *key*. Output representasi dari *token* [CLS] pada lapisan terakhir BERT digunakan sebagai representasi kalimat secara keseluruhan dan diteruskan ke lapisan klasifikasi. Proses klasifikasi dilakukan dengan menerapkan fungsi *softmax* untuk menghasilkan probabilitas setiap kelas, yang dirumuskan sebagai:

$$P(y|x) = softmax(W h_{[CLS]} + b) \quad (4)$$

di mana $h_{[CLS]}$ merupakan vektor keluaran *token* [CLS], sedangkan W dan b adalah bobot dan bias pada lapisan klasifikasi. Untuk proses pelatihan model, digunakan fungsi *loss* Cross-Entropy sebagai objektif optimasi, yang dirumuskan sebagai:

$$L = - \sum_{i=1}^c y_i \log(\hat{y}_i) \quad (5)$$

dengan y_i merupakan label sebenarnya dan $\hat{y}_i$ adalah probabilitas prediksi untuk kelas ke- i . Proses optimasi dilakukan menggunakan *optimizer* berbasis *gradient descent*, di mana parameter model diperbarui melalui mekanisme *backpropagation* selama proses *fine-tuning* pada dataset berita hoaks.

3.1.5 Evaluasi Model

Proses evaluasi diterapkan guna mengukur seberapa efektif model BERT dalam melakukan klasifikasi berita hoaks secara akurat. Beberapa metrik evaluasi umum pada klasifikasi biner digunakan, di mana **akurasi (accuracy)** berperan dalam mengukur proporsi prediksi yang benar dibandingkan dengan total data yang diuji. [15].

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN} \quad (6)$$

Eq. 6 menunjukkan rumus akurasi, di mana TP adalah jumlah *true positive* (berita hoaks diklasifikasi benar), TN adalah *true negative* (berita valid diklasifikasi benar), FP adalah *false positive* (berita valid salah diklasifikasi sebagai hoaks), dan FN adalah *false negative* (berita hoaks salah diklasifikasi sebagai valid). Akurasi digunakan untuk melihat performa model secara umum, namun pada kondisi data yang tidak seimbang diperlukan metrik tambahan. Salah satunya adalah **presisi (precision)**, yang mengukur seberapa banyak prediksi positif yang benar-benar relevan. [16].

$$Precision = \frac{TP}{TP + FP} \quad (7)$$

Eq. 7 menunjukkan rumus presisi. Presisi penting jika tujuan utamanya adalah meminimalkan false positive, misalnya ketika sangat penting untuk tidak mengklasifikasikan berita valid sebagai berita hoaks. Selanjutnya ada **Recall (Sensitivity)** yang mengukur seberapa banyak data positif yang berhasil diidentifikasi oleh model [17].

$$Recall = \frac{TP}{TP + FN} \quad (8)$$

Eq. 8 menunjukkan rumus *recall*. *Recall* sangat berguna ketika fokus utama adalah meminimalkan false negative, misalnya agar berita hoaks tidak lolos klasifikasi. **F1-score** merupakan metrik evaluasi yang diperoleh dari rata-rata harmonik presisi dan *recall*, sehingga mampu menggambarkan performa model secara lebih seimbang [18].

$$F1 - Score = 2 \times \frac{Precision \times Recall}{Precision + Recall} \quad (9)$$

Eq. 9 menunjukkan rumus *F1-Score*. Metrik ini ideal digunakan ketika terdapat trade-off antara presisi dan *recall*, serta dalam kasus dataset yang tidak seimbang.

Penelitian lain menunjukkan pentingnya evaluasi yang mendalam menggunakan metrik seperti *F1-Score* untuk menilai model dalam klasifikasi berita hoaks [19].

Confusion matrix digunakan sebagai metode evaluasi tambahan untuk mengkaji kesalahan klasifikasi secara lebih mendalam, dengan menyajikan perbandingan antara prediksi model dan kelas sebenarnya dalam bentuk tabel. Tabel ini terdiri atas empat elemen penting, yaitu *True Positive* (TP), yang menunjukkan kasus ketika berita hoaks berhasil diklasifikasi dengan benar sebagai hoaks; *False Positive* (FP), yaitu kasus ketika berita valid salah diklasifikasi sebagai hoaks; *True Negative* (TN), yaitu kasus ketika berita valid berhasil diklasifikasi dengan benar sebagai valid; dan *False Negative* (FN), yaitu kasus ketika berita hoaks salah diklasifikasi sebagai berita valid. Analisis terhadap keempat komponen ini membantu dalam memahami jenis kesalahan yang paling sering dilakukan oleh model, serta dalam meningkatkan performa model melalui penyesuaian lebih lanjut.

Dengan menggunakan *confusion matrix*, peneliti dapat memahami jenis kesalahan apa yang paling sering dilakukan model, sehingga dapat menjadi dasar untuk peningkatan performa model di tahap selanjutnya.

3.2 Arsitektur dan Konfigurasi Eksperimen (Spesifikasi)

Dalam penelitian ini digunakan arsitektur BERT (*Bidirectional Encoder Representations from Transformers*) dengan model dasar bert-base-uncased yang disediakan oleh pustaka Transformers dari Hugging Face. Model ini memiliki 12 layer encoder, 768 dimensi hidden size, dan 12 *self-attention* heads [20]. Model bertipe uncased digunakan agar tidak membedakan huruf kapital dan huruf kecil, yang dinilai sesuai untuk tugas klasifikasi umum seperti deteksi berita hoaks [21]. Proses tokenisasi dilakukan menggunakan BertTokenizer.from_pretrained('bert-base-uncased'), di mana setiap input teks akan diproses dengan menambahkan *token* khusus seperti [CLS] di awal dan [SEP] di akhir kalimat, serta penyesuaian panjang input menggunakan padding dan truncation.

