



Klasifikasi Persepsi Publik Terhadap Perang Dagang Amerika Serikat Menggunakan Algoritma Naïve Bayes Classifier

Bunga Nurul Manisa*, Aidil Halim Lubis

Fakultas Sains dan Teknologi, Program Studi Ilmu Komputer, Universitas Islam Negeri Sumatera Utara, Medan, Indonesia

Email: ^{1,*}bnurulmanisa@gmail.com, ²aidilhalimlubis@uinsu.ac.id

Email Penulis Korespodensi: bnurulmanisa@gmail.com

Abstrak—Kebijakan tarif impor yang diterapkan oleh Presiden Amerika Serikat pada 2 April 2025 memicu ketegangan dalam perdagangan global serta menimbulkan berbagai reaksi dari masyarakat. Perbedaan persepsi publik terhadap kebijakan tersebut menghasilkan opini yang beragam, baik berupa dukungan, kritik, maupun sikap netral, sehingga diperlukan analisis sentimen untuk memahami kecenderungan opini masyarakat secara lebih terukur. Penelitian ini bertujuan untuk mengklasifikasikan persepsi publik terhadap perang dagang Amerika Serikat melalui analisis sentimen berbasis data Twitter dengan menggunakan algoritma *Naïve Bayes Classifier* (NBC). Dataset yang digunakan terdiri dari 2.000 tweet yang dikumpulkan menggunakan kata kunci "perang dagang" dan "kenaikan tarif impor" dalam periode 3 hingga 30 April 2025. Tahapan preprocessing dilakukan melalui enam proses, yaitu cleaning, case folding, tokenizing, slangword normalization, stopword removal, dan stemming untuk meningkatkan kualitas data teks. Pelabelan data dilakukan secara otomatis menggunakan metode lexicon-based dengan kamus InSet, menghasilkan distribusi sentimen sebesar 83,5% negatif, 12,8% positif, dan 3,8% netral. Representasi fitur menggunakan metode TF-IDF, kemudian data dibagi menjadi data latih dan data uji dengan perbandingan 80:20. Untuk mengatasi ketidakseimbangan distribusi kelas, diterapkan metode *Synthetic Minority Over-sampling Technique* (SMOTE). Hasil pengujian menunjukkan bahwa model NBC tanpa SMOTE memperoleh akurasi sebesar 83,5%, namun model cenderung bias terhadap kelas mayoritas. Setelah penerapan SMOTE, distribusi data menjadi seimbang dengan masing-masing kelas sebanyak 1.335 data. Meskipun akurasi menurun menjadi 76%, nilai Macro F1-Score meningkat dari 0,30 menjadi 0,45, yang menunjukkan peningkatan kemampuan model dalam menangani klasifikasi multi-kelas secara lebih adil. Model juga mampu mendeteksi kelas positif dengan recall 43% dan kelas netral sebesar 13%, sehingga memberikan gambaran persepsi publik yang lebih representatif terhadap isu perang dagang Amerika Serikat.

Kata Kunci: Naïve Bayes Classifier; Analisis Sentimen; Perang Dagang; SMOTE; TF-IDF

Abstract—The import tariff policy implemented by the President of the United States on April 2, 2025 triggered tensions in global trade and provoked various public reactions. Differences in public perceptions of the policy generated diverse opinions, including support, criticism, and neutral responses, making sentiment analysis necessary to understand public opinion trends more systematically. This study aims to classify public perceptions of the U.S. trade war through sentiment analysis of Twitter data using the *Naïve Bayes Classifier* (NBC) algorithm. The dataset consists of 2,000 tweets collected using the keywords "trade war" and "import tariff increase" during April 3–30, 2025. Six preprocessing stages were applied: cleaning, case folding, tokenizing, slangword normalization, stopword removal, and stemming to improve data quality and consistency. Automatic labeling was conducted using a lexicon-based method with the InSet dictionary, yielding sentiment distributions of 83.5% negative, 12.8% positive, and 3.8% neutral. Feature representation was performed using TF-IDF, followed by an 80:20 train-test split. To address class imbalance, the *Synthetic Minority Over-sampling Technique* (SMOTE) was applied. Experimental results show that the NBC model without SMOTE achieved an accuracy of 83.5% but exhibited bias toward the majority class. After applying SMOTE, the dataset became balanced with 1,335 samples per class. Although overall accuracy decreased to 76%, the Macro F1-Score improved from 0.30 to 0.45, indicating improved model performance in handling multi-class classification more fairly. Additionally, the model achieved a recall of 43% for the positive class and 13% for the neutral class, providing a more representative evaluation of public sentiment toward the U.S. trade war issue.

Keywords: Naïve Bayes Classifier; Sentiment Analysis; Trade War; SMOTE; TF-IDF

1. PENDAHULUAN

Perkembangan teknologi informasi yang pesat telah mendorong masyarakat untuk semakin aktif dalam menyampaikan pendapat melalui media sosial. Platform X (Twitter) menjadi salah satu media utama yang digunakan publik untuk mengekspresikan opini, kritik, maupun respons terhadap berbagai isu yang berkembang secara global. Kehadiran media sosial tidak hanya berfungsi sebagai sarana komunikasi, tetapi juga sebagai sumber data yang sangat besar dan dinamis untuk memahami persepsi publik secara *real-time*. Salah satu isu global yang memicu perhatian luas masyarakat adalah kebijakan tarif impor Amerika Serikat yang diterapkan pada 2 April 2025, yang menimbulkan berbagai reaksi dari masyarakat di berbagai negara [1].

Pertumbuhan ekonomi suatu negara sangat dipengaruhi oleh aktivitas perdagangan internasional. Melalui perdagangan internasional, negara dapat memperluas pasar, memperoleh akses terhadap teknologi yang lebih maju, serta meningkatkan investasi asing. Namun, kebijakan proteksionisme seperti kenaikan tarif impor dapat menghambat arus perdagangan global dan menciptakan ketidakseimbangan ekonomi. Presiden Amerika Serikat, Donald Trump, menerapkan kebijakan kenaikan tarif impor secara signifikan dengan tujuan mengurangi ketergantungan terhadap barang impor serta mendorong peningkatan produksi domestik. Kebijakan ini memicu ketegangan dalam perdagangan global dan meningkatkan potensi terjadinya perang dagang akibat adanya kebijakan balasan dari negara mitra dagang [2].

Dampak dari kebijakan tersebut tidak hanya dirasakan oleh negara-negara besar, tetapi juga memengaruhi negara berkembang seperti Indonesia. Berdasarkan data Badan Pusat Statistik (BPS), terjadi penurunan total nilai ekspor Indonesia dari Maret ke April 2025 sebesar 10,77% [3]. Penurunan ini mencakup sektor migas sebesar 19,52% dan sektor nonmigas sebesar 10,19%. Selain itu, ketidakpastian ekonomi global akibat kebijakan tersebut juga berdampak pada



melemahnya nilai tukar rupiah terhadap dolar Amerika Serikat. Kondisi ini menyebabkan kenaikan harga barang impor dan memberikan tekanan terhadap inflasi domestik. Hal ini menunjukkan bahwa kebijakan ekonomi suatu negara dapat memberikan dampak yang luas terhadap stabilitas ekonomi global maupun nasional [4].

