



Model Deteksi Berita Palsu Menggunakan Pendekatan Bidirectional Long Short-Term Memory (BiLSTM)

Ari Muzakir¹, Uci Suryani^{2*}

¹Program Studi Teknik Informatika, Fakultas Sains dan Teknologi, Universitas Bina Darma, Palembang, Indonesia

²Program Studi Teknologi Rekayasa Multimedia, Politeknik Darussalam, Palembang, Indonesia

Email: ¹arimuzakir@binadarma.ac.id, ²uci.suryani1@gmail.com

Abstrak

Isu berita palsu telah menarik perhatian masyarakat dan akademisi. Penyebaran informasi yang tidak akurat berpotensi mengubah pandangan publik dan memungkinkan manipulasi opini. Dalam konteks data yang melimpah, kami mengembangkan model untuk mendeteksi berita palsu dengan mengklasifikasikan fitur linguistik murni. Dengan pendekatan pembelajaran mendalam, kami mengevaluasi respons terhadap artikel tertentu menggunakan model Jaringan Saraf Rekuren *Bidirectional Long Short-Term Memory* (BiLSTM) dan representasi kata dari embeddings GloVe. Hasil evaluasi menunjukkan adaptabilitas model pada data latih dengan kerugian terendah 0.30% dan akurasi tinggi 99.14%. Gabungan antara embeddings GloVe dan Bi-LSTM memunculkan hasil yang positif. Penelitian ini memiliki potensi untuk memberikan kontribusi dalam penanggulangan penyebaran berita palsu yang semakin meresahkan di berbagai bidang.

Kata Kunci: Berita palsu, Embeddings GloVe, *Bidirectional Long Short-Term Memory* (BiLSTM)

1. PENDAHULUAN

Penyebaran berita palsu di platform media sosial telah mengalami lonjakan perhatian yang signifikan karena situasi politik akhir-akhir ini dan meningkatnya kekhawatiran terhadap dampak negatif yang mungkin timbul. Sebagai contoh, pada bulan Januari 2017, juru bicara pemerintah Jerman mengungkapkan bahwa mereka menghadapi suatu fenomena yang belum pernah mereka alami sebelumnya, merujuk pada penyebaran berita palsu [3]. Lebih dari sekadar menjadi sumber informasi yang mengganggu



dalam kehidupan kita, berita palsu juga berpotensi untuk memanipulasi persepsi dan kesadaran publik secara massal.

Mendeteksi berita palsu merupakan tantangan kritis dalam bidang *Natural Language Processing* (NLP). Pertumbuhan pesat platform jejaring sosial tidak hanya menghasilkan peningkatan besar dalam akses informasi, tetapi juga mempercepat penyebaran berita palsu [1]. Fenomena berita palsu, desas-desus, penyebaran informasi yang keliru, dan upaya mendeteksi informasi yang menyesatkan saat ini menjadi perhatian utama karena potensinya untuk merusak tatanan sosial kita [2]. Kehadiran berita saat ini dapat menyebar dengan kecepatan tinggi melalui internet. Karena berita palsu dirancang untuk menarik perhatian pembaca, mereka cenderung menyebar dengan cepat. Bagi banyak pembaca, mengenali berita palsu bisa menjadi tugas yang rumit, dan pembaca semacam itu seringkali akhirnya meyakini bahwa berita palsu tersebut benar adanya [3].

Berita palsu memiliki dampak yang sangat merugikan terutama pada individu yang cenderung mudah percaya, bahkan bisa merusak reputasi seseorang secara permanen. Berita palsu memiliki potensi untuk memanipulasi pikiran individu atau kelompok dalam jangka waktu yang panjang, tanpa disadari bahwa persepsi masyarakat sedang terbentuk menuju pemahaman yang keliru. Penelitian sebelumnya telah mencatat bahwa media sosial memiliki pengaruh yang signifikan dalam menyebarkan rasa takut dan kepanikan terkait wabah COVID-19, dan ini berpotensi menyebabkan dampak negatif pada kesehatan mental dan psikologis masyarakat. Adanya korelasi statistik yang kuat dan positif antara penggunaan media sosial dan penyebaran kepanikan terkait COVID-19 mengindikasikan bahwa mayoritas individu, terutama remaja dan orang dewasa berusia 18-35 tahun, menghadapi tingkat kecemasan psikologis yang tinggi [4], [5].

