

ANALISIS SENTIMEN KOMENTAR MASYARAKAT TERHADAP PELAYANAN PUBLIK PEMERINTAH DKI JAKARTA DENGAN ALGORITMA SUPER VECTOR MACHINE DAN NAIVE BAYES

Raniya Rakarahayu Putri, Nuri Cahyono*

Program Studi S1 Informatika, Universitas Amikom Yogyakarta

nuricahyono@amikom.ac.id

ABSTRAK

Sebagai pelayan masyarakat, pemerintah memiliki peran yang penting dan tanggung jawab untuk menyediakan layanan yang memadai. Dengan meningkatnya jumlah pengguna Instagram di Indonesia memberikan kemudahan bagi masyarakat dalam berkomunikasi. Ekspresi masyarakat di Instagram terhadap pelayanan Pemerintah DKI Jakarta menunjukkan kompleksitas. Komentar-komentar mencerminkan ketidakpuasan dan kritik tajam terhadap layanan public. Penelitian ini mencoba melakukan klasifikasi komentar menjadi dua kelas positive dan negative dengan menerapkan dua metode klasifikasi, yaitu *Support Vector Machine* (SVM) dan *Multinomial Naive Bayes* (MNB), dalam analisis sentimen terhadap komentar masyarakat pada akun media sosial DKI Jakarta. Tahap pemodelan mencakup pembagian data latih dan uji dengan variasi rasio, dan metode SVM dievaluasi menggunakan kernel linear dan radial basis function (RBF) dengan grid search cross-validation. Hasil menunjukkan bahwa SVM memberikan akurasi yang sedikit lebih tinggi daripada MNB, mencapai 82%. Parameter optimal untuk SVM adalah $C=100$ dan $\gamma=0.1$ pada kernel RBF. Pada MNB, parameter $\alpha=2.0$ dan $\text{fit_prior}=\text{True}$ memberikan kinerja optimal dengan akurasi 80%. Evaluasi dilakukan dengan confusion matrix dan 10-folds cross validation. Meskipun SVM sedikit lebih unggul, penelitian selanjutnya direkomendasikan untuk eksplorasi teknik-teknik baru dalam machine learning dan pengembangan model klasifikasi yang lebih kompleks. Penelitian ini memberikan kontribusi penting dalam memahami sentimen masyarakat terhadap pelayanan publik DKI Jakarta, membuka pintu untuk pengembangan analisis sentimen lebih lanjut dengan metode-metode machine learning.

Kata kunci : *sentiment analisis, multinomial naïve bayes, super vector machine, grid search*

1. PENDAHULUAN

Perkembangan masyarakat modern di era teknologi informasi telah mengubah dinamika interaksi dan ekspresi pandangan terhadap kehidupan. Media sosial, khususnya Instagram, menjadi kanal utama bagi masyarakat untuk menyuarakan opini. Menurut data Napoleon Cat, jumlah pengguna Instagram di Indonesia mencapai 97,38 juta orang pada Oktober 2022, menunjukkan popularitas *platform* ini, terutama karena fokus pada aspek visual[1].

Pentingnya memberikan layanan yang baik dan profesional kepada masyarakat merupakan aspek fundamental baik bagi instansi pemerintahan maupun swasta. Sebagai pelayan masyarakat, pemerintah memiliki peran yang penting dan tanggung jawab untuk menyediakan layanan yang memadai. Layanan ini sangat dibutuhkan oleh masyarakat untuk memenuhi berbagai kebutuhan mereka, yang merupakan bagian tak terpisahkan dari kehidupan manusia[2].

Namun, di balik kemudahan berkomunikasi, ekspresi masyarakat di Instagram terhadap pelayanan Pemerintah DKI Jakarta menunjukkan kompleksitas. Komentar-komentar mencerminkan ketidakpuasan dan kritik tajam terhadap layanan publik, tetapi juga terdapat apresiasi terhadap upaya pemerintah. Fenomena ini menciptakan kebutuhan untuk pendekatan analisis sentimen yang holistic[16].

