

## KLASIFIKASI BERITA HOAKS DAN FAKTA BERDASARKAN REPRESENTASI TF-IDF MENGGUNAKAN SUPPORT VECTOR MACHINE

Dinda Aulika Elfarin Aritonang, Putri Maharani,  
Nabila Ramadhani, Ken Ditha Tania, Allsela Meiriza

Sistem Informasi, Universitas Sriwijaya  
Jl. Raya Palembang - Prabumulih No.KM. 32, Indralaya Indah, Kec. Indralaya,  
Kabupaten Ogan Ilir, Sumatera Selatan, Indonesia  
09031382328131@student.unsri.ac.id

### ABSTRAK

Perkembangan media digital yang pesat menyebabkan penyebaran informasi berlangsung sangat cepat, termasuk munculnya berita hoaks yang dapat menimbulkan misinformasi di masyarakat. Permasalahan dalam penelitian ini adalah sulitnya membedakan berita hoaks dan fakta secara otomatis karena tingginya volume informasi yang beredar. Penelitian ini dilakukan untuk menganalisis kinerja algoritma Support Vector Machine (SVM) dalam membedakan berita hoaks dan fakta berbahasa Indonesia menggunakan fitur Term Frequency–Inverse Document Frequency (TF-IDF). Dataset yang dipakai berjumlah 1.896 berita yang terdiri dari 948 berita hoaks dan 948 berita fakta. Langkah-langkah penelitian mencakup pembersihan teks, ekstraksi fitur dengan TF-IDF, pembagian data menjadi data pelatihan (80%) dan data pengujian (20%), serta proses pengklasifikasian dengan SVM. Penilaian kinerja model dilakukan dengan confusion matrix dengan metrik *accuracy*, *precision*, *recall*, dan *F1-score*. Hasil penelitian menunjukkan bahwa model ini berhasil mencapai akurasi 97% dengan nilai *precision*, *recall*, dan *F1-score* masing-masing sebesar 0,97. Temuan ini menandakan bahwa metode SVM yang menggunakan TF-IDF efektif dalam mengenali berita hoaks dan berita yang benar.

**Kata kunci :** Deteksi Hoaks, Klasifikasi Berita, Text Mining, TF-IDF, Support Vector Machine

### 1. PENDAHULUAN

Pesatnya perkembangan teknologi digital mendorong peningkatan produksi serta distribusi informasi melalui berbagai platform media daring. Informasi dapat diakses secara cepat dan luas oleh masyarakat, namun sering kali tidak melalui proses verifikasi yang memadai. Kondisi tersebut menyebabkan meningkatnya penyebaran berita hoaks yang berpotensi menimbulkan misinformasi, memengaruhi stabilitas sosial, serta menurunkan tingkat kepercayaan masyarakat terhadap informasi digital [1].

Di Indonesia, tingginya intensitas penggunaan internet dan media sosial turut mempercepat penyebaran informasi yang belum terverifikasi. Ini berisiko menciptakan persepsi masyarakat yang tidak sesuai dengan kenyataan. Oleh karena itu, dibutuhkan suatu pendekatan yang mampu mengidentifikasi serta mengklasifikasikan berita secara otomatis agar dapat membedakan informasi hoaks dan fakta secara lebih akurat.

Dalam kajian text mining, proses klasifikasi teks umumnya dilakukan melalui beberapa tahapan, mulai dari pembersihan data (*preprocessing*), pembentukan fitur, hingga penerapan algoritma klasifikasi. Pendekatan machine learning banyak dimanfaatkan karena kemampuannya dalam mengolah data teks dalam skala besar serta mengenali pola secara otomatis [10].

Sejumlah penelitian terdahulu menunjukkan bahwa algoritma Support Vector Machine (SVM)

memiliki performa yang baik dalam klasifikasi teks berbahasa Indonesia. Tri Wahyudi menyatakan bahwa kombinasi SVM dengan metode TF-IDF mampu menghasilkan kinerja klasifikasi yang optimal [9].

Selain itu, penelitian oleh Dewi menunjukkan efektivitas TF-IDF dan SVM dalam mendeteksi berita hoaks [2], sementara Sani mengungkapkan bahwa SVM menunjukkan kinerja yang lebih baik dibandingkan Naive Bayes [3].

Penelitian lain oleh Tandiano juga menegaskan bahwa penggunaan TF-IDF dapat meningkatkan kualitas representasi fitur dalam proses klasifikasi teks [4].

Berdasarkan permasalahan tersebut, penelitian ini menerapkan algoritma Support Vector Machine (SVM) dengan bantuan metode Term Frequency–Inverse Document Frequency (TF-IDF) untuk mengklasifikasikan berita hoaks dan fakta. TF-IDF berfungsi untuk mengkonversi teks menjadi representation numerik yang didasarkan pada tingkat kepentingan kata, sedangkan SVM dipilih karena kemampuannya dalam mengelola data dengan dimensi tinggi serta menciptakan model yang konsisten.

Penelitian ini bertujuan untuk mengimplementasikan dan mengevaluasi kinerja metode SVM berbasis TF-IDF dalam klasifikasi berita hoaks dan fakta berbahasa Indonesia. Hasil dari penelitian ini diharapkan mampu memberikan kontribusi dalam pengembangan metode klasifikasi teks serta mendukung pembangunan sistem deteksi hoaks secara otomatis di media digital.