Dalam situasi tersebut, penting untuk memahami bagaimana persepsi publik terbentuk terhadap kebijakan perang dagang Amerika Serikat. Persepsi publik merupakan respon kognitif dan emosional masyarakat terhadap suatu fenomena yang dipengaruhi oleh berbagai faktor, seperti informasi yang diterima, pengalaman, dan kondisi sosial ekonomi. Oleh karena itu, analisis terhadap opini publik menjadi penting untuk memberikan gambaran yang objektif mengenai respons masyarakat. Salah satu pendekatan yang dapat digunakan adalah analisis sentimen berbasis data media sosial, yang memungkinkan pengolahan data teks dalam jumlah besar secara otomatis untuk mengidentifikasi kecenderungan opini publik [5].

Analisis sentimen merupakan bagian dari text mining yang bertujuan untuk mengklasifikasikan teks ke dalam kategori sentimen tertentu, seperti positif, negatif, atau netral. Dalam penelitian ini, algoritma *Naïve Bayes Classifier* (NBC) digunakan sebagai metode klasifikasi karena memiliki keunggulan dalam hal kesederhanaan, kecepatan proses, serta kemampuan dalam menangani data teks berdimensi tinggi. Selain itu, NBC juga dikenal memiliki performa yang cukup baik dalam berbagai penelitian sebelumnya, sehingga menjadi salah satu metode yang banyak digunakan dalam analisis sentimen [6][7].

Beberapa penelitian terdahulu telah membahas penerapan analisis sentimen menggunakan algoritma *Naïve Bayes Classifier* (NBC) pada berbagai kasus. Penelitian yang dilakukan oleh Pratama dkk. menggunakan algoritma NBC untuk analisis sentimen ulasan aplikasi layanan digital dan memperoleh tingkat akurasi yang cukup tinggi dalam mengklasifikasikan opini pengguna. Penelitian tersebut menunjukkan bahwa NBC memiliki kemampuan yang baik dalam pengolahan data teks dengan proses komputasi yang relatif sederhana [8]. Selanjutnya, penelitian oleh Ramadhan dan Putri menerapkan metode NBC pada analisis sentimen media sosial terkait kebijakan ekonomi pemerintah dan menghasilkan performa klasifikasi yang stabil pada data teks berbahasa Indonesia [9]. Penelitian lain oleh Sari dkk. mengombinasikan TF-IDF dan NBC dalam analisis sentimen Twitter terkait isu publik, yang menunjukkan bahwa representasi fitur TF-IDF mampu meningkatkan kualitas klasifikasi sentiment [10]. Selain itu, penelitian oleh Wijaya dan Nugroho menerapkan metode SMOTE untuk mengatasi ketidakseimbangan data pada klasifikasi sentimen dan berhasil meningkatkan nilai recall pada kelas minoritas secara signifikan [11]. Meskipun penelitian-penelitian sebelumnya telah menunjukkan efektivitas algoritma NBC dan metode SMOTE dalam analisis sentimen, masih terdapat beberapa kesenjangan penelitian (research gap). Pertama, sebagian besar penelitian terdahulu berfokus pada topik ulasan aplikasi, layanan digital, atau kebijakan lokal, sedangkan penelitian terkait persepsi publik terhadap perang dagang Amerika Serikat masih sangat terbatas. Kedua, penelitian sebelumnya umumnya hanya berfokus pada peningkatan akurasi model tanpa menganalisis secara mendalam performa klasifikasi pada data multi-kelas yang tidak seimbang. Ketiga, sebagian penelitian masih menggunakan pelabelan manual yang membutuhkan waktu dan tenaga lebih besar, sedangkan penelitian ini menggunakan pendekatan lexicon-based dengan kamus InSet untuk melakukan pelabelan data secara otomatis. Selain itu, penelitian ini juga mengombinasikan preprocessing yang lebih lengkap dengan penerapan SMOTE untuk meningkatkan kemampuan model dalam mendeteksi kelas minoritas, sehingga diharapkan mampu menghasilkan klasifikasi sentimen yang lebih representatif terhadap opini publik mengenai isu perang dagang Amerika Serikat [12][13].

Namun, dalam implementasinya, analisis sentimen sering menghadapi permasalahan ketidakseimbangan distribusi kelas (*class imbalance*). Kondisi ini terjadi karena jumlah data pada satu kelas lebih dominan dibandingkan kelas lainnya, sehingga model cenderung bias terhadap kelas mayoritas dan mengabaikan kelas minoritas. Untuk mengatasi permasalahan tersebut, digunakan metode *Synthetic Minority Over-sampling Technique* (SMOTE), yang bekerja dengan menghasilkan data sintetis pada kelas minoritas melalui proses interpolasi antar data yang ada. Dengan demikian, distribusi data menjadi lebih seimbang dan model dapat belajar secara lebih optimal [14][15].

Penelitian ini menggunakan *dataset* sebanyak 2.000 tweet yang dikumpulkan dari platform X dengan kata kunci “perang dagang” dan “kenaikan tarif impor”. Data tersebut kemudian diproses melalui enam tahapan *preprocessing*, yaitu *cleaning*, *case folding*, *tokenizing*, *slangword normalization*, *stopword removal*, dan *stemming*, untuk meningkatkan kualitas data sebelum dilakukan analisis lebih lanjut. Selain itu, pelabelan data dilakukan secara otomatis menggunakan metode *lexicon-based* dengan kamus *InSet*, sehingga proses pelabelan dapat dilakukan secara efisien tanpa memerlukan anotasi manual [5][16].

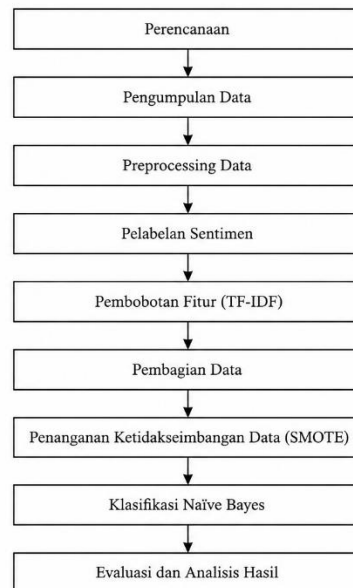
Berdasarkan latar belakang tersebut, penelitian ini bertujuan untuk mengidentifikasi distribusi persepsi publik terhadap perang dagang Amerika Serikat berdasarkan data Twitter, serta mengevaluasi performa algoritma *Naïve Bayes Classifier* dalam melakukan klasifikasi sentimen. Selain itu, penelitian ini juga bertujuan untuk menganalisis pengaruh penerapan metode SMOTE dalam meningkatkan kualitas klasifikasi pada data yang tidak seimbang. Diharapkan hasil penelitian ini dapat memberikan kontribusi dalam pengembangan analisis sentimen serta menjadi referensi bagi pembuat kebijakan dalam memahami dinamika opini publik terhadap isu ekonomi global.

2. METODOLOGI PENELITIAN

2.1 Tahapan Penelitian

Penelitian ini dilakukan melalui beberapa tahapan utama yang ringkas dan terstruktur. Tahap pertama adalah Penelitian ini dilaksanakan melalui beberapa tahapan yang disusun secara sistematis agar setiap proses dapat berjalan secara

terstruktur dan sesuai dengan tujuan penelitian. Tahapan penelitian dimulai dari identifikasi permasalahan hingga evaluasi hasil klasifikasi. Setiap tahapan saling berkaitan dan menjadi dasar bagi tahapan berikutnya sehingga menghasilkan proses analisis sentimen yang terarah dan dapat dipertanggungjawabkan. Kerangka penelitian yang digunakan pada penelitian ini ditunjukkan pada Gambar 1.



