

Analisis Sentimen Berbasis Jaringan LSTM dan BERT terhadap Diskusi Twitter tentang Pemilu 2024

Muammar Khadapi¹, Victor Maruli Pakpahan²

STMIK KAPUTAMA, Jl. Veteran No.4A, Tangsi, Kec. Binjai Kota, Kota Binjai, Sumatera Utara^{1,2}
khdafi5@gmail.com

Abstrak. Pemilihan Umum (Pemilu) merupakan peristiwa politik penting yang memicu banyak diskusi di media sosial, terutama di platform seperti Twitter. Analisis sentimen dari diskusi ini dapat memberikan wawasan mengenai pandangan masyarakat terhadap calon, partai, serta isu-isu yang terkait. Penelitian ini berfokus pada penerapan dua model deep learning, yaitu Long Short-Term Memory (LSTM) dan Bidirectional Encoder Representations from Transformers (BERT), untuk menganalisis sentimen diskusi Twitter tentang Pemilu 2024. Kedua model ini dipilih karena kemampuan mereka dalam menangani data teks yang kompleks dan konteks bahasa alami. Dataset yang digunakan dalam penelitian ini terdiri dari ribuan tweet terkait Pemilu 2024, yang diklasifikasikan ke dalam tiga kategori sentimen, yaitu positif, negatif, dan netral. Data terlebih dahulu diproses melalui tahap pembersihan teks dan tokenisasi. Model LSTM dan BERT dilatih menggunakan dataset ini untuk memprediksi sentimen dengan fokus pada peningkatan akurasi prediksi. Hasil eksperimen menunjukkan bahwa model BERT secara konsisten memberikan performa yang lebih baik dibandingkan dengan LSTM. Model BERT berhasil mencapai akurasi validasi sebesar 76.48% pada epoch kedua, sedangkan model LSTM hanya mencapai akurasi maksimal 87 %. Meskipun demikian, model BERT mulai menunjukkan gejala overfitting pada epoch ketiga, dengan peningkatan nilai loss pada data validasi. Hal ini menunjukkan bahwa tuning lebih lanjut pada hyperparameter seperti jumlah epoch dan learning rate diperlukan untuk meningkatkan generalisasi model. Sementara itu, model LSTM menunjukkan stabilitas yang lebih baik, meskipun akurasinya lebih rendah, terutama dalam menangani dependensi konteks yang lebih sederhana. Secara keseluruhan, penelitian ini menegaskan bahwa model BERT lebih efektif dalam menangkap konteks kompleks pada teks Twitter terkait Pemilu 2024 dibandingkan dengan LSTM. Namun, tantangan seperti overfitting dan optimasi hyperparameter tetap menjadi perhatian utama. Untuk meningkatkan performa lebih lanjut, perlu dipertimbangkan teknik augmentasi data dan tuning hyperparameter yang lebih optimal. Penelitian ini juga membuka peluang untuk pengembangan model hibrida yang menggabungkan keunggulan LSTM dan BERT dalam analisis sentimen berbasis teks.

Kata Kunci : Analisis Sentimen, BERT, Deep Learning, Hyperparameter Tuning, LSTM, Overfitting, Pemilu 2024, Twitter.

Abstract. General Elections (Pemilu) are important political events that trigger a lot of discussions on social media, especially on platforms like Twitter. Sentiment analysis of these discussions can provide insights into people's views on candidates, parties, and related issues. This study focuses on the application of two deep learning models, namely Long Short-Term Memory (LSTM) and Bidirectional Encoder Representations from Transformers (BERT), to analyze the sentiment of Twitter discussions about the 2024 Election. These two models were chosen because of their ability to handle complex text data and natural language context. The dataset used in this study consists of thousands of tweets related to the 2024 Election, which are classified into three sentiment categories, namely positive, negative, and neutral. The data is first processed through text cleaning and tokenization stages. LSTM and BERT models are trained using this dataset to predict sentiment with a focus on improving prediction accuracy. Experimental results show that the BERT model consistently performs better than LSTM. The BERT model successfully achieved a validation accuracy of 76.48% in the second epoch, while the LSTM model only achieved a maximum accuracy of 87%. However, the BERT model began to show signs of overfitting in the third epoch, with an increase in the loss value on the validation data. This suggests that further tuning of hyperparameters such as the number of epochs and learning rate is needed to improve model generalization. Meanwhile, the LSTM model showed better stability, despite its lower accuracy, especially in handling simpler context dependencies. Overall, this study confirms that the BERT model is more effective in capturing complex contexts in Twitter texts related to the 2024 Election compared to LSTM. However, challenges such as overfitting and hyperparameter optimization remain major concerns. To further improve performance, data augmentation techniques and more optimal hyperparameter tuning need



