

Klasifikasi Sentimen Kebohongan Berita Menggunakan Metode Indobert

Muhammad Diky Fadhilahsyah Ramadhan*, Fajri Rakhmat Umbara, Ridwan Ilyas

Universitas Jenderal Achmad Yani, Indonesia

Email: mdikyfadhilahsr20@if.unjani.ac.id

ABSTRAK

Dalam era digital, penyebaran informasi melalui berita daring berkembang pesat, tetapi ancaman disinformasi atau berita palsu menjadi tantangan signifikan. Penelitian ini bertujuan untuk mengembangkan sistem klasifikasi sentimen kebohongan berita menggunakan algoritma IndoBERT guna membedakan berita hoaks dan fakta secara akurat. Dataset yang digunakan mencakup berita hoaks dan fakta dari sumber-sumber terpercaya seperti Turnbackhoax dan Cek Fakta, untuk mengembangkan sistem klasifikasi sentimen kebohongan berita menggunakan algoritma Bidirectional Encoder Representations from Transformers (BERT) yang disesuaikan untuk bahasa Indonesia, yaitu IndoBERT. Tahapan penelitian meliputi imputasi data, pengolahan data atau pre-processing, yang meliputi pembersihan data untuk menangani masalah data yang tidak bersih, penyeimbangan data menggunakan random oversampler dan random undersampler, pembagian data (80% data latih, 20% data uji). Hasil menunjukkan bahwa model IndoBERT dengan random oversampler dan random undersampler menunjukkan akurasi yang cukup tinggi dalam klasifikasi berita palsu yaitu sebesar 99.35% berdasarkan atribut yang digunakan pada data. Penelitian ini diharapkan dapat memberikan kontribusi pada pengembangan sistem deteksi hoaks yang efektif, mendukung validasi informasi, dan mencegah dampak negatif dari penyebaran berita palsu.

Kata Kunci: IndoBERT; klasifikasi berita; hoaks; random oversampler; random undersampler

ABSTRACT

In the digital age, the spread of information through online news is rapidly growing, but the threat of misinformation or fake news poses a significant challenge. This study aims to develop a sentiment classification system for fake news using the IndoBERT algorithm to accurately distinguish between hoax and factual news. The dataset includes hoax and factual news from trusted sources such as Turnbackhoax and Cek Fakta to develop a sentiment classification system for fake news using the Bidirectional Encoder Representations from Transformers (BERT) algorithm, customized for the Indonesian language, known as IndoBERT. The research process includes data imputation, data processing or pre-processing, which involves data cleaning to handle issues with unclean data, data balancing using random oversampling and random undersampling, and data splitting (80% training data, 20% testing data). The results show that the IndoBERT model with random oversampler dan random undersampler achieves a high accuracy of 99.35% in classifying fake news based on the attributes used in the data. This research is expected to contribute to the development of an effective hoax detection system, support information validation, and prevent the negative impact of fake news dissemination.

Keyword: IndoBERT; news classification; hoax; random oversampler; random undersampler

PENDAHULUAN

Kebohongan berita sering kali menggunakan judul sensasional, kutipan palsu, atau manipulasi gambar untuk menarik perhatian pembaca (Tandoc et al., 2018). Data dari Digital 2023 Global Overview Report menunjukkan bahwa 73% populasi Indonesia aktif di media

sosial, menjadikannya salah satu pasar digital terbesar di dunia. Namun, tingginya aktivitas ini tidak diimbangi dengan literasi digital yang memadai, sehingga masyarakat rentan terpapar informasi palsu (Kemp, 2023).

Seriusnya dampak berita palsu dalam konteks media sosial dan online. Oleh karena itu, mengidentifikasi dan mendeteksi keberadaan kebohongan dalam berita menjadi semakin mendesak untuk menjaga kredibilitas informasi (Pennycook et al., 2020). Pengembangan sistem deteksi kebohongan dalam berita memerlukan pendekatan yang lebih canggih dengan memanfaatkan algoritma machine learning agar terdapatnya peningkatan akurasi dalam mengenali pola-pola kebohongan pada penerapan Natural Language Processing (NLP) memungkinkan sistem untuk memahami konteks dan makna di balik kata-kata (Shu et al., 2017).

Dalam penelitian ini, akan menggunakan algoritma Bidirectional Encoders Representational from Transformers (BERT) untuk mengidentifikasi kebohongan dalam berita. Penelitian ini dilakukan untuk bahasa Indonesia dan menggunakan dataset berita dari beberapa sumber online (Mahmud et al., 2023). Salah satu penelitian terdahulu menggunakan dataset Kaggle Berita Palsu Indonesia (Hoax) dan Deteksi Berita Hoax Indonesia untuk melakukan pemrosesan data dengan melakukan tokenisasi, pembersihan data, dan pembagian data menjadi data latih dan data uji. Metode yang digunakan dalam penelitian ini adalah algoritma klasifikasi seperti Naive Bayes, Decision Tree, Random Forest, Support Vector Machine (SVM), dan K-Nearest Neighbor (KNN). Namun, penelitian ini memiliki kekurangan dalam hal ukuran dataset yang digunakan dan kurangnya variasi jenis berita hoax yang dianalisis. Selain itu, penelitian ini hanya menggunakan teknik-teknik machine learning yang relatif sederhana dan belum mempertimbangkan teknik-teknik machine learning yang lebih canggih seperti deep learning (Triyono et al., 2022).

Perkembangan pesat dalam pemrosesan bahasa alami (NLP) dimungkinkan melalui penerapan model teknologi tinggi seperti BERT (Devlin et al., 2018). BERT, yang merupakan singkatan dari Representasi Encoder Dua Arah dari Transformers, dan RoBERTa, pendekatan BERT yang dioptimalkan, telah membuka pintu baru dalam memahami konteks linguistik dan memproses teks yang kompleks.

Kemampuan BERT dalam memahami hubungan kontekstual antar kata menjadi dasar utama penggunaannya dalam berbagai tugas NLP, termasuk klasifikasi teks dan deteksi kebohongan. BERT, menghadirkan peningkatan lebih lanjut dalam mengoptimalkan representasi vektor kata, memberikan pendekatan yang lebih kuat untuk tugas pemrosesan bahasa.

Dalam hal ini, metode evaluasi tidak hanya berfokus pada akurasi tetapi juga melibatkan pemeriksaan perilaku model, menjadi penting. Pemahaman menyeluruh tentang perilaku model NLP dapat memberikan wawasan tambahan mengenai pengembangan dan peningkatan model, terutama dalam tugas klasifikasi teks yang melibatkan berita palsu (Triyono et al., 2022).

Oleh karena itu, penelitian ini akan mengeksplorasi potensi penerapan BERT dalam klasifikasi kebohongan, mengkaji aspek penilaian pola perilaku untuk meningkatkan keandalan deteksi kebohongan berbohong berbasis teks.

