

BAB I

PENDAHULUAN

1.1 Latar Belakang

Perkembangan teknologi informasi dan komunikasi pada era digital mendorong percepatan akses dan penyebaran informasi melalui internet dan media sosial. Namun, kondisi ini juga memicu meningkatnya penyebaran berita palsu atau hoaks, khususnya pada isu politik. Berdasarkan data Kementerian Komunikasi dan Digital, sepanjang tahun 2024 teridentifikasi sebanyak 1.923 konten hoaks, berita bohong, dan informasi palsu yang beredar di ruang siber Indonesia (Komdigi, 2025). Temuan Tim AIS Subdirektorat Pengendalian Konten Direktorat Jenderal Aplikasi Informatika menunjukkan bahwa jumlah hoaks bersifat fluktuatif, dengan puncak pada Oktober 2024 sebanyak 215 konten dan terendah pada Februari 2024 sebanyak 131 konten. Dari keseluruhan temuan tersebut, hoaks bertema politik mencapai 237 konten, yang menegaskan bahwa isu politik menjadi salah satu sasaran utama penyebaran informasi palsu di ruang digital. Tingginya hoaks politik berpotensi membentuk opini publik yang keliru serta menurunkan kepercayaan masyarakat terhadap institusi negara dan proses demokrasi.

Hoaks politik memiliki karakteristik yang kompleks karena sering disampaikan melalui bahasa provokatif, informasi yang dipelintir, serta pembingkaihan berita yang menyesatkan. Intensitas penyebarannya cenderung meningkat pada periode pemilihan umum, ketika arus informasi politik di ruang digital semakin tinggi. Situasi tersebut menjadikan proses identifikasi hoaks secara manual tidak lagi efektif, karena memerlukan waktu yang lama dan rentan terhadap subjektivitas. Oleh karena itu, diperlukan pendekatan otomatis yang mampu mengklasifikasikan berita politik secara akurat dan konsisten berdasarkan tingkat kebenaran informasinya.

Sejalan dengan kebutuhan tersebut, klasifikasi teks dalam *machine learning* menjadi solusi yang relevan untuk mengelompokkan dokumen berdasarkan pola yang dipelajari dari data latih, sehingga memungkinkan pengolahan data teks dalam jumlah besar secara objektif. Salah satu algoritma yang banyak digunakan dalam klasifikasi teks adalah *Support Vector Machine* (SVM), yang bekerja dengan membentuk *hyperplane* optimal

dengan margin maksimum untuk memisahkan kelas data. SVM memiliki keunggulan dalam menangani data berdimensi tinggi dengan banyak fitur bernilai nol, seperti data teks yang direpresentasikan menggunakan pembobotan *Term Frequency-Inverse Document Frequency* (TF-IDF), sehingga mampu menghasilkan model yang efisien dan stabil. Pada penelitian ini, pendekatan klasifikasi teks berbasis SVM digunakan untuk membedakan berita politik berbahasa Indonesia ke dalam dua kategori, yaitu berita valid dan berita hoaks, melalui analisis karakteristik linguistik serta fitur kata pada setiap dokumen.

Penelitian terdahulu yang dilakukan oleh (Maulana et al., 2025) menunjukkan bahwa algoritma *Support Vector Machine* (SVM) memiliki performa yang unggul dalam klasifikasi teks berita hoaks. Penerapan pendekatan *text mining* berbasis SVM pada klasifikasi berita politik mampu menghasilkan kinerja yang konsisten dengan tingkat akurasi yang tinggi. Hasil evaluasi menunjukkan bahwa model yang dikembangkan mencapai akurasi sebesar 94,24%, serta didukung oleh nilai *precision*, *recall*, dan *f1-score* yang seimbang pada setiap kelas. Temuan ini mengindikasikan bahwa SVM efektif dalam mengidentifikasi pola linguistik pada teks berita dan mampu membangun model klasifikasi yang andal, sehingga metode tersebut dinilai relevan dan layak diterapkan dalam penelitian untuk mengklasifikasikan berita politik ke dalam kategori hoaks dan valid.

Selain performa, aspek interpretabilitas juga menjadi perhatian penting, terutama pada klasifikasi berita politik yang bersifat sensitif. Model yang tidak dapat dijelaskan berisiko menurunkan kepercayaan pengguna terhadap hasil klasifikasi. Untuk mengatasi hal tersebut, penelitian ini menerapkan pendekatan *Explainable Artificial Intelligence* (XAI) menggunakan metode *Local Interpretable Model-Agnostic Explanations* (LIME), yang memberikan penjelasan lokal dengan mengidentifikasi kata atau frasa yang paling berpengaruh terhadap keputusan model.

