

ANALISIS SENTIMEN TWEET PERANG DAGANG TRUMP-CHINA MENGUNAKAN METODE SVM BERBASIS SMOTE

Heriyanto¹, Fachri Amsury², Alisya Shafa Salsabila³

^{1,2,3}Universitas Bina Sarana Informatika, Indonesia

Penulis Korespondensi: heriyanto.hio@bsi.ac.id.

ABSTRAK

Perang dagang antara Amerika Serikat dan China pada masa pemerintahan Donald Trump memberikan dampak ekonomi global, termasuk bagi Indonesia, dan menjadi topik hangat di media sosial X. Penelitian ini bertujuan mengklasifikasi sentimen masyarakat Indonesia terhadap isu perang tarif tersebut serta mengevaluasi kinerja algoritma Support Vector Machine (SVM) yang dikombinasikan dengan Synthetic Minority Oversampling Technique (SMOTE) dalam menangani ketidakseimbangan data. Pendekatan Knowledge Discovery in Databases (KDD) digunakan dengan mengumpulkan 1.021 tweet berbahasa Indonesia periode Februari–Mei 2025 menggunakan kata kunci “tarif trump indonesia” dan “perang dagang amerika china indonesia”. Setelah preprocessing dan pelabelan, diperoleh 595 tweet negatif dan 426 positif. Evaluasi menunjukkan bahwa model SVM tanpa SMOTE menghasilkan akurasi 84,39%, presisi 85,37%, dan AUC 0,9102. Setelah penerapan SMOTE, kinerja meningkat dengan akurasi 85,71%, presisi 95,74%, dan AUC 0,9449, meskipun recall sedikit menurun. F1-score juga meningkat dari 81,38% menjadi 84,38%. Hasil ini menunjukkan bahwa penerapan SMOTE efektif meningkatkan performa SVM pada data tidak seimbang dan bahwa sentimen publik Indonesia terhadap isu perang tarif cenderung negatif.

Kata Kunci: *Klasifikasi Sentimen, Support Vector Machine (SVM), SMOTE, Perang Dagang, Media Sosial X*

Riwayat Artikel :

Tanggal diterima : 04-06-2026

Tanggal terbit : 01-07-2026

Kutipan :

Heriyanto, Fachri Amsury, & Salsabila, A. S. (2026). ANALISIS SENTIMEN TWEET PERANG DAGANG TRUMP-CHINA MENGGUNAKAN METODE SVM BERBASIS SMOTE. *INFOTECH Journal*, 12(1), 158–163. <https://doi.org/10.31949/infotech.v12i1.18510>

1. PENDAHULUAN

Perang dagang antara Amerika Serikat dan China yang dimulai pada tahun 2018 telah menjadi salah satu konflik ekonomi paling signifikan dalam sejarah perdagangan internasional modern sehingga memberikan dampak yang besar bagi perdagangan (Ratna Sari & Khaldun, 2023). Ketegangan ini dipicu oleh kebijakan yang diambil oleh Presiden Donald Trump, yang menaikkan tarif impor terhadap produk-produk asal China sebagai upaya untuk mengurangi defisit perdagangan dan menekan dominasi ekonomi Tiongkok di pasar global (A.Nurmamurti et al., 2022). Sebagai respons atas tindakan tersebut, China juga memberlakukan tarif balasan terhadap produk-produk asal Amerika Serikat. Konflik ini tidak hanya melibatkan dua negara adidaya, tetapi juga menimbulkan efek domino yang signifikan terhadap ekonomi negara-negara lain termasuk Indonesia (Cenderawasih et al., 2022).

Indonesia, sebagai negara berkembang yang memiliki ketergantungan yang cukup tinggi terhadap perdagangan internasional, turut merasakan dampak dari ketegangan yang sedang berlangsung. Beberapa sektor industri, seperti elektronik, manufaktur, dan ekspor bahan mentah, mengalami tekanan akibat ketidakpastian yang melanda pasar global (Rahmi & Sari, 2024). Selain dampak ekonomi yang bersifat langsung, konflik ini juga memicu diskusi publik yang luas di kalangan masyarakat Indonesia, khususnya di media sosial seperti X yang sebelumnya dikenal sebagai Twitter (A. Sitanggang et al., 2024). Platform tersebut menjadi ruang virtual yang sangat aktif untuk menampung opini, kritik, dan diskusi publik mengenai isu-isu internasional, termasuk perang tarif antara Amerika Serikat dan China (Wilantari & Bawono, 2021).