Sebagai respons terhadap isu ini, penelitian sebelumnya telah dilakukan menggunakan metode *Naive Bayes* untuk menganalisis sentimen terkait pelayanan pemerintah[2]. Hasilnya, akurasi sebesar 69.23% untuk KemenPUPR dan 64.10% untuk Kemenkeu diakui, namun penelitian menyarankan perlunya eksplorasi lebih lanjut dengan algoritma lain. Penelitian kedua yang dilakukan oleh Rani Yunita dan Mia Kamayani pada tahun 2023 yaitu membandingkan hasil evaluasi algoritma *Support Vector Machine* (SVM) dengan *Naive Bayes* menggunakan 80% data latih dan 20% data uji. Hasilnya menunjukkan bahwa terdapat 331 sentimen positif dan 369 sentimen negatif. Dari hasil tersebut, ditarik kesimpulan bahwa *Support Vector Machine* (SVM) menjadi algoritma yang terbaik dengan akurasi 80%, recall 83%, presisi 76%, dan F1-Score 79% [3].

Penelitian terakhir oleh Akbar Zaiem Praghakusma dan Novrido Charibaldi pada tahun 2021 melakukan Komparasi Fungsi Kernel Metode *Support Vector Machine* untuk Analisis Sentimen Instagram dan Twitter. Hasilnya menunjukkan bahwa kernel linier memiliki akurasi tertinggi 83.06%, kernel sigmoid memiliki presisi tertinggi 91.93%, dan kernel polinomial memiliki recall tertinggi 91.17%. Penelitian ini memberikan wawasan pada kinerja SVM dengan berbagai kernel dalam konteks analisis sentimen terhadap lembaga pemberantasan korupsi di Indonesia [15].

Melihat hasil dan temuan dari penelitian-penelitian sebelumnya, terlihat bahwa baik algoritma *Super Vector Machine* maupun *Naïve Bayes* menunjukkan tingkat akurasi yang bervariasi dalam setiap kasus. Oleh karena itu, penelitian ini juga bertujuan untuk membandingkan kinerja kedua algoritma tersebut dalam menganalisis sentimen masyarakat terhadap layanan publik DKI Jakarta, dengan harapan dapat memberikan kontribusi pada pemahaman lebih mendalam terhadap permasalahan ini. Selain itu, penelitian ini akan menerapkan teknik *Hyperparameter tuning* untuk meningkatkan konsistensi dan prediksi model sesuai dengan target yang diharapkan.

2. TINJAUAN PUSTAKA

2.1. Text Mining

Text Mining merupakan penemuan pengetahuan di database dalam bentuk tekstual (*knowledge discovery in textual database* atau disingkat dengan KDT) [4]. bisa disebut juga penggalian atau pencarian data-data yang berbentuk teks[5], yang di antaranya adalah ketertarikan terhadap pengetahuan yang baru dibuat, diartikan sebagai bagian dari proses penggalian atau pencarian data text yang sebelumnya tidak diketahui, sehingga dapat dimengerti, mempunyai potensi dan pola praktis atau pengetahuan dari koleksi data teks atau *corpus* masif (kumpulan teks yang menangkap penggunaan bahasa dalam bentuk tertulis atau lisan secara utuh dan padat) dan tidak terstruktur. Text mining merupakan lingkup penelitian yang komprehensif, yang masuk hampir disetiap lini kehidupan kita. Tujuan menyeluruh dalam text mining, pada dasarnya adalah untuk mengubah teks menjadi data untuk dianalisis melalui aplikasi dengan suatu metode analisis [6].

2.2. Analisis Sentimen

Analisis sentimen merupakan salah satu contoh dari bidang *Natural Language Processing* (NLP) yang paling populer. *Natural Language Processing* adalah bidang ilmiah yang membahas tentang bagaimana caranya agar komputer bisa bekerja dan berpikir seperti manusia. *Natural Language Processing* merupakan bagian dari *Artificial Intelligence* atau kecerdasan buatan. Dalam perkembangan data mining, *Artificial Intelligence* (AI) merupakan salah satu dari empat cabang ilmu data mining, yaitu statistika, database, dan pencarian informasi. Dalam penerapannya, AI juga memerlukan *machine learning* sebagai algoritma penyelesaian. Adanya *machine learning* digunakan untuk menggantikan manusia dalam mengambil keputusan. Machine learning tidak mempunyai perasaan seperti manusia sehingga keputusan yang diambil berdasarkan data yang sudah diolah [7].