Karena data terkait berita palsu telah tersedia dalam jumlah yang besar, membangun model dengan menerapkan algoritme berdasarkan sumber berita menjadi suatu tugas yang relatif mudah. Saat ini, belum ada platform umum yang dikenal luas yang telah menerapkan algoritma yang terfokus pada sumber berita tertentu. Meskipun terdapat metode untuk mendeteksi berita palsu, hingga saat ini belum ada yang berhasil mengintegrasikan berbagai model ini menjadi satu kesatuan yang utuh. Dalam upaya tersebut, beberapa pendekatan telah diusulkan untuk

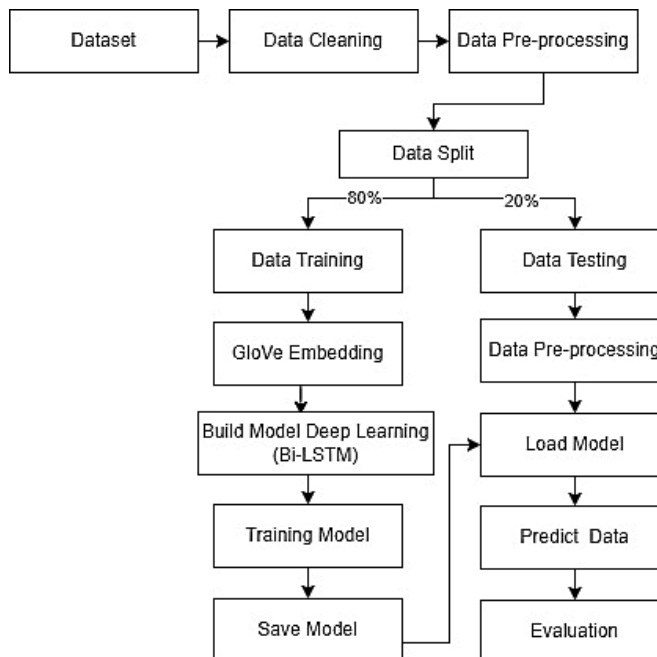
mengotomatisasi deteksi *respons*, sering kali menggambarkan penyebaran berita palsu sebagai epidemi dalam jaringan sosial [6], [7]. Pendekatan lain mencakup penggunaan fitur yang dirancang secara manual, seperti jumlah suka di Facebook, yang diintegrasikan dengan metode klasifikasi konvensional [8]. Namun, terdapat kendala, seperti keterbatasan akses ke data jaringan sosial secara praktis, dan kesulitan dalam memilih fitur secara manual [9].

Dalam penelitian ini, tujuan utama dari penelitian ini yaitu untuk mengatasi isu yang berkaitan dengan penyebaran berita palsu di berbagai kategori umum, termasuk namun tidak terbatas pada politik, media sosial, olahraga, industri media, pendidikan, kejahatan, sains, dan teknologi. Kami melakukan klasifikasi berita palsu berdasarkan fitur-fitur linguistik murni dalam penelitian ini. Untuk meningkatkan ketepatan deteksi berita palsu, kami memanfaatkan tiga elemen kunci yang ada dalam suatu berita: isi teks, tanggapan dari pengguna, dan sumber berita itu sendiri. Kami tidak hanya mengandalkan seleksi fitur manual, melainkan kami mengusulkan suatu model berbasis jaringan saraf yang mendalam berbasis BiLSTM. Model ini memiliki kemampuan untuk secara otomatis mengidentifikasi fitur-fitur yang relevan. Pendekatan jaringan saraf ini juga memungkinkan model untuk memanfaatkan informasi dari berbagai aspek dan domain yang berbeda, serta menangkap hubungan temporal dalam cara pengguna berinteraksi dengan artikel. Dalam rancangan kami, tahap deteksi berita palsu dipecah menjadi beberapa modul yang memerhatikan tiga aspek: aktivitas pengguna, artikel berita, dan respons yang diterima oleh artikel tersebut. Kami mengembangkan model pembelajaran mendalam yang dirancang untuk menilai *respons* yang muncul terhadap artikel tertentu. Model yang dibangun berdasarkan Jaringan Saraf Rekuren (BiLSTM) memanfaatkan informasi spesifik artikel, termasuk jarak temporal antara aktivitas pengguna dan artikel, serta representasi teks menggunakan pendekatan doc2vec [10], yang mencerminkan teks yang dihasilkan dalam interaksi tersebut (seperti cuitan).

