

Klasifikasi Berita Hoaks pada Media *Online* Menggunakan *Open Source Intelligence* dan Algoritma *Support Vector Machine*

Muh. Darmawan Aryadinata¹ | Fahrin Irhamna Rachman^{*2} | Ida Mulyadi²

¹ Mahasiswa Program Studi Informatika,
Fakultas Teknik, Universitas
Muhammadiyah Makassar, Indonesia.

Email:

05841109621@student.unismuh.ac.id

² Program Studi Informatika, Fakultas Teknik,
Universitas Muhammadiyah Makassar,
Indonesia.

Email:

farachim141020@unismuh.ac.id

idamulyadi@unismuh.ac.id ;

Korespondensi:

^{*}Fahrin Irhamna Rachman

farachim141020@unismuh.ac.id

ABSTRAK

Perkembangan media *online* yang pesat memudahkan masyarakat dalam memperoleh informasi secara cepat dan luas. Namun, kemudahan tersebut juga meningkatkan penyebaran berita hoaks yang dapat menimbulkan dampak negatif bagi masyarakat. Oleh karena itu, diperlukan suatu sistem yang mampu mengklasifikasikan berita hoaks secara otomatis dan akurat. Penelitian ini bertujuan untuk membangun sistem klasifikasi berita hoaks pada media *online* menggunakan metode *Open Source Intelligence* (OSINT) dan algoritma *Support Vector Machine* (SVM). Dataset diperoleh dari sumber terbuka berbasis web, yaitu situs klarifikasi hoaks dan portal berita *online*, yang kemudian melalui tahapan *Preprocessing* meliputi pembersihan teks, normalisasi, dan tokenisasi. Proses ekstraksi fitur dilakukan menggunakan metode *Term Frequency-Inverse Document Frequency* (TF-IDF) untuk merepresentasikan teks dalam bentuk numerik. Model klasifikasi dibangun menggunakan algoritma SVM dengan *kernel linear* karena efektif dalam menangani data teks berdimensi tinggi. Hasil pengujian menunjukkan bahwa model yang dikembangkan mampu mengklasifikasikan berita hoaks dan non-hoaks dengan tingkat akurasi sebesar 96,20%, serta didukung oleh nilai *precision*, *recall*, dan *F1-score* yang tinggi. Hal ini menunjukkan bahwa kombinasi metode OSINT, TF-IDF, dan SVM efektif dalam membangun sistem klasifikasi berita hoaks berbasis web dengan performa yang baik.

Kata kunci: Klasifikasi Teks, Berita Hoaks, Media *Online*, OSINT, TF-IDF, *Support Vector Machine*.

ABSTRACT

The rapid development of online media makes it easier for people to obtain information quickly and widely. However, this convenience also increases the spread of hoax news, which can have negative impacts on society. Therefore, a system capable of automatically and accurately classifying hoax news is needed. This research aims to develop a hoax news classification system in online media using *Open Source Intelligence* (OSINT) methods and the *Support Vector Machine* (SVM) algorithm. The Dataset was obtained from web-based open sources, namely hoax clarification websites and online news portals. The data were then subjected to *Preprocessing* stages including text cleaning, normalization, and tokenization. The feature extraction process was carried out using the *Term Frequency-Inverse Document Frequency* (TF-IDF) method to represent text in numerical form. The classification model was built using the SVM algorithm with a linear kernel because it is effective in handling high-dimensional text data. Test results show that the developed model is capable of classifying hoax and non-hoax news with an accuracy rate of 96.20%, supported by high *precision*, *recall*, and *F1-score* values. This demonstrates that the combination of OSINT, TF-IDF, and SVM methods is effective in building a web-based hoax news classification system with good performance.

Keywords: Text Classification, Hoax News, Online Media, OSINT, TF-IDF, *Support Vector Machine*.

1 | PENDAHULUAN

Perkembangan teknologi informasi dan komunikasi telah mendorong penyebaran informasi secara masif di berbagai platform digital, terutama media *online*. Transformasi digital yang cepat menyebabkan arus informasi menjadi semakin tidak terbatas oleh ruang dan waktu. Seiring dengan hal tersebut, Kementerian Komunikasi dan Informatika (Kominfo, 2025) menyatakan bahwa peredaran konten hoaks di Indonesia masih menjadi permasalahan serius yang tersebar di berbagai platform digital,

dengan isu dominan meliputi bidang politik, kesehatan, dan ekonomi. Kondisi ini menunjukkan urgensi pengembangan sistem deteksi berita hoaks yang efisien, adaptif, dan mampu mengikuti dinamika penyebaran informasi di era digital.