2.3. Text Preprocessing

Text Preprocessing adalah suatu tahap pertama dalam mengolah teks untuk digunakan dalam

pengubahan dokumen menjadi data yang terstruktur sesuai dengan kebutuhannya agar dapat diolah lebih lanjut dalam proses text mining [8]. Tahapan dalam pra-proses teks adalah sebagai berikut.

- Cleansing*, merupakan proses membersihkan tweet dari kata yang tidak diperlukan untuk mengurangi noise.
- Case Folding*, merupakan proses untuk mengubah semua karakter teks menjadi huruf kecil serta menghilangkan tanda baca dan angka.
- Stemming*, merupakan proses untuk mendapatkan kata dasar dengan cara menghilangkan awalan, akhiran, sisipan, dan *confixes* (kombinasi dari awalan dan akhiran).
- Stopwords*, merupakan kosakata yang bukan termasuk kata unik atau ciri pada suatu dokumen atau tidak menyampaikan pesan apapun secara signifikan pada teks atau kalimat.
- Tokenizing*, merupakan proses memecah yang semula berupa kalimat menjadi kata-kata.

2.4. Pembobotan Kata TF-IDF

Term Frequency Inverse Document Frequency (TF-IDF) adalah sebuah cara yang digunakan untuk menghitung nilai frekuensi sebuah term atau kata pada sebuah dokumen atau banyak dokumen [9]. Semakin sering kata tersebut muncul, semakin tinggi nilai TF-nya. Document Frequency (DF) merupakan variabel yang bertujuan untuk menampung jumlah kemunculan suatu term pada keseluruhan dokumen [10]. Rumus yang digunakan dalam memperoleh nilai TF-IDF dituliskan dengan persamaan (1), (2) dan persamaan (3).

$$tf_{ij} = \frac{f_{d(i)}}{\max f_{d(i)}} \quad (1)$$

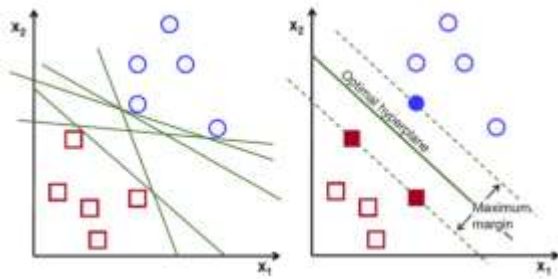
$$idf(t, D) = \log \left(\frac{N}{|\{d \in D : t \in d\}|} \right) \quad (2)$$

$$idf(t, D) = \log \left(\frac{N}{df(t)+1} \right) \quad (3)$$

Pada rumus (3) Penambahan 1 untuk menghindari pembagian terhadap 0 jika $df(t)$ tidak ditemukan pada corpus.

2.5. Super Vector Machine

Support Vector Machine (SVM) diperkenalkan pertama kali pada akhir 1979 oleh Vapnik dan kemudian mendapat perhatian lebih lanjut pada tahun 1992 bersama dengan Boser [4]. SVM merupakan metode dalam machine learning yang digunakan untuk prediksi dan klasifikasi. Ide dasar SVM adalah menggabungkan konsep-konsep sebelumnya dalam penyelesaian masalah klasifikasi dan prediksi. Konsep dasar SVM adalah menemukan hyperplane pemisah yang dapat memisahkan dua kelas secara optimal, dengan persamaan hyperplane dianggap baik jika memiliki margin terbesar [14].



