

ANALISIS SENTIMEN MASYARAKAT TERHADAP LARANGAN PENGECEK MENJUAL LPG 3 KG BERSUBSIDI MENGGUNAKAN ALGORITMA NAÏVE BAYES

Muhammad Hidayatul Arifin, Ruth Amelia Vega S. Meliala,
Fadilla Amanah, Windy Auli, Arnita, Fanny Ramadhani

Ilmu Komputer, Universitas Negeri Medan
Jl. William Iskandar Ps. V, Kenangan Baru, Kec. Percut Sei Tuan,
Kabupaten Deli Serdang, Sumatera Utara 20221, Indonesia
arif.4233250005@mhs.unimed.ac.id

ABSTRAK

Kebijakan pemerintah yang melarang pengecer untuk menjual gas LPG 3 kg bersubsidi menghasilkan berbagai komentar di media sosial, terutama YouTube, di mana banyak komentar mengungkapkan kekhawatiran tentang kesulitan akses dan dampak ekonomi terhadap masyarakat kecil. Dengan menggunakan algoritma *Naive Bayes*, penelitian ini mengkaji sentimen publik terhadap kebijakan tersebut. Data dikumpulkan melalui proses scraping komentar YouTube; setelah itu, IndoBERT digunakan untuk melakukan pelabelan otomatis dan pembersihan data. Teknik SMOTE digunakan untuk mengatasi ketidakseimbangan kelas dalam dataset, sedangkan metode TF-IDF digunakan untuk mengekstraksi fitur teks. Hasil analisis menunjukkan bahwa 70% komentar bersentimen negatif, dan model *multinomial Naive Bayes* mencapai akurasi 85,67%. Penggunaan *oversampling* juga terbukti meningkatkan *recall* untuk kelas sentimen positif dan netral. Hasilnya menunjukkan bahwa untuk membuat kebijakan yang diterapkan lebih mudah dipahami dan diterima oleh masyarakat, pemerintah harus merevisi strategi komunikasi dan distribusi LPG bersubsidi.

Kata kunci : *Naive Bayes, Analisis Sentimen, LPG 3 kg, YouTube, Media Sosial*

1. PENDAHULUAN

Kebijakan yang dikeluarkan oleh pemerintah sering kali mendapatkan perhatian publik, terutama ketika kebijakan tersebut berhubungan langsung dengan kebutuhan pokok masyarakat, seperti energi. Salah satu kebijakan yang baru-baru ini diterapkan adalah pembatasan penjualan gas LPG 3 Kg yang hanya boleh dijual oleh agen resmi. Tujuan utama dari kebijakan ini adalah untuk memastikan bahwa distribusi gas LPG 3 Kg bersubsidi tepat sasaran dan hanya digunakan oleh masyarakat yang membutuhkan. Namun, kebijakan ini menimbulkan beragam respons dari masyarakat, yang dapat berupa reaksi positif maupun negatif.

Pada masa kini, dengan maraknya penggunaan platform digital, media sosial seperti YouTube, Twitter, dan Instagram telah menjadi tempat utama bagi masyarakat untuk mengungkapkan pandangan dan opini mereka. Komentar yang diposting di media sosial tersebut sering kali mencerminkan sentimen yang dimiliki masyarakat terhadap berbagai kebijakan yang diterapkan oleh pemerintah. Oleh karena itu, penting untuk melakukan analisis sentimen untuk memahami secara lebih mendalam bagaimana masyarakat merespons kebijakan terkait pembatasan penjualan gas LPG 3 Kg bersubsidi ini.

Penelitian ini sangat penting karena LPG 3 kg bersubsidi sangat penting bagi berbagai kalangan, terutama rumah tangga dan UMKM yang menggunakan gas bersubsidi, atau pada keluarga berpenghasilan rendah. Kebijakan tersebut dan dampaknya berlangsung terhadap aksesibilitas ekonomi dan masyarakat yang mana menimbulkan berbagai reaksi di media sosial, yang kemudian

memicu diskusi dan opini publik yang luas. Analisis sentimen sekarang menjadi alat penting untuk mempelajari persepsi masyarakat secara menyeluruh berkat banyaknya platform digital. Informasi yang diperoleh dari ini dapat membantu pemerintah membuat kebijakan yang lebih transparan, inklusif, dan berkeadilan.

Dalam konteks ini, analisis sentimen menjadi alat yang penting untuk mengukur opini masyarakat mengenai kebijakan tersebut. Salah satu metode yang digunakan untuk analisis sentimen adalah *Naive Bayes*, yang telah terbukti efektif dalam klasifikasi teks dan menganalisis sentimen dengan akurasi yang tinggi. Hal ini relevan dengan penggunaan *Naive Bayes* dalam menganalisis komentar-komentar yang ada di media sosial, di mana data yang terlibat sangat besar dan tersebar.

Penelitian sebelumnya telah menunjukkan bahwa algoritma *Naive Bayes* efektif dalam analisis sentimen komentar YouTube. Menurut penelitian Harpizon dkk., komentar video dakwah memiliki distribusi sentimen sekitar 67% positif, 27% netral, dan 6% negatif, dengan akurasi 87% [1].

Selain itu, penelitian yang dilakukan oleh Muhamad Taufik Sugandi, Martanto, dan Umi Hayati tentang pendapat masyarakat tentang kebijakan pelayanan kesehatan memiliki akurasi 96% dari hampir 3.000 data komentar [2].

Selain itu, penelitian yang dilakukan oleh Dhaifa Farah Zhafira, Bayu Rahayudi, dan Indriati, yang menggunakan pembobotan TF-IDF untuk analisis sentimen juga menemukan bahwa akurasi tertinggi mencapai 97% [3].

Hasil-hasil ini menunjukkan bahwa Naive Bayes adalah algoritma yang dapat diandalkan dan efektif untuk mengklasifikasikan sentimen dari data komentar YouTube, dan mereka mendukung penerapannya dalam penelitian ini untuk mengungkap opini publik secara menyeluruh.

Melalui penelitian ini, diharapkan dapat diperoleh pemahaman yang lebih mendalam tentang bagaimana masyarakat memandang kebijakan yang diterapkan, yang akan sangat berguna untuk membantu pemerintah dalam mengevaluasi dan memperbaiki kebijakan tersebut di masa mendatang [4].

Oleh karena itu, penelitian ini tidak hanya memberikan kontribusi terhadap perkembangan ilmu pengetahuan dalam bidang analisis data, tetapi juga memiliki nilai praktis dalam membantu pengambilan keputusan kebijakan publik yang lebih baik.