X bukan hanya media komunikasi, tetapi juga sumber data sosial yang dapat dipergunakan dalam banyak hal seperti politik, text mining, data analisis, analisis sentimen dan prediksi (E. D. Sitanggang et al., 2023). Hal ini menunjukkan bahwa X tidak hanya mencerminkan opini publik, tetapi juga mampu memengaruhi wacana politik dan kebijakan global (Wijaya et al., 2022). Oleh karena itu, dalam penelitian ini, analisis sentimen pengguna X terhadap perang tarif antara Trump dan China menjadi relevan untuk memahami respons dan persepsi publik terhadap isu tersebut serta dampak pada kehidupan mereka secara langsung maupun tidak langsung (Andriawan & Ernawati, 2024).

Untuk memperoleh informasi yang bernilai dari data sosial yang sangat besar dan tidak terstruktur tersebut, dibutuhkan suatu pendekatan sistematis. Pendekatan yang digunakan untuk menggali pola, tren, dan sentimen dari data semacam ini adalah melalui Knowledge Discovery in Databases (KDD) (Amsury et al., 2023). Pemilihan KDD sebagai landasan metodologis penelitian ini didasarkan pada sifatnya sebagai proses yang bertujuan untuk menggali dan menganalisis sejumlah besar

himpunan data dan mengekstrak informasi serta pengetahuan yang berguna (Ependi & Putra, 2019). Pendekatan ini memastikan bahwa seluruh alur penelitian, dari data mentah hingga penemuan sentimen, tercakup secara terstruktur. Dalam konteks penemuan pengetahuan tersebut, Support Vector Machine (SVM) dipilih sebagai algoritma klasifikasi utama. Pemilihan metode SVM dalam penelitian ini didasarkan pada keunggulannya dalam mengklasifikasikan data teks dengan akurasi tinggi (Ariansyah & Kusmira, 2021). Beberapa studi telah menunjukkan bahwa SVM mampu mengungguli Naïve Bayes dalam analisis sentimen. Penelitian oleh Atmajaya yang membandingkan SVM dan Naive Bayes dalam analisis sentimen pengguna Twitter terhadap ChatGPT, menunjukkan bahwa SVM mencapai akurasi 59% (menggunakan label Vader), secara signifikan melampaui Naive Bayes yang hanya sebesar 47% pada konteks yang sama (Atmajaya et al., 2023). Selain itu, penelitian oleh Zain yang membandingkan SVM dengan Naive Bayes dan K-Nearest Neighbors (KNN) dalam analisis sentimen tweet terkait calon presiden 2024, menunjukkan bahwa SVM memberikan kinerja superior dengan akurasi total sebesar 0.88, serta konsistensi tinggi dalam precision dan recall untuk semua kategori sentiment (Zain et al., 2025). Hal ini semakin menegaskan potensi SVM sebagai metode yang efektif dan kuat dalam analisis sentimen di X.

Meskipun SVM efektif, dataset sentimen sering kali dihadapkan pada masalah ketidakseimbangan kelas, di mana jumlah kelas data dengan sentimen tertentu (misalnya, positif atau negatif) jauh lebih banyak dibandingkan sentimen lainnya. Ketidakseimbangan ini dapat menyebabkan model klasifikasi menjadi bias dan kurang akurat dalam memprediksi kelas minoritas (Amsury et al., 2020). Untuk mengatasi permasalahan ini, penelitian ini akan mengintegrasikan teknik SMOTE (Synthetic Minority Over-sampling Technique). SMOTE merupakan prinsip oversampling yaitu menambah data dari kelas minor agar jumlahnya seimbang dengan data kelas mayor (Mansurifar & Shi, 2020). Teknik ini relevan karena kemampuannya untuk menyeimbangkan distribusi data dan memungkinkan algoritma SVM untuk belajar lebih efektif dari semua kelas pada dataset yang mengalami distribusi kelas tidak seimbang (Ariansyah & Kusmira, 2021).

