

Perbandingan Naïve Bayes dan K-Nearest Neighbor dalam Klasifikasi Polaritas Tweet pada Cross-Domain #MakanBergiziGratis dan #Danantara

Ahmad Fadly^a, Purnawansyah^b, Herdianti Darwis^c

Universitas Muslim Indonesia, Makassar, Indonesia

^a13020210285@umi.ac.id; ^bpurnawansyah@umi.ac.id; ^cherdianti.darwis@umi.ac.id

Received: 08-04-2026 | Revised: 20-05-2026 | Accepted: 24-06-2026 | Published: 29-06-2026

Abstrak

Media sosial seperti Twitter (X) menghasilkan data opini publik dalam jumlah besar yang dapat dimanfaatkan untuk analisis sentimen. Namun, perbedaan konteks antar topik atau hashtag sering menyebabkan terjadinya *domain shift* yang dapat menurunkan performa model klasifikasi. Penelitian ini bertujuan untuk menganalisis performa algoritma *Naïve Bayes* dan *K-Nearest Neighbors* (KNN) dalam skenario *cross-domain* menggunakan dua representasi fitur, yaitu TF-IDF dan FastText. Dataset diperoleh dari hashtag #MakanBergiziGratis sebagai domain sumber dan #Danantara sebagai domain target. Metode yang digunakan meliputi preprocessing teks, ekstraksi fitur, pemodelan, serta evaluasi menggunakan *accuracy*, *precision*, *recall*, dan *F1-score*. Hasil penelitian menunjukkan bahwa kedua model memiliki performa tinggi pada internal test dengan *accuracy* sebesar 0.94. Namun, pada pengujian lintas domain, Naïve Bayes dengan TF-IDF menunjukkan performa yang lebih stabil dengan *accuracy* sebesar 0.75, sedangkan KNN dengan FastText mengalami penurunan signifikan, terutama pada kelas negatif dengan nilai *F1-score* sebesar 0.00. Temuan ini menunjukkan bahwa pemilihan algoritma dan representasi fitur sangat mempengaruhi kemampuan generalisasi model dalam menghadapi *domain shift*.

Kata kunci: Analisis Sentimen, Naïve Bayes, KNN, TF-IDF, *Cross-domain*.

Pendahuluan

Platform media sosial telah mengubah cara manusia berinteraksi dengan memungkinkan komunikasi tanpa batas geografis serta mempercepat distribusi informasi secara global [1]. Salah satu platform yang banyak digunakan untuk mengekspresikan opini publik adalah Twitter, yang kini dikenal sebagai X [2], [3]. Di Indonesia, jumlah pengguna X terus mengalami peningkatan setiap tahunnya, sehingga menghasilkan volume data opini publik yang sangat besar dan berpotensi untuk dianalisis guna memahami persepsi masyarakat terhadap berbagai isu, termasuk kebijakan dan program pemerintah seperti Program Makan Bergizi Gratis (MBG) serta lembaga Daya Anagata Nusantara (Danantara). Analisis sentimen menjadi salah satu teknik yang banyak digunakan untuk mengidentifikasi kecenderungan opini publik, apakah bersifat positif atau negatif [4]. Namun, karakteristik data pada media sosial, khususnya tweet yang cenderung singkat, mengandung singkatan, bahasa informal, serta perubahan konteks antar hashtag, menimbulkan tantangan tersendiri dalam proses analisis. Hal ini menyebabkan performa model klasifikasi menjadi sangat bergantung pada metode yang digunakan.

Sejumlah penelitian sebelumnya menunjukkan bahwa performa algoritma klasifikasi teks tidak selalu konsisten. Pada analisis sentimen terkait transportasi publik di Indonesia, algoritma Naïve Bayes dibandingkan dengan *K-Nearest Neighbors* (KNN) menunjukkan hasil yang bervariasi tergantung pada dataset yang digunakan [5]. Sementara itu, penelitian lain menunjukkan bahwa *Support Vector Machine* (SVM) mampu memberikan performa terbaik dibandingkan Naïve Bayes dan KNN pada kasus analisis sentimen vaksinasi COVID-19 [6]. Variasi hasil tersebut mengindikasikan bahwa performa algoritma sangat dipengaruhi oleh karakteristik dataset, ukuran data, serta metode pelabelan yang digunakan. Temuan serupa juga ditunjukkan pada penelitian lain di media sosial X yang mengkaji variasi algoritma machine learning dalam klasifikasi topik sensitif seperti *body shaming*, di mana performa model menunjukkan perbedaan signifikan antar algoritma yang digunakan [7]. Selain itu, implementasi *Naïve Bayes* pada analisis sentimen isu politik di Twitter juga menunjukkan bahwa algoritma ini cukup efektif, meskipun sangat bergantung pada kualitas data dan tahapan pra-proses yang dilakukan [8].

