

ANALISIS SENTIMEN PUBLIK TERHADAP DANANTARA DI MEDIA SOSIAL X MENGGUNAKAN NLP DAN PEMBELAJARAN MESIN

Muhamad Fazri¹⁾, Apriade Voutama²⁾

¹Prodi Sistem Informasi, Fakultas Ilmu Komputer, Universitas Singaperbangsa Karawang

²Universitas Singaperbangsa Karawang

email: 1muhamadfazry758@gmail.com, 2apriade.voutama@staff.unsika.ac.id

Abstract

Danantara is a national issue currently gaining significant public attention on social media platform X (formerly Twitter) due to its links to strategic projects and broad socio-economic impacts. This study aims to analyze public sentiment toward Danantara using a Natural Language Processing (NLP) approach based on semi-supervised learning, combining K-Means Clustering for automatic labeling and classification using Support Vector Machine (SVM) and Bidirectional Encoder Representations from Transformers (BERT). A total of 2,009 Indonesian-language tweets were collected and preprocessed through cleaning, tokenization, stopword removal, and stemming. Sentiment labeling showed a highly imbalanced distribution: 87.0% negative, 3.6% neutral, and 9.4% positive. A word cloud visualization revealed dominant critical terms such as “bumn,” “corruption,” and “investment.” Model performance evaluation indicated that SVM achieved the highest accuracy at 98%, while BERT reached only 25%. BERT’s low performance was attributed to the small dataset size, class imbalance, and the informal nature of the tweet language, which challenged contextual understanding. These findings suggest that the TF-IDF + SVM combination is more effective for sentiment analysis of short, informal Indonesian texts in low-resource settings. The study’s main contribution lies in applying a semi-supervised approach to analyze a local policy issue using a novel combination of unsupervised and supervised techniques rarely explored in Indonesian NLP research.

Keywords: *Danantara, Sentiment Analysis, Social Media, K-Means, SVM*

Abstrak

Danantara merupakan isu nasional yang tengah menjadi sorotan publik di media sosial X (sebelumnya Twitter) karena keterkaitannya dengan proyek strategis dan dampak sosial-ekonomi yang luas. Penelitian ini bertujuan untuk menganalisis sentimen publik terhadap Danantara menggunakan pendekatan Natural Language Processing (NLP) berbasis semi-supervised learning dengan menggabungkan metode K-Means Clustering untuk pelabelan otomatis dan model klasifikasi Support Vector Machine (SVM) serta Bidirectional Encoder Representations from Transformers (BERT). Sebanyak 2.009 tweet berbahasa Indonesia dikumpulkan dan diproses melalui tahapan preprocessing yang mencakup pembersihan, tokenisasi, stopword removal, dan stemming. Hasil pelabelan menunjukkan distribusi sentimen yang tidak seimbang, dengan dominasi negatif sebesar 87,0%, netral 3,6%, dan positif 9,4%. Visualisasi word cloud memperlihatkan kata-kata dominan bernuansa kritik seperti “bumn”, “korupsi”, dan “investasi”. Evaluasi performa model menunjukkan bahwa SVM memberikan akurasi tertinggi sebesar 98%, sedangkan BERT hanya mencapai 25%. Rendahnya performa BERT disebabkan oleh keterbatasan jumlah data, distribusi kelas yang tidak seimbang, serta kompleksitas bahasa informal dalam tweet. Temuan ini menunjukkan bahwa pendekatan TF-IDF + SVM lebih tepat digunakan untuk analisis opini publik di media sosial berbahasa Indonesia, terutama ketika data terbatas dan tidak berlabel. Kontribusi utama penelitian ini terletak pada penerapan pendekatan semi-supervised untuk menganalisis isu lokal menggunakan kombinasi teknik unsupervised dan supervised learning yang belum banyak dieksplorasi dalam konteks NLP Indonesia.

Kata Kunci: *Danantara, Analisis Sentimen, Media Sosial, K-Means, SVM*

1. PENDAHULUAN

Perkembangan teknologi digital telah mengubah cara masyarakat dalam menyampaikan opini dan menanggapi berbagai isu sosial, ekonomi, dan politik (Ardiansyah & others, 2024). Media sosial menjadi ruang publik yang dinamis, di mana opini dapat terbentuk, berkembang, dan menyebar secara cepat. Salah satu platform yang paling berpengaruh adalah X (sebelumnya Twitter), yang memungkinkan pengguna untuk menyampaikan pandangan secara real-time dalam format cuitan singkat (Paramita & Ibrahim, 2023). Platform

<https://doi.org/10.35145/joisie.v9i1.4924>

JOISIE licensed under a Creative Commons Attribution-ShareAlike 4.0 International License (CC BY-SA 4.0)

ini memiliki potensi besar untuk menangkap kecenderungan opini publik terhadap isu-isu strategis seperti proyek Danantara, yang sedang menjadi perhatian nasional (Putri & Pratama, 2022).