2. TINJAUAN PUSTAKA

2.1. Data Mining

Data mining merupakan proses esensial dalam menghadapi big data, di mana kumpulan data yang sangat besar dan kompleks sering kali tidak dapat diurai dengan metode analisis tradisional. Dalam konteks rekayasa pengeboran, data mining berperan penting untuk mengungkap pola dan korelasi yang tersembunyi, yang nantinya dapat mendukung pengambilan keputusan yang lebih tepat dalam operasi pengeboran.

Big data dianggap sebagai kumpulan data dalam jumlah yang sangat besar, sangat kompleks sehingga metode analisis tradisional bahkan tidak dapat mulai memecahkannya. Oleh karena itu, analitik big data masuk untuk memproses dan dari situ menghasilkan pola dan korelasi yang konsisten yang dapat digunakan untuk menemukan dan mengkomunikasikan tips dan tren yang penting dan berguna [5].

2.2. Analisis Sentimen

Analisis sentimen merupakan sebuah metode berbasis komputer yang berfungsi untuk mengidentifikasi, mengekstrak, dan mengklasifikasikan pandangan masyarakat dari teks ke dalam kategori positif, negatif, atau netral [6] atau dari opini-opini, sentimen, serta emosi yang diekspresikan dalam teks [7], contoh pada suatu perusahaan untuk mendalami bagaimana konsumen merasakan produk mereka dan bagaimana mereka meresponsnya [8].

Metode ini sering digunakan di berbagai bidang, seperti dalam studi kebijakan publik, untuk menganalisis reaksi masyarakat terhadap suatu kebijakan. Dalam hal kebijakan energi, contohnya larangan distribusi LPG 3 Kg oleh pengecer, analisis sentimen bisa menjadi instrumen yang efektif untuk menilai tingkat penerimaan atau penolakan dari masyarakat terhadap kebijakan ini. Analisis sentimen juga memberi kesempatan kepada para pemangku kebijakan untuk membuat keputusan yang lebih baik

berdasarkan informasi yang diperoleh dari opini publik [3].

2.3. Algoritma Naïve Bayes

Dalam analisis sentimen, *Naïve Bayes* adalah salah satu algoritma yang paling umum digunakan. Algoritma ini bekerja berdasarkan teorema Bayes, yang menyatakan bahwa setiap fitur dalam dataset saling independen. Kemampuannya untuk menghasilkan akurasi yang tinggi meskipun dengan dataset yang relatif kecil adalah keunggulan utama *Naïve Bayes*. Selain itu, algoritma ini dikenal karena kemudahan implementasinya dan kecepatan komputasinya yang tinggi [9].

Dalam penelitian sebelumnya, *Naïve Bayes* telah digunakan dengan baik untuk sentimen masyarakat tentang larangan pemerintah terhadap penjualan iPhone 16, dengan tingkat akurasi 71,82% [10].

Ini menunjukkan bahwa *Naïve Bayes* adalah alat yang tepat untuk menganalisis sentimen pada kebijakan publik.

2.4. Preprocessing Data dan Pembobotan Fitur

Beberapa langkah penting dalam analisis sentimen termasuk *preprocessing* data, pembobotan kata, dan klasifikasi. *Preprocessing* meliputi *case folding* (mengubah teks menjadi huruf kecil), *removal of stopwords* (menghapus kata-kata umum yang tidak penting untuk analisis), dan *stemming* (mengembalikan kata ke bentuk aslinya). Untuk memastikan bahwa data yang akan dianalisis bersih dan siap untuk diproses lebih lanjut, tahap *preprocessing* sangat penting [11].

2.5. Labeling dengan IndoBERT

Pelabelan data adalah langkah penting dalam analisis sentimen untuk menentukan kategori sentimen yang tepat. *Labeling* dengan IndoBERT adalah penggunaan model BERT monolingual yang telah dioptimalkan untuk bahasa Indonesia untuk mengubah teks mentah menjadi label sentimen seperti positif, netral, dan negatif. IndoBERT memiliki kemampuan untuk menghasilkan representasi embedding konteks yang kaya, yang memungkinkannya mengidentifikasi lebih akurat aspek bahasa Indonesia dibandingkan dengan model multibahasa. Untuk menetapkan label akhir, skor probabilitas yang dihasilkan dari embedding tersebut kemudian dibandingkan dengan ambang batas tertentu. Untuk memudahkan analisis lebih lanjut, data ini kemudian diubah menjadi format numerik [12].

2.6. Pembobotan Kata dengan TF-IDF

Metode *Term Frequency-Inverse Document Frequency* (TF-IDF) adalah metode yang sering digunakan untuk membobot kata. Metode ini memberikan bobot yang lebih tinggi pada kata-kata yang lebih relevan berdasarkan frekuensi kata-kata tersebut muncul dalam dokumen dan seberapa unik

kata-kata tersebut dalam kumpulan dokumen. Metode ini membantu klasifikasi menjadi lebih akurat [3].

Penggunaan TF-IDF meningkatkan akurasi klasifikasi dari 91% menjadi 96% dalam penelitian tentang analisis sentimen kebijakan Kampus Merdeka [3], menunjukkan bahwa pembobotan kata dengan TF-IDF dapat meningkatkan hasil analisis sentimen secara signifikan.

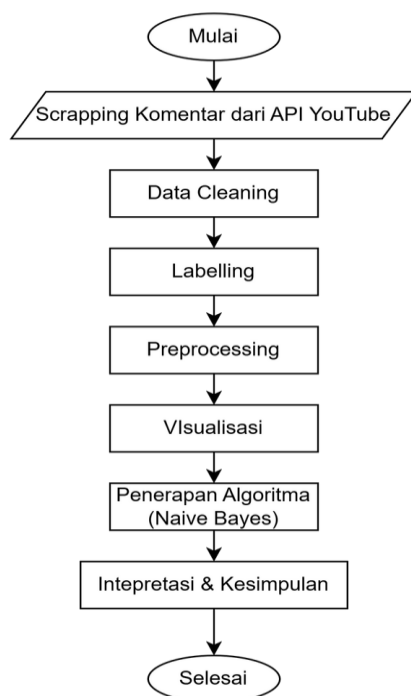
2.7. Evaluasi Model

Untuk menilai kemampuan model analisis sentimen untuk mengklasifikasikan data, confusion matrix digunakan dengan metrik seperti *accuracy*, *precision*, *recall*, dan skor F1. *Accuracy* mengukur persentase prediksi yang benar, sementara *precision* dan *recall* mengukur ketepatan dan kelengkapan prediksi. F1-score adalah rata-rata harmonik dari *precision* dan *recall*, yang memberikan gambaran keseimbangan antara kedua metrik tersebut [4].

