

BAB I

PENDAHULUAN

1.1 Latar Belakang Masalah

Perkembangan teknologi *Natural Language Processing* (NLP) dalam beberapa tahun terakhir mengalami kemajuan yang sangat pesat, terutama dengan hadirnya pendekatan berbasis *deep learning*. NLP memungkinkan komputer untuk memahami, menganalisis, serta memproses bahasa manusia secara otomatis dalam berbagai bentuk, khususnya data teks yang dihasilkan dari aktivitas digital seperti media sosial, forum diskusi, dan aplikasi pesan singkat. Seiring meningkatnya penggunaan internet dan media sosial, jumlah data teks yang dihasilkan juga semakin besar, sehingga membutuhkan metode pengolahan yang efektif dan mampu menangkap makna secara akurat (Islam dkk., 2024).

Salah satu jenis data yang paling banyak dihasilkan dalam lingkungan digital saat ini adalah teks pendek, seperti komentar, tweet, dan ulasan singkat. Teks pendek memiliki karakteristik yang berbeda dibandingkan teks panjang karena keterbatasan konteks dan informasi yang terkandung di dalamnya. Hal ini menyebabkan proses klasifikasi menjadi lebih kompleks karena model harus mampu memahami makna dari informasi yang terbatas dan sering kali ambigu. Selain itu, teks pendek juga cenderung tidak mengikuti struktur bahasa yang baku, sehingga semakin meningkatkan tingkat kesulitan dalam proses analisis (Lin & Nuha, 2023).

Keterbatasan konteks pada teks pendek tidak hanya menyulitkan proses analisis oleh model, tetapi juga berdampak pada pemahaman pengguna terhadap isi teks tersebut. Banyak teks pendek seperti *tweet* tidak secara eksplisit menunjukkan topik utama, sehingga pembaca sering kali kesulitan dalam menangkap maksud yang ingin disampaikan. Padahal, kemampuan untuk mengetahui topik dari suatu teks menjadi penting dalam membantu proses penyaringan informasi serta memahami opini publik yang berkembang di media sosial (Islam dkk., 2024). Oleh karena itu, diperlukan suatu pendekatan yang mampu mengidentifikasi topik secara otomatis melalui proses klasifikasi teks. Klasifikasi sentimen memungkinkan teks

pendek dikelompokkan berdasarkan kesamaan makna sehingga informasi menjadi lebih terstruktur dan mudah dipahami. Dalam hal ini, pendekatan berbasis *deep learning* banyak digunakan karena mampu mempelajari pola kompleks dalam data teks serta meningkatkan performa klasifikasi dibandingkan metode konvensional (Widiantoro dkk., 2024).

Salah satu topik yang banyak dibahas masyarakat Indonesia di media sosial, khususnya *platform X/Twitter*, adalah judi *online*. Topik judi *online* sering menjadi perhatian publik karena berkaitan dengan dampak sosial dan meningkatnya aktivitas perjudian digital di masyarakat. Julianti dkk. (2024) menjelaskan bahwa media sosial menjadi sarana penting bagi masyarakat untuk menyampaikan komentar dan pendapat terkait pengalaman mengenai judi *online*. Tingginya aktivitas pengguna dalam membahas topik tersebut menghasilkan data teks pendek dalam jumlah besar berupa opini, komentar, dan tanggapan masyarakat di media sosial.

Pemilihan topik judi online dalam penelitian ini didasarkan pada tingginya intensitas diskusi masyarakat mengenai fenomena tersebut di *platform X/Twitter* serta karakteristik teks yang umumnya singkat, informal, dan kontekstual sehingga sesuai dengan karakteristik teks pendek (*short text*). Julianti dkk. (2024) juga menyatakan bahwa metode *Natural Language Processing* (NLP) memungkinkan proses analisis teks secara otomatis untuk memperoleh pemahaman terhadap opini masyarakat mengenai judi *online* di media sosial. Oleh karena itu, topik judi *online* dinilai relevan untuk digunakan dalam penelitian klasifikasi teks pendek berbasis NLP, khususnya dalam menguji kemampuan model Recurrent Neural Network (RNN), Long Short-Term Memory (LSTM), dan BiDirectional Long Short-Term Memory (BiLSTM) dalam memahami konteks teks berbahasa Indonesia pada media sosial.