Selain pemilihan algoritma, representasi teks juga berperan penting dalam menentukan hasil analisis. Berbagai pendekatan *machine learning* dan representasi fitur telah digunakan dalam analisis sentimen, seperti yang

ditunjukkan dalam studi perbandingan model pada data skala besar [9], [10] serta penerapan TF-IDF dan Word2Vec dalam klasifikasi teks [11]. Studi lain juga menunjukkan pemanfaatan model berbasis pembelajaran mendalam seperti BERT dalam analisis sentimen pada data Twitter [12]. Hal ini menunjukkan bahwa pemilihan representasi teks sangat memengaruhi performa model klasifikasi. Studi komparatif menunjukkan bahwa *Naïve Bayes* cenderung stabil ketika dikombinasikan dengan representasi berbasis frekuensi seperti TF-IDF, sedangkan KNN lebih optimal ketika menggunakan *embedding* berdimensi rendah karena sensitivitasnya terhadap skala dan dimensi fitur. Oleh karena itu, penting untuk membandingkan dua kombinasi metode yang umum digunakan, yaitu *Naïve Bayes* + TF-IDF dan KNN + FastText. Penggunaan algoritma *Naïve Bayes* dan *K-Nearest Neighbors* juga telah banyak diterapkan pada berbagai kasus klasifikasi di luar analisis sentimen, seperti pengenalan pola tulisan tangan [13], sistem pakar berbasis KNN [14], serta klasifikasi citra menggunakan pendekatan KNN dan metode lainnya [15]. Hal ini menunjukkan bahwa kedua algoritma tersebut memiliki fleksibilitas tinggi dan dapat diadaptasi pada berbagai jenis permasalahan.

Tantangan lain dalam analisis sentimen adalah kemampuan model dalam menghadapi perbedaan konteks antar domain. Model yang dilatih pada satu domain sering kali mengalami penurunan performa ketika diterapkan pada domain lain akibat fenomena *domain shift*, yaitu perbedaan distribusi kosakata dan gaya bahasa. Penelitian sebelumnya menekankan pentingnya pendekatan *cross-domain* untuk meningkatkan kemampuan generalisasi model [16], [17]. Meskipun demikian, sebagian besar penelitian masih dilakukan pada domain yang sama, sehingga belum mampu mengevaluasi kemampuan generalisasi model secara optimal ketika dihadapkan pada perbedaan konteks antar hashtag [18], [19].

Berdasarkan permasalahan tersebut, penelitian ini bertujuan untuk mengevaluasi efektivitas algoritma *Naïve Bayes* dengan representasi TF-IDF serta *K-Nearest Neighbors* dengan *embedding FastText* dalam skenario *cross-domain*, yaitu pelatihan pada hashtag #MakanBergiziGratis dan pengujian pada hashtag #Danantara. Adapun kontribusi penelitian ini meliputi: (1) membandingkan performa kedua algoritma dalam skenario transfer antar hashtag, (2) menganalisis kemampuan generalisasi antara TF-IDF dan FastText, serta (3) memberikan rekomendasi praktis bagi analis opini publik dalam memilih metode yang cepat, efisien, dan robust terhadap perubahan konteks.

Metode

Metode Penelitian ini menggunakan pendekatan kuantitatif dengan metode eksperimen dan komparatif untuk mengevaluasi performa algoritma klasifikasi dalam skenario analisis sentimen lintas domain. Fokus utama penelitian adalah membandingkan kombinasi *Naïve Bayes* dengan TF-IDF dan *K-Nearest Neighbors* (KNN) dengan FastText pada Twitter berbahasa Indonesia.

A. Sumber dan Pengumpulan Data

Dataset dalam penelitian ini diperoleh dari *platform* Twitter (X) melalui teknik *crawling* dengan memanfaatkan dua *hashtag* utama, yaitu #MakanBergiziGratis sebagai *domain* sumber dan #Danantara sebagai *domain* target. Data yang dikumpulkan berupa *tweet* berbahasa Indonesia dalam periode Februari hingga Oktober 2025. Jumlah data yang diperoleh sebanyak 3.553 *tweet* pada *domain* #MakanBergiziGratis dan 1.355 *tweet* pada *domain* #Danantara, yang selanjutnya digunakan dalam proses analisis sentimen pada skenario *cross-domain*.