3. METODE PENELITIAN

Penelitian ini menggunakan pendekatan kuantitatif yang berfokus pada analisis sentimen. Untuk tujuan ini, peneliti mengaplikasikan algoritma *Naïve Bayes* yang bertujuan untuk mengklasifikasikan pendapat masyarakat terkait kebijakan larangan penjualan gas LPG 3 Kg oleh pengecer. Dengan pendekatan kuantitatif, penelitian ini akan mengandalkan data numerik yang dapat dihitung dan dianalisis secara objektif untuk menggali wawasan tentang reaksi publik terhadap kebijakan tersebut. Metode ini cocok karena memungkinkan pengumpulan dan analisis data dalam jumlah besar yang terdapat dalam platform digital.

Tahapan penelitian akan dijelaskan pada gambar *flowchart* berikut.



Gambar 1. *Flowchart* Metodologi

3.1. Pengumpulan Data

Data yang digunakan dalam penelitian ini diperoleh dari berbagai sumber yang relevan dengan topik penelitian, yaitu sentimen masyarakat terhadap kebijakan larangan penjualan gas LPG 3 Kg. Data yang digunakan dalam penelitian ini bersumber dari komentar salah satu video di platform YouTube yang berjudul “Mulai 1 Februari 2025, Gas LPG 3 Kg Tidak Dijual Pengecer | IDX CHANNEL” dengan banyak data 3048 baris data atau komentar.

3.2. Data Cleaning

Setelah data dikumpulkan, langkah selanjutnya adalah pembersihan data untuk memastikan komentar yang dianalisis relevan dan berkualitas. Pertama, menghapus duplikasi, yaitu menghilangkan komentar yang sama menggunakan `drop_duplicates`. Kedua, menghapus komentar kosong dengan `dropna(subset=['comments'])` agar hanya komentar berisi teks yang diproses. Selanjutnya, membersihkan teks dengan menghapus tanda baca, karakter khusus, URL, *mention* (@username), dan *lowercase* (mengubah menjadi huruf kecil) agar tidak mengganggu analisis.

Dengan langkah-langkah ini, data yang digunakan menjadi lebih bersih dan siap untuk analisis.

3.3. Labeling

Pada tahap ini, proses pelabelan dilakukan dengan menggunakan pipeline Transformers untuk klasifikasi sentimen menggunakan model IndoBERT. Menurut Nurhasiyah, "IndoBERT menunjukkan performa unggul dalam memahami konteks bahasa Indonesia, sehingga meningkatkan akurasi penentuan sentimen." Setiap teks yang dihasilkan dari proses preprocessing dianalisis dan diberi skor sentimen, yang disebut *sentiment-analysis*. Label akhir adalah skor tertinggi yang dihasilkan oleh model [12].

IndoBERT menggunakan arsitektur Transformer dengan mekanisme *self-attention* untuk menangkap hubungan antar kata dalam konteks kalimat maupun antar kalimat secara mendalam. Model ini memulai proses dengan tahap pra-pelatihan menggunakan kumpulan besar teks berbahasa Indonesia, sehingga dapat memahami pola bahasa, struktur sintaksis, dan makna semantik yang khas dalam konteks bahasa tersebut. Setelah proses pra-pelatihan, IndoBERT kemudian dioptimalkan untuk tugas-tugas spesifik pemrosesan bahasa alami seperti analisis sentimen, klasifikasi teks, atau deteksi entitas, melalui ekstraksi fitur penting dari teks dan pemrosesan melalui lapisan-lapisan neural untuk menghasilkan output yang sesuai. Dengan pendekatan ini, IndoBERT mampu mengolah dan memahami konteks bahasa Indonesia dengan akurat, berkat pengalaman yang diperoleh dari beragam jenis teks dan situasi komunikasi.

IndoBERT secara teknis adalah versi BERT yang dilatih khusus untuk bahasa Indonesia. Ini memungkinkannya mengidentifikasi aspek semantik

husus dalam korpus lokal. Menurut Nurhasiyah, IndoBERT memanfaatkan mekanisme self-attention dua arah (bidirectional) untuk mendapatkan pemahaman yang lebih baik tentang konteks kata. Tokenisasi dilakukan secara internal oleh model; kemudian, setiap token diproses pada lapisan-lapisan encoder untuk menghasilkan representasi vektor. Dengan menggunakan representasi ini, IndoBERT dapat mengekstraksi ciri bahasa yang terkait dengan tugas klasifikasi sentimen [12].

3.4. Preprocessing

Setelah data diberi label, langkah berikutnya adalah *preprocessing* teks. Ini memungkinkan model klasifikasi untuk lebih mudah mengidentifikasi pola sentimen.

Pertama Normalisasi, kata-kata yang tidak baku atau slang harus dinormalisasi dengan menggunakan kamus yang telah didefinisikan. Untuk ilustrasi, "gak" berubah menjadi "tidak" dan "gmn" berubah menjadi "bagaimana".

Selanjutnya *Stopword*, melakukan stopwords, yaitu proses memilih kata-kata penting dari hasil token [13], kata-kata umum seperti "dan" atau "atau" dihapus, yang tidak mempengaruhi emosi secara signifikan.

Setelah itu melakukan Tokenisasi, yaitu memecah teks menjadi kata-kata [14].

Tokenisasi dilakukan dengan memecah teks menjadi kata-kata kecil, yang dikenal sebagai token, sehingga lebih mudah untuk menganalisis frekuensi atau pola kata.

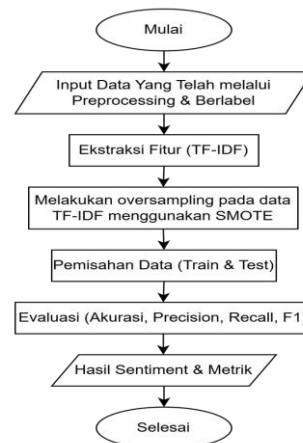
Pada langkah terakhir *Stemming*, yang berarti mengubah kata ke bentuk aslinya [14], misalnya, "memakan" atau "dimakan" menjadi "makan". Agar model tidak dibingungkan oleh variasi kata, proses ini biasanya menggunakan library seperti sastrawi.