Penelitian sebelumnya berfokus pada upaya mendeteksi dan mengurangi penyebaran berita palsu dari sumber tradisional seperti situs berita online atau portal berita yang memiliki celah yaitu mempertimbangkan pola pemilihan judul yang kurang baik dan pelabelan leksikal kurang spesifik dari penyebaran berita palsu di platform tersebut sehingga mempengaruhi benchmark (Murugesan, 2022).

Dalam penelitian ini, dataset yang digunakan menghadapi tantangan ketidakseimbangan data, yaitu kondisi di mana distribusi jumlah pengamatan pada label kelas dalam data latih tidak merata. Kelas mayoritas memiliki lebih banyak data dibandingkan kelas minoritas. Ketidakseimbangan ini dapat menyebabkan interpretabilitas yang rendah dan penurunan nilai akurasi dan recall (Yu et al., 2024). Sebagai alternatif, disarankan menggunakan teknik random oversampling dan random undersampling untuk menambah jumlah data dari kelas yang minoritas, sedangkan undersampling digunakan untuk menyeimbangkan data mayoritas sehingga model tidak mengalami overfitting. Kombinasi dari kedua dataset juga dilakukan untuk meningkatkan akurasi klasifikasi (Muliono et al., 2022).

Berdasarkan uraian diatas maka penelitian ini akan melakukan klasifikasi pada data berita palsu dengan menggunakan Bidirectional Encoders Representational from Transformers (BERT) model IndoBERT dan mengingat terdapat ketidakseimbangan kelas pada dataset yang digunakan maka untuk menangani masalah ini akan diambil dengan menerapkan teknik random oversampling dan random undersampling.

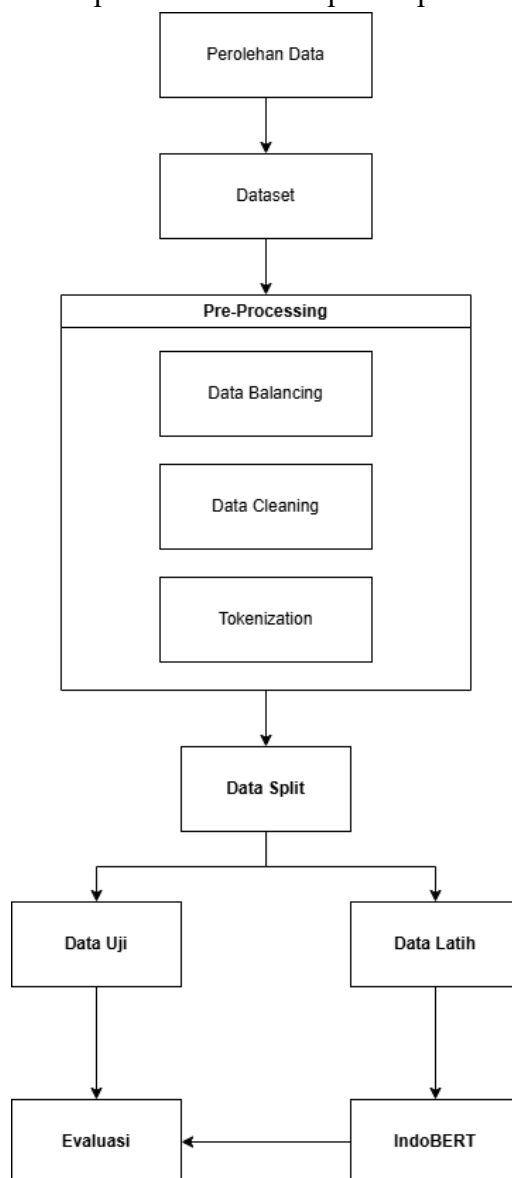
Penelitian ini mengangkat permasalahan mengenai penanganan ketidakseimbangan distribusi kelas pada dataset deteksi berita palsu agar model pembelajaran mesin dapat melakukan klasifikasi yang lebih akurat. Pendekatan dengan menerapkan teknik resampling melalui dua pendekatan yaitu random oversampling dan random undersampling yang digunakan bersama metode IndoBERT. Evaluasi komprehensif akan dilakukan untuk mengukur tingkat keberhasilan kombinasi metode tersebut dalam menyelesaikan masalah ketidakseimbangan data.

Penelitian ini memiliki tujuan yang ingin dicapai yaitu: Mengevaluasi hasil kinerja model IndoBERT dalam melakukan klasifikasi sentiment data berita hoaks berdasarkan artikel TurnBackHoax dan CekFakta. Melakukan perhitungan jumlah sentimen berita hoaks pada artikel TurnBackHoax dan CekFakta. Mengatasi ketidakseimbangan kelas pada dataset dengan menggunakan teknik Random Oversampling dan Random Undersampling dan mengevaluasi dampaknya terhadap kinerja model.

Penelitian memiliki luaran dan manfaat seperti: Pengembangan model klasifikasi teks menggunakan metode BERT dengan jenis IndoBERT untuk berita hoaks berdasarkan artikel TurnBackHoax dan CekFakta. Model ini dapat melakukan klasifikasi teks berita hoaks berhubungan dengan maraknya tersebar berita hoaks. Mendapatkan pemahaman yang lebih baik mengenai berita hoaks dan berita fakta terhadap penyebaran berita. Model klasifikasi berita yang sudah dikembangkan pada penelitian ini dapat menjadi permulaan awal untuk pengembangan model NLP lebih jauh.

METODE PENELITIAN

Metode penelitian menampilkan alur model proses penelitian yang dilakukan seperti:



Gambar 1. Diagram Alir Klasifikasi Berita Hoaks dengan IndoBERT

Sumber: Dokumen penelitian ini

1. Perolehan Data

Pengumpulan data dilakukan untuk mengumpulkan dataset untuk training dan akan diklasifikasikan untuk memprediksi perolehan berita untuk masing masing kategori berita

2. Pra-proses Data

Pada tahap preprocessing, dilakukan pengolahan dataset untuk mempersiapkannya agar dapat dengan mudah diolah oleh algoritma machine learning. Beberapa langkah yang dilakukan dalam tahap ini adalah sebagai berikut:

1) 2Data Cleaning

Pada tahapan ini dilakukan proses penyiapan data sebelum dilakukan analisis atau pembelajaran mesin dengan menghapus data yang tidak relevan atau mengganggu. Data cleaning merupakan salah satu tahap penting dalam preprocessing data.

