

BAB 2

LANDASAN TEORI

2.1 Ketidakpastian Kebijakan Perdagangan (TPU) dan Tarif Trump

Trade Policy Uncertainty (TPU) merujuk pada ketidakpastian ekonomi yang muncul akibat perubahan mendadak atau ancaman perubahan dalam regulasi perdagangan internasional oleh pemerintah, yang sering kali digunakan sebagai instrumen tekanan politik dan negosiasi [1, 2]. Kebijakan tarif agresif yang diterapkan selama masa pemerintahan Donald Trump terbukti memicu lonjakan TPU yang signifikan secara global. Strategi "Trump 2.0" yang direncanakan pada April 2025 diprediksi menjadi ekspansi proteksionisme terbesar Amerika Serikat sejak tahun 1930-an, dengan kombinasi tarif dasar 10% dan tambahan tarif "timbal balik" berdasarkan surplus perdagangan negara mitra [2, 4].

Secara empiris, peningkatan TPU ini berdampak buruk pada operasional perusahaan melalui penurunan investasi sebesar 2,3%, pemotongan biaya riset (R&D) sebesar 2,3%, hingga memangkas laba bersih perusahaan secara drastis mencapai 11,5% [1, 6]. Bagi Indonesia, kebijakan proteksionisme ini merupakan ancaman serius karena Amerika Serikat membeli sekitar satu dari sepuluh dolar dari total ekspor barang Indonesia, atau sekitar \$23,3 miliar pada tahun 2021 [4, 6]. Rumus tarif timbal balik tersebut berpotensi menaikkan bea masuk barang ekspor Indonesia ke AS menjadi sebesar 32%, yang dapat menyebabkan kekurangan (*shortfall*) nilai ekspor mencapai USD 1 miliar atau setara 0,3% dari PDB nasional [4].

Sektor-sektor strategis seperti tekstil, elektronik, dan pertanian menjadi yang paling rentan terhadap guncangan tarif tersebut [6, 7]. Ancaman tarif tersebut tidak semata-mata berimbas pada hubungan bilateral antarnegara, melainkan turut merusak kestabilan rantai pasok global dan memicu tindakan balasan (*retaliation*) dari blok ekonomi lain, yang memperburuk instabilitas ekonomi internasional secara sistemik [2]. Pemantauan terhadap respon publik menjadi krusial bagi pemerintah guna mendeteksi potensi keresahan ekonomi dan merancang strategi mitigasi terhadap hambatan dagang non-tarif [6, 12].

2.2 Media Sosial

Media sosial telah bertransformasi menjadi arena primer bagi masyarakat untuk mengekspresikan pemikiran, emosi, dan persepsi terhadap berbagai isu strategis, termasuk kebijakan ekonomi global dan ketidakpastian kebijakan perdagangan (*Trade Policy Uncertainty*) [9, 12]. Secara kolektif, platform digital ini berfungsi sebagai repositori data emosional individu yang sangat masif, di mana volume informasi yang dihasilkan dapat mencapai 12 GB setiap harinya [8]. Di Indonesia, platform media sosial berperan sebagai saluran utama bagi masyarakat untuk menyuarakan respon terhadap kebijakan publik secara *real-time*, didukung oleh basis pengguna aktif platform X (sebelumnya Twitter) yang mencapai kurang lebih 27,5 juta orang pada tahun 2024 [9, 25].

Twitter/X secara khusus diakui sebagai sensor emosional pasar secara *real-time* yang mampu memengaruhi keputusan investor dan volatilitas keuangan global [8, 9]. Dalam ranah ekonomi, pernyataan tokoh berpengaruh atau unggahan yang mengandung kata kunci sensitif seperti “*Tariff*” atau “*Trade War*” terbukti memicu respons instan pada pasar modal, termasuk penurunan harga pada indeks S&P 500 dan lonjakan indeks volatilitas (VIX) dalam waktu singkat setelah dipublikasikan [8]. Fenomena ini sangat relevan mengingat kebijakan proteksionisme “Trump 2.0” pada April 2025 yang mengenakan tarif universal diprediksi akan berdampak signifikan pada sektor manufaktur padat karya Indonesia [4]. Analisis terhadap percakapan digital ini memberikan gambaran spektrum emosional yang luas, mulai dari antusiasme positif hingga kritik negatif yang mencerminkan tingkat kekhawatiran publik terhadap guncangan ekonomi [9, 10].

Selain platform berbasis teks pendek, YouTube telah berkembang menjadi wadah diskusi publik yang signifikan melalui fitur komentarnya yang merepresentasikan persepsi masyarakat terhadap konten visual berita ekonomi makro [10, 12]. Berbeda dengan X yang bersifat spontan dan singkat, komentar di YouTube cenderung lebih naratif dan mendalam, yang memberikan dimensi tambahan bagi analisis kebijakan karena sering kali memicu interaksi berantai melalui fitur balasan (*replies*) [10, 12]. Penggabungan data dari kedua platform ini dianggap esensial untuk mendapatkan analisis yang komprehensif, mengingat YouTube mampu menangkap diskursus publik yang lebih terperinci mengenai dampak jangka panjang dari sebuah kebijakan ekonomi dibandingkan pesan singkat di Twitter [10, 12].