Setelah proses ini selesai, data yang sudah siap dikonversi menjadi vektor numerik, misalnya dengan metode TF-IDF.

3.5. Visualisasi

Sebelum pelatihan model dimulai, visualisasi data dilakukan untuk mendapatkan pemahaman awal tentang fitur dataset yang digunakan. Dalam proses ini, diagram batang dibuat untuk menunjukkan distribusi sentimen. Ini dilakukan untuk melihat berapa banyak komentar yang dikategorikan sebagai positif, negatif, atau netral. Untuk menunjukkan kata atau frasa yang paling sering muncul dalam setiap kelas sentimen, juga dibuat histogram frekuensi kata. Visualisasi ini sangat membantu peneliti memahami pola dan kecenderungan umum dalam data. Selain itu, memberikan dasar untuk prosedur analisis dan klasifikasi lanjutan.

3.6. Penerapan Naïve Bayes



Gambar 2. Flowchart Alur Eksperimen (Penerapan Naïve Bayes, TF-IDF, SMOTE, Evaluasi)

Pada tahap ini, data komentar yang telah dibersihkan dan diberi label diklasifikasikan menggunakan algoritma *Naive Bayes*, terutama model MultinomialNB. Algoritma ini sangat efektif untuk analisis TF-IDF pada teks berbasis frekuensi, dan untuk memastikan proporsi sentimen tetap seimbang, data teks diubah menjadi vektor numerik menggunakan TF-IDF.

Dalam pemrosesan bahasa alami, teknik TF-IDF digunakan untuk memberikan bobot pada kata-kata dalam dokumen berdasarkan frekuensi kemunculannya dan seberapa unik kata tersebut dalam dokumen secara keseluruhan [10] [15].

Pada penelitian ini akan menggunakan teknik SMOTE. Ketika distribusi sentimen tidak merata, metode seperti SMOTE dapat digunakan untuk menyeimbangkan kelas. Menurut Pulungan dkk., SMOTE mensintesis data baru dari sampel yang ada untuk meningkatkan jumlah sampel pada kelas minoritas. Ini dilakukan dengan menggunakan rumus [16]:

$$X_{syn} = X_i + (X_{knn} - X_i) \times \delta \quad (1)$$

Keterangan:

X_{syn} : Data sintesis (sampel baru) yang dihasilkan.

X_i : Sampel minoritas yang dipilih untuk di-oversampling.

X_{knn} : Salah satu tetangga terdekat dari X_i , diperoleh dari perhitungan K-Nearest Neighbors.

δ : Nilai acak antara 0 dan 1 yang menentukan jarak sampel sintesis antara X_i dan X_{knn} .

Metode ini membantu mengurangi bias terhadap kelas mayoritas. Selain itu, kinerja model dalam mengklasifikasikan data dari kelas minoritas secara signifikan lebih baik.

Setelah model dilatih, evaluasi dilakukan untuk menilai performanya dengan menggunakan akurasi, *precision*, *recall*, dan skor F1. Pengukuran performa dilakukan dengan menghitung *statistic*, yaitu *True Positives* (TP), *True Negative* (TN), *False Positive* (FP), dan *False Negative* (FN) [9].

Berikut ini adalah rumus pada masing masing metrik [6]:

$$Accuracy = \frac{TP+TN}{TP+TN+FP+FN} \quad (2)$$

$$Precision = \frac{TP}{TP+FP} \quad (3)$$

$$Recall = \frac{TP}{TP+FN} \quad (4)$$

$$F1 - Score = 2 \times \frac{Precision \times Recall}{Precision + Recall} \quad (5)$$

Pola kesalahan klasifikasi juga diamati dengan menggunakan *confusion matrix*, adalah metode untuk mengevaluasi bagaimana model kategori biner atau lebih berfungsi pada suatu dataset [11].

3.7. Penerapan Naïve Bayes

Untuk memastikan bahwa model berfungsi dengan benar, tahap terakhir adalah menganalisis hasil klasifikasi. Untuk menilai kinerja, metrik seperti *accuracy*, *precision*, *recall*, dan skor F1. Selain itu, hasil visualisasi diperiksa untuk mengetahui apakah kata-kata kunci dalam kelas negatif menunjukkan masalah utama atau apakah komentar netral mungkin mengandung kata-kata umum. Analisis ini tidak hanya melihat bagaimana model bekerja, tetapi juga memiliki dampak pada kebijakan. Dominasi komentar negatif dapat membuat pemangku kepentingan mempertimbangkan kembali kebijakan atau strategi komunikasi mereka.

4. HASIL DAN PEMBAHASAN

4.1. Gambaran Umum Dataset

Penelitian ini mengumpulkan 3.048 komentar YouTube terkait kebijakan larangan penjualan LPG 3 kg oleh pengecer. Setelah proses *scraping* menggunakan YouTube API dan pembersihan awal, komentar yang diperoleh mencerminkan berbagai opini masyarakat, baik dukungan, penolakan, maupun netral.

Tabel 1. Dataset Hasil Scraping

username	comments
@snalpili2317	SOAL menyengsarakan rakyat, Pejabat Negara ini tidak kekurangan ide...
@HahaHa-vx2pj	Nah gitu jaga ketat
@sriusman192	Pdhl pajak unt sma rakyat, mending harga di naikin drpd dipersulit
@????????????	Dulu katanya pakai gas karena melimpah di negara Konoha, kok sekarang kayak bebanan gitu. Pakai ktp juga nanti yang mampu mendadak miskin. Yang usaha mikro beli banyak terus di jual lagi lebih mahal.
@snalpili2317	SOAL menyengsarakan rakyat, Pejabat Negara ini tidak kekurangan ide...
....	

4.2. Data Cleaning

Pada tahap pembersihan data, penelitian ini menghilangkan duplikasi, baris kosong, dan elemen teks yang tidak relevan, seperti mention, angka, dan tanda baca, serta mengubah teks menjadi huruf kecil untuk konsistensi. Meskipun jumlah data berkurang menjadi 3.027, kualitas informasi meningkat, yang membuatnya lebih cocok untuk analisis sentimen.