B. Preprocessing Data

Tahapan *preprocessing* dilakukan untuk meningkatkan kualitas data sebelum proses klasifikasi [20]. Proses ini mencakup *case folding* untuk mengubah seluruh teks menjadi huruf kecil, *cleaning* untuk menghapus URL, emoji, dan karakter khusus, serta *normalization* untuk mengubah kata tidak baku menjadi bentuk standar. Selanjutnya, dilakukan *tokenization* untuk memecah teks menjadi token kata, diikuti dengan *stopword removal* untuk menghilangkan kata-kata umum yang tidak informatif, serta *stemming* menggunakan Sastrawi untuk memperoleh bentuk dasar kata. Tahapan ini bertujuan menghasilkan data yang bersih dan konsisten sehingga dapat meningkatkan kinerja model klasifikasi [21].

C. Pelabelan Data

Data *tweet* diberi label secara manual ke dalam dua kelas, yaitu positif dan negatif. Pelabelan ini bertujuan untuk menghasilkan *ground truth* yang akurat sebagai acuan dalam proses pelatihan dan evaluasi model, sehingga dapat meningkatkan validitas hasil klasifikasi sentimen [22].

D. Pembagian Data (*Train-Test Split*)

Pada tahap pembagian data (*train-test split*), dataset diatur menggunakan proporsi standar untuk keperluan pelatihan dan evaluasi model. Data dari *domain* #MakanBergiziGratis digunakan sebagai basis pelatihan, di mana 80% dialokasikan sebagai *training set*, kemudian 10% sebagai *validation set*, dan 10% sisanya sebagai *internal test* untuk mengukur performa model pada *domain* yang sama.

Sementara itu, seluruh data dari *domain* #Danantara digunakan sebagai *final test* atau pengujian lintas *domain* (*cross-domain testing*). Pembagian ini dirancang untuk mensimulasikan skenario *cross-domain*, yaitu ketika model dilatih pada satu *domain* namun diuji pada *domain* yang berbeda, sehingga dapat menilai kemampuan generalisasi model secara lebih menyeluruh [23], [24].

E. Ekstraksi Fitur

Dalam penelitian ini, digunakan dua metode representasi teks, yaitu *TF-IDF* (*Term Frequency-Inverse Document Frequency*) dan *FastText*. Metode *TF-IDF* digunakan pada model *Naïve Bayes* untuk merepresentasikan frekuensi kata dalam dokumen, sedangkan *FastText* digunakan pada model *K-Nearest Neighbors* (KNN) untuk menghasilkan representasi berbasis *embedding* dengan mempertimbangkan *subword information* [25].

F. Pemodelan

Selanjutnya, pada tahap pemodelan digunakan dua algoritma klasifikasi, yaitu *Naïve Bayes* (*Multinomial*) dan (KNN). *Naïve Bayes* dipilih karena efisien dalam menangani data teks berbasis frekuensi, sementara KNN digunakan dengan metrik *cosine similarity* untuk mengukur kedekatan antar dokumen [26]. Pemilihan kedua algoritma ini didasarkan pada perbandingan antara metode probabilistik dan metode berbasis jarak dalam analisis sentimen [27].

G. Evaluasi Model

Performa model dievaluasi menggunakan beberapa metrik, yaitu *accuracy*, *precision*, *recall*, dan *F1-score*. *Accuracy* digunakan untuk mengukur proporsi prediksi yang benar terhadap seluruh data uji, yang dihitung berdasarkan perbandingan jumlah prediksi benar (*true positive* dan *true negative*) terhadap total data [28]. Selanjutnya, *precision* digunakan untuk menilai sejauh mana model tepat dalam memprediksi kelas positif, yaitu perbandingan antara *true positive* dan seluruh prediksi positif. Sementara itu, *recall* mengukur kemampuan model dalam menemukan seluruh data yang benar-benar positif, yaitu perbandingan antara *true positive* dan total data positif aktual. Adapun *F1-score* merupakan rata-rata harmonis antara *precision* dan *recall*, yang digunakan untuk memberikan keseimbangan antara kedua metrik tersebut, terutama ketika terdapat ketidakseimbangan kelas, dan biasanya dihitung menggunakan pendekatan *macro averaging* [29].