Sejumlah penelitian ilmiah terbaru mendukung pentingnya penggunaan algoritma pembelajaran mesin dalam deteksi hoaks. Misalnya, penelitian oleh (Nugroho et al., 2025) menunjukkan bahwa algoritma seperti *Support Vector Machine* (SVM), *Logistic Regression*, dan *XGBoost* dapat dikombinasikan dengan representasi fitur *Term Frequency-Inverse Document Frequency* (TF-IDF) untuk mengklasifikasikan berita palsu dan nyata secara efektif dalam skenario konten digital berbasis teks. Selain itu, studi oleh (Desriansyah & Utma Sari, 2024) mengemukakan bahwa pendekatan SVM dan *Multilayer Perceptron* (MLP) mampu mencapai performa tinggi dalam mendeteksi berita hoaks berbahasa Indonesia setelah melalui tahap pra-pemrosesan teks dan ekstraksi fitur TF-IDF. Penelitian lain yang menggunakan algoritma SVM pada platform media sosial juga memperlihatkan bahwa model ini memberikan hasil klasifikasi yang sangat baik ketika dikombinasikan dengan teknik pemrosesan teks yang benar. Tantangan utama dalam menangani hoaks adalah kemampuannya meniru karakteristik berita kredibel, baik dari sisi struktur, gaya bahasa, maupun penyajian konten, sehingga sulit dideteksi secara manual. Verifikasi fakta yang dilakukan oleh lembaga seperti Masyarakat Anti Fitnah Indonesia (MAFINDO, 2024) belum mampu sepenuhnya mengimbangi kecepatan penyebaran informasi di platform seperti Facebook dan WhatsApp. Dalam konteks ini, pendekatan *Open Source Intelligence* (OSINT) menjadi semakin relevan karena memungkinkan pengumpulan dan analisis data secara otomatis dari berbagai sumber publik.

Sejumlah penelitian telah mengkaji deteksi hoaks dengan berbagai metode. (Dewi et al., 2025) menerapkan *Support Vector Machine* (SVM) untuk klasifikasi berita palsu berbasis teks dan memperoleh akurasi sebesar 92,4% pada *Dataset* berita Indonesia dari platform X. Penelitian (Rizal & Kurniawan, 2024) menggunakan kombinasi TF-IDF dan SVM untuk mendeteksi hoaks dengan hasil yang konsisten di atas 90%, menunjukkan efektivitas metode ini dalam konteks bahasa Indonesia. Sementara itu, (Prastyapradipta & Nurhadi, 2025) membandingkan *model Deep Learning* dengan metode tradisional seperti SVM dan menemukan bahwa SVM tetap kompetitif dalam kondisi *Dataset* terbatas. Namun, penelitian-penelitian tersebut cenderung mengabaikan aspek intelijen sumber terbuka (OSINT) yang dapat memperkaya fitur data dari dimensi nontekstual seperti metadata akun, pola penyebaran, dan waktu publikasi, yang terbukti penting dalam mendeteksi pola hoaks yang terorganisir (Alshuwaier & Alsulaiman, 2025). Keterbatasan penelitian sebelumnya terletak pada fokus yang sempit terhadap aspek tekstual berita, tanpa memperhatikan konteks sosial dan jaringan penyebaran informasi. Sebagian besar penelitian hanya menggunakan *Dataset* statis tanpa integrasi terhadap sumber data dinamis seperti media *online real-time* yang dapat diakses melalui OSINT. Selain itu, penggunaan SVM sering kali terbatas pada optimasi parameter dasar tanpa eksplorasi terhadap *feature engineering* berbasis data intelijen terbuka. Oleh karena itu, terdapat kebutuhan akan pendekatan integratif yang menggabungkan *Open Source Intelligence* dengan algoritma *Support Vector Machine*, yang tidak hanya menganalisis isi berita tetapi juga memperhitungkan konteks distribusi informasi. Pendekatan ini diharapkan dapat menghasilkan model klasifikasi yang lebih adaptif, komprehensif, dan kontekstual dalam mendeteksi hoaks di media *online* Indonesia.

Penelitian ini bertujuan untuk mengembangkan model klasifikasi berita hoaks pada media *online* dengan memanfaatkan kombinasi *Open Source Intelligence* dan algoritma *Support Vector Machine* (SVM) untuk meningkatkan akurasi dan efisiensi deteksi berita palsu. Secara teoritis, penelitian ini diharapkan dapat memperkaya literatur tentang penerapan OSINT dalam bidang analisis data dan *Machine Learning*, khususnya pada konteks bahasa dan sosial Indonesia. Secara praktis, hasil penelitian dapat memberikan manfaat bagi lembaga pemeriksa fakta, pemerintah, dan platform media *online* dalam mendeteksi dan menekan penyebaran informasi palsu secara otomatis dan *real-time*. Dengan demikian, penelitian ini berkontribusi dalam memperkuat ekosistem informasi digital yang sehat, terpercaya, dan berintegritas di Indonesia.

2 | METODE

Penelitian ini menggunakan pendekatan kuantitatif dengan metode klasifikasi teks untuk mengidentifikasi berita hoaks pada media *online*. Proses penelitian dilakukan secara sistematis yang meliputi tahapan pengumpulan data, *Preprocessing* teks, ekstraksi fitur, pembangunan model klasifikasi, serta evaluasi performa model. Setiap tahapan dirancang untuk memastikan bahwa sistem yang dikembangkan mampu mengklasifikasikan berita secara otomatis dan akurat berdasarkan pola teks yang terkandung dalam dokumen.