2.3 Natural Language Processing (NLP)

Dalam ranah kecerdasan buatan, terdapat satu cabang ilmu yang secara spesifik berupaya menjembatani kesenjangan antara cara manusia berkomunikasi dan kemampuan mesin dalam memahaminya, yang dikenal dengan istilah *Natural Language Processing* (NLP) [9, 20, 21, 26]. Tujuan utamanya adalah menggali makna serta pandangan subjektif yang tersembunyi di balik kumpulan teks yang belum tersusun secara sistematis. Pada tahap berikutnya, NLP berfungsi sebagai fondasi komputasional yang mengubah satuan-satuan bahasa tertulis menjadi representasi berbentuk angka, sehingga data tersebut dapat diproses lebih lanjut oleh model-model klasifikasi pembelajaran mesin [9, 20, 27].

NLP bertindak sebagai jembatan yang memungkinkan mesin untuk mengenali pola bahasa alami melalui berbagai teknik pemrosesan teks, seperti klasifikasi sentimen dan pengenalan entitas [26]. Aplikasi NLP dalam pemrosesan teks mencakup berbagai tahapan teknis, mulai dari pembersihan data mentah hingga ekstraksi fitur yang bermakna secara semantik untuk digunakan oleh algoritma pembelajaran mesin [20].

Tantangan utama dalam NLP untuk media sosial adalah sifat penulisan yang sering kali tidak formal, variatif dalam penggunaan bahasa, serta banyaknya penggunaan singkatan dan *slang*, terutama dalam bahasa Indonesia informal [7, 12, 20, 28]. Teknik dasar NLP seperti *preprocessing* teks dan ekstraksi fitur statistik tetap menjadi fondasi krusial dalam menangani kompleksitas data media sosial yang dipenuhi oleh variasi bahasa tidak baku [7, 12, 20, 28].

2.4 Analisis Sentimen

Opinion mining, atau yang lebih populer disebut analisis sentimen, adalah salah satu kajian dalam bidang komputasi yang berfokus pada upaya mengidentifikasi sekaligus mengklasifikasikan kecenderungan sikap dan nuansa emosional yang terdapat pada suatu teks, lalu mengelompokkannya menjadi tiga kutub utama yaitu positif, negatif, dan netral [22, 23, 26, 28]. Secara teknis, analisis ini berupaya memetakan gambaran umum persepsi masyarakat terhadap kualitas layanan, produk, maupun kebijakan pemerintah [7, 10, 25, 28].

Bidang penelitian ini telah berkembang pesat seiring dengan meluasnya penggunaan web dan media sosial sebagai wadah publik untuk mengekspresikan sikap terhadap suatu produk, layanan, maupun kebijakan publik [10, 28]. Penerapan

analisis sentimen otomatis menjadi sangat penting untuk memproses volume data opini digital yang besar secara cepat dan efisien [12, 21]. Hasil klasifikasi emosi massa ini memberikan wawasan strategis bagi pengambil keputusan untuk memantau respon publik terhadap isu-isu strategis nasional seperti dampak ekonomi akibat kebijakan perdagangan luar negeri [10, 12, 29].

Penelitian ini mengadopsi pendekatan pembelajaran mesin menggunakan *Naïve Bayes* karena efisiensi dan kemampuannya bekerja optimal dengan data latih terbatas [22, 29].

2.5 Preprocessing Teks

Sebelum data mentah dapat diproses oleh algoritma klasifikasi, diperlukan serangkaian langkah awal yang disebut *preprocessing*, yang berperan penting dalam menyaring berbagai gangguan (*noise*) sehingga kualitas hasil pengklasifikasian dapat meningkat [7, 20, 21, 27, 28]. Tahapan ini sangat krusial karena keputusan yang dibuat selama pemrosesan awal akan berdampak langsung pada akurasi dan validitas hasil analisis sentimen akhir [28].

1. *Cleaning*

Cleaning mencakup eliminasi elemen teks yang dianggap tidak bermakna seperti URL, simbol khusus, angka, *emoji*, *hashtag*, serta penyebutan pengguna (*@username*) [7, 10, 20, 21]. Penghapusan karakter *noise* ini diperlukan agar algoritma dapat fokus sepenuhnya pada kosakata yang membawa informasi sentimen yang signifikan [20].