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN} \quad (1)$$

$$Precision = \frac{TP}{TP + FP} \quad (2)$$

$$Recall = \frac{TP}{TP + FN} \quad (3)$$

$$F1 = \frac{2 \times Precision \times recall}{Precision + Recall} \quad (4)$$

Tabel 1. *Confusion Matrix*

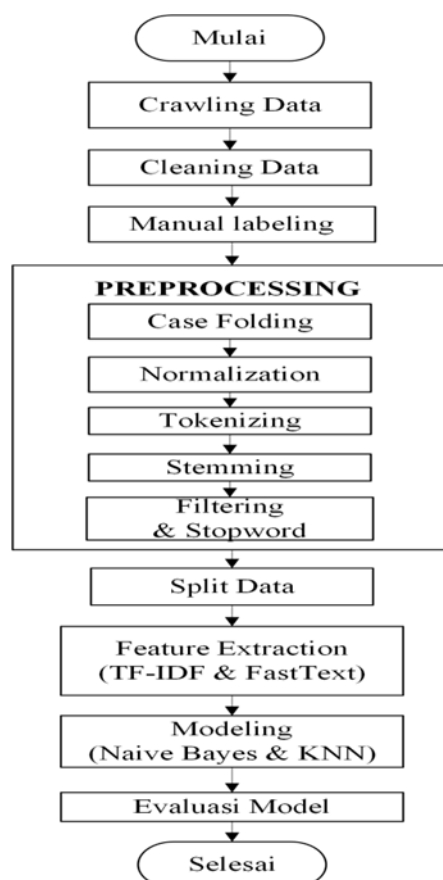
		<i>predicted</i>	
		<i>positives</i>	<i>negatives</i>
<i>actual</i>	<i>positives</i>	TP	FN
	<i>negatives</i>	FP	TN

Pengujian model dianalisis menggunakan *confusion matrix* untuk melihat distribusi prediksi benar dan salah pada setiap kelas. Komponen utama dalam *confusion matrix* meliputi *true positive* (TP), yaitu jumlah

data positif yang berhasil diprediksi dengan benar; *true negative* (TN), yaitu data negatif yang diprediksi dengan benar; *false positive* (FP), yaitu data negatif yang salah diprediksi sebagai positif; serta *false negative* (FN), yaitu data positif yang salah diprediksi sebagai negatif [30]. Keempat komponen ini digunakan sebagai dasar perhitungan metrik evaluasi seperti *accuracy*, *precision*, *recall*, dan *F1-score*.

Perancangan

Perancangan sistem dalam penelitian ini bertujuan untuk membangun alur proses analisis sentimen secara terstruktur dalam skenario *cross-domain*. Sistem dirancang sebagai *pipeline* pemrosesan data teks yang dimulai dari tahap pengumpulan data hingga evaluasi model klasifikasi. Secara umum, arsitektur sistem terdiri dari beberapa tahapan utama, yaitu pengumpulan data, pembersihan data, pelabelan, *preprocessing*, pembagian dataset, ekstraksi fitur, pemodelan, serta evaluasi. Setiap tahapan dirancang saling terintegrasi untuk memastikan proses pembelajaran model berjalan optimal dalam menghadapi perbedaan distribusi data antar *domain*. Untuk memberikan gambaran alur sistem secara menyeluruh, arsitektur penelitian ditunjukkan pada Gambar 1.



Gambar 1. Arsitektur sistem analisis sentimen *cross-domain*

Berdasarkan Gambar 1, proses dimulai dari pengumpulan data *tweet* berdasarkan *hashtag* yang telah ditentukan, yaitu #MakanBergiziGratis sebagai *domain* sumber dan #Danantara sebagai *domain* target. Data yang diperoleh kemudian melalui tahap penghapusan duplikasi untuk menjaga kualitas dataset.

Selanjutnya, data diberi label secara manual untuk menghasilkan *ground truth* yang valid. Tahap *preprocessing* dilakukan untuk menstandarkan teks sehingga dapat diproses oleh algoritma *machine learning*. Proses ini mencakup *case folding*, *cleaning*, *normalization*, *tokenization*, *stopword removal*, dan *stemming*.

Setelah *preprocessing*, dataset dibagi menjadi data latih, validasi, dan pengujian internal, serta data uji lintas *domain*. Pembagian ini dirancang untuk mensimulasikan kondisi nyata di mana model harus mampu melakukan generalisasi pada data dari domain yang berbeda.

Pada tahap ekstraksi fitur, digunakan dua pendekatan representasi teks, yaitu TF-IDF dan *FastText*. Perbedaan pendekatan ini bertujuan untuk menganalisis pengaruh representasi fitur terhadap performa model dalam menghadapi *domain shift*.

Selanjutnya, proses pemodelan dilakukan menggunakan dua algoritma, yaitu *Naïve Bayes* dan *K-Nearest Neighbors* (KNN). Kedua algoritma tersebut digunakan untuk membandingkan pendekatan probabilistik dan berbasis jarak dalam klasifikasi sentimen. Tahap akhir adalah evaluasi model menggunakan metrik klasifikasi untuk mengukur performa pada pengujian internal maupun pada skenario *cross-domain*. Untuk menjelaskan setiap tahapan secara lebih rinci, alur sistem yang digunakan dalam penelitian ini disajikan pada Tabel 2.