Berdasarkan uraian tersebut, penelitian ini bertujuan mengembangkan model klasifikasi berita politik berbahasa Indonesia menggunakan algoritma *Support Vector Machine* (SVM) dengan representasi fitur TF-IDF serta menerapkan metode LIME untuk meningkatkan interpretabilitas model. Pendekatan ini diharapkan mampu menghasilkan

sistem klasifikasi yang akurat, transparan, dan dapat mendukung upaya peningkatan literasi digital masyarakat terkait hoaks politik.

1.2 Rumusan Masalah

Penelitian ini berfokus pada penerapan algoritma *Support Vector Machine* (SVM) dan metode *Local Interpretable Model-Agnostic Explanations* (LIME) dalam klasifikasi berita politik berbahasa Indonesia secara akurat dan transparan. Berdasarkan fokus tersebut, rumusan masalah dalam penelitian ini dirumuskan sebagai berikut:

1. Bagaimana kinerja algoritma *Support Vector Machine* (SVM) dalam mengklasifikasikan berita politik ke dalam kategori hoaks dan valid?
2. Bagaimana hasil evaluasi model klasifikasi SVM dalam mengelompokkan berita politik berdasarkan metrik akurasi, *precision*, *recall*, dan *f1-score*?
3. Bagaimana penerapan pendekatan *Explainable Artificial Intelligence* (XAI) menggunakan metode LIME dapat memberikan penjelasan terhadap hasil klasifikasi yang dihasilkan oleh model SVM?

1.3 Tujuan

Berdasarkan rumusan masalah yang telah dirumuskan, penelitian ini bertujuan sebagai berikut:

1. Menganalisis kinerja algoritma *Support Vector Machine* (SVM) dalam mengklasifikasikan berita politik berbahasa Indonesia ke dalam kategori hoaks dan valid dengan menggunakan metrik evaluasi akurasi, *precision*, *recall*, dan *f1-score*.
2. Menerapkan pendekatan *Explainable Artificial Intelligence* (XAI) melalui metode *Local Interpretable Model-Agnostic Explanations* (LIME) untuk memberikan interpretasi terhadap hasil klasifikasi yang dihasilkan oleh model SVM.
3. Mengidentifikasi kata, frasa, serta pola linguistik yang berpengaruh signifikan dalam membedakan berita politik hoaks dan berita politik valid berdasarkan hasil interpretasi model XAI.

1.4 Manfaat

Hasil penelitian ini diharapkan memberikan manfaat sebagai berikut:

1. Memberikan kontribusi terhadap pengembangan keilmuan di bidang *Natural Language Processing* (NLP) dan *Explainable Artificial Intelligence* (XAI), khususnya dalam penerapan algoritma *Support Vector Machine* (SVM) dan metode *Local Interpretable Model-Agnostic Explanations* (LIME) untuk klasifikasi berita berbahasa Indonesia.
2. Menjadi bahan rujukan bagi penelitian selanjutnya dalam pengembangan model klasifikasi misinformasi yang transparan, interpretatif, dan akuntabel.
3. Memberikan pemahaman serta dukungan konseptual bagi praktisi media dan pemeriksa fakta (*fact-checker*) dalam menafsirkan hasil klasifikasi berita politik, sekaligus membantu masyarakat meningkatkan literasi digital dalam membedakan informasi yang valid dan tidak valid.

1.5 Ruang Lingkup

Ruang lingkup penelitian ini ditetapkan untuk memperjelas batasan, fokus, dan tahapan penelitian agar pelaksanaannya tetap terarah serta selaras dengan tujuan yang telah ditetapkan. Adapun ruang lingkup penelitian ini meliputi hal-hal sebagai berikut:

1. Penelitian ini dibatasi pada klasifikasi berita politik, sehingga analisis hanya difokuskan pada konten yang terkait dengan isu politik.
2. Data yang digunakan diperoleh melalui proses web scraping dari situs CNN Indonesia dan Turnbackhoax.id pada tahun 2024, dengan variabel yang dianalisis meliputi *timestamp*, *title*, *full text*, dan *category* dari setiap berita.
3. Metode yang diterapkan dalam penelitian ini adalah algoritma *Support Vector Machine* (SVM) untuk klasifikasi, yang dipadukan dengan pendekatan *Explainable Artificial Intelligence* (XAI) melalui metode *Local Interpretable Model-Agnostic Explanations* (LIME) untuk mempermudah interpretasi hasil model.
4. Pelaksanaan penelitian meliputi beberapa tahapan, dimulai dari persiapan data, yaitu pemberian label, penggabungan variabel dan dataset, pembersihan awal, serta pengacakan data. Selanjutnya dilakukan pra-pemrosesan data yang mencakup *casefolding*, *cleaning*, *tokenizing*, *stopword removal* dan *stemming*. Tahap

pemodelan mencakup pembagian data ke dalam jenis tertentu, uji coba berbagai kernel SVM, pemilihan jenis terbaik, pengujian hasil model menggunakan *confusion matrix*, serta penerapan metode LIME untuk memberikan interpretasi yang jelas dan transparan terhadap hasil klasifikasi.