2) Words Tokenization

Pada tahapan tokenisasi BERT teks diubah menjadi token sebelum diproses oleh model BERT. Token merupakan unit dasar dari teks yang dapat diproses oleh model BERT. Model BERT menggunakan tokenisasi subword, yaitu tokenisasi yang membagi kata-kata menjadi subword. Subword adalah unit kecil dari kata yang memiliki makna. Misalnya, kata "sleeping" dapat dipecah menjadi subword "sleep" dan "ing". Tokenisasi BERT dilakukan dengan menggunakan algoritma wordpiece. Algoritma wordpiece menggunakan kamus yang berisi subword yang umum digunakan. Algoritma ini membagi kata menjadi subword berdasarkan kamus tersebut.

3. Data Split

Penelitian ini menggunakan rasio pembagian 80:10:10, di mana 80% data digunakan untuk pelatihan model, 10% data digunakan untuk pengujian model dan 10% data digunakan untuk mengukur kinerja model. Sebelum pembagian data dilakukan, pengacakan data dilakukan untuk mengurangi varians dan memastikan generalisasi dari model-model tersebut. Pengacakan juga membantu membuat data pelatihan lebih representatif terhadap distribusi keseluruhan data dan menghindari overfitting pada model.

4. IndoBERT Modelling

IndoBERT merupakan suatu model dari Representasi Encoder Bidireksional dari Transformers (BERT) yang disiapkan dengan menggunakan data bahasa Indonesia. IndoBERT menggunakan mekanisme transformer yang mempelajari hubungan antar kata dalam sebuah teks/kalimat, yang dilatih murni sebagai masked language model (MLM) training menggunakan framework dari Huggingface (Rahmawati et al., 2022).

1) Pre-Training

Pada proses ini dilakukan proses Pre-Training menggunakan model IndoBERT. Model dilatih menggunakan lebih dari 220 juta kata yang dikumpulkan dari tiga sumber utama: (1) Wikipedia bahasa Indonesia, (2) artikel berita Kompas, Tempo dan Liputan6 serta (3) Korpus web Indonesia (Medved dan Suchomel).

2) Fine-Tuning

Pada proses ini dilakukan penyempurnaan model. *Fine-tuning* dilaksanakan dengan mempertimbangkan hyperparameter dan juga memilih lapisan output IndoBERT. Hal ini bertujuan untuk menciptakan model yang dapat mengklasifikasikan tweet secara optimal.

3) Klasifikasi Sentimen

Klasifikasi dilakukan menggunakan model IndoBERT yang telah di-fine-tune. Dari sejumlah data berlabel yang diklasifikasikan, terlihat jumlah berita yang dianggap sebagai hoax dan yang dianggap sebagai fakta.

5. Evaluasi

Tujuan evaluasi dilakukan adalah untuk mengevaluasi seberapa baik sistem dapat bekerja. Dalam penelitian ini, evaluasi dilakukan menggunakan metode cross validation. Dalam cross validation, dataset dibagi menjadi k subset, di mana salah satu subset digunakan sebagai data validasi dan subset lainnya sebagai data pelatihan. Data pelatihan digunakan untuk melatih model, sementara data validasi digunakan untuk mengukur kinerja model pada data yang tidak pernah dilihat sebelumnya. Selain itu, beberapa metrik digunakan untuk mengukur performa model, seperti akurasi dan kerugian. Akurasi mengukur sejauh mana model dapat memprediksi dengan benar, sementara kerugian mengukur sejauh mana perbedaan antara

prediksi model dan nilai sebenarnya. Dengan menggunakan metrik-metrik ini, evaluasi membantu dalam mengevaluasi sejauh mana model dapat melakukan klasifikasi dengan akurasi dan efisiensi yang baik.

HASIL DAN PEMBAHASAN

1. Implementasi Sistem

Sistem yang diimplementasikan dalam penelitian ini menggabungkan dua pendekatan, yaitu backend dan frontend. Backend dibangun menggunakan Python dan framework Flask, yang menangani proses pembelajaran, pengolahan data, dan klasifikasi. Di sisi frontend, digunakan HTML, CSS, dan JavaScript untuk merancang antarmuka pengguna. Berikut adalah tampilan antarmuka yang telah dirancang.

1) Halaman Antarmuka Dashboard

Halaman Dashboard berfungsi sebagai halaman utama sistem ini, yang menampilkan informasi mengenai judul penelitian, serta identitas mahasiswa dan pembimbing yang terlibat. Berikut adalah tampilan antarmuka dari dashboard yang dapat dilihat pada Gambar 4.1.

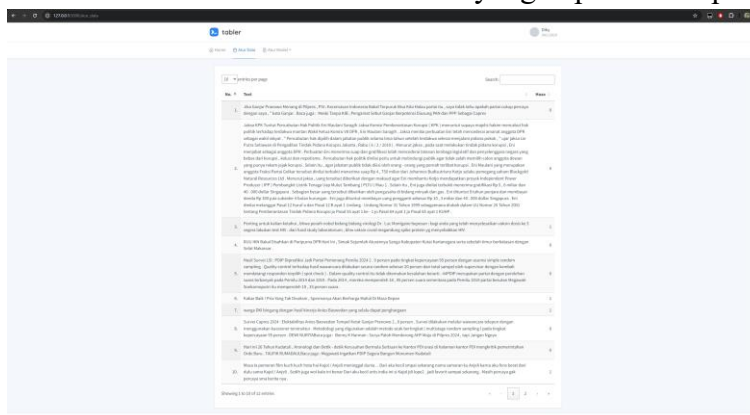


Gambar 2. Halaman Antarmuka Dashboard

Sumber: Pengembangan sistem penelitian

2) Halaman Antarmuka Atur Data

Halaman Atur Data memiliki dua fungsi utama, yaitu "Upload Data CSV" untuk mengunggah dataset, dan "Lihat Data" untuk melihat dataset yang telah diunggah. Berikut adalah tampilan antarmuka dari halaman Kelola Data yang dapat dilihat pada Gambar 4.2.

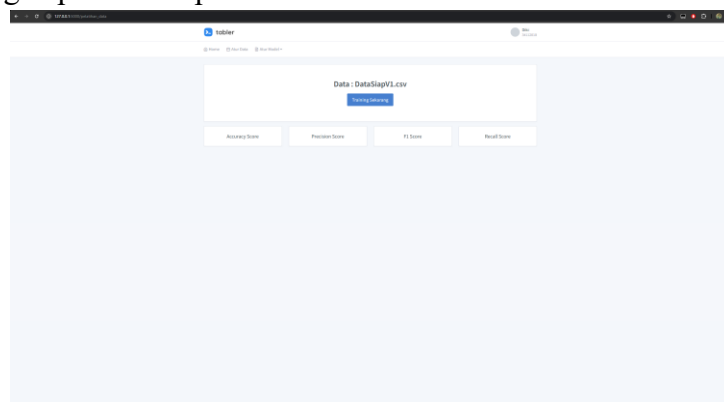


Gambar 3. Halaman Antarmuka Atur Data

Sumber: Pengembangan sistem penelitian

3) Halaman Antarmuka Pelatihan Data

Halaman Training Data memungkinkan pengguna untuk melatih model menggunakan data yang telah diunggah sebelumnya. Berikut adalah tampilan antarmuka dari halaman Training Data yang dapat dilihat pada Gambar 4.