Tabel 2. Tahapan sistem analisis sentimen

No	Tahapan	Input	Proses	Output
1	<i>Crawling data</i>	<i>Hashtag</i> Twitter	Pengambilan data tweet	Dataset mentah
2	<i>Remove duplicate</i>	Dataset mentah	Menghapus data duplikat	Dataset unik
3	<i>Manual labelling</i>	Datraset unik	Pelabelan sentimen (positif/negatif)	Data berlabel
4	<i>Preprocessing</i>	Data berlabel	<i>Case folding, normalization. Tokenizing. Stemming. Filtering & stopword</i>	Data bersih
5	<i>Split data</i>	Data bersih	Pembagian data (<i>train, validation, test, cross-domain</i>)	Data terpisah
6	<i>Feature extraction</i> (TF-IDF)	Data <i>train</i>	Ekstraksi fitur berbasis frekuensi kata	Vektor TF-IDF
7	<i>Feature extraction</i> (FastText)	Data <i>train</i>	Ekstraksi fitur berbasis embedding	Vektor FastText
8	<i>Modelling</i> (Naïve bayes)	Vektor TF-IDF	Klasifikasi multinomial Naïve Bayes	Model NB
9	<i>Modelling</i> (KNN)	Vektor FastText	Klasifikasi K-Nearst Neighbor	Model KNN
10	Evaluasi model	Model & data uji	Perhitungan metrik (<i>accuracy, precision, Recall, F1-Score</i>)	Nilai performa
11	Hasil akhir	Nilai performa	Analisis hasil	Kesimpulan penelitian

Tabel 2 menunjukkan bahwa proses analisis dimulai dari pengumpulan data hingga evaluasi model secara sistematis. Setiap tahapan memiliki peran penting dalam memastikan kualitas data dan akurasi hasil klasifikasi. Urutan proses tersebut dirancang untuk mendukung skenario pembelajaran lintas *domain*, sehingga model dapat diuji kemampuannya dalam menghadapi perbedaan konteks antar dataset.

Pemodelan

Pada tahap ini dilakukan pembangunan model klasifikasi sentimen menggunakan dua pendekatan, yaitu *Naïve Bayes* dan *K-Nearest Neighbors* (KNN). Kedua model digunakan untuk membandingkan performa metode probabilistik dan metode berbasis jarak dalam skenario *cross-domain*.

A. *Naïve Bayes* (Multinomial)

Naïve Bayes merupakan algoritma klasifikasi berbasis probabilitas yang menggunakan *Teorema Bayes* dengan asumsi independensi antar fitur. Pada penelitian ini digunakan Multinomial *Naïve Bayes* yang umum diterapkan pada data teks. Probabilitas suatu dokumen d termasuk ke dalam kelas c dihitung menggunakan persamaan:

$$P(c|d) = \frac{P(c) \times P(d|c)}{P(d)} \quad (5)$$

Dalam implementasinya nilai $P(d|c)$ dihitung berdasarkan frekuensi kemunculan kata yang dipresentasikan menggunakan TF-IDF. Metode ini memiliki beberapa keunggulan yaitu efisien pada data teks berdimensi tinggi, stabil terhadap variasi data, serta tidak memerlukan komputasi yang kompleks.

B. *K-Nearest Neighbors* (KNN)

KNN merupakan algoritma klasifikasi berbasis jarak yang menentukan kelas suatu data berdasarkan kedekatan dengan tetangga terdekat dalam ruang fitur. Pada penelitian ini, representasi fitur menggunakan *FastText embedding*, sedangkan pengukuran jarak menggunakan *cosine similarity*. Rumus *cosine similarity* :

$$\text{Similarity}(A, B) = \frac{A \cdot B}{\|A\| \|B\|} \quad (5)$$

Nilai kedekatan digunakan untuk menentukan k tetangga terdekat, kemudian kelas ditentukan berdasarkan mayoritas label dari tetangga tersebut. Metode *K-Nearest Neighbors* (KNN) memiliki beberapa karakteristik yaitu sensitif terhadap distribusi data, bergantung pada pemilihan nilai k , serta tidak melakukan proses eksplisit (*lazy learning*).

C. Skema Pemodelan

Skema Pemodelan *Cross-Domain* dalam penelitian ini dilakukan dengan melatih model menggunakan data dari domain sumber (#MakanBergiziGratis) dan mengujinya pada data dari domain target (#Danantara). Pendekatan *cross-domain* ini digunakan untuk mengukur kemampuan generalisasi model terhadap data yang memiliki distribusi berbeda.

D. Parameter dan Implementasi

Parameter yang digunakan dalam penelitian ini meliputi Naïve Bayes dengan pendekatan Multinomial menggunakan parameter *default*, K-Nearest Neighbors (KNN) yang menggunakan *cosine similarity* dengan nilai K tertentu, *TF-IDF* dengan skema unigram, serta *FastText* dengan *pretrained embedding* bahasa Indonesia. Untuk menghindari *data leakage*, proses ekstraksi fitur hanya dilakukan pada data *training*, kemudian diterapkan pada data uji.