Gambar 1. *Hyperplane Supervector Machine*

2.6. Multinomial Naïve Bayes

Multinomial naïve bayes digunakan dalam klasifikasi teks, di mana frekuensi penggunaan kata memiliki peran krusial dalam memprediksi atau memberi label pada data. Metode ini dapat mengklasifikasikan data ke dalam dua kelas atau lebih, dan khususnya efektif ketika diterapkan pada data kategori. *Multinomial naïve bayes* menunjukkan kinerja yang lebih baik dibanding jenis lainnya dalam kasus data teks, terutama ketika teknik vektorisasi seperti tf-idf dan countvectorize digunakan untuk membentuk variabel bebas. Multinomial Naïve Bayes dapat diformulasikan sebagai berikut[11].

$$P(p|n) \propto P_{(p)} \prod_{1 \leq k \leq nd} P(t_k | p) \quad (4)$$

$$P(t_k | p) = \frac{\text{count}(t_k | p) + 1}{\text{count}(t_p) + |V|} \quad (5)$$

Keterangan:

$P(t_k | p)$ = probabilitas munculnya dokumen text (t_k), n adalah jumlah dokumen dan p adalah polaritas.

$(t_k | p)$ = jumlah t_k muncul di dokumen text yang memiliki polaritas p dan jumlah (tp) berarti jumlah token yang ada di artikel berita dengan polaritas p .

2.7. Grid Search Validation

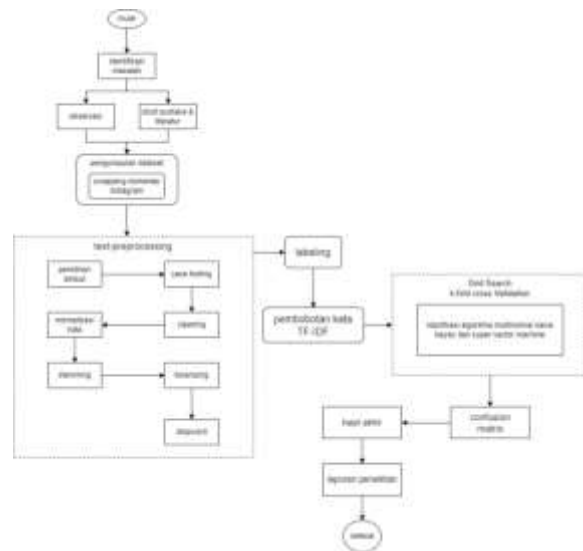
Grid Search Cross Validation (GridSearchCV) adalah metode untuk memilih kombinasi model dan hyperparameter dengan menguji setiap kombinasi dan melakukan validasi. Metode ini secara otomatis memvalidasi setiap kombinasi model dan hyperparameter, menghemat waktu pemrosesan. Hyperparameter yang menghasilkan akurasi terbaik dengan nilai error terkecil dianggap sebagai hyperparameter optimal [12]. Grid Search membagi jangkauan parameter ke dalam grid, menjelajahi setiap titik untuk mendapatkan parameter optimal. Dalam konteks SVM, Grid Search mengoptimalkan hyperparameter menggunakan teknik cross validation sebagai metrik kinerja. Tujuannya adalah mengidentifikasi kombinasi hyperparameter yang optimal untuk classifier agar dapat memprediksi data yang tidak diketahui dengan akurat.

2.8. Confusion Matrix

Cofusion Matrix adalah tabel yang dapat menggambarkan kinerja algoritma dalam mengklasifikasikan kelas-kelas [13].

3. METODE PENELITIAN

Penelitian ini memuat diagram alur yang mengilustrasikan langkah-langkah penelitian secara visual, memudahkan pemahaman dengan menampilkan tahapan-tahapan yang akan dilakukan.



Gambar 2. Alur Penelitian

Metode penelitian dimulai dengan merumuskan masalah, melakukan studi literatur, dan mengumpulkan data komentar masyarakat dari akun Instagram Pemerintah DKI Jakarta. Setelah itu, dilakukan seleksi data berdasarkan kriteria tertentu. Data tekstual kemudian diproses melalui tahap preprocessing untuk membersihkan data, dan dilakukan analisis sentimen menggunakan lexicon-based. Hasilnya diberi label pada setiap dokumen. Langkah selanjutnya melibatkan pembobotan token menggunakan metode tfidf untuk persiapan analisis menggunakan metode *Multinomial Naïve Bayes* dan *Support Vector Machine*.