Tabel 2. Data Sebelum dan Sesudah Cleaning

Sebelum Cleaning	Sesudah Cleaning
SOAL menyengsarakan rakyat, Pejabat Negara ini tidak kekurangan ide...	soal menyengsarakan rakyat pejabat negara ini tidak kekurangan ide

4.3. Hasil Labeling

Sebelum proses klasifikasi, penelitian ini melakukan *labeling* otomatis dengan IndoBERT. Setiap komentar memperoleh label positif, negatif, atau netral. Karena data relatif besar, tidak dilakukan validasi manual. Dengan demikian, keakuratan label sangat bergantung pada kemampuan IndoBERT dalam memahami konteks kalimat.

Tabel 3. Data Berlabel Sentimen

comments	label
soal menyengsarakan rakyat pejabat negara ini tidak kekurangan ide	negative
nah gitu jaga ketat	negative
pdhl pajak unt sma rakyat mending harga di naikin drpd dipersulit	negative
subsidi untuk rakyat kecil dipersulit terus kenapa pemerintah tidak mencari dana dari uang yg dikorupsi para koruptor yg jumlahnya ribuan triliun dlm setahun yg lebih besar dr subsidi gas dan minyak untuk rakyat kecil	negative
yg bikin susah itu kalangan atas merampas hak kalangan bawah di gas melon sudah di tempel khusus orang miskin masih saja oknum pns dan pegawai bumh pengusaha makan beli gas melon	negative
mesti di awasi dengan ketat	neutral

Distribusi label setelah pelabelan :

- Negatif: 2.116 komentar
- Netral: 508 komentar
- Positif: 403 komentar

4.4. Hasil Preprocessing

Tahap *preprocessing* meliputi normalisasi, stopword removal, tokenisasi, dan stemming. Hasil akhirnya adalah teks yang lebih terstruktur dan singkat, memudahkan model dalam mengenali pola sentimen.

Pada normalisasi, kata-kata yang tidak baku atau slang harus dinormalisasi dengan menggunakan kamus yang telah didefinisikan.

Tabel 4. Normalisasi

Sebelum Normalisasi	Sebelum Normalisasi
soal menyengsarakan rakyat pejabat negara ini tidak kekurangan ide	soal menyengsarakan rakyat pejabat negara ini tidak kekurangan ide

Dalam melakukan stopwords, dilakukan proses memilih kata-kata penting dari hasil token

Tabel 5. Stopword

Sebelum Stopword	Sesudah Stopword
soal menyengsarakan rakyat pejabat negara ini tidak kekurangan ide	soal menyengsarakan rakyat pejabat negara tidak kekurangan ide

Dalam melakukan tokenisasi, dilakukan proses memecah teks menjadi kata-kata.

Tabel 6. Tokenisasi

Sebelum Tokenisasi	Sesudah Tokenisasi
soal menyengsarakan rakyat pejabat negara tidak kekurangan ide	['soal', 'menyengsarakan', 'rakyat', 'pejabat', 'negara', 'tidak', 'kekurangan', 'ide']

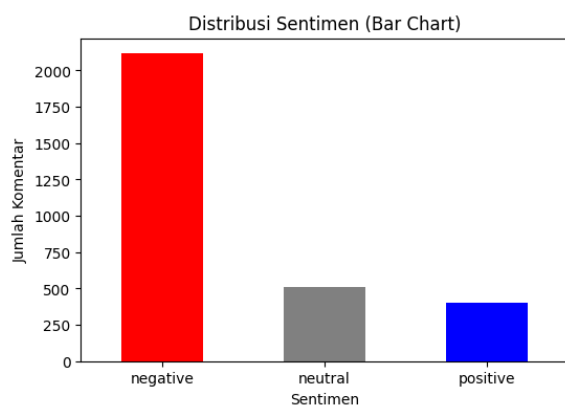
Terakhir, stemming dilakukan, yang berarti mengubah kata ke bentuk aslinya.

Tabel 7. Stemming

Sebelum Stemming	Setelah Stemming
soal menyengsarakan rakyat pejabat negara tidak kekurangan ide	Soal sengsara rakyat jabat negara tidak kurang ide

4.5. Visualisasi Data

Setelah data dibersihkan dan dilabeli, penelitian ini melakukan beberapa teknik visualisasi untuk memahami karakteristik dan persebaran sentimen pada komentar yang dikumpulkan. Berikut hasil-hasil utama visualisasi:

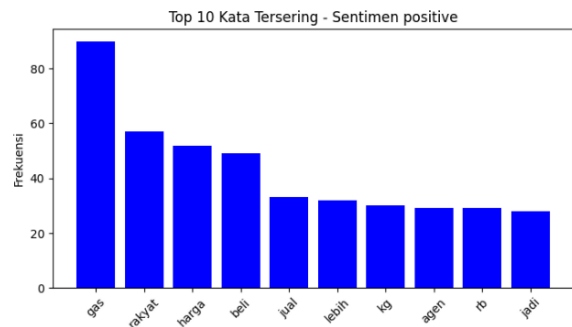


Gambar 3. Bar Chart Distribusi Sentimen

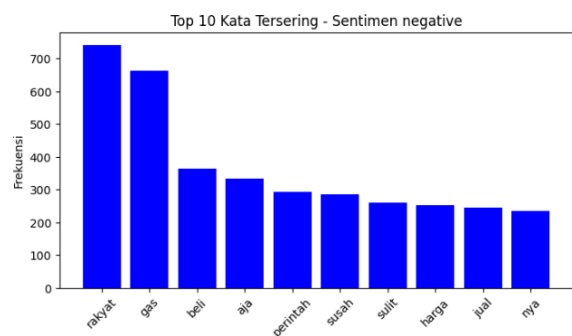
Bar chart pada Gambar 3 menunjukkan bahwa komentar negatif terhadap kebijakan yang melarang pengecer menjual LPG 3 kg dominan. Sekitar 70% komentar berisi keluhan atau kritik, 17% netral, dan hanya 13% positif. Banyak komentar negatif menyoroti kesulitan yang dihadapi masyarakat kecil

dalam mendapatkan LPG, terutama terkait dengan akses ke agen resmi dan pertanyaan tentang bagaimana distribusi yang adil.