2.1 | Pengumpulan Data (OSINT)

Open Source Intelligence (OSINT) adalah metode pengumpulan, analisis, dan pemanfaatan informasi dari sumber publik yang dapat diakses secara bebas. Konsep ini berasal dari dunia intelijen militer dan keamanan, namun dalam dua dekade terakhir telah diadaptasi secara luas di bidang keamanan siber, analisis data, serta penelitian sosial digital (Lowenthal, 2022). Pengumpulan data pada penelitian ini dilakukan menggunakan metode *Open Source Intelligence* (OSINT), yaitu teknik pengambilan informasi dari sumber terbuka seperti portal berita daring dan media *online*. Proses *scraping* dilakukan untuk memperoleh teks berita yang relevan dengan topik hoaks dan fakta.

Pendekatan OSINT mengandalkan kombinasi antara *data mining*, *web scraping*, dan *API integration* untuk mengekstraksi data dalam jumlah besar secara otomatis. Data yang dikumpulkan dapat berupa teks, gambar, video, metadata akun, serta pola waktu publikasi. Dalam penelitian deteksi hoaks, data tersebut digunakan sebagai bahan mentah untuk pelatihan algoritma klasifikasi (Anwar et al., 2023). Data penelitian dikumpulkan menggunakan pendekatan *Open Source Intelligence* (OSINT) dari media *online* dan situs pemeriksa fakta. *Dataset* yang diperoleh kemudian melalui proses validasi dan pelabelan berdasarkan sumber terpercaya seperti *turnbackhoax.id*, *komdigi.go.id*, *kompas.com*. *Dataset* yang diperoleh kemudian melalui proses pelabelan menjadi dua kategori, yaitu:

- a. Kelas 1 : Hoaks
- b. Kelas 0 : Fakta

Total data yang digunakan dalam penelitian ini sebanyak ± 780 data, yaitu:

- a. ± 340 Data hoaks
- b. ± 440 Data fakta

Dataset tersebut selanjutnya diproses ke tahap *Preprocessing* untuk meningkatkan kualitas teks sebelum dilakukan ekstraksi fitur. Berikut merupakan contoh beberapa data yang digunakan dalam penelitian ini, khususnya mengenai sampel *Dataset* fakta yang telah dikumpulkan dan diberi label.

2.2 | Preprocessing Data

Sebelum dilakukan proses klasifikasi, data teks melalui tahap *Preprocessing* untuk meningkatkan kualitas data. Tahapan yang dilakukan meliputi:

- a. *Cleaning* – Menghapus tanda baca, angka, simbol, dan karakter khusus.
- b. *Case Folding* – Mengubah seluruh teks menjadi huruf kecil.
- c. *Tokenization* – Memecah teks menjadi token/kata.
- d. *Stopword Removal* – Menghapus kata-kata umum yang tidak memiliki makna signifikan.

Tahapan ini bertujuan untuk mengurangi noise dan meningkatkan performa model klasifikasi.

2.3 | Pembagian Data (Data Splitting)

Sebelum proses pelatihan model, *Dataset* dibagi menjadi dua bagian, yaitu data latih (*training data*) dan data uji (*testing data*). Pembagian data dilakukan menggunakan metode *hold-out* dengan rasio 80:20, di mana 80% data digunakan untuk pelatihan model dan 20% digunakan untuk pengujian. Data latih digunakan untuk membangun model klasifikasi menggunakan algoritma *Support Vector Machine* (SVM), sedangkan data uji digunakan untuk mengevaluasi kemampuan generalisasi model terhadap data yang belum pernah dilihat sebelumnya.

Pembagian data dilakukan secara acak untuk menghindari bias distribusi kelas. Dengan metode ini, model dapat diuji secara objektif dalam mengklasifikasikan berita hoaks dan non-hoaks. *Dataset* yang telah diproses kemudian dibagi menjadi data *training* dan data *testing* dengan rasio 80:20.

- a. 80% digunakan sebagai *data training*
- b. 20% digunakan sebagai *data testing*

Pembagian dilakukan menggunakan metode *random splitting* dengan parameter *random_state* untuk menjaga konsistensi hasil. Apabila distribusi kelas tidak seimbang, digunakan teknik *stratified sampling* untuk mempertahankan proporsi kelas pada kedua subset data. Tahap ini bertujuan untuk mengukur kemampuan generalisasi model terhadap data yang belum pernah dilatih sebelumnya.

Contoh hasil *splitting* data :

Label 0 (Fakta) : 449 (56.91%)

Label 1 (Hoaks) : 340 (43.09%)

HASIL SPLITTING DATA (80:20)

Jumlah data latih : 631 Jumlah data uji : 158

Distribusi data latih: Fakta (0) : 359 - Hoaks (1) : 272

Distribusi data uji: Fakta (0) : 90 - Hoaks (1) : 68

2.4 | Ekstraksi Fitur TF-IDF .

Metode *Term Frequency-Inverse Document Frequency* (TF-IDF) digunakan untuk mengubah data teks menjadi representasi numerik dalam bentuk vektor. Proses yang dilakukan:

- a. TF-IDF *fit_transform* pada data *training*
- b. TF-IDF *transform* pada data *testing*

Hasil ekstraksi menghasilkan matriks fitur berdimensi:

a. Data *training* : (631, 1810)

b. Data *testing* : (158, 1810)

Representasi ini memungkinkan algoritma SVM untuk melakukan proses klasifikasi dalam ruang vektor berdimensi tinggi.

1) Contoh Data Training

Pada tahap *training*, teks terlebih dahulu melalui proses *Preprocessing* seperti:

- a) *Case folding* (semua huruf menjadi kecil)
- b) *Tokenizing* (memecah kalimat menjadi kata)
- c) *Stopword removal* (jika digunakan)
- d) *Stemming* (jika digunakan)

Setelah itu dilakukan proses ekstraksi fitur menggunakan TF-IDF (*Term Frequency – Inverse Document Frequency*).

- (1) TF (*Term Frequency*) → seberapa sering kata muncul dalam dokumen
- (2) IDF (*Inverse Document Frequency*) → seberapa unik kata tersebut dalam seluruh *Dataset*

Semakin sering kata muncul dalam satu dokumen tetapi jarang muncul di dokumen lain, maka bobotnya semakin besar.

CONTOH (1) OUTPUT DATA TRAINING

Teks Asli:

→ hoaks video raffi ahmad promosi situs judol menang

Bobot TF-IDF (*Training*):

ahmad hoaks judol menang promosi raffi situs video 0.3808942 0.3808 0.4243 0.3668 0.4243 0.3668 0.2076

(1) Kata “raffi” dan “menang” memiliki bobot tertinggi (0.4243) → artinya kata tersebut cukup spesifik dan jarang muncul di dokumen lain.

(2) Kata “hoaks” memiliki bobot lebih kecil (0.1942) → kemungkinan karena sering muncul di banyak dokumen dalam *Dataset* sehingga nilai IDF-nya rendah.

(3) Kata yang sering muncul di banyak berita akan memiliki bobot lebih kecil karena kurang diskriminatif.

Bobot-bobot ini kemudian menjadi vektor numerik yang digunakan sebagai input ke algoritma *Support Vector Machine* (SVM) untuk proses pembelajaran pola.

2) Contoh Data Testing

- a) Pada tahap *testing*, prosesnya hampir sama:
- b) Teks dibersihkan
- c) Diubah menjadi token
- d) Dikonversi menjadi bobot TF-IDF, *Vocabulary* (daftar kata) yang digunakan pada *testing* harus mengikuti *vocabulary* dari data *training*, bukan membuat *vocabulary* baru.

CONTOH (1) OUTPUT DATA TESTING

Teks Asli:

→ hoaks video menkeu Purbaya marah sri mulyani Bobot TF-IDF (*Testing*):

hoaks marah menkeu mulyani Purbaya sri video 0.2434219 0.3969 0.4764 0.2898 0.4764 0.2598

- (1) Kata “sri” dan “mulyani” memiliki bobot tertinggi (0.4764) → menunjukkan kata tersebut cukup unik dalam *Dataset*.
- (2) Kata “hoaks” dan “video” memiliki bobot lebih rendah → kemungkinan karena sering muncul dalam banyak dokumen.
- (3) Nilai bobot berbeda dengan *training* karena perhitungan IDF didasarkan pada seluruh *Dataset training*.

3) Perbedaan Training dan testing

Setelah proses ekstraksi fitur menggunakan TF-IDF dilakukan, tahapan selanjutnya dalam penelitian ini adalah membagi *Dataset* menjadi data *training* dan data *testing*. Pembagian ini bertujuan untuk memastikan bahwa model yang dibangun tidak hanya mampu mengenali pola pada data yang telah dipelajari, tetapi juga mampu melakukan generalisasi terhadap data baru yang belum pernah dilihat sebelumnya. Data *training* digunakan untuk membentuk pola klasifikasi berdasarkan fitur yang telah diekstraksi, sedangkan data *testing* digunakan untuk mengevaluasi kinerja model dalam melakukan prediksi secara objektif. Perbedaan fungsi dan peran kedua jenis data tersebut dapat dilihat pada tabel 6 berikut.

Representasi TF-IDF menghasilkan matriks fitur berdimensi tinggi yang menggambarkan bobot setiap kata pada setiap dokumen. Matriks inilah yang digunakan sebagai input dalam proses klasifikasi menggunakan algoritma SVM.

2.5 | Evaluasi Model

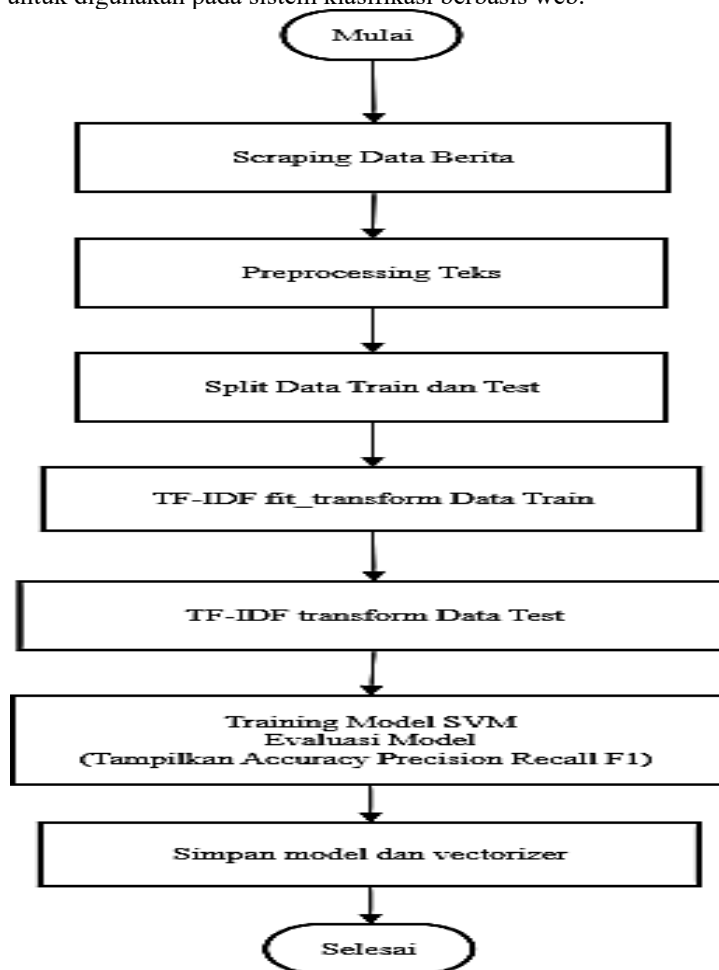
Evaluasi kinerja model dilakukan menggunakan *Confusion matrix* dan beberapa metrik evaluasi, yaitu:

- *Accuracy*
- *Precision*
- *Recall*
- *F1-score*

Metrik tersebut digunakan untuk menilai kemampuan model dalam mengklasifikasikan berita hoaks dan non-hoaks secara tepat. Setelah seluruh tahapan penelitian dijelaskan secara sistematis, mulai dari proses pengumpulan data menggunakan pendekatan *Open Source Intelligence* (OSINT), *Preprocessing teks*, pembagian data, ekstraksi fitur menggunakan TF-IDF, hingga pelatihan dan evaluasi model *Support Vector Machine* (SVM), diperlukan representasi visual untuk memperjelas alur

kerja sistem yang dikembangkan. Visualisasi dalam bentuk flowchart bertujuan untuk memberikan gambaran menyeluruh mengenai tahapan proses yang dilakukan secara berurutan, sehingga memudahkan pembaca dalam memahami struktur metodologi penelitian.

Flowchart berikut menggambarkan alur proses mulai dari tahap *scraping* data berita hingga penyimpanan model dan vectorizer yang telah dilatih untuk digunakan pada sistem klasifikasi berbasis web.



GAMBAR 1. Flowchart Metodologi Sistem Klasifikasi Berita Hoaks

3 | HASIL

Hasil Berdasarkan hasil pengujian terhadap data uji, model klasifikasi yang dibangun menunjukkan performa yang sangat baik. Hasil evaluasi model menunjukkan tingkat akurasi sebesar 96,20%. Selain itu, nilai precision, recall, dan *F1-score* yang diperoleh juga menunjukkan performa yang tinggi pada kedua kelas. *Confusion matrix* menunjukkan bahwa sebagian besar data berhasil diklasifikasikan dengan benar, dengan jumlah kesalahan klasifikasi yang relatif kecil.

Untuk mengukur kinerja model klasifikasi yang telah dibangun, dilakukan evaluasi menggunakan *Confusion matrix* dan beberapa metrik performa, yaitu *accuracy*, *precision*, *recall*, dan *F1-score*. *Metrik-metrik* tersebut digunakan untuk menilai kemampuan model dalam membedakan berita hoaks dan non-hoaks secara tepat. *Accuracy* digunakan untuk mengukur tingkat ketepatan prediksi secara keseluruhan, sedangkan *precision* dan *recall* digunakan untuk mengevaluasi keseimbangan antara kesalahan positif dan kemampuan model dalam mendeteksi kelas tertentu. *F1-score* digunakan sebagai indikator harmonisasi antara *precision* dan *recall*. Hasil evaluasi model terhadap data uji disajikan pada Tabel 1 berikut.

TABEL 1 Nilai Evaluasi

Metrik	Nilai
<i>Accuracy</i>	96.20%
<i>Precision</i>	96.97%
<i>Recall</i>	94.12%
<i>F1-score</i>	95.52%

Hasil tersebut menunjukkan bahwa kombinasi metode OSINT, TF-IDF, dan SVM mampu menghasilkan sistem klasifikasi berita hoaks dengan tingkat akurasi dan konsistensi yang baik.

4 | PEMBAHASAN

Berdasarkan hasil penelitian yang telah dilakukan, dapat dianalisis bahwa kombinasi metode *Open Source Intelligence* (OSINT), *Natural Language Processing* (NLP), TF-IDF, dan algoritma *Support Vector Machine* (SVM) mampu menghasilkan sistem klasifikasi berita hoaks dengan performa yang sangat baik. Integrasi beberapa metode tersebut memberikan kontribusi penting dalam setiap tahapan proses klasifikasi, mulai dari pengumpulan data, *Preprocessing* teks, ekstraksi fitur, hingga proses prediksi berita hoaks dan fakta.

Metode *Open Source Intelligence* (OSINT) terbukti efektif dalam proses pengumpulan data berita dari berbagai sumber terbuka di internet. Di Indonesia, penelitian berbasis OSINT masih tergolong baru. (Setiawan et al., 2023) mengimplementasikan OSINT untuk pengumpulan data terkait hoaks politik dari Facebook dan Telegram menggunakan *web scraping* dan API monitoring. Pendekatan ini memungkinkan *Dataset* diperoleh secara lebih cepat, luas, dan relevan terhadap kondisi aktual penyebaran informasi di media *online*. Penggunaan OSINT memberikan keunggulan dibanding pengumpulan *Dataset* secara manual karena data yang diperoleh bersifat dinamis dan dapat mengikuti perkembangan isu yang sedang viral di masyarakat. Selain itu, penggunaan sumber terpercaya seperti turnbackhoax.id, kompas.com, dan komdigi.go.id membantu meningkatkan validitas data yang digunakan dalam proses pelatihan model.

Tahapan *Preprocessing* juga memiliki pengaruh yang sangat besar terhadap kualitas model klasifikasi. Pada penelitian ini dilakukan beberapa tahapan *Preprocessing*, yaitu *case folding*, *tokenizing*, *stopword removal*, dan *stemming*. Tahapan tersebut bertujuan untuk membersihkan teks dari karakter atau kata-kata yang tidak relevan sehingga informasi penting dalam teks dapat dipertahankan. Proses *case folding* membantu menyeragamkan seluruh huruf menjadi lowercase agar tidak terjadi perbedaan representasi kata akibat penggunaan huruf kapital. Selanjutnya, *tokenizing* digunakan untuk memecah kalimat menjadi token-token kata sehingga lebih mudah dianalisis oleh sistem. Tahapan *stopword removal* berfungsi untuk menghapus kata-kata umum yang tidak memiliki kontribusi besar terhadap proses klasifikasi, seperti kata penghubung atau kata yang sering muncul dalam berbagai dokumen. Setelah itu, *stemming* dilakukan untuk mengubah kata menjadi bentuk dasarnya sehingga variasi kata dapat direpresentasikan secara lebih konsisten. Dengan adanya *Preprocessing* tersebut, teks menjadi lebih terstruktur dan kualitas data meningkat sehingga mampu membantu algoritma SVM mengenali pola klasifikasi dengan lebih baik. Metode TF-IDF (*Term Frequency–Inverse Document Frequency*) yang digunakan dalam penelitian ini juga terbukti efektif dalam merepresentasikan teks ke dalam bentuk numerik. TF-IDF memberikan bobot lebih besar terhadap kata-kata yang sering muncul pada suatu dokumen tetapi jarang muncul pada keseluruhan *Dataset*. *Natural Language Processing* (NLP) merupakan cabang dari kecerdasan buatan yang berfokus pada bagaimana komputer dapat memahami, memproses, dan menganalisis bahasa manusia secara otomatis (Jurafsky & Martin, 2023). Hal ini memungkinkan sistem mengenali kata-kata yang memiliki karakteristik kuat terhadap suatu kategori berita. Sebagai contoh, kata-kata seperti “hadiah”, “viral”, “klik”, atau “sebar” cenderung memiliki hubungan dengan berita hoaks sehingga bobot TF-IDF menjadi lebih tinggi. Sebaliknya, kata-kata umum yang sering muncul pada semua dokumen akan memiliki bobot lebih rendah karena dianggap kurang penting dalam proses klasifikasi.

Hasil ekstraksi fitur TF-IDF menghasilkan representasi data berdimensi tinggi, yaitu sebanyak 1810 fitur. Dalam kondisi tersebut, algoritma *Support Vector Machine* (SVM) dengan kernel linear menunjukkan performa yang sangat optimal. SVM bekerja dengan membangun *hyperplane* optimal yang mampu memisahkan data ke dalam kelas-kelas berbeda dengan margin maksimum, sehingga menghasilkan performa klasifikasi yang stabil dan akurat (Awad & Khanna, 2022). Kombinasi TF-IDF dan SVM terbukti efektif dalam berbagai penelitian terkini untuk klasifikasi berita hoaks karena mampu merepresentasikan karakteristik teks secara numerik dan meningkatkan akurasi prediksi model (Hanif et al., 2025). Hasil evaluasi model menunjukkan nilai *accuracy* sebesar 96.20%, *precision* sebesar 96.97%, *recall* sebesar 94.12%, dan *F1-score* sebesar 95.52%. Nilai *accuracy* menunjukkan bahwa model mampu mengklasifikasikan sebagian besar data uji dengan benar. Dari total 158 data *testing*, model berhasil memprediksi 152 data secara tepat. Hal ini membuktikan bahwa kombinasi *Preprocessing*, TF-IDF, dan SVM sangat efektif dalam mendeteksi perbedaan pola antara berita hoaks dan fakta.

Nilai *precision* sebesar 96.97% menunjukkan bahwa ketika model memberikan label “Hoaks”, maka prediksi tersebut hampir selalu benar. Tingginya *precision* sangat penting dalam sistem deteksi hoaks karena dapat mengurangi kemungkinan berita fakta diklasifikasikan sebagai hoaks. Kesalahan semacam ini dapat menyebabkan misinformasi baru apabila sistem digunakan secara luas oleh masyarakat. Oleh karena itu, tingginya *precision* menjadi indikator bahwa model memiliki tingkat kepercayaan yang baik dalam proses klasifikasi.

Sementara itu, nilai *recall* sebesar 94.12% menunjukkan bahwa model mampu menemukan sebagian besar berita hoaks yang terdapat pada *Dataset* pengujian. *Recall* menjadi metrik penting dalam deteksi hoaks karena berkaitan langsung dengan kemampuan sistem dalam menangkap berita palsu yang beredar di masyarakat. Semakin tinggi *recall*, maka semakin sedikit berita hoaks yang lolos dari sistem deteksi. Meskipun nilainya sedikit lebih rendah dibanding *precision*, nilai *recall* pada penelitian ini tetap menunjukkan performa yang sangat baik. Nilai *F1-score* sebesar 95.52% menunjukkan bahwa model memiliki keseimbangan yang baik antara *precision* dan *recall*. Hal ini membuktikan bahwa model tidak hanya fokus pada satu

aspek evaluasi saja, tetapi mampu mempertahankan performa klasifikasi secara konsisten pada kedua kelas data. Dengan demikian, model dinilai cukup stabil dan layak untuk diimplementasikan pada sistem berbasis web.

Apabila dibandingkan dengan penelitian terdahulu, hasil penelitian ini menunjukkan performa yang kompetitif bahkan lebih tinggi pada beberapa aspek evaluasi. Penelitian oleh Dewi et al. (2025) memperoleh akurasi sebesar 92,4% menggunakan metode TF-IDF dan SVM pada data media sosial, sedangkan penelitian ini mampu mencapai akurasi 96,20%. Peningkatan performa tersebut dipengaruhi oleh kualitas *Preprocessing*, penggunaan *Dataset* hasil OSINT, serta pemanfaatan sumber data yang telah diverifikasi terlebih dahulu. Selain itu, integrasi OSINT memberikan keunggulan tambahan karena data yang digunakan lebih dinamis dan relevan dengan kondisi aktual di media *online*.

Implementasi model ke dalam sistem berbasis web juga menunjukkan hasil yang baik. Sistem mampu memberikan hasil klasifikasi secara otomatis dan *real-time* dengan waktu respons yang relatif cepat. Penggunaan *framework Streamlit* membantu proses pengembangan aplikasi menjadi lebih sederhana karena integrasi model *Machine Learning* dengan antarmuka pengguna dapat dilakukan secara langsung menggunakan bahasa pemrograman Python. Antarmuka sistem yang sederhana juga memudahkan pengguna umum dalam melakukan pengecekan berita secara mandiri tanpa memerlukan pemahaman teknis terkait *Machine Learning*. Meskipun demikian, penelitian ini masih memiliki beberapa keterbatasan. Salah satu keterbatasan utama adalah jumlah *Dataset* yang masih relatif terbatas sehingga kemampuan generalisasi model dapat ditingkatkan lebih lanjut apabila jumlah data diperbesar. Selain itu, sistem yang dibangun masih menggunakan klasifikasi biner, yaitu hanya membedakan berita hoaks dan fakta. Sistem belum mampu mengidentifikasi jenis hoaks yang lebih spesifik seperti satire, propaganda, *clickbait*, atau *misleading information*.

Selain itu, keterbatasan teknis pada model berbasis TF-IDF dan SVM yang berfokus pada frekuensi kata membuat sistem ini belum mampu memahami konteks semantik bahasa alami secara mendalam. Akibatnya, potensi misklasifikasi masih dapat terjadi ketika sistem dihadapkan pada kalimat ambigu, sarkasme, atau ironi. Oleh karena itu, penelitian selanjutnya disarankan untuk menerapkan metode *Deep Learning* seperti *Long Short-Term Memory (LSTM)*, *Bidirectional Encoder Representations from Transformers (BERT)*, atau arsitektur *Transformer* lainnya guna meningkatkan sensitivitas model terhadap konteks bahasa. Kendati memiliki beberapa keterbatasan tersebut, secara keseluruhan penelitian ini telah memberikan kontribusi penting melalui keberhasilan integrasi OSINT, *preprocessing* NLP, TF-IDF, dan SVM dalam menghasilkan sistem klasifikasi berita hoaks berbasis web yang stabil, responsif, dan aplikatif bagi masyarakat luas.

Penelitian ini memiliki beberapa keterbatasan. Keterbatasan utama terletak pada skala dataset, sehingga perluasan data di masa depan diharapkan dapat meningkatkan generalisasi model. Dari aspek fungsionalitas, sistem saat ini terbatas pada klasifikasi biner (hoaks dan fakta) tanpa mengidentifikasi subjenis hoaks seperti satire, propaganda, *clickbait*, atau *misleading information*. Secara teknis, kombinasi TF-IDF dan SVM yang berbasis frekuensi kata juga belum optimal dalam menangkap konteks semantik yang kompleks, membuat sistem ini sensitif terhadap ambiguitas, sarkasme, dan ironi. Oleh karena itu, penerapan metode *Deep Learning* seperti LSTM, BERT, atau arsitektur *Transformer* sangat direkomendasikan untuk penelitian selanjutnya. Meski demikian, penelitian ini secara umum memberikan kontribusi penting melalui keberhasilan integrasi OSINT, *preprocessing* NLP, TF-IDF, dan SVM dalam membangun arsitektur web klasifikasi hoaks yang berkinerja stabil.