2. METODE

Penelitian ini melibatkan dua tahap inti. Tahap pertama merupakan tahap pelatihan, dimulai dengan pengambilan data pelatihan yang tersimpan dalam basis data. Data pelatihan kemudian menjalani proses pembersihan guna menghilangkan data yang tidak berkualitas. Selanjutnya, dilakukan proses augmentasi data untuk menjaga keseimbangan antar kelas. Data

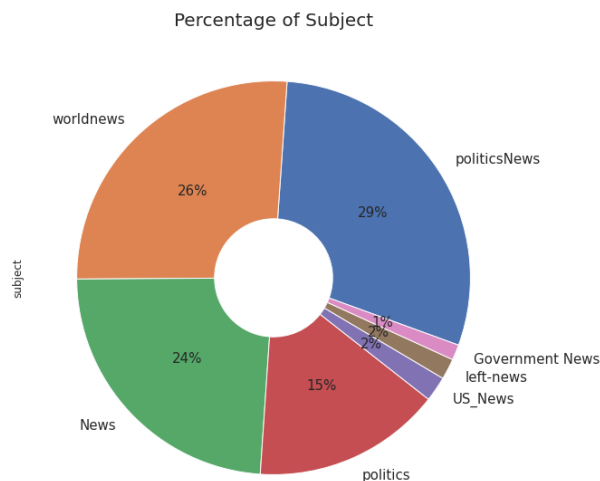
yang telah diperbesar kemudian diolah dan diubah menjadi vektor kata menggunakan penyematan kata yang telah sebelumnya dilatih. Vektor kata yang dihasilkan dari penyematan GloVe yang sudah terlatih digunakan untuk melatih model pembelajaran mendalam BiLSTM. Akhirnya, model yang telah dilatih disimpan dalam basis data guna penggunaan pada tahap pengujian (Gambar 1).



Gambar 1. Tahapan penelitian pada *data training* dan *data testing*

2.1. Dataset

Dataset yang kami gunakan dalam penelitian ini terdiri dari dua kategori data, yaitu berita palsu dan berita asli, dengan distribusi sebagai berikut: berita palsu (23481 baris) dan berita asli (21417 baris). Kumpulan data pelatihan ini mencakup atribut utama seperti judul (judul artikel berita), teks (teks artikel), subjek (subjek artikel berita), dan tanggal (tanggal artikel diposting). Subjek dari artikel beragam, mencakup berita dunia (26%), berita umum (24%), berita politik (29%), politik (15%), serta informasi lainnya (6%) (lihat Gambar 2).



Gambar 2. Distribusi data palsu dan data asli dalam hal subjek

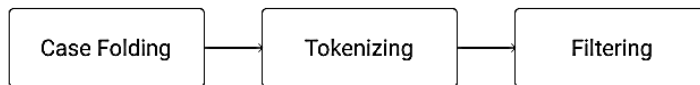
2.2. Data Cleaning

Sumber informasi berita utamanya berasal dari platform media sosial seperti Twitter, yang seringkali memuat kalimat dengan simbol dan karakter yang tidak relevan. Oleh karena itu, kami melakukan tahap pembersihan data dengan mengeliminasi simbol dan karakter tersebut, termasuk URL, tag HTML, tanda baca, dan tanda pagar (#).

2.3. Data Pre-processing

Pra-pemrosesan teks merupakan serangkaian langkah penting untuk membersihkan dan mempersiapkan data teks sebelum dilakukan analisis lebih lanjut. Seperti yang ditunjukkan dalam Gambar 2, beberapa tahap umum yang biasanya digunakan dalam pra-pemrosesan data termasuk: penyesuaian huruf (*case folding*), tokenisasi, dan penyaringan atau penghapusan kata-kata umum (*stop word removing*). Tahap pertama, yaitu "penyesuaian huruf," mengubah semua huruf dalam teks menjadi huruf kecil guna menghindari variasi hasil dari kapitalisasi. Selanjutnya, dalam langkah "tokenisasi," teks dipecah menjadi unit-unit yang lebih kecil seperti kata-kata atau simbol, sehingga mempermudah pemahaman dan analisis. Terakhir, langkah "penyaringan" atau penghapusan "kata berhenti" dilakukan untuk menghilangkan kata-kata umum yang tidak memiliki makna signifikan, sehingga memungkinkan penekanan pada

kata-kata kunci yang relevan dalam teks. Proses-proses ini membantu meningkatkan kualitas data, mengurangi dimensi, dan menyiapkan teks untuk langkah selanjutnya, seperti klasifikasi atau analisis sentimen.