3.1. Pengumpulan Data

Data diambil dari Instagram menggunakan skrip Instagram Post Extractor pada website Phantombuster, kemudian disimpan dalam file csv. Fokus analisis sentimen adalah pada komentar masyarakat di akun Instagram @dkijakarta, untuk menilai opini terkait pelayanan publik Pemprov DKI Jakarta. Langkah awal adalah memilih atribut komentar menggunakan operator select attributes untuk diproses pada tahap selanjutnya.

3.2. Text-Preprocessing

Pre-processing data untuk analisis sentimen melibatkan langkah-langkah seperti pembersihan teks untuk menghilangkan karakter tidak relevan,

tokenisasi untuk memecah teks menjadi unit-unit yang lebih kecil, dan normalisasi untuk memastikan konsistensi dalam representasi kata. Tujuan pre-processing adalah meningkatkan kualitas data agar hasil analisis sentimen lebih akurat dan dapat diandalkan

3.2.1. Case Folding

Case folding adalah proses di mana semua karakter dalam suatu teks atau komentar diubah menjadi huruf kecil atau huruf kapital. Tujuannya adalah untuk menghindari perbedaan penulisan yang tidak konsisten, sehingga kata-kata dengan penulisan yang berbeda tetapi memiliki arti yang sama dapat dikenali dengan benar dalam analisis.

3.2.2. Cleaning

Pembersihan data, atau data cleaning, adalah proses krusial dalam pengolahan data yang melibatkan penghapusan simbol atau karakter yang tidak relevan untuk memastikan ketepatan klasifikasi data[6]. Keseluruhan keandalan dataset dapat dicapai melalui tahap pembersihan data yang dilakukan pada awal proses. Kegiatan dalam pembersihan data mencakup penghapusan simbol, kata atau simbol yang tidak terkait dengan bahasa alami, seperti @mention, #hashtag, URL, dan karakter lainnya. Selain itu, penanganan data kosong menjadi fokus penting dengan menghapus baris yang mengandung data kosong, menghindari potensi outlier yang dapat mempengaruhi akurasi analisis. Penghapusan data duplikat juga menjadi langkah esensial untuk mencegah dampak data ganda terhadap tingkat akurasi hasil analisis sentimen.

3.2.3. Normalisasi Kata

Normalisasi kata yang disingkat dan diperbaiki bertujuan untuk mengatasi variasi dalam penulisan kata yang umumnya terjadi dalam media sosial atau teks informal. Proses normalisasi ini dilakukan untuk memastikan bahwa kata-kata yang sebenarnya memiliki makna yang sama atau serupa dapat dikenali dan diproses dengan benar dalam analisis sentimen.

3.2.4. Stemming

Pada tahap ini, dilakukan pengubahan kata-kata dalam komentar menjadi bentuk dasarnya (stem) dengan menghilangkan imbuhan atau akhiran kata. Tujuannya adalah untuk mengurangi variasi kata yang sebenarnya memiliki akar kata yang sama.

3.2.5. Tokenizing

Tokenizing, merupakan proses memecahkan data menggunakan spasi untuk dijadikan token-token. Proses ini merupakan proses akhir untuk mendapatkan data yang akan diolah untuk sentimen analisis.

3.2.6. Stopwords

Stopword, merupakan proses menghilangkan kata-kata yang tidak mendeskripsikan sesuatu yang semestinya dihilangkan.

3.2.7. Lexicon Based Labeling

Pada tahap ini pelabelan dataset dilakukan dengan cara mengecek setiap kata (token) pada dokumen berdasarkan kamus atau lexicon (kumpulan kata-kata yang memiliki bobot sentimen positif atau negatif).

3.3. Labeling

Supervised Learning membutuhkan *class* atau target untuk setiap *record* pada *dataset*. Dalam konteks penelitian ini, kelas merujuk pada label sentimen yang terkait dengan setiap ulasan. Proses penentuan label dilakukan secara otomatis menggunakan metode auto-labelling dengan menerapkan kamus kata (lexicon) untuk mendapatkan nilai sentimen dari setiap ulasan. Nilai sentimen dari setiap elemen menentukan arah sentimen yang terungkap, yakni apakah teks tersebut menyampaikan sentimen positif, negatif, atau netral.