## 2. TINJAUAN PUSTAKA

Penelitian mengenai data mining pada data teks tidak terstruktur terus berkembang dalam lima tahun terakhir, khususnya pada permasalahan klasifikasi berita hoaks dan fakta di media digital. Penyebaran informasi yang cepat melalui media sosial dan portal berita meningkatkan risiko disinformasi yang dapat mempengaruhi opini publik. Dalam konteks text mining, prosedur klasifikasi dokumen umumnya mengikuti siklus Knowledge Discovery in Databases (KDD), yaitu seleksi data, preprocessing, transformasi, pemodelan, dan evaluasi [5].

Tahapan ini menjadi krusial karena kualitas preprocessing akan sangat mempengaruhi performa model klasifikasi.

Representasi teks merupakan tahap penting dalam klasifikasi berita. Salah satu cara yang paling sering dipakai adalah Term Frequency–Inverse Document Frequency (TF-IDF). Cara ini memberi nilai pada kata berdasarkan seberapa sering ia muncul di sebuah dokumen. Nilai tersebut juga memperhitungkan seberapa jarang kata itu muncul di seluruh kumpulan dokumen [2].

Pendekatan ini mampu menonjolkan kata yang bersifat diskriminatif antar kelas. Ini membuat kualitas fitur lebih baik untuk proses klasifikasi. Riset membuktikan bahwa pembobotan TF-IDF memberi hasil yang lebih stabil dibanding representasi frekuensi kata secara sederhana pada teks bahasa Indonesia [6].

Beberapa penelitian di Indonesia telah menerapkan kombinasi TF-IDF dan SVM pada deteksi hoaks. Penelitian oleh Indra membandingkan algoritma SVM, K-Nearest Neighbor, dan Random Forest untuk mendeteksi hoaks pada konteks Pemilu Indonesia 2024 dan membuktikan bahwa SVM dengan fitur TF-IDF memperoleh performa akurasi yang kompetitif [5].

Penelitian oleh Desriansyah juga menunjukkan bahwa kombinasi preprocessing yang optimal dengan TF-IDF dapat meningkatkan performa klasifikasi berita digital berbahasa Indonesia [6].

Penelitian oleh Mustofa menemukan bahwa klasifikasi berbasis machine learning dapat digunakan untuk mengidentifikasi berita hoaks melalui tahapan text preprocessing, pembobotan kata, dan klasifikasi dokumen [11].

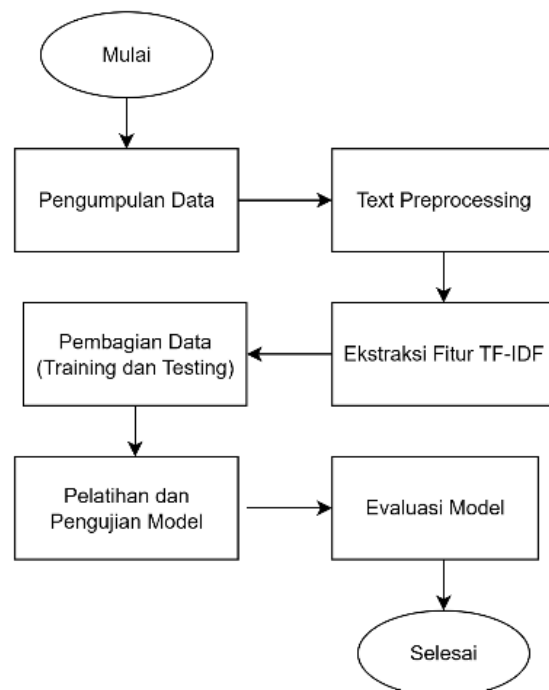
Penelitian lain oleh Ula juga menyatakan bahwa pendekatan text mining mampu membantu proses deteksi berita hoaks pada media digital berbahasa Indonesia [12].

Berdasarkan penelitian lima tahun terakhir tersebut, dapat diidentifikasi bahwa gabungan representasi TF-IDF dan Support Vector Machine adalah pendekatan yang relevan dan efektif dalam klasifikasi berita hoaks dan fakta, khususnya pada korpus berbahasa Indonesia. Oleh sebab itu, penelitian ini memfokuskan pada penggunaan TF-IDF sebagai metode ekstraksi fitur dan SVM sebagai algoritma klasifikasi dalam kerangka data mining guna

menghasilkan model yang akurat dan stabil dalam membedakan berita hoaks dan fakta.