Gambar 2. Tahapan *data pre-processing*

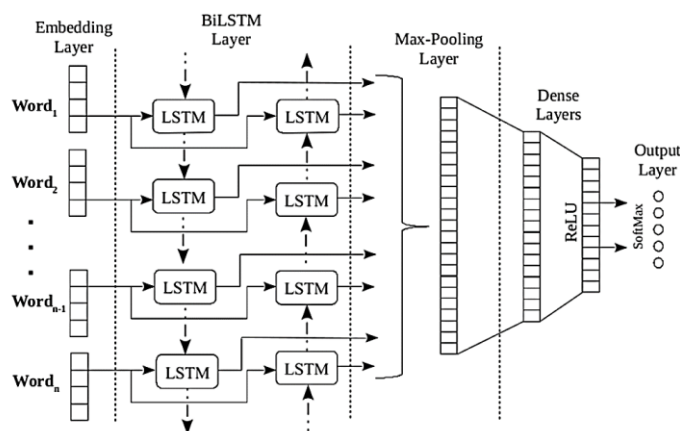
2.4. Data Split

Pembagian data dilakukan menjadi dua bagian: *dataset* pelatihan (80%) dan *dataset* pengujian (20%) yang merupakan pendekatan standar dalam pembelajaran mesin [11]. Dalam pendekatan ini, sebagian besar data digunakan untuk melatih model, memungkinkan model untuk memahami pola dan fitur-fitur dari *dataset* tersebut. Hal ini membantu mengurangi resiko *overfitting* dan memungkinkan model untuk menerapkan apa yang telah dipelajari pada data baru. Di sisi lain, *dataset* pengujian yang terpisah digunakan untuk mengukur kinerja model pada data yang belum pernah dilihat sebelumnya. Dengan memisahkan *dataset* pengujian dari dataset pelatihan, validitas dan kemampuan generalisasi model dapat dievaluasi dengan objektif, sehingga memastikan bahwa model yang dihasilkan mampu berkinerja baik dalam situasi apapun.

2.5. Word Embedding

Penyematan kata atau representasi kata terdistribusi adalah teknik yang mengubah kata-kata menjadi vektor angka, di mana kata-kata yang memiliki makna serupa akan saling dekat dalam representasi visual [12]. Hal ini tercapai dengan memanfaatkan fakta bahwa penyematan kata mampu menangkap aspek semantik dan sintaksis dari kata-kata dalam kumpulan teks yang luas [13]. Metode ini semakin populer dalam berbagai penelitian seperti analisis sentimen, identifikasi entitas, pemberian label pada bagian-bagian kalimat, dan penelitian lain yang berfokus pada analisis teks, karena memberikan hasil yang menjanjikan [14]. Contoh-contoh penyematan kata yang telah dilatih sebelumnya yang umum digunakan termasuk Word2Vec, GloVe, dan fastText.

Dalam penelitian ini, kami memilih menggunakan penyematan kata GloVe (*Global Vectors for Word Representation*). Pendekatan ini menggabungkan statistik matriks kemunculan kata dalam berbagai konteks dengan matriks kemunculan bersama kata-kata dalam korpus teks. *Global Vectors* (GloVe) juga merupakan penyematan kata sebelumnya yang populer, menggabungkan informasi matriks kemunculan kata dalam korpus dengan faktorisasi matriks untuk menghasilkan representasi vektor dari kata-kata [15]. Pendekatan GloVe berawal dari pembentukan matriks kemunculan kata dari data korpus yang luas. Setelahnya, dilakukan proses faktorisasi pada matriks tersebut untuk menghasilkan representasi vektor kata. Dengan memanfaatkan informasi ini, vektor representasi kata yang dihasilkan mampu menangkap hubungan semantik dan sintaksis antara kata-kata. Salah satu keunggulan utama dari GloVe adalah kemampuannya dalam memahami hubungan sinonim, antonim, analogi, serta struktur sintaksis dalam teks.