Gambar 1. Tahapan Penelitian

Berdasarkan Gambar 1, penelitian ini terdiri atas sembilan tahapan sebagai berikut.

a. Perencanaan

Tahap awal dilakukan dengan mengidentifikasi permasalahan penelitian, menentukan tujuan penelitian, melakukan studi literatur, serta menyusun metode yang digunakan dalam penelitian. Tahap ini menjadi dasar dalam menentukan arah penelitian yang akan dilakukan [17].

b. Pengumpulan Data

Tahap ini dilakukan dengan mengumpulkan data dari platform X (Twitter) menggunakan teknik crawling melalui library *tweet-harvest*. Pengambilan data menggunakan kata kunci "perang dagang" dan "kenaikan tarif impor" pada periode 3–30 April 2025 sehingga diperoleh sebanyak 2.000 tweet yang relevan dengan topik penelitian.

c. Preprocessing Data

Data hasil crawling kemudian dibersihkan agar siap diproses lebih lanjut. Tahapan preprocessing meliputi *cleaning*, *case folding*, *tokenizing*, *slangword normalization*, *stopword removal*, dan *stemming* sehingga diperoleh data teks yang lebih bersih, konsisten, dan representatif.

d. Pelabelan Sentimen

Data yang telah diproses selanjutnya diberi label sentimen menggunakan pendekatan lexicon-based dengan kamus InSet. Berdasarkan skor sentimen yang diperoleh, setiap tweet diklasifikasikan ke dalam kategori positif, negatif, atau netral.

e. Pembobotan Fitur (TF-IDF)

Tahap berikutnya adalah mengubah data teks menjadi bentuk numerik menggunakan metode *Term Frequency -Inverse Document Frequency* (TF-IDF). Pembobotan ini bertujuan menghasilkan representasi fitur yang dapat diproses oleh algoritma klasifikasi [18].

f. Pembagian Data

Dataset kemudian dibagi menjadi data latih (80%) dan data uji (20%) menggunakan metode Stratified Split agar proporsi masing-masing kelas tetap terjaga pada kedua kelompok data.

g. Penanganan Ketidakseimbangan Data (SMOTE)

Karena distribusi data sentimen tidak seimbang, diterapkan metode *Synthetic Minority Over-sampling Technique* (SMOTE) pada data latih untuk menghasilkan data sintesis pada kelas minoritas sehingga distribusi data menjadi lebih seimbang [19].

h. Klasifikasi Naïve Bayes

Data yang telah melalui proses pembobotan dan penyeimbangan kelas kemudian diklasifikasikan menggunakan algoritma *Naïve Bayes Classifier* (NBC). Model dibangun menggunakan data latih dan selanjutnya digunakan untuk memprediksi sentimen pada data uji.



i. Evaluasi dan Analisis Hasil

Tahap terakhir dilakukan evaluasi terhadap performa model menggunakan *Confusion Matrix* dengan metrik *Accuracy*, *Precision*, *Recall*, dan *F1-Score*. Selain itu dilakukan analisis terhadap distribusi sentimen dan pengaruh penerapan SMOTE terhadap peningkatan performa model klasifikasi [20].

3. HASIL DAN PEMBAHASAN

3.1 Hasil Pengumpulan Data

Data penelitian bersumber dari tweet masyarakat Indonesia di *platform X* yang membahas kebijakan kenaikan tarif impor oleh Presiden Amerika Serikat tahun 2025. Proses *crawling* dilakukan menggunakan bahasa pemrograman Python dengan bantuan Node.js dan library *tweet-harvest* untuk mengambil data langsung dari *platform X*. Data yang diambil merupakan tweet dari masyarakat dalam kurun waktu satu bulan, yaitu dari 3 April 2025 sampai 30 April 2025 dengan kata kunci "perang dagang dan kenaikan tarif impor" serta limit 2.000 data. Hasil *crawling* menghasilkan *dataset* dengan 15 kolom, namun hanya kolom *username* dan *full_text* yang digunakan sebagai bahan analisis, sehingga kolom lainnya dibersihkan pada tahap *preprocessing*.

3.2 Hasil Preprocessing

Proses *preprocessing* meliputi pengubahan *dataset* menjadi format yang lebih teratur sehingga dapat diproses oleh sistem pada tahap berikutnya. Pada tahap ini data mentah yang diperoleh dari hasil *crawling* akan diolah agar menjadi data yang bersih dan lebih konsisten untuk digunakan pada tahap selanjutnya. Tahapan mencakup: *cleaning*, *case folding*, *tokenizing*, *slangword*, *stopword*, serta *stemming*.

Pada Tabel 1 ditampilkan sampel hasil proses *cleaning* yang bertujuan menghapus karakter yang tidak diperlukan seperti mention, simbol, tanda baca, URL, dan karakter khusus lainnya. Proses ini dilakukan agar teks menjadi lebih bersih dan fokus pada informasi utama yang akan dianalisis.

Tabel 1. Sampel Hasil Proses *Cleaning*

Text	Cleaning
@DokterTifa Dunia lagi perang dagang	Dunia lagi perang dagang
@BosPurwa Salahkan perang Russia-Ukraine salahkan global salahkan Donald Trump salahkan dinamika politik salahkan perang dagang fix ya wanita emang tidak mau disalahkan!,	Salahkan perang Russia Ukraine salahkan global salahkan Donald Trump salahkan dinamika politik salahkan perang dagang fix ya wanita emang tidak mau disalahkan
@ltoedjoean Pemerintah dengan cermat memitigasi dampak perang dagang dengan mendorong konsumsi domestik dan meningkatkan ekspor produk bernilai tambah	Pemerintah dengan cermat memitigasi dampak perang dagang dengan mendorong konsumsi domestik dan meningkatkan ekspor produk bernilai tambah
@NarasiNewsroom Kemungkinan bagian dari perang dagang...memecah belah	Kemungkinan bagian dari perang dagang memecah belah

Pada Tabel 2 ditampilkan hasil proses *case folding*, yaitu mengubah seluruh huruf menjadi huruf kecil (lowercase). Tahapan ini bertujuan untuk menyeragamkan format teks sehingga kata yang sama tidak dianggap berbeda hanya karena perbedaan penggunaan huruf kapital.

Tabel 2. Sampel Hasil Proses *Case Folding*

Text	Case Folding
Dunia lagi perang dagang	dunia lagi perang dagang
Salahkan perang Russia Ukraine salahkan global salahkan Donald Trump salahkan dinamika politik salahkan perang dagang fix ya wanita emang tidak mau disalahkan	salahkan perang russia ukraine salahkan global salahkan donald trump salahkan dinamika politik salahkan perang dagang fix ya wanita emang tidak mau disalahkan
Pemerintah dengan cermat memitigasi dampak perang dagang dengan mendorong konsumsi domestik	pemerintah dengan cermat memitigasi dampak perang dagang dengan mendorong konsumsi domestik

Selanjutnya, Tabel 3 menunjukkan hasil proses *tokenizing*, yaitu memecah kalimat menjadi kumpulan token atau kata-kata tunggal. Proses tokenisasi memudahkan sistem dalam melakukan analisis terhadap setiap kata yang terkandung dalam teks.