2. *Case Folding*

Case folding merupakan tahap yang berfungsi mengubah keseluruhan karakter pada dokumen teks menjadi format huruf kecil (*lowercase*), dengan maksud mencegah munculnya duplikasi kata yang sebenarnya sama namun berbeda dari segi penulisan huruf besar-kecilnya [7, 10, 20, 27]. Langkah ini menjamin bahwa sistem bersifat *case-insensitive*, sehingga kata yang identik namun memiliki perbedaan kapitalisasi tidak akan diperlakukan sebagai dua entitas yang berbeda [26].

3. *Tokenization*

Tokenization merupakan tahapan yang berfungsi memilah suatu rangkaian kalimat menjadi bagian-bagian terkecil berupa kata atau *token*, di mana hasil

pemilahan ini menjadi acuan inti pada keseluruhan proses pengolahan data berbasis teks [7, 10, 20, 21]. Hasil pemilahan berupa *token-token* tersebut nantinya tidak berhenti pemanfaatannya di tahap ini saja, melainkan terus digunakan hingga proses ekstraksi fitur dilaksanakan [28].

4. *Normalization*

Tahap normalisasi difokuskan pada upaya mengubah kata-kata tidak baku, bentuk singkatan, maupun ungkapan gaul ke dalam wujud kata baku yang sesuai dengan acuan KBBI, sehingga makna semantik antarkata tetap konsisten [7, 10, 20, 25]. Langkah semacam ini menjadi krusial, terutama pada data yang bersumber dari media sosial, mengingat banyaknya penggunaan bahasa gaul (*slang*) serta berbagai kesalahan pengetikan (*typos*) yang dapat mengganggu keseragaman makna jika tidak ditangani [28].

5. *Stopword Removal*

Stopword removal adalah tahap yang difungsikan untuk menyaring kosakata-kosakata yang lazim dijumpai berulang kali di dalam teks, namun pada dasarnya tidak memberikan sumbangan informasi yang signifikan bagi proses analisis sentimen [7, 20, 26, 27]. Melalui tahap ini, dimensi data dapat dipersempit sehingga efisiensi dari sisi komputasi turut meningkat [10].

6. *Stemming*

Stemming mengembalikan kata-kata berimbuhan ke bentuk dasarnya (*root word*). Dalam konteks bahasa Indonesia, proses ini membantu menyederhanakan variasi morfologi kata agar algoritma dapat mengenalinya sebagai unit makna yang sama [10, 20, 21, 27]. Meskipun teknik alternatif seperti *lemmatization* juga tersedia, *stemming* tetap populer karena kesederhanaannya dalam menyederhanakan variasi morfologi kata [28].

2.6 Term Frequency-Inverse Document Frequency (TF-IDF)

TF-IDF merupakan suatu skema pembobotan berbasis statistik yang dirancang untuk mengkuantifikasi seberapa signifikan peranan sebuah kata dalam suatu dokumen apabila dibandingkan dengan keseluruhan dokumen yang terhimpun dalam sebuah korpus [7, 11, 20, 30]. Melalui pendekatan ini, representasi teks yang semula bersifat kualitatif dikonversi ke dalam bentuk vektor numerik (*Vector Space Model*), di mana masing-masing kata memperoleh nilai bobot tersendiri yang

proporsional terhadap derajat relevansi dan keunikannya dalam konteks keseluruhan korpus [20].

2.6.1 Term Frequency (TF)

Term Frequency (TF) mengukur intensitas atau seberapa sering sebuah kata muncul dalam satu dokumen spesifik [11, 20, 23]. Rumus TF adalah sebagai berikut:

$$TF(t, d) = \frac{f_{t,d}}{\sum_{t' \in d} f_{t',d}} \quad (2.1)$$

Di mana:

- $TF(t, d)$: *Term Frequency* kata t dalam dokumen d
- $f_{t,d}$: Frekuensi kemunculan kata t dalam dokumen d
- $\sum_{t' \in d} f_{t',d}$: Total seluruh kata dalam dokumen d

2.6.2 Inverse Document Frequency (IDF)

Inverse Document Frequency (IDF) menilai tingkat kepentingan kata di seluruh korpus, di mana kata yang jarang muncul mendapatkan bobot lebih tinggi karena dianggap lebih informatif [11, 20, 27]. Rumus IDF adalah:

$$IDF(t) = \log \frac{N}{df(t)} \quad (2.2)$$

Di mana:

- $IDF(t)$: *Inverse Document Frequency* kata t
- N : Total jumlah dokumen dalam korpus
- $df(t)$: Jumlah dokumen yang mengandung kata t

2.6.3 Perhitungan TF-IDF

Nilai akhir TF-IDF diperoleh dari perkalian antara skor TF dan IDF guna menghasilkan fitur numerik yang optimal bagi model klasifikasi [7, 20, 21]:

$$TF-IDF(t, d) = TF(t, d) \times IDF(t) \quad (2.3)$$

Semakin tinggi nilai *TF-IDF* suatu istilah, semakin besar peran istilah tersebut dalam membedakan satu dokumen dengan dokumen lainnya, sehingga sering kali memberikan hasil yang lebih baik dibandingkan metode ekstraksi fitur sederhana seperti *n-gram* [20].