E. Hasil Penelitian

Tabel 3. Hasil pada domain sumber (#MakanBergiziGratis)

<i>Model</i>	<i>Accuracy</i>	<i>Precision</i>	<i>Recall</i>	<i>F1-Score</i>
<i>Naive Bayes</i> (TF-IDF)	0.94	0.93	0.94	0.93
KNN (FastText)	0.94	0.94	0.94	0.94

Tabel 3 menunjukkan bahwa kedua model mampu mengenali pola data dengan baik ketika data pelatihan dan pengujian berasal dari domain yang sama.

Tabel 4. Hasil pada *cross-domain* (#Danantara)

<i>Model</i>	<i>Accuracy</i>	<i>Precision</i>	<i>Recall</i>	<i>F1-Score</i>
<i>Naive Bayes</i> (TF-IDF)	0.75	0.74	0.75	0.74
KNN (FastText)	0.60	0.58	0.60	0.57

Hasil pengujian pada *domain* target (#Danantara) menunjukkan adanya penurunan performa. Penurunan ini terjadi akibat perbedaan distribusi data antara domain sumber dan domain target (*domain shift*).

Kesimpulan

Berdasarkan hasil penelitian yang telah dilakukan, dapat disimpulkan bahwa algoritma *Naïve Bayes* dengan representasi TF-IDF memiliki performa yang lebih stabil dibandingkan *K-Nearest Neighbors* (KNN) dengan *FastText* dalam skenario analisis sentimen *cross-domain*. Meskipun kedua model menunjukkan akurasi yang tinggi pada pengujian internal (0.94), performa mengalami penurunan pada pengujian lintas *domain* akibat adanya perbedaan distribusi data.

Naïve Bayes mampu mempertahankan kinerja dengan akurasi sebesar 0.75 dan tetap dapat mengklasifikasikan kedua kelas sentimen dengan baik. Sebaliknya, KNN mengalami penurunan performa yang signifikan, terutama dalam mendeteksi kelas negatif, yang menunjukkan keterbatasan metode berbasis jarak dalam menghadapi *domain shift*.

Hasil penelitian ini menegaskan bahwa pemilihan algoritma dan representasi fitur memiliki pengaruh yang signifikan terhadap kemampuan generalisasi model. Oleh karena itu, pendekatan probabilistik seperti *Naïve Bayes* dengan TF-IDF lebih direkomendasikan untuk kasus analisis sentimen lintas *domain*.

Adapun saran untuk penelitian selanjutnya adalah menggunakan metode yang lebih kompleks seperti *deep learning* (LSTM atau BERT), menerapkan teknik *domain adaptation*, serta menambahkan jumlah dan variasi dataset agar model dapat belajar representasi yang lebih umum dan meningkatkan performa pada skenario lintas *domain*.

Daftar Pustaka

- [1] P. K. Tetteh dan P. K. Kankam, "The role of social media in information dissemination to improve youth interactions," *Cogent Soc. Sci.*, vol. 10, no. 1, Des 2024, doi: 10.1080/23311886.2024.2334480.
- [2] F. Gaisbauer, A. Pournaki, S. Banisch, dan E. Olbrich, "Ideological differences in engagement in public debate on Twitter," *PLoS One*, vol. 16, no. 3, hlm. e0249241, Mar 2021, doi: 10.1371/journal.pone.0249241.
- [3] E. Purnama Harahap, H. Dwi Purnomo, A. Iriani, I. Sembiring, dan T. Nurtino, "Trends in sentiment of Twitter users towards Indonesian tourism: analysis with the k-nearest neighbor method," *Computer Science and Information Technologies*, vol. 5, no. 1, hlm. 19–28, Mar 2024, doi: 10.11591/csit.v5i1.p19-28.
- [4] N. M. Damayanti, "Analisis Sentimen Publik Pada Tagar #BTSCOMEBACK di Platform X Menggunakan IndoBERTweet," *Jurnal Informatika dan Teknik Elektro Terapan*, vol. 13, no. 3, Jul 2025, doi: 10.23960/jitet.v13i3.7176.
- [5] I. Iwandini, A. Triayudi, dan G. Soepriyono, "Analisa Sentimen Pengguna Transportasi Jakarta Terhadap Transjakarta Menggunakan Metode Naives Bayes dan K-Nearest Neighbor," *Journal of Information System Research (JOSH)*, vol. 4, no. 2, hlm. 543–550, Jan 2023, doi: 10.47065/josh.v4i2.2937.
- [6] M. A. Arsyad, T. Hasanuddin, dan M. Hasnawi, "Implementasi K-Nearest Neighbor (KNN) pada Sistem Pakar Diagnosa Penyakit Untuk Menentukan Kelayakan Sapi sebagai Hewan Qurban Berbasis Web," *Buletin Sistem Informasi dan Teknologi Islam*, vol. 3, no. 3, hlm. 167–173, Agu 2022, doi: 10.33096/busiti.v3i3.632.
- [7] K. Munawaroh dan A. Alamsyah, "Performance Comparison of SVM, Naïve Bayes, and KNN Algorithms for Analysis of Public Opinion Sentiment Against COVID-19 Vaccination on Twitter," *Journal of Advances in Information Systems and Technology*, vol. 4, no. 2, hlm. 113–125, Mar 2023, doi: 10.15294/jaist.v4i2.59493.
- [8] S. Fi. Nurul Fitri H, F. Fattah, dan H. Azis, "Comparative Analysis of Machine Learning Algorithm Variations in Classifying Body Shaming Topics on Social Media X," *Indonesian Journal of Data and Science*, vol. 5, no. 2, Jul 2024, doi: 10.56705/ijodas.v5i2.82.
- [9] F. P. Syah, T. Hasanuddin, dan N. Kurniati, "Implementasi Naive Bayes Untuk Analisis Sentimen Pada data Twitter Tentang Isu Politik di Indonesia," *LINIER: Literatur Informatika dan Komputer*, vol. 2, no. 3, hlm. 302–316, Okt 2025, doi: 10.33096/linier.v2i3.3142.
- [10] C. G. Özmen dan S. Gündüz, "Comparison of Machine Learning Models for Sentiment Analysis of Big Turkish Web-Based Data," *Applied Sciences*, vol. 15, no. 5, hlm. 2297, Feb 2025, doi: 10.3390/app15052297.
- [11] L. Xiao, Q. Li, Q. Ma, J. Shen, Y. Yang, dan D. Li, "Text classification algorithm of tourist attractions subcategories with modified TF-IDF and Word2Vec," *PLoS One*, vol. 19, no. 10, hlm. e0305095, Okt 2024, doi: 10.1371/journal.pone.0305095.
- [12] A. R. Sembiring dan C. K. Dewa, "Sentiment Analysis On Indonesian Tweets about the 2024 Election," *Sinkron*, vol. 9, no. 1, hlm. 413–422, Jan 2025, doi: 10.33395/sinkron.v9i1.14481.
- [13] E. Zuo, A. Aysa, M. Muhammat, Y. Zhao, dan K. Ubul, "Context aware semantic adaptation network for cross domain implicit sentiment classification," *Sci. Rep.*, vol. 11, no. 1, hlm. 22038, Nov 2021, doi: 10.1038/s41598-021-01385-1.
- [14] I. I. Saputri, P. Purnawansyah, dan H. Herman, "Implementasi Metode Naive Bayes Pada Pengenalan Tulisan Tangan Lontara," *Buletin Sistem Informasi dan Teknologi Islam*, vol. 2, no. 3, hlm. 167–175, Agu 2021, doi: 10.33096/busiti.v2i3.845.
- [15] H. Basri, P. Purnawansyah, H. Darwis, dan F. Umar, "Klasifikasi Daun Herbal Menggunakan K-Nearest Neighbor dan Convolutional Neural Network dengan Ekstraksi Fourier Descriptor," *Jurnal Teknologi dan Manajemen Informatika*, vol. 9, no. 2, hlm. 79–90, Des 2023, doi: 10.26905/jtmi.v9i2.10350.