Dalam konteks ini, analisis sentimen menjadi pendekatan yang krusial untuk memahami kecenderungan opini publik terhadap suatu isu strategis. Salah satu teknologi utama yang digunakan adalah Natural Language Processing (NLP), yakni cabang dari kecerdasan buatan yang memungkinkan komputer untuk memproses, memahami, dan mengklasifikasikan teks secara otomatis. NLP sangat efektif digunakan pada konten berformat pendek dan informal seperti cuitan di media social (Syafrianto, 2022). Proses ini mencakup analisis struktur kalimat dan makna kata untuk mendeteksi emosi atau sikap pengguna secara sistematis (Tanggraeni & Sitokdana, 2022).

Penerapan NLP dalam konteks sosial digital di Indonesia juga terus berkembang. Salah satu inisiatif penting adalah pengembangan benchmark nasional IndoNLU yang diperkenalkan oleh (Aji & others, 2021). Benchmark ini digunakan untuk mengevaluasi kinerja berbagai model NLP yang dirancang khusus untuk Bahasa Indonesia, serta memperluas cakupan riset NLP yang kontekstual dan relevan. Sementara itu, studi terbaru menunjukkan bahwa NLP juga mampu digunakan untuk menangkap dinamika sentimen publik terhadap isu ekonomi dan social (Permana et al., 2023).

Berbagai studi sebelumnya telah memanfaatkan NLP untuk analisis sentimen dalam konteks komersial. (Saputra et al., 2023) mengkaji opini pengguna Telegram menggunakan Naïve Bayes, dengan akurasi mencapai 85,31%. (Agustina et al., 2022) meneliti ulasan produk di Shopee dan menunjukkan bahwa sentimen pelanggan dapat memberikan insight terhadap layanan. Sementara itu, (Ibrahim & Siregar, 2021) menerapkan SVM dan Naïve Bayes untuk klasifikasi opini pengguna Twitter, memperkuat relevansi pendekatan ini dalam konteks media sosial berbahasa Indonesia. Namun, penelitian tentang analisis sentimen pada platform X masih memiliki tantangan tersendiri, seperti keterbatasan jumlah karakter dalam cuitan, penggunaan bahasa informal, serta keberagaman gaya penulisan yang dapat memengaruhi akurasi klasifikasi sentimen. (Wulandari & Fauzan, 2023) menyoroti bahwa gaya bahasa informal yang khas di media sosial sering kali menyulitkan proses tokenisasi dan pemahaman konteks oleh model NLP. Selain itu, (Girsang & Simanjuntak, 2023) menunjukkan bahwa panjang teks dan kompleksitas struktur kalimat juga berdampak signifikan terhadap performa classifier, terutama pada platform seperti X yang memiliki format teks pendek dan padat makna.

Dari sisi teknis, analisis sentimen telah banyak memanfaatkan berbagai algoritma machine learning, seperti Naïve Bayes, Support Vector Machine (SVM), Random Forest, hingga metode berbasis Deep Learning. Misalnya, (Efendi et al., 2022) menunjukkan bahwa algoritma Naïve Bayes mampu mengklasifikasikan data dengan tingkat akurasi yang tinggi dalam konteks klasifikasi pembayaran. Sementara itu, (Gifari et al., 2022) membuktikan efektivitas SVM dalam menganalisis ulasan film berbasis teks.

Selain algoritma, pemilihan teknik representasi teks juga memengaruhi hasil klasifikasi. TF-IDF diketahui lebih efektif untuk dataset berukuran kecil, seperti yang dilaporkan dalam studi (Kusumawardani & Hidayatullah, 2020). Sebaliknya, Word Embedding cenderung lebih unggul ketika diterapkan pada korpus teks besar. Di sisi lain, model sekuensial seperti Long Short-Term Memory (LSTM) mampu menangkap pola urutan kata dengan baik, namun membutuhkan kapasitas komputasi yang lebih besar (Lestandy & Kurniawan, 2021).

Dalam beberapa tahun terakhir, pendekatan berbasis Transformer seperti BERT (Bidirectional Encoder Representations from Transformers) telah menunjukkan performa unggul dalam menangani teks pendek dan konteks kompleks, seperti dalam analisis cuitan (Rahman & Hasanah, 2025). Namun, efektivitas BERT sangat tergantung pada proses fine-tuning terhadap domain dan bahasa tertentu. (Lesmana & Pramudito, 2022) membuktikan bahwa fine-tuning BERT pada data berbahasa Indonesia dapat meningkatkan akurasi klasifikasi secara signifikan, khususnya dalam konteks media sosial yang cenderung informal.

