

BAB I

PENDAHULUAN

1.1 Latar Belakang

Perkembangan teknologi informasi di Indonesia telah mengubah portal berita online menjadi pilar utama pemenuhan kebutuhan informasi masyarakat. Namun, kemudahan ini diikuti oleh fenomena "infodemik", di mana informasi tidak akurat menyebar lebih cepat daripada verifikasi faktanya. Berdasarkan data laporan dari Kementerian Komunikasi dan Informatika (Kominfo), sepanjang tahun 2024 saja, teridentifikasi ribuan isu *Hoax* yang tersebar di berbagai platform digital, yang didominasi oleh isu politik, kesehatan, dan penipuan. Hal ini sejalan dengan temuan Muhammad Salim Albana *et al.* (2024) yang menyatakan bahwa interaksi komunikasi *Hoax* telah menjadi ancaman serius bagi kesejahteraan sosial di Indonesia.

Masalah utama yang dihadapi saat ini adalah rendahnya kemampuan filterisasi informasi pada level pengguna akhir. Jika terus dibiarkan tanpa adanya sistem deteksi otomatis yang andal, *Hoax* ini akan memicu polarisasi sosial, ketidakpercayaan pada institusi publik, hingga ancaman stabilitas nasional. *Hoax* bukan lagi sekadar kesalahan data, melainkan ancaman terhadap demokrasi dan keamanan digital.

Secara hukum, Pemerintah Indonesia telah memberikan perhatian serius melalui Undang-Undang No. 19 Tahun 2016 tentang Perubahan atas UU No. 11 Tahun 2008 tentang ITE, khususnya Pasal 28 ayat (1) yang melarang penyebaran berita bohong yang merugikan konsumen. Meskipun regulasi telah ada, implementasi penegakan hukum seringkali bersifat reaktif (setelah kejadian), sehingga diperlukan pendekatan teknologi preventif melalui otomatisasi deteksi.

Secara teoritis, disiplin ilmu *Natural Language Processing* (NLP) dan *Machine Learning* menawarkan solusi melalui klasifikasi teks biner. Implementasi NLP dalam mengelola dokumen teks dalam jumlah besar secara sistematis telah

terbukti efektif dalam mempermudah pengelompokan informasi yang kompleks (Khaidar *et al.*, 2025). Dua algoritma yang paling sering digunakan dalam deteksi *Hoax* adalah *Naive Bayes* dan *Support Vector Machine (SVM)*, yang masing-masing memiliki pendekatan komputasi yang berbeda.

Naive Bayes merupakan algoritma klasifikasi berbasis probabilitas yang bekerja dengan menghitung kemungkinan suatu artikel termasuk kategori *Hoax* atau valid berdasarkan distribusi kata-kata di dalamnya. Algoritma ini menggunakan Teorema Bayes dengan asumsi independensi fitur, di mana setiap kata dianggap tidak saling mempengaruhi dalam menentukan kelas akhir. Keunggulan utama *Naive Bayes* terletak pada efisiensi komputasinya yang tinggi dan kemampuannya menangani *dataset* berdimensi besar dengan waktu pelatihan yang cepat (Peretz *et al.*, 2024). Relevansi algoritma ini dalam menangani klasifikasi data berbasis teks di Indonesia masih sangat tinggi, sebagaimana dibuktikan oleh Nasyira *et al.* (2025) dan Khaidar *et al.* (2025) yang menunjukkan bahwa *Naive Bayes* tetap menjadi pilihan handal untuk mengklasifikasikan kualitas informasi dan sentimen opini pada media sosial dengan tingkat performa yang stabil.

Di sisi lain, *Support Vector Machine (SVM)* menggunakan pendekatan geometris dengan mencari *hyperplane* (bidang pemisah) optimal yang dapat memisahkan dua kelas data dengan margin maksimum. Dalam konteks klasifikasi teks, SVM mengubah artikel berita menjadi vektor numerik, kemudian mencari garis pemisah terbaik antara berita valid dan *Hoax* di ruang multidimensi. SVM dikenal unggul dalam menangani data dengan dimensi tinggi dan mampu menghasilkan model yang *robust* meskipun dengan *dataset* yang relatif kecil (Nurwanda & Rizkiani, 2023). Penelitian terbaru oleh Alaiya *et al.* (2025) memperkuat posisi SVM sebagai metode yang sangat kompetitif dalam analisis tekstual, di mana penggunaan kernel dan parameter yang tepat mampu memberikan hasil akurasi yang signifikan pada data bahasa Indonesia.

Meskipun kedua algoritma memiliki pendekatan yang berbeda, keduanya sama-sama memerlukan tahap ekstraksi fitur menggunakan metode TF-IDF (*Term Frequency-Inverse Document Frequency*) sebelum dapat memproses teks. TF-IDF

bekerja dengan memberikan bobot numerik pada setiap kata berdasarkan frekuensi kemunculannya dalam dokumen (TF) dan keunikannya di seluruh korpus (IDF). Metode ini efektif mengidentifikasi kata-kata kunci yang membedakan struktur bahasa berita valid dengan bahasa *Hoax* yang cenderung provokatif dan emosional (Rahmadani *et al.*, 2024).