Dalam proses pengolahan teks, data perlu direpresentasikan dalam bentuk numerik agar dapat diproses oleh model *deep learning*. Representasi teks tradisional seperti *bag-of-words* dan TF-IDF memiliki keterbatasan karena tidak mampu menangkap hubungan semantik antar kata maupun konteks penggunaannya dalam kalimat, sehingga kurang efektif terutama pada teks pendek yang memiliki

informasi terbatas. Untuk mengatasi hal tersebut, digunakan teknik *word embedding* seperti *Word2Vec* yang mampu merepresentasikan kata dalam bentuk vektor berdasarkan konteks kemunculannya, sehingga kata-kata yang memiliki makna serupa akan berada pada posisi yang berdekatan dalam ruang vektor. Pendekatan ini memungkinkan model untuk memahami makna kata secara lebih kontekstual dibandingkan metode konvensional, sehingga dapat meningkatkan kualitas representasi data sebelum dilakukan proses klasifikasi. Penelitian yang dilakukan oleh Widhi dkk. (2023) menunjukkan bahwa penggunaan *Word2Vec* mampu meningkatkan performa model dalam tugas klasifikasi teks karena kemampuannya dalam menangkap hubungan semantik antar kata secara lebih baik.

Dalam konteks bahasa Indonesia, tantangan tersebut menjadi lebih kompleks karena adanya variasi morfologi yang cukup tinggi, seperti penggunaan imbuhan, reduplikasi, serta variasi bentuk kata tidak baku yang sering digunakan dalam komunikasi digital. Selain itu, ketersediaan sumber daya NLP untuk bahasa Indonesia masih relatif terbatas dibandingkan bahasa Inggris, sehingga penelitian di bidang ini masih memiliki banyak ruang untuk dikembangkan (Cahyawijaya dkk., 2022).

Untuk mengatasi berbagai tantangan tersebut, pendekatan berbasis *deep learning* telah banyak digunakan dalam berbagai tugas NLP, termasuk klasifikasi teks. Salah satu arsitektur yang banyak digunakan adalah *Recurrent Neural Network* (RNN), yang dirancang khusus untuk memproses data sekuensial dan mampu mempertahankan informasi dari langkah sebelumnya. Namun, RNN memiliki keterbatasan dalam menangkap ketergantungan jangka panjang akibat masalah *vanishing gradient*, sehingga performanya cenderung menurun pada data yang memiliki hubungan konteks yang kompleks (Alizadegan dkk., 2024).

Sebagai pengembangan dari RNN, *Long Short-Term Memory* (LSTM) diperkenalkan untuk mengatasi permasalahan tersebut dengan menambahkan mekanisme memori yang mampu mempertahankan informasi dalam jangka waktu yang lebih lama. LSTM terbukti mampu meningkatkan performa dalam berbagai tugas NLP karena kemampuannya dalam menangkap dependensi jangka panjang dalam teks (Chauhan dkk., 2024). Selanjutnya, *Bidirectional Long Short-Term*

Memory (BiLSTM) dikembangkan untuk meningkatkan kemampuan model dengan memproses data dari dua arah, yaitu dari masa lalu ke masa depan dan sebaliknya, sehingga mampu memahami konteks secara lebih menyeluruh.

Meskipun berbagai arsitektur tersebut telah menunjukkan performa yang baik, tidak terdapat satu model yang secara konsisten unggul dalam semua kondisi. Performa suatu model sangat bergantung pada karakteristik data yang digunakan, termasuk panjang teks, kompleksitas bahasa, serta jenis tugas yang dilakukan (Islam dkk., 2024). Hal ini menunjukkan bahwa pemilihan algoritma yang tepat menjadi faktor penting dalam keberhasilan penerapan NLP.