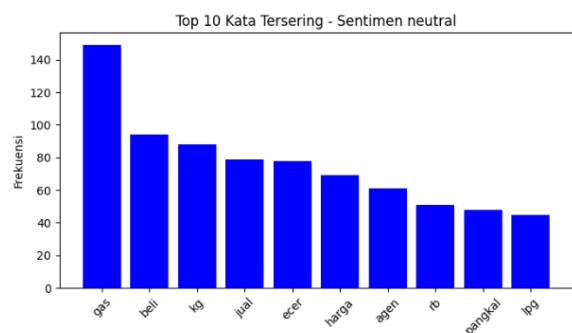
Reaksi ini menunjukkan bahwa media sosial telah menjadi cara cepat bagi masyarakat untuk menyuarakan pendapat mereka. Pemerintah harus lebih baik mengkomunikasikan kebijakannya, terutama yang berkaitan dengan kebutuhan dasar, agar tidak ditolak. Suara publik harus diperhatikan melalui visualisasi ini, yang bukan hanya data.



Gambar 4. Kata Tersering pada Sentimen Positif



Gambar 5. Kata Tersering pada Sentimen Negatif



Gambar 6. Kata Tersering pada Sentimen Netral

Pada Gambar 4, 5, dan 6 di atas menunjukkan "gas" dan "beli" sebagai kata kunci utama ketiganya. Dalam sentimen negatif pada gambar 4, kata "perintah" muncul sebagai hasil dari memotong "pemerintahan" selama *preprocessing*, disertai dengan "susah" dan "sulit" yang menunjukkan keluhan atau hambatan yang terkait dengan gas LPG 3 kg.

4.6. Hasil Penerapan Naïve Bayes

Algoritma Naïve Bayes diterapkan untuk klasifikasi sentimen. Langkah-langkahnya adalah:

4.7. Ekstraksi Fitur (TF-IDF)

Setiap komentar diubah menjadi vektor TF-IDF, menitikberatkan kata-kata yang jarang namun bermakna penting.

```
# Pembagian Data Train dan Test (80:20)
from sklearn.model_selection import train_test_split

X_train, X_test, y_train, y_test = train_test_split(
    X_res,
    y_res,
    test_size=0.2,
    random_state=42,
    stratify=y_res
)

print("Train size:", X_train.shape, y_train.shape)
print("Test size :", X_test.shape, y_test.shape)

Train size: (5078, 5969) (5078,)
Test size : (1270, 5969) (1270,)
```

Gambar 7. Ekstraksi TF-IDF

Pada Gambar 7, kode tersebut mengubah teks komentar menjadi representasi numerik menggunakan TF-IDF. Kolom 'comments' dijadikan fitur (X), sedangkan kolom 'label' adalah target (y). Setelah objek `TfidfVectorizer` dibuat, pemanggilan `fit_transform(X)` akan mempelajari kata-kata unik dan menghitung bobot TF-IDF setiap kata. Hasilnya adalah matriks (ditunjukkan oleh `X_tfidf.shape`), di mana setiap baris merepresentasikan satu komentar dan setiap kolom mewakili satu kata dalam korpus.

4.8. Penanganan Ketidakseimbangan Data

Jika hasil label menunjukkan data negatif jauh lebih banyak, penelitian ini dapat menerapkan oversampling seperti SMOTE atau model lain dengan `class_weight='balanced'`.

```
# Melakukan oversampling pada data TF-IDF menggunakan SMOTE
!pip install imblearn
from imblearn.over_sampling import SMOTE

sm = SMOTE(random_state=42)

X_res, y_res = sm.fit_resample(X_tfidf, y)

print("Distribusi label sebelum SMOTE:")
print(y.value_counts())

print("\nDistribusi label setelah SMOTE:")
print(y_res.value_counts())
```

Gambar 8. Penerapan Oversampling dengan SMOTE

Pada Gambar 8, kode tersebut memanfaatkan SMOTE untuk menyeimbangkan jumlah data di tiap kelas sentimen. Pertama, TF-IDF yang sudah dibuat (`X_tfidf`) dan label (`y`) di-fit_resample oleh SMOTE. Hasilnya adalah `X_res` dan `y_res`, di mana kelas minoritas telah ditambah secara sintetis. Output `y.value_counts()` sebelum dan sesudah SMOTE menunjukkan bahwa kelas negatif, netral, dan positif kini memiliki jumlah yang sama, sehingga model tidak lagi condong ke kelas mayoritas.

Untuk mengatasi ketidakseimbangan data, teknik *oversampling* SMOTE digunakan. Sebelum penyeimbangan, distribusi data menunjukkan ketidakseimbangan yang signifikan, dengan 2.116 komentar negatif, sementara hanya 508 komentar netral dan 403 komentar positif. SMOTE mengintesis sampel baru pada kelas minoritas sehingga jumlah data untuk setiap kelas positif, netral, dan negatif seimbang. Proses penyeimbangan ini menurunkan bias terhadap kelas mayoritas. Itu juga meningkatkan kemampuan model untuk menemukan pola dalam kelas yang sebelumnya tidak representatif. Oleh karena itu, data menjadi lebih seimbang setelah SMOTE diterapkan. Ini membantu model lebih baik dalam klasifikasi sentimen.

```
Distribusi label sebelum SMOTE:
label
negative    2116
neutral      508
positive     403
Name: count, dtype: int64

Distribusi label setelah SMOTE:
label
negative    2116
neutral    2116
positive    2116
Name: count, dtype: int64
```

Gambar 9. Jumlah Data Sebelum dan Setelah SMOTE

Gambar 9 menunjukkan perbandingan jumlah sampel untuk masing-masing label (negatif, netral, dan positif) sebelum dan sesudah teknik oversampling minoritas sintetis yang dikenal sebagai SMOTE. Sebelum SMOTE, data tidak seimbang karena jumlah sampel "negatif" lebih besar daripada "netral" dan "positif". Namun, setelah SMOTE diterapkan, jumlah sampel untuk masing-masing label menjadi 2116. Tujuannya adalah untuk menyeimbangkan data sehingga model yang dilatih dapat mempelajari pola dengan lebih adil pada setiap label dan tidak terlalu berat pada satu kelas.

4.9. Train-Test Split

Data dibagi menjadi 80% latih (train) dan 20% uji (test) secara stratified, agar proporsi setiap kelas di kedua subset tetap sama.

```
# Ekstraksi TF-IDF
from sklearn.feature_extraction.text import TfidfVectorizer

x = df['comments']
y = df['label']

vectorizer = TfidfVectorizer()
X_tfidf = vectorizer.fit_transform(x)

print("TF-IDF shape:", X_tfidf.shape)

TF-IDF shape: (3027, 5969)
```

Gambar 10. Pembagian Dataset

Pada Gambar 9, kode memecah data hasil SMOTE (`X_res`, `y_res`) menjadi train (80%) dan test (20%) menggunakan `train_test_split`. Parameter

stratify=y_res memastikan proporsi setiap kelas sentimen tetap terjaga di kedua subset. Dengan random_state=42, pembagian data konsisten di setiap eksekusi. Hasilnya, data latih berjumlah 5.078 baris, sedangkan data uji berjumlah 1.270 baris.

4.10. Proses Pelatihan

Model MultinomialNB dilatih menggunakan data latih. Semakin sering suatu kata muncul di komentar negatif, maka bobotnya akan lebih tinggi untuk memprediksi negatif, dan begitu pula untuk positif/netral.

```
# Membuat Model Naive Bayes
from sklearn.naive_bayes import MultinomialNB

model = MultinomialNB()

model.fit(X_train, y_train)

print("Model training selesai.")
```

Model training selesai.

Gambar 11. Membuat Model Naive Bayes

Pada Gambar 11, MultinomialNB lalu melatihnya dengan data latih (X_train, y_train). Setelah pemanggilan model.fit(), model telah mempelajari pola kata di setiap kelas sentimen dan siap digunakan untuk memprediksi data baru.

4.11. Evaluasi

Pada Evaluasi, dilakukan dengan mengukur akurasi, precision, recall, dan F1-score, dan menampilkan confusion matrix untuk melihat distribusi prediksi terhadap label asli di setiap kelas.

```
MultinomialNB Accuracy: 0.8567
MultinomialNB Precision: 0.8568
MultinomialNB Recall: 0.8566
MultinomialNB F1-Score: 0.8550
```

```
=== Classification Report ===
```

	precision	recall	f1-score	support
negative	0.85	0.75	0.80	423
neutral	0.83	0.89	0.86	423
positive	0.89	0.93	0.91	424
accuracy			0.86	1270
macro avg	0.86	0.86	0.85	1270
weighted avg	0.86	0.86	0.86	1270

```
=== Confusion Matrix ===
[[316  67  40]
 [ 37 377  9]
 [ 17  12 395]]
```

Gambar 12. Hasil Evaluasi Metrik dan Confusion Matrix

Pada Gambar 12, Output dihasilkan dengan memanfaatkan berbagai metrik untuk menilai kinerja model. Pertama, model.predict(X_test) menghasilkan prediksi pada data uji. Lalu, fungsi seperti accuracy_score, precision_score, recall_score, dan f1_score menghitung nilai akurasi, precision, recall, dan F1-score secara terpisah. Parameter average='macro' artinya tiap kelas dinilai sama bobot.

- Accuracy: Persentase keseluruhan prediksi yang benar.
- Precision, Recall, F1-score: Mengukur ketepatan dan kelengkapan model dalam mengenali setiap kelas sentimen.
- classification_report: Menampilkan ringkasan metrik per kelas (negatif, netral, positif).
- confusion_matrix: Menunjukkan distribusi prediksi terhadap label asli, sehingga bisa dilihat mana kelas yang sering salah diprediksi.

Dari hasil yang output, terlihat bahwa model mencapai akurasi sekitar 85,6% (MultinomialNB Accuracy: 0.8567), dengan precision, recall, dan F1-score yang cukup seimbang di ketiga kelas. Confusion matrix menggambarkan detail jumlah komentar tiap kelas yang berhasil diprediksi benar atau salah.

Berdasarkan confusion matrix, dari total 423 data yang memiliki label negatif, 316 diidentifikasi dengan akurat sebagai negatif, 67 keliru diidentifikasi sebagai netral, dan 40 salah diidentifikasi sebagai positif. Untuk 423 data yang dilabeli netral, 379 teridentifikasi dengan benar, sementara 37 terklasifikasi salah sebagai negatif dan 9 sebagai positif. Terakhir, dari 424 data yang memiliki label positif, 395 diidentifikasi dengan benar, 17 salah diidentifikasi sebagai negatif, dan 12 salah diidentifikasi sebagai netral.

5. KESIMPULAN DAN SARAN

Hasil analisis menunjukkan bahwa sentimen-sentimen negatif masih mendominasi, menunjukkan ketidakpuasan masyarakat terhadap kebijakan pemerintah yang membatasi penjualan LPG 3 kg bersubsidi, yang mana gas bersubsidi hanya dijual di pangkalan gas. Mayoritas komentar negatif menunjukkan keluhan nyata dan perbedaan makna sentimen yang ditemukan melalui visualisasi kata, sementara SMOTE membantu mengatasi ketidakseimbangan data dengan meningkatkan recall pada kelas positif dan netral. Akurasi 85,67% dicapai oleh model Naive Bayes dengan TF-IDF dan oversampling SMOTE. Berdasarkan temuan ini, pemerintah disarankan meninjau ulang jalur distribusi LPG dengan menambah titik penjualan resmi dan meningkatkan pengawasan.

Diharapkan penelitian mendatang akan menyelidiki penggunaan model berbasis transformasi (misalnya BERT atau SVM) atau algoritma klasifikasi lainnya dengan tujuan mengatasi keterbatasan Naive Bayes dalam mengidentifikasi variasi bahasa seperti slang dan sarkasme. Untuk meningkatkan kinerja dalam kelas minoritas dengan mengoptimalkan parameter dan teknik oversampling model. Selain itu, penting untuk mempertimbangkan perbaikan kualitas dataset dengan mengintegrasikan data dari berbagai sumber. Terakhir, dengan melibatkan pemangku kepentingan dan masyarakat sejak tahap perumusan kebijakan, transparansi pengambilan keputusan dapat ditingkatkan. Hasil analisis sentimen dapat lebih akurat mencerminkan kebutuhan dan aspirasi masyarakat.