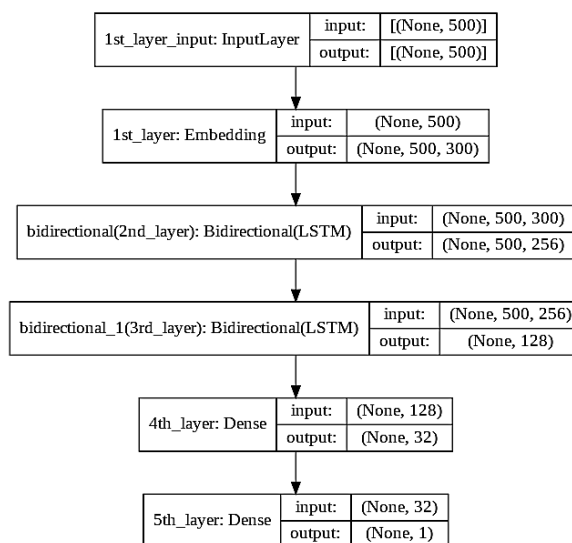
Gambar 3. Jaringan BiLSTM dengan *embeddings Glove*

2.6. Building Model Deep Learning

Untuk membangun model klasifikasi guna mendeteksi berita palsu, kami mengadopsi arsitektur BiLSTM. BiLSTM merupakan bentuk arsitektur saraf berulang [16] yang terdiri dari dua bagian LSTM: satu maju (*forward*) dan satu mundur (*backward*). *Forward* LSTM membaca teks input dalam urutan kronologis, memanfaatkan konteks informasi dari masa lalu. Sebaliknya, *backward* LSTM membaca teks dalam urutan terbalik dan menyimpan informasi kontekstual dari masa depan. Kedua LSTM ini

menghasilkan dua vektor keluaran urutan yang independen. Kami memperoleh vektor keluaran untuk setiap kata dengan menggabungkan vektor maju dan mundur tersebut (lihat Gambar 3).

Deskripsi teks masukan dipecah menjadi *token* dan diberikan ke lapisan *embedding*. Lapisan *embedding* memetakan urutan masukan ke dalam matriks berbentuk $n \times d$, di mana n adalah jumlah kata dalam deskripsi teks, dan d adalah dimensi vektor Glove yang digunakan. Matriks keluaran dari lapisan *embedding* ini kemudian diberikan sebagai input ke dalam lapisan BiLSTM. *Dropout* dengan tingkat 0,2 diterapkan pada lapisan masukan BiLSTM, sedangkan *dropout* dengan tingkat 0,1 digunakan untuk lapisan output BiLSTM. Setelah lapisan BiLSTM, dilakukan lapisan *global max-pooling*, menghasilkan keluaran berbentuk $d \times 1$. Keluaran ini kemudian diteruskan melalui dua lapisan *Dense* dengan aktivasi *Rectifier Linear Unit (ReLU)*. Keluaran dari lapisan padat ini kemudian dilalui melalui lapisan *SoftMax* yang memiliki m unit. Dimensi dari m tergantung pada jumlah kelas dalam tugas yang ada.



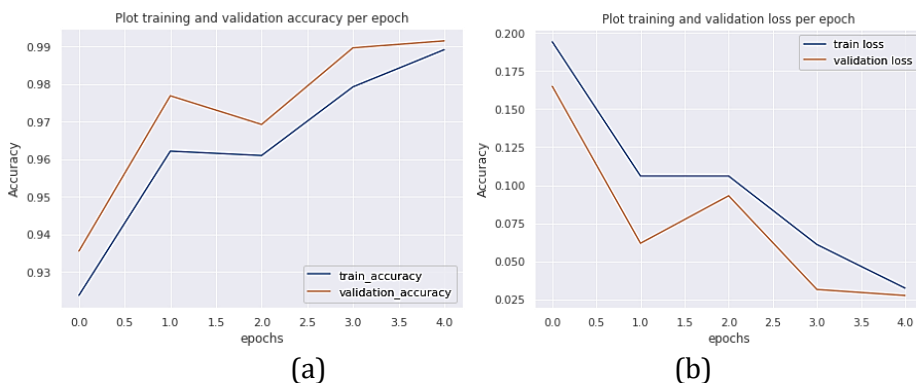
Gambar 4. Arsitektur BiLSTM dengan 5 layer