to be considered. This study also opens up opportunities for the development of hybrid models that combine the advantages of LSTM and BERT in text-based sentiment analysis.

Keyword : 2024 Election, BERT, Deep Learning, Hyperparameter Tuning, LSTM, Overfitting, Sentiment Analysis, Twitter.

PENDAHULUAN

Pemilihan Presiden 2024 (Pilpres 2024) menjadi salah satu topik politik yang paling diperbincangkan, terutama dalam era digital saat ini. *Twitter*, sebagai salah satu platform media sosial terpopuler, memunculkan beragam opini dan diskusi terkait Pilpres 2024. Dalam konteks ini, analisis sentimen menjadi teknik yang semakin berkembang, memungkinkan peneliti untuk memahami dan mengukur opini serta perasaan yang terkandung dalam teks secara otomatis. Hal ini sangat penting untuk memperoleh gambaran yang jelas mengenai pandangan publik terhadap isu-isu yang berkaitan dengan pemilihan umum[1],[2].

Pengolahan data opini atau *tweet* di *Twitter* memerlukan pendekatan komputasional yang efisien untuk mengidentifikasi dan menganalisis sentimen yang terdapat dalam tweet-tweet tersebut. Oleh karena itu, penelitian ini bertujuan untuk mengidentifikasi beberapa permasalahan yang dihadapi dalam analisis sentimen, termasuk:

1. Volume dan Keanekaragaman Data: Jumlah tweet yang dihasilkan setiap harinya sangat besar, dengan variasi bahasa, gaya, dan konteks yang luas[3],[4].
2. Sentimen yang Beragam: Opini publik mengenai Pemilu 2024 dapat berkisar dari dukungan penuh hingga ketidakpuasan dan keberatan[5].
3. Kesulitan dalam Pemrosesan Bahasa Alami: Memahami dan menganalisis bahasa manusia dalam konteks yang beragam merupakan tantangan utama dalam pemrosesan bahasa alami (*Natural Language Processing/NLP*)[6].

Dalam menghadapi tantangan ini, penelitian ini menggunakan jaringan LSTM (*Long Short-Term Memory*) dan model BERT (*Bidirectional Encoder Representations from Transformers*) sebagai alat analisis sentimen yang canggih. Jaringan LSTM efektif untuk memodelkan urutan data, seperti teks dalam *tweet*, sementara BERT dapat menangani konteks yang lebih kompleks dan mendalam dalam pemahaman teks[7],[8].

Pendekatan yang diambil dalam penelitian ini mencakup penggunaan dua model NLP yang saling melengkapi:

1. Penggunaan Jaringan LSTM: LSTM mampu "mengingat" informasi dalam jangka panjang, sehingga dapat mengidentifikasi pola sentimen dalam teks[9].
2. Penggunaan Model BERT: Model BERT digunakan untuk memahami konteks dan hubungan antara kata-kata dalam teks secara lebih mendalam, yang telah terbukti sangat efektif dalam NLP[10].