- [16] L. Geni, E. Yulianti, dan D. I. Sensuse, "Sentiment Analysis of Tweets Before the 2024 Elections in Indonesia Using Bert Language Models," *Jurnal Ilmiah Teknik Elektro Komputer dan Informatika*, vol. 9, no. 3, hlm. 746–757, Agu 2023, doi: 10.26555/jiteki.v9i3.26490.
- [17] S. Wang, J. Zhou, Q. Chen, Q. Zhang, T. Gui, dan X. Huang, "Domain Generalization via Causal Adjustment for Cross-Domain Sentiment Analysis," dalam *Proceedings of the 2024 Joint International Conference on Computational Linguistics, Language Resources and Evaluation (LREC-COLING 2024)*, N. Calzolari, M.-Y. Kan, V. Hoste, A. Lenci, S. Sakti, dan N. Xue, Ed., Torino, Italia: ELRA and ICCL, Mei 2024, hlm. 5286–5298. [Daring]. Tersedia pada: <https://aclanthology.org/2024.lrec-main.470/>
- [18] M. Rostami, D. Bose, S. Narayanan, dan A. Galstyan, "Domain Adaptation for Sentiment Analysis Using Robust Internal Representations," dalam *Findings of the Association for Computational Linguistics: EMNLP 2023*, H. Bouamor, J. Pino, dan K. Bali, Ed., Singapore: Association for Computational Linguistics, Des 2023, hlm. 11484–11498. doi: 10.18653/v1/2023.findings-emnlp.769.
- [19] Y. Zhang, Y. Zhang, W. Guo, X. Cai, dan X. Yuan, "Learning Disentangled Representation for Multimodal Cross-Domain Sentiment Analysis," *IEEE Trans. Neural Netw. Learn. Syst.*, vol. 34, no. 10, hlm. 7956–7966, Okt 2023, doi: 10.1109/TNNLS.2022.3147546.
- [20] A. Novanto, D. Indra, dan W. Astuti, "Analisis Pre-processing Sentimen Terhadap Komentar Layanan Indihome Pada Twitter," *Buletin Sistem Informasi dan Teknologi Islam*, vol. 5, no. 1, hlm. 30–36, Feb 2024, doi: 10.33096/busiti.v5i1.2066.
- [21] Muh. R. Zulkifli, Purnawansyah, dan H. Darwis, "Indonesian Cross-Platform Sentiment Analysis: DANN Transfer from General Applications to TradingView," *Indonesian Journal of Data and Science*, vol. 6, no. 3, hlm. 481–491, Des 2025, doi: 10.56705/ijodas.v6i3.318.
- [22] M. Mujahid *dkk.*, "Data oversampling and imbalanced datasets: an investigation of performance for machine learning and feature engineering," *J. Big Data*, vol. 11, no. 1, hlm. 87, Jun 2024, doi: 10.1186/s40537-024-00943-4.
- [23] K. Pawluszek-Filipiak dan A. Borkowski, "On the Importance of Train–Test Split Ratio of Datasets in Automatic Landslide Detection by Supervised Classification," *Remote Sens. (Basel)*, vol. 12, no. 18, hlm. 3054, Sep 2020, doi: 10.3390/rs12183054.
- [24] N. Marchelo, T. N. Fatyanosa, dan R. S. Perdana, "Klasifikasi Sentimen Lintas Domain pada Komentar Media Sosial Menggunakan IndoBERT dengan Fine-Tuning dan Data Augmentation," *Jurnal Pengembangan Teknologi Informasi dan Ilmu Komputer*, vol. 9, no. 12, 2025.
- [25] L. Afuan, "Sentiment Analysis of the Kampus Merdeka Program on Twitter Using Support Vector Machine and a Feature Extraction Comparison: TF-IDF vs. FastText," *Journal of Applied Data Sciences*, vol. 5, no. 4, hlm. 1738–1753, Des 2024, doi: 10.47738/jads.v5i4.436.
- [26] H. Basri, P. Purnawansyah, H. Darwis, dan F. Umar, "Klasifikasi Daun Herbal Menggunakan K-Nearest Neighbor dan Convolutional Neural Network dengan Ekstraksi Fourier Descriptor," *Jurnal Teknologi dan Manajemen Informatika*, vol. 9, no. 2, hlm. 79–90, Des 2023, doi: 10.26905/jtmi.v9i2.10350.
- [27] A. Amelia, L. N. Hayati, dan H. Darwis, "Analisis Sentimen Masyarakat Terhadap Sistem Pembayaran MyPertamina dengan Metode Random Forest, SVM, dan Naïve Bayes," *LINIER: Literatur Informatika dan Komputer*, vol. 1, no. 1, hlm. 28–44, Nov 2024, doi: 10.33096/linier.v1i1.2269.
- [28] R. Puspitasari, T. Amaliah, dan H. Darwis, "Comparative Machine Learning Models for Dementia Prediction Using SMOTE," *International Journal of Artificial Intelligence in Medical Issues*, vol. 3, no. 2, hlm. 79–90, Nov 2025, doi: 10.56705/ijaimi.v3i2.351.
- [29] D. Pratmanto, F. F. D. Imanawan, dan V. Maarif, "Analisis Sentimen Pada Ulasan Pengguna Aplikasi Identitas Kependudukan Digital Dengan Metode Naive Bayes Dan K-Nearest," *Computatio : Journal of Computer Science and Information Systems*, vol. 7, no. 2, hlm. 155–166, Des 2023, doi: 10.24912/computatio.v7i2.26322.
- [30] H. Azis, M. S. Andini, Herman, L. Syafie, A. R. Manga', dan L. B. Ilmawan, "A Comparative Study of SGD, Adam, AdamW, and RMSprop Optimizers for VGG19-Based Rupiah Banknote Classification," dalam *2025 9th International Conference On Electrical, Electronics And Information Engineering (ICEEIE)*, IEEE, Sep 2025, hlm. 1–6. doi: 10.1109/ICEEIE66203.2025.11251942.