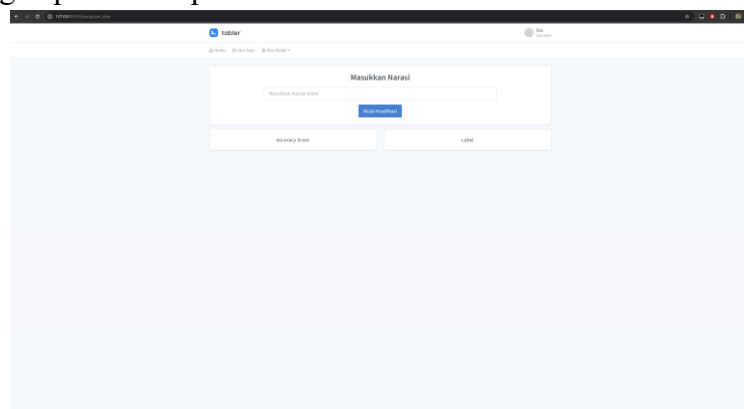


Gambar 4. Halaman Antarmuka Pelatihan Data

Sumber: Pengembangan sistem penelitian

4) Halaman Antarmuka Pengujian Data

Halaman Pengujian Data menyediakan menu untuk melakukan klasifikasi dengan memasukkan data baru melalui formulir. Hal ini memungkinkan pengguna untuk menguji performa model dan melihat hasil klasifikasi. Berikut adalah tampilan antarmuka dari halaman Testing Data yang dapat dilihat pada Gambar 4.4.



Gambar 5. Halaman Antarmuka Pengujian Data

Sumber: Pengembangan sistem penelitian

2. Pengujian Perangkat Lunak

Pengujian perangkat lunak merupakan tahap penting untuk memastikan kesesuaian antara hasil yang diharapkan dengan respons sistem yang diberikan. Pengujian dilakukan menggunakan metode yang telah ditentukan, dengan tujuan mendeteksi kesalahan sejak awal dan memperbaikinya secepat mungkin, sehingga tidak ada kesalahan yang tersisa.

1) Metode Pengujian

Metode pengujian yang digunakan dalam penelitian ini adalah black box testing. Black box testing adalah teknik pengujian perangkat lunak di mana penguji tidak perlu mengetahui kode atau struktur internal perangkat lunak yang diuji. Fokus utama dari metode ini adalah untuk memverifikasi fungsionalitas perangkat lunak dengan memberikan input

tertentu dan memeriksa output yang dihasilkan, berdasarkan spesifikasi yang telah ditetapkan.

2) Tahapan Pengujian

Tahapan pengujian perangkat lunak dalam penelitian ini terdiri dari beberapa langkah yang dirancang untuk memastikan perangkat lunak berfungsi sesuai dengan spesifikasi dan memenuhi kebutuhan pengguna. Langkah-langkah tersebut adalah sebagai berikut:

- a. 1. Pengelompokan Proses.
- b. 2. Menentukan Tujuan dari Setiap Proses Pengujian.
- c. 3. Menetapkan Kategori Keberhasilan.
- d. 4. Menyusun Skenario Pengujian.
- e. 5. Evaluasi Hasil Pengujian.

3) Pengelompokan Proses

Pada tahap pengelompokan proses, fokus utama adalah mengorganisir berbagai fitur dan fungsi perangkat lunak ke dalam dua kelompok utama, yaitu:

- a. 1. Kelola Data.
- b. 2. Kelola Model.

4) Tujuan Tiap Tahapan Pengujian

Tujuan utama dari pengujian adalah untuk memastikan bahwa perangkat lunak memenuhi semua persyaratan dan berfungsi dengan baik sesuai yang diharapkan. Berikut adalah tujuan dari setiap tahapan pengujian yang dapat dilihat dalam tabel 1.

Table 1. Tujuan Tiap Tahap Pengujian

No.	Proses	Tujuan
1.	Kelola Data	Tahapan pengujian Kelola Data bertujuan untuk memastikan bahwa fitur-fitur pengelolaan data berfungsi dengan baik. Pengujian ini mencakup dua aspek utama, yaitu menguji kemampuan sistem dalam menampilkan data untuk memastikan bahwa data ditampilkan dengan benar, serta menguji fitur unggah data dalam format CSV untuk memastikan proses unggah berjalan lancar dan data terintegrasi dengan baik ke dalam sistem.
2.	Kelola Model	Tahapan pengujian Kelola Model bertujuan untuk menguji efektivitas dan akurasi fitur model machine learning. Pengujian ini meliputi dua aspek utama, yaitu: pertama, menguji fitur pelatihan data (training data) untuk melatih model dan memastikan model mampu membuat prediksi; kedua, menguji fitur pengujian data (testing data) untuk mengevaluasi kinerja model menggunakan dataset yang berbeda dan memastikan hasil prediksinya akurat.

Sumber: Dokumen penelitian

5) Kategori Keberhasilan

Dalam proses pengujian, kategori keberhasilan dapat diidentifikasi sebagai berikut:

1. Sesuai

Pengujian "sesuai" jika fitur berfungsi sebagaimana mestinya, data ditampilkan dengan benar, unggah data CSV berjalan lancar, model berhasil dilatih, dan klasifikasi menghasilkan prediksi yang akurat saat diuji.

2. Tidak Sesuai

Pengujian "tidak sesuai" jika terdapat masalah pada fitur, tampilan data, atau proses unggah data CSV. Selain itu, model yang gagal dilatih atau menghasilkan prediksi yang tidak akurat saat diuji juga masuk dalam kategori ini.

6) Skenario Pengujian

Skenario pengujian menggambarkan serangkaian langkah yang dilakukan untuk menguji berbagai fitur dan fungsi sistem. Rincian skenario pengujian yang lebih terperinci dapat dilihat dalam Table 2.

Table 2. Skenario Pengujian

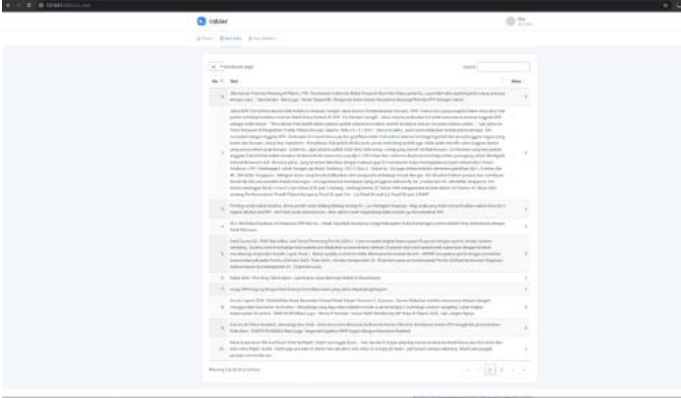
No.	Proses	Nama Fungsi	Kode Uji	Skenario
1.	Kelola Data	Lihat Data	KU-1	Menguji apakah sistem dapat menampilkan data
		Unggah Data	KU-2	Menguji kemampuan sistem untuk mengunggah data dalam format CSV dan memastikan data terintegrasi dengan baik ke dalam sistem.
2.	Kelola Model	Pelatihan Data	KU-3	Menguji fitur pelatihan model dengan menggunakan dataset yang disediakan.
		Pengujian Data	KU-4	Menguji fitur pengujian model dengan data yang diinputkan di form untuk menilai akurasi dan kinerja model.

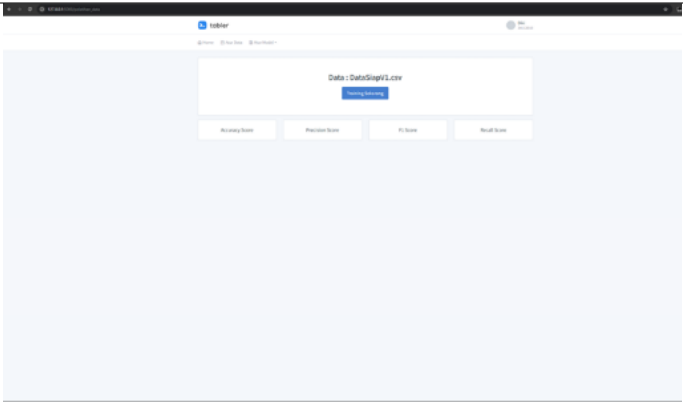
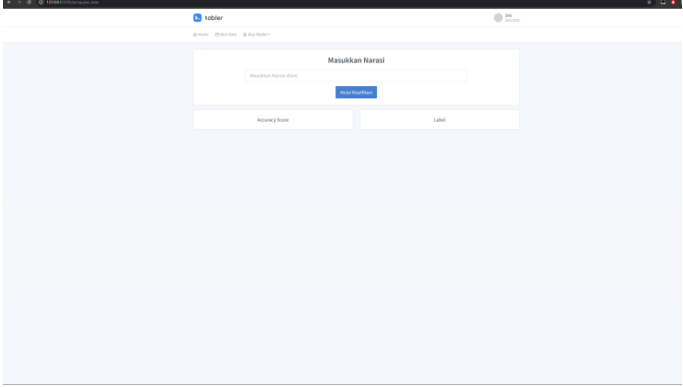
Sumber: Dokumen penelitian

7) Pelaksanaan Pengujian

Tahap pelaksanaan pengujian mencakup proses untuk menentukan apakah fungsi sistem sesuai dengan hasil yang diharapkan berdasarkan respons yang diberikan. Rincian pelaksanaan pengujian fungsionalitas dapat dilihat pada Table 3.

Table 3. Pelaksanaan Pengujian

Kode Uji	Nama Kode Uji	Hasil yang Diharapkan	Hasil yang Terjadi	Keterangan
KU-1	Lihat Data	Sistem dapat menampilkan data dengan sesuai dan lengkap.		Sesuai[✓]
KU-2	Unggah Data	Model dilatih dengan dataset yang disediakan dan menunjukkan akurasi pengujian yang		Sesuai[✓]

KU-3	Pelatihan Data	Model dilatih dengan dataset yang disediakan dan menunjukkan akurasi pengujianannya.		Sesuai[✓]
KU-4	Pengujian Data	Model memberikan hasil prediksi yang akurat dan relevan saat diuji dengan dataset yang berbeda.		Sesuai[✓]

Sumber: Dokumen penelitian

8) Evaluasi Hasil Pengujian

Berdasarkan hasil pengujian perangkat lunak yang tercantum dalam Tabel 4.3, persentase kesesuaian sistem dengan fungsi-fungsinya dapat dihitung sebagai berikut:

Jumlah Kode Uji : 4

Jumlah Kode Uji Sesuai : 4

Jumlah Kode Uji Tidak Sesuai : 0

$$Presentase = \frac{Jumlah\ Kode\ Uji - Kode\ Uji\ Tidak\ Sesuai}{Jumlah\ Kode\ Uji} \times 100\%$$

$$Presentase = \frac{4 - 0}{4} \times 100\% = 100\%$$

Berdasarkan hasil pengujian perangkat lunak yang tercantum dalam Tabel 4.3, seluruh 4 kode uji yang dilakukan menunjukkan hasil sesuai dengan harapan, tanpa adanya kode uji yang tidak sesuai. Hal ini menunjukkan bahwa sistem berhasil melewati pengujian dengan baik dan memenuhi semua fungsi yang telah ditetapkan. Dengan demikian, persentase kesesuaian sistem terhadap fungsi-fungsinya adalah 100%. Kesimpulan ini mengindikasikan bahwa sistem siap digunakan dan memiliki kinerja yang sesuai dengan ekspektasi.

3. Pengujian Model IndoBERT

Pada tahap pengujian model algoritma menggunakan IndoBERT, dilakukan dua pendekatan yaitu IndoBERT dengan Random Oversampler dan Random Undersampler dan tanpa Random Oversampler dan Random Undersampler untuk mengatasi ketidakseimbangan kelas dalam dataset. Evaluasi kinerja model dilakukan menggunakan Confusion Matrix, dengan metrik-metrik seperti accuracy, recall, precision, dan F1-score yang dihitung dari true

labels dan predicted labels. Pengujian dilakukan dengan rasio data latih 80% dan data uji 20%. Hasil pengujian menunjukkan performa model pada kedua pendekatan, memberikan gambaran tentang efektivitas Random Oversampler dan Random Undersampler dalam meningkatkan kinerja model IndoBERT. Proses Random Oversampler dan Random Undersampler, yang menggabungkan Random Oversampling dan Random Undersampling, bertujuan untuk menangani ketidakseimbangan kelas dalam dataset dan mengevaluasi kinerja model IndoBERT secara lebih optimal.

Table 4.4 Jumlah Hasil Splitting tanpa Random Oversampler dan Random Undersampler

Rasio	Proses		Total Data
	Data Latih	Data Uji	
80% : 20%	25082	6271	31353

Sumber: Hasil eksperimen penelitian

Berikut adalah jumlah pembagian data setelah dilakukan proses pembagian data tanpa menggunakan Random Oversampler dan Random Undersampler.

Table 4.5 Jumlah Hasil Splitting dengan Random Oversampler dan Random Undersampler

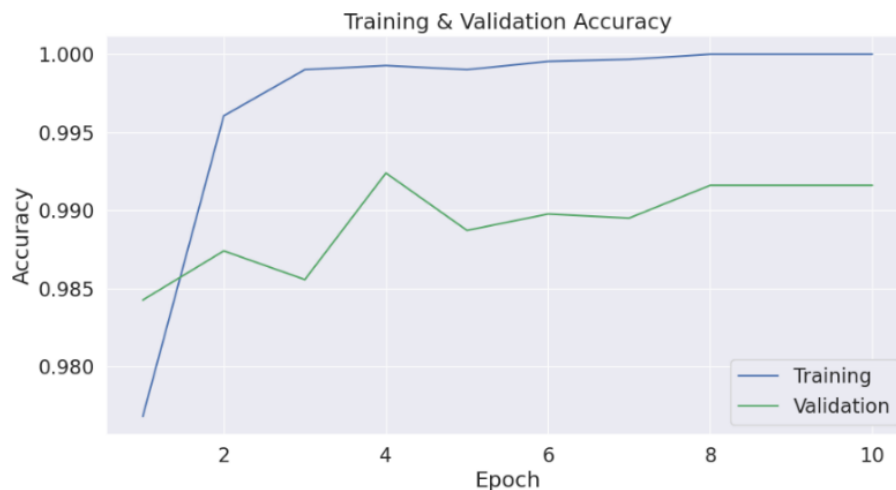
Rasio	Proses		Total Data
	Data Latih	Data Uji	
80% : 20%	19.200	4800	24000

Sumber: Hasil eksperimen penelitian

Dalam penggunaan IndoBERT dengan penggabungan Random Oversampler dan Random Undersampler berupa Random Oversampling dan Random Undersampling dapat sangat bermanfaat untuk membangun model yang optimal. Penelitian sebelumnya menunjukkan bahwa efektivitas metode IndoBERT sangat tergantung pada pengaturan parameter yang dilakukan oleh pengguna. Dengan penggunaan Random Oversampler dan Random Undersampler, pengguna dapat menyesuaikan parameter kunci secara efisien untuk meningkatkan kinerja model IndoBERT. Pendekatan ini memungkinkan peneliti untuk menemukan kombinasi parameter yang paling sesuai, sehingga model dapat beradaptasi dengan baik terhadap tugas-tugas pemrosesan bahasa alami yang spesifik.

1) Pengujian Model IndoBERT Tanpa Random Oversampler dan Random Undersampler

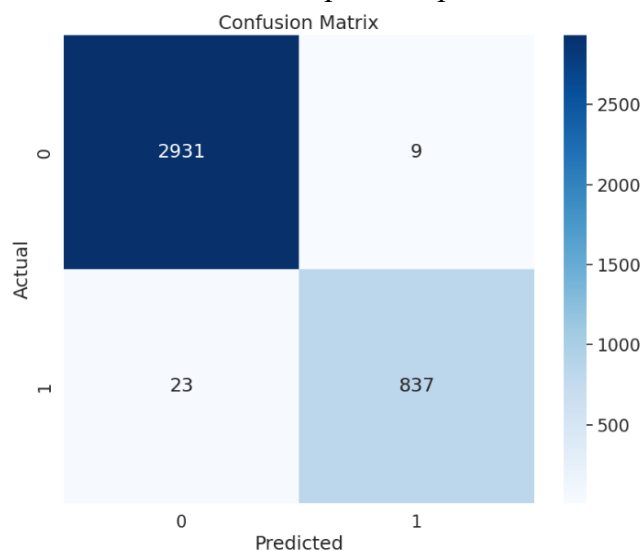
Pengujian model IndoBERT dilakukan untuk mengevaluasi kinerja model dalam pemrosesan bahasa alami tanpa menggunakan teknik resampling seperti Random Oversampler dan Random Undersampler. Hasil pengujian ini kemudian dianalisis menggunakan metrik evaluasi seperti Confusion Matrix dan Learning Curve untuk menilai kemampuan model dalam mengklasifikasikan data secara akurat. Dengan pendekatan ini, peneliti dapat memahami efektivitas IndoBERT dalam menangani tugas-tugas tertentu serta mengidentifikasi area yang perlu ditingkatkan.



Gambar 6. Learning Rate Tanpa Random Oversampler dan Random Undersampler
Sumber: Hasil eksperimen penelitian



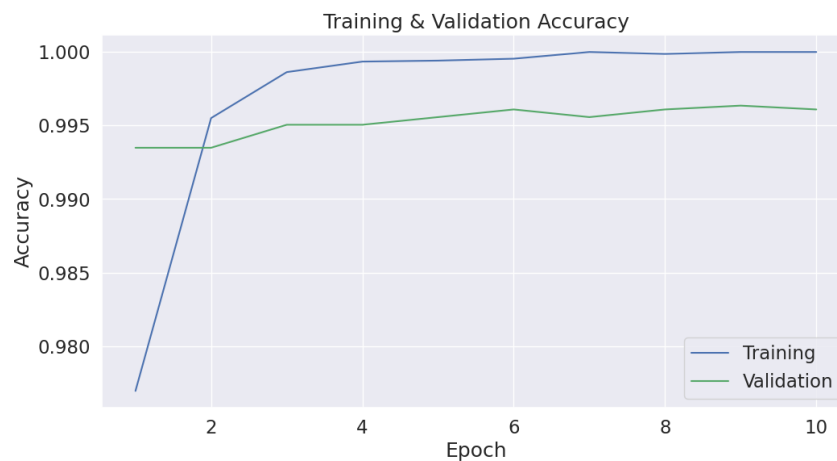
Gambar 7. Learning Rate Tanpa Random Oversampler dan Random Undersampler 2
Sumber: Hasil eksperimen penelitian



Gambar 8. Confusion Matrix Tanpa Random Oversampler dan Random Undersampler
Sumber: Hasil eksperimen penelitian

- 2) Pengujian Model IndoBERT Dengan Random Oversampler dan Random Undersampler
Pengujian model IndoBERT dilakukan untuk mengevaluasi kinerja model dalam

pemrosesan bahasa alami dengan menggunakan teknik resampling seperti Random Oversampler dan Random Undersampler. Hasil pengujian ini kemudian dianalisis menggunakan metrik evaluasi seperti Confusion Matrix dan Learning Curve untuk menilai kemampuan model dalam mengklasifikasikan data secara akurat, terutama dalam mendeteksi kelas minoritas. Dengan pendekatan ini, peneliti dapat memahami sejauh mana Random Oversampler dan Random Undersampler meningkatkan efektivitas IndoBERT dalam menangani tugas-tugas tertentu, serta mengidentifikasi area yang memerlukan perbaikan lebih lanjut.

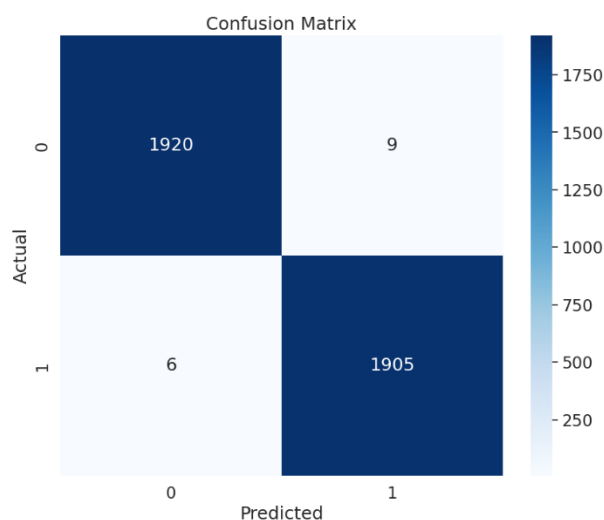


Gambar 9. Learning Rate Dengan Random Oversampler dan Random Undersampler

Sumber: Hasil eksperimen penelitian



Gambar 4.9 Learning Rate Dengan Random Oversampler dan Random Undersampler



Gambar 11. Confusion Matrix dengan Random Oversampler dan Random Undersampler

Sumber: Hasil eksperimen penelitian

4. Analisis Hasil Pengujian IndoBERT

Dari hasil perhitungan evaluation metric, dapat dilihat bahwa kedua model menunjukkan kinerja yang sangat baik dalam klasifikasi data. Model pertama memiliki akurasi sekitar 99.35%, yang menunjukkan bahwa dari keseluruhan prediksi, hampir semua klasifikasi dilakukan dengan benar. Namun, recall model pertama berada di angka 97.35%, yang berarti hanya 97.35% dari semua contoh kelas positif yang berhasil terdeteksi. Ini menunjukkan bahwa meskipun model pertama memiliki akurasi yang tinggi, masih terdapat beberapa contoh kelas positif yang tidak terdeteksi, yang dapat berpengaruh pada performa model dalam situasi di mana deteksi kelas positif sangat penting.

Sebaliknya, model kedua menunjukkan akurasi yang lebih tinggi, yaitu 99.91%, dan recall yang juga lebih baik pada angka 99.69%. Hal ini menandakan bahwa model kedua tidak hanya lebih akurat dalam klasifikasi secara keseluruhan, tetapi juga jauh lebih efektif dalam mendeteksi kelas positif. Dengan recall yang lebih tinggi, model kedua menunjukkan kemampuannya untuk mengidentifikasi hampir semua contoh kelas positif yang ada, sehingga mengurangi risiko kesalahan dalam pengambilan keputusan yang dapat terjadi jika kelas positif tidak terdeteksi.

Selain itu, precision pada model kedua mencapai 99.53%, sementara model pertama berada pada 98.93%. Precision yang lebih tinggi pada model kedua menunjukkan bahwa di antara semua prediksi kelas positif yang dihasilkan, proporsi yang benar jauh lebih besar. Ini sangat penting dalam aplikasi nyata, terutama dalam konteks di mana kesalahan klasifikasi kelas positif dapat berdampak signifikan, seperti dalam deteksi penyakit atau deteksi penipuan. Secara keseluruhan, meskipun kedua model menunjukkan hasil yang memuaskan, model kedua dengan penggunaan Random Oversampler dan Random Undersampler memberikan performa yang lebih baik dalam hal akurasi, precision, dan recall. Temuan ini menandakan bahwa teknik resampling dapat membantu meningkatkan kemampuan model dalam menangani ketidakseimbangan kelas, sehingga menjadikannya lebih dapat diandalkan untuk aplikasi yang

memerlukan deteksi yang akurat terhadap kelas minoritas.

KESIMPULAN

Penelitian ini berfokus pada klasifikasi sentimen kebohongan dalam berita dengan memanfaatkan metode IndoBERT, yang merupakan pengembangan dari model BERT yang dirancang khusus untuk bahasa Indonesia. Dalam konteks ini, IndoBERT berfungsi sebagai alat untuk menganalisis dan mengidentifikasi berita hoaks secara efektif, mengingat maraknya penyebaran informasi yang tidak akurat di era digital saat ini. Hasil analisis yang dilakukan menunjukkan bahwa model IndoBERT memiliki kemampuan yang sangat baik dalam membedakan antara berita hoaks dan fakta, dengan evaluasi yang menunjukkan akurasi dan metrik F1 yang memuaskan. Ini menegaskan bahwa IndoBERT dapat diandalkan sebagai solusi dalam mendeteksi kebohongan dalam berita. Salah satu aspek penting yang ditemukan dalam penelitian ini adalah peran krusial dari pra-pemrosesan data. Proses ini mencakup beberapa tahapan, seperti pembersihan dan tokenisasi, yang terbukti sangat signifikan dalam meningkatkan kualitas data yang digunakan untuk pelatihan model. Dengan melakukan pembersihan data, informasi yang tidak relevan atau mengganggu dapat dihilangkan, sehingga model dapat belajar dari set data yang lebih bersih dan terstruktur. Tokenisasi, di sisi lain, memungkinkan model untuk memecah teks menjadi unit-unit yang lebih kecil yang dapat diproses, yang pada gilirannya meningkatkan pemahaman model terhadap konteks linguistik yang ada. Hasil penelitian menunjukkan bahwa model IndoBERT yang dilatih model dengan tidak menggunakan random oversampling dan random undersampling menunjukkan akurasi sebesar 99.35% dan recall sebesar 97.35%, sementara untuk model yang dilatih dengan random oversampling dan random undersampling memberikan performa yang lebih baik dibandingkan dengan model tanpa pendekatan tersebut hingga mencapai akurasi sebesar 99.91% dan recall sebesar 99.69%, menunjukkan peningkatan yang signifikan dalam kemampuan model untuk mendeteksi berita palsu. Penggunaan dataset yang beragam juga merupakan faktor penentu dalam keberhasilan penelitian ini. Dataset yang diambil dari sumber-sumber terpercaya seperti TurnBackHoax dan CekFakta memberikan perspektif yang lebih luas mengenai pola-pola kebohongan yang muncul dalam berita. Dengan menggunakan data dari berbagai sumber, bias yang mungkin muncul akibat ketergantungan pada satu sumber dapat diminimalkan. Hal ini tidak hanya meningkatkan validitas hasil, tetapi juga memperkaya pemahaman tentang dinamika penyebaran informasi hoaks di masyarakat. Kontribusi penelitian ini terhadap pengembangan sistem deteksi berita hoaks di Indonesia sangat signifikan. Dengan mengimplementasikan model IndoBERT, penelitian ini tidak hanya memberikan solusi praktis untuk masalah yang semakin mendesak, tetapi juga membuka jalan bagi penelitian lebih lanjut dalam bidang pemrosesan bahasa alami dan klasifikasi teks. Dalam hal ini, hasil penelitian ini dapat menjadi referensi bagi para peneliti dan praktisi yang ingin mengembangkan alat serupa untuk mendeteksi berita hoaks atau melakukan analisis sentimen dalam konteks lain.