2.7 Bag of Words (BoW) dan CountVectorizer

Bag of Words (BoW) adalah model ekstraksi fitur yang merepresentasikan teks berdasarkan frekuensi kemunculan kata tanpa memperhatikan urutan maupun struktur gramatikalnya [11, 20]. Dalam model ini, setiap dokumen dipandang sebagai kumpulan kata unik (*vocabulary*) beserta jumlah kemunculannya masing-masing [11]. Dengan demikian, dua kalimat yang memiliki susunan kata berbeda namun menggunakan komposisi kata yang identik akan menghasilkan representasi vektor yang sama [11].

CountVectorizer merupakan implementasi teknis dari konsep BoW yang tersedia pada pustaka *scikit-learn* untuk bahasa pemrograman Python [31]. Cara kerjanya dimulai dengan membangun kosakata dari seluruh korpus data, kemudian melakukan transformasi teks menjadi matriks dokumen-kata (*document-term matrix*) [31]. Setiap sel dalam matriks tersebut berisi nilai bilangan bulat (*integer count*) yang mencerminkan frekuensi kemunculan kata *t* dalam dokumen *d* [31, 32]. Berbeda dengan BoW yang bersifat konseptual, *CountVectorizer* bertindak sebagai alat (*tool*) yang mengotomatisasi konversi teks mentah menjadi fitur numerik diskrit yang siap diolah oleh model klasifikasi [31].

Karakteristik keluaran *CountVectorizer* yang menghasilkan nilai *integer* (bilangan cacah) inilah yang membuatnya sangat sesuai untuk dipasangkan dengan algoritma *Multinomial Naive Bayes Classifier* (MNBC) [31, 32]. Secara matematis, MNBC diturunkan dari distribusi multinomial yang dirancang khusus untuk memodelkan data berupa hitungan kejadian diskrit, seperti frekuensi kata dalam teks, dan bukan untuk nilai kontinu [31, 32]. Hal ini memastikan bahwa perhitungan probabilitas pada model didasarkan pada distribusi data yang tepat [9, 20].

2.8 Synthetic Minority Over-sampling Technique (SMOTE)

SMOTE merupakan teknik penyeimbangan data yang dirancang untuk menanggulangi persoalan ketidakseimbangan kelas (*class imbalance*) pada *dataset*, di mana model *machine learning* bertendensi condong ke kelas mayoritas dan mengabaikan kelas minoritas [23, 27]. SMOTE beroperasi dengan cara membangkitkan sampel sintetis baru melalui interpolasi linear antara sampel minoritas yang tersedia dengan tetangga terdekatnya (*k-nearest neighbors*), sehingga lebih bervariasi dibandingkan sekadar menduplikasi data yang sudah ada [23]. Penting untuk dicatat bahwa SMOTE hanya diterapkan pada data latih guna mencegah terjadinya kebocoran informasi (*data leakage*) yang dapat membuat evaluasi model menjadi tidak objektif [20, 27].

2.9 Algoritma Naïve Bayes Classifier (NBC)

Naïve Bayes Classifier (NBC) tergolong ke dalam keluarga algoritma klasifikasi probabilistik pada pembelajaran mesin terawasi yang penerapannya berlandaskan Teorema Bayes. Penyebutan "naif" pada algoritma ini muncul dari penyederhanaan asumsi yang diterapkan, yakni menganggap seluruh fitur saling lepas (*independen*) tanpa adanya korelasi antarfitur dalam memprediksi suatu kelas [11, 22, 31, 32].

Teorema Bayes dimanfaatkan untuk menghitung probabilitas *posterior* dari suatu kelas (C) berdasarkan fitur (d) yang diamati melalui perbandingan *likelihood*, *prior*, dan *evidence* [21, 22, 25, 31]. Secara matematis, teorema ini dinyatakan sebagai:

$$P(C|d) = \frac{P(d|C) \cdot P(C)}{P(d)} \quad (2.4)$$

Di mana:

- $P(C|d)$: Probabilitas *posterior* (probabilitas kelas C setelah mengamati dokumen d)
- $P(d|C)$: *Likelihood* (probabilitas dokumen d muncul dalam kelas C)
- $P(C)$: *Prior* (probabilitas awal kelas C)
- $P(d)$: *Evidence* (probabilitas dokumen d secara keseluruhan)

Algoritma ini menggunakan estimasi *Maximum A Posteriori* (MAP) untuk menentukan kelas dengan probabilitas tertinggi bagi sekumpulan data input tertentu [31].

2.9.1 Varian Algoritma Naïve Bayes

Berdasarkan Teorema Bayes tersebut, algoritma *Naïve Bayes Classifier* memiliki beberapa varian yang berbeda menurut asumsi distribusi data terhadap fitur input yang digunakan, di antaranya *Multinomial Naïve Bayes*, *Bernoulli Naïve Bayes*, dan *Gaussian Naïve Bayes* [20].

Multinomial Naïve Bayes dirancang untuk memodelkan data hitungan (*count data*), seperti frekuensi kemunculan kata pada hasil ekstraksi fitur berbasis *CountVectorizer* maupun *TF-IDF*, sehingga umum digunakan pada permasalahan klasifikasi teks [20, 31].

Bernoulli Naïve Bayes mengasumsikan setiap fitur bersifat biner, yaitu hanya merepresentasikan ada atau tidaknya suatu kata dalam dokumen tanpa memperhitungkan frekuensi kemunculannya [31]. Berdasarkan studi pembandingan model *event Naïve Bayes* untuk klasifikasi teks, *Bernoulli Naïve Bayes* cenderung unggul pada kondisi ukuran kosakata (*vocabulary*) yang kecil, sedangkan *Multinomial Naïve Bayes* umumnya memberikan performa lebih baik pada ukuran kosakata yang lebih besar [33].

Gaussian Naïve Bayes dirancang untuk menangani fitur bertipe data kontinu yang diasumsikan mengikuti distribusi normal, sehingga lebih umum diterapkan pada permasalahan klasifikasi dengan atribut numerik non-tekstual dibandingkan data hasil ekstraksi fitur teks [20, 31].

Berdasarkan karakteristik ketiga varian tersebut, penelitian ini menggunakan *Multinomial Naïve Bayes* karena data hasil ekstraksi fitur (*CountVectorizer* maupun *TF-IDF*) berbentuk hitungan frekuensi kata, sehingga sesuai dengan asumsi distribusi multinomial sebagaimana dijelaskan lebih rinci pada Subbab 2.9.2.

2.9.2 Multinomial Naïve Bayes

Varian *Multinomial Naïve Bayes* sangat efektif untuk klasifikasi teks karena bekerja berdasarkan frekuensi kemunculan kata atau *token* di dalam dokumen sebagai dasar perhitungan probabilitas. Varian ini sangat ideal untuk data diskrit dan memodelkan frekuensi kemunculan kata sebagai hitungan (*counts*) dengan asumsi

distribusi multinomial untuk setiap fitur [9, 20, 22, 32].

Rumus Multinomial Naïve Bayes adalah:

$$P(C|d) \propto P(C) \prod_{i=1}^{n_d} P(t_i|C) \quad (2.5)$$

Di mana:

- $P(C|d)$: Probabilitas dokumen d termasuk dalam kelas C
- $P(C)$: Probabilitas prior kelas C
- $P(t_i|C)$: Probabilitas kata t_i muncul dalam kelas C
- n_d : Jumlah total kata dalam dokumen d
- $\propto$: Sebanding dengan (karena nilai probabilitas akhir hanya dibandingkan antarkelas untuk menentukan klasifikasi)

2.9.3 Contoh Penerapan Naïve Bayes pada Analisis Sentimen

Sebagai ilustrasi penerapan Persamaan 2.5, berikut disajikan contoh perhitungan sederhana menggunakan data hipotetis (bukan data aktual penelitian) yang telah melalui tahap *preprocessing*. Misalkan terdapat tiga dokumen latihan sebagai berikut:

- Kelas Positif: "untung ekspor"
- Kelas Negatif: "rugikan ekspor", "rugikan industri"

Kosakata (*vocabulary*) yang terbentuk dari seluruh korpus berjumlah 4 kata: {untung, ekspor, rugikan, industri}.

Probabilitas *prior* untuk masing-masing kelas dihitung dari proporsi dokumen latihan:

$$P(\text{Positif}) = \frac{1}{3}, \quad P(\text{Negatif}) = \frac{2}{3}$$

Probabilitas *likelihood* tiap kata dihitung menggunakan *Laplace smoothing* untuk menghindari nilai probabilitas nol:

$$P(t_i|C) = \frac{\text{jumlah kemunculan } t_i \text{ di kelas } C + 1}{\text{total kata di kelas } C + |V|}$$

Dengan total kata pada kelas Positif = 2 dan kelas Negatif = 4, serta $|V| = 4$, diperoleh $P(\text{untung}|\text{Positif}) = \frac{2}{6}$, $P(\text{rugikan}|\text{Positif}) = \frac{1}{6}$, $P(\text{untung}|\text{Negatif}) = \frac{1}{8}$, dan $P(\text{rugikan}|\text{Negatif}) = \frac{3}{8}$.

Misalkan terdapat data uji berupa kalimat "untung rugikan". Skor masing-masing kelas dihitung dengan mengalikan *prior* dan seluruh *likelihood* kata penyusunnya sesuai Persamaan 2.5:

$$\text{Skor(Positif)} = \frac{1}{3} \times \frac{2}{6} \times \frac{1}{6} \approx 0,0185$$

$$\text{Skor(Negatif)} = \frac{2}{3} \times \frac{1}{8} \times \frac{3}{8} \approx 0,0313$$

Karena Skor (Negatif) lebih besar dibandingkan Skor (Positif), maka berdasarkan estimasi *Maximum A Posteriori* (MAP), kalimat "untung rugikan" diklasifikasikan ke dalam kelas Negatif. Ilustrasi ini menggambarkan bagaimana MNBC memanfaatkan frekuensi kemunculan kata hasil ekstraksi fitur untuk menentukan kecenderungan sentimen suatu opini publik pada penelitian ini.

2.10 Evaluasi Model Klasifikasi

Confusion Matrix digunakan sebagai instrumen evaluasi untuk menilai performa algoritma klasifikasi melalui perbandingan antara label yang diprediksi sistem dan label aktual (*ground truth*). Instrumen ini merangkum hasil klasifikasi ke dalam empat kategori, yaitu *True Positive* (TP), *True Negative* (TN), *False Positive* (FP), dan *False Negative* (FN) [20, 21, 25, 27]. Keempat nilai tersebut selanjutnya menjadi dasar penghitungan beragam metrik statistik sehingga performa model dapat diukur secara objektif [20].

1. Accuracy

Accuracy menghitung proporsi keseluruhan prediksi yang tepat, baik positif maupun negatif, terhadap total keseluruhan data uji [7, 20, 27]. Rumus *accuracy* adalah:

$$\text{Accuracy} = \frac{TP + TN}{TP + TN + FP + FN} \quad (2.6)$$

2. Precision

Precision mengukur rasio prediksi benar positif dibandingkan dengan total hasil prediksi positif yang dihasilkan model [7, 20, 25]. Rumus *precision* adalah:

$$Precision = \frac{TP}{TP + FP} \quad (2.7)$$

3. *Recall*

Recall mengukur kapasitas model dalam menjangkau kembali seluruh data aktual yang berlabel positif, atau keberhasilan sistem dalam mengambil kembali data aktual yang relevan [7, 10, 20, 27]. Rumus *recall* adalah:

$$Recall = \frac{TP}{TP + FN} \quad (2.8)$$

4. *F1-Score*

F1-Score merupakan rata-rata harmonik dari presisi dan *recall* yang menyajikan gambaran performa secara berimbang, khususnya pada kondisi data yang tidak seimbang (*imbalanced*) [7, 12, 20]. Rumus *F1-Score* adalah:

$$F1-Score = 2 \times \frac{Precision \times Recall}{Precision + Recall} \quad (2.9)$$

2.11 Penelitian terdahulu

Penelitian mengenai analisis sentimen berbasis algoritma *Naïve Bayes Classifier* (NBC) telah menjadi tolak ukur fundamental dalam mengevaluasi respon publik di platform media sosial, khususnya di Indonesia [19, 20]. Tinjauan literatur ini bertujuan untuk memetakan perkembangan teknik pemrosesan bahasa alami (NLP), membandingkan efektivitas NBC terhadap berbagai isu kebijakan publik, serta mengidentifikasi kekosongan penelitian pada domain dampak ekonomi internasional bagi ketahanan nasional. Melalui pemahaman terhadap posisi penelitian terdahulu, studi ini dapat memperkuat landasan metodologis guna menganalisis dampak guncangan kebijakan perdagangan global secara lebih akurat.

Ringkasan perbandingan antarstudi terdahulu tersebut dirangkum dan disajikan pada tabel berikut ini:



UMN
UNIVERSITAS
MULTIMEDIA
NUSANTARA

Tabel 2.1. Tabel Perbandingan Penelitian Terdahulu

No	Referensi	Fokus	Dataset	Metode	Hasil	Rangkuman Penelitian	Kesamaan	Perbedaan
1	Benguria <i>et al.</i> (2022): Anxiety or pain? The impact of tariffs and uncertainty on Chinese firms in the trade war [1]	Dampak tarif Trump terhadap investasi firma.	Laporan tahunan firma terdaftar di Tiongkok.	Regresi Kuantitatif	Turun 2.3%	Peningkatan guncangan TPU Trump menurunkan investasi firma dan memangkas laba bersih hingga 11.5%.	Fokus pada dampak ekonomi proteksionisme Trump.	Menggunakan data laporan perusahaan, bukan opini media sosial.
2	Sabila (2025): A Qualitative Policy Analysis of the Trump 2.0 Universal Tariff Title [4]	Analisis kebijakan "Trump 2.0" terhadap ekspor RI.	Statistik ekspor & dokumen kebijakan.	Kualitatif (SUT Tracing)	N/A	Memprediksi shortfall ekspor Indonesia mencapai USD 1 miliar akibat tarif universal 10-32%.	Meneliti dampak spesifik kebijakan Trump terhadap Indonesia.	Bersifat kualitatif deskriptif tanpa pengolahan NLP.
3	Andrian <i>et al.</i> (2023): Implementation Of <i>Naïve Bayes</i> Algorithm In Sentiment Analysis Of Twitter Social Media Users Regarding Their Interest To Pay The Tax [20]	Sentimen minat bayar pajak di Indonesia.	2.617 twit Twitter/X.	NBC (Bernoulli) & SMOTE	91,03%	Varian Bernoulli NBC dengan teknik SMOTE efektif menangani data bising dan tidak seimbang di Twitter.	Implementasi <i>Naïve Bayes</i> pada teks Twitter bahasa Indonesia.	Objek penelitian adalah kebijakan pajak domestik.
4	Sihombing <i>et al.</i> (2024): Comparison Of Machine Learning Algorithms In Public Sentiment Analysis Of TAPERA Policy [19]	Analisis opini publik kebijakan TAPERA.	1.189 twit Twitter/X.	NBC, SVM, & Random Forest	69,17% (NBC)	NBC terbukti paling efektif dibanding SVM dalam mengklasifikasi opini kebijakan perumahan.	Memvalidasi performa NBC untuk analisis kebijakan publik di Indonesia.	Fokus pada kebijakan domestik sektoral (TAPERA).

Lanjut ke halaman berikutnya

Tabel 2.1 – Lanjutan

No	Referensi	Fokus	Dataset	Metode	Hasil	Rangkuman Penelitian	Kesamaan	Perbedaan
5	Riswandi (2025): Sentiment Analysis of Public Comments on the YouTube Video “Trump Unveils Sweeping Global Tariffs in Watershed Moment for World Trade” by BBC News Using the Long Short-Term Memory Method [10]	Respon publik global atas tarif Trump.	3.543 komentar YouTube.	LSTM (<i>Deep Learning</i>)	76,02%	Dominasi sentimen negatif (64,6%) terhadap pengumuman tarif global Trump di YouTube.	Analisis sentimen publik terhadap guncangan kebijakan proteksionisme Trump.	Menggunakan metode LSTM dan platform YouTube.
6	Aini <i>et al.</i> (2023): Economic Impact Due Covid-19 Pandemic: Sentiment Analysis on Twitter Using <i>Naive Bayes Classifier</i> and <i>Support Vector Machine</i> [25]	Dampak ekonomi pandemi Covid-19 di RI.	255 twit Twitter/X.	Lexicon, SVM, & NBC	83% (NBC)	NBC mampu menangani <i>dataset</i> kecil dan mendeteksi keresahan ekonomi publik secara akurat.	Analisis sentimen dampak ekonomi di RI menggunakan NBC.	Berfokus pada dampak krisis kesehatan (Pandemi).
7	Cam <i>et al.</i> (2024): Sentiment analysis of financial Twitter posts on Twitter with the <i>machine learning</i> classifiers [22]	Sentimen finansial Twitter Bursa Istanbul.	17.189 twit finansial (Turki).	NB, SVM, MLP, KNN, DT	78,5% (NBC)	Memantau respon pasar modal terhadap opini digital untuk memprediksi volatilitas pasar.	Fokus pada klasifikasi sentimen data finansial dari media sosial.	Menggunakan <i>dataset</i> bahasa asing (Turki).

Lanjut ke halaman berikutnya

Tabel 2.1 – Lanjutan

No	Referensi	Fokus	Dataset	Metode	Hasil	Rangkuman Penelitian	Kesamaan	Perbedaan
8	Atimi <i>et al.</i> (2025): Enhancing the Accuracy of a Sentiment Analysis Model for Election Prediction Using a Hybrid Approach: An Experiment on Indonesian Tweets [9]	Prediksi pemilu di Indonesia di Twitter/X.	3.361 twit Twitter/X.	Hybrid NBC & IndoBERTweet	82%	Integrasi <i>contextual embeddings</i> dengan penalaran probabilistik NBC meningkatkan akurasi klasifikasi.	Implementasi NLP pada data Twitter bahasa Indonesia menggunakan NBC.	Berfokus pada domain politik (Pemilu).
9	Atmadji <i>et al.</i> (2025): The Use of Naïve Bayes Classifier in Sentiment Analysis at Indonesia's Super Priority Tourism Destinations Based on User Reviews [23]	Analisis ulasan super prioritas.	6.720 ulasan Google Maps.	Naïve Bayes & SMOTE	75%	Implementasi NBC untuk mendukung pengambilan keputusan kebijakan pariwisata berbasis data.	Penggunaan NBC untuk evaluasi program pemerintah Indonesia.	Menggunakan data ulasan tempat wisata, bukan teks kebijakan.
10	Nandini (2025): Sentiment Analysis of Economic Policy Comments on YouTube Using Ensemble Machine Learning [12]	Sentimen kebijakan ekonomi Tom Lembong.	1.029 komentar YouTube.	Ensemble ML (SVM variants)	88,25% (NBC)	Pendekatan augmentasi data meningkatkan performa klasifikasi sentimen kebijakan ekonomi.	Analisis respon publik terhadap figur pembuat kebijakan ekonomi.	Menggunakan sumber data YouTube dan metode <i>Ensemble</i> .