DAFTAR PUSTAKA

- [1] H. A. R. Haprizon, R. Kurniawan, I. Iskandar, R. Salambue, E. Budianita, and F. Syafria, "Analisis Sentimen Komentar Di YouTube Tentang Ceramah Ustadz Abdul Somad Menggunakan Algoritma Naïve Bayes," *Jurnal Nasional Komputasi dan Teknologi Informasi*, vol. 5, no. 1, pp. 131–140, 2022, doi: <https://doi.org/10.32672/jnkti.v5i1.4008>.
- [2] M. T. Sugandi, Martanto, and U. Hayati, "Analisis Sentimen Komentar Pengguna Youtube terhadap Kebijakan Baru Badan Penyelenggara Jaminan Kesehatan Sosial Menggunakan Naïve Bayes," *Jurnal Informatika dan Rekayasa Perangkat Lunak*, vol. 6, no. 1, pp. 218–227, 2024.
- [3] D. F. Zhafira, B. Rahayudi, and Indriati, "ANALISIS SENTIMEN KEBIJAKAN KAMPUS MERDEKA MENGGUNAKAN NAIVE BAYES DAN PEMBOBOTAN TF-IDF BERDASARKAN KOMENTAR PADA YOUTUBE," *Jurnal Sistem Informasi, Teknologi Informasi, dan Edukasi Sistem Informasi (JUST-SI)*, vol. 2, no. 1, pp. 55–63, 2021, doi: <https://doi.org/10.25126/justsi.v2i1.24>.
- [4] R. Aztin and K. Adiyarta, "Penerapan Text Mining Dengan Algoritma Naïve Bayes Untuk Mengklasifikasikan Sentimen Rakyat Terhadap Minyak Goreng Subsidi Pemerintah," in *SENAFTI*, Jakarta, 2022, pp. 645–652. [Online]. Available: <https://senafiti.budiluhur.ac.id/index.php/senafiti>
- [5] Q. Xue, *Data Analytics for Drilling Engineering*. Cham: Springer Nature Switzerland AG, 2020. doi: <https://doi.org/10.1007/978-3-030-34035-3>.
- [6] I. B. Setiawan, J. Maulindar, and Nurchim, "Penerapan Algoritma Naïve Bayes untuk Analisis Sentimen Pada Aplikasi Kesehatan Digital," *G-Tech: Jurnal Teknologi Terapan*, vol. 8, no. 4, pp. 2301–2312, Oct. 2024, doi: <https://doi.org/10.70609/gtech.v8i4.5020>.
- [7] A. Nurian and B. N. Sari, "ANALISIS SENTIMEN ULASAN PENGGUNA APLIKASI GOOGLE PLAY MENGGUNAKAN NAÏVE BAYES," *JITET (Jurnal Informatika dan Teknik Elektro Terapan)*, vol. 11, no. 3, pp. 829–835, 2023, doi: [10.23960/jitet.v11i3%20s1.3348](https://doi.org/10.23960/jitet.v11i3%20s1.3348).
- [8] Kevin, M. Enjeli, and A. Wijaya, "Analisis Sentimen Penggunaan Aplikasi Kinemaster Menggunakan Metode Naive Bayes," *JURNAL ILMIAH COMPUTER SCIENCE (JICS)*, vol. 2, no. 2, pp. 89–98, Jan. 2024, doi: [10.58602/jics.v2i2.24](https://doi.org/10.58602/jics.v2i2.24).
- [9] F. S. Mufidah, S. Winarno, F. Al Zami, E. D. Udayanti, and R. R. Sani, "Analisis Sentimen Masyarakat Terhadap Layanan Shopeefood Melalui Media Sosial Twitter Dengan Algoritma Naïve Bayes Classifier," *JOINS (Journal of Information System)*, vol. 7, no. 1, pp. 14–25, May 2022, doi: [10.33633/joins.v7i1.5883](https://doi.org/10.33633/joins.v7i1.5883).
- [10] F. Kartika Ilmi, S. Shahira Julyinda, E. M. Kiswana, and M. Y. Fathoni, "Analisis Sentimen Emosi Publik Terhadap Kebijakan iPhone 16 Menggunakan Algoritma Naive Bayes," in *Proceedings of the National Conference on Electrical Engineering, Informatics, Industrial Technology, and Creative Media (CENTIVE)*, Purwokerto, 2024, pp. 1141–1153.
- [11] B. P. Salsabila, P. B. C. T. Putri, N. Ramadhani, and A. P. Sari, "PENERAPAN ALGORITMA NAIVE BAYES TERHADAP KUALITAS UDARA DI JAKARTA DAN REKOMENDASI AKTIVITAS MASYARAKAT," *JATI (Jurnal Mahasiswa Teknik Informatika)*, vol. 8, no. 6, pp. 11732–11738, 2024, doi: <https://doi.org/10.36040/jati.v8i6.11592>.
- [12] Nurhasiyah, R. Dwiyanaputra, S. I. Murpratiwi, and A. Aranta, "ANALISIS SENTIMEN PENGGUNA PLATFORM MEDIA SOSIAL X PADA TOPIK PEMILIHAN PRESIDEN 2024 MENGGUNAKAN PERBANDINGAN MODEL MONOLINGUAL DAN MULTILINGUAL BERT," *JATI (Jurnal Mahasiswa Teknik Informatika)*, vol. 9, no. 1, pp. 626–634, 2025, doi: <https://doi.org/10.36040/jati.v9i1.12430>.
- [13] D. Darwis, N. Siskawati, and Z. Abidin, "Penerapan Algoritma Naive Bayes untuk Analisis Sentimen Review Data Twitter BMKG Nasional," *Jurnal TEKNO KOMPAK*, vol. 15, no. 1, pp. 131–145, 2021, doi: <https://doi.org/10.33365/jtk.v15i1.744>.
- [14] V. Fazrian, T. Suprpti, and R. Narasati, "PENERAPAN ALGORITMA NAIVE BAYES TERHADAP ANALISIS SENTIMEN APLIKASI GAME MULTIPLAYER ONLINE BATTLE ARENA (STUDI KASUS : MOBILE LEGEND)," *JATI (Jurnal Mahasiswa Teknik Informatika)*, vol. 8, no. 1, pp. 1005–1012, 2024, doi: <https://doi.org/10.36040/jati.v8i1.8432>.
- [15] D. Alita and R. A. Shodiqin, "Sentimen Analisis Vaksin Covid-19 Menggunakan Naive Bayes Dan Support Vector Machine," *Journal of Artificial Intelligence and Technology Information (JAITI)*, vol. 1, no. 1, pp. 1–12, Feb. 2023, doi: [10.58602/jaiti.v1i1.20](https://doi.org/10.58602/jaiti.v1i1.20).
- [16] M. P. Pulungan, A. Purnomo, and A. Kurniasih, "PENERAPAN SMOTE UNTUK MENGATASI IMBALANCE CLASS DALAM KLASIFIKASI KEPRIBADIAN MBTI MENGGUNAKAN NAIVE BAYES CLASSIFIER," *Jurnal Teknologi Informasi dan Ilmu Komputer (JTIK)*, vol. 11, no. 5, pp. 1033–1042, 2024, doi: [10.25126/jtiik.2024117989](https://doi.org/10.25126/jtiik.2024117989)