3.4. Pembobotan Kata TF-IDF

Feature-Weighting (Pembobotan Fitur)

- Term Frequency* (TF) mencerminkan sejauh mana suatu kata muncul dalam suatu kalimat.
- Inverse Document Frequency* (IDF) mengukur seberapa umum kemunculan suatu kata dalam seluruh korpus kalimat.

3.5. Grid Search

Dalam konteks machine learning, grid search digunakan untuk menemukan parameter terbaik bagi suatu model. Misalnya, jika kita memiliki dua parameter yang ingin dioptimalkan, seperti alpha dan beta, maka grid search akan menciptakan suatu grid dengan berbagai kombinasi nilai alpha dan beta yang akan diuji pada model. Setelah model diuji dengan setiap kombinasi, kita dapat menilai performa model (misalnya, akurasi) dan memilih kombinasi parameter yang memberikan hasil terbaik.

3.6. K-Fold Cross Validation

Pada tahap ini dataset akan dikelompokkan menjadi k-subset secara acak, setiap kelompok disebut dengan fold, salah satu dari fold akan dijadikan data uji dan sisanya dijadikan sebagai data latih, dilakukan proses iterasi sebanyak k-kali setiap iterasi akan dicatat performanya untuk diambil nilai rata-rata.

3.7. Klasifikasi Multinomial Naïve Bayes dan Super Vector Machine

Metode klasifikasi yang digunakan dalam penelitian ini adalah multinomial naïve bayes dan SVM dengan penerapan *hyperparameter tuning* menggunakan *GridSearchCV*. Selain itu, pembuatan

model tanpa *hyperparameter tuning* juga dilakukan sebagai pembandingan *GridSearchCV* adalah sebuah metode dalam *library sklearn* yang digunakan untuk mencari kombinasi *hyperparameter* terbaik menggunakan teknik *cross validation*.

Algoritma Naïve bayes yang digunakan dalam penelitian ini adalah multinomial naïve bayes dengan parameter α . Terdapat 4 kombinasi nilai α , yaitu (0.1, 0.5, 1.0, 1.5). Lalu algoritma SVM, peneliti akan menggunakan dua kernel, yaitu linear, dan RBF. Untuk parameter yang digunakan ada empat kombinasi nilai C (0.1, 1, 10, 100) dan 4 kombinasi nilai γ (0.001, 0.01, 0.1, 1). Semua parameter tersebut akan dilatih dengan data training menggunakan 3 *cross validation* agar mendapatkan model dengan parameter terbaik.

3.8. Confusion Matrix

Penilaian efektivitas suatu model klasifikasi dapat ditentukan melalui parameter evaluasi kinerja, seperti akurasi, recall, f1-score, dan presisi. Perhitungan parameter-parameter tersebut memerlukan sebuah matriks yang umumnya dikenal sebagai *confusion matrix*.

4. HASIL DAN PEMBAHASAN

Adapun hasil dan pembahasan yang akan dijelaskan menjadi beberapa sub bab.

4.1. Pengumpulan Data

Dalam penelitian ini, data komentar masyarakat dari akun media sosial seperti @dkijakarta dikumpulkan menggunakan tools seperti Phantombuster. Proses dimulai dengan menetapkan

tujuan penelitian, seperti analisis sentimen atau evaluasi pandangan masyarakat terhadap topik tertentu. Setelah itu, parameter pencarian, termasuk rentang waktu komentar relevan, ditentukan (dalam kasus ini, 10 Juni - 17 September). Phantombuster digunakan untuk membuat skrip scraping khusus, dengan Instagram Profile Post Extractor mengekstrak data postingan dan Instagram Post Commenters Export mengekstrak informasi komentator dari akun @dkijakarta. Skrip ini memanfaatkan API Instagram untuk mengambil data publik, memungkinkan pengumpulan data secara otomatis.