BiLSTM merupakan varian dari arsitektur *Recurrent Neural Network (RNN)* yang melibatkan dua *Long Short-Term Memory (LSTM)* berorientasi ke arah berlawanan. Tujuan di balik arsitektur ini adalah meningkatkan kapasitas memori LSTM dengan memberikan konteks informasi dari masa

lalu dan masa mendatang [17]. Gambar 4 mengilustrasikan rancangan arsitektur Bidirectional LSTM yang diadopsi dalam penelitian ini, dengan lima lapisan total dan input layer berdimensi 500. Desain arsitektur melibatkan langkah-langkah seperti lapisan BiLSTM, di mana LSTM yang serupa diterapkan pada kedua arah. Dilanjutkan dengan lapisan *global max-pooling*, diikuti oleh lapisan *fully connected*, dan lapisan *output* yang digunakan untuk memprediksi kelas. Pengaturan parameter untuk arsitektur model BiLSTM meliputi *max_features* = 10000, *embedding_vector* = 300, *EPOCHS* = 5, *BATCH_SIZE* = 128, *max_len* = 500, *dropout* = 0.1, serta *learning_rate* = 0.00001.

3. HASIL DAN PEMBAHASAN

Dalam konteks penelitian ini, sebuah model yang menggabungkan *embeddings* GloVe dengan arsitektur jaringan rekuren BiLSTM telah dikembangkan dan dievaluasi untuk tugas klasifikasi. Pencapaian terpenting adalah peningkatan signifikan dalam validasi kinerja model, yang naik dari 0.98% menjadi 0.99%. Ini menunjukkan peningkatan yang substansial dalam akurasi prediksi model, mengindikasikan kemampuannya untuk memahami dan mengklasifikasikan data secara lebih tepat.



Gambar 5. Grafik pengujian. (a) akurasi data *training* dan *validation*. (b) *training* dan *validation loss*

Menariknya, proses perjalanan pada tahap pelatihan model tidak berjalan mulus. Meskipun awalnya dirancang untuk melatih model selama 9 *epoch*, tantangan muncul ketika sesi pertama terhenti pada *epoch* pertama karena masalah kelebihan penggunaan memori. Solusi yang diambil adalah

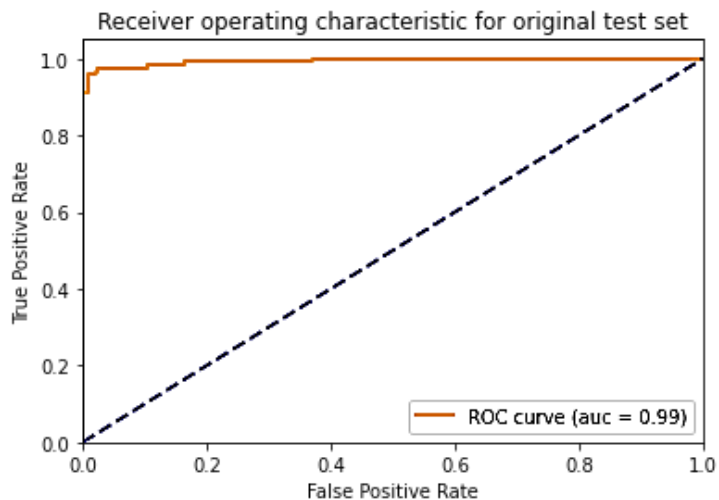
mengunduh titik pemeriksaan model (*checkpoint*) setelah *epoch* ke-4, yang kemudian digunakan sebagai dasar untuk melanjutkan pelatihan dengan 5 *epoch* tambahan. Meskipun terjadi gangguan, hasil akhir yang berhasil dicapai menunjukkan bahwa adaptasi dan ketekunan dalam menghadapi hambatan teknis dapat menghasilkan peningkatan kinerja yang signifikan.

Hasil evaluasi kinerja model (Gambar 5) menunjukkan prestasi yang mengesankan. Model mencapai nilai *loss* terendah sebesar 0.30%, menggambarkan sejauh mana model dapat menyesuaikan diri dengan data latihan. Di samping itu, akurasi pada data pelatihan dan evaluasi mencapai 99.14%, yang menegaskan kemampuan model dalam mengklasifikasikan data dengan tingkat keakuratan yang sangat tinggi. Kombinasi antara representasi kata yang kaya dari embeddings GloVe dan kemampuan arsitektur rekuren Bi-LSTM dalam menangkap konteks jarak jauh dan dekat sepertinya memberikan hasil yang sangat positif.

Hasil yang diperoleh menunjukkan bahwa penggabungan metode GloVe dan BiLSTM mampu menghasilkan model yang tangguh dalam tugas klasifikasi. Meskipun terdapat kendala teknis dalam perjalanan pelatihan, akhirnya model berhasil mencapai tingkat akurasi yang sangat baik dan tingkat *loss* yang rendah. Meskipun prestasi ini patut diapresiasi, disarankan untuk melanjutkan penelitian dengan lebih banyak variasi dataset dan analisis lebih dalam guna memastikan generalisasi yang lebih luas dan konsisten dari model yang dihasilkan.

Hasil evaluasi model pada data pengujian yang berisi berita asli menunjukkan performa yang sangat mengesankan. Dalam hal akurasi, model mencapai skor sebesar 97.5%. Ini menunjukkan bahwa model memiliki kemampuan yang sangat baik dalam mengklasifikasikan berita asli dengan benar. Kinerja ini secara signifikan melampaui tingkat kebetulan dan mengindikasikan bahwa model mampu mengenali dan memprediksi berita asli dengan presisi yang tinggi.

Selain itu, kami menggunakan metrik evaluasi AUC (*Area Under the Curve*) sebagai metrik evaluasi juga menghasilkan skor yang mengesankan, yaitu 99.4%. Skor AUC mencerminkan kemampuan model dalam membedakan antara kelas positif dan negatif. Dengan skor AUC yang mendekati 1, model berhasil mencapai tingkat diskriminasi yang sangat tinggi, memperkuat bukti bahwa model memiliki kemampuan untuk mengklasifikasikan dengan tepat (Gambar 6).



Gambar 6. Metrik evaluasi dengan kurva AUC

Selanjutnya, hasil evaluasi juga melibatkan presisi, recall, dan skor F1. Presisi, yang mengukur jumlah positif yang benar dibagi dengan total positif yang diprediksi, mencapai skor 97.5%. Ini menunjukkan bahwa dari semua berita yang diprediksi sebagai asli oleh model, 97.5% di antaranya benar-benar merupakan berita asli. Recall, yang mengukur jumlah positif yang benar dibagi dengan total positif yang sebenarnya, juga mencapai 97.5%. Ini menggambarkan bahwa model berhasil mendeteksi 97,5% dari seluruh berita asli yang ada dalam dataset. Selanjutnya, skor F1, yang menggabungkan presisi dan recall, juga mencapai 97.5%. Skor F1 ini mencerminkan keseimbangan antara presisi dan recall, dan hasilnya yang tinggi mengindikasikan bahwa model memiliki performa yang baik dalam klasifikasi berita asli.

Secara keseluruhan, evaluasi hasil ini menggambarkan performa yang sangat baik dari model dalam mengenali dan mengklasifikasikan berita asli. Skor akurasi, AUC, presisi, recall, dan skor F1 yang tinggi semakin memvalidasi kemampuan model yang dikembangkan dengan menggunakan embeddings GloVe dan arsitektur Bi-LSTM. Namun, penting untuk mengevaluasi model pada berbagai dataset yang lebih luas dan beragam untuk memahami sejauh mana generalisasi yang mungkin dicapai oleh model ini dalam berbagai konteks.

4. KESIMPULAN

Dalam penelitian ini, model yang menggabungkan embeddings GloVe dengan arsitektur BiLSTM telah dievaluasi dengan hasil yang mengesankan. Terjadi peningkatan signifikan dalam validasi kinerja dari 98% menjadi 99%, menunjukkan akurasi prediksi yang lebih baik. Meskipun menghadapi kendala teknis, model berhasil mencapai akurasi tes berita asli sebesar 97.5%, dan nilai AUC, presisi, recall, dan skor F1 semuanya mencapai 97.5%, menegaskan kinerja yang konsisten. Penggabungan GloVe dan BiLSTM terbukti efektif dalam meningkatkan kemampuan model dalam mengklasifikasikan data dan memahami konteks. Meskipun hasil ini mengesankan, pengujian lebih lanjut pada dataset yang lebih luas dan variasi masih diperlukan untuk memastikan generalisasi yang lebih baik.