Sejumlah penelitian terbaru telah mencoba membandingkan performa model-model berbasis RNN dan variannya. Misalnya, penelitian oleh Airlangga (2024) menunjukkan bahwa terdapat perbedaan performa yang signifikan antara RNN, LSTM, dan BiLSTM dalam tugas klasifikasi berita palsu, di mana masing-masing model memiliki keunggulan pada kondisi tertentu. Penelitian lain oleh Suhaeni dkk. (2024) juga menunjukkan bahwa BiLSTM cenderung menghasilkan performa yang lebih baik dibandingkan LSTM dalam klasifikasi sentimen pada data media sosial berbahasa Indonesia. Selain itu, penelitian oleh Aditya dkk. (2024) menunjukkan bahwa LSTM mampu meningkatkan akurasi dibandingkan RNN dalam klasifikasi teks bahasa Indonesia.

Namun demikian, sebagian besar penelitian tersebut memiliki keterbatasan, baik dari segi jenis data, jumlah algoritma yang dibandingkan, maupun fokus penelitian yang belum secara spesifik menargetkan teks pendek berbahasa Indonesia. Banyak penelitian hanya membandingkan dua algoritma atau menggunakan dataset yang berbeda-beda, sehingga sulit untuk menarik kesimpulan yang komprehensif mengenai model mana yang paling optimal dalam konteks tertentu (Nugroho & Nugroho, 2024).

Selain itu, belum adanya standar atau acuan yang jelas dalam pemilihan model menyebabkan proses pengembangan sistem NLP sering kali dilakukan secara *trial and error*. Kondisi ini dapat mengakibatkan pemilihan model yang kurang optimal, sehingga berdampak pada rendahnya performa sistem yang dihasilkan. Padahal, perbedaan performa antar model dapat memberikan dampak

yang signifikan, terutama dalam aplikasi nyata seperti analisis sentimen, klasifikasi opini publik, dan sistem rekomendasi berbasis teks.

Berdasarkan permasalahan tersebut, diperlukan suatu penelitian yang secara khusus melakukan analisis komparatif terhadap algoritma RNN, LSTM, dan BiLSTM pada teks pendek berbahasa Indonesia. Melalui pendekatan ini, diharapkan dapat diperoleh gambaran yang lebih jelas mengenai keunggulan dan keterbatasan masing-masing algoritma dalam menangani karakteristik teks pendek. Selain itu, hasil penelitian ini diharapkan dapat memberikan kontribusi berupa bukti empiris yang dapat dijadikan sebagai acuan dalam pemilihan model yang lebih tepat dan optimal dalam pengembangan aplikasi NLP berbahasa Indonesia

1.2 Identifikasi Masalah

Berdasarkan latar belakang yang telah diuraikan, dapat diidentifikasi beberapa masalah yang menjadi fokus dalam penelitian ini:

1. Belum adanya penelitian yang secara spesifik membandingkan performa algoritma RNN, LSTM, dan BiLSTM pada teks pendek berbahasa Indonesia secara komprehensif.
2. Keterbatasan pemahaman mengenai karakteristik dan keunggulan masing-masing algoritma dalam menangani struktur linguistik dan morfologi bahasa Indonesia yang kompleks pada teks pendek.
3. Tidak tersedianya panduan empiris yang dapat dijadikan acuan dalam pemilihan algoritma yang optimal untuk aplikasi NLP teks pendek berbahasa Indonesia.

1.3 Pembatasan Masalah

Untuk memfokuskan penelitian dan memastikan hasil yang dapat diinterpretasi dengan baik, penelitian ini dibatasi pada aspek-aspek berikut:

1. Algoritma yang dibandingkan terbatas pada RNN, LSTM, dan BiLSTM tanpa melibatkan algoritma *deep learning* lainnya seperti *Transformer* atau *Convolutional Neural Network* (CNN).