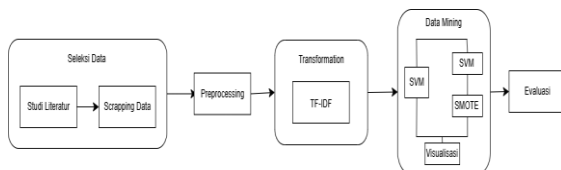
Dengan menggabungkan kekuatan data X dan pendekatan machine learning berbasis KDD, serta memanfaatkan efektivitas SVM dalam klasifikasi dan Ketidakseimbangan kelas merupakan tantangan utama dalam membangun model prediktif pada penelitian ini maka perlu menerapkan teknik SMOTE dalam penanganan ketidakseimbangan data, penelitian ini bertujuan untuk mengidentifikasi dan mengklasifikasikan opini publik Indonesia terhadap kebijakan tarif Amerika terhadap China. Melalui analisis sentimen ini, diharapkan penelitian dapat memberikan kontribusi nyata bagi pemahaman sosial masyarakat digital Indonesia dalam

menghadapi perang tarif Trump dan China. Dengan demikian, penelitian ini tidak hanya relevan secara ilmiah, tetapi juga secara praktis, dalam memahami persepsi publik terhadap fenomena global yang kompleks dan dinamis, serta memberikan ruang eksplorasi baru dalam penggabungan antara ilmu sosial, teknologi informasi, dan hubungan internasional melalui pendekatan data-driven yang kini.

Keburuan penelitian ini bukan sekedar pada penggabungan dua metode, melainkan pada pembuktian empiris bagaimana generasi sampel sintesis SMOTE mampu mengoreksi bias hyperplane SVM pada dataset kebijakan tarif Amerika terhadap China yang memiliki tingkat ketidakseimbangan, Urgensi karya ini terletak pada kebutuhan mendesak untuk meningkatkan sensitivitas Recall deteksi kelas minoritas tanpa mengorbankan presisi, yang mana kegagalan deteksi pada kasus ini akan berimplikasi pada algoritma SVM yang menjadi bias, sehingga gagal mengenali pola penting pada data minoritas. Dengan demikian, kontribusi penelitian ini memberikan perspektif baru dan solusi yang lebih andal untuk permasalahan klasifikasi tidak seimbang pada dataset kebijakan tarif Amerika terhadap China.

2. METODE

Metodologi yang digunakan dalam penelitian ini menggunakan pendekatan Knowledge Discovery in Database (KDD), yaitu proses sistematis untuk menemukan pengetahuan atau pola dari kumpulan data besar. Adapun proses Knowledge Discovery in Database (KDD) (Pranata & Utomo, 2020), sebagai berikut



Gambar 1. Tahapan penelitian metode KDD

1. Selection

Pengumpulan data merupakan langkah awal dalam proses data mining, yang bertujuan untuk memperoleh informasi mentah dari sumber yang relevan. Pada tahapan ini, data diambil menggunakan alat bantu digital dengan bahasa pemrograman tertentu dan disimpan dalam format yang mudah diolah, seperti XLSX. Data yang dikumpulkan biasanya berkaitan dengan topik atau isu tertentu, sesuai dengan tujuan penelitian yang dilakukan.

2. Preprocessing

Tahap ini berfungsi untuk menyiapkan data agar layak digunakan dalam proses analisis. Pembersihan data mencakup penghapusan data kosong, duplikat, ketidaksesuaian format, serta unsur-unsur yang tidak relevan seperti retweet, mention, hashtag, atau tautan. Setelah dibersihkan, data kemudian melalui proses transformasi dasar seperti pengubahan huruf,

pemisahan kata, dan penyaringan elemen umum. Tujuannya adalah untuk menyederhanakan data sebelum masuk ke tahap berikutnya.

3. Transformation

Transformasi data adalah proses mengonversi data mentah menjadi bentuk numerik atau representasi tertentu yang dapat dikenali oleh algoritma analisis. Tahap ini membantu sistem dalam mengenali pola atau struktur dalam data teks, sehingga mempermudah proses klasifikasi atau prediksi yang akan dilakukan di tahap selanjutnya.

4. Pengolahan Data (Data Mining)

Pengolahan data merupakan inti dari tahapan data mining, di mana data yang sudah bersih dan terstruktur dianalisis untuk mendapatkan informasi bermakna. Analisis ini bertujuan untuk mengidentifikasi pola atau kecenderungan tertentu dalam data, seperti kecenderungan opini atau sentimen dalam suatu isu. Dalam tahap ini, pendekatan analisis dilakukan dengan algoritma tertentu sesuai kebutuhan dan tujuan penelitian.