4.2. Preprocessing Data

Dalam tahap pre-processing data, dataset komentar masyarakat dari akun Instagram @dkijakarta, yang awalnya berjumlah 2526, mengalami pemilihan hingga tersisa 2458 data yang relevan. Proses ini melibatkan beberapa langkah, dimulai dengan Case Folding untuk menyamakan huruf menjadi huruf kecil. Tahap selanjutnya adalah Cleaning, di mana data dibersihkan dari elemen-elemen tidak relevan seperti URL, username, hashtag, tanda baca, simbol, emoji, dan angka, serta menghapus data duplikat. Normalisasi kata dilakukan untuk menangani variasi penulisan kata, dan stemming untuk mengubah kata-kata ke bentuk dasarnya. Terakhir, tokenizing digunakan untuk memecah teks komentar menjadi unit-unit kecil, memfasilitasi analisis sentimen dan penghitungan frekuensi kata. Proses ini menghasilkan dataset final sebanyak 2149 entri setelah pre-processing.

Tabel 1. Hasil Preprocessing

Proses	Hasil
Case Folding	untuk transparansi, bagaimana kalau @dkijakarta @humaspajakjakarta menempelkan stiker/label/pengumuman ditempat parkir yang dipajaki. sehingga pengguna dapat membedakan mana parkir resmi, dan mana parkir tidak resmi. lebih keren lagi kalau pembayarannya bisa menggunakan qris (seperti jukir e-parkir di @pemko.medan @dishubmedan.bidparkir). kayaknya tidak sulit diterapkan kalau memang mau tertib toh? cc: @herubudihartono @dishubdkijakarta @saberpungli_r
Cleaning	untuk transparansi bagaimana kalau menempelkan stiker label pengumuman ditempat parkir yang dipajaki sehingga pengguna dapat membedakan mana parkir resmi dan mana parkir tidak resmi lebih keren lagi kalau pembayarannya bisa menggunakan qris seperti jukir eparkir di medan bidparkir kayaknya tidak sulit diterapkan kalau memang mau tertib toh cc
Normalisasi Kata	untuk transparansi bagaimana kalau menempelkan stiker label pengumuman ditempat parkir yang dipajaki sehingga pengguna dapat membedakan mana parkir resmi dan mana parkir tidak resmi lebih keren lagi kalau pembayarannya bisa menggunakan qris seperti jukir eparkir di medan bidparkir kayaknya tidak sulit diterapkan kalau memang mau tertib toh cc
Stemming	untuk transparansi bagaimana kalau tempel stiker label umum tempat parkir yang pajak sehingga guna dapat beda mana parkir resmi dan mana parkir tidak resmi lebih keren lagi kalau bayar bisa guna qris seperti jukir eparkir di medan bidparkir kayak tidak sulit terap kalau memang mau tertib toh cc
Tokenizing	['untuk', 'transparansi', 'bagaimana', 'kalau', 'tempel', 'stiker', 'label', 'umum', 'tempat', 'parkir', 'yang', 'pajak', 'sehingga', 'guna', 'dapat', 'beda', 'mana', 'parkir', 'resmi', 'dan', 'mana', 'parkir', 'tidak', 'resmi', 'lebih', 'keren', 'lagi', 'kalau', 'bayar', 'bisa', 'guna', 'qris', 'seperti', 'jukir', 'eparkir', 'di', 'medan', 'bidparkir', 'kayak', 'tidak', 'sulit', 'terap', 'kalau', 'memang', 'mau', 'tertib', 'toh', 'cc']
Stopwords	['transparansi', 'tempel', 'stiker', 'label', 'parkir', 'pajak', 'beda', 'parkir', 'resmi', 'parkir', 'resmi', 'keren', 'bayar', 'qris', 'jukir', 'eparkir', 'medan', 'bidparkir', 'kayak', 'sulit', 'terap', 'tertib', 'cc']

4.3. Labeling

Pelabelan pada penelitian ini menggunakan pendekatan *lexicon based*, dengan cara menghitung bobot setiap kata pada ulasan. Lexicon atau kamus kata yang digunakan adalah Indonesia Setimen Lexicon yang populer dengan sebutan InSet Lexicon.

Tabel 2. Hasil Labeling

Hasil Stopword	Total Bobot	Orientasi Sentimen
['transparansi', 'tempel', 'stiker', 'label', 'parkir', 'pajak', 'beda', 'parkir', 'resmi', 'parkir', 'resmi', 'keren', 'bayar', 'