

BAB I

PENDAHULUAN

1.1 Latar Belakang

Perkembangan teknologi informasi dan komunikasi yang sangat pesat telah mengubah pola masyarakat dalam mengakses dan menyebarkan informasi, khususnya melalui media digital dan media sosial [1]. Kemudahan distribusi informasi tersebut meningkatkan penyebaran berita palsu atau hoaks yang sulit diverifikasi kebenarannya [2]. Hoaks tidak hanya menimbulkan kesalahpahaman informasi, tetapi juga berpotensi memicu keresahan sosial, konflik, serta menurunkan kepercayaan publik terhadap institusi resmi [3]. Dalam konteks Indonesia, tingginya konsumsi informasi digital berbahasa Indonesia menyebabkan penyebaran hoaks menjadi salah satu permasalahan yang perlu mendapatkan perhatian lebih lanjut [4].

Seiring meningkatnya penyebaran hoaks digital, pendekatan berbasis *Machine Learning* dan *Natural Language Processing* (NLP) mulai banyak digunakan untuk membantu proses identifikasi teks hoaks secara otomatis [4]. Berbagai penelitian sebelumnya menunjukkan bahwa klasifikasi teks dapat digunakan untuk membedakan teks hoaks dan non-hoaks berdasarkan karakteristik linguistik tertentu [1], [2], [3], [5]. Algoritma *Support Vector Machine* (SVM) banyak dipilih karena kemampuannya menangani data teks berdimensi tinggi dan bersifat *sparse* secara efisien, sehingga sesuai apabila dikombinasikan dengan *Term Frequency-Inverse Document Frequency* (TF-IDF) yang juga menghasilkan representasi vektor berdimensi tinggi dan *sparse* berdasarkan frekuensi kemunculan kata. Kombinasi keduanya terbukti menghasilkan performa klasifikasi yang baik pada teks hoaks berbahasa Indonesia [4], [9]. Di sisi lain, *Random Forest* dipilih untuk dipasangkan dengan *Word Embedding FastText* karena kemampuan *ensemble learning* pada *Random Forest* lebih sesuai untuk menangkap hubungan non-linear antar dimensi vektor semantik padat (*dense*) yang dihasilkan oleh *Word Embedding FastText*, berbeda dengan representasi *sparse* berbasis frekuensi pada *Term Frequency-Inverse Document Frequency* (TF-IDF) [8], [12].

Meskipun berbagai penelitian telah menunjukkan performa yang baik, hasil evaluasi antar penelitian masih menunjukkan variasi yang cukup signifikan [2], [3], [9]. Sebagai contoh, penelitian klasifikasi berita hoaks Covid-19 menggunakan *Support Vector Machine* (SVM) dengan kernel RBF hanya memperoleh *Accuracy* sebesar 90,46%, dengan

Precision dan *Recall* yang jauh lebih rendah yaitu masing-masing 66,86% dan 64,53% [19], sedangkan penelitian klasifikasi sentimen menggunakan *Random Forest* dengan *Word Embedding FastText* memperoleh *Accuracy* 70,27%, *Precision* 80%, *Recall* 54%, dan *F1-Score* 54% [20]. Perbedaan tersebut dipengaruhi oleh penggunaan dataset, metode representasi fitur, teknik *preprocessing*, serta skema pengujian yang berbeda pada setiap penelitian [5]. Kondisi ini menyebabkan hasil performa antar algoritma belum dapat dibandingkan secara langsung untuk menentukan pendekatan yang lebih optimal dalam klasifikasi teks hoaks berbahasa Indonesia. Oleh karena itu, diperlukan penelitian yang melakukan pengujian secara konsisten menggunakan dataset, skema validasi, dan metrik evaluasi yang sama agar perbedaan performa model dapat dianalisis secara lebih objektif.

Berdasarkan kondisi tersebut, penelitian ini melakukan evaluasi komparatif antara algoritma *Support Vector Machine* berbasis *Term Frequency-Inverse Document Frequency* (TF-IDF) dan *Random Forest* berbasis *Word Embedding FastText* dalam klasifikasi teks hoaks berbahasa Indonesia. Berbeda dengan penelitian-penelitian sebelumnya yang diuji pada kondisi dataset dan skema pengujian yang berbeda-beda, kedua algoritma pada penelitian ini dievaluasi menggunakan dataset yang sama, tahapan *preprocessing* yang sama, serta skema validasi *Stratified 5-Fold Cross Validation* yang identik, sehingga perbedaan performa yang diperoleh dapat dikaitkan secara langsung dengan perbedaan metode representasi fitur dan algoritma klasifikasi yang digunakan, bukan disebabkan oleh perbedaan kondisi pengujian. Perbandingan kinerja kedua model dilakukan berdasarkan metrik evaluasi *Accuracy*, *Precision*, *Recall*, dan *F1-Score*.

1.2 Rumusan Masalah

Berdasarkan latar belakang yang telah diuraikan, rumusan masalah dalam penelitian ini adalah bagaimana perbandingan kinerja algoritma *Support Vector Machine* (SVM) berbasis (TF-IDF) dan *Random Forest* berbasis *Word Embedding FastText* dalam klasifikasi teks hoaks berbahasa Indonesia menggunakan dataset dan skema pengujian yang sama berdasarkan metrik *accuracy*, *precision*, *recall*, dan *F1-score*.

1.3 Tujuan

Tujuan penelitian ini adalah melakukan evaluasi komparatif terhadap algoritma *Support Vector Machine* berbasis TF-IDF dan *Random Forest* berbasis *Word Embedding FastText* dalam klasifikasi teks hoaks berbahasa Indonesia menggunakan dataset dan skema

pengujian yang sama untuk menganalisis perbedaan performa model berdasarkan metrik *accuracy*, *precision*, *recall*, dan *F1-score*.

1.4 Manfaat

Manfaat dari penelitian ini adalah:

1. Memberikan kontribusi akademik dalam bentuk analisis komparatif terhadap metode klasifikasi teks hoaks berbasis *Machine Learning*.
2. Menjadi referensi bagi penelitian selanjutnya dalam pengembangan metode deteksi hoaks berbahasa Indonesia.
3. Menyediakan gambaran empiris mengenai pengaruh representasi fitur terhadap performa model klasifikasi teks.

1.5 Ruang Lingkup

1. Penelitian ini menggunakan data berupa teks berbahasa Indonesia yang bersifat publik dan berkaitan dengan konten hoaks serta non-hoaks.
2. Data hoaks diperoleh dari portal *TurnBackHoax* sebagai sumber verifikasi (*fact-checking*), sedangkan data non-hoaks diperoleh dari portal berita *Detik.com*, dengan pelabelan mengikuti hasil verifikasi resmi dari masing-masing sumber. Pengumpulan data dilakukan pada rentang periode publikasi tahun 2020 hingga 2024 menggunakan metode *web scraping* [20].

Unit analisis dan penetapan label:

- a. Label hoaks: konten yang klaimnya telah dinyatakan salah/menyesatkan oleh *TurnBackHoax* atau lembaga cek fakta resmi *Detik.com*.
 - b. Label non-hoaks: klarifikasi/fakta/berita terpercaya atau konten yang telah dibantah secara resmi.
3. Tahapan *pre-processing* yang diterapkan meliputi normalisasi huruf kecil (*case folding*), deteksi bahasa Indonesia, dan penghapusan kata henti (*stopword removal*), dengan *stemming* digunakan hanya jika terbukti meningkatkan kinerja model.
 4. Model yang dibandingkan dalam penelitian ini adalah *Support Vector Machine* (SVM) dengan representasi fitur *Term Frequency-Inverse Document Frequency* (TF-IDF) dan *Random Forest* dengan representasi fitur *Word Embedding FastText*.
 5. Pembagian data dilakukan menggunakan skema *Stratified 5-Fold Cross Validation* dengan proporsi data latih (*training*) sebesar 80% dan data uji (*testing*) sebesar 20% pada setiap *fold*, diterapkan secara konsisten untuk seluruh model yang dibandingkan.

6. Performa model dievaluasi menggunakan metrik *Accuracy*, *Precision*, *Recall*, dan *F1-Score* sebagai dasar perbandingan kinerja antar model.
7. Penelitian tidak membahas data multimodal seperti gambar, video, dan audio, tidak menganalisis jejaring atau faktor penyebaran, dan tidak membangun sistem waktu nyata. Melainkan berfokus pada analisis komparatif performa metode klasifikasi dalam konteks penelitian eksperimental.