5 | KESIMPULAN

Berdasarkan hasil penelitian yang telah dilakukan, maka dapat ditarik kesimpulan sebagai berikut:

1. Sistem klasifikasi berita hoaks pada media *online* berhasil dibangun dengan memanfaatkan pendekatan *Open Source Intelligence (OSINT)* untuk pengumpulan data, tahapan *Preprocessing* menggunakan *Natural Language Processing (NLP)*, ekstraksi fitur TF-IDF, serta algoritma *Support Vector Machine (SVM)* sebagai metode klasifikasi. Sistem yang dikembangkan telah berhasil diimplementasikan dalam bentuk aplikasi berbasis web yang mampu melakukan klasifikasi berita secara otomatis.

2. Algoritma *Support Vector Machine (SVM)* menunjukkan kinerja yang sangat baik dalam mengklasifikasikan berita hoaks berbahasa Indonesia. Berdasarkan hasil pengujian menggunakan data *testing*, model memperoleh nilai *accuracy* sebesar 96,20%, *precision* sebesar 96,97%, *recall* sebesar 94,12%, dan *F1-score* sebesar 95,52%. Hasil tersebut menunjukkan bahwa kombinasi *Preprocessing* NLP, TF-IDF, dan SVM efektif dalam membedakan berita hoaks dan fakta dengan tingkat akurasi yang tinggi.

Namun keterbatasan teknis pada model berbasis TF-IDF dan SVM yang berfokus pada frekuensi kata membuat sistem ini belum mampu memahami konteks semantik bahasa alami secara mendalam. Akibatnya, potensi misklasifikasi masih dapat terjadi ketika sistem dihadapkan pada kalimat ambigu, sarkasme, atau ironi. Oleh karena itu, penelitian selanjutnya disarankan untuk menerapkan metode *Deep Learning* seperti *Long Short-Term Memory (LSTM)*, *Bidirectional Encoder Representations from Transformers (BERT)*, atau arsitektur *Transformer* lainnya guna meningkatkan sensitivitas model terhadap konteks bahasa.

Daftar Pustaka

- Alshuwaier, F. A., & Alsulaiman, H. A. (2025). Integrating *Open Source Intelligence* with *Machine Learning* for Social Media Threat Detection. *Journal of Information Security and Intelligence*, 14(2), 88–97.
- Anwar, S., Putra, H., & Rahman, F. (2023). OSINT-Based Data Collection Framework for Fake News Detection. *Indonesian Journal of Computer Science*, 18(3), 112–122.
- Awad, M., & Khanna, R. (2022). *Efficient Learning Machines: Theories, Concepts, and Applications for Engineers and System Designers*. Springer. <https://doi.org/10.1007/978-3-030-84499-2>
- Desriansyah, D., & Utna Sari, E. (2024). Detection of Indonesian Fake News Using *Support Vector Machine* and Multilayer Perceptron. *Jurnal Ilmiah Sistem Dan Komputer*, 9(1), 23–34. <https://jurnal.unidha.ac.id/index.php/jiska/article/view/2024>
- Dewi, L. A., Sari, M., & Nugraha, D. (2025). SVM-Based Fake News Detection on Indonesian Social Media Data. *Journal of Data Science and Artificial Intelligence*, 7(1), 45–54.
- Dewi, R., Prasetyo, D., & Hasanah, U. (2025). Fake News Detection in Indonesian Texts Using SVM with TF-IDF Features. *Jurnal Teknologi Informasi Indonesia*, 12(2), 65–73.
- Hanif, R., Santoso, A., & Puspita, D. (2025). Hybrid SVM and OSINT-Driven Feature Enrichment for COVID-19 Hoax Detection in Indonesian Social Media. *Asian Journal of Data Analytics*, 5(2), 78–90.
- Jurafsky, D., & Martin, J. H. (2023). *Speech and Language Processing* (3rd ed.). Prentice Hall.
- Kominfo. (2025). *Laporan Statistik Hoaks dan Disinformasi Tahun 2024*. Kementerian Komunikasi dan Informatika Republik Indonesia. <https://kominfo.go.id>
- Lowenthal, M. M. (2022). *Intelligence: From Secrets to Policy* (9th ed.). CQ Press
- MAFINDO. (2024). *Laporan Tahunan Cek Fakta Indonesia 2024*. Masyarakat Anti Fitnah Indonesia. <https://turnbackhoax.id>
- Nugroho, S., & Andini, P. (2024). Hoax detection system on social media using text mining. *Jurnal Sistem Informasi*, 20(1), 45–54.
- Prastyapradipta, D., & Nurhadi, A. (2025). Comparative Analysis of *Deep Learning* and SVM for Indonesian Fake News Detection. *Procedia Computer Science*, 230, 521–529.
- Rizal, A., & Kurniawan, B. (2024). Enhancing Fake News Detection in Indonesian Language Using TF-IDF and N-Gram Features. *Indonesian Journal of Information Systems*, 10(4), 233–241.
- Setiawan, E., Rahmawati, D., & Hidayat, F. (2023). OSINT Implementation for Political Hoax Data Collection on Social Media. *Journal of Cyber Intelligence and Data Mining*, 9(2), 102–110.