DAFTAR PUSTAKA

- E. C. Tandoc, Z. W. Lim, and R. Ling, "Defining 'Fake News': A typology of scholarly definitions," Feb. 07, 2018, Routledge. doi: 10.1080/21670811.2017.1360143.
- Simon Kemp, "Digital 2023: Indonesia," <https://datareportal.com/reports/digital-2023-indonesia>.

- G. Pennycook, A. Bear, E. T. Collins, and D. G. Rand, “The implied truth effect: Attaching warnings to a subset of fake news headlines increases perceived accuracy of headlines without warnings,” *Manage Sci*, vol. 66, no. 11, pp. 4944–4957, Nov. 2020, doi: 10.1287/mnsc.2019.3478.
- K. Shu, A. Sliva, S. Wang, J. Tang, and H. Liu, “Fake News Detection on Social Media: A Data Mining Perspective,” Aug. 2017, [Online]. Available: <http://arxiv.org/abs/1708.01967>
- M. A. I. Mahmud et al., “Toward News Authenticity: Synthesizing Natural Language Processing and Human Expert Opinion to Evaluate News,” *IEEE Access*, vol. 11, pp. 11405–11421, 2023, doi: 10.1109/ACCESS.2023.3241483.
- L. Triyono, R. Gernowo, M. Rahaman, and T. R. Yudiantoro, “International Journal On Informatics Visualization journal homepage : www.joiv.org/index.php/joiv International Journal On Informatics Visualization Indonesian Fake News Detection Using Various Machine Learning Technique.” [Online]. Available: www.joiv.org/index.php/joiv
- J. Devlin, M.-W. Chang, K. Lee, and K. Toutanova, “BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding,” Oct. 2018, [Online]. Available: <http://arxiv.org/abs/1810.04805>
- S. Murugesan, “Estimation Of Precision In Fake News Detection Using Novel Bert Algorithm And Comparison With Random Forest,” 2022, doi: 10.22541/au.165237518.82791368/v1.
- B. Yu, F. Tang, D. Ergu, R. Zeng, B. Ma, and F. Liu, “Efficient Classification of Malicious URLs: M-BERT - A Modified BERT Variant for Enhanced Semantic Understanding,” *IEEE Access*, vol. 12, pp. 13453–13468, 2024, doi: 10.1109/ACCESS.2024.3357095.
- Y. Muliono, F. L. Gaol, B. Soewito, and H. L. H. S. Warnars, “Hoax Classification in Imbalanced Datasets Based on Indonesian News Title using RoBERTa,” in 2022 3rd International Conference on Artificial Intelligence and Data Sciences: Championing Innovations in Artificial Intelligence and Data Sciences for Sustainable Future, AiDAS 2022 - Proceedings, Institute of Electrical and Electronics Engineers Inc., 2022, pp. 264–268. doi: 10.1109/AiDAS56890.2022.9918747.
- Giulio Alfarano; Prof. Elena Maria Baralis; Prof. Raphaël Troncy, “Detecting Fake News Using Natural Language Processing,” 2021.
- H. Xie, Z. Qin, G. Y. Li, and B. H. Juang, “Deep Learning Enabled Semantic Communication Systems,” *IEEE Transactions on Signal Processing*, vol. 69, pp. 2663–2675, 2021, doi: 10.1109/TSP.2021.3071210.
- A. Rahmawati, A. Alamsyah, and A. Romadhony, “Hoax News Detection Analysis using IndoBERT Deep Learning Methodology,” in 2022 10th International Conference on Information and Communication Technology, ICoICT 2022, Institute of Electrical and Electronics Engineers Inc., 2022, pp. 368–373. doi: 10.1109/ICoICT55009.2022.9914902.
- M. Q. Alnabhan and P. Branco, “BERTGuard: Two-Tiered Multi-Domain Fake News Detection with Class Imbalance Mitigation,” *Big Data and Cognitive Computing*, vol. 8, no. 8, Aug. 2024, doi: 10.3390/bdcc8080093.
- F. Koto, A. Rahimi, J. H. Lau, and T. Baldwin, “IndoLEM and IndoBERT: A Benchmark

- Dataset and Pre-trained Language Model for Indonesian NLP,” Online. [Online]. Available: <https://huggingface.co/>
- I. Budiman and R. Ramadina, “Penerapan Fungsi Data Mining Klasifikasi untuk Prediksi Masa Studi Mahasiswa Tepat Waktu pada Sistem Informasi Akademik Perguruan Tinggi,” *IJCCS*, vol. x, No.x, no. 1, pp. 1–5.
- S. Hendrian, “Algoritma Klasifikasi Data Mining Untuk Memprediksi Siswa Dalam Memperoleh Bantuan Dana Pendidikan,” *Faktor Exacta*, vol. 11, no. 3, Oct. 2018, doi: 10.30998/faktorexacta.v11i3.2777.
- A. Kurniawan et al., “Pengaruh Transformasi Digital Terhadap Kinerja Bank Pembangunan Daerah Di Indonesia Program Studi Manajemen Fakultas Ekonomi Dan Bisnis Universitas Komputer Indonesia Bandung,” 2021.
- X. Yang, Q. Kuang, W. Zhang, and G. Zhang, “AMDO: An Over-Sampling Technique for Multi-Class Imbalanced Problems,” *IEEE Trans Knowl Data Eng*, vol. 30, no. 9, pp. 1672–1685, Sep. 2018, doi: 10.1109/TKDE.2017.2761347.
- T. Hasanin and T. M. Khoshgoftaar, “The effects of random undersampling with simulated class imbalance for big data,” in *Proceedings - 2018 IEEE 19th International Conference on Information Reuse and Integration for Data Science, IRI 2018, Institute of Electrical and Electronics Engineers Inc.*, Aug. 2018, pp. 70–79. doi: 10.1109/IRI.2018.00018.
- T. Quatrani, “Introduction to the Unified Modeling Language,” 2003. [Online]. Available: <http://www-106.ibm.com/developerworks/rational/library/998.html>
- A. Aggarwal, A. Chauhan, D. Kumar, M. Mittal, and S. Verma, “Classification of Fake News by Fine-tuning Deep Bidirectional Transformers based Language Model,” *EAI Endorsed Transactions on Scalable Information Systems*, vol. 7, no. 27, pp. 1–12, 2020, doi: 10.4108/eai.13-7-2018.163973.