## 3. METODE PENELITIAN

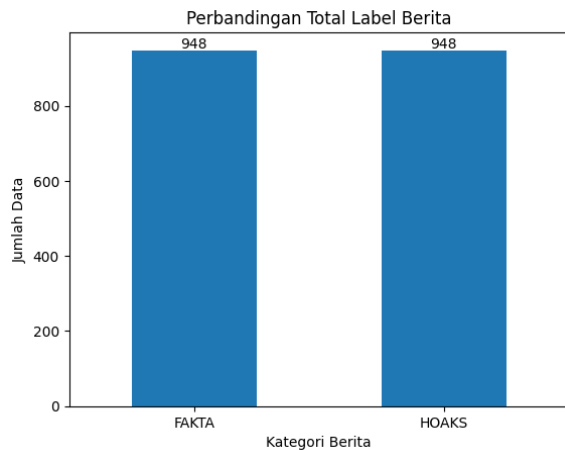
Penelitian klasifikasi teks merupakan salah satu penerapan data mining yang bertujuan mengidentifikasi pola informasi dari data teks secara terstruktur. Pada penelitian ini dilakukan klasifikasi berita hoaks dan fakta dengan menerapkan pendekatan text mining dan algoritma Support Vector Machine (SVM), serta metode pembobotan Term Frequency–Inverse Document Frequency (TF-IDF) yang efektif dalam merepresentasikan fitur teks serta menghasilkan performa klasifikasi yang baik. Proses penelitian mengikuti langkah-langkah text mining yang mencakup pengumpulan data, pra-pemrosesan teks, ekstraksi fitur, pelatihan model, dan penilaian kinerja [8]. Seluruh tahapan penelitian berbasis supervised learning karena dataset telah memiliki label hoaks dan fakta, alur penelitian ini bisa dilihat sebagaimana ditunjukkan pada Gambar 1.



Gambar 1. Alur penelitian

### 3.1. Pengumpulan Data

Data didapat melalui dataset berita dari platform Kaggle yang menyediakan sekumpulan data berita dengan kategori hoaks dan fakta. Dataset yang dipakai dalam penelitian ini berjumlah 1.896 data berita, yang terdiri atas 948 berita hoaks dan 948 berita fakta. Data tersebut termasuk dalam pendekatan supervised learning, karena setiap dokumen telah dilengkapi dengan label kelas yang digunakan sebagai referensi dalam proses *training* serta *testing* model klasifikasi. Distribusi label pada dataset dapat dilihat pada Gambar 2.



Gambar 2. Distribusi Label Berita dalam Dataset

### 3.2. Text Preprocessing

Tahap text preprocessing memiliki tujuan untuk membersihkan dan mempersiapkan data teks sebelum fitur diekstraksi. Proses ini penting untuk meminimalkan noise serta meningkatkan mutu data, sehingga berpotensi mempengaruhi performa model klasifikasi. Tahapan preprocessing yang dilakukan meliputi:

- Cleansing** : Menghilangkan karakter yang tidak perlu seperti tanda baca, angka, URL, simbol, emoji, dan karakter khusus lainnya.
- Case Folding** : Mengkonversi semua huruf dalam teks menjadi huruf kecil untuk mencegah perbedaan dalam penulisan kata akibat penggunaan huruf kapital.
- Normalisasi** : Menyeragamkan bentuk kata tidak baku menjadi kata baku sesuai dengan kaidah bahasa Indonesia, termasuk memperbaiki singkatan, kosakata informal, dan variasi penulisan kata.
- Tokenize** : Memecah teks menjadi bagian-bagian kata (token) sehingga kata dapat dianalisis secara terpisah pada pemrosesan teks.
- Stopword Removal** : Mengeliminasi kata yang sering muncul tetapi tidak memberikan dampak signifikan pada proses klasifikasi, seperti kata penghubung atau kata yang tidak memberikan makna penting dalam dokumen.
- Stemming** : Kata yang memiliki imbuhan diubah menjadi bentuk dasar, sehingga variasi kata dengan makna yang serupa dapat direpresentasikan dengan lebih sederhana.

Proses preprocessing tersebut dilakukan dengan menggunakan bahasa pemrograman Python berkolaborasi Natural Language Toolkit (NLTK) serta Sastrawi Stemmer untuk proses stemming bahasa Indonesia. Tahapan preprocessing diterapkan secara berurutan sebelum dilakukan ekstraksi fitur.

### 3.3. Ekstraksi Fitur TF-IDF

Melakukan tahapan ekstraksi fitur dengan memanfaatkan metode Term Frequency-Inverse Document Frequency (TF-IDF) agar bisa mengonversi

data teks ke suatu bentuk nilai numerik. Metode ini memberikan bobot pada setiap kata dengan mempertimbangkan frekuensi kemunculannya dalam suatu dokumen, serta seberapa unik kemunculannya dibandingkan dengan dokumen lain dalam keseluruhan kumpulan data. Kata yang muncul secara dominan pada dokumen, namun jarang terdapat pada dokumen lainnya akan memperoleh bobot yang lebih tinggi.

Tahapan ini menghasilkan matriks fitur numerik yang menjadi masukan bagi algoritma pengelompokan SVM. Representasi TF-IDF dihitung menggunakan `TfidfVectorizer` dari pustaka `Scikit-learn`. Parameter yang digunakan meliputi pembatasan jumlah fitur maksimum serta penghilangan kata dengan frekuensi sangat rendah. Proses ini menghasilkan matriks dokumen-kata berdimensi tinggi sebagai representasi numerik dalam proses pelatihan model SVM.

