

KLASIFIKASI BERITA HOAKS MENGGUNAKAN METODE RANDOM FOREST BERBASIS TF-IDF

Eneng Salwa Khoerunisa¹⁾, Desti²⁾, M Revan Fatkhurezi Deviansyah³⁾

Donie Makapeli⁴⁾ Somantri⁵⁾

Teknik Informatika, Fakultas Teknik Komputer dan Desain, Universitas Nusa Putra Jl. Raya Cibolang No.21, Cibolang Kaler, Kec.Cisaat, Kabupaten Sukabumi, Jawa Barat 43152

e-mail: eneng.salwa_ti23@nusaputra.ac.id¹⁾, desti.ti23@nusaputra.ac.id²⁾, revan.fatkhurezi_ti23@nusaputra.ac.id³⁾, donie.makapeli_ti23@nusaputra.ac.id⁴⁾, somantri@nusaputra.ac.id⁵⁾

* Korespondensi: e-mail: eneng.salwa_ti23@nusaputra.ac.id

ABSTRAK

Perkembangan teknologi informasi menyebabkan penyebaran berita melalui media digital meningkat pesat. Namun, kondisi ini juga memicu maraknya penyebaran berita hoaks yang dapat menyesatkan masyarakat. Oleh karena itu, diperlukan metode otomatis untuk mengklasifikasikan berita hoaks dan berita valid. Penelitian ini bertujuan untuk menerapkan metode machine learning Random Forest dalam mengklasifikasikan berita berbahasa Indonesia menjadi kategori hoaks dan valid. Dataset yang digunakan berjumlah 1000 data berita yang telah diberi label. Tahapan penelitian meliputi preprocessing teks menggunakan stopword removal dan stemming Sastrawi, ekstraksi fitur menggunakan TF-IDF, serta pelatihan model Random Forest. Hasil pengujian menunjukkan bahwa model Random Forest mampu mencapai tingkat akurasi sebesar 86,00%, dengan nilai precision 85,56%, recall 83,70%, dan F1-score 84,62%. Hasil tersebut menunjukkan bahwa metode Random Forest memiliki kinerja yang baik dan cukup efektif dalam mendeteksi berita hoaks.

Kata Kunci: berita hoaks, klasifikasi teks, TF-IDF, Random Forest, machine learning

ABSTRACT

The development of information technology has led to a rapid increase in the dissemination of news through digital media. However, this condition has also triggered the widespread spread of hoax news that can mislead the public. Therefore, an automated method is needed to classify hoax news from valid news. This study aims to apply the Random Forest machine learning method to classify Indonesian-language news into hoax and valid categories. The dataset used consists of 1,000 labeled news items. The research stages included text preprocessing using stopword removal and Sastrawi stemming, feature extraction using TF-IDF, and training a Random Forest model. Test results showed that the Random Forest model achieved an accuracy rate of 86.00%, with a precision of 85.56%, a recall of 83.70%, and an F1-score of 84.62%. These results indicate that the Random Forest method performs well and is quite effective in detecting hoax news.

Keywords: hoax news, text classification, TF-IDF, Random Forest, machine learning

I. PENDAHULUAN

Perkembangan teknologi informasi dan komunikasi yang semakin pesat telah mengubah cara masyarakat dalam memperoleh dan menyebarkan informasi, terutama melalui media sosial. Platform seperti Facebook, Twitter, dan Instagram memungkinkan pengguna untuk berinteraksi serta membagikan berita secara cepat dan luas tanpa proses verifikasi yang memadai. Kondisi ini menyebabkan penyebaran berita *hoax* semakin sulit dikendalikan dan dapat menimbulkan dampak negatif di masyarakat, seperti perpecahan sosial, kesalahpahaman, hingga menurunnya kepercayaan publik terhadap media digital [1].

Berita *hoax* merupakan informasi palsu atau menyesatkan yang sengaja dibuat dengan tujuan untuk memengaruhi opini publik atau memperoleh keuntungan tertentu. Penyebaran *hoax* di *social media* sering kali didorong oleh faktor emosional pengguna yang dengan mudah membagikan informasi tanpa

memastikan kebenarannya terlebih dahulu. Akibatnya, proses verifikasi manual menjadi tidak efisien mengingat besarnya volume data dan kecepatan arus informasi di dunia digital [2].

Untuk mengatasi permasalahan tersebut, diperlukan sistem otomatis berbasis *machine learning* yang mampu melakukan klasifikasi teks guna membedakan antara berita *hoax* dan berita *valid* secara cepat dan akurat. *Machine learning* memungkinkan sistem untuk mempelajari pola serta karakteristik tertentu pada teks berdasarkan data pelatihan yang telah diberi label, sehingga mampu melakukan klasifikasi berita baru secara otomatis [3].

Berbagai penelitian sebelumnya telah menerapkan algoritma *machine learning* seperti *Naive Bayes* dan *Random Forest* dalam mendeteksi berita *hoax* dan menunjukkan hasil akurasi yang cukup baik [4][5]. Namun, perbedaan karakteristik data, jumlah dataset, serta metode ekstraksi fitur dapat memengaruhi performa model yang dihasilkan. Berdasarkan keterbatasan penelitian sebelumnya, diperlukan kajian lebih lanjut mengenai algoritma yang efisien dan tetap stabil pada jumlah data yang terbatas. Salah satu algoritma yang memenuhi kriteria tersebut adalah *Random Forest*, karena menggunakan pendekatan *ensemble* dengan menggabungkan banyak pohon keputusan sehingga mampu mengurangi risiko *overfitting* dan menghasilkan performa yang stabil, khususnya pada klasifikasi data teks.

Penelitian ini bertujuan untuk menerapkan dan menganalisis performa metode *Random Forest* dalam mengklasifikasikan berita *hoax* dan berita *valid* berbahasa Indonesia. Evaluasi kinerja model dilakukan menggunakan metrik *accuracy*, *precision*, *recall*, dan *F1-score* untuk mengukur kemampuan model dalam mendeteksi berita *hoax* secara efektif.

II. KAJIAN TEORI

A. Social Media

Social media merupakan media berbasis internet yang memungkinkan pengguna untuk saling berkomunikasi, berinteraksi, serta berbagi informasi dalam berbagai bentuk seperti teks, gambar dan video. Perkembangan *social media* didukung oleh teknologi Web 2.0 yang memungkinkan pengguna tidak hanya berperan sebagai penerima informasi, tetapi juga sebagai pembuat dan penyebar konten. *Social media* memungkinkan terjadinya pertukaran *user-generated content* secara luas dan interaktif.

Perkembangan *social media* membawa dampak signifikan terhadap cara masyarakat memperoleh informasi. Informasi dapat diakses dan disebar dalam waktu yang sangat singkat tanpa batasan ruang dan waktu. Namun, kemudahan tersebut juga menimbulkan permasalahan serius karena tidak semua informasi yang beredar telah melalui proses verifikasi. Minimnya pengawasan serta rendahnya literasi digital sebagian pengguna menyebabkan media sosial menjadi sarana yang sangat efektif dalam penyebaran berita *hoax*, terutama ketika informasi disajikan secara provokatif dan emosional[6].

B. Berita Hoax

Berita *hoax* merupakan informasi palsu atau menyesatkan yang disusun sedemikian rupa sehingga menyerupai berita yang benar dan mudah dipercaya oleh masyarakat. Beberapa penelitian menjelaskan bahwa berita *hoax* umumnya dibuat dengan tujuan memengaruhi opini publik, menciptakan keresahan, atau memperoleh keuntungan tertentu, baik secara ekonomi maupun politik[7]. Berita *hoax* sering disajikan dengan judul sensasional dan bahasa yang menggugah emosi pembaca.

Penyebaran berita *hoax* di *social media* kerap terjadi karena pengguna cenderung membagikan informasi berdasarkan dorongan emosional tanpa membaca isi berita secara menyeluruh atau memeriksa keabsahan sumbernya. Akibatnya, berita *hoax* dapat menyebar lebih cepat dibandingkan berita yang telah diverifikasi kebenarannya. Dampak dari penyebaran *hoax* tidak hanya menimbulkan kesalahpahaman informasi, tetapi juga berpotensi memicu konflik sosial serta menurunkan kepercayaan masyarakat terhadap media digital.

C. Text Mining

Text mining merupakan proses pengolahan data teks untuk mengekstraksi informasi dan pengetahuan yang memiliki makna. *Text mining* bertujuan untuk menemukan pola, hubungan dan struktur tersembunyi dalam kumpulan dokumen teks yang tidak terstruktur[8]. Dalam penelitian deteksi berita *hoax*, *text mining* menjadi tahap penting karena data yang digunakan umumnya berupa teks bebas dengan tingkat variasi bahasa yang tinggi.

Melalui proses prapemrosesan, teks diubah ke dalam bentuk yang lebih terstruktur dan relevan untuk dianalisis. Proses ini membantu mengurangi kompleksitas data serta meningkatkan kualitas fitur yang digunakan dalam klasifikasi. Dengan *text mining*, sistem dapat mengenali pola penggunaan kata yang membedakan berita *hoax* dan *non-hoax* sehingga mendukung proses pengambilan keputusan oleh algoritma klasifikasi.

D. Machine Learning

Machine learning merupakan cabang dari kecerdasan buatan yang memungkinkan sistem komputer untuk belajar dari data dan meningkatkan kinerjanya tanpa pemrograman eksplisit. Suatu sistem dikatakan belajar apabila kinerjanya dalam menyelesaikan tugas tertentu meningkat seiring bertambahnya pengalaman atau data yang digunakan. Pendekatan ini sangat efektif dalam menangani permasalahan klasifikasi teks yang melibatkan data dalam jumlah besar dan kompleks[9].

Dalam konteks deteksi berita *hoax*, *machine learning* digunakan untuk mempelajari pola bahasa, frekuensi kata, serta karakteristik teks dari data latih yang telah diberi label. Model yang telah dilatih kemudian digunakan untuk mengklasifikasikan berita baru secara otomatis. Kemampuan generalisasi ini menjadikan *machine learning* sebagai solusi yang efisien dalam menghadapi dinamika bahasa dan variasi konten di *social media*.

E. Random Forest

Random Forest merupakan algoritma *machine learning* berbasis *ensemble* yang digunakan untuk tugas klasifikasi dan regresi. Algoritma ini bekerja dengan membangun sejumlah pohon keputusan (*decision tree*) dan menggabungkan hasil prediksi dari masing-masing pohon untuk menentukan keputusan akhir. Pendekatan *ensemble* ini bertujuan untuk meningkatkan akurasi model serta mengurangi risiko *overfitting* yang sering terjadi pada pohon keputusan tunggal.

Pada *Random Forest*, setiap pohon keputusan dibangun menggunakan subset data latih yang dipilih secara acak melalui teknik *bootstrap sampling*. Selain itu, pada setiap proses pemisahan node hanya sebagian fitur yang dipilih secara acak untuk menentukan percabangan. Mekanisme ini membuat setiap pohon memiliki struktur yang berbeda sehingga meningkatkan kemampuan generalisasi model.

Dalam proses klasifikasi, *Random Forest* menggunakan metode *majority voting*, yaitu kelas yang paling banyak diprediksi oleh seluruh pohon keputusan akan dijadikan sebagai hasil akhir. Metode ini membuat *Random Forest* lebih stabil dan *robust* terhadap *noise* serta mampu menangani data berdimensi tinggi, seperti data teks hasil ekstraksi fitur *TF-IDF*[10]. Berdasarkan karakteristik tersebut, *Random Forest* banyak digunakan dalam penelitian klasifikasi teks, termasuk pada deteksi berita *hoax*, karena mampu memberikan performa yang stabil dan konsisten pada dataset berukuran kecil hingga menengah.

F. Klasifikasi Teks dan Evaluasi Kinerja Model

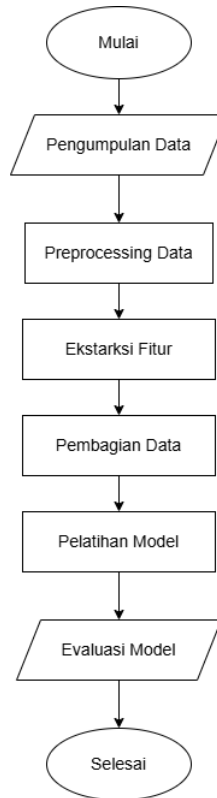
Evaluasi kinerja model klasifikasi bertujuan untuk mengukur kemampuan model dalam membedakan berita *hoax* dan *non-hoax* secara akurat. Evaluasi dilakukan menggunakan metrik *accuracy*, *precision*, *recall* dan *F1-score*. *Accuracy* digunakan untuk mengukur tingkat ketepatan model secara keseluruhan, sedangkan *precision* dan *recall* digunakan untuk menilai kemampuan model dalam mengidentifikasi berita *hoax* secara tepat dan menyeluruh.

Selain itu, *confusion matrix* digunakan untuk memberikan gambaran yang lebih rinci mengenai jumlah prediksi yang benar dan salah pada masing-masing kelas. Evaluasi yang komprehensif diperlukan agar

model yang dihasilkan tidak hanya memiliki tingkat akurasi yang tinggi, tetapi juga keseimbangan yang baik dalam mendeteksi berita *hoax*, khususnya pada dataset berukuran kecil.

III. METODE

Tahapan metode yang dilakukan dalam penelitian deteksi berita *hoax* dapat dilihat dari gambar I berikut:



Gambar I Tahapan Penelitian

Sumber : Dokumentasi Peneliti

Penelitian ini menggunakan pendekatan *eksperimen* dengan memanfaatkan teknik *text mining* dan algoritma *machine learning* untuk mengklasifikasikan berita *hoax* dan berita *valid*. Tahapan penelitian disusun secara sistematis mulai dari pengumpulan data hingga evaluasi performa model agar hasil yang diperoleh dapat dianalisis secara objektif. Penelitian ini menggunakan pendekatan eksperimental dengan beberapa tahapan utama sebagai berikut:

A. Pengumpulan Data

Data yang digunakan dalam penelitian ini berupa dataset berita berbahasa Indonesia yang telah diberi label, yaitu berita *hoax* dan berita *valid*. Dataset terdiri dari 1000 data berita yang dikumpulkan dari berbagai sumber dan telah melalui proses pelabelan sebelumnya. Keberadaan label ini memungkinkan penelitian dilakukan dengan pendekatan *supervised learning*, di mana model dilatih menggunakan data yang telah diketahui kelasnya.

B. Preprocessing Data

Tahap *preprocessing* merupakan tahapan penting dalam penelitian ini karena kualitas data teks sangat mempengaruhi kinerja model klasifikasi. Proses *preprocessing* dilakukan untuk membersihkan teks dari unsur-unsur yang tidak relevan serta menyeragamkan bentuk kata. Tahapan *preprocessing* yang dilakukan meliputi:

1. Mengubah seluruh teks menjadi huruf kecil (*case folding*) untuk menghindari perbedaan makna akibat perbedaan kapitalisasi.

2. Menghapus *URL*, angka dan tanda baca yang tidak memiliki kontribusi terhadap makna teks.
3. Menghapus kata-kata umum (*stopword*) dalam Bahasa Indonesia menggunakan *library NLTK*, karena kata-kata tersebut sering muncul namun tidak memiliki nilai diskriminatif yang tinggi.
4. Melakukan proses *stemming* menggunakan *library* sastrawi untuk mengubah kata ke bentuk dasarnya, sehingga variasi kata dapat direduksi dan memperkaya representasi fitur.

Tahapan *preprocessing* ini bertujuan untuk menghasilkan teks yang lebih bersih, ringkas dan siap digunakan pada tahap ekstraksi fitur.

C. Ekstraksi Fitur

Setelah *preprocessing*, teks berita direpresentasikan ke dalam bentuk numerik menggunakan metode *Term Frequency–Inverse Document Frequency (TF-IDF)*. Metode ini digunakan untuk mengukur tingkat kepentingan suatu kata dalam sebuah dokumen dibandingkan dengan keseluruhan dokumen dalam dataset. Pada penelitian ini, ekstraksi fitur dilakukan dengan membatasi jumlah fitur maksimum sebanyak 5.000 serta menggunakan kombinasi unigram dan bigram. Pendekatan ini diharapkan mampu menangkap konteks kata tunggal maupun pasangan kata yang sering muncul dalam berita *hoax* maupun *valid*.

D. Pembagian Data

Dataset yang telah direpresentasikan dalam bentuk fitur numerik kemudian dibagi menjadi dua bagian, yaitu data latih dan data uji. Pembagian data dilakukan dengan perbandingan 80% untuk data latih dan 20% untuk data uji menggunakan metode *train-test split*. Selain itu, pembagian data dilakukan secara *stratified* untuk menjaga proporsi kelas *hoax* dan *valid* tetap seimbang pada kedua bagian data.

E. Pelatihan Model Random Forest

Model klasifikasi yang digunakan dalam penelitian ini adalah *Random Forest*. Algoritma ini membangun sejumlah pohon keputusan dan menggabungkan hasil prediksinya untuk menghasilkan keputusan akhir. Dalam penelitian ini, *Random Forest* dikonfigurasi dengan jumlah 200 pohon keputusan. Selain itu, parameter *class weight* digunakan untuk mengatasi ketidakseimbangan jumlah data antara kelas *hoax* dan *valid*. Proses pelatihan dilakukan menggunakan data latih yang telah diekstraksi fitur *TF-IDF*-nya.

F. Evaluasi Model

Evaluasi performa model dilakukan menggunakan data uji yang sebelumnya tidak dilibatkan dalam proses pelatihan. Beberapa metrik evaluasi digunakan untuk mengukur kinerja model, yaitu akurasi, *precision*, *recall*, dan *F1-score*. Selain itu, *confusion matrix* digunakan untuk melihat secara lebih detail jumlah prediksi yang benar dan salah pada masing-masing kelas. Penggunaan beberapa metrik evaluasi ini bertujuan untuk memberikan gambaran yang lebih komprehensif terhadap kemampuan model dalam mengklasifikasikan berita *hoax* dan *valid*.

Evaluasi didasarkan pada *confusion matrix* yang terdiri dari empat komponen utama, yaitu:

- *True Positive (TP)*: jumlah berita *hoax* yang diprediksi sebagai *hoax*.
- *True Negative (TN)*: jumlah berita *valid* yang diprediksi sebagai *valid*.
- *False Positive (FP)*: jumlah berita *valid* yang salah diprediksi sebagai *hoax*.
- *False Negative (FN)*: jumlah berita *hoax* yang salah diprediksi sebagai *valid*.

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN}$$

$$Precision = \frac{TP}{TP + FP}$$

$$Recall = \frac{TP}{TP + FN}$$

$$F1 - Score = \frac{2 \times (Precision \times Recall)}{(Precision + Recall)}$$

Keterangan:

- *Accuracy*: Mengukur tingkat ketepatan model secara keseluruhan dalam mengklasifikasikan data, baik kelas positif maupun negatif.
- *Precision*: Menunjukkan tingkat ketepatan prediksi kelas positif, yaitu seberapa banyak data yang diprediksi positif benar-benar positif.
- *Recall*: Mengukur kemampuan model dalam mendeteksi seluruh data yang termasuk kelas positif.
- *F1-Score*: Merupakan rata-rata harmonik dari *precision* dan *recall*, digunakan untuk menilai keseimbangan antara ketepatan dan kelengkapan prediksi model.

IV. HASIL DAN PEMBAHASAN

A. Hasil Pengujian Model Random Forest

Pada penelitian ini, metode *Random Forest* digunakan untuk melakukan klasifikasi berita *hoax* dan berita *valid* berbasis teks. Dataset yang digunakan berjumlah 1000 data berita, yang terdiri dari 538 berita *valid* dan 462 berita *hoax*. Data dibagi menjadi data latih dan data uji menggunakan metode *train-test split* dengan proporsi 80% data latih dan 20% data uji secara *stratified* untuk menjaga keseimbangan distribusi kelas.

Data teks kemudian direpresentasikan ke dalam bentuk numerik menggunakan metode *TF-IDF*. Model *Random Forest* dilatih menggunakan 200 pohon keputusan ($n_estimators = 200$) dengan pengaturan *class_weight balanced* untuk menangani perbedaan jumlah data pada masing-masing kelas.

Hasil evaluasi model *Random Forest* terhadap data uji ditunjukkan melalui nilai *accuracy*, *precision*, *recall* dan *F1-score*. Berdasarkan hasil pengujian, model menghasilkan akurasi sebesar 86%, yang menunjukkan bahwa model mampu mengklasifikasikan berita *hoax* dan berita *valid* dengan tingkat ketepatan yang cukup baik.

B. Confusion Matrix

Pada tahap evaluasi kinerja model, digunakan *confusion matrix* untuk mengetahui sejauh mana model *Random Forest* mampu mengklasifikasikan data berita ke dalam kelas *valid* dan *hoax* secara tepat. *Confusion matrix* memberikan gambaran rinci mengenai jumlah prediksi yang benar maupun salah pada masing-masing kelas, sehingga dapat digunakan untuk menganalisis kesalahan klasifikasi yang terjadi.

Sebelum proses pelatihan (*training*) dilakukan, dataset yang digunakan merupakan data murni (*pure data*) yang terdiri dari **538 data berita valid** dan **462 data berita hoax**, sehingga total keseluruhan data berjumlah **1.000 data**. Dataset ini kemudian melalui tahap pra-proses dan pembagian data sebelum digunakan untuk membangun model klasifikasi.

Hasil pengujian model *Random Forest* terhadap data uji disajikan dalam bentuk *confusion matrix* pada Tabel I. Berdasarkan tabel tersebut, dapat diketahui bahwa model mampu mengklasifikasikan sebagian besar data dengan benar, meskipun masih terdapat sejumlah kecil kesalahan prediksi pada kedua kelas.

Tabel I Confusion Matrix Hasil Klasifikasi Random Forest

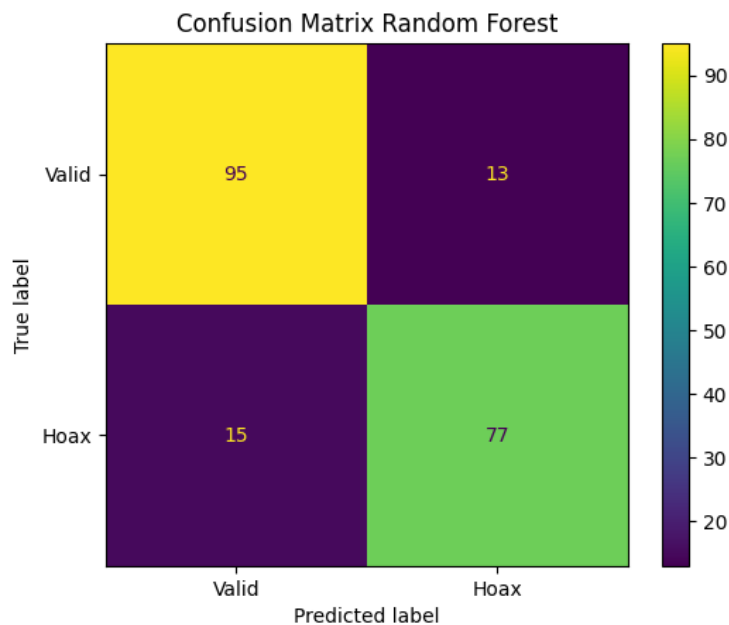
<i>Class Actual</i>	<i>Prediksi Valid</i>	<i>Prediksi Hoax</i>	<i>Total Data Aktual</i>
<i>Valid</i>	95	13	538
<i>Hoax</i>	15	77	462

Berdasarkan Tabel I, dapat dilihat bahwa:

Total data aktual merupakan jumlah data murni sebelum proses training.

- Sebanyak **95 data berita valid** berhasil diklasifikasikan dengan benar sebagai *valid*.
- Terdapat **13 data valid** yang salah diklasifikasikan sebagai *hoax*.
- Sebanyak **77 data hoax** berhasil diprediksi dengan benar sebagai *hoax*.
- Terdapat **15 data hoax** yang salah diklasifikasikan sebagai *valid*.

Hasil ini menunjukkan bahwa model *Random Forest* memiliki performa yang cukup baik dalam membedakan berita *valid* dan *hoax*. Meskipun demikian, kesalahan klasifikasi yang terjadi mengindikasikan bahwa masih terdapat kemiripan karakteristik antara beberapa berita *valid* dan *hoax*, sehingga diperlukan pengembangan lebih lanjut, baik dari sisi pemilihan fitur maupun penyesuaian parameter model, untuk meningkatkan akurasi klasifikasi.


Gambar II Visualisasi Confusion Matrix

Sumber : Dokumentasi Peneliti

C. Hasil Evaluasi Model Random Forest

Selain *confusion matrix*, performa model dievaluasi menggunakan metrik *precision*, *recall* dan *F1-score* untuk masing-masing kelas. Hasil evaluasi ditampilkan pada Tabel II

Tabel II Hasil Pengujian Model Random Forest

<i>Class</i>	<i>Precision</i>	<i>Recall</i>	<i>F1-Score</i>	<i>Support</i>
<i>Valid</i>	0.86	0.88	0.87	108

<i>Hoax</i>	0.86	0.84	0.85	92
<i>Accuracy</i>			0.86	200

Berdasarkan Tabel II diperoleh nilai:

- *Precision* sebesar 86% untuk kelas valid dan 86% untuk kelas *hoax*, yang menunjukkan tingkat ketepatan prediksi model pada kedua kelas.
- *Recall* sebesar 88% pada kelas *valid* dan 84% pada kelas *hoax*,
- yang menandakan kemampuan model dalam mengenali data aktual, dengan performa yang lebih baik pada kelas *valid*.
- *F1-score* sebesar 87% untuk kelas *valid* dan 85% untuk kelas *hoax*, yang menunjukkan keseimbangan yang baik antara *precision* dan *recall*.
- Akurasi keseluruhan sebesar 86%, yang menandakan bahwa model *Random Forest* cukup efektif dalam melakukan klasifikasi berita *hoax* dan berita *valid*.

Hasil pengujian menunjukkan bahwa metode *Random Forest* yang dikombinasikan dengan *TF-IDF* mampu memberikan performa klasifikasi yang baik dalam mendeteksi berita *hoax* pada media digital. Akurasi sebesar 86% menunjukkan bahwa mayoritas data uji dapat diklasifikasikan dengan benar oleh model.

Kesalahan klasifikasi yang terjadi diduga disebabkan oleh kemiripan kata dan struktur kalimat antara berita *hoax* dan *valid*, sehingga model mengalami kesulitan dalam membedakan konteks berita tertentu. Selain itu, penggunaan *TF-IDF* berbasis kata masih memiliki keterbatasan dalam menangkap makna semantik secara mendalam.

Meskipun demikian, hasil ini membuktikan bahwa *Random Forest* merupakan algoritma yang cukup andal untuk tugas klasifikasi teks, khususnya pada kasus deteksi berita *hoax*. Ke depannya, performa model dapat ditingkatkan dengan menambahkan jumlah data, melakukan optimasi parameter, atau mengombinasikan metode ini dengan pendekatan *deep learning* atau *word embedding*.

V. KESIMPULAN

Berdasarkan hasil penelitian yang telah dilakukan, dapat disimpulkan bahwa metode *Random Forest* yang dikombinasikan dengan teknik *TF-IDF* mampu digunakan secara efektif untuk melakukan klasifikasi berita hoaks dan valid. Proses pra-pemrosesan teks yang meliputi pembersihan data, penghapusan *stopword*, serta *stemming* menggunakan Sastrawi terbukti membantu meningkatkan kualitas fitur teks yang digunakan dalam proses klasifikasi.

Hasil pengujian model menunjukkan bahwa *Random Forest* menghasilkan akurasi sebesar 86%, dengan nilai *precision*, *recall*, dan *F1-score* yang relatif seimbang pada kedua kelas, yaitu berita *valid* dan berita *hoax*. *Confusion matrix* menunjukkan bahwa sebagian besar data uji berhasil diklasifikasikan dengan benar, meskipun masih terdapat beberapa kesalahan klasifikasi yang disebabkan oleh kemiripan pola kata dan konteks antar berita.

Secara keseluruhan, hasil penelitian ini membuktikan bahwa algoritma *Random Forest* memiliki performa yang cukup baik dan stabil dalam mendeteksi berita *hoax* berbasis teks. Model ini dapat dijadikan sebagai salah satu solusi pendukung dalam membantu penyaringan informasi di media digital. Untuk penelitian selanjutnya, disarankan untuk meningkatkan jumlah dan variasi data, melakukan optimasi parameter model, serta mengombinasikan metode ini dengan pendekatan lain seperti *word embedding* atau model *deep learning* guna memperoleh hasil klasifikasi yang lebih optimal.

VI. DAFTAR PUSTAKA

- [1] R. Wati, "Penerapan Algoritma Naive Bayes dan Particle Swarm Optimization untuk Klasifikasi Berita Hoax pada Media Sosial," *JITK (Jurnal Ilmu Pengetahuan dan Komputer)*, vol. 5, no. 2, pp. 9–14, Feb. 2020.
- [2] H. Yuliani, "Literasi Digital dalam Menangkal Berita Hoax di Media Sosial (Studi pada Mahasiswa FISIP Komunikasi Universitas Muhammadiyah Bengkulu)," *Jurnal Humas dan Media Kontemporer (MADIA)*, vol. 2, no. 1, Jan. 2021.
- [3] M. N. Raza, "Sistem Deteksi Berita Hoax Menggunakan Algoritma Naïve Bayes dan Random Forest pada Machine Learning," *Pondasi: Journal of Applied Science Engineering*, vol. 3, no. 1, pp. 43–57, 2024.
- [4] M. Z. Haq, C. S. Otiva, A. Ayuliana, U. W. Nuryanto, and D. Suryadi, "Algoritma Naïve Bayes untuk Mengidentifikasi Hoaks di Media Sosial," *Jurnal Minfo Polgan*, vol. 13, no. 1, pp. 1079–1084, Jul. 2024.
- [5] R. Astuti and A. T. Sipahutar, "Perbandingan klasifikasi berita hoax politik pada media sosial X menggunakan algoritma Naïve Bayes dan Random Forest," *Jurnal Ilmu Komputer, Sistem Informasi dan Teknologi Informasi (Innotech)*, vol. 2, no. 1, pp. 61–67, Jan. 2025.
- [6] D. N. Aliya and N. Yuliana, "Persepsi Publik di Indonesia Mengenai Berita Hoaks di Media Sosial," *Sindoro: Cendikia Pendidikan*, vol. 5, no. 1, pp. 1–10, 2024.
- [7] N. R. Anggraini, M. R. Juwita, M. S. Ulum, M. D. Prananta, M. F. Hidayat, and N. A. Purnama, "Hoaxes in the Digital Era: An Analysis of Social Media Users' Perceptions and Attitudes," *Jurnal Lemhannas RI (JLRI)*, vol. 12, no. 4, pp. 567–580, Dec. 2024.
- [8] M. Ula, "Analisa dan deteksi konten hoax pada media berita Indonesia menggunakan machine learning," *Jurnal Teknologi Terapan & Sains*, vol. 1, no. 2, pp. 1–8, Dec. 2020.
- [9] 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," *DSI: Jurnal Data Science Indonesia*, vol. 5, no. 1, pp. 33–46, Jul. 2025.
- [10] A. M. Wahid, Turino, K. A. Nugroho, T. Safitri, Darmono, and F. S. Utomo, "Optimasi Logistic Regression dan Random Forest untuk Deteksi Berita Hoax Berbasis TF-IDF," *Jurnal Pendidikan dan Teknologi Indonesia (JPTI)*, vol. 4, no. 8, pp. 381–392, Aug. 2024.
- [11] H. A. Santoso, E. H. Rachmawanto, A. Nugraha, A. A. Nugroho, D. R. I. M. Setiadi, and R. S. Basuki, "Hoax Classification and Sentiment Analysis of Indonesian News Using Naive Bayes Optimization," *TELKOMNIKA Telecommunication, Computing, Electronics and Control*, vol. 18, no. 2, pp. 799–806, Apr. 2020.
- [12] N. Agustina, Adrian, and M. Hermawati, "Implementasi Algoritma Naïve Bayes Classifier untuk Mendeteksi Berita Palsu pada Sosial Media," *Faktor Exacta*, vol. 14, no. 4, pp. 206–213, Dec. 2021.
- [13] R. R. Sani, Y. A. Pratiwi, S. Winarno, E. D. Udayanti, and F. A. Zami, "Analisis Perbandingan Algoritma Naive Bayes Classifier dan Support Vector Machine untuk Klasifikasi Hoax pada Berita Online Indonesia," *Jurnal Masyarakat Informatika*, vol. 13, no. 2, pp. 85–96, 2022.
- [14] R. Q. Putra and R. A. Saputra, "Classification of Hoax News Using the Naïve Bayes Method," *International Journal of Software Engineering and Computer Science (IJSECS)*, vol. 4, no. 1, pp. 40–48, Apr. 2024.
- [15] S. Karima and A. B. Mutiara, "Comparison of Classification Algorithms for Predicting Indonesian Fake News Using Balanced and Imbalanced Datasets," *Faktor Exacta*, vol. 16, no. 1, pp. 57–69, Mar. 2023.