

BAB I

PENDAHULUAN

Bab ini membahas mengenai latar belakang, rumusan masalah, manfaat dan tujuan, ruang lingkup, dan sistematika penulisan pada skripsi yang berjudul Klasifikasi Berita Palsu Berbahasa Indonesia Menggunakan Metode BERT-CNN.

1.1 Latar Belakang

Dalam era informasi digital yang semakin berkembang pesat, kemudahan akses terhadap berita dan informasi menjadi salah aspek penting dalam kehidupan sehari-hari. Masyarakat saat ini memiliki opsi yang praktis untuk memperoleh dan bertukar informasi melalui sosial media, yang sebelumnya hanya terbatas melalui email, blog, dan web (Istriyani & Widian, 2017). Namun, kemudahan tersebut juga menimbulkan berbagai masalah, seperti *cyber bullying*, SARA, ujaran kebencian, dan penyebaran berita palsu (*hoax*) (Parhan dkk., 2021).

Berita palsu atau *hoax* merupakan informasi atau berita yang disebar untuk menyesatkan atau memanipulasi orang lain (Pennycook & Rand, 2019). *Hoax* bertujuan untuk menggiring opini, membuat opini publik, dan membentuk persepsi (Rahadi, 2017). Penyebaran berita *hoax* memiliki tujuan sebagai bahan lelucon, menjatuhkan pesaing, promosi penipuan, dan menyebarkan sesuatu hal yang belum diverifikasi kebenarannya (Rahadi, 2017).

Kominfo (2023) mengungkapkan terdapat 11.642 temuan konten *hoax* sejak periode Agustus 2018 sampai dengan Mei 2023. Dari 11.642 temuan tersebut mayoritas kategori berita palsu adalah kesehatan, pemerintahan, dan penipuan. Hal ini menunjukkan bahwa sulitnya menyaring informasi yang ada di internet. Hal ini didukung dengan survei yang dilakukan Katadata Insight Center (KIC) yang mengungkapkan 30% sampai 60% masyarakat Indonesia terpapar *hoax* saat mengakses dan berkomunikasi melalui sosial media (KIC, 2020). Kondisi ini diperburuk di mana hanya 21% sampai 36% dari data tersebut yang mampu mengenali berita *hoax* (KIC, 2020). Hal ini disebabkan karena lemahnya literasi informasi digital masyarakat (Solihah & Yasir, 2022).

Pemerintah telah melakukan beberapa cara untuk penanganan penyebaran berita palsu, dimulai dengan literasi media dan kerjasama dengan pers, edukasi masyarakat,

dan membentuk Gerakan Literasi Nasional (Arwendria & Oktavia, 2019). Namun penerapan hal ini akan memerlukan banyak waktu dan sumber daya untuk pelaksanaannya. Permasalahan lain yang muncul adalah deteksi berita palsu secara manual membutuhkan keterlibatan ahli/pakar untuk dapat menghasilkan data yang kredibel dan hal ini juga membutuhkan waktu yang lama dan berbiaya mahal (Zhou & Zafarani, 2018). Dalam hal ini maka dibutuhkan sebuah sistem yang dapat menangani masalah ini dengan cepat dan dengan sumber daya seminimal mungkin. Oleh karena itu, penerapan *machine learning* pada masalah ini dipilih untuk memenuhi kebutuhan tersebut.

Berbagai penelitian mengenai deteksi berita palsu berbahasa Indonesia telah banyak diterapkan menggunakan metode *classical machine learning*, seperti *Naïve Bayes* (Rahutomo dkk., 2019), *Support Vector Machine* (Ropikoh dkk., 2021), *K-Nearest Neighbor* (Azhar, 2021), dan lainnya. Namun, terdapat beberapa permasalahan pada penerapan metode *classical machine learning*, yaitu keterbatasan dalam representasi kata, keterbatasan memahami konteks, keterbatasan menangani data teks yang panjang, dan membutuhkan ekstraksi fitur manual (Kowsari dkk., 2019).

Permasalahan tersebut dapat diatasi dengan menggunakan model *deep learning*, seperti *Convolutional Neural Network* (CNN) dan *Long Short Term Network* (LSTM) (Kowsari dkk., 2019). Kurniawan (2020) meneliti deteksi berita palsu berbahasa Indonesia menggunakan CNN dan LSTM dengan *word2vec* sebagai representasi teks, di mana CNN menghasilkan akurasi 88% dan LSTM menghasilkan akurasi 84%. Berdasarkan hasil penelitian tersebut, dapat disimpulkan bahwa algoritma CNN lebih baik dalam melakukan tugas klasifikasi berita palsu berbahasa Indonesia daripada algoritma LSTM. Hal ini dikarenakan CNN memiliki kelebihan untuk dapat melakukan pengenalan pola lokal dalam data dengan mendeteksi fitur dalam teks (Widhiyasana dkk., 2021). Pada penelitian yang dilakukan Kurniawan (2020) sebelumnya masih menggunakan *word embedding* sebagai teks representasinya, sedangkan Bao dkk (2021) meneliti *Bidirectional Encoder Representations from Transformers* (BERT) sebagai teks representasi pada klasifikasi teks menggunakan metode *deep learning* dan menghasilkan peningkatan akurasi, presisi, dan *recall* jika dibandingkan dengan metode *deep learning* menggunakan *word embedding* sebagai teks representasi. Hasil penelitian yang dilakukan oleh Bao dkk (2021) menjelaskan

bahwa akurasi model CNN dengan menggunakan BERT sebagai teks representasinya memiliki performa yang lebih baik, di mana model BERT-CNN memiliki akurasi 0,9402. Hasil ini jauh meningkat jika dibandingkan dengan akurasi model CNN tanpa menggunakan BERT yaitu 0,9104.

Berdasarkan data di atas, penelitian ini bertujuan menyelesaikan masalah klasifikasi berita palsu berbahasa Indonesia dengan menggunakan metode BERT-CNN. Penelitian ini diharapkan membantu perkembangan deteksi berita palsu berbahasa Indonesia serta dapat digunakan untuk pencegahan penyebaran berita palsu.

1.2 Rumusan Masalah

Berdasarkan latar belakang di atas, rumusan masalah pada penelitian ini adalah bagaimana mengimplementasi metode BERT-CNN pada klasifikasi berita palsu berbahasa Indonesia.

1.3 Tujuan dan Manfaat

Tujuan dari penelitian skripsi ini adalah dihasilkan sebuah model BERT-CNN yang dapat mengklasifikasikan berita palsu berbahasa Indonesia.

Manfaat dari penelitian skripsi ini adalah diperolehnya model klasifikasi berita palsu berbahasa Indonesia menggunakan metode BERT-CNN yang dapat digunakan untuk mengklasifikasikan berita asli dan palsu yang nantinya dapat diterapkan sebagai *filter* dalam tindakan pencegahan penyebaran berita palsu.

1.4 Ruang Lingkup

Ruang lingkup penelitian skripsi ini dibatasi oleh aspek-aspek sebagai berikut:

1. Pengumpulan *dataset* menggunakan metode *scraping* pada website detik.com dan turnbackhoax.id.
2. *Dataset* yang digunakan berupa data teks berukuran 12.624 data dengan rincian 10.769 data berita palsu dan 1.855 data berita benar.
3. Implementasi dilakukan menggunakan bahasa pemrograman *python* di *Jupyter Notebook* pada *platform Google Colaboratory*.

1.5 Sistematika Penulisan

Sistem penulisan dibuat untuk memberikan gambaran umum dan penjelasan mengenai penulisan serta pembahasan mengenai Klasifikasi Berita Palsu Berbahasa

Indonesia menggunakan Metode BERT-CNN. Berikut adalah sistematika penulisan yang digunakan:

BAB I PENDAHULUAN

Bab ini menyajikan latar belakang masalah, rumusan masalah, tujuan dan manfaat penelitian, ruang lingkup, serta sistematika penulisan pada skripsi yang berjudul Klasifikasi Berita Palsu Berbahasa Indonesia Menggunakan Metode BERT-CNN.

BAB II LANDASAN TEORI

Bab ini menyajikan hasil studi pustaka mengenai dasar teori yang berhubungan dengan pelaksanaan dan penyusunan penulisan penelitian Klasifikasi Berita Palsu Berbahasa Indonesia menggunakan Metode BERT-CNN seperti *state of the art*, berita palsu, pra-pemrosesan data, *imbalance dataset*, pembagian data, *transfer learning*, BERT, CNN, fungsi aktivasi, *adam optimizer*, *loss function*, dan *confusion matrix*.

BAB III METODOLOGI PENELITIAN

Bab ini menyajikan tentang metodologi penelitian yang digunakan dalam penulisan penelitian Klasifikasi Berita Palsu Berbahasa Indonesia menggunakan Metode BERT-CNN.

BAB IV HASIL DAN PEMBAHASAN

Bab ini berisi pembahasan hasil dari penulisan penelitian Klasifikasi Berita Palsu Berbahasa Indonesia menggunakan Metode BERT-CNN.

BAB V PENUTUP

Bab ini berisi tentang kesimpulan dari bab-bab yang telah diuraikan pada bab-bab sebelumnya dan berisi saran sebagai bahan masukan untuk pengembangan lebih lanjut.