

Komparasi Model Klasifikasi Teks dalam Mendeteksi Berita Hoaks Berbahasa Indonesia

Duwi Purnama Sidik*¹

¹Universitas Brawijaya, Indonesia

e-mail: duwi763@gmail.com

Received: 23-06-2026

Accepted: 30-06-2026

Published: 30-06-2026

Abstrak

Penyebaran berita hoaks di era digital menjadi tantangan yang dapat memengaruhi opini publik dan menimbulkan disinformasi. Penelitian ini bertujuan membandingkan kinerja algoritma Naïve Bayes dan Support Vector Machine (SVM) dalam mendeteksi berita hoaks berbahasa Indonesia. Metode penelitian menggunakan pendekatan eksperimen yang meliputi prapemrosesan teks, ekstraksi fitur menggunakan TF-IDF, pelatihan model, dan evaluasi kinerja. Dataset yang digunakan terdiri atas 3.772 berita yang diklasifikasikan ke dalam kategori hoaks dan valid. Evaluasi dilakukan menggunakan metrik accuracy, precision, recall, dan F1-score. Hasil penelitian menunjukkan bahwa SVM memperoleh performa lebih baik dengan accuracy 70,87%, precision 72%, dan F1-score 82%, sedangkan Naïve Bayes memperoleh accuracy 66,52%, precision 67%, dan F1-score 80%. Namun, Naïve Bayes memiliki recall lebih tinggi sebesar 99% dibandingkan SVM sebesar 97%. Hasil ini menunjukkan bahwa SVM lebih efektif dalam klasifikasi berita hoaks berbahasa Indonesia.

Kata Kunci: Berita Hoaks; Klasifikasi Teks; Machine Learning; Naïve Bayes; Support Vector Machine

How to Cite:

Sidik, D. P. (2026). Komparasi Model Klasifikasi Teks dalam Mendeteksi Berita Hoaks Berbahasa Indonesia. *JINTEN: Journal of Intelligent Systems and Digital Science*, 1(1), 1-12.

Copyright ©2026 to the Author. Published by CV. Ihsan Cahaya Pustaka

This is an open access article under the [CC-BY-SA](https://creativecommons.org/licenses/by-sa/4.0/) license



1. PENDAHULUAN

Perkembangan teknologi informasi dan komunikasi telah mengubah cara masyarakat memperoleh, mengakses, dan menyebarluaskan informasi melalui berbagai platform digital. Kemudahan distribusi informasi melalui media sosial, portal berita daring, dan aplikasi pesan instan memungkinkan masyarakat memperoleh informasi secara cepat tanpa batas ruang dan waktu. Namun, kondisi tersebut juga memunculkan tantangan baru berupa meningkatnya penyebaran berita hoaks yang dapat menjangkau masyarakat dalam waktu singkat tanpa melalui proses verifikasi yang memadai. Mustafa dan Alfianti (2025) menjelaskan bahwa tingginya arus informasi digital menyebabkan masyarakat semakin rentan menerima informasi yang tidak valid. Fenomena serupa juga diungkapkan oleh Kesuma (2026) yang menyatakan bahwa berita hoaks masih menjadi permasalahan serius dalam ekosistem informasi digital karena mampu memengaruhi persepsi dan perilaku masyarakat secara luas.

Peningkatan penyebaran berita hoaks semakin terlihat pada isu-isu yang memiliki sensitivitas tinggi, seperti kesehatan, politik, ekonomi, dan kebijakan publik. Informasi yang tidak akurat sering kali dikemas dengan judul yang menarik sehingga mudah dipercaya dan disebarluaskan kembali oleh pengguna internet. Menurut Munir dan Indra (2024), penyebaran hoaks pada media sosial meningkat secara signifikan ketika masyarakat dihadapkan pada isu yang memicu perhatian publik. Imran et

al. (2024) juga menemukan bahwa informasi palsu terkait isu politik dan pemilu memiliki potensi besar dalam membentuk opini publik secara tidak objektif. Kondisi tersebut menyebabkan proses verifikasi fakta secara manual menjadi semakin sulit dilakukan karena tingginya volume informasi yang beredar setiap hari.

Sebagai respons terhadap permasalahan tersebut, berbagai pendekatan berbasis kecerdasan buatan mulai dikembangkan untuk membantu proses identifikasi berita hoaks secara otomatis. Salah satu pendekatan yang banyak digunakan adalah Machine Learning karena mampu mempelajari pola dari data historis dan menghasilkan model klasifikasi yang dapat digunakan untuk memprediksi kategori suatu informasi. Desriansyah et al. (2025) menjelaskan bahwa algoritma Machine Learning memiliki kemampuan yang baik dalam mengenali karakteristik berita palsu berdasarkan pola linguistik yang terdapat dalam teks. Sejalan dengan hal tersebut, Hafizh dan Juanita (2026) menunjukkan bahwa penerapan algoritma klasifikasi pada berita berbahasa Indonesia mampu menghasilkan performa yang cukup baik dalam membedakan berita hoaks dan berita valid. Oleh karena itu, Machine Learning menjadi salah satu solusi yang menjanjikan dalam pengembangan sistem deteksi hoaks otomatis.

Implementasi Machine Learning pada klasifikasi teks tidak dapat dipisahkan dari pemanfaatan Natural Language Processing (NLP). NLP berfungsi untuk mengubah data teks yang tidak terstruktur menjadi bentuk yang dapat dipahami oleh algoritma pembelajaran mesin melalui berbagai tahapan pengolahan bahasa alami. Ledjap et al. (2024) menjelaskan bahwa proses NLP meliputi tokenisasi, normalisasi, dan pengolahan struktur bahasa yang bertujuan meningkatkan kualitas data teks. Selain itu, Ramadhan (2025) menyatakan bahwa kualitas hasil klasifikasi sangat dipengaruhi oleh tahapan prapemrosesan karena proses tersebut mampu mengurangi noise dan mempertahankan informasi penting yang terkandung dalam dokumen. Dengan demikian, NLP memiliki peran penting dalam meningkatkan efektivitas sistem klasifikasi berita hoaks berbasis teks.

Setelah melalui proses prapemrosesan, teks perlu direpresentasikan ke dalam bentuk numerik agar dapat diproses oleh algoritma klasifikasi. Salah satu metode representasi fitur yang paling banyak digunakan adalah Term Frequency–Inverse Document Frequency (TF-IDF), yang memberikan bobot pada kata berdasarkan tingkat kemunculannya dalam dokumen dan keseluruhan korpus. Syah et al. (2025) menjelaskan bahwa TF-IDF mampu menghasilkan representasi teks yang efektif karena dapat menonjolkan kata-kata yang memiliki kontribusi tinggi terhadap suatu kategori dokumen. Hasil penelitian Suryanti dan Prasetyaningrum (2025) juga menunjukkan bahwa penggunaan TF-IDF dapat meningkatkan kualitas fitur dan mendukung peningkatan performa algoritma klasifikasi berbasis teks. Oleh karena itu, TF-IDF banyak digunakan dalam penelitian klasifikasi dokumen, termasuk pada deteksi berita hoaks.

Di antara berbagai algoritma klasifikasi yang tersedia, Naïve Bayes dan Support Vector Machine (SVM) merupakan dua metode yang paling sering digunakan dalam penelitian deteksi berita hoaks. Menurut Soleman (2021), Naïve Bayes memiliki keunggulan dalam efisiensi komputasi karena menggunakan pendekatan probabilistik yang sederhana dan mampu bekerja dengan baik pada data teks. Sementara itu, Hasan Dalimunthe et al. (2024) menjelaskan bahwa SVM memiliki kemampuan yang baik dalam menangani data berdimensi tinggi melalui pembentukan hyperplane optimal sebagai batas pemisah antar kelas. Penelitian Pangesti et al. (2023) juga menunjukkan bahwa Naïve Bayes memiliki performa yang cukup kompetitif pada sistem deteksi hoaks, sedangkan Sani et al. (2022) menemukan bahwa SVM cenderung menghasilkan tingkat akurasi yang lebih tinggi pada beberapa kasus klasifikasi berita online.

Berbagai penelitian sebelumnya telah mengimplementasikan maupun membandingkan kedua algoritma tersebut dalam mendeteksi berita hoaks berbahasa Indonesia. Mustafa dan Alfianti (2025) menerapkan Naïve Bayes pada klasifikasi berita palsu berbasis web dan memperoleh hasil yang cukup baik. Di sisi lain, Putri Ramadani et al. (2026) menunjukkan bahwa SVM mampu memberikan performa yang kompetitif dalam klasifikasi berita hoaks. Imran et al. (2024) menemukan adanya perbedaan kinerja antara Naïve Bayes dan SVM pada dataset berita politik, sedangkan Hasan Dalimunthe et al. (2024) menunjukkan bahwa efektivitas kedua algoritma dipengaruhi oleh karakteristik data dan teknik ekstraksi fitur yang digunakan. Temuan-temuan tersebut menunjukkan bahwa belum terdapat kesimpulan yang konsisten mengenai algoritma yang paling efektif dalam mendeteksi berita hoaks berbahasa Indonesia.

Selain pendekatan Machine Learning konvensional, beberapa penelitian terbaru mulai mengembangkan metode berbasis Deep Learning untuk meningkatkan akurasi deteksi hoaks. Ripa'i et al. (2024) membuktikan bahwa model BERT mampu memahami konteks bahasa dengan lebih baik dibandingkan pendekatan tradisional berbasis frekuensi kata. Sementara itu, Andryani et al. (2025) menunjukkan bahwa model Deep Learning pada beberapa kasus mampu menghasilkan performa yang lebih tinggi dibandingkan algoritma Machine Learning konvensional. Meskipun demikian, model seperti Naïve Bayes dan SVM masih banyak digunakan karena memiliki kompleksitas komputasi yang lebih rendah, proses pelatihan yang lebih cepat, serta lebih mudah diimplementasikan pada berbagai lingkungan sistem.

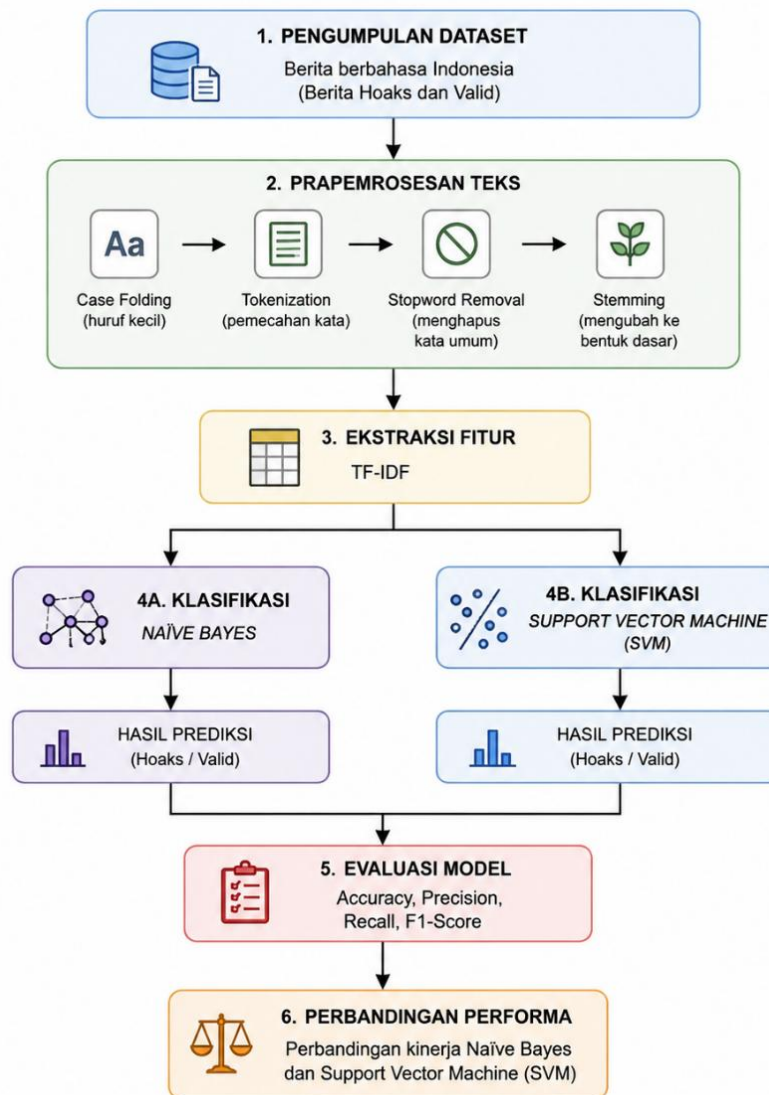
Berdasarkan uraian tersebut, penelitian ini bertujuan untuk melakukan komparasi model klasifikasi teks menggunakan algoritma Naïve Bayes dan Support Vector Machine dalam mendeteksi berita hoaks berbahasa Indonesia. Penelitian dilakukan dengan menerapkan tahapan prapemrosesan teks menggunakan pendekatan NLP serta ekstraksi fitur menggunakan TF-IDF sebelum proses klasifikasi. Pemilihan kedua algoritma didasarkan pada hasil penelitian Sani et al. (2022) dan Imran et al. (2024) yang menunjukkan bahwa Naïve Bayes dan SVM merupakan metode yang paling banyak digunakan dalam klasifikasi berita hoaks. Hasil penelitian diharapkan dapat memberikan gambaran yang lebih komprehensif mengenai performa kedua algoritma sekaligus menjadi referensi bagi pengembangan sistem deteksi hoaks yang lebih efektif dan akurat.

2. METODE

Desain Penelitian

Penelitian ini menggunakan pendekatan kuantitatif dengan metode eksperimen untuk membandingkan kinerja algoritma Naïve Bayes dan Support Vector Machine (SVM) dalam mendeteksi berita hoaks berbahasa Indonesia. Metode eksperimen dipilih karena memungkinkan pengujian dan evaluasi performa beberapa algoritma klasifikasi pada dataset yang sama sehingga hasil yang diperoleh dapat dibandingkan secara objektif. Menurut Sani et al. (2022), pendekatan komparatif pada algoritma klasifikasi dapat digunakan untuk mengidentifikasi model yang memiliki kemampuan terbaik dalam mengolah data teks. Pendekatan serupa juga digunakan oleh Imran et al. (2024) dalam penelitian klasifikasi berita hoaks menggunakan Naïve Bayes dan SVM.

Tahapan penelitian terdiri atas pengumpulan dataset, prapemrosesan teks, ekstraksi fitur menggunakan TF-IDF, pelatihan model klasifikasi, pengujian model, serta evaluasi kinerja berdasarkan beberapa metrik pengukuran. Seluruh tahapan penelitian dirancang secara sistematis untuk memperoleh hasil evaluasi yang objektif dan dapat digunakan sebagai dasar dalam membandingkan performa kedua algoritma klasifikasi.



Gambar 1. Alur Penelitian

Berdasarkan Gambar 1, penelitian diawali dengan proses pengumpulan dataset berita berbahasa Indonesia yang telah diberi label hoaks dan valid. Selanjutnya, data melalui tahap prapemrosesan yang meliputi case folding, tokenization, stopwords removal, dan stemming untuk meningkatkan kualitas data teks (Ledjap et al., 2024; Ramadhan, 2025). Data yang telah diproses kemudian direpresentasikan menggunakan metode TF-IDF sebelum dilakukan pelatihan model menggunakan algoritma Naïve Bayes dan Support Vector Machine. Tahap berikutnya adalah pengujian model dan evaluasi menggunakan metrik accuracy, precision, recall, dan F1-score untuk menentukan algoritma yang memiliki performa terbaik dalam mendeteksi berita hoaks berbahasa Indonesia (Desriansyah et al., 2025; Kesuma, 2026).

Dataset Penelitian

Dataset yang digunakan berupa kumpulan berita berbahasa Indonesia yang telah diberi label hoaks dan valid. Data diperoleh dari sumber dataset publik yang tersedia untuk kepentingan penelitian klasifikasi teks dan deteksi hoaks. Setiap data terdiri atas teks berita dan label kategori yang digunakan sebagai target klasifikasi. Dataset kemudian dibagi menjadi data latih dan data uji menggunakan rasio 80:20. Pembagian data dilakukan secara stratified sampling untuk menjaga

proporsi masing-masing kelas pada data latih dan data uji sehingga distribusi data tetap representatif (Hasan Dalimunthe et al., 2024; Putri Ramadani et al., 2026).

Prapemrosesan Teks

Sebelum proses klasifikasi dilakukan, seluruh data teks melalui tahap prapemrosesan untuk meningkatkan kualitas data dan mengurangi informasi yang tidak relevan. Tahapan prapemrosesan meliputi case folding, tokenization, stopwords removal, dan stemming. Case folding dilakukan dengan mengubah seluruh huruf menjadi huruf kecil untuk menghindari perbedaan representasi akibat penggunaan huruf kapital. Selanjutnya, tokenization digunakan untuk memecah kalimat menjadi unit kata yang lebih mudah diproses oleh sistem. Stopword removal dilakukan untuk menghilangkan kata-kata umum yang tidak memiliki kontribusi signifikan terhadap proses klasifikasi, sedangkan stemming digunakan untuk mengubah kata ke bentuk dasar sehingga variasi kata yang memiliki makna sama dapat direpresentasikan secara seragam (Ledjap et al., 2024; Ramadhan, 2025).

Tahapan prapemrosesan berperan penting dalam meningkatkan kualitas representasi data karena mampu mengurangi noise dan mempertahankan informasi yang relevan untuk proses klasifikasi. Dengan demikian, model yang dibangun dapat bekerja secara lebih efektif dalam mengenali pola pada teks berita.

Ekstraksi Fitur Menggunakan TF-IDF

Setelah proses prapemrosesan selesai, data teks diubah menjadi representasi numerik menggunakan metode Term Frequency-Inverse Document Frequency (TF-IDF). Metode ini digunakan untuk memberikan bobot pada setiap kata berdasarkan frekuensi kemunculannya dalam dokumen dan keseluruhan korpus. Kata yang sering muncul pada suatu dokumen tetapi jarang muncul pada dokumen lain akan memperoleh bobot yang lebih tinggi sehingga dianggap lebih representatif terhadap isi dokumen tersebut.

Penggunaan TF-IDF pada penelitian klasifikasi teks terbukti mampu meningkatkan kualitas fitur dan mendukung peningkatan performa algoritma klasifikasi. Syah et al. (2025) menjelaskan bahwa TF-IDF merupakan salah satu metode representasi teks yang efektif untuk berbagai tugas klasifikasi dokumen. Temuan serupa juga dilaporkan oleh Suryanti dan Prasetyaningrum (2025) yang menunjukkan bahwa kombinasi TF-IDF dengan algoritma Machine Learning menghasilkan performa klasifikasi yang baik.

Algoritma Klasifikasi

Penelitian ini menggunakan dua algoritma klasifikasi, yaitu Naïve Bayes dan Support Vector Machine (SVM). Naïve Bayes merupakan algoritma probabilistik yang bekerja berdasarkan Teorema Bayes dengan asumsi bahwa setiap fitur bersifat independen terhadap fitur lainnya. Algoritma ini dikenal memiliki proses pelatihan yang cepat dan efisien pada data teks (Mustafa & Alfianti, 2025; Pangesti et al., 2023).

Sementara itu, Support Vector Machine (SVM) merupakan algoritma yang bekerja dengan mencari hyperplane optimal untuk memisahkan dua kelas data. SVM memiliki kemampuan yang baik dalam menangani data berdimensi tinggi dan sering digunakan dalam berbagai penelitian klasifikasi teks karena mampu menghasilkan performa yang stabil dan akurat (Hasan Dalimunthe et al., 2024; Putri Ramadani et al., 2026). Pada penelitian ini, implementasi SVM menggunakan kernel linear karena sesuai dengan karakteristik data teks yang direpresentasikan menggunakan TF-IDF.

Evaluasi Kinerja Model

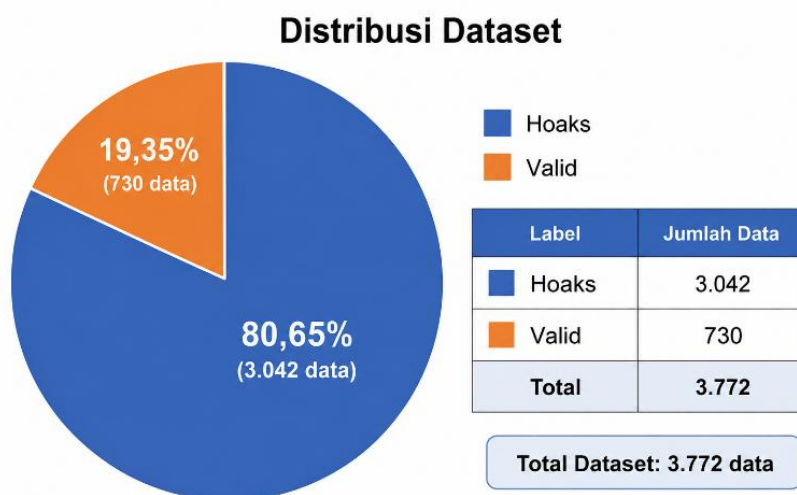
Kinerja model dievaluasi menggunakan confusion matrix dan beberapa metrik evaluasi, yaitu accuracy, precision, recall, dan F1-score. Accuracy digunakan untuk mengukur tingkat ketepatan model dalam mengklasifikasikan data secara keseluruhan. Precision menunjukkan kemampuan model dalam menghasilkan prediksi positif yang benar, sedangkan recall mengukur kemampuan model dalam mendeteksi seluruh data positif yang tersedia. F1-score digunakan untuk mengukur keseimbangan antara precision dan recall sehingga memberikan gambaran performa model yang lebih komprehensif (Desriansyah et al., 2025; Kesuma, 2026).

Hasil evaluasi dari kedua algoritma kemudian dibandingkan untuk menentukan model yang memiliki performa terbaik dalam mendeteksi berita hoaks berbahasa Indonesia berdasarkan dataset yang digunakan pada penelitian ini.

3. HASIL DAN PEMBAHASAN

3.1. Distribusi Dataset

Dataset yang digunakan dalam penelitian ini terdiri atas berita berbahasa Indonesia yang telah diklasifikasikan ke dalam kategori hoaks dan valid. Sebelum dilakukan proses klasifikasi, dataset dibagi menjadi data latih dan data uji menggunakan rasio 80:20. Pembagian data dilakukan dengan pendekatan stratified sampling untuk menjaga proporsi masing-masing kelas sehingga distribusi data pada proses pelatihan dan pengujian tetap representatif. Menurut Hasan Dalimunthe et al. (2024), distribusi data yang seimbang berperan penting dalam menghasilkan model klasifikasi yang lebih stabil. Imran et al. (2024) juga menjelaskan bahwa kualitas distribusi dataset dapat memengaruhi kemampuan model dalam mengenali pola antar kelas secara optimal.



Gambar 2. Distribusi Dataset

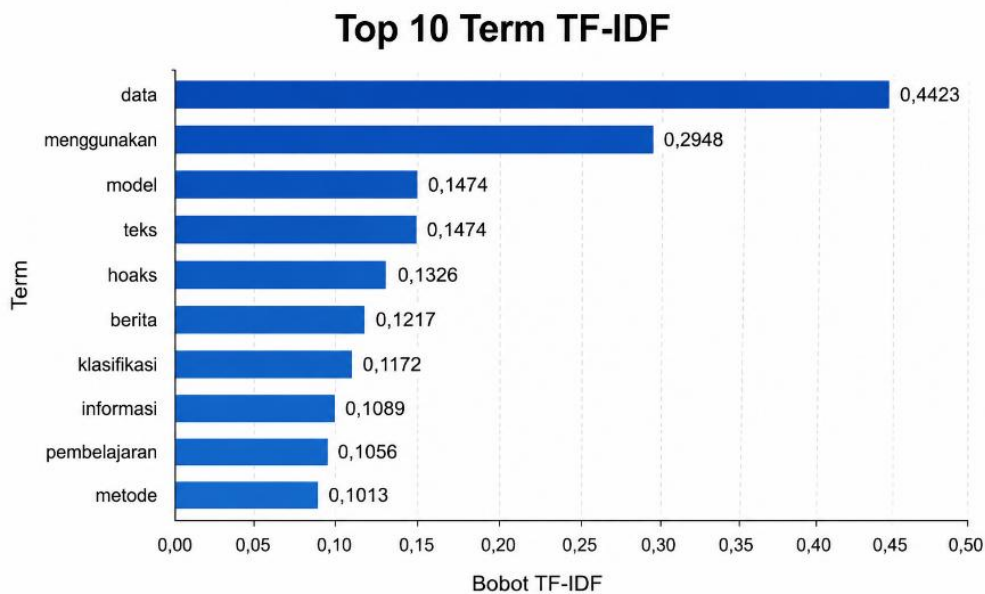
Berdasarkan Gambar 2 terlihat bahwa jumlah data hoaks lebih dominan dibandingkan data valid. Ketidakseimbangan data seperti ini umum ditemukan pada penelitian deteksi hoaks dan dapat memengaruhi performa model klasifikasi, terutama pada metrik precision dan recall (Desriansyah et al., 2025; Kesuma, 2026).

3.2. Hasil Prapemrosesan dan Ekstraksi Fitur TF-IDF

Tahap prapemrosesan dilakukan untuk meningkatkan kualitas data teks sebelum proses klasifikasi. Tahapan yang diterapkan meliputi *case folding*, *tokenization*, *stopword removal*, dan *stemming*. Proses ini bertujuan untuk menghilangkan elemen-elemen yang tidak relevan serta

menyederhanakan bentuk kata sehingga informasi penting dalam dokumen dapat dipertahankan. Ledjap et al. (2024) menyatakan bahwa prapemrosesan merupakan tahap penting dalam NLP karena mampu mengurangi noise yang dapat mengganggu proses pembelajaran model. Ramadhan (2025) juga menjelaskan bahwa kualitas hasil klasifikasi sangat dipengaruhi oleh representasi data yang dihasilkan pada tahap awal pengolahan teks.

Setelah proses prapemrosesan selesai dilakukan, tahap berikutnya adalah ekstraksi fitur menggunakan metode TF-IDF (*Term Frequency–Inverse Document Frequency*). Metode ini digunakan untuk mengubah data teks menjadi representasi numerik yang dapat diproses oleh algoritma klasifikasi. TF-IDF memberikan bobot pada setiap kata berdasarkan tingkat kemunculannya dalam dokumen dan keseluruhan korpus sehingga kata yang dianggap lebih representatif akan memperoleh bobot yang lebih tinggi.



Gambar 3. Top 10 Term TF-IDF

Berdasarkan Gambar 3, terlihat bahwa distribusi bobot TF-IDF tidak merata pada seluruh term yang muncul dalam korpus. Beberapa term memiliki bobot yang lebih tinggi dibandingkan term lainnya sehingga dianggap lebih representatif terhadap isi dokumen. Kemunculan term seperti data, menggunakan, model, teks, dan klasifikasi menunjukkan bahwa korpus penelitian memiliki keterkaitan yang kuat dengan konteks pengolahan data dan klasifikasi berita. Hasil ini menunjukkan bahwa proses ekstraksi fitur berhasil mempertahankan informasi penting yang dapat digunakan oleh algoritma klasifikasi dalam membedakan berita hoaks dan berita valid. Temuan tersebut sejalan dengan penelitian Suryanti dan Prasetyaningrum (2025) yang menyatakan bahwa penggunaan TF-IDF mampu meningkatkan kualitas representasi teks dengan menonjolkan kata-kata yang memiliki daya diskriminatif tinggi.

3.3. Hasil Klasifikasi *Naïve Bayes* dan *Support Vector Machine*

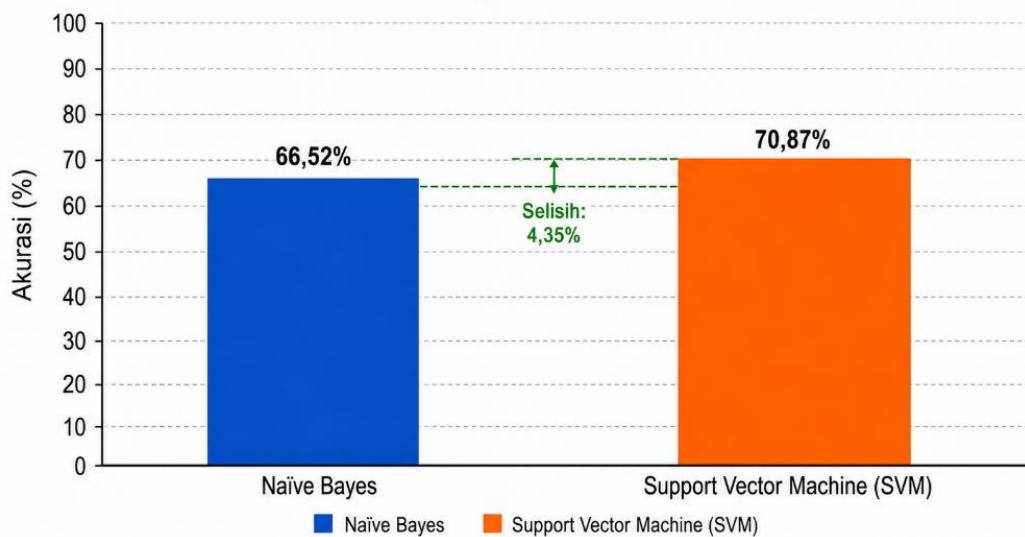
Setelah proses ekstraksi fitur selesai dilakukan, data digunakan untuk melatih model *Naïve Bayes* dan *Support Vector Machine*. Kedua model kemudian diuji menggunakan data uji yang sebelumnya tidak digunakan selama proses pelatihan. Evaluasi dilakukan menggunakan metrik accuracy, precision, recall, dan F1-score untuk memperoleh gambaran yang lebih lengkap mengenai performa masing-masing algoritma.

Tabel 1. Hasil Evaluasi Model

Metrik	Naïve Bayes	SVM
Accuracy	66,52%	70,87%
Precision	67%	72%
Recall	99%	97%
F1-Score	80%	82%

Berdasarkan Tabel 3, kedua algoritma menunjukkan performa yang cukup baik dalam mendeteksi berita hoaks berbahasa Indonesia. Namun demikian, terdapat perbedaan hasil pada setiap metrik evaluasi yang digunakan. Support Vector Machine memperoleh nilai accuracy sebesar 70,87%, lebih tinggi dibandingkan Naïve Bayes yang memperoleh nilai accuracy sebesar 66,52%. Hasil tersebut menunjukkan bahwa SVM memiliki kemampuan yang lebih baik dalam mengklasifikasikan data secara keseluruhan. Pada metrik precision, SVM juga menunjukkan performa yang lebih tinggi dengan nilai 72%, sedangkan Naïve Bayes memperoleh nilai 67%. Kondisi ini menunjukkan bahwa SVM lebih mampu menghasilkan prediksi positif yang tepat dibandingkan Naïve Bayes.

Perbandingan Akurasi Model

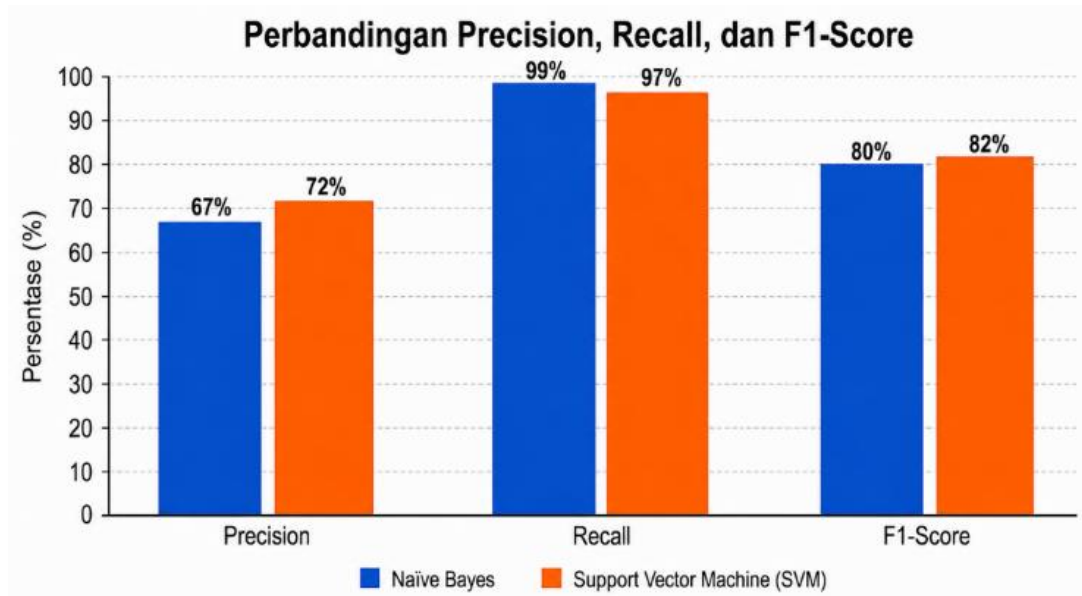


Gambar 4. Perbandingan Akurasi Model

Gambar 4 memperlihatkan perbandingan tingkat akurasi antara kedua algoritma yang digunakan dalam penelitian. Terlihat bahwa SVM memperoleh tingkat akurasi yang lebih tinggi dibandingkan Naïve Bayes dengan selisih sekitar 4,35%. Menurut Hasan Dalimunthe et al. (2024), kemampuan SVM dalam membentuk hyperplane optimal menjadi salah satu faktor yang menyebabkan algoritma ini mampu menghasilkan performa yang lebih baik pada data teks berdimensi tinggi. Temuan ini juga sejalan dengan penelitian Sani et al. (2022) yang menunjukkan bahwa SVM cenderung memberikan tingkat akurasi yang lebih tinggi dibandingkan Naïve Bayes pada klasifikasi berita hoaks.

3.4. Analisis Perbandingan Performa Model

Analisis lebih lanjut dilakukan dengan membandingkan nilai precision, recall, dan F1-score yang diperoleh oleh masing-masing algoritma. Ketiga metrik tersebut digunakan untuk memberikan gambaran yang lebih rinci mengenai kemampuan model dalam mendeteksi berita hoaks.



Gambar 5. Perbandingan *Precision*, *Recall*, dan *F1-Score*

Berdasarkan Gambar 5, SVM memperoleh nilai precision sebesar 72%, lebih tinggi dibandingkan Naïve Bayes yang memperoleh nilai sebesar 67%. Hasil tersebut menunjukkan bahwa SVM lebih efektif dalam mengurangi kesalahan klasifikasi positif (*false positive*). Kemampuan ini penting dalam sistem deteksi hoaks karena dapat mengurangi kemungkinan berita valid diklasifikasikan sebagai berita palsu. Putri Ramadani et al. (2026) menjelaskan bahwa kemampuan pemisahan kelas yang baik menjadi salah satu keunggulan utama SVM pada tugas klasifikasi teks.

Sebaliknya, Naïve Bayes memperoleh nilai recall sebesar 99%, sedikit lebih tinggi dibandingkan SVM yang memperoleh nilai 97%. Tingginya nilai recall menunjukkan bahwa Naïve Bayes mampu mendeteksi hampir seluruh berita hoaks yang terdapat dalam dataset. Pangesti et al. (2023) menjelaskan bahwa pendekatan probabilistik yang digunakan oleh Naïve Bayes membuat algoritma ini lebih sensitif terhadap kelas positif sehingga menghasilkan nilai recall yang tinggi. Meskipun demikian, nilai F1-score menunjukkan bahwa SVM tetap memiliki performa yang lebih seimbang dengan nilai sebesar 82%, sedangkan Naïve Bayes memperoleh nilai sebesar 80%.

Secara keseluruhan, hasil penelitian menunjukkan bahwa Support Vector Machine memiliki performa yang lebih unggul dibandingkan Naïve Bayes dalam mendeteksi berita hoaks berbahasa Indonesia. Keunggulan tersebut terlihat pada metrik accuracy, precision, dan F1-score yang lebih tinggi. Namun demikian, Naïve Bayes tetap memiliki kelebihan pada aspek recall dan efisiensi komputasi. Oleh karena itu, pemilihan algoritma perlu disesuaikan dengan kebutuhan sistem yang akan dikembangkan. Jika sistem lebih mengutamakan ketepatan klasifikasi dan keseimbangan performa, maka SVM merupakan pilihan yang lebih tepat. Sebaliknya, jika sistem membutuhkan proses pelatihan yang cepat dan sensitivitas tinggi terhadap deteksi berita hoaks, maka Naïve Bayes masih menjadi alternatif yang layak digunakan. Selain itu, penelitian selanjutnya dapat mengeksplorasi pendekatan berbasis Deep Learning seperti BERT dan Word2Vec yang telah

menunjukkan performa menjanjikan pada berbagai penelitian deteksi hoaks (Ripa'i et al., 2024; Andryani et al., 2025).

4. KESIMPULAN

Penelitian ini bertujuan untuk membandingkan kinerja algoritma Naïve Bayes dan Support Vector Machine (SVM) dalam mendeteksi berita hoaks berbahasa Indonesia menggunakan pendekatan Natural Language Processing (NLP) dan ekstraksi fitur Term Frequency–Inverse Document Frequency (TF-IDF). Proses penelitian meliputi tahapan prapemrosesan teks, ekstraksi fitur, pelatihan model, serta evaluasi menggunakan metrik accuracy, precision, recall, dan F1-score. Hasil penelitian menunjukkan bahwa kedua algoritma mampu melakukan klasifikasi berita hoaks dengan performa yang cukup baik, namun menghasilkan karakteristik kinerja yang berbeda.

Berdasarkan hasil pengujian, algoritma Support Vector Machine memperoleh nilai accuracy sebesar 70,87%, precision sebesar 72%, dan F1-score sebesar 82%, yang lebih tinggi dibandingkan Naïve Bayes dengan accuracy 66,52%, precision 67%, dan F1-score 80%. Sementara itu, Naïve Bayes memperoleh nilai recall sebesar 99%, sedikit lebih tinggi dibandingkan SVM yang memperoleh nilai 97%. Hasil tersebut menunjukkan bahwa SVM memiliki kemampuan yang lebih baik dalam menghasilkan klasifikasi yang akurat dan seimbang, sedangkan Naïve Bayes lebih unggul dalam mendeteksi sebagian besar data hoaks yang tersedia pada dataset.

Secara keseluruhan, Support Vector Machine dapat direkomendasikan sebagai model yang lebih efektif untuk klasifikasi berita hoaks berbahasa Indonesia karena menghasilkan performa terbaik pada sebagian besar metrik evaluasi. Temuan ini menunjukkan bahwa pendekatan berbasis hyperplane yang digunakan oleh SVM lebih mampu menangani karakteristik data teks yang direpresentasikan menggunakan TF-IDF dibandingkan pendekatan probabilistik pada Naïve Bayes. Penelitian selanjutnya dapat mengembangkan metode deteksi hoaks dengan memanfaatkan dataset yang lebih beragam, teknik seleksi fitur yang lebih optimal, serta algoritma berbasis Deep Learning seperti BERT, Word2Vec, atau Transformer untuk meningkatkan akurasi dan kemampuan generalisasi model.

UCAPAN TERIMA KASIH

Penulis mengucapkan terima kasih kepada seluruh pihak yang telah memberikan dukungan dan kontribusi dalam pelaksanaan penelitian ini. Semoga penelitian ini dapat memberikan manfaat bagi pengembangan sistem informasi pariwisata berbasis mobile web serta menjadi referensi bagi penelitian selanjutnya di bidang teknologi informasi dan pariwisata digital.

DAFTAR PUSTAKA

- Andryani, R., Julian, D., Syaputra, R., Syazili, A., Rusli, A., Ramadan, R., & Negara, E. S. (2025). Comparing deep learning and machine learning for detecting fake news on social media. *Jurnal Ilmiah Ilmu Terapan Universitas Jambi*, 9(3), 1091–1103. <https://doi.org/10.22437/jiituj.v9i3.46370>
- Desriansyah, M. D., Sari, I., & Zulfahmi, Z. (2025). Analisis efektivitas algoritma machine learning dalam deteksi hoaks pada berita digital berbahasa Indonesia. *Jurnal Sistem Informasi dan Informatika*, 3(2), 63–69. <https://doi.org/10.47233/jiska.v3i1.2024>

- Hafiizh, H., & Juanita, S. (2026). Model deteksi berita hoaks bahasa Indonesia menggunakan Multinomial Naïve Bayes dan AdaBoost classifier. *Bulletin of Computer Science Research*, 6(2), 577–584. <https://doi.org/10.47065/bulletincsr.v6i2.927>
- Hasan Dalimunthe, A., Ula, M., & Meiyanti, R. (2024). Unjuk kerja algoritma Support Vector Machine (SVM) dan Naïve Bayes dalam pengklasifikasian berita hoaks pada Twitter tentang Aksi Cepat Tanggap (ACT). *Jurnal Elektronika dan Teknologi Informasi*, 5(2), 11–22. <https://doi.org/10.5201/jet.v5i2.400>
- Huda, B. S. N., Ayni, F. N., & Angelina, D. A. (2026). Sistem pendeteksi berita hoax berbasis Word2Vec dan logistic regression. In *Seminar Nasional Teknologi & Sains* (Vol. 5, No. 1, pp. 620–625). <https://doi.org/10.29407/y6et8s42>
- Imran, B., Karim, M. N., & Ningsih, N. I. (2024). Klasifikasi berita hoax terkait pemilihan umum presiden Republik Indonesia tahun 2024 menggunakan Naïve Bayes dan SVM. *Dinamika Rekayasa*, 20(1), 1–9. <https://doi.org/10.20884/1.dinarek.2024.20.1.27>
- Karimah, U., & Fatah, Z. (2025). Klasifikasi berita hoaks di media sosial menggunakan algoritma Naive Bayes dan RapidMiner. *JISCO: Journal of Information System and Computing*, 3(2), 55–65. <https://doi.org/10.30631/jisco.v3i2.4028>
- Kesuma, C. (2026). Online news hoax detection using machine learning classification algorithms. *International Journal of Informatics, Economics, Management and Science*, 5(1), 82–94. <https://doi.org/10.52362/ijiems.v5i1.2287>
- Ledjap, A. M. B., Rochmawati, F. P., Marsanda, D. A. E., & Sari, A. P. (2024). Pemanfaatan natural language processing untuk pengecekan ejaan sesuai KBBI. *Jurnal Mahasiswa Teknik Informatika*, 3(2), 46–56. <https://doi.org/10.35473/jamastika.v3i2.3255>
- Munir, C., & Indra, I. (2024). Penerapan algoritma Naïve Bayes pada proses deteksi berita hoax pada pemberitaan media sosial: Studi kasus pemilu 2024. *Telematika MKOM*, 16(1), 1–9. <https://doi.org/10.36080/telematikamkom.3836>
- Mustafa, T. F., & Alfianti, H. (2025). Klasifikasi berita palsu berbahasa Indonesia menggunakan algoritma Naive Bayes berbasis web. *Jurnal Sains Informatika Terapan*, 4(3), 657–663. <https://doi.org/10.62357/jsit.v4i3.564>
- Pangesti, I., Zen, N. A., & Kurnianto, D. (2023). Evaluasi kinerja algoritma Naive Bayes pada sistem deteksi berita hoax. *Journal of Informatics and Communication Technology (JICT)*, 5(2), 103–110. https://doi.org/10.52661/j_ict.v5i2.238
- Putri Ramadani, N. A., Pandia, N. A., & Siregar, S. P. H. (2026). Klasifikasi berita hoaks menggunakan algoritma Support Vector Machine. *Prosiding Seminar Nasional Ilmu Teknik*, 2(2), 249–255. <https://doi.org/10.61132/prosemnasproit.v2i2.201>
- Ramadhan, M. F. (2025). Klasifikasi topik dan sentimen judul berita dengan augmentasi dan TF-IDF. *RIGGS: Journal of Artificial Intelligence and Digital Business*, 4(2), 6732–6741. <https://doi.org/10.31004/riggs.v4i2.1692>
- Ripa'i, A., Santoso, F., & Lazim, F. (2024). Hoax news detection with website comparison using deep learning approach BERT algorithm. *G-Tech: Jurnal Teknologi Terapan*, 8(3), 1749–1758. <https://doi.org/10.33379/gtech.v8i3.4541>
- Sani, R. R., Pratiwi, Y. A., Winarno, S., Udayanti, E. D., & Alzami, F. (2022). Analisis perbandingan algoritma Naive Bayes classifier dan Support Vector Machine untuk klasifikasi berita hoax pada berita online Indonesia. *Jurnal Masyarakat Informatika*, 13(2), 85–98. <https://doi.org/10.14710/jmasif.13.2.47983>

- Soleman, S. (2021). Pemanfaatan metode klasifikasi Naïve Bayes untuk pendeteksi berita hoax pada artikel berbahasa Indonesia. *Jurnal CoreIT: Jurnal Hasil Penelitian Ilmu Komputer dan Teknologi Informasi*, 7(2), 83–93. <http://dx.doi.org/10.24014/coreit.v7i2.14290>
- Suryanti, R., & Prasetyaningrum, P. (2025). Perbandingan metode TF-IDF dan Bag of Words dalam analisis sentimen diet kopi Americano di media sosial Twitter menggunakan Naïve Bayes. *Building of Informatics, Technology and Science (BITS)*, 7(1), 104–115. <https://doi.org/10.47065/bits.v7i1.7244>
- Syah, M. P., Wardani, A. P., & Idhom, M. (2025). Perbandingan representasi teks TF-IDF dan BERT terhadap akurasi cosine similarity dalam penilaian otomatis jawaban berbasis teks. *Data Sciences Indonesia (DSI)*, 5(1), 47–59. <https://doi.org/10.47709/dsi.v5i1.6021>