Meskipun model deep learning menawarkan potensi besar, model klasik seperti SVM masih menunjukkan keunggulan dalam konteks dataset kecil dan tidak seimbang. (Maulana & Sari, 2021) menunjukkan bahwa SVM dapat mengungguli model LSTM dalam situasi data terbatas, terutama ketika fitur yang digunakan berbasis TF-IDF. Sementara itu, (Fathiarahma & Prasetya, 2023) berhasil mengimplementasikan Naïve Bayes dalam klasifikasi kegiatan keluarga menggunakan Orange Data Mining dan memperoleh akurasi tinggi, (Veronica & Voutama, 2023) juga membuktikan efektivitas Naïve Bayes dalam klasifikasi penyakit berbasis

<https://doi.org/10.35145/joisie.v9i1.4924>

JOISIE licensed under a Creative Commons Attribution-ShareAlike 4.0 International License (CC BY-SA 4.0)

data riil, dengan tingkat akurasi yang kompetitif, sehingga relevan diterapkan dalam konteks klasifikasi opini publik.

Namun, hingga saat ini belum ada studi yang secara spesifik menganalisis sentimen publik terhadap Danantara dengan pendekatan gabungan unsupervised learning (K-Means) dan supervised learning (SVM dan BERT) pada platform X. Hal ini menciptakan kesenjangan penelitian dalam pemanfaatan pendekatan semi-supervised untuk isu kebijakan lokal di Indonesia.

Penelitian ini bertujuan untuk mengisi kesenjangan tersebut dengan menganalisis sentimen publik terhadap Danantara melalui pendekatan NLP, menggunakan K-Means untuk pelabelan otomatis dan SVM serta BERT untuk klasifikasi sentimen. Penelitian ini juga bertujuan untuk membandingkan performa model dalam menangani karakteristik data media sosial Indonesia yang informal, pendek, dan tidak seimbang secara label.

2. METODE PENELITIAN

Penelitian ini menggunakan pendekatan semi-supervised learning untuk menganalisis sentimen publik terhadap Danantara di media sosial X (Twitter). Proses penelitian mengikuti alur Knowledge Discovery in Databases (KDD), yang mencakup pengumpulan data, pra-pemrosesan, pelabelan sentimen, klasifikasi, dan evaluasi performa model.

2.1. Crawling Data

Data dikumpulkan dari platform X (Twitter) menggunakan teknik crawling berbasis Python. Tools yang digunakan dalam proses ini meliputi snsrape, pandas, dan Node.js untuk mendukung akuisisi data secara otomatis. Data diambil menggunakan kata kunci “Dantara”, dengan filter berbahasa Indonesia, serta mencakup periode diskusi yang aktif mengenai isu tersebut. Total sebanyak 2.009 tweet berhasil dikumpulkan dan digunakan sebagai korpus utama dalam penelitian ini (Nugroho & Kurniawan, 2023)

 Jumlah tweet yang didapatkan 2009.

Gambar 1. Jumlah Tweet

2.2. Preprocessing Data

Langkah preprocessing dilakukan untuk meningkatkan kualitas teks sebelum dianalisis. Proses ini mencakup beberapa tahapan:

1. Cleaning: Menghapus URL, angka, tanda baca, mention, dan emoji dari tweet.

Tabel 1. Cleaning

Sebelum Cleaning	Setelah Cleaning
@BigAlphaID Ada isu negatif dikit aja tentang Danantara Makin ambrol. Siap siap 2025 bakal lebih berat dari 2024	ada isu negatif dikit aja tentang danantara makin ambrol siap siap bakal lebih berat dari
@IndoPopBase Kalau Danantara apa bakal bersih? trio IKN Foodestate Hambalang belum lagi yg lain ngelola dana ribuan trilyun udah jelas bancakan semua negri isinya bandit kok optimis	kalau danantara apa bakal bersih trio ikn foodestate hambalang belum lagi yg lain ngelola dana ribuan trilyun udah jelas bancakan semua negri isinya bandit kok optimis
@dheppylovely Wujudkan investasi Danantara yang berintegritas sama-sama. Ayo kita bantu Indonesia melangkah maju!	wujudkan investasi danantara yang berintegritas samasama ayo kita bantu indonesia melangkah maju

2. Tokenisasi: Memisahkan teks menjadi token atau kata-kata individual.

Tabel 2. Tokenisasi

Text
['Lagi', 'rame', 'Dantara', 'tiba', '-', 'tiba', 'muncul', 'Pertamina', '193', 'Triliun', 'RON', '90', 'ke', 'RON', '92', '!', '#', 'DantaraKaramkanNKRI', '#', 'DantaraKaramkanNKRI']
['bank', 'yg', 'gak', 'pemerintah', 'dantara', 'itu', 'apa', 'ya', 'yg', 'bagus', 'buat', 'nabunt']

['Danantara', 'Kuasai', '15', '%', 'Market', 'Cap', 'IHSG', 'Bos', 'Bursa', 'Bilang', 'Ini', 'Dampaknya', 'https://t.co/iaWjwCp4Zt']

3. Stopword Removal: Menghapus kata-kata tidak penting dengan pustaka Sastrawi.

Tabel 3. Stopword

Text
BCA error apakah ada hubungannya dengan danantara ? Wkwk jelek bgt pikiran guee
- IHSG ambruk - Rupiah terbakar di Rp 16.500 - Yamaha Sanken Sritex PHK Masal
- BBRI BBNI ambrol pasca gabung danantara - Pertamina kena skandal oplosan pertamax Closingan ramadhan kali ini ditutup dgn ekonomi oke gas turun
https://t.co/WV5QDiQPbK
Pasar bergejolak ekstrem pasar gak bisa boong IHSG jumat 28 feb 6270 6 turun- 214 85=3 31 % danantara bank Emas di didirikan reaksi pasar negative

4. Stemming: Mengubah kata ke bentuk dasar menggunakan algoritma Nazief-Adriani (Sastrawi).

Tabel 4. Stemming & Lemmatization

Text
soon danantara 14k t
danantara berpontensi menjadi katalis
positif bagi pertumbuhan ekonomi
belum danantara aja udah serem bgt ni bank

Penggunaan TF-IDF dalam tahapan selanjutnya didasarkan pada hasil penelitian Kusumawardani & Hidayatullah (2020), yang menyatakan bahwa TF-IDF efektif dalam menangani data teks pendek seperti tweet, terutama pada dataset dengan ukuran terbatas.

Contoh hasil preprocessing ditampilkan dalam bentuk tabel pada bagian lampiran untuk memperjelas perbedaan sebelum dan sesudah pembersihan data.

2.3. Pelebelan Sentimen

Data tweet yang telah dibersihkan belum memiliki label sentimen, sehingga dilakukan pelabelan awal menggunakan metode K-Means Clustering. K digunakan sebesar 3, yaitu: positif (2), netral (1), dan negatif (0).

Penggunaan K-Means bertumpu pada asumsi bahwa opini publik terhadap isu tertentu dapat terbagi secara alami ke dalam tiga sentimen tersebut. Pendekatan ini selaras dengan strategi semi-supervised, di mana hasil clustering digunakan sebagai label awal dalam model supervised learning. Studi (Nahjan & Nursidik, 2023) juga menunjukkan bahwa K-Means efektif dalam mengidentifikasi pola dalam data tidak berlabel, seperti transaksi penjualan, sehingga mendukung pemanfaatannya dalam analisis awal sentimen publik berbasis teks.

Tabel 5. Pelebelan Sentimen Dengan K-Means

Text	Sentimen
dengan danantara kita dapat mengharapkan pertumbuhan ekonomi yang merata dan berkelanjutan	Positif(2)
danantara indonesia kekuatan ekonomi baru nusantara	Neutral(1)
entah apa lagi alasan mereka untuk membenarkan kehadiran danantara sudah jelas investor ragu tidak yakin kalau danantara ini mengelolah investasi mereka	Negatif(0)

2.4. Klasifikasi Sentimen

Tweet yang telah diberi label melalui K-Means selanjutnya diklasifikasikan menggunakan dua model supervised, yaitu Support Vector Machine (SVM) dan BERT.

a. Support Vector Machine (SVM)

1. Ekstraksi fitur dilakukan dengan menggunakan metode TF-IDF (Term Frequency-Inverse Document Frequency) dari pustaka Scikit-learn.
2. Parameter yang digunakan:
3. Kernel: linear
4. C: 1.0 (nilai default regularisasi)
5. Model dilatih pada 80% data latih, dan diuji pada 20% data uji.

b. BERT (Bidirectional Encoder Representations from Transformers)

1. Model yang digunakan adalah IndoBERT-base (bukan multilingual), yang dioptimalkan untuk bahasa Indonesia.
2. Tokenisasi dilakukan dengan tokenizer IndoBERTTokenizer.
3. Embedding dan fine-tuning dilakukan menggunakan Huggingface Transformers dan PyTorch.
4. Optimizer yang digunakan: AdamW dengan learning rate default dari Huggingface.
5. Data dibagi dengan proporsi yang sama: 80% data latih dan 20% data uji.

SVM dipilih karena terbukti efektif menangani dataset kecil dan tidak seimbang, serta mampu mengklasifikasikan opini teks pendek dengan akurat.

2.5. Evaluasi Model

Evaluasi performa model dilakukan dengan menggunakan metrik berikut:

1. Akurasi: Persentase prediksi benar terhadap total data uji.
2. Precision: Tingkat ketepatan model dalam memprediksi suatu kelas.
3. Recall: Kemampuan model menangkap semua sampel yang relevan.
4. F1-Score: Harmonik rata-rata precision dan recall.
5. Confusion Matrix: Visualisasi hasil klasifikasi tiap kelas.

Pembagian data menggunakan teknik stratified sampling untuk menjaga proporsi distribusi kelas. Namun, karena jumlah data negatif dominan, performa model seperti BERT bisa menurun akibat ketidakseimbangan kelas. Temuan ini juga diperkuat oleh (Siregar & Hidayat, 2020) yang menyatakan bahwa model berbasis deep learning rentan terhadap bias jika distribusi kelas tidak seimbang.

Evaluasi ini membantu menilai efektivitas metode semi-supervised yang digunakan dalam menangkap sentimen publik terhadap Danantara secara objektif dan kuantitatif (Anggriyani, 2024).

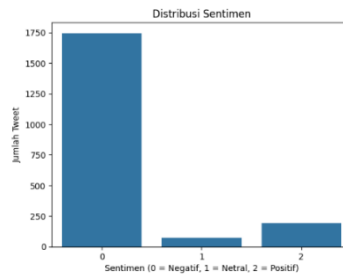
3. HASIL DAN PEMBAHASAN

Bagian ini menyajikan hasil dari proses analisis yang dilakukan berdasarkan data tweet yang telah dikumpulkan dan diproses. Penekanan diberikan pada distribusi sentimen hasil pelabelan, visualisasi kata kunci dominan, serta perbandingan performa model klasifikasi sentimen SVM dan BERT.

3.1. Hasil Pengolahan Data

Setelah melalui tahap pra-pemrosesan, sebanyak 2.009 tweet berbahasa Indonesia dianalisis menggunakan metode K-Means Clustering dengan jumlah kluster (K) ditentukan sebanyak tiga, yang merepresentasikan kategori negatif (0), netral (1), dan positif (2). Hasil clustering menunjukkan distribusi yang sangat tidak seimbang:

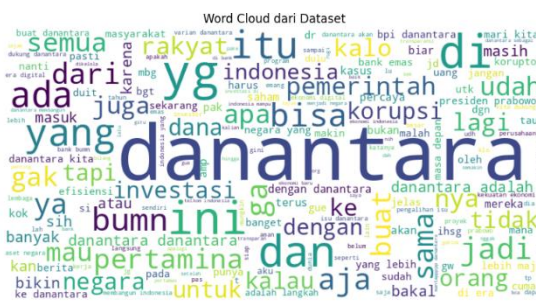
1. Negatif: 1.745 tweet (87.0%)
2. Netral: 73 tweet (3.6%)
3. Positif: 191 tweet (9.4%)



Gambar 2. Distribusi Sentimen Hasil K-Means

Hasil ini menunjukkan bahwa sebagian besar tweet mengandung sentimen negatif terhadap isu Danantara. Temuan ini mencerminkan opini publik yang cenderung kritis terhadap proyek tersebut. Namun, mengingat K-Means merupakan metode unsupervised, keakuratan pelabelan perlu divalidasi melalui uji sampling manual serta diperkuat dengan evaluasi model supervised di tahap berikutnya.

Sebagai bagian dari eksplorasi tematik, dilakukan visualisasi kata-kata yang paling sering muncul dalam tweet dengan menggunakan word cloud.



Gambar 3. Word Cloud

Visualisasi word cloud memperlihatkan bahwa kata “danantara” merupakan kata yang paling dominan, diikuti oleh kata-kata seperti “bumn”, “rakyat”, “investasi”, “pemerintah”, “korupsi”, dan “pertamina”. Selain itu, muncul juga frasa informal seperti “kalo aja”, “yg”, dan “udah”, yang mencerminkan gaya bahasa khas pengguna media sosial X. Keberadaan kata-kata tersebut menunjukkan bahwa diskusi publik terkait Danantara tidak hanya mencakup aspek ekonomi dan kelembagaan, tetapi juga dipenuhi nuansa skeptisisme, kecurigaan terhadap tata kelola, serta kekhawatiran terhadap dampak proyek ini terhadap rakyat dan negara. Temuan ini mendukung interpretasi bahwa sentimen publik lebih banyak diarahkan pada aspek negatif, sebagaimana terlihat dalam distribusi hasil pelabelan sebelumnya.

3.2. Analisis Performa Model

Secara keseluruhan, hasil evaluasi menunjukkan bahwa SVM memberikan performa klasifikasi yang sangat baik, dengan skor F1 yang tinggi dan konsisten di semua kelas. Hal ini menunjukkan bahwa kombinasi TF-IDF dan SVM mampu menangani data teks pendek seperti tweet secara efektif, meskipun data tidak seimbang. Temuan ini selaras dengan hasil penelitian Maulana & Sari (2021) yang menyatakan bahwa SVM unggul pada skenario data terbatas dan tidak seimbang.

Sebaliknya, BERT menunjukkan performa yang buruk dengan akurasi hanya 25%. Ketimpangan performa ini dapat dijelaskan oleh beberapa faktor:

1. Jumlah Data Terbatas

Model berbasis deep learning seperti BERT membutuhkan jumlah data yang besar agar dapat melakukan generalisasi secara optimal. Dataset berukuran kecil, seperti dalam penelitian ini, menyebabkan BERT kesulitan belajar representasi fitur yang efektif.

2. Kompleksitas Struktur Tweet

Tweet sering kali mengandung bahasa informal, singkatan, simbol, dan emoji yang tidak selalu dikenali oleh tokenizer BERT, sehingga memengaruhi proses embedding dan pemahaman konteks.

3. Ketidak Seimbangan Kelas

Sebagian besar data memiliki label negatif (87%), sementara kelas netral dan positif jauh lebih sedikit. Ketidakseimbangan ini menyulitkan model deep learning untuk mengenali dan mengklasifikasikan minoritas kelas secara akurat.

4. Tanpa Fine-Tuning

Penggunaan IndoBERT tanpa fine-tuning domain-spesifik membuat model tidak optimal menangkap konteks unik dalam tweet berbahasa Indonesia. Padahal, studi Lesmana & Pramudito (2022) menunjukkan bahwa fine-tuning sangat krusial untuk meningkatkan akurasi dalam teks lokal dan informal.

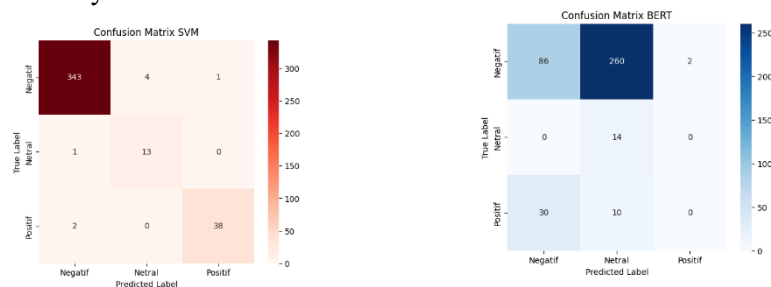
Sebaliknya, SVM dengan fitur TF-IDF tetap mampu bekerja secara stabil meskipun dengan data terbatas dan distribusi kelas tidak seimbang. Hal ini menunjukkan bahwa model klasik seperti SVM masih relevan digunakan dalam konteks analisis sentimen pada media sosial dengan karakteristik teks pendek dan informal.

3.3. Evaluasi Model dan Interpretasi

Evaluasi model dilakukan untuk menilai kualitas klasifikasi yang dihasilkan oleh SVM dan BERT. Metrik evaluasi yang digunakan meliputi:

1. Akurasi
Mengukur seberapa banyak prediksi yang benar terhadap keseluruhan data uji.
2. Precision
Menilai seberapa tepat model dalam mengklasifikasikan data ke dalam kelas tertentu.
3. Recall
Mengukur kemampuan model dalam menemukan semua data yang relevan pada masing-masing kelas.
4. F1-Score
Rata-rata harmonik dari precision dan recall, berguna dalam kondisi distribusi kelas tidak seimbang.
5. Confusion Matrix
Menampilkan distribusi prediksi benar dan salah untuk masing-masing kelas.

Pembagian data dilakukan dengan proporsi 80% untuk pelatihan dan 20% untuk pengujian, menggunakan teknik stratified split agar distribusi antar kelas tetap terjaga. Evaluasi dilakukan pada set data uji yang tidak pernah dilihat model sebelumnya.



Gambar 4. (a) Confusion Matix SVM (b) Confusion Matrix BERT

Tabel 6. Evaluasi Kinerja Model SVM

Kelas	Precision	Recall	F1-Score	Support
0	0.99	0.99	0.99	348
1	0.76	0.93	0.84	14
2	0.97	0.95	0.96	40
Akurasi	-	-	0.98	402
MAcro	0.91	0.95	0.93	402
Avg				
Weight	0.98	0.98	0.98	402
d Avg				

Tabel 7. Evaluasi Evaluasi Kinerja Model BERT

Kelas	Precision	Recall	F1-Score	Support
0	0.74	0.25	0.37	348
1	0.05	1.00	0.09	14
2	0.00	0.93	0.00	40
Akurasi	-	-	0.25	402
MAcro	0.26	0.42	0.13	402
Avg				
Weight	0.64	0.25	0.32	402
d Avg				

Secara keseluruhan, hasil evaluasi menunjukkan bahwa model SVM menghasilkan performa yang konsisten dan seimbang pada ketiga kelas sentimen, baik negatif, netral, maupun positif. Hal ini menunjukkan kemampuan SVM dalam mengelola fitur teks berbasis TF-IDF secara efisien, bahkan pada dataset yang terbatas dan memiliki distribusi kelas yang tidak merata.

Sebaliknya, model BERT menunjukkan ketimpangan ekstrem dalam performanya, di mana ia hanya mampu mengenali kelas minoritas tertentu dengan baik, tetapi gagal mengklasifikasikan mayoritas tweet secara akurat. Kondisi ini diduga disebabkan oleh kombinasi faktor, seperti jumlah data yang terbatas, dominasi kelas negatif dalam dataset, serta kompleksitas bahasa informal dalam tweet yang tidak sepenuhnya ditangani oleh tokenizer BERT.

Temuan ini memperkuat pemahaman bahwa metode berbasis TF-IDF dan SVM lebih sesuai dan efisien untuk analisis opini publik pada data teks pendek, berbahasa Indonesia, serta tidak berlabel dalam jumlah besar, seperti yang umum ditemukan pada media sosial X.

4. KESIMPULAN

Penelitian ini bertujuan untuk menganalisis sentimen publik terhadap isu Danantara melalui cuitan di media sosial X (sebelumnya dikenal sebagai Twitter) dengan pendekatan Natural Language Processing (NLP) berbasis kombinasi unsupervised dan supervised learning. Berdasarkan hasil pelabelan menggunakan K-Means Clustering, ditemukan bahwa mayoritas tweet (87.0%) mengandung sentimen negatif, sedangkan sentimen netral dan positif masing-masing hanya sebesar 3.6% dan 9.4%. Temuan ini mengindikasikan bahwa diskursus publik terhadap Danantara didominasi oleh opini yang bersifat kritis dan penuh ketidakpuasan.

Dari segi pemodelan, Support Vector Machine (SVM) terbukti lebih unggul dibandingkan dengan BERT, dengan akurasi mencapai 98%, jauh di atas performa BERT yang hanya 25%. Keunggulan SVM ini menunjukkan bahwa metode klasik berbasis TF-IDF lebih efektif dalam kondisi data terbatas dan tidak seimbang, serta lebih tahan terhadap kompleksitas bahasa informal yang umum pada media sosial. Sebaliknya, kegagalan BERT menunjukkan bahwa model deep learning memerlukan jumlah data yang lebih besar, representasi bahasa yang sesuai domain, dan distribusi kelas yang lebih seimbang agar dapat berfungsi secara optimal.

Dari sisi kontribusi ilmiah, penelitian ini memberikan pendekatan semi-supervised yang relevan untuk kondisi di mana data tidak memiliki label, khususnya dalam konteks media sosial Indonesia. Kombinasi K-Means, SVM, dan BERT dalam satu kerangka penelitian belum banyak dijumpai dalam kajian lokal, sehingga penelitian ini memberikan sumbangan metodologis bagi pengembangan riset NLP berbahasa Indonesia.

Namun demikian, penelitian ini memiliki beberapa keterbatasan. Pertama, data yang digunakan terbatas pada satu periode waktu dan satu platform sosial media. Kedua, proses pelabelan menggunakan K-Means tidak menjamin kualitas label seakurat pelabelan manual. Ketiga, distribusi data yang tidak seimbang sangat memengaruhi kinerja model berbasis deep learning. Oleh karena itu, hasil ini perlu ditafsirkan secara hati-hati, terutama dalam generalisasi terhadap populasi yang lebih luas.

Untuk penelitian selanjutnya, disarankan untuk:

1. Menggunakan dataset yang lebih besar dan beragam, baik dari segi waktu maupun platform (misalnya Facebook, TikTok, YouTube).
2. Menggunakan metode validasi silang (cross-validation) agar hasil evaluasi lebih stabil.
3. Melibatkan manusia (annotator) dalam proses pelabelan untuk meningkatkan kualitas data.
4. Menyempurnakan arsitektur deep learning seperti fine-tuning IndoBERT pada domain media sosial.

Secara praktis, penelitian ini dapat digunakan oleh pemangku kebijakan dan analis media untuk memahami persepsi masyarakat secara cepat terhadap isu strategis nasional. Selain itu, pendekatan ini dapat direplikasi pada isu-isu lain guna mendukung pengambilan keputusan berbasis data.

5. DAFTAR PUSTAKA

- Agustina, N., Citra, D. H., Purnama, W., Nisa, C., & Kurnia, A. R. (2022). Implementasi Algoritma Naive Bayes untuk Analisis Sentimen Ulasan Shopee pada Google Play Store. *MALCOM: Indonesian Journal of Machine Learning and Computer Science*, 2(1), 47–54.
- Aji, A. F., & others. (2021). IndoNLU: Benchmark Natural Language Understanding untuk Bahasa Indonesia. *ArXiv Preprint ArXiv:2101.00390*.
- Anggriyani, R. (2024). Analisis Sentimen Publik terhadap Kebijakan Pemerintah Menggunakan Twitter. *Jurnal Penelitian Komunikasi Dan Opini Publik*, 8(1), 12–19.
- Ardiansyah, A., & others. (2024). Implementasi Analisis Sentimen Masyarakat Mengenai Kenaikan Harga Bbm Dengan Metode Naive Bayes. *JOISIE (Journal Of Information Systems And Informatics Engineering)*, 8(1), 1–9.
- Efendi, D. M., Mintoro, S., Lubis, S. H., & Lestari, S. (2022). Klasifikasi Kinerja Pembayaran Angsuran dengan Algoritma Naive Bayes. *Jurnal Informasi Dan Komputer*, 10(1), 57–61.
- Fathiarahma, D., & Prasetya, A. (2023). Penerapan Naive Bayes pada Klasifikasi Kegiatan Keluarga Menggunakan Orange Data Mining. *Jurnal Sistem Informasi*, 11(2), 78–85.
- Gifari, R., Wibowo, R., & Iskandar, D. (2022). Analisis Sentimen pada Review Film Menggunakan Support Vector Machine dan Random Forest. *Jurnal Teknologi Dan Sistem Komputer*, 10(2), 183–190.
- Girsang, P., & Simanjuntak, J. (2023). Pengaruh Panjang Teks dan Struktur Kalimat terhadap Akurasi Sentiment Classifier di Media Sosial. *Jurnal Teknik Informatika*, 19(1), 40–47.
- Ibrahim, A., & Siregar, F. (2021). Penerapan SVM dan Naive Bayes untuk Analisis Sentimen Twitter dalam Bahasa Indonesia. *Jurnal Teknologi Informasi Dan Ilmu Komputer*, 8(1).
- Kusumawardani, R., & Hidayatullah, M. R. (2020). Perbandingan Metode TF-IDF dan Word2Vec untuk Representasi Teks pada Analisis Sentimen. *Jurnal Informatika: Jurnal Pengembangan IT (JPIT)*, 5(1), 1–6.
- Lesmana, B., & Pramudito, M. (2022). Fine-Tuning IndoBERT untuk Analisis Sentimen Bahasa Indonesia pada Media Sosial. *Jurnal Ilmu Komputer Dan Informasi*, 15(2), 101–108.
- Lestandy, R., & Kurniawan, A. (2021). Analisis Sentimen Komentar YouTube Menggunakan LSTM dan Word Embedding. *Jurnal RESTI (Rekayasa Sistem Dan Teknologi Informasi)*, 5(6), 1093–1099.
- Maulana, D., & Sari, R. (2021). Perbandingan Kinerja Algoritma SVM dan LSTM untuk Klasifikasi Sentimen Bahasa Indonesia. *Jurnal Teknologi Dan Sistem Komputer*, 9(3), 273–280.

- Nahjan, M., & Nursidik, M. (2023). Penerapan K-Means Clustering untuk Analisis Pola Transaksi Penjualan pada E-Commerce. *Jurnal Teknologi Informasi Dan Data*, 8(2), 111–118.
- Nugroho, D., & Kurniawan, A. (2023). Strategi Semi-Supervised Learning dalam Analisis Sentimen Twitter: Studi Kasus Isu Nasional. *Jurnal Artificial Intelligence Dan Data Science*, 3(1), 25–34.
- Paramita, P., & Ibrahim, A. (2023). Analisis Sentimen Terhadap Pengguna QRIS pada Twitter Menggunakan Naive Bayes Classifier. *JOISIE*, 7(1), 1–6.
- Permana, A. A., Noviyanto, W. A., & Kristiyanti, D. A. (2023). Sentimen Analisis Opini Masyarakat terhadap UMKM pada Media Sosial Twitter. *Jurnal Minfo Polgan*, 12(1), 163–170.
- Putri, A. D., & Pratama, R. Y. (2022). Sentiment Analysis terhadap Isu Nasional Menggunakan Twitter Data. *Jurnal Ilmu Komputer Dan Teknologi Informasi*, 10(2), 155–161.
- Rahman, F., & Hasanah, N. (2025). Penerapan BERT untuk Analisis Sentimen pada Media Sosial Bahasa Indonesia. *Jurnal Data Sains Dan Analitik*, 3(1).
- Saputra, N. A., Alexandra, J., & Trisno, I. B. (2023). Analisis Sentimen Obrolan Telegram Menggunakan Naive Bayes. *JATI*, 7(2), 1321–1327.
- Siregar, H., & Hidayat, S. (2020). Pengaruh Ketidakseimbangan Data pada Kinerja Model Deep Learning dalam Analisis Sentimen. *Jurnal Informatika Dan Sistem Informasi*, 6(1), 23–30.
- Syafrianto, A. (2022). Perbandingan Algoritma Naive Bayes dan Decision Tree Pada Sentimen Analisis. *The Indonesian Journal of Computer Science Research*, 1(2).
- Tanggraeni, A. I., & Sitokdana, M. N. N. (2022). Analisis Sentimen Aplikasi E-Government pada Google Play Menggunakan Naive Bayes. *JATISI*, 9(2), 785–795.
- Veronica, A., & Voutama, D. (2023). Analisis Klasifikasi Penyakit Diabetes pada Perempuan Menggunakan Algoritma Naive Bayes. *Jurnal Informatika Dan Komputer*, 6(1), 22–28.
- Wulandari, S., & Fauzan, M. (2023). Kendala Bahasa Informal dalam Klasifikasi Sentimen Twitter. *Jurnal Linguistik Komputasional*, 2(1), 45–51.