Dalam tataran praktis, efektivitas sistem deteksi berita bohong masih menunjukkan hasil yang beragam. Pengujian yang dilakukan oleh Nurwanda & Rizkiani (2023) serta Ropikoh *et al.* (2021) mengonfirmasi bahwa algoritma *Support Vector Machine* (SVM) memiliki keunggulan dalam hal akurasi saat diuji pada *dataset* berskala global. Namun, temuan ini tidak bersifat mutlak. Penelitian terbaru dari Rahmadani *et al.* (2024) justru membuktikan bahwa algoritma *Naive Bayes*, apabila dikombinasikan dengan metode ekstraksi fitur TF-IDF, mampu mencapai tingkat akurasi yang sangat tinggi, yakni sebesar 94,12%. Kemampuan *Naive Bayes* dalam menangani karakteristik berita lokal di Indonesia juga telah divalidasi oleh Rahutomo *et al.* (2019) yang menunjukkan bahwa algoritma ini memiliki performa yang sangat andal untuk *dataset* berbahasa Indonesia.

Ketidakkonsistenan hasil antara beberapa penelitian tersebut mengindikasikan bahwa performa sebuah algoritma sangat bergantung pada karakteristik data dan periode waktu pengambilan informasi. Hal ini memicu perdebatan mengenai metode mana yang paling optimal untuk digunakan dalam konteks portal berita nasional saat ini. Berdasarkan urgensi tersebut, diperlukan pengujian ulang yang lebih komprehensif untuk membandingkan kembali performa *Naive Bayes* dan SVM menggunakan data yang lebih segar, guna menemukan solusi deteksi yang paling presisi di tengah evolusi pola *Hoax* tahun 2025.

Oleh karena itu, penelitian berjudul "**Perbandingan Akurasi Algoritma *Naive Bayes* dan *Support Vector Machine* dalam Mendeteksi *Hoax* pada Portal Berita Nasional di Indonesia**" ini mendesak untuk dilakukan guna memberikan tolok ukur (*benchmark*) baru yang lebih relevan bagi perkembangan teknologi deteksi *Hoax* di Indonesia.

1.2 Rumusan Masalah

Berdasarkan identifikasi masalah tersebut, rumusan masalah dalam penelitian ini adalah:

1. Bagaimana efektivitas tahap *preprocessing* data teks dan ekstraksi fitur TF-IDF pada *dataset Hoax* berita nasional periode 2025?
2. Berapa tingkat akurasi, presisi, dan *recall* yang dihasilkan oleh algoritma *Naive Bayes* dalam mengklasifikasikan *Hoax* pada portal berita nasional Indonesia?
3. Berapa tingkat akurasi, presisi, dan *recall* yang dihasilkan oleh algoritma *Support Vector Machine* (SVM) dalam mengklasifikasikan *Hoax* pada portal berita nasional Indonesia?
4. Apakah terdapat perbedaan performa yang signifikan antara algoritma *Naive Bayes* dan *Support Vector Machine* (SVM), dan manakah yang lebih unggul dalam mendeteksi *Hoax* pada portal berita nasional Indonesia?

1.3 Batasan Masalah

1. Data penelitian dikumpulkan secara manual. Data *Hoax* (label 0) diambil dari situs verifikasi fakta TurnbackHoax.id. Data valid (label 1) diambil dari artikel berita yang diterbitkan oleh portal berita nasional terverifikasi (seperti Detik.com dan Antara News).
2. Periode pengumpulan data dibatasi pada berita yang dipublikasikan antara 1 Januari 2025 hingga Desember 2025 untuk menjamin relevansi dan kemutakhiran data.
3. Penelitian hanya berfokus pada analisis teks (isi artikel), tidak termasuk analisis gambar, video, atau metadata lainnya.
4. Penelitian hanya membandingkan dua algoritma klasifikasi, yaitu *Naive Bayes Classifier* (spesifiknya *Multinomial Naive Bayes*) dan *Support Vector Machine* (SVM) (spesifiknya dengan *kernel linear*).
5. Metode ekstraksi fitur yang digunakan adalah *Term Frequency-Inverse Document Frequency* (TF-IDF).

6. Kinerja model diukur berdasarkan metrik Akurasi, Presisi, dan *Recall* yang dihitung dari *Confusion Matrix*.
7. Implementasi model menggunakan bahasa pemrograman Python dan *library* utama Scikit-learn, Pandas, serta NLTK/Sastrawi untuk pra-pemrosesan.

1.4 Tujuan Penelitian

Tujuan dari penelitian ini adalah sebagai berikut:

1. Mengimplementasikan teknik *preprocessing* teks (*Case Folding, Cleaning, Tokenizing, Stopword Removal, Stemming*) dan pembobotan TF-IDF pada *dataset* berita.
2. Membangun model klasifikasi menggunakan algoritma *Naive Bayes* untuk mendeteksi *Hoax* berdasarkan *dataset* 2025.
3. Membangun model klasifikasi menggunakan algoritma *Support Vector Machine* (SVM) untuk mendeteksi *Hoax* berdasarkan *dataset* 2025.
4. Menganalisis dan membandingkan kinerja kedua algoritma untuk menentukan metode yang paling akurat dalam mengidentifikasi berita *Hoax* di Indonesia.

1.5 Manfaat Penelitian

Hasil dari penelitian ini diharapkan dapat memberikan manfaat sebagai berikut:

1. Memberikan kontribusi empiris dalam memecahkan kontradiksi hasil penelitian sebelumnya terkait efektivitas NB vs SVM pada data bahasa Indonesia.
2. Sebagai referensi pengembangan sistem filter berita otomatis bagi penyedia portal media atau pengembang aplikasi deteksi *Hoax*.
3. Membantu dalam memitigasi dampak buruk penyebaran informasi palsu melalui pengembangan teknologi yang lebih presisi.
4. Memberikan kontribusi awal dalam pengembangan alat bantu (*tools*) otomatis yang lebih akurat untuk membantu masyarakat, jurnalis, dan platform digital dalam menyaring dan mengidentifikasi *Hoax*.