2. Fokus tugas penelitian adalah klasifikasi teks pendek untuk mengklasifikasikan sentimen pada setiap teks dengan melihat kemampuan algoritma dalam mempelajari pola data sekuensial.
3. Data yang digunakan diperoleh melalui teknik *scraping* pada *platform* media sosial *X* (*Twitter*) dengan memanfaatkan *filter* bahasa (*lang=in*) dan menggunakan kata kunci yang berkaitan dengan judi online, seperti “judi online”, dan “judol” untuk memperoleh data teks pendek berbahasa Indonesia yang relevan dengan topik penelitian.
4. Panjang teks dibatasi maksimal 280 karakter sesuai dengan karakteristik standar *platform X*.
5. Evaluasi performa model dilakukan menggunakan metrik *accuracy*, *precision*, *recall*, dan *F1-score*, serta analisis efisiensi komputasi berdasarkan waktu pelatihan dan penggunaan sumber daya sistem.
6. Proses *word embedding*, yaitu transformasi teks menjadi vektor numerik untuk diolah oleh model *deep learning*, akan menggunakan metode *Word2Vec*.
7. Penelitian membahas aspek komputasi dan efisiensi waktu training, dan juga pada performa akurasi model.
8. Penelitian ini fokus pada klasifikasi biner karena sentimen Netral pada data hanya dibawah 10%.

1.4 Perumusan Masalah

Berdasarkan latar belakang, identifikasi masalah, dan pembatasan masalah tersebut, maka didapatkan sebuah rumusan masalah yaitu "Bagaimana hasil dan perbandingan performa model RNN, LSTM, dan BiLSTM mengklasifikasikan sentimen judi *online* di tahun 2025 pada teks pendek berbahasa Indonesia"?

1.5 Tujuan Penelitian

Berdasarkan uraian di atas, maka tujuan dari penelitian ini adalah untuk menganalisis hasil dan perbandingan performa dari model RNN, LSTM, dan

BILSTM saat diterapkan untuk klasifikasi sentimen judi *online* di 2025 pada teks pendek berbahasa Indonesia.

1.6 Manfaat Penelitian

Penelitian ini diharapkan dapat memberikan manfaat sebagai berikut:

1.6.1 Manfaat Teoritis

1. Menambah khazanah ilmu pengetahuan di bidang *Natural Language Processing*, khususnya dalam pemrosesan teks berbahasa Indonesia yang memiliki kompleksitas morfologi tinggi.
2. Memberikan kontribusi empiris terhadap pemahaman karakteristik dan performa algoritma RNN, LSTM, dan BiLSTM pada teks pendek, terutama dalam konteks bahasa Indonesia.

1.6.2 Manfaat Praktis

1. Memberikan panduan bagi pengembang aplikasi dalam memilih algoritma *deep learning* yang tepat untuk mengolah data teks media sosial yang berkarakter singkat dan dinamis.
2. Membantu perusahaan teknologi dan *startup* dalam mengoptimalkan sistem NLP mereka untuk pasar Indonesia yang memiliki jumlah pengguna internet sebesar 229,43 juta jiwa.
3. Memberikan dasar empiris untuk pengambilan keputusan teknis dalam proyek-proyek yang melibatkan pemrosesan bahasa alami, khususnya untuk teks pendek.
4. Berkontribusi pada peningkatan kualitas layanan digital berbahasa Indonesia melalui pemilihan algoritma yang optimal, mengingat tingginya penggunaan internet di Indonesia yang terus meningkat konsisten.
5. Penelitian ini diharapkan dapat memberikan informasi mengenai *trade-off* antara performa klasifikasi dan efisiensi

komputasi pada model RNN, LSTM, dan BiLSTM sehingga dapat menjadi pertimbangan dalam pemilihan model NLP untuk teks pendek.



Intelligentia - Dignitas