Tabel 2. Spesifikasi dan Proses Pra-pemrosesan Input pada Model BERT

Komponen / Tahapan	Deskripsi
Model BERT	bert-base-uncased
Jumlah layer encoder	12 Layer
Dimensi hidden size	768
Jumlah attention heads	12
Tokenizer	BertTokenizer.from_pretrained('bert-base-uncased')
Token khusus	[CLS] (awal input), [SEP] (akhir kalimat)
Panjang input	Maksimal 256 token (dengan padding & truncation)

Komponen / Tahapan	Deskripsi
Output tokenisasi	Input IDs, Attention Mask
Format data akhir	Tensor PyTorch (<code>torch.tensor(...)</code>)
Tujuan tokenisasi	Menyesuaikan format input dengan struktur model BERT untuk klasifikasi

Setelah tahap tokenisasi, data teks yang telah diubah ke dalam bentuk numerik disusun sebagai tensor PyTorch, sehingga dapat langsung digunakan dalam pelatihan model BERT. Komponen input IDs dan attention mask digunakan untuk memastikan bahwa hanya *token-token* bermakna yang dihitung selama proses forward pass. Dengan mengikuti struktur input yang konsisten ini, model BERT dapat melakukan klasifikasi secara efektif dan efisien terhadap teks berita yang dianalisis, khususnya dalam membedakan antara berita valid dan berita hoaks.

4. Hasil dan Pembahasan

Hasil dan pembahasan pada bab ini berfokus pada proses analisis serta evaluasi model BERT dalam mengklasifikasikan berita politik antara hoaks dan fakta. Uraian mencakup pengamatan terhadap dataset, efektivitas preprocessing, kontribusi tiap parameter pelatihan, serta performa akhir model berdasarkan berbagai metrik evaluasi.

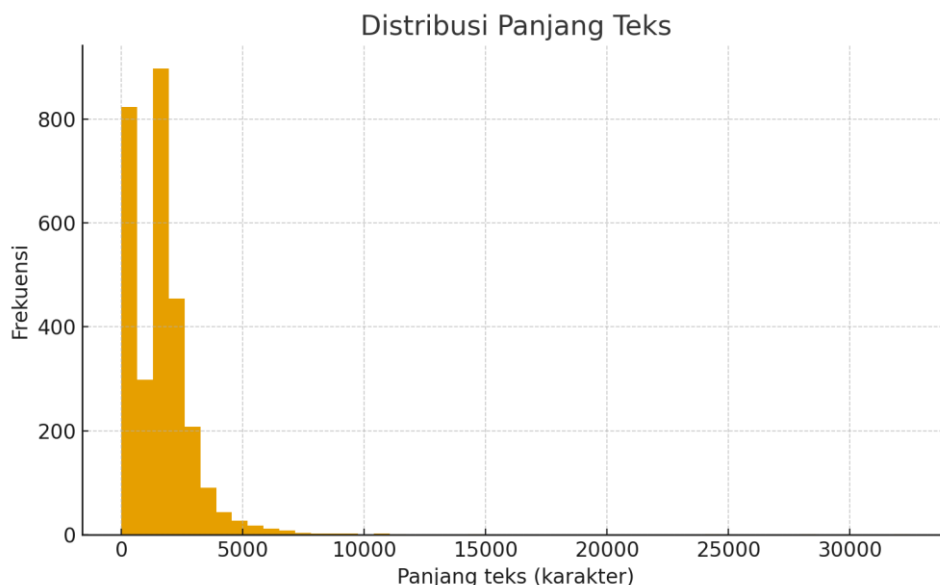
4.1. Analisis Dataset

Dataset yang digunakan dalam penelitian ini terdiri dari 3.807 data berita politik yang berasal dari *Indonesian Fact and Hoax Dataset* (IFAHD). Data tersebut terbagi ke dalam dua kategori utama, yaitu berita fakta dan berita hoaks, yang masing-masing telah diverifikasi oleh tim pemeriksa fakta. Dari total keseluruhan data, terdapat 1.500 data berita fakta dan 2.307 data berita hoaks, sehingga distribusinya tidak seimbang dengan proporsi hoaks yang lebih tinggi. Ketidakseimbangan kelas ini menjadi perhatian khusus karena dapat memengaruhi performa model selama proses pelatihan, terutama dalam mendeteksi kelas dengan jumlah data lebih sedikit.

Tabel 3. Distribusi Data Berdasarkan Kategori Berita

No	Kategori Berita	Jumlah Data
1	Fakta	1.500
2	Hoaks	2.307
Total	—	3.807

Distribusi panjang teks pada dataset dianalisis untuk memahami karakteristik konten berita yang digunakan dalam pelatihan model. Berdasarkan hasil pengukuran, panjang teks pada dataset menunjukkan variasi yang cukup besar antar sampel. Mayoritas teks memiliki panjang antara 100 hingga 350 kata, yang mencerminkan karakteristik umum berita politik yang cenderung lebih panjang dan informatif. Selain itu, ditemukan pula sejumlah teks yang sangat pendek, dengan panjang kurang dari 50 kata, serta beberapa teks yang sangat panjang melebihi 500 kata. Variasi panjang teks ini penting untuk diperhatikan karena dapat memengaruhi proses tokenisasi dan pembatasan panjang input pada model BERT, mengingat model memiliki batas maksimum *token*. Oleh karena itu, analisis distribusi panjang teks memberikan gambaran awal mengenai kompleksitas dan keragaman konten yang akan diproses oleh model sebelum dilakukan tahapan preprocessing.



Gambar 2. Distribusi panjang teks berita pada dataset sebelum dilakukan proses preprocessing.