Tabel 3. Sampel Hasil Proses *Tokenize*

Text	Tokenize
dunia lagi perang dagang	['dunia', 'lagi', 'perang', 'dagang']



Text	Tokenize
pemerintah dengan cermat memitigasi dampak perang dagang dengan mendorong konsumsi domestik	['pemerintah', 'dengan', 'cermat', 'memitigasi', 'dampak', 'perang', 'dagang', 'dengan', 'mendorong', 'konsumsi', 'domestik', ...]

Pada Tabel 4 ditampilkan hasil proses stopwords removal yang bertujuan menghapus kata-kata umum yang tidak memiliki pengaruh signifikan terhadap proses klasifikasi sentimen, seperti “dan”, “yang”, atau “dengan”. Setelah proses ini dilakukan, teks menjadi lebih fokus pada kata-kata penting yang memiliki makna sentimen.

Tabel 4. Sampel Hasil Proses *Stopword*

Text	Stopword
['dunia', 'lagi', 'perang', 'dagang']	['dunia', 'perang', 'dagang']
['pemerintah', 'dengan', 'cermat', 'memitigasi', 'dampak', 'perang', 'dagang', 'dengan', 'mendorong', 'konsumsi', 'domestik', 'dan', 'meningkatkan', 'ekspor', 'produk', 'bernilai', 'tambah']	['pemerintah', 'cermat', 'memitigasi', 'dampak', 'perang', 'dagang', 'mendorong', 'konsumsi', 'domestik', 'meningkatkan', 'ekspor', 'produk', 'bernilai']
['perang', 'dagang', 'memecah', 'belah']	['perang', 'dagang', 'memecah', 'belah']

Kemudian, Tabel 5 memperlihatkan hasil proses slangword normalization, yaitu proses mengubah kata tidak baku atau bahasa sehari-hari menjadi bentuk kata baku. Contohnya kata “emang” diubah menjadi “memang”. Tahapan ini bertujuan meningkatkan konsistensi data teks agar lebih mudah dipahami oleh sistem klasifikasi.

Tabel 5. Sampel Hasil Proses *Slangword*

Text	Slangword
['salahkan', 'perang', 'russia', 'ukraine', 'salahkan', 'global', 'salahkan', 'donald', 'trump', 'fix', 'wanita', 'emang', 'disalahkan']	['salahkan', 'perang', 'russia', 'ukraine', 'salahkan', 'global', 'salahkan', 'donald', 'trump', 'fix', 'wanita', 'memang', 'disalahkan']
['perang', 'dagang', 'memecah', 'belah']	['perang', 'dagang', 'pecah', 'belah']
['perang', 'dagang', 'ditambah', 'suhu', 'geopolitik', 'bikin', 'kebijakan', 'akibatnya', 'barang', 'mahal']	['perang', 'dagang', 'tambah', 'suhu', 'geopolitik', 'bikin', 'kebijakan', 'akibat', 'barang', 'mahal']

Selanjutnya, Tabel 6 menunjukkan hasil stemming, yaitu proses mengubah kata berimbuhan menjadi kata dasar. Sebagai contoh, kata “meningkatkan” diubah menjadi “tingkat” dan “mendorong” menjadi “dorong”. Proses stemming dilakukan untuk mengurangi variasi kata sehingga dapat meningkatkan efektivitas proses klasifikasi sentimen.

Tabel 6. Sampel Hasil Proses *Stemming*

Text	Stemming
['dunia', 'perang', 'dagang']	dunia perang dagang
['salahkan', 'perang', 'russia', 'ukraine', 'salahkan', 'global', 'salahkan', 'donald', 'trump', 'fix', 'wanita', 'memang', 'disalahkan']	salah perang russia ukraine salah global salah donald trump fix wanita memang salah
['pemerintah', 'cermat', 'memitigasi', 'dampak', 'perang', 'dagang', 'mendorong', 'konsumsi', 'domestik', 'meningkatkan', 'ekspor', 'produk']	pemerintah cermat mitigasi dampak perang dagang dorong konsumsi domestik tingkat ekspor produk
['perang', 'dagang', 'pecah', 'belah']	perang dagang pecah belah

Seluruh data yang telah melalui *preprocessing* lengkap dikompilasi dan disimpan dalam format CSV untuk memudahkan penyimpanan, pengelolaan, serta analisis pada tahap selanjutnya.

3.3 Hasil Pemodelan

Tahapan pemodelan diawali dengan proses pelabelan *dataset* guna menentukan kelas sentimen pada data. Pelabelan dilakukan menggunakan pendekatan *lexicon-based* dengan kamus *InSet* yang terdiri dari 3.069 kata positif dan 6.609 kata negatif dengan pembobotan dari -5 hingga +5. Jika total skor > 0 maka teks dikategorikan Positif; jika < 0 maka Negatif; dan jika = 0 maka Netral.

Pada Tabel 7 ditampilkan sampel kamus *InSet* yang digunakan dalam proses pelabelan otomatis. Kamus tersebut berisi daftar kata positif dan negatif beserta nilai bobot masing-masing kata. Bobot ini digunakan untuk menentukan kecenderungan sentimen pada setiap tweet berdasarkan jumlah skor total yang diperoleh.

Tabel 7. Sampel Kamus *InSet*

Kamus Positif	Bobot	Kamus Negatif	Bobot
hai	+3	sesal	-4
terimakasih	+5	pengen	-2
maaf	+2	sakit	-5



Kamus Positif	Bobot	Kamus Negatif	Bobot
hello	+2	malas	-4
mohon	+2	katanya	-1

Selanjutnya, Tabel 8 memperlihatkan hasil pelabelan data berdasarkan perhitungan bobot sentimen dari setiap kata. Total bobot yang dihasilkan menentukan apakah suatu data termasuk kategori positif, negatif, atau netral. Proses ini dilakukan secara otomatis menggunakan pendekatan lexicon-based sehingga lebih efisien dibandingkan pelabelan manual.

Tabel 8. Sampel Hasil Pelabelan Data

Hasil <i>Preprocessing</i>	Bobot per Kata	Total Bobot	Label Sentimen
dunia perang dagang	[1, -3, 0]	-2	Negatif
pemerintah cermat mitigasi dampak perang dagang dorong konsumsi domestik tingkat ekspor produk	[0, 0, 0, -3, -3, 0, 0, 0, 3, 1, 3]	+1	Positif
dampak tidak main perang dagang harga barang tani as rugi industri ganggu konsumen korban	[-3, -5, 2, -3, 0, 3, 0, 0, 0, -5, 0, -4, 0, -4]	-19	Negatif
imbas perang dagang lesu harga minyak sebab faktor temu celah korupsi	[-4, -3, 0, -4, 3, 0, 3, -4, 2, -3, -4]	-11	Negatif

Pada Tabel 9 ditampilkan distribusi label sentimen dari keseluruhan dataset sebanyak 2.000 data. Hasil distribusi menunjukkan bahwa sentimen negatif mendominasi dengan persentase sebesar 83,5%, sedangkan sentimen positif sebesar 12,8% dan netral sebesar 3,8%. Dominasi sentimen negatif menunjukkan bahwa sebagian besar masyarakat memberikan respons negatif terhadap kebijakan perang dagang Amerika Serikat.

Tabel 9. Distribusi Label Sentimen dari 2.000 Data

Label Sentimen	Jumlah Data	Persentase (%)
Negatif	1.669	83,5%
Positif	256	12,8%
Netral	75	3,8%
Total	2.000	100%

Hasil pelabelan menunjukkan ketidakseimbangan kelas yang sangat signifikan dengan sentimen negatif mendominasi 83,5%. Dominasi sentimen negatif ini mencerminkan kekhawatiran publik yang kuat terhadap dampak kebijakan tarif impor AS, yang sejalan dengan data penurunan ekspor Indonesia sebesar 10,77% pada periode yang sama.

Sebelum pembobotan TF-IDF, data yang telah dilabeli dibagi terlebih dahulu dengan rasio 80:20 menggunakan stratified split, menghasilkan 1.600 data latih dan 400 data uji. Alasan pembagian data dilakukan setelah pelabelan adalah untuk menjaga keseimbangan distribusi label antara data latih dan data uji.

Pada Tabel 10 ditampilkan sampel data latih yang digunakan dalam proses pembentukan model klasifikasi. Data tersebut terdiri dari hasil preprocessing yang telah diberi label sentimen sehingga dapat digunakan sebagai data pembelajaran model NBC.

Tabel 10. Sampel Data Latih

Sentimen Latih	Kelas
['tegang', 'perang', 'dagang', 'sebab', 'lemah', 'ekonomi', 'dunia', 'tekan', 'inflasi', 'global']	Positif
['jujur', 'seru', 'ngulik', 'perang', 'dagang']	Positif
['prabowo', 'waspada', 'perang', 'dagang', 'trump', 'picu', 'phk', 'massal', 'indonesia']	Netral
['perang', 'dagang', 'ancam', 'resesi']	Negatif
['orang', 'negara', 'pusing', 'pikir', 'perang', 'dagang', 'global', 'orang', 'ngeributin', 'ijazah', 'mantan', 'presiden', 'parah']	Negatif

Selanjutnya, Tabel 11 menunjukkan nilai Document Frequency (DF) dari beberapa term pada data latih. Nilai DF digunakan dalam proses pembobotan TF-IDF untuk mengetahui jumlah dokumen yang mengandung kata tertentu. Semakin tinggi nilai DF suatu kata, maka semakin sering kata tersebut muncul dalam dokumen.

Tabel 11. Nilai DF dari Sampel Data Latih

Term	D1	D2	D3	D4	D5	DF
tegang	1	0	0	0	0	1
perang	1	1	1	1	1	5
dagang	1	1	1	1	1	5
global	1	0	0	0	1	2
ancam	0	0	0	1	0	1
resesi	0	0	0	1	0	1



Term	D1	D2	D3	D4	D5	DF
orang	0	0	0	0	2	2

3.4 Hasil Klasifikasi *Naïve Bayes Classifier* Tanpa SMOTE

Klasifikasi menggunakan model MultinomialNB dengan $\alpha=1,0$. Setelah proses pelatihan model NBC selesai, tahap selanjutnya adalah evaluasi model untuk mengukur performa klasifikasi sentimen menggunakan data uji yang telah dipisahkan sebelumnya. Berdasarkan hasil evaluasi, model cenderung mengklasifikasikan seluruh data ke dalam kelas negatif. Dari total 400 data uji, sebanyak 334 data negatif berhasil diprediksi dengan benar (TP = 334). Namun, terdapat 66 data non-negatif (15 netral dan 51 positif) yang salah diklasifikasikan sebagai negatif (FP = 66). Sementara itu, seluruh data netral dan positif tidak berhasil dikenali dengan benar, sehingga FN = 0 dan TN = 0.

Pada Tabel 12 ditampilkan confusion matrix hasil klasifikasi NBC tanpa SMOTE. Tabel tersebut menunjukkan bahwa model cenderung memprediksi seluruh data ke dalam kelas negatif. Sebanyak 334 data negatif berhasil diprediksi dengan benar, sedangkan seluruh data netral dan positif gagal dikenali oleh model.

Tabel 12. Confusion Matrix NBC Tanpa SMOTE

Aktual / Prediksi	Negatif	Netral	Positif
Negatif	334	0	0
Netral	15	0	0
Positif	51	0	0

Selanjutnya, Tabel 13 menunjukkan nilai precision, recall, f1-score, dan accuracy dari model NBC tanpa SMOTE. Hasil evaluasi memperlihatkan bahwa model memperoleh akurasi sebesar 83,5%, namun nilai tersebut dipengaruhi oleh dominasi kelas negatif. Nilai recall untuk kelas netral dan positif sebesar 0% menunjukkan bahwa model belum mampu mengenali kelas minoritas secara baik.

Tabel 13. Nilai Precision dan Recall NBC Tanpa SMOTE

Kelas	Precision	Recall	F1-Score	Support
Negatif	0,835 (83,5%)	1,000 (100%)	0,91	334
Netral	0,000 (0%)	0,000 (0%)	0,00	15
Positif	0,000 (0%)	0,000 (0%)	0,00	51
Accuracy	-	-	83,5%	400
Macro Avg	0,28	0,33	0,30	400
Weighted Avg	0,70	0,83	0,76	400

Hasil evaluasi menunjukkan bahwa model memiliki bias yang sangat kuat terhadap kelas negatif dan belum mampu membedakan kelas netral maupun positif. Meskipun pengujian menghasilkan akurasi sebesar 83,5%, nilai tersebut sebagian besar disebabkan oleh dominasi data pada kelas negatif serta kecenderungan model yang selalu memprediksi ke kelas tersebut. Dengan demikian, akurasi yang tinggi tidak mencerminkan kemampuan model dalam mengklasifikasikan semua kelas secara seimbang. Mengingat kinerja yang kurang memuaskan pada kelas netral dan positif, kemudian diterapkan metode SMOTE untuk menyeimbangkan distribusi data sehingga diharapkan model dapat meningkatkan kemampuan klasifikasi pada seluruh kelas.

3.5 Penerapan SMOTE dan Evaluasi

Distribusi data latih sebelum SMOTE menunjukkan ketidakseimbangan yang signifikan: kelas negatif mencapai 83,44% (1.335 data), sementara kelas positif dan netral hanya masing-masing 12,81% (205 data) dan 3,75% (60 data). Kondisi ini menyebabkan model cenderung bias ke kelas mayoritas. SMOTE diterapkan dengan $k=5$ untuk membangkitkan data sintesis pada kelas minoritas melalui interpolasi: $x_{new} = x_i + \lambda \times (x_{nn} - x_i)$. Kelas positif ditambahkan sebanyak 1.130 data sintesis (1.335 – 205), sedangkan kelas netral bertambah sebanyak 1.275 data sintesis (1.335 – 60). Sementara itu, kelas negatif tidak mengalami penambahan karena sudah memiliki jumlah data terbanyak.

Pada Tabel 14 ditampilkan distribusi data sebelum dan sesudah penerapan SMOTE. Hasil tersebut menunjukkan bahwa setelah proses oversampling, jumlah data pada masing-masing kelas menjadi seimbang dengan total 1.335 data untuk setiap kelas sentimen.

Tabel 14. Distribusi Data Sebelum dan Sesudah SMOTE

Kelas	Sebelum SMOTE	Data Sintetis	Sesudah SMOTE	Persentase
Negatif	1.335	0	1.335	33,33%
Positif	205	+1.130	1.335	33,33%
Netral	60	+1.275	1.335	33,33%
Total	1.600	+2.405	4.005	100%



Setelah diterapkan SMOTE, distribusi data menjadi seimbang dengan jumlah yang sama untuk setiap kelas, yaitu 1.335 data masing-masing sebesar 33,33%. Hasil pengujian setelah SMOTE menunjukkan bahwa dari 400 data uji, sebanyak 304 data berhasil diklasifikasikan dengan benar (*True Positive*), sedangkan 96 data lainnya mengalami kesalahan klasifikasi.

Selanjutnya, Tabel 15 menunjukkan confusion matrix hasil klasifikasi NBC setelah penerapan SMOTE. Berdasarkan tabel tersebut, model mulai mampu mengenali kelas positif dan netral meskipun masih terdapat beberapa kesalahan klasifikasi pada masing-masing kelas.

Tabel 15. *Confusion Matrix* NBC Setelah SMOTE

Aktual / Prediksi	Negatif	Netral	Positif
Negatif	280	19	35
Netral	5	2	8
Positif	24	5	22

Pada Tabel 16 ditampilkan hasil evaluasi precision, recall, dan f1-score setelah penerapan SMOTE. Hasil tersebut menunjukkan peningkatan performa model pada kelas minoritas, terutama pada kelas positif yang memperoleh recall sebesar 43% dan kelas netral sebesar 13%.

Tabel 16. Nilai *Precision* dan *Recall* NBC Setelah SMOTE

Kelas	Precision	Recall	F1-Score	Support
Negatif	0,906 (91%)	0,838 (84%)	0,871 (0,80)	334
Netral	0,077 (8%)	0,133 (13%)	0,097 (0,10)	15
Positif	0,338 (34%)	0,431 (43%)	0,379 (0,38)	51
Accuracy	-	-	76,0%	400
Macro Avg	0,44	0,47	0,45	400
Weighted Avg	0,80	0,76	0,78	400

Nilai *macro average* dan *weighted average* digunakan untuk memberikan gambaran kinerja model klasifikasi pada seluruh kelas. *Macro average* menghitung rata-rata nilai *precision*, *recall*, dan *f1-score* dari setiap kelas tanpa memperhatikan jumlah data pada masing-masing kelas. Berdasarkan hasil pengujian setelah penerapan SMOTE, nilai *macro precision*, *macro recall*, dan *macro f1-score* berada pada kisaran 0,44–0,47, yang mengindikasikan bahwa performa model belum merata di semua kelas, khususnya pada kelas netral yang memiliki performa rendah. Sementara itu, nilai *weighted precision*, *weighted recall*, dan *weighted f1-score* berada pada kisaran 0,76–0,80, yang menunjukkan bahwa secara keseluruhan model NBC memiliki kinerja yang cukup baik setelah proses SMOTE.

3.6 Perbandingan Performa Sebelum dan Sesudah SMOTE

Pada Tabel 17 ditampilkan perbandingan performa model NBC sebelum dan sesudah penerapan SMOTE. Hasil tersebut menunjukkan bahwa meskipun akurasi menurun dari 83,5% menjadi 76,0%, nilai Macro F1-Score meningkat dari 0,30 menjadi 0,45. Peningkatan tersebut menunjukkan bahwa model menjadi lebih baik dalam mengenali seluruh kelas sentimen secara lebih seimbang setelah penerapan SMOTE.

Tabel 17. Perbandingan Performa NBC Tanpa SMOTE vs NBC + SMOTE

Metrik	NBC Tanpa SMOTE	NBC + SMOTE	Perubahan
Accuracy	83,5%	76,0%	-7,5% ↓
Macro Precision	0,28	0,44	+0,16 ↑
Macro Recall	0,33	0,47	+0,14 ↑
Macro F1-Score	0,30	0,45	+0,15 ↑
Weighted Precision	0,70	0,80	+0,10 ↑
Weighted Recall	0,83	0,76	-0,07 ↓
Weighted F1-Score	0,76	0,78	+0,02 ↑
Recall Kelas Negatif	100%	84%	-16% ↓
Recall Kelas Positif	0%	43%	+43% ↑
Recall Kelas Netral	0%	13%	+13% ↑
F1-Score Kelas Positif	0,00	0,38	+0,38 ↑
F1-Score Kelas Netral	0,00	0,10	+0,10 ↑

Penerapan SMOTE meningkatkan *Macro F1-Score* sebesar 0,15 poin dari 0,30 menjadi 0,45, dan model kini mampu mendeteksi kelas positif (*recall* 43%) serta kelas netral (*recall* 13%) yang sebelumnya sama sekali tidak terdeteksi. Meskipun akurasi keseluruhan turun 7,5%, penurunan ini justru mencerminkan evaluasi yang lebih realistis karena akurasi tinggi sebelumnya hanya disebabkan bias terhadap kelas mayoritas negatif. Peningkatan *Weighted F1-Score* dari 0,76 menjadi 0,78 juga menunjukkan perbaikan kinerja keseluruhan yang nyata. Hasil ini membuktikan bahwa SMOTE efektif dalam menangani ketidakseimbangan kelas pada analisis sentimen.

3.7 Analisis Wordcloud per Kelas Sentimen

Pada sub bab ini disajikan visualisasi *wordcloud* untuk masing-masing kategori sentimen, yang bertujuan menampilkan kata-kata yang paling dominan dan sering muncul pada setiap kelas sentimen.



Gambar 2. Wordcloud Sentimen Positif

Gambar 1 merupakan *WordCloud* sentimen positif yang menampilkan kata-kata paling sering muncul pada data yang dikategorikan sebagai sentimen positif. Kata dengan ukuran paling besar adalah "perang" dan "dagang", yang menunjukkan bahwa topik utama dalam sentimen ini berkaitan dengan isu perang dagang. Selain itu, kata seperti "china", "amerika", dan "indonesia" juga cukup dominan, menandakan bahwa pembahasan banyak terkait hubungan perdagangan antarnegara. Kemunculan kata "tarif", "impor", "harga", dan "emas" menunjukkan fokus pada aspek ekonomi dan kebijakan perdagangan. Secara keseluruhan, sentimen positif didominasi oleh isu ekonomi global dan kebijakan perdagangan yang dipersepsikan secara baik dalam data yang dianalisis.



Gambar 3. Wordcloud Sentimen Negatif

Gambar 2 memperlihatkan kata-kata yang sering muncul dalam data yang bernada negatif. Kata seperti "perang", "dagang", dan "china" tampak paling dominan, menunjukkan bahwa isu konflik perdagangan menjadi fokus utama. Selain itu, muncul kata-kata seperti "tarif", "ancam", "turun", dan "dampak", yang mencerminkan kekhawatiran terhadap kondisi ekonomi dan kemungkinan kerugian yang ditimbulkan. Secara umum, sentimen negatif menggambarkan pandangan yang menyoroti risiko dan tekanan akibat situasi perang dagang tersebut [4].



Gambar 4. Wordcloud Sentimen Netral

Gambar 3 menampilkan kata-kata yang sering muncul pada data yang bersifat informatif atau deskriptif tanpa kecenderungan opini tertentu. Kata seperti "perang", "dagang", dan "china" masih mendominasi, menunjukkan bahwa topik utama tetap berkaitan dengan isu perang dagang. Namun, kemunculan kata "lihat", "laku", "naik", "tingkat", dan "ekonomi" menunjukkan bahwa sentimen ini lebih banyak berisi informasi, analisis, atau laporan kondisi ekonomi. Secara umum, wordcloud ini menggambarkan penyampaian fakta dan perkembangan situasi tanpa menonjolkan penilaian positif maupun negatif.



3.8 Pembahasan

Hasil penelitian menunjukkan bahwa distribusi sentimen publik terhadap kebijakan perang dagang Amerika Serikat didominasi oleh sentimen negatif sebesar 83,5%, sedangkan sentimen positif dan netral masing-masing hanya sebesar 12,8% dan 3,8%. Dominasi sentimen negatif tersebut mengindikasikan bahwa sebagian besar masyarakat memberikan respons yang cenderung kritis terhadap kebijakan kenaikan tarif impor Amerika Serikat karena dianggap berpotensi memberikan dampak terhadap perdagangan internasional maupun kondisi ekonomi Indonesia. Temuan ini sejalan dengan penelitian Maghriby dan Irawan [16], yang menyatakan bahwa isu ekonomi global pada media sosial umumnya memperoleh dominasi sentimen negatif karena masyarakat lebih banyak menyoroti risiko ekonomi dibandingkan potensi manfaatnya.

Hasil penelitian juga menunjukkan bahwa algoritma *Naïve Bayes Classifier* (NBC) tanpa penanganan ketidakseimbangan data menghasilkan nilai akurasi sebesar 83,5%, namun model gagal mengenali kelas sentimen positif dan netral. Kondisi tersebut menunjukkan bahwa tingginya nilai akurasi tidak selalu mencerminkan kualitas model ketika distribusi kelas tidak seimbang. Fenomena ini sesuai dengan penelitian Rahmatulloh dkk. [5] dan Palupi [17], yang menjelaskan bahwa model *Naïve Bayes* cenderung memberikan prediksi yang bias terhadap kelas mayoritas apabila diterapkan pada dataset dengan distribusi kelas yang tidak proporsional.

Setelah diterapkan metode *Synthetic Minority Over-sampling Technique* (SMOTE), performa model mengalami perubahan yang cukup signifikan. Walaupun nilai akurasi menurun menjadi 76%, nilai Macro F1-Score meningkat dari 0,30 menjadi 0,45, sedangkan recall kelas positif meningkat dari 0% menjadi 43% dan recall kelas netral meningkat dari 0% menjadi 13%. Peningkatan tersebut menunjukkan bahwa SMOTE berhasil mengurangi bias model terhadap kelas mayoritas sehingga kemampuan model dalam mengenali seluruh kategori sentimen menjadi lebih seimbang. Hasil ini konsisten dengan penelitian Wijaya dan Nugroho [11], yang melaporkan bahwa penerapan SMOTE mampu meningkatkan kemampuan klasifikasi pada kelas minoritas meskipun sering diikuti oleh penurunan akurasi secara keseluruhan.

Apabila dibandingkan dengan beberapa penelitian sebelumnya yang menggunakan algoritma *Naïve Bayes* pada analisis sentimen, penelitian ini memiliki karakteristik yang berbeda. Sebagian besar penelitian terdahulu, seperti penelitian Insan dkk. [9] mengenai ulasan aplikasi BRI Mo maupun penelitian Putri dkk. [12] mengenai sentimen produk skincare, menggunakan dataset dengan distribusi kelas yang relatif lebih seimbang sehingga akurasi menjadi indikator utama keberhasilan model. Sebaliknya, penelitian ini menggunakan data Twitter mengenai isu perang dagang yang memiliki distribusi kelas sangat tidak seimbang, sehingga evaluasi menggunakan Macro Precision, Macro Recall, dan Macro F1-Score menjadi lebih relevan dibandingkan hanya menggunakan accuracy. Oleh karena itu, peningkatan nilai Macro F1-Score setelah penerapan SMOTE menunjukkan peningkatan kualitas model yang lebih representatif dalam menangani klasifikasi multi-kelas.

Selain itu, hasil visualisasi wordcloud memperlihatkan bahwa kata-kata dominan pada ketiga kategori sentimen masih berkaitan dengan istilah *perang, dagang, China, Amerika, tarif, dan ekonomi*. Hal tersebut menunjukkan bahwa percakapan publik lebih banyak berfokus pada dampak ekonomi akibat kebijakan perdagangan internasional dibandingkan aspek politik semata. Temuan ini memperkuat hasil distribusi sentimen yang menunjukkan dominasi persepsi negatif masyarakat terhadap kebijakan perang dagang Amerika Serikat.

Secara keseluruhan, penelitian ini memberikan kontribusi dalam dua aspek. Pertama, penelitian ini memperlihatkan bahwa kombinasi preprocessing yang lengkap, pelabelan otomatis menggunakan kamus InSet, representasi fitur TF-IDF, serta penerapan SMOTE mampu menghasilkan model klasifikasi sentimen yang lebih seimbang dibandingkan penggunaan *Naïve Bayes* tanpa penanganan ketidakseimbangan data. Kedua, penelitian ini memperluas penerapan analisis sentimen pada isu ekonomi global, khususnya persepsi masyarakat Indonesia terhadap perang dagang Amerika Serikat, yang hingga saat ini masih relatif sedikit dibahas pada penelitian sebelumnya.

4. KESIMPULAN

Berdasarkan keseluruhan rangkaian penelitian yang telah dilakukan, dapat disimpulkan bahwa analisis sentimen terhadap kebijakan kenaikan tarif impor Amerika Serikat tahun 2025 yang memicu isu perang dagang berhasil diimplementasikan secara sistematis melalui tahapan pengumpulan data, preprocessing, pelabelan, hingga klasifikasi. Proses crawling menghasilkan data tweet yang kemudian melalui tahap preprocessing seperti cleaning, case folding, tokenizing, stopword removal, normalisasi slangword, dan stemming sehingga diperoleh data yang lebih bersih dan siap dianalisis. Hasil pelabelan berbasis leksikon menunjukkan bahwa opini publik cenderung didominasi oleh sentimen negatif, yang mengindikasikan adanya respons kritis dan kekhawatiran masyarakat terhadap dampak kebijakan tersebut. Pada tahap klasifikasi menggunakan algoritma *Naïve Bayes*, model awal tanpa penanganan ketidakseimbangan data menghasilkan akurasi yang cukup tinggi, namun tidak mampu mengenali kelas minoritas secara optimal karena dominasi kelas negatif. Setelah diterapkan teknik SMOTE, distribusi data menjadi lebih seimbang sehingga model mampu meningkatkan kemampuan dalam mengenali kelas positif dan netral secara lebih proporsional, meskipun terjadi penurunan akurasi keseluruhan. Namun demikian, peningkatan nilai recall pada kelas minoritas serta kenaikan Macro F1-Score menunjukkan bahwa model menjadi lebih adil dan representatif dalam melakukan klasifikasi. Dengan demikian, penelitian ini membuktikan bahwa kombinasi preprocessing yang tepat dan penanganan data tidak seimbang sangat



berpengaruh terhadap kualitas hasil analisis sentimen, serta memberikan gambaran yang lebih objektif terhadap persebaran opini publik terkait isu perang dagang tersebut.

REFERENCES

- [1] A. Rahman, F. Rahmat, M. Y. Fariqi, S. Adi, and P. S. Informatika, "Metode Naive Bayes untuk Menganalisis Akurasi Sentimen Komentar di Youtube," *J. EECCIS*, vol. 14, no. 1, pp. 31–34, 2020, [Online]. Available: <http://bit.ly/2u802Pe>
- [2] B. Syabani Sabrawera and Farid Hirji Badruzzaman, "Peran Jaringan Bisnis dalam Upaya Pengembangan Usaha (Studi Kasus pada Moriska Baby)," *JEMSI (Jurnal Ekon. Manajemen, dan Akuntansi)*, vol. 10, no. 4, pp. 2413–2423, 2024, doi: 10.35870/jemsi.v10i4.2621.
- [3] Y. Ariyani and D. H. Perkasa, "Bagaimana Kecerdasan Budaya Memengaruhi Keterikatan Kerja dan Retensi Karyawan dalam Tim Multikultural: Literature Review," *J. Ekon. Manaj. Sist. Inf.*, vol. 6, no. 3, pp. 1411–1422, 2025, doi: 10.38035/jemsi.v6i3.3663.
- [4] M. Andriana and T. Sumarlin, "Analisis Sistem Informasi Anggaran," *J. Manaj. Sos. Ekon.*, vol. 3, no. 2, pp. 158–163, 2023.
- [5] R. Rahmatulloh, M. Iqbal Ibrahim, M. R. Handayani, K. Umam, and N. C. H. Wibowo, "Model Klasifikasi Naive Bayes untuk Pemetaan Persepsi Publik Secara Real-Time pada Media Sosial: Studi Kasus RUU TNI 2025," *Decod. J. Pendidik. Teknol. Inf.*, vol. 5, no. 2, pp. 365–379, 2025, [Online]. Available: <https://journal.umkendari.ac.id/decode/article/view/1139>
- [6] T. I. Solihati, N. Hidayanti, and R. Kania, "Implementasi Data Mining Evaluasi Kinerja Penelitian Mahasiswa Dengan Menggunakan Algoritma Naive Bayes," *J. Theorems (The Orig. Research ...)*, vol. 6, no. 2, pp. 135–147, 2022.
- [7] T. Setiadi, F. Noviyanto, H. Hardianto, A. Tarmuji, A. Fadlil, and M. Wibowo, "Implementation of naïve bayes method in food crops planting recommendation," *Int. J. Sci. Technol. Res.*, vol. 9, no. 2, pp. 4750–4755, 2020.
- [8] A. Pebdika, R. Herdiana, and D. Solihudin, "Klasifikasi Menggunakan Metode Naive Bayes Untuk Menentukan Calon Penerima Pip," *JATI (Jurnal Mhs. Tek. Inform.)*, vol. 7, no. 1, pp. 452–458, 2023, doi: 10.36040/jati.v7i1.6303.
- [9] M. K. Insan, U. Hayati, and O. Nurdian, "Analisis Sentimen Aplikasi Brimo Pada Ulasan Pengguna Di Google Play Menggunakan Algoritma Naive Bayes," *J. Mhs. Tek. Inform.*, vol. 7, no. 1, pp. 478–483, 2023.
- [10] A. Misbachudin Riyadi, H. Sibyan, I. Ahmad Ihsanuddin, and M. Alif Muwafiq Baihaqi, "Klasifikasi Penerima Beasiswa Menggunakan Metode Naive Bayes (Studi Kasus SMP Negeri 3 Selomerto)," *J. Eng. Inform.*, vol. 1, no. 2, pp. 53–59, 2023, doi: 10.56854/jei.v1i2.61.
- [11] Heliyanti Susana, "Penerapan Model Klasifikasi Metode Naive Bayes Terhadap Penggunaan Akses Internet," *J. Ris. Sist. Inf. dan Teknol. Inf.*, vol. 4, no. 1, pp. 1–8, 2022, doi: 10.52005/jursistekni.v4i1.96.
- [12] K. S. Putri, I. R. Setiawan, and A. Pambudi, "Analisis Sentimen Terhadap Brand Skincare Lokal Menggunakan Naive Bayes Classifier," *Technol. J. Ilm.*, vol. 14, no. 3, p. 227, 2023, doi: 10.31602/tji.v14i3.11259.
- [13] B. F. Haikal *et al.*, "Perbandingan Algoritma Naive Bayes Dan Dempster Shafer Untuk Diagnosis Penyakit ISPA," vol. 04, no. 3, pp. 147–157, 2025.
- [14] M. A. Mokoagow and A. S. Purnomo, "Penerapan Metode Naive Bayes Pada Sistem Pakar Untuk Mendiagnosis Penyakit Ibu Hamil," vol. 4, no. 2, 2024.
- [15] J. Kecerdasan and T. Informasi, "Sistem Pakar Diagnosis Penyakit Ispa Menggunakan Metode Naive Expert System Diagnosis Of Ari Disease Using Naive Bayes Method Based On Web Based Puskesmas Teratak," vol. 2, no. 1, pp. 32–42, 2023.
- [16] M. A. Maghriby and H. Irawan, "Analisis Persepsi Publik Mengenai Resesi Ekonomi Global 2023 Sektor Bisnis di Media Sosial Twitter Menggunakan Algoritma Naive Bayes dan Topic Modelling," *Widya Cipta J. Sekr. dan Manaj.*, vol. 7, no. 2, pp. 74–85, 2023, doi: 10.31294/widyacipta.v7i2.15577.
- [17] E. Sri Palupi, "Klasifikasi Sentimen Netizen Di Media Sosial X Terhadap Ancaman Resesi Ekonomi Indonesia Dengan Menggunakan Algoritma Naive Bayes," *JATI (Jurnal Mhs. Tek. Inform.)*, vol. 9, no. 5, pp. 8058–8064, 2025, doi: 10.36040/jati.v9i5.14878.
- [18] M. Khanna, M. Kulshrestha, L. K. Singh, S. Thawkar, and K. Shrivastava, "Performance Evaluation of Machine Learning Algorithms for Stock Price and Stock Index Movement Prediction Using Trend Deterministic Data Prediction," *Int. J. Appl. Metaheuristic Comput.*, vol. 13, no. 1, pp. 1–30, 2022, doi: 10.4018/ijamc.292511.
- [19] N. Nazifah, "Analisis Perbandingan Decision Tree Algoritma C4.5 dengan algoritma lainnya: Sistematis Literature Review," *J. Inform. dan Teknol. Komput. (J-ICOM)*, vol. 4, no. 2, pp. 57–64, 2023, doi: 10.55377/j-icom.v4i2.7719.
- [20] D. Suryani, A. Yulianti, E. L. Maghfiroh, and ..., "Quality Classification of Palm Oil Products Using Naive Bayes Method," *Sist. J. Sist. ...*, 2021.