Lanjut ke halaman berikutnya

Tabel 2.1 – Lanjutan

No	Referensi	Fokus	Dataset	Metode	Hasil	Rangkuman Penelitian	Kesamaan	Perbedaan
11	Neogi <i>et al.</i> (2021): Sentiment analysis and classification of indian farmers' protest using twitter data [11]	Analisis sentimen protes petani di India.	18.000 twit Twitter/X.	<i>Naive Bayes</i> & <i>Random Forest</i>	89% Presisi	Mengklasifikasi opini publik terhadap gejolak kebijakan pemerintah menggunakan volume data besar.	Menggunakan platform Twitter dan algoritma <i>Naive Bayes</i> untuk isu sosial-ekonomi.	Objek penelitian adalah gejolak agraria di India.
12	Gunawan <i>et al.</i> (2023): Sentiment analysis of twitter towards the 2024 indonesian presidential candidates using the <i>naive bayes</i> algorithms [29]	Sentimen Capres RI 2024 di Twitter.	892 twit Twitter/X.	<i>Naive Bayes</i>	Akurasi 82,74%	NBC mampu menangkap pola dukungan pemilih terhadap kandidat presiden dengan presisi tinggi.	Implementasi <i>Naive Bayes</i> pada teks Twitter bahasa Indonesia untuk tokoh publik.	Berfokus pada domain politik lektoral.
13	Suhag & Pandya (2024): Comparative analysis of large language models and traditional methods for sentiment analysis of tweets <i>dataset</i> [34]	Komparasi LLM vs metode tradisional.	31.000 twit Twitter/X.	XLNet, KNN, RF, XGBoost	XLNet 99,54%	Menunjukkan keunggulan model bahasa besar (LLM) dibandingkan algoritma tradisional seperti KNN dan RF.	Evaluasi performa berbagai klasifikator pada <i>dataset</i> Twitter.	Menggunakan model XLNet sebagai metode utama.
14	Martiti & Juliane (2021): Implementation of <i>naive bayes</i> algorithm on sentiment analysis application [21]	Kepuasan layanan Grab Indonesia.	457 komentar Grab Indonesia.	<i>Naive Bayes</i>	Akurasi 86,66%	Membangun aplikasi klasifikasi otomatis untuk mengukur tingkat kepuasan layanan publik digital.	Menerapkan algoritma <i>Naive Bayes</i> untuk analisis teks sentimen bahasa Indonesia.	Berfokus pada ulasan kualitas layanan aplikasi komersial.

Analisis terhadap performa akurasi pada penelitian terdahulu menunjukkan pola yang konsisten di mana algoritma *Naïve Bayes Classifier* (NBC) memberikan hasil yang sangat kompetitif dan handal dalam klasifikasi teks bahasa Indonesia [19, 20, 21]. NBC secara konsisten memberikan hasil yang stabil di rentang 69% hingga 91% pada *dataset* media sosial yang bising dan informal [9, 19, 20, 29]. *Insight* krusial membuktikan bahwa performa NBC dapat dioptimalkan melalui teknik SMOTE untuk menangani *class imbalance* hingga mencapai 91,03% [20, 23, 27].

Secara metodologis, literatur menekankan bahwa kualitas tahapan *preprocessing* merupakan pondasi utama validitas klasifikasi [12, 21, 28, 35]. Tahapan sistematis seperti normalisasi sesuai KBBI dan *stemming* terbukti krusial mereduksi *noise* data informal [21, 23, 25, 36]. Penggunaan *TF-IDF* mendominasi ekstraksi fitur karena kemampuannya memberikan pembobotan kata secara kontekstual yang efektif untuk algoritma pembelajaran mesin klasik [11, 12, 23].

Analisis terhadap gap domain mengonfirmasi keterbatasan literatur yang mengintegrasikan sentimen digital dengan risiko kebijakan perdagangan universal [4, 6]. Mayoritas penelitian di Indonesia masih terfokus pada kebijakan domestik [19, 20, 21]. Belum ditemukan studi yang membedah respon publik terhadap ancaman tarif universal "Trump 2.0" yang diprediksi menyebabkan kerugian ekspor USD 1 miliar [4].

Penelitian ini mengambil posisi strategis mengisi celah tersebut menggunakan NBC untuk membedah dampak sentimen tarif Trump terhadap ekonomi Indonesia [6, 8, 22]. Kontribusi utama studi ini adalah pemanfaatan platform Twitter/X dan YouTube sebagai sensor emosional publik secara *real-time* terhadap kebijakan ekonomi global [8, 9, 10]. Dengan menyinergikan teori guncangan TPU dan *Naïve Bayes*, penelitian ini memberikan wawasan berbasis data bagi strategi ketahanan ekonomi nasional [2, 4].