4.2. Hasil Preprocessing Data

Pada tahap ini dilakukan serangkaian proses preprocessing untuk memastikan bahwa teks berada dalam bentuk yang bersih, konsisten, dan siap digunakan sebagai input model BERT. Proses ini meliputi pembersihan karakter-karakter yang tidak relevan seperti angka, simbol, URL, tanda baca berlebihan, serta karakter non-alfabet yang dapat menimbulkan noise selama pelatihan model. Selanjutnya, dilakukan case folding untuk menyeragamkan seluruh huruf menjadi huruf kecil, sehingga perbedaan kapitalisasi tidak memengaruhi makna *token*. Proses dilanjutkan dengan penghapusan stopwords agar kata-kata yang tidak memberikan kontribusi semantik dapat dieliminasi. Selain itu, teks diubah menjadi *token* menggunakan *tokenizer* BERT yang menghasilkan *token* berbasis sub-kata. Tahap terakhir adalah stemming untuk mengembalikan kata ke bentuk dasarnya. Seluruh tahapan preprocessing ini bertujuan untuk meningkatkan kualitas representasi teks sehingga model dapat mempelajari konteks dengan lebih optimal. Hasil preprocessing menunjukkan bahwa teks yang awalnya tidak terstruktur berhasil disederhanakan menjadi bentuk yang lebih bersih dan konsisten untuk kebutuhan pelatihan model.

Tabel 4. Contoh Data Sebelum dan Sesudah Preprocessing

Sebelum	Sesudah	Label
Anies di Milad BKMT: Pengajian Menghasilkan Ibu-ibu Berpengetahuan. Mantan	anies milad bkmt pengajian hasil ibu ibu pengetahuan mantan gubernur dki jakarta anies	0
BISA DILIHAT SI ONTA YAMAN NGGAK PEDULI ITU APA YANG PENTING DAPAT SUARA SEGALA CARA DEMI KEKUASAAN	lihat onta yaman tidak peduli penting dapat suara cara demi kuasa	1

4.3. Pembagian Data

Setelah seluruh data melalui proses preprocessing, langkah berikutnya adalah melakukan pembagian dataset menjadi data latih dan data uji. Pembagian ini bertujuan untuk memberikan pemisahan yang jelas antara data yang digunakan dalam proses pelatihan model

dan data yang digunakan untuk mengukur kemampuan generalisasi model terhadap teks yang belum pernah dilihat sebelumnya.

Tabel 5. Pembagian Data Latih dan Data Uji

No	Jenis Data	Jumlah Data	Presentase
1	Data Latih	3.045	80%
2	Data Uji	762	20%
Total	-	3.807	100%

Rasio pembagian data sebesar 80:20 dipilih karena merupakan standar yang banyak digunakan dalam penelitian berbasis machine learning, khususnya pada tugas klasifikasi teks. Dengan memberikan 80% data sebagai data latih, model memperoleh jumlah sampel yang cukup untuk mempelajari pola, struktur, dan karakteristik teks secara lebih komprehensif. Sementara itu, porsi 20% digunakan sebagai data uji untuk menghasilkan evaluasi yang objektif terhadap kemampuan model dalam melakukan generalisasi terhadap data baru. Rasio ini dinilai optimal karena mampu menjaga keseimbangan antara kebutuhan pelatihan dan evaluasi. Penggunaan rasio lain, seperti 70:30 atau 90:10, tidak dipilih karena dapat menimbulkan permasalahan. Rasio 70:30 berpotensi mengurangi kualitas pembelajaran karena data latih yang lebih sedikit, sedangkan rasio 90:10 dapat menyebabkan evaluasi yang kurang representatif akibat jumlah data uji yang terlalu terbatas. Oleh karena itu, rasio 80:20 dipandang paling sesuai untuk penelitian ini.

4.4. Hasil Pemodelan BERT

Pada bagian ini disajikan hasil pemodelan menggunakan arsitektur BERT yang telah melalui proses pelatihan dengan dataset berita politik yang telah dipraproses sebelumnya. Proses pemodelan mencakup konfigurasi parameter pelatihan, performa model selama proses *training*, serta visualisasi perkembangan nilai loss dan akurasi pada setiap epoch. Hasil ini digunakan untuk menilai sejauh mana model mampu mempelajari pola teks dan membedakan berita fakta serta berita hoaks secara efektif.

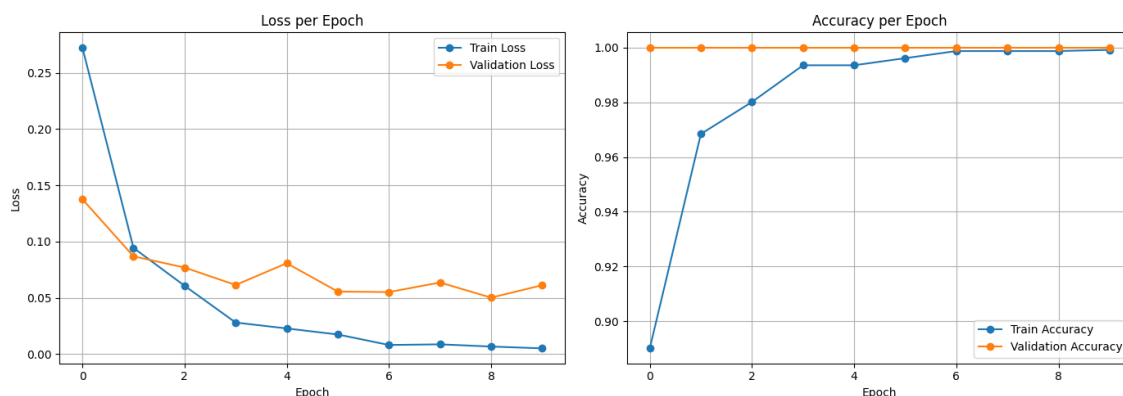
Tabel 6. Parameter Pelatihan Model BERT

Parameter	Nilai
Model	BERT Base Uncased
Maksimal Panjang Input	256 Token
Epoch	10
Batch Size	16
Learning Rate	2e-5
Optimizer	AdamW
Loss Function	Cross-Entropy Loss
Train-Test Split	80:20
Warmup Steps	0
Dropout Probability	0.1

Tabel parameter pelatihan pada model BERT menunjukkan konfigurasi yang digunakan selama proses *training*. Model yang diterapkan adalah **BERT Base Uncased**, yaitu versi standar BERT dengan 12 layer encoder dan tanpa membedakan huruf kapital, sehingga sesuai untuk tugas klasifikasi teks umum. Panjang input ditetapkan sebesar **256 token** untuk menyeimbangkan kebutuhan menangkap konteks informasi dalam teks berita dan efisiensi komputasi. Pelatihan dilakukan selama **10 epoch** dengan **batch size 16**, yang dipilih agar proses pembaruan parameter berlangsung stabil tanpa memerlukan memori komputasi yang besar. Nilai

learning rate 2e-5 merupakan nilai yang umum digunakan pada *fine-tuning* BERT, karena cukup kecil untuk mencegah overshooting namun masih efektif dalam memperbaiki bobot model.

Optimizer yang digunakan adalah **AdamW**, yang dirancang khusus untuk meningkatkan stabilitas *training* model berbasis Transformer. Fungsi loss yang digunakan adalah **Cross-Entropy Loss**, sesuai untuk permasalahan klasifikasi biner antara berita fakta dan hoaks. Pembagian data **80:20** antara data latih dan uji memastikan keseimbangan antara kemampuan pembelajaran dan evaluasi generalisasi model. Selain itu, konfigurasi **warmup steps 0** dan **dropout 0.1** juga digunakan untuk mencegah overfitting dan menjaga stabilitas model selama proses pelatihan. Secara keseluruhan, parameter-parameter ini dipilih untuk menghasilkan proses *fine-tuning* yang optimal pada dataset berita politik.



Gambar 3. Grafik perkembangan nilai loss dan akurasi selama proses pelatihan model BERT.

Dari grafik *loss*, dapat diamati bahwa nilai *training loss* mengalami penurunan yang signifikan dan berkelanjutan, dari sekitar **0,27** pada epoch pertama hingga mendekati nol pada epoch kesembilan, yang menunjukkan proses pembelajaran model berjalan sangat baik. Sementara itu, *validation loss* juga mengalami tren penurunan dari sekitar **0,13** pada epoch awal menjadi sekitar **0,05–0,07** pada epoch akhir, meskipun dengan fluktuasi kecil di beberapa titik. Pola ini menandakan bahwa model tidak mengalami overfitting signifikan, karena *validation loss* tetap berada pada rentang yang stabil dan menurun secara umum.

Pada grafik akurasi, terlihat bahwa *train accuracy* meningkat secara drastis dari sekitar **0,95** pada epoch 0 hingga mendekati **1,00 (100%)** pada epoch ke-3, kemudian stabil di angka tersebut hingga epoch terakhir. Sementara itu, *validation accuracy* menunjukkan konsistensi yang sangat tinggi, yaitu berada di **1,00 (100%)** sejak epoch awal hingga epoch ke-9 tanpa adanya penurunan. Hal ini menunjukkan bahwa model mencapai performa maksimum sangat cepat dan mampu mempertahankannya secara konsisten pada data validasi.

Secara keseluruhan, hasil ini menunjukkan bahwa model BERT tidak hanya berhasil mempelajari pola linguistik secara efektif, tetapi juga mencapai performa yang sangat tinggi dengan stabilitas yang baik. Kombinasi penurunan *loss* yang konsisten dan akurasi validasi yang tetap sempurna menunjukkan bahwa model berada pada kondisi optimal tanpa gejala overfitting. Efektivitas ini kemungkinan berasal dari kemampuan pretrained BERT yang sudah kuat dalam memahami struktur bahasa Indonesia, sehingga proses *fine-tuning* hanya memperhalus parameter pada domain berita politik.

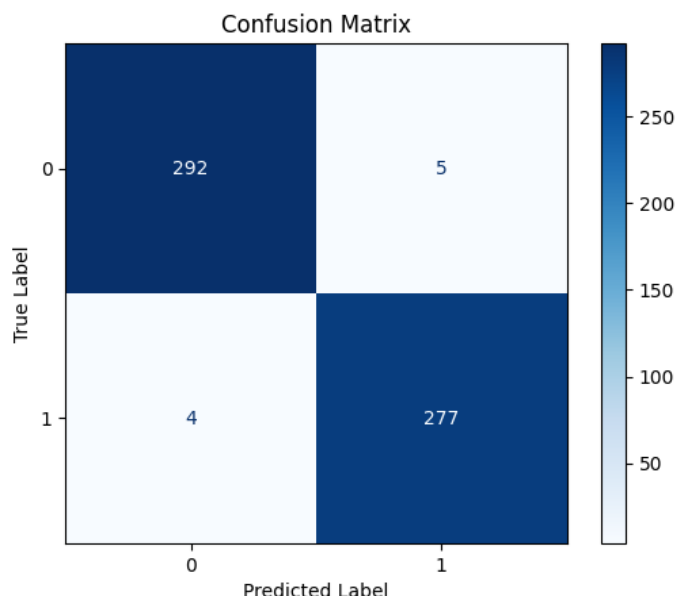
Dengan performa yang sangat tinggi ini, BERT terbukti mampu menangani variasi teks kompleks, termasuk gaya bahasa politik yang berubah-ubah. Meski demikian, untuk eksplorasi lebih lanjut, penelitian selanjutnya dapat mempertimbangkan penggunaan dataset yang lebih besar, *augmentation* teks, modifikasi arsitektur, atau pengaturan *hyperparameter* lebih lanjut untuk menguji apakah performa model dapat ditingkatkan atau lebih stabil pada data yang lebih menantang.

4.5. Hasil Evaluasi Model

Pada bagian ini disajikan hasil evaluasi model BERT setelah melalui proses pelatihan. Evaluasi dilakukan untuk mengukur kemampuan model dalam melakukan klasifikasi berita hoaks dan berita fakta pada data uji. Analisis meliputi *confusion matrix*, *classification report*, serta

interpretasi metrik evaluasi guna mengetahui tingkat akurasi, presisi, *recall*, dan *F1-score* yang dicapai oleh model.

4.5.1. Confusion Matrics



Gambar 4. *Confusion Matrix* dari proses pelatihan model BERT.

Gambar *confusion matrix* menunjukkan performa model BERT dalam mengklasifikasikan data uji ke dalam dua kelas, yaitu berita fakta (kelas 0) dan berita hoaks (kelas 1). Pada sel diagonal terlihat bahwa model berhasil mengklasifikasikan **292 data fakta** secara benar sebagai fakta (*True Negative*) dan **277 data hoaks** secara benar sebagai hoaks (*True Positive*). Nilai yang tinggi pada kedua sel ini menunjukkan bahwa model memiliki kemampuan yang sangat baik dalam mengenali pola linguistik dari masing-masing kelas.

Pada sisi kesalahan prediksi, terdapat **5 data fakta** yang salah diklasifikasikan sebagai hoaks (*False Positive* = 5). Jumlah ini terbilang sangat kecil, yang menunjukkan bahwa model sangat jarang menganggap berita valid sebagai hoaks. Sementara itu, terdapat **4 data hoaks** yang salah diprediksi sebagai fakta (*False Negative* = 4). Angka FN yang rendah ini mengindikasikan bahwa model mampu mengenali mayoritas berita hoaks meskipun masih ada beberapa kasus yang terlewat.

Distribusi nilai yang ditampilkan pada *confusion matrix* secara keseluruhan memperlihatkan bahwa model BERT menunjukkan performa klasifikasi yang optimal dan konsisten pada kedua kelas. Tingginya nilai TP dan TN serta rendahnya FP dan FN menegaskan bahwa model mampu membedakan berita fakta dan hoaks secara efektif. Hasil ini menunjukkan bahwa model telah berfungsi optimal pada dataset politik yang digunakan dan memiliki stabilitas yang kuat dalam proses prediksinya.

4.5.2. Report

Tabel 7. *Classification Report* Model BERT

Kelas	Precision	Recall	F1-Score	Support
Valid (0)	0.99	0.98	0.98	297
Hoaks (1)	0.98	0.99	0.98	281
Accuracy			0.98	578
Macro Avg	0.98	0.98	0.98	578
Weighted Avg	0.98	0.98	0.98	578

Tabel *classification report* memberikan gambaran menyeluruh mengenai performa model BERT dalam mengklasifikasikan berita fakta (kelas 0) dan berita hoaks (kelas 1). Pada kelas Valid/Fakta (0), model memperoleh nilai **precision sebesar 0.99**, yang menunjukkan bahwa hampir seluruh prediksi model terhadap kelas fakta memang benar-benar berasal dari kelas tersebut. Nilai **recall sebesar 0.98** mengindikasikan bahwa model berhasil mengenali hampir seluruh data fakta dengan sangat baik, hanya sedikit sampel yang terlewat. Kombinasi kedua metrik tersebut menghasilkan **F1-score 0.98**, mencerminkan performa yang sangat kuat dan seimbang dalam klasifikasi berita fakta.

Pada kelas Hoaks (1), model juga menunjukkan performa tinggi dengan **precision sebesar 0.98**, yang berarti bahwa sebagian besar prediksi hoaks oleh model benar-benar sesuai dengan label aslinya. Nilai **recall sebesar 0.99** menunjukkan bahwa model mampu menangkap hampir seluruh berita hoaks tanpa melewatkan banyak sampel. Hal ini menghasilkan **F1-score sebesar 0.98**, menandakan bahwa model memiliki kinerja yang sangat baik dalam mendeteksi hoaks, bahkan sedikit lebih baik dari performa pada kelas fakta.

Secara keseluruhan, model mencapai **akurasi sebesar 0.98 (98%)**, yang menunjukkan bahwa hampir seluruh prediksi model sesuai dengan label sebenarnya. Nilai **macro average dan weighted average**, masing-masing sebesar **0.98**, menguatkan bahwa performa model stabil pada kedua kelas tanpa adanya ketimpangan yang berarti. Hasil ini memperlihatkan bahwa model BERT sangat efektif dan mampu mengklasifikasikan berita politik baik fakta maupun hoaks dengan tingkat akurasi yang sangat tinggi, sehingga dapat diandalkan dalam tugas deteksi hoaks pada domain ini.

4.5.3. Pola dan Temuan Hasil

Berdasarkan hasil *confusion matrix*, model BERT menunjukkan kemampuan klasifikasi yang sangat baik pada kedua kelas, baik berita fakta maupun berita hoaks. Nilai **True Negative sebesar 292** menunjukkan bahwa sebagian besar data berita fakta berhasil diklasifikasikan dengan benar sebagai fakta. Namun, terdapat **5 data fakta** yang salah diprediksi sebagai hoaks (*False Positive*), meskipun jumlah ini relatif kecil dibanding keseluruhan data. Pada kelas hoaks, model berhasil mengidentifikasi **277 berita hoaks** dengan benar (*True Positive*), sementara **4 data hoaks** lainnya salah diklasifikasikan sebagai fakta (*False Negative*). Secara umum, distribusi ini memperlihatkan bahwa model memiliki akurasi tinggi dan kesalahan prediksi yang sangat minimal.

Pola kesalahan ini menunjukkan bahwa model memiliki sensitivitas yang baik terhadap kedua kelas. Hal ini terlihat dari nilai **recall** untuk kelas fakta (0) yang mencapai **0.98**, serta **recall** kelas hoaks (1) yang mencapai **0.99**, mencerminkan kemampuan model dalam mengenali hampir seluruh data pada masing-masing kelas. **Precision** pada kedua kelas juga sangat tinggi, yaitu **0.99 untuk kelas fakta** dan **0.98 untuk kelas hoaks**, yang berarti bahwa ketika model memberikan suatu prediksi, keputusannya hampir selalu tepat. Kombinasi **precision** dan **recall** yang sama-sama tinggi menghasilkan **F1-score 0.98** pada kedua kelas, menandakan performa yang sangat seimbang.

Temuan ini mengindikasikan bahwa model BERT telah bekerja dengan performa yang sangat kuat dalam mendeteksi berita fakta maupun hoaks. Rendahnya nilai **false positive** dan **false negative** memperlihatkan bahwa model mampu meminimalkan kesalahan prediksi, baik dalam menghindari pelabelan hoaks yang salah maupun dalam mengenali berita hoaks secara akurat. Dengan hasil ini, model dapat dianggap stabil dan sangat efektif, meskipun penelitian selanjutnya tetap dapat mengeksplorasi optimasi tambahan untuk meningkatkan robustness terhadap variasi teks yang lebih beragam.

4.6. Diskusi

Hasil penelitian memperlihatkan bahwa penerapan model BERT menghasilkan performa yang optimal dalam mengklasifikasikan berita politik, baik untuk kelas fakta maupun kelas hoaks. Keseimbangan distribusi data antar kelas berperan penting dalam mendukung proses pelatihan yang stabil, sehingga model tidak mengalami bias signifikan terhadap salah satu kelas. Selain itu, tahapan praproses yang mencakup pembersihan teks, normalisasi, dan tokenisasi turut membantu model dalam memahami struktur kalimat serta variasi gaya bahasa yang umum ditemukan pada teks berita politik.

Pada tahap pelatihan, tren penurunan nilai loss yang konsisten serta peningkatan nilai akurasi pada setiap epoch menunjukkan bahwa proses pembelajaran berlangsung secara efektif.

Nilai train accuracy yang meningkat hingga mendekati 1.00 dan *validation accuracy* yang tetap stabil pada nilai yang sama mengindikasikan bahwa model memiliki kemampuan generalisasi yang baik terhadap data yang tidak digunakan selama pelatihan. Stabilitas performa validasi tersebut menunjukkan bahwa model tidak mengalami overfitting, sehingga hasil pelatihan dapat merepresentasikan kemampuan model secara objektif dalam melakukan klasifikasi.

Pemilihan parameter pelatihan, seperti learning rate sebesar $2e-5$, batch size 16, serta penggunaan optimizer AdamW, terbukti mendukung stabilitas dan efektivitas proses *fine-tuning*. Kombinasi parameter ini memungkinkan model untuk menyesuaikan bobot secara optimal sehingga mampu mempelajari pola semantik dan konteks linguistik yang kompleks, khususnya pada berita politik berbahasa Indonesia yang memiliki karakteristik bahasa formal dan narasi argumentatif.

Hasil evaluasi menggunakan *confusion matrix* menunjukkan bahwa model mampu mengklasifikasikan sebagian besar data secara tepat pada kedua kelas. Model berhasil mengklasifikasikan 292 dari 297 berita fakta dan 277 dari 281 berita hoaks dengan benar. Nilai *precision*, *recall*, dan *F1-score* yang berada pada kisaran 0.98–0.99 menunjukkan bahwa model memiliki performa yang seimbang dalam mengenali berita fakta maupun mendeteksi berita hoaks. Tingginya nilai *recall* pada kedua kelas mengindikasikan bahwa model jarang melewatkan sampel penting, sehingga risiko kesalahan klasifikasi dapat diminimalkan.

Temuan dalam penelitian ini memperkuat hasil penelitian terdahulu yang menunjukkan bahwa metode pembelajaran mesin klasik masih memiliki keterbatasan dalam memahami konteks linguistik secara mendalam pada tugas klasifikasi berita hoaks. Penelitian Qubra dan Saputra [22] menunjukkan bahwa penggunaan metode Naïve Bayes mampu mencapai akurasi sebesar 94,73%, namun model tersebut masih menghasilkan kesalahan klasifikasi yang relatif signifikan, terutama pada false negative, yang mengindikasikan keterbatasan model dalam memahami makna kontekstual teks berita. Hal serupa juga ditunjukkan oleh Ferawaty et al. [23] yang menggunakan metode Support Vector Machine (SVM) dan Logistic Regression dengan akurasi sebesar 96,4%. Meskipun menghasilkan performa yang lebih tinggi dibandingkan Naïve Bayes, pendekatan ini tetap bergantung pada fitur statistik dan leksikal sehingga kurang optimal dalam menangkap hubungan semantik yang kompleks, khususnya pada berita politik yang sarat dengan konteks, implikatur, dan variasi makna.

Keterbatasan metode klasik tersebut sejalan dengan temuan Wibowo et al. [8] yang menyatakan bahwa pendekatan berbasis fitur manual tidak mampu merepresentasikan makna kontekstual secara mendalam. Pendekatan pembelajaran mendalam awal seperti CNN-LSTM dan LSTM berbasis *embedding* statis yang diusulkan oleh Tan et al. [9] dan Sunan et al. [10] memang menunjukkan peningkatan performa dibandingkan metode statistik murni, namun masih memproses informasi secara sekuensial dan satu arah, sehingga kurang adaptif terhadap perbedaan makna kata dalam konteks yang berbeda. Berbeda dengan pendekatan tersebut, penelitian ini mengimplementasikan BERT yang mampu memodelkan konteks dua arah secara simultan. Hasil penelitian menunjukkan bahwa model BERT mencapai akurasi sebesar 98%, yang secara signifikan lebih tinggi dibandingkan metode klasik maupun pendekatan pembelajaran mendalam awal. Hal ini membuktikan bahwa BERT lebih unggul dalam menangkap makna kontekstual teks, sehingga mampu mengatasi kesenjangan (research gap) dari metode sebelumnya dan menghasilkan performa klasifikasi yang lebih stabil dan seimbang pada klasifikasi berita hoaks politik berbahasa Indonesia. Dengan demikian, temuan ini menunjukkan bahwa pemodelan konteks dua arah yang dimiliki BERT merupakan faktor kunci dalam meningkatkan akurasi dan keseimbangan klasifikasi, khususnya pada teks berita politik yang memiliki struktur bahasa kompleks dan makna implisit.

Secara keseluruhan, diskusi ini menunjukkan bahwa penerapan model BERT tidak hanya menghasilkan performa klasifikasi yang tinggi, tetapi juga memberikan kontribusi terhadap penguatan temuan-temuan sebelumnya dalam bidang klasifikasi hoaks berbasis pemrosesan bahasa alami. Penelitian ini berkontribusi pada pengembangan kajian klasifikasi berita hoaks dengan menunjukkan bahwa BERT merupakan pendekatan yang efektif dan andal untuk menangani teks berita politik berbahasa Indonesia, serta memiliki potensi untuk dikembangkan lebih lanjut pada penelitian selanjutnya.

5. Simpulan

Penelitian ini menunjukkan bahwa model BERT mampu melakukan klasifikasi berita politik ke dalam kategori fakta dan hoaks dengan tingkat akurasi yang sangat tinggi, yaitu **98%**.

Model menunjukkan performa yang sangat baik pada kedua kelas. Pada kelas berita fakta, model mencapai **precision 0.99**, **recall 0.98**, dan **F1-score 0.98**. Performa pada kelas hoaks juga sama kuatnya, dengan **precision 0.98**, **recall 0.99**, dan **F1-score 0.98**. Hasil ini menunjukkan bahwa model mampu mengenali baik berita fakta maupun berita hoaks secara akurat, sekaligus menjaga keseimbangan performa pada kedua kelas tanpa adanya kecenderungan bias tertentu.

Temuan ini memenuhi tujuan penelitian, yaitu membangun model berbasis NLP yang akurat dan andal untuk deteksi hoaks dalam domain politik. Performa yang sangat stabil selama proses pelatihan, didukung oleh pemilihan parameter yang optimal serta tahapan preprocessing yang tepat, turut berkontribusi terhadap kemampuan model dalam memahami pola linguistik dan konteks semantik pada teks berita politik.

Untuk pengembangan ke depan, penelitian dapat diarahkan pada perluasan dan variasi dataset dengan menambahkan lebih banyak contoh dari sumber yang lebih beragam, termasuk platform media sosial. Eksplorasi terhadap model Transformer lain seperti IndoBERT, RoBERTa, atau model generatif modern juga dapat dilakukan untuk menguji peningkatan performa. Selain itu, pendekatan seperti augmentasi teks, penyeimbangan kelas, atau penyesuaian threshold prediksi dapat digunakan untuk semakin memperkuat sensitivitas model. Dengan langkah-langkah tersebut, sistem deteksi hoaks berbasis NLP berpotensi menjadi lebih akurat, adaptif, dan reliabel dalam menghadapi dinamika penyebaran misinformasi digital yang semakin kompleks.

Daftar Referensi

- [1] R. K. Kaliyar, A. Goswami, and P. Narang, "FakeBERT: Fake news detection in social media with a BERT-based deep learning approach," *Multimed Tools Appl*, vol. 80, no. 8, pp. 11765–11788, Mar. 2021, doi: 10.1007/s11042-020-10183-2.
- [2] A. J. Keya, Md. A. H. Wadud, M. F. Mridha, M. Alatiyyah, and Md. A. Hamid, "AugFakeBERT: Handling Imbalance through Augmentation of Fake News Using BERT to Enhance the Performance of Fake News Classification," *Applied Sciences*, vol. 12, no. 17, p. 8398, Aug. 2022, doi: 10.3390/app12178398.
- [3] N. Rai, D. Kumar, N. Kaushik, C. Raj, and A. Ali, "Fake News Classification using transformer based enhanced LSTM and BERT," *International Journal of Cognitive Computing in Engineering*, vol. 3, pp. 98–105, June 2022, doi: 10.1016/j.ijcce.2022.03.003.
- [4] S. Minaee, N. Kalchbrenner, E. Cambria, N. Nikzad, M. Chenaghlu, and J. Gao, "Deep Learning--based Text Classification: A Comprehensive Review," *ACM Comput. Surv.*, vol. 54, no. 3, pp. 1–40, Apr. 2022, doi: 10.1145/3439726.
- [5] 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.
- [6] J. Y. Khan, Md. T. I. Khondaker, S. Afroz, G. Uddin, and A. Iqbal, "A benchmark study of machine learning models for online fake news detection," *Machine Learning with Applications*, vol. 4, p. 100032, June 2021, doi: 10.1016/j.mlwa.2021.100032.
- [7] R. K. Rahardi and S. Firda Utami, "Socio-political hoaxes in virtual public space: Assertive illocutionary manifestations of dishonesty in a critical pragmatics perspective," *BAHASTRA*, vol. 44, no. 2, pp. 205–223, Oct. 2024, doi: 10.26555/bs.v44i2.998.
- [8] Technology Information and Industry Faculty, Stikubank University, Jl. Tri Lomba Juang No. 1, Semarang, Indonesia, J. S. Wibowo, E. N. Wahyudi, Vocational Faculty, Stikubank University, Jl. Kendeng V Bendan Ngisor, Semarang, Indonesia, H. Listiyono, and Vocational Faculty, Stikubank University, Jl. Kendeng V Bendan Ngisor, Semarang, Indonesia, "Performance Comparison of SVM, Naive Bayes, and Random Forest Models in Fake News Classification," *ETJ*, vol. 09, no. 08, Aug. 2024, doi: 10.47191/etj/v9i08.27.
- [9] K. L. Tan, C. Poo Lee, and K. M. Lim, "Fake News Detection with Hybrid CNN-LSTM," in *2021 9th International Conference on Information and Communication Technology (ICoICT)*, Yogyakarta, Indonesia: IEEE, Aug. 2021, pp. 606–610. doi: 10.1109/ICoICT52021.2021.9527469.
- [10] R. A. Sunan, H. F. E. K., and C. S. K. Aditya, "Klasifikasi Hoax Berita Politik Menggunakan Algoritma Long Short-Term Memory (LSTM) dengan Penambahan Fitur Embedding Global Vector (GloVe)," *JEPIN*, vol. 10, no. 2, p. 287, Aug. 2024, doi: 10.26418/jp.v10i2.76042.
- [11] L. Deping, W. Hongjuan, L. Mengyang, and L. Pei, "News text classification based on Bidirectional Encoder Representation from Transformers," in *2021 International Conference*

- on Artificial Intelligence, Big Data and Algorithms (CAIBDA), Xi'an, China: IEEE, May 2021, pp. 137–140. doi: 10.1109/CAIBDA53561.2021.00036.
- [12] B. Sabiri, A. Khtira, B. El Asri, and M. Rhanoui, "Analyzing BERT's Performance Compared to Traditional Text Classification Models:," in *Proceedings of the 25th International Conference on Enterprise Information Systems*, Prague, Czech Republic: SCITEPRESS - Science and Technology Publications, 2023, pp. 572–582. doi: 10.5220/0011983100003467.
- [13] M. Guo, "Text classification by BERT-Capsules," *TE*, vol. 1, no. 5, Feb. 2024, doi: 10.61173/wcg0nf17.
- [14] M. A. B. Al-Tarawneh, O. Al-irr, K. S. Al-Maaitah, H. Kanj, and W. H. F. Aly, "Enhancing Fake News Detection with Word Embedding: A Machine Learning and Deep Learning Approach," *Computers*, vol. 13, no. 9, p. 239, Sept. 2024, doi: 10.3390/computers13090239.
- [15] G. Grunau and S. Linn, "Detection and Diagnostic Overall Accuracy Measures of Medical Tests," *Rambam Maimonides Med J*, vol. 9, no. 4, p. e0027, Oct. 2018, doi: 10.5041/RMMJ.10351.
- [16] T. F. Monaghan *et al.*, "Foundational Statistical Principles in Medical Research: Sensitivity, Specificity, Positive Predictive Value, and Negative Predictive Value," *Medicina*, vol. 57, no. 5, p. 503, May 2021, doi: 10.3390/medicina57050503.
- [17] R. Yacouby and D. Axman, "Probabilistic Extension of Precision, Recall, and F1 Score for More Thorough Evaluation of Classification Models," in *Proceedings of the First Workshop on Evaluation and Comparison of NLP Systems*, Online: Association for Computational Linguistics, 2020, pp. 79–91. doi: 10.18653/v1/2020.eval4nlp-1.9.
- [18] Manuel Molina, "Un intruso de otro mundo: F1-score.," *Revista Electrónica AnestesiaR*, vol. 16, no. 4, May 2024, doi: 10.30445/rear.v16i4.1258.
- [19] N. G. Ramadhan, F. D. Adhinata, A. J. T. Segara, and D. P. Rakhmadani, "Deteksi Berita Palsu Menggunakan Metode Random Forest dan Logistic Regression," *Jur. Ris. Kom.*, vol. 9, no. 2, p. 251, Apr. 2022, doi: 10.30865/jurikom.v9i2.3979.
- [20] Tushar Agarwal, Jitender Jangid, Gaurav Kumar, "Transformer and Natural language processing; A recent development.," *tjjpt*, vol. 44, no. 1, pp. 140–143, Nov. 2023, doi: 10.52783/tjjpt.v44.i1.2225.
- [21] B. Juarto, "Classification of Hoax News Using Machine Learning and Neural Networks with BERT Embeddings," in *2023 3rd International Conference on Electronic and Electrical Engineering and Intelligent System (ICE3IS)*, Yogyakarta, Indonesia: IEEE, Aug. 2023, pp. 311–316. doi: 10.1109/ICE3IS59323.2023.10335413.
- [22] R. Qubra and R. A. Saputra, "Classification of Hoax News Using the Naïve Bayes Method," *ijsecs*, vol. 4, no. 1, pp. 40–48, Apr. 2024, doi: 10.35870/ijsecs.v4i1.2068.
- [23] F. Ferawaty, M. Manullang, and R. Hanitio, "Fake News Classification Using SVM and Logistic Regression Methods," in *2024 2nd International Conference on Technology Innovation and Its Applications (ICTIIA)*, Medan, Indonesia: IEEE, Sept. 2024, pp. 1–5. doi: 10.1109/ICTIIA61827.2024.10761435.