REFERENCES

- [1] R. Oshikawa, J. Qian, and W. Y. Wang, "A survey on natural language processing for fake news detection," *arXiv Prepr. arXiv1811.00770*, 2018.
- [2] A. Roy, K. Basak, A. Ekbal, and P. Bhattacharyya, "A deep ensemble framework for fake news detection and classification," *arXiv Prepr. arXiv1811.04670*, 2018.
- [3] H. Jwa, D. Oh, K. Park, J. M. Kang, and H. Lim, "exbake: Automatic fake news detection model based on bidirectional encoder representations from transformers (bert)," *Appl. Sci.*, vol. 9, no. 19, p. 4062, 2019.
- [4] R. Yunanto, A. P. Purfini, and A. Prabuwisesa, "Survei Literatur: Deteksi Berita Palsu Menggunakan Pendekatan Deep Learning," *J. Manaj. Inform.*, vol. 11, no. 2, pp. 118–130, 2021.
- [5] A. R. Ahmad and H. R. Murad, "The impact of social media on panic during the COVID-19 pandemic in Iraqi Kurdistan: online questionnaire study," *J. Med. Internet Res.*, vol. 22, no. 5, p. e19556, 2020.
- [6] F. Jin, E. Dougherty, P. Saraf, Y. Cao, and N. Ramakrishnan, "Epidemiological modeling of news and rumors on twitter," in *Proceedings of the 7th workshop on social network mining and analysis*, 2013, pp. 1–9.
- [7] S. Kumar, R. West, and J. Leskovec, "Disinformation on the web: Impact, characteristics, and detection of wikipedia hoaxes," in

- Proceedings of the 25th international conference on World Wide Web*, 2016, pp. 591–602.
- [8] S. Kwon, M. Cha, and K. Jung, "Rumor detection over varying time windows," *PLoS One*, vol. 12, no. 1, p. e0168344, 2017.
 - [9] N. Ruchansky, S. Seo, and Y. Liu, "Csi: A hybrid deep model for fake news detection," in *Proceedings of the 2017 ACM on Conference on Information and Knowledge Management*, 2017, pp. 797–806.
 - [10] Q. Le and T. Mikolov, "Distributed representations of sentences and documents," in *International conference on machine learning*, 2014, pp. 1188–1196. doi: 10.48550/arXiv.1405.4053.
 - [11] A. Muzakir, H. Syaputra, and F. Panjaitan, "A Comparative Analysis of Classification Algorithms for Cyberbullying Crime Detection: An Experimental Study of Twitter Social Media in Indonesia," *Sci. J. Informatics; Vol 9, No 2 Novemb. 2022, Doi: 10.15294/sji.v9i2.35149*.
 - [12] W. Lopez, J. Merlino, and P. Rodriguez-Bocca, "Learning semantic information from Internet Domain Names using word embeddings," *Eng. Appl. Artif. Intell.*, vol. 94, p. 103823, 2020.
 - [13] S. Lai, K. Liu, S. He, and J. Zhao, "How to generate a good word embedding," *IEEE Intell. Syst.*, vol. 31, no. 6, pp. 5–14, 2016.
 - [14] A. B. Soliman, K. Eissa, and S. R. El-Beltagy, "Aravec: A set of arabic word embedding models for use in arabic nlp," *Procedia Comput. Sci.*, vol. 117, pp. 256–265, 2017.
 - [15] S. M. Rezaeinia, A. Ghodsi, and R. Rahmani, "Improving the accuracy of pre-trained word embeddings for sentiment analysis," *arXiv Prepr. arXiv1711.08609*, 2017.
 - [16] M. Schuster and K. K. Paliwal, "Bidirectional recurrent neural networks," *IEEE Trans. Signal Process.*, vol. 45, no. 11, pp. 2673–2681, 1997.
 - [17] G. Rao, W. Huang, Z. Feng, and Q. Cong, "LSTM with sentence representations for document-level sentiment classification," *Neurocomputing*, vol. 308, pp. 49–57, 2018.