5. Evaluation/Interpretation

Tahapan ini bertujuan untuk menilai hasil dari proses analisis, apakah sesuai dengan harapan dan akurat dalam menjawab permasalahan penelitian. Evaluasi dilakukan menggunakan tolok ukur tertentu seperti tingkat akurasi dan ketepatan hasil. Hasil evaluasi inilah yang kemudian menjadi dasar untuk menarik kesimpulan atau interpretasi terhadap fenomena yang sedang diteliti.

3. PEMBAHASAN

Penelitian ini menganalisis sentimen pengguna platform X (Twitter) di Indonesia terhadap isu perang tarif antara Amerika Serikat dan China pada masa pemerintahan Donald Trump, , menggunakan pendekatan Knowledge Discovery in Database (KDD) yang terdiri dari lima tahapan utama, yaitu selection, preprocessing, transformation, data mining, dan evaluation.

1. Selection

Tahap selection berfokus pada proses pengumpulan data mentah dari sumber relevan. Data dikumpulkan melalui proses scraping menggunakan alat Tweet-Harvest selama periode Februari–Mei 2025 dengan kata kunci seperti “Trump China Tariff”, “perang dagang Amerika China”, dan “Trump trade war” yang menghasilkan menghasilkan 1.021 tweet berbahasa Indonesia yang menjadi dataset utama penelitian.

2. Preprocessing

Tahap preprocessing dilakukan untuk mempersiapkan data agar dapat diolah secara komputasional dan mengurangi noise pada teks. Proses ini meliputi pembersihan teks, case folding, tokenizing, stopword removal, dan stemming.

Pada proses pembersihan data, elemen yang tidak relevan dari teks komentar twitter dihapuskan.

Contoh pembersihan tersebut digambarkan pada Tabel 1, yang menunjukkan elemen yang dihapus.

Tabel 1. Preprocessing

Clean data Twitter	Hasil
Hapus Mention @username	@prabowo Pak Vietnam sudah GerCep nego total dengan native Amerika... Ntar lagi Vietnam tidak kena tarif oleh si Trump alias tarifnya 0%. Kapan Indonesia GerCep nya... #masa kalah dgn Vietnam. Biar kita gak kena imbasnya perang dagangnya si Trump.... https://t.co/mR6w0IvAL3
Hapus hashtag	Pak Vietnam sudah GerCep nego total dengan native Amerika... Ntar lagi Vietnam tidak kena tarif oleh si Trump alias tarifnya 0%. Kapan Indonesia GerCep nya... #masa kalah dgn Vietnam. Biar kita gak kena imbasnya perang dagangnya si Trump.... https://t.co/mR6w0IvAL3
Hapus link (http/https)	Pak Vietnam sudah GerCep nego total dengan native Amerika... Ntar lagi Vietnam tidak kena tarif oleh si Trump alias tarifnya 0%. Kapan Indonesia GerCep nya... masa kalah dgn Vietnam. Biar kita gak kena imbasnya perang dagangnya si Trump.... https://t.co/mR6w0IvAL3
Hapus emoji	Pak Vietnam sudah GerCep nego total dengan native Amerika... Ntar lagi Vietnam tidak kena tarif oleh si Trump alias tarifnya 0%. Kapan Indonesia GerCep nya... masa kalah dgn Vietnam. Biar kita gak kena imbasnya perang dagangnya si

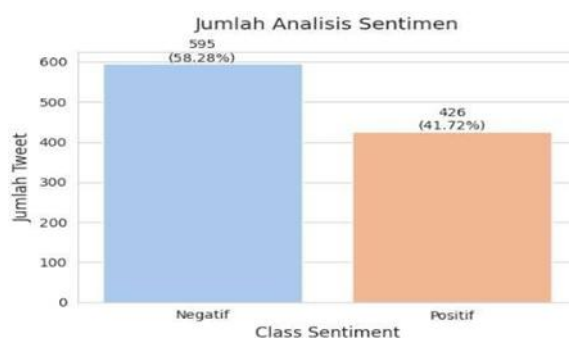
Clean data Twitter	Hasil
	Trump....
Hapus karakter aneh (non-alfanumerik, selain spasi)	Pak Vietnam sudah GerCep nego total dengan native Amerika... Ntar lagi Vietnam tidak kena tarif oleh si

Setelah pembersihan teks, dilakukan tahapan preprocessing selanjutnya seperti case folding, tokenizing, stopwords removal, dan stemming yang hasilnya dapat dilihat pada table 2.

Tabel 2. Hasil Preprocessing

Clean data Twitter	Hasil
Pembersihan teks	Vietnam sudah GerCep nego total dengan Amerika Ntar lagi Vietnam tidak kena tarif oleh Trump alias tarifnya Kapan Indonesia GerCep masa kalah Vietnam Biar kita kena imbasnya perang dagangnya Trump
Casefolding	vietnam sudah gercep nego total dengan amerika ntar lagi vietnam tidak kena tarif oleh trump alias tarifnya kapan indonesia gercep masa kalah vietnam biar kita kena imbasnya perang dagangnya trump
Normalisasi	vietnam sudah gerak cepat nego total dengan amerika entar lagi vietnam tidak kena tarif oleh trump alias tarifnya kapan indonesia gerak cepat masa kalah vietnam biar kita kena imbasnya perang dagangnya trump
Tokenisasi	['vietnam', 'sudah', 'gerak', 'cepat', 'nego', 'total', 'dengan', 'amerika', 'entar', 'lagi', 'vietnam', 'tidak', 'kena', 'tarif', 'oleh', 'trump', 'alias', 'tarifnya', 'kapan', 'indonesia', 'gerak', 'cepat', 'masa', 'kalah', 'vietnam', 'biar', 'kita', 'kena', 'imbasnya',

Clean data Twitter	Hasil
	'perang', 'dagangnya', 'trump']
Stopword	['vietnam', 'gerak', 'cepat', 'nego', 'total', 'amerika', 'entar', 'vietnam', 'kena', 'tarif', 'trump', 'alias', 'tarifnya', 'indonesia', 'gerak', 'cepat', 'kalah', 'vietnam', 'biar', 'kena', 'imbasnya', 'perang', 'dagangnya', 'trump']
Stemming	vietnam gerak cepat nego total amerika entar vietnam kena tarif trump alias tarif indonesia gerak cepat kalah vietnam biar kena imbas perang dagang trump



Gambar 2. Hasil Pelabelan Sentimen

Seperti pada gambar 2, Distribusi hasil pelabelan menunjukkan ketidakseimbangan data, dengan 595 tweet (58,28%) berlabel negatif dan 426 tweet (41,72%) berlabel positif. Hal ini menunjukkan bahwa persepsi publik Indonesia terhadap perang tarif Trump–China cenderung negatif.

Untuk memperjelas representasi opini publik, dilakukan visualisasi kata dominan dalam setiap kategori sentimen menggunakan wordcloud. Hasil menunjukkan bahwa pada sentimen positif, kata-kata seperti “peluang”, “ekonomi”, dan “kerjasama” sering muncul, sedangkan sentimen negatif didominasi oleh kata “krisis”, “inflasi”, dan “penurunan”.



Gambar 3. Wordcloud Sentimen Positif



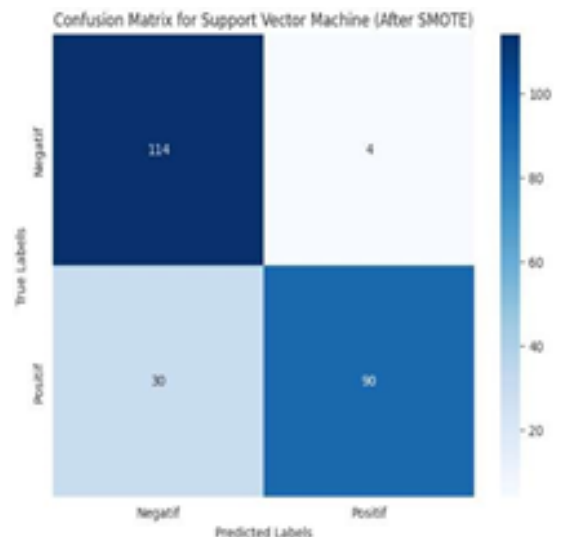
Gambar 4. Wordcloud Sentimen Negatif

4. Data Mining

Proses data mining dilakukan menggunakan algoritma Support Vector Machine (SVM) untuk klasifikasi sentimen, dan Synthetic Minority Over-sampling Technique (SMOTE) untuk menyeimbangkan dataset yang tidak seimbang. Model dilatih dengan rasio 80:20 untuk data latih dan uji. Sebelum penerapan SMOTE, model menghasilkan akurasi sebesar 84,39%, precision 85,37%, recall 77,78%, dan F1-score 81,38% dengan nilai AUC 0,9102.

Setelah penerapan Synthetic Minority Over-sampling Technique (SMOTE) untuk menyeimbangkan data, performa model meningkat menjadi 85,71% akurasi, precision 95,74%, recall 75,00%, F1-score 84,38%, dan AUC 0,9449.

Perbandingan metrik ini divisualisasikan pada Tabel 2, yang menunjukkan peningkatan nilai akurasi dan AUC setelah penerapan SMOTE.

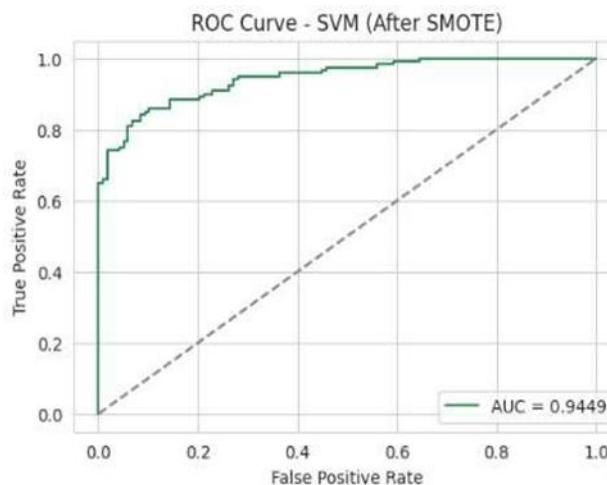


Gambar 5. Confusion Matrixl SVM Sesudah SMOTE

Selain itu, hasil confusion matrix pada Gambar 5 menunjukkan bahwa setelah penerapan SMOTE, model lebih seimbang dalam mengklasifikasikan data positif dan negatif, dengan pengurangan kesalahan klasifikasi pada kelas minoritas.

Sebanyak 114 tweet dengan sentimen negatif berhasil diprediksi dengan benar sebagai

negatif (True Negative). Sebanyak 4 tweet dengan sentimen negatif salah diprediksi sebagai positif (False Positive). Sebanyak 30 tweet dengan sentimen positif salah diprediksi sebagai negatif (False Negative). Sebanyak 90 tweet dengan sentimen positif berhasil diprediksi dengan benar sebagai positif (True Positive).



Gambar 6. Kurva ROC SVM dengan SMOTE

Pada gambar 6 memperlihatkan bahwa model memiliki AUC (Area Under the Curve) sebesar 0.9449, yang mengindikasikan bahwa model memiliki kemampuan yang sangat baik dalam membedakan antara kelas Positif dan Negatif. Kurva yang berada jauh di atas garis diagonal baseline menegaskan bahwa model memiliki performa klasifikasi yang kuat. Secara keseluruhan, penerapan SMOTE mampu membantu mengatasi ketidakseimbangan kelas pada data latih dan meningkatkan kemampuan model dalam melakukan prediksi secara lebih adil terhadap kedua label kelas.

Tabel 3. Perbandingan Kinerja Model SVM Sebelum dan Sesudah SMOTE

Metrik Evaluasi	SVM Sebelum SMOTE	SVM Sesudah SMOTE
Akurasi	84.39 %	85.71%
Precision	85.37%	95.74%
Recall	77.78%	75.00%
F1-score	81,38%	84.38%
True Positive (TP)	70	90
True Negative (TN)	103	114
False Positive (FP)	12	4
False Negative (FN)	20	30
AUC (ROC)	0.9102	0.9449

Setelah dilakukan perbandingan antara model SVM sebelum dan sesudah penerapan SMOTE, terlihat bahwa performa model mengalami peningkatan pada tabel 3. Akurasi model meningkat dari 84,39% menjadi 85,71%, dengan kenaikan signifikan pada precision kelas positif dari 85,37% menjadi 95,74%. Meskipun recall sedikit menurun dari 77,78% menjadi 75,00% ini disebabkan efek samping dari pergeseran decision boundary, karena setelah penerapan teknik SMOTE algoritma (SVM)

sekarang lebih ketat/conservative dalam memprediksi kelas positif (untuk menekan False Positive), beberapa data positif yang berada di area abu-abu (borderline) atau noise yang dulunya tertangkap, sekarang terlewat dan diklasifikasikan sebagai negatif (False Negative naik sedikit), nilai F1-score justru naik dari 81,38% menjadi 84,38%. Selain itu, nilai AUC juga meningkat dari 0,9102 menjadi 0,9449, yang menunjukkan bahwa model semakin baik dalam membedakan antara kelas positif dan negatif. Hal ini membuktikan bahwa penerapan SMOTE efektif dalam mengatasi ketidakseimbangan data dan meningkatkan kinerja klasifikasi model.

4. KESIMPULAN

Berdasarkan penelitian yang dilakukan untuk analisis sentimen terhadap tweet pengguna platform X di Indonesia terkait isu perang tarif antara Donald Trump dan China dengan metode SVM dan SMOTE maka dapat disimpulkan beberapa hal berikut, distribusi sentimen pengguna terhadap isu perang tarif Trump dan China didominasi oleh sentimen negatif, dengan Sentimen Negatif mendominasi dengan 595 tweet (58,28%), sedangkan sentimen Positif hanya sebanyak 426 tweet (41,72%).

Kinerja Model SVM tanpa SMOTE menghasilkan akurasi sebesar 84,39% dengan f1-score 81,38%, namun cenderung bias terhadap kelas mayoritas karena recall untuk sentimen positif lebih rendah. Setelah diterapkan SMOTE, akurasi meningkat menjadi 85,71% dan f1-score menjadi 84,10%, dengan precision kelas positif mencapai 95,74%. Hal ini menunjukkan bahwa SMOTE efektif dalam menyeimbangkan data dan meningkatkan performa klasifikasi, khususnya untuk sentimen positif. Nilai AUC naik dari 0,9102 menjadi 0,9449, menunjukkan peningkatan kemampuan model dalam membedakan dua kelas.

Berdasarkan hasil penelitian tersebut dapat disimpulkan bahwa Algoritma Support Vector Machine (SVM) terbukti mampu melakukan klasifikasi sentimen dengan cukup baik dan penerapan SMOTE menunjukkan hasil yang baik dalam peningkatan performa model. Oleh karena itu, Hipotesis nol (H_0) ditolak, dan hipotesis alternatif (H_1) diterima, yang berarti bahwa terdapat perbedaan yang signifikan dalam distribusi dan hasil klasifikasi sentimen antara sebelum dan sesudah penerapan SMOTE pada data tweet terkait perang tarif Trump dan China.

Berdasarkan hasil penelitian yang telah dilakukan, masih terdapat sejumlah keterbatasan. Oleh karena itu, peneliti berharap agar penelitian ini dapat dikembangkan lebih lanjut di masa mendatang. Beberapa saran yang dapat penulis berikan antara lain:

Penggunaan teknik balancing data seperti SMOTE terbukti efektif dalam meningkatkan

performa model pada data yang tidak seimbang. Oleh karena itu, disarankan bagi peneliti selanjutnya untuk selalu memperhatikan distribusi kelas sebelum melakukan klasifikasi, dan mempertimbangkan penggunaan teknik balancing lainnya seperti ADASYN atau Borderline-SMOTE sebagai perbandingan.

Penelitian ini menggunakan algoritma Support Vector Machine (SVM) sebagai metode klasifikasi. Untuk penelitian mendatang, disarankan untuk membandingkan performa beberapa algoritma lain seperti Random Forest, XGBoost, atau LSTM agar diperoleh model klasifikasi yang lebih optimal untuk data teks dari media sosial.

Dalam tahap preprocessing, penelitian ini hanya menggunakan kamus normalisasi dan stopword standar. Disarankan agar peneliti berikutnya mempertimbangkan penggunaan stopword tambahan dan library sastrawi, untuk menangani kata tidak baku dan istilah populer dalam bahasa Indonesia di kalangan pengguna X.