### 3.4. Pembagian Data

Setelah seluruh tahapan preprocessing diselesaikan, dataset selanjutnya dipisahkan menjadi dua bagian, yaitu data pelatihan dan data pengujian. Pemisahan data dilakukan dengan perbandingan 80% ditujukan pada data pelatihan dan 20% ditujukan pada data pengujian. Data pelatihan digunakan dalam proses pembelajaran Support Vector Machine (SVM) agar model dapat memahami pola serta karakteristik dari teks berita yang dianalisis. Di sisi lain, data pengujian digunakan untuk menilai kemampuan model dalam mengklasifikasikan data yang belum pernah dipakai selama pelatihan, sehingga kinerja model dapat diukur dengan lebih objektif.

### 3.5. Pelatihan dan Pengujian Model

Pada tahap ini, model dibuat dengan memanfaatkan algoritma Support Vector Machine (SVM) dengan kernel linear (`LinearSVC`). Pemilihan kernel linear dilakukan berdasarkan pada karakteristik data teks berdimensi tinggi hasil representasi TF-IDF yang umumnya lebih efektif dipisahkan menggunakan hyperplane linear. Model dilatih menggunakan data training untuk mempelajari pola kata yang membedakan berita hoaks dan fakta, dengan parameter default dari pustaka `Scikit-learn`.

Setelah proses pelatihan, model dievaluasi menggunakan data testing melalui proses prediksi terhadap label berita. Label prediksi yang diperoleh kemudian dibandingkan dengan label sebenarnya untuk menilai kinerja model dalam mengklasifikasikan berita hoaks dan fakta secara akurat.

### 3.6. Evaluasi Model

Proses penilaian kinerja model dilakukan dengan menggunakan Confusion Matrix sebagai dasar penilaian klasifikasi. Dari matrix tersebut kemudian dihitung beberapa metrik evaluasi, yaitu *accuracy*, *precision*, *recall*, dan *F1-score*. *Accuracy*

dimanfaatkan untuk menilai seberapa tepat prediksi model tersebut terhadap data yang diuji. *Precision* menunjukkan sejauh mana model dapat mengklasifikasikan data ke kategori yang tepat. *Recall* menggambarkan sejauh mana model dapat mengidentifikasi semua data dalam kategori yang sama. *F1-score* memberikan ukuran keseimbangan antara *precision* dan *recall*. Nilai-nilai evaluasi tersebut dipergunakan sebagai dasar untuk menilai seberapa baik performa model Support Vector Machine dalam mengklasifikasikan berita hoaks dan berita fakta.

#### 4. HASIL DAN PEMBAHASAN

Hasil penelitian menunjukkan penggunaan text mining untuk mengelompokkan berita hoaks dan fakta menggunakan algoritma Support Vector Machine (SVM) dengan metode ekstraksi fitur Term

Frequency-Inverse Document Frequency (TF-IDF). Data yang digunakan berupa dataset berita berbahasa Indonesia yang telah dikategorikan sebagai berita fakta dan hoaks, yang kemudian diproses dengan Python.

##### 4.1. Pengumpulan Data

Berdasarkan proses pengumpulan data, diperoleh dataset berita sebanyak 1.896 data yang terdiri dari 948 berita fakta dan 948 berita hoaks. Distribusi data yang seimbang antar kelas menunjukkan tidak adanya permasalahan class imbalance, sehingga model klasifikasi dapat dilatih tanpa bias terhadap salah satu kategori berita. Rincian distribusi jumlah data pada masing-masing kategori ditampilkan pada Tabel 1. Dataset tersebut kemudian digunakan pada tahap preprocessing teks dan ekstraksi fitur sebagai input dalam pelatihan model Support Vector Machine.

Tabel 1. Dataset awal hasil web crawling

| No | Judul                                                                  | Content                                                                                                                                                                                                                                              | Label |
|----|------------------------------------------------------------------------|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|-------|
| 1. | Serang Gereja saat Malam Natal                                         | Orang Palestina membenci Natal. Mereka menyerang sebuah gereja saat latihan Natal. Warga Palestina di Betlehem melepaskan tembakan, melemparkan batu, dan mencoba masuk ke dalam gereja pada malam Natal.                                            | Hoaks |
| 2. | Kapolri Sangat Menyesalkan dan Minta Maaf Tim Pengawalan Pukul Jurnali | Kapolri Jenderal Listyo Sigit Prabowo meminta maaf atas terjadinya insiden HOAKS seorang anggota tim pengawalan melakukan kekerasan fisik terhadap jurnalis. Jenderal Sigit menyesalkan insiden tersebut dan berjanji mengusut tuntas peristiwa itu. | Fakta |

##### 4.2. Data Pre-Processing

Dalam *Pre-Processing*, terdapat beberapa tahapan yang perlu dilakukan, yaitu ;

##### 4.3. Cleansing

Sebelum dilakukan *cleansing*, kata masih terdapat tanda titik dan pemisah kalimat, seperti pada kalimat "Orang Palestina membenci Natal." Setelah proses cleansing, tanda baca dihilangkan sehingga teks menjadi rangkaian kalimat bersih tanpa simbol tambahan.

Tabel 2. Cleansing

| Sebelum Cleaning                                                                                                                                                                                          | Sesudah Cleaning                                                                                                                                                                                                                                     |
|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| Orang Palestina membenci Natal. Mereka menyerang sebuah gereja saat latihan Natal. Warga Palestina di Betlehem melepaskan tembakan, melemparkan batu, dan mencoba masuk ke dalam gereja pada malam Natal. | Dia cuma butuh di dengar dan dia pun sebenarnya Orang Palestina membenci Natal Mereka menyerang sebuah gereja saat latihan Natal Warga Palestina di Betlehem melepaskan tembakan melemparkan batu dan mencoba masuk ke dalam gereja pada malam Natal |

##### 4.4. Case Folding

Tahapan ini dilakukan proses *case folding*, yaitu proses normalisasi teks dengan mengubah huruf yang ada menjadi huruf kecil (*lowercase*).

Tabel 3. Case Folding

| Sebelum Case Folding                                                                                                                                                                                 | Sesudah Case Folding                                                                                                                                                                                 |
|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| Orang Palestina membenci Natal Mereka menyerang sebuah gereja saat latihan Natal Warga Palestina di Betlehem melepaskan tembakan melemparkan batu dan mencoba masuk ke dalam gereja pada malam Natal | orang palestina membenci natal mereka menyerang sebuah gereja saat latihan natal warga palestina di betlehem melepaskan tembakan melemparkan batu dan mencoba masuk ke dalam gereja pada malam natal |

##### 4.5. Normalisasi

Normalisasi dilakukan menyeragamkan bentuk kata tidak baku menjadi kata baku sesuai dengan kaidah bahasa Indonesia.

Tabel 4. Normalisasi

| Sebelum normalisasi                                                                                                                                                                                  | Sesudah normalisasi                                                                                                                                                                                  |
|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| orang palestina membenci natal mereka menyerang sebuah gereja saat latihan natal warga palestina di betlehem melepaskan tembakan melemparkan batu dan mencoba masuk ke dalam gereja pada malam natal | orang palestina membenci natal mereka menyerang sebuah gereja saat latihan natal warga palestina di betlehem melepaskan tembakan melemparkan batu dan mencoba masuk ke dalam gereja pada malam natal |

#### 4.6. Tokenize

*Tokenize* melakukan pemecahan teks menjadi unit-unit kata tunggal, seperti “orang”, “palestina”, “membenci”, “natal”, dan seterusnya.

Tabel 5. *Tokenize*

| Sebelum tokenize                                                                                                                                                                                                                | Sesudah tokenize                                                                                                                                                                                                                                                                                                        |
|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| orang palestina<br>membenci natal mereka<br>menyerang sebuah<br>gereja saat latihan natal<br>warga palestina di<br>betlehem melepaskan<br>tembakan melemparkan<br>batu dan mencoba<br>masuk ke dalam gereja<br>pada malam natal | ['orang', 'palestina',<br>'membenci', 'natal', 'mereka',<br>'menyerang', 'sebuah',<br>'gereja', 'saat', 'latihan',<br>'natal', 'warga', 'palestina',<br>'di', 'betlehem', 'melepaskan',<br>'tembakan', 'melemparkan',<br>'batu', 'dan', 'mencoba',<br>'masuk', 'ke', 'dalam',<br>'gereja', 'pada', 'malam',<br>'natal'] |

#### 4.7. Stopword Removal

Pada tahap ini, kata-kata seperti “dan”, “di”, “ke”, dan “pada” dihapus karena tidak memberikan kontribusi dalam membedakan kategori berita. Setelah tahap ini, tersisa kata-kata yang lebih substantif dan relevan terhadap isi informasi.

Tabel 6. *Stopword Removal*

| Sebelum stopwords removal                                                                                                                                                                                                                                                                                               | Sesudah stopwords removal                                                                                                                                                                                                                         |
|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| ['orang', 'palestina',<br>'membenci', 'natal', 'mereka',<br>'menyerang', 'sebuah',<br>'gereja', 'saat', 'latihan',<br>'natal', 'warga', 'palestina',<br>'di', 'betlehem',<br>'melepaskan', 'tembakan',<br>'melemparkan', 'batu', 'dan',<br>'mencoba', 'masuk', 'ke',<br>'dalam', 'gereja', 'pada',<br>'malam', 'natal'] | ['orang', 'palestina',<br>'membenci', 'natal',<br>'menyerang', 'gereja',<br>'latihan', 'natal', 'warga',<br>'palestina', 'betlehem',<br>'melepaskan', 'tembakan',<br>'melemparkan', 'batu',<br>'mencoba', 'masuk',<br>'gereja', 'malam', 'natal'] |

#### 4.8. Stemming

*Stemming* dibuat dengan mengubah kata yang memiliki imbuhan menjadi bentuk dasar, seperti “membenci” menjadi “benci”, “menyerang” menjadi “serang”, dan “melemparkan” menjadi “lempar”. Proses ini menyederhanakan variasi morfologis kata sehingga representasi fitur menjadi lebih ringkas dan mampu meningkatkan efektivitas model Support Vector Machine dalam membedakan berita hoaks dan fakta.

Tabel 7. *Stemming*

| Sebelum Stemming                                                                                                                                                                                                                               | Sesudah Stemming                                                                                                                                        |
|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|---------------------------------------------------------------------------------------------------------------------------------------------------------|
| ['orang', 'palestina', 'membenci',<br>'natal', 'menyerang', 'gereja',<br>'latihan', 'natal', 'warga',<br>'palestina', 'betlehem',<br>'melepaskan', 'tembakan',<br>'melemparkan', 'batu',<br>'mencoba', 'masuk', 'gereja',<br>'malam', 'natal'] | orang palestina benci<br>natal serang gereja<br>latih natal warga<br>palestina betlehem<br>lepas tembak lempar<br>batu coba masuk<br>gereja malam natal |

#### 4.9. Ekstraksi Fitur TF-IDF

Proses ini dilakukan dengan metode Term Frequency–Inverse Document Frequency (TF-IDF) agar mengubah data teks menjadi bentuk numerik. Berdasarkan hasil transformasi menggunakan *TfidfVectorizer*, diperoleh matriks fitur berdimensi  $1327 \times 15016$ , yang menunjukkan bahwa sebanyak 1327 dokumen berita berhasil direpresentasikan ke dalam 15.016 fitur kata unik.

Jumlah fitur yang relatif besar menunjukkan keberagaman kosakata pada dataset berita hoaks dan fakta. Representasi ini membuat kemungkinan model memberikan bobot lebih tinggi pada kata yang bersifat spesifik terhadap suatu dokumen, sehingga membantu algoritma Support Vector Machine dalam membedakan karakteristik linguistik antar kategori berita.

#### 4.10. Pembagian Data

Setelah tahap preprocessing dan ekstraksi fitur TF-IDF, dataset dibagi menjadi dua subset, yaitu data *training* dan data *testing* dengan rasio 80:20. Pembagian data dilakukan agar model bisa mempelajari pola dari data *training* dan dievaluasi pada data yang sebelumnya belum digunakan.

Dari total dataset, sebanyak 1.327 data berperan menjadi data *training*, sedangkan 569 data berperan menjadi data *testing*. Dua bagian ini bertujuan untuk menghindari overfitting serta mengukur kemampuan model Support Vector Machine dalam melakukan klasifikasi berita hoaks dan fakta.

#### 4.11. Pelatihan dan Pengujian Model

Setelah tahap preprocessing dan ekstraksi fitur selesai, model Support Vector Machine (SVM) kemudian dilatih menggunakan data *training* dan diuji dengan data *testing*. Cara kerja model akan diuji menggunakan metrik, yaitu *accuracy*, *precision*, *recall*, dan *F1-score*, serta ditampilkan dalam bentuk *classification report*.

Hasilnya menunjukkan bahwa model mencapai tingkat akurasi 97% pada data *testing*, yang menunjukkan kemampuan model dalam membedakan berita fakta dan hoaks secara baik. Evaluasi performa pada masing-masing kelas (FAKTA dan HOAKS) ditampilkan pada *classification report* di Gambar 3.

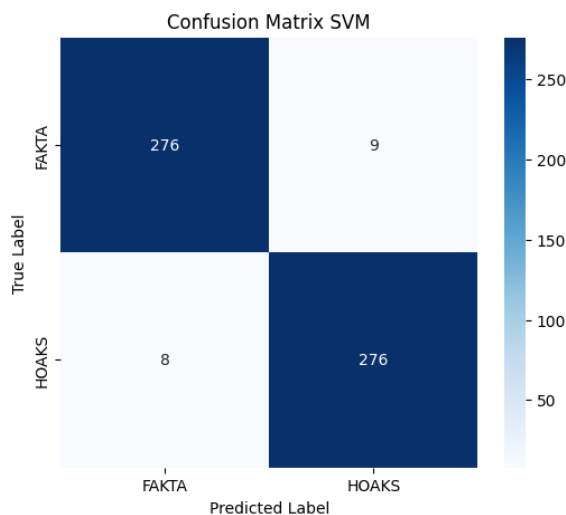
|              | precision | recall | f1-score | support |
|--------------|-----------|--------|----------|---------|
| FAKTA        | 0.97      | 0.97   | 0.97     | 285     |
| HOAKS        | 0.97      | 0.97   | 0.97     | 284     |
| accuracy     |           |        | 0.97     | 569     |
| macro avg    | 0.97      | 0.97   | 0.97     | 569     |
| weighted avg | 0.97      | 0.97   | 0.97     | 569     |

Gambar 3. Hasil program klasifikasi SVM

Berdasarkan hasilnya, nilai *precision*, *recall*, dan *F1-score* untuk kedua kelas memiliki nilai identik dan seimbang, yaitu 0,97. Hal ini menunjukkan bahwa model memiliki kemampuan generalisasi yang kuat

dan tidak menunjukkan bias dalam mengenali pola berita fakta maupun berita hoaks.

Confusion matrix juga dihasilkan untuk memvisualisasikan ada berapa banyak dari tiap *True Positive* (TP), *True Negative* (TN), *False Positive* (FP), dan *False Negative* (FN). Secara keseluruhan, hasil ini membuktikan kalau gabungan fitur TF-IDF dan algoritma SVM sangat relevan untuk mendeteksi hoaks pada konteks berita politik. Visualisasi Confusion Matrix ada pada Gambar 4.



Gambar 4. Confusion Matrix Model LinearSVC

Hasil klasifikasi model LinearSVC dapat diinterpretasikan sebagai berikut:

- True Positive* (TP): Sebanyak 276 data berita hoaks berhasil diidentifikasi dengan benar sebagai hoaks oleh model.
- True Negative* (TN): Sebanyak 276 data berita fakta berhasil diidentifikasi dengan benar sebagai fakta oleh model.
- False Positive* (FP): Terdapat 9 data berita fakta yang salah dikelompokkan oleh model sebagai berita hoaks.
- False Negative* (FN): Terdapat 8 data berita hoaks yang salah dikelompokkan oleh model sebagai berita fakta.

#### 4.12. Evaluasi Model

Tahap evaluasi model dilakukan untuk mengukur efektivitas model dalam membedakan berita FAKTA dan HOAKS. Pada gambar di atas menampilkan hasil evaluasi dari model Support Vector Machine (SVM) dengan menggunakan tiga metrik utama, yaitu *precision*, *recall*, dan *F1-score* pada setiap kelas. Adapun metrik evaluasi yang dijadikan tolak ukur di penelitian ini meliputi:

##### a. Precision (Presisi)

*Precision* dimanfaatkan untuk mengukur seberapa tepat model ini dalam melakukan pengelompokan suatu kelas. Nilai *precision* yang tinggi menunjukkan bahwa ketika model mengklasifikasikan berita sebagai FAKTA atau

HOAKS, sebagian besar hasil prediksi tersebut benar. Pada gambar terlihat bahwa nilai *precision* untuk kelas FAKTA dan HOAKS berada di 0,97 yang menandakan bahwa model sangat jarang menghasilkan kesalahan prediksi positif (*false positive*).

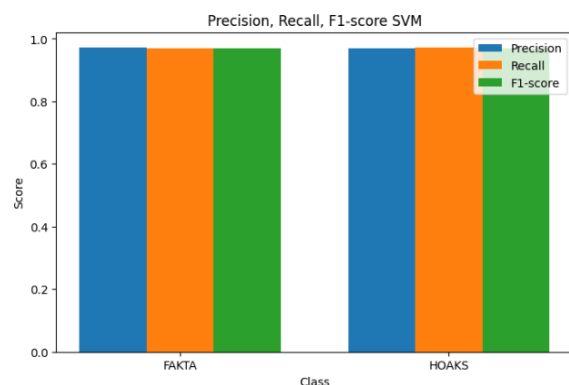
##### b. Recall

*Recall* menunjukkan model mampu untuk mengidentifikasi semua data yang memang termasuk ke dalam suatu kelas. Semakin besar nilai *recall*, maka semakin banyak berita yang benar-benar FAKTA atau HOAKS yang berhasil diidentifikasi oleh model. Berdasarkan grafik, nilai *recall* untuk kedua kelas juga berada di 0,97. Model menunjukkan bahwa mampu mendeteksi hampir seluruh data pada masing-masing kategori dengan sangat baik dan meminimalkan kesalahan *false negative*.

##### c. F1-score

*F1-score* adalah ukuran yang didapat dari kombinasi antara *precision* dan *recall* melalui rata-rata harmonis. Metrik ini digunakan untuk memberikan keseimbangan antara kedua nilai tersebut, contohnya ketika keduanya sama-sama penting dalam sistem klasifikasi. Pada grafik terlihat bahwa *F1-score* untuk kelas FAKTA dan HOAKS juga mendekati 0,98. Hal ini menunjukkan bahwa model mempertahankan keseimbangan antara ketepatan prediksi dan kemampuan mendeteksi data pada masing-masing kelas.

Model SVM memberikan performa yang sangat memuaskan untuk membedakan berita FAKTA dan HOAKS. Hal ini dapat dilihat dari nilai *precision*, *recall*, dan *F1-score* yang hampir mendekati 1,0, yang menandakan bahwa tingkat kesalahan klasifikasi sangat kecil. Selain itu, model juga menunjukkan kemampuan generalisasi yang baik terhadap data *testing*. Visualisasi hasil evaluasi terdapat pada Gambar 5.



Gambar 5. Precision, Recall, F1-Score Per Class

## 5. KESIMPULAN DAN SARAN

Penelitian ini menunjukkan bahwa penggunaan algoritma Support Vector Machine (SVM) dengan representasi fitur Term Frequency–Inverse Document Frequency (TF-IDF) efektif dalam melakukan

klasifikasi berita hoaks dan fakta berbahasa Indonesia. Berdasarkan hasil pengujian terhadap dataset yang terdiri dari 1.896 berita, model ini mencatat akurasi hingga 97%, dengan *precision*, *recall*, dan *F1-score* sebesar 0,97 pada kedua kelas. Hasil itu menunjukkan bahwa kombinasi TF-IDF dan SVM mampu membedakan berita hoaks dan fakta dengan tingkat ketepatan yang tinggi. Proses text *preprocessing* yang mencakup tahap *cleansing*, *case folding*, normalisasi, *tokenisasi*, *stopword removal*, dan *stemming* penting untuk menghasilkan data yang berkualitas sehingga model dapat mengenali pola kata dengan lebih efektif. Untuk berikutnya disarankan menggunakan dataset dalam jumlah yang besar dan lebih memiliki variasi supaya model mampu mempelajari lebih banyak variasi bahasa. Penelitian berikutnya juga dapat mencoba metode atau algoritma lain sebagai pembanding guna meningkatkan kinerja klasifikasi. Diharapkan hasil ini bisa dijadikan acuan untuk pengembangan metode klasifikasi teks serta mendukung pengembangan sistem deteksi hoaks otomatis di media digital.

#### DAFTAR PUSTAKA

- [1] M. Berrondo-otermine and A. Sarasa-cabezuelo, "Application of Artificial Intelligence Techniques to Detect Fake News: A Review," *Electronics*, vol. 12, no. 24, pp. 1–12, 2023, doi: <https://doi.org/10.3390/electronics12245041>.
- [2] A. K. Dewi, N. F. Rahmadani, R. Syahputri, L. R. Nasution, and M. Furqan, "Deteksi Berita Hoax pada Platform X Menggunakan Pendekatan Text Mining dan Algoritma Machine Learning," *Data Sci. Indones.*, vol. 5, no. 1, pp. 33–46, 2025, doi: <https://doi.org/10.47709/dsi.v5i1.6011>.
- [3] R. R. Sani, Y. A. Pratiwi, S. Winarno, E. D. Udayanti, and F. Al Zami, "Analisis Perbandingan Algoritma Naive Bayes Classifier dan Support Vector Machine untuk Klasifikasi Hoax pada Berita Online Indonesia," *J. Masy. Inform.*, vol. 13, no. 2, pp. 85–98, 2022, doi: <https://doi.org/10.14710/jmasif.13.2.47983>.
- [4] A. H. Tandiano and D. Jollyta, "Classification of fake news in Indonesian language using support vector machine method," *J. Teknol. dan Sist. Inf.*, vol. 10, no. 2, pp. 315–322, 2024, doi: <http://dx.doi.org/10.33330/jurteks.v10i2.2895>.
- [5] I. Indra, A. U. Hamdani, S. Setiawati, Z. D. Mentari, and M. H. Purnomo, "Comparison of k-  
nn , svm , and random forest algorithm for detecting hoax on Indonesian election 2024," *J. Nas. Pendidik. Tek. Inform.*, vol. 13, no. 1, pp. 166–179, 2024, doi: <https://doi.org/10.23887/janapati.v13i1.76079>.
- [6] M. D. Desriansyah, I. U. Sari, and Z. Zulfahmi, "Analisis Efektivitas Algoritma Machine Learning dalam Deteksi Hoaks: Pada Berita Digital Berbahasa Indonesia," *J. Sist. Inf. Dan Inform.*, vol. 3, no. 2, pp. 63–69, 2025, doi: <https://doi.org/10.47233/jiska.v3i1.2024>.
- [7] M. F. Ramadhan, "Klasifikasi Topik dan Sentimen Judul Berita dengan Augmentasi dan TF-IDF," *RIGGS*, vol. 4, no. 2, pp. 6732–6741, 2025, doi: <https://doi.org/10.31004/riggs.v4i2.1692>.
- [8] K. Kowsari, K. J. Meimandi, M. Heidarysafa, S. Mendu, L. Barnes, and D. Brown, "Text Classification Algorithms: A Survey," *Information*, vol. 10, no. 4, pp. 1–68, 2019, doi: [10.3390/info10040150](https://doi.org/10.3390/info10040150).
- [9] T. Wahyudi and N. Azmi, "Klasifikasi Sentimen X-Twitter Perihal Pemindahan Ibu Kota Indonesia Menggunakan Ekstraksi Fitur TF-IDF dan Metode Support Vector Machine (SVM)," *Jurnal Teknologi Informasi: Jurnal Keilmuan dan Aplikasi Bidang Teknik Informatika*, vol. 18, no. 2, pp. 185–199, 2024, doi: [10.47111/jti.v18i2.15015](https://doi.org/10.47111/jti.v18i2.15015).
- [10] A. H. Tandiano and D. Jollyta, "Classification of Fake News in Indonesian Language Using Support Vector Machine Method," *Jurnal Teknologi dan Sistem Informasi*, vol. 10, no. 2, pp. 315–322, 2024, doi: <http://dx.doi.org/10.33330/jurteks.v10i2.2895>.
- [11] R. R. Sani, Y. A. Pratiwi, S. Winarno, E. D. Udayanti, and F. Alzami, "Analisis Perbandingan Algoritma Naive Bayes Classifier dan Support Vector Machine untuk Klasifikasi Berita Hoax pada Berita Online Indonesia," *Jurnal Masyarakat Informatika*, vol. 13, no. 2, pp. 85–98, 2022, doi: <https://doi.org/10.14710/jmasif.13.2.47983>.
- [12] M. Ula, M. Mahendra, Alvanof, and R. Trisandi, "Analisa dan Deteksi Konten Hoax pada Media Berita Indonesia Menggunakan Machine Learning," *Jurnal Teknologi Terapan & Sains*, vol. 1, no. 2, 2020, doi: [10.29103/tts.v1i2.32](https://doi.org/10.29103/tts.v1i2.32).