Dengan menerapkan kedua model ini, analisis sentimen diharapkan dapat mengidentifikasi pola atau faktor yang memengaruhi opini publik, serta memberikan rekomendasi kebijakan dan strategi komunikasi yang lebih efektif kepada para pemangku kepentingan terkait Pemilu 2024. Klasifikasi sentimen dalam penelitian ini mencakup:

1. Sentimen Positif: Tweet yang menunjukkan dukungan, kepuasan, atau pandangan positif terhadap Pemilu 2024.
2. Sentimen Negatif: Tweet yang mencerminkan ketidakpuasan, keberatan, atau pandangan negatif terhadap Pemilu 2024.
3. Sentimen Netral: Tweet yang bersifat informatif tanpa menunjukkan sikap positif atau negatif.

Metodologi yang diterapkan dalam penelitian ini meliputi beberapa langkah kunci:

1. Pengumpulan Data: Sumber data berasal dari tweet yang mengandung *hashtag* atau kata kunci terkait Pemilu 2024.



2. Preprocessing Data: Data yang terkumpul akan dibersihkan dari karakter khusus, diubah menjadi huruf kecil, dihapus dari stopwords, dan dilakukan tokenisasi untuk mempersiapkan data untuk analisis.
3. Pemodelan dengan LSTM dan BERT: LSTM digunakan untuk memodelkan urutan data, sedangkan BERT untuk memahami konteks teks.
4. Analisis Sentimen: Melakukan analisis untuk mengidentifikasi pola-pola sentimen yang dominan dalam diskusi Twitter.
5. Interpretasi dan Evaluasi: Hasil analisis akan dievaluasi dan diinterpretasikan, dengan menghasilkan kesimpulan yang memberikan wawasan berharga[11].

Rekomendasi kebijakan dan strategi komunikasi akan disusun berdasarkan pemahaman yang lebih baik tentang sentimen masyarakat dan faktor-faktor yang memengaruhi opini publik dalam diskusi di Twitter[10].

METODOLOGI PENELITIAN

Metodologi penelitian ini mengacu pada tahapan CRISP-DM yang diawali dengan *Business Understanding*. Pada tahap ini, peneliti mengidentifikasi tujuan analisis sentimen yang ingin dicapai, seperti memahami pola sentimen masyarakat terkait Pemilu 2024, faktor-faktor yang memengaruhi sentimen tersebut, serta menghasilkan rekomendasi kebijakan yang relevan. Dengan menetapkan tujuan yang jelas, penelitian dapat diarahkan untuk memberikan wawasan yang berharga bagi para pemangku kepentingan.

Setelah tujuan ditetapkan, penelitian berlanjut ke tahap *Data Understanding*. Di sini, peneliti mengumpulkan data dari platform *Twitter* dengan menggunakan kata kunci yang relevan untuk diskusi seputar Pemilu 2024. Analisis dilakukan untuk memahami karakteristik data, termasuk volume *tweet*, variasi bahasa yang digunakan, dan konteks diskusi. Peneliti juga melakukan analisis sentimen awal untuk mendapatkan gambaran umum mengenai sentimen yang berkembang di kalangan pengguna Twitter.

Setelah memahami data, peneliti melanjutkan ke tahap *Data Preparation*, di mana data yang telah dikumpulkan disiapkan untuk analisis lebih lanjut. Langkah-langkah yang dilakukan mencakup pemilihan atribut yang relevan, transformasi data, pembersihan dataset dari noise, serta penyaringan data untuk memastikan relevansi informasi. Teknik pengolahan teks, seperti *casefolding*, *stemming*, dan penghapusan *stopword*, diterapkan untuk mempersiapkan data teks. Proses ini diakhiri dengan pembagian dataset menjadi set pelatihan dan pengujian untuk pelatihan model.

Pada tahap *Modeling*, peneliti menerapkan dua model utama, yaitu LSTM dan BERT, yang digunakan untuk menganalisis sentimen dari teks yang telah diproses. Model LSTM dipilih karena kemampuannya dalam menangkap pola sentimen dengan mempertahankan informasi jangka panjang, sementara BERT digunakan untuk memahami konteks dan hubungan antar kata dengan lebih mendalam. Kedua model dilatih menggunakan data pelatihan dan dievaluasi menggunakan set pengujian untuk mengukur kinerja dan akurasi mereka dalam klasifikasi sentimen.

Terakhir, pada tahap *Evaluation*, peneliti melakukan pengukuran performa model menggunakan metrik seperti akurasi dan F1-score. Model divalidasi dengan data baru untuk memastikan bahwa hasil analisis dapat diandalkan. Setelah mengevaluasi kedua model, hasil analisis disajikan dalam bentuk wawasan yang mencakup tren sentimen, faktor-faktor yang memengaruhi, serta rekomendasi kebijakan dan strategi komunikasi untuk para pemangku kepentingan terkait Pemilu 2024. Dengan demikian, penelitian ini diharapkan dapat memberikan kontribusi signifikan dalam pemahaman opini publik dan mendukung pengambilan keputusan yang lebih baik.



HASIL DAN PEMBAHASAN

Dalam penelitian ini, kami berhasil melakukan analisis sentimen terhadap opini publik mengenai Pemilu 2024 dengan menggunakan dua pendekatan berbasis model: LSTM dan BERT. Dataset yang digunakan terdiri dari 4.376 data tweet yang dikumpulkan dari Twitter tentang Pemilu 2024, dengan memperhatikan kata kunci yang telah ditentukan, yaitu "pemilu 2024", "Sirekap", dan "tolak pemilu curang" dalam periode 7 maret 2024 hingga 8 maret 2024. Setelah proses pembersihan dan praproses data, kami mendapatkan dataset akhir yang siap untuk dianalisis

Hasil Kinerja Model LSTM

Model Long Short-Term Memory (LSTM) yang diimplementasikan dalam penelitian ini menghasilkan performa yang signifikan dalam menyelesaikan tugas klasifikasi sentimen. Pada tahap pelatihan, model berhasil meningkatkan akurasi secara signifikan seiring bertambahnya epoch. Di *epoch* pertama, model mencatatkan akurasi 41.73% dengan *loss* 0.9547. Peningkatan yang mencolok terlihat pada *epoch* kedua, di mana akurasi meningkat menjadi 80.99% dengan *loss* berkurang menjadi 0.5036. Ini menunjukkan bahwa model mulai belajar dari pola yang terdapat dalam data. Puncak performa dicapai pada *epoch* kedelapan, di mana akurasi pelatihan mencapai 99.93% dengan *loss* yang sangat rendah, yaitu 0.0036.

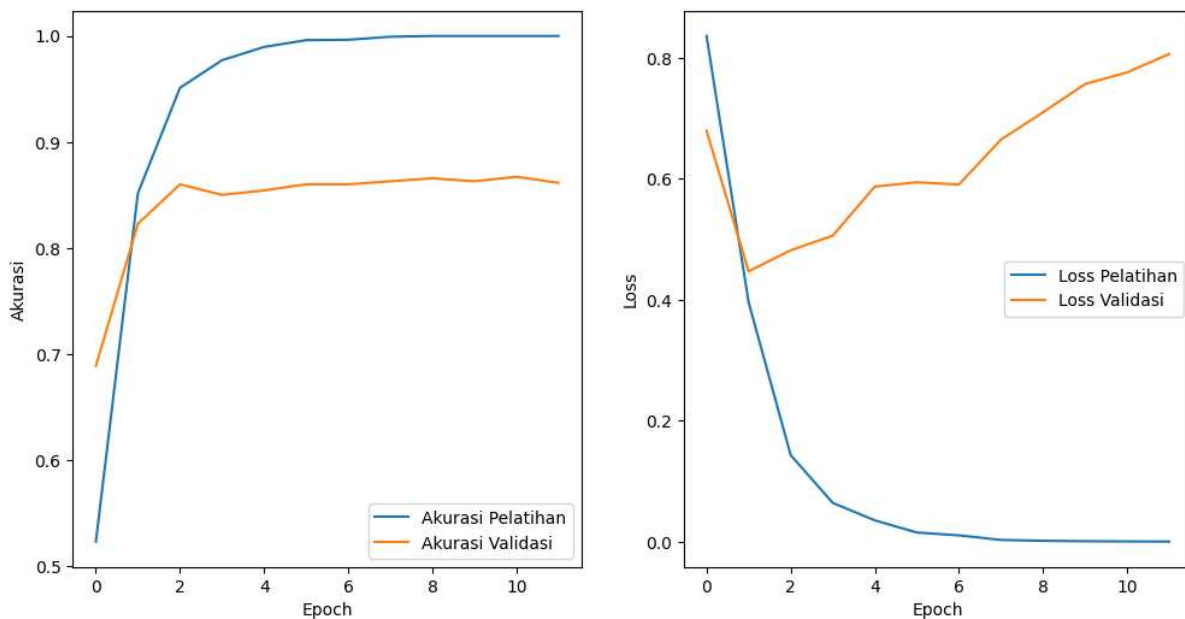
Meskipun hasil pelatihan menunjukkan performa yang sangat baik, model LSTM mengalami fluktuasi pada akurasi validasi. Akurasi validasi tertinggi diperoleh pada *epoch* kesembilan, yaitu 86.57%, sementara nilai *loss* validasi menunjukkan pola yang bervariasi. Pada akhir pelatihan, model diuji pada dataset yang terpisah dan menghasilkan akurasi 85.16%. Ini menunjukkan bahwa meskipun model sangat baik dalam mengenali pola dalam data pelatihan, ada tantangan dalam generalisasi terhadap data baru.

Analisis lebih mendalam terhadap performa model menunjukkan nilai *precision*, *recall*, dan F1-score yang bervariasi untuk masing-masing kelas sentimen. Kelas negatif, misalnya, memiliki *precision* 0.89, *recall* 0.87, dan F1-score 0.88, yang menandakan bahwa model cukup efektif dalam mengidentifikasi sentimen negatif. Untuk kelas netral, nilai *precision* tercatat di angka 0.77 dengan *recall* 0.86, menghasilkan F1-score sebesar 0.82. Kelas positif juga menunjukkan hasil yang baik, dengan *precision* 0.83, *recall* 0.80, dan F1-score 0.81. Secara keseluruhan, model LSTM mencapai akurasi 85% dan nilai rata-rata F1-score di atas 0.80 untuk semua kelas, yang mencerminkan kinerjanya yang baik.

Namun, meskipun hasil yang dicapai cukup menjanjikan, penelitian ini masih menyisakan ruang untuk pengoptimalan lebih lanjut. Baik dari sisi arsitektur model maupun *hyperparameter tuning*, masih terdapat peluang untuk meningkatkan performa model, terutama pada data yang lebih bervariasi. Penggunaan teknik regularisasi, seperti dropout, dan eksplorasi berbagai konfigurasi hyperparameter dapat membantu memperbaiki kestabilan dan akurasi model dalam menghadapi tantangan yang lebih kompleks.

Ke depan, penelitian lebih lanjut dapat mengeksplorasi penggunaan model yang lebih kompleks atau hybrid, yang menggabungkan keunggulan LSTM dengan teknik pembelajaran mendalam lainnya. Pendekatan ini diharapkan mampu mencapai hasil yang lebih baik dalam klasifikasi sentimen. Selain itu, integrasi model dengan metode lain, seperti BERT, dapat dijajaki untuk memanfaatkan kekuatan masing-masing model dalam pengolahan bahasa alami, sehingga menghasilkan model yang lebih robust dan efisien.





Gambar 1. Grafik Hasil Pelatihan Model LSTM

Grafik akurasi dan loss dari model LSTM memberikan wawasan yang mendalam mengenai performa model dalam tugas klasifikasi sentimen. Pada grafik akurasi, terdapat dua garis yang menunjukkan akurasi pelatihan dan akurasi validasi. Akurasi pelatihan, yang ditunjukkan oleh garis biru, menunjukkan peningkatan yang sangat cepat pada epoch pertama dan kedua, mencapai hampir 100% pada epoch keenam dan seterusnya setelah stabil di sekitar epoch keempat. Sebaliknya, akurasi validasi, yang ditunjukkan oleh garis oranye, awalnya mengalami kenaikan signifikan, namun stagnasi terjadi pada rentang epoch kedua hingga keempat, dengan nilai akurasi yang tetap berada di kisaran 85% hingga 87%. Temuan ini menunjukkan bahwa meskipun model berhasil mempelajari pola dari data pelatihan dengan baik, perbedaan yang mencolok antara akurasi pelatihan dan akurasi validasi menandakan adanya fenomena *overfitting*, di mana model terlalu beradaptasi dengan data pelatihan dan tidak mampu mengeneralisasi pada data baru.

Selanjutnya, analisis grafik loss memperkuat temuan tersebut. Loss pelatihan, yang ditunjukkan oleh garis biru, mengalami penurunan signifikan, terutama pada beberapa epoch pertama, dan hampir mencapai nol pada epoch keenam. Hal ini mengindikasikan bahwa model berhasil memahami pola dalam data pelatihan secara efektif. Namun, loss validasi yang ditunjukkan oleh garis oranye menunjukkan penurunan di awal, tetapi mulai meningkat secara bertahap setelah epoch ketiga. Peningkatan loss validasi ini mengindikasikan bahwa meskipun model terus belajar dari data pelatihan, performanya pada data validasi mulai menurun. Kondisi ini merupakan indikasi lebih lanjut dari *overfitting*, di mana model yang terlalu kompleks mungkin tidak mampu mempertahankan generalisasi yang baik pada data yang belum pernah dilihat sebelumnya.

Hasil Kinerja Model BERT

Hasil penelitian ini menunjukkan performa model BERT yang dilatih selama tiga epoch untuk klasifikasi sentimen pada data diskusi Twitter terkait Pemilu 2024.



```

Epoch 1/3
219/219 [=====] - 4766s 22s/step - loss:
0.8944 - accuracy: 0.6043 - val_loss: 0.5992 - val_accuracy: 0.7534
Epoch 2/3
219/219 [=====] - 4751s 22s/step - loss:
0.8153 - accuracy: 0.6329 - val_loss: 0.5923 - val_accuracy: 0.7648
Epoch 3/3
219/219 [=====] - 4705s 21s/step - loss:
0.8286 - accuracy: 0.6231 - val_loss: 0.8174 - val_accuracy: 0.7511

```

Gambar 2. Hasil Pelatihan Data Menggunakan BERT

Pada epoch pertama, model mencapai akurasi pelatihan sebesar 60,43% dengan nilai loss yang tinggi, yaitu 0,8944. Sementara itu, akurasi validasi mencapai 75,34%, yang menunjukkan bahwa meskipun model baru mulai mempelajari representasi dari data, BERT mampu menangkap konteks dari teks yang diberikan. Hal ini menegaskan salah satu kekuatan utama arsitektur berbasis *transformer*, yaitu kemampuannya dalam memahami pola bahasa alami.

Pada epoch kedua, terdapat peningkatan yang signifikan dalam kinerja model, di mana akurasi pelatihan meningkat menjadi 63,29% dan nilai loss menurun menjadi 0,8153. Selain itu, akurasi validasi juga meningkat menjadi 76,48% dengan nilai loss validasi yang turun menjadi 0,5923. Peningkatan ini menunjukkan bahwa model semakin mampu menggeneralisasi pola dari data, dengan BERT secara efektif menangani dependensi konteks antar kata dalam kalimat. Hal ini penting karena menunjukkan bahwa model tidak hanya mengingat data pelatihan, tetapi juga dapat memprediksi sentimen pada data yang belum pernah dilihat sebelumnya.

Namun, pada epoch ketiga, hasil yang dicapai tidak sebaik dua epoch sebelumnya. Akurasi pelatihan sedikit menurun menjadi 62,31% dan nilai loss pelatihan meningkat menjadi 0,8286. Yang lebih signifikan, loss validasi juga meningkat menjadi 0,8174, diiringi dengan penurunan akurasi validasi menjadi 75,11%. Fenomena ini mungkin menunjukkan awal dari *overfitting*, di mana model terlalu menyesuaikan diri dengan data pelatihan sehingga kinerjanya menurun pada data validasi yang tidak dikenal. *Overfitting* ini sering terjadi pada model yang kompleks seperti BERT, terutama jika jumlah epoch terlalu banyak atau data pelatihan tidak cukup beragam.

```

# Evaluasi model
loss, accuracy = model.evaluate(inputs['input_ids'], labels)
print(f'Accuracy: {accuracy}')
137/137 [=====] - 1639s 12s/step - loss: 0.9069 -
accuracy: 0.6453
Accuracy: 0.645338237285614

```

Gambar 2. Hasil Evaluasi Model BERT

Hasil evaluasi model setelah proses pelatihan menunjukkan akurasi yang dicapai adalah sebesar 64,53% dengan nilai loss sebesar 0,9069 pada data uji. Meskipun hasil ini mencerminkan kemampuan model dalam mengklasifikasikan sentimen dengan benar, nilai akurasi ini menunjukkan bahwa masih terdapat ruang untuk perbaikan dalam kinerja model. Nilai loss yang tinggi mengindikasikan bahwa model belum cukup optimal dalam menyesuaikan parameter-parameter jaringan untuk menghasilkan prediksi yang akurat. Analisis lebih lanjut menunjukkan bahwa keterbatasan dalam jumlah data pelatihan dan kompleksitas data teks dapat berkontribusi terhadap kinerja yang kurang optimal.

Untuk melanjutkan perbaikan, beberapa pendekatan yang dapat diambil termasuk eksperimen dengan tuning hyperparameter seperti learning rate, ukuran batch, dan jumlah epoch pelatihan.



Penerapan teknik augmentasi data dan fine-tuning model lebih lanjut dengan dataset yang lebih besar dan beragam juga sangat dianjurkan. Selain itu, langkah-langkah seperti *cross-validation* dapat diterapkan untuk memastikan bahwa model tidak mengalami *overfitting* pada subset tertentu dari data pelatihan. Dalam analisis sentimen yang kompleks, kombinasi model seperti BERT dengan teknik lain, seperti mekanisme perhatian yang lebih khusus, juga dapat dipertimbangkan untuk meningkatkan kemampuan model dalam menangani data yang lebih kompleks.

KESIMPULAN

Dalam penelitian ini, algoritma LSTM dan BERT telah dievaluasi dalam konteks klasifikasi sentimen, masing-masing dengan keunggulan dan tantangan yang berbeda. LSTM, sebagai model yang dirancang untuk mengatasi urutan data, menunjukkan kemampuannya dalam menangkap pola temporal dan dependensi jangka panjang dalam teks. Hasil pelatihan dan evaluasi model LSTM menunjukkan akurasi yang kompetitif, namun juga dihadapkan pada tantangan seperti *overfitting* ketika tidak diimbangi dengan data yang cukup beragam dan teknik regularisasi yang tepat.

Di sisi lain, model BERT, yang berbasis pada arsitektur transformer, telah terbukti sangat efektif dalam memahami konteks yang kompleks dari data teks berbahasa alami. Dengan kemampuannya untuk menangkap representasi kontekstual dari kata-kata, BERT menunjukkan performa yang lebih baik dalam klasifikasi sentimen pada data diskusi Twitter terkait Pemilu 2024. Meskipun hasil evaluasi BERT menunjukkan akurasi yang lebih tinggi dibandingkan dengan LSTM, model ini juga mengalami tanda-tanda *overfitting* pada epoch ketiga, yang menekankan pentingnya pengaturan hyperparameter dan teknik pengendalian model yang tepat.

Perbandingan antara kedua model menunjukkan bahwa meskipun BERT memiliki keunggulan dalam hal akurasi dan pemahaman konteks, LSTM tetap relevan dalam aplikasi di mana data urutan dan temporal menjadi fokus utama. Selain itu, hasil dari kedua model menyoroti perlunya eksplorasi lebih lanjut terkait teknik augmentasi data, *fine-tuning hyperparameter*, dan strategi regularisasi untuk meningkatkan kinerja model.

Secara keseluruhan, penelitian ini menegaskan bahwa baik LSTM maupun BERT memiliki potensi yang signifikan dalam analisis sentimen, tetapi setiap model perlu disesuaikan dengan karakteristik data dan tujuan spesifik dari aplikasi yang diinginkan. Melalui pengembangan lebih lanjut dan integrasi teknik yang sesuai, kedua model ini dapat terus berkontribusi dalam peningkatan akurasi klasifikasi sentimen serta aplikasi dalam berbagai bidang lain yang memerlukan pemrosesan bahasa alami.

DAFTAR PUSTAKA

- [1] M. Khadapi, A. M. H. Pardede, and N. Novriyenni, "Providing Recommendations to New ILMCI Edu Voucher Customers Using the Market Basket Analyst Algorithm," in *2023 International Conference of Computer Science and Information Technology (ICOSNIKOM)*, 2023, pp. 1–6. doi: 10.1109/ICoSNiKOM60230.2023.10364536.
- [2] R. Vindua and A. U. Zailani, "Analisis Sentimen Pemilu Indonesia Tahun 2024 Dari Media Sosial Twitter Menggunakan Python," *JURIKOM (Jurnal Ris. Komputer)*, vol. 10, no. 2, p. 479, 2023, doi: 10.30865/jurikom.v10i2.5945.
- [3] Y. Cao, Z. Sun, L. Li, and W. Mo, "A Study of Sentiment Analysis Algorithms for Agricultural Product Reviews Based on Improved BERT Model," *Symmetry (Basel)*, vol. 14, no. 8, 2022, doi: 10.3390/sym14081604.
- [4] A. S. Talaat, "Sentiment analysis classification system using hybrid BERT models," *J. Big Data*, vol. 10, no. 1, 2023, doi: 10.1186/s40537-023-00781-w.
- [5] A. C. M. V. Srinivas, C. Satyanarayana, C. Divakar, and K. P. Sirisha, "Sentiment Analysis



- using Neural Network and LSTM,” *IOP Conf. Ser. Mater. Sci. Eng.*, vol. 1074, no. 1, p. 012007, 2021, doi: 10.1088/1757-899x/1074/1/012007.
- [6] A. Kurniasih and L. P. Manik, “On the Role of Text Preprocessing in BERT Embedding-based DNNs for Classifying Informal Texts,” *Int. J. Adv. Comput. Sci. Appl.*, vol. 13, no. 6, pp. 927–934, 2022, doi: 10.14569/IJACSA.2022.01306109.
- [7] G. S. . Murthy, S. R. Allu, B. Andhavarapu, M. Bgadi, and M. Belusonti, “Text based Sentiment Analysis using Long Short Term Memory (LSTM),” *Int. J. Eng. Res. Technol.*, vol. 9, no. 05, pp. 299–303, 2020.
- [8] M. P. Geetha and D. Karthika Renuka, “Improving the performance of aspect based sentiment analysis using fine-tuned Bert Base Uncased model,” *Int. J. Intell. Networks*, vol. 2, no. March, pp. 64–69, 2021, doi: 10.1016/j.ijin.2021.06.005.
- [9] B. BİLEN and F. HORASAN, “LSTM Network based Sentiment Analysis for Customer Reviews,” *Politek. Derg.*, vol. 25, no. 3, pp. 959–966, 2022, doi: 10.2339/politeknik.844019.
- [10] A. Bello, S. C. Ng, and M. F. Leung, “A BERT Framework to Sentiment Analysis of Tweets,” *Sensors*, vol. 23, no. 1, 2023, doi: 10.3390/s23010506.
- [11] A. M. Mantika, A. Triayudi, and R. T. Aldisa, “Sentiment Analysis on Twitter Using Naïve Bayes and Logistic Regression for the 2024 Presidential Election,” vol. 2, no. 1, pp. 44–55, 2024.