PUSTAKA

- A.Nurmamurti, R., Faradilla, A. Y., Rismanto, A. I., Afifah, S. N., Hamida, A., & Sari, K. H. (2022). *Analisis Kebijakan Luar Negeri Trump : Studi Kasus*. 2(1), 62–70.
- Amsury, F., Kurniawati, I., & Rizki Fahdia, M. (2023). Implementasi Association Rules Menentukan Pola Pemilihan Menu Di the Gade Coffee & Gold Menggunakan Algoritma Apriori. *INFOTECH Journal*, 9(1), 279–286. <https://doi.org/10.31949/infotech.v9i1.5357>
- Amsury, F., Ruhjana, N., Saputra, I., & Sulistyowati, D. N. (2020). Classification of Customer Complaints on Instagram Comments Using Naïve Bayes Algorithm With N-Gram Feature Extension. *Jurnal Techno Nusa Mandiri*, 17(2), 109–116. <https://doi.org/10.33480/techno.v17i2.1632>
- Andriawan, M. G., & Ernawati, T. (2024). *PENGUNAAN ALGORITMA NAÏVE BAYES DAN SUPPORT VECTOR MACHINE UNTUK ANALISIS SENTIMEN KONFLIK*. 12(3), 3222–3230.
- Ariansyah, A., & Kusmira, M. (2021). Analisis Sentimen Pengaruh Pembelajaran Daring Terhadap Motivasi Belajar Di Masa Pandemi Menggunakan Naive Bayes Dan Svm. *Faktor Exacta*, 14(3), 100. <https://doi.org/10.30998/faktorexacta.v14i3.10325>
- Atmajaya, D., Febrianti, A., & Darwis, H. (2023). *Indonesian Journal of Computer Science*. 12(1), 2173–2181.
- Cenderawasih, U., Menufandu, D. N., & Cenderawasih, U. (2022). *DAMPAK PERANG DAGANG TERHADAP NERACA PERDAGANGAN AMERIKA SERIKAT-CHINA*. 2(4), 627–636. <https://doi.org/10.53866/jimi.v2i4.175>
- Ependi, U., & Putra, A. (2019). Solusi Prediksi Persediaan Barang dengan Menggunakan Algoritma Apriori (Studi Kasus: Regional Part Depo Auto 2000 Palembang). *Jurnal Edukasi Dan Penelitian Informatika (JEPIN)*, 5(2), 139. <https://doi.org/10.26418/jp.v5i2.32648>
- Mansourifar, H., & Shi, W. (2020). Deep synthetic minority over-sampling technique. *ArXiv*, 16, 321–357.
- Pranata, B. S., & Utomo, D. P. (2020). Penerapan Data Mining Algoritma FP-Growth Untuk Persediaan Sparepart Pada Bengkel Motor (Study Kasus Bengkel Sinar Service). *Bulletin of Information Technology (BIT)*, 1(2), 83–91.
- Rahmi, C., & Sari, A. E. (2024). *Dampak Perang Dagang Amerika Serikat Dengan China Terhadap Ekonomi Indonesia Studi Kasus : Dalam bidang Ekspor Kakao*. 1(3), 580–591.
- Ratna Sari, A. I., & Khaldun, R. I. (2023). *Retaliiasi China terhadap Amerika Serikat dalam Konteks Perang Dagang*. 3(2).
- Sitanggang, A., Umaidah, Y., Adam, R. I., Karawang, U. S., & Timur, T. (2024). *ANALISIS SENTIMEN MASYARAKAT TERHADAP PROGRAM MAKAN SIANG GRATIS PADA MEDIA*. 12(3).
- Sitanggang, E. D., Pinem, A., Perangin-angin, J., Sembiring, M., & Saroha Simanjuntak. (2023). Pembangunan dan Pelatihan Penggunaan Website SMK Swasta Teknik Dairi. *ULINA: Jurnal Pengabdian Kepada Masyarakat*, 1(1), 23–27. <https://doi.org/10.58918/ulina.v1i1.191>
- Wijaya, N., Irsyad, H., & Taqwiym, A. (2022). Pelatihan Pemanfaatan Canva Dalam Mendesain Poster. *Fordicate*, 1(2), 192–199. <https://doi.org/10.35957/fordicate.v1i2.2418>
- Wilantari, R. N., & Bawono, S. (2021). *TANTANGAN DOMINASI AMERIKA SERIKAT OLEH TIONGKOK DALAM PERANG DAGANG*. 13(1), 38–42.
- Zain, A. F., Azies, H. Al, & Ananda, I. K. (2025). *Analisis Sentimen Ulasan Pengguna iPhone dengan Pendekatan Hibrida*. 1039–1049. <https://doi.org/10.33364/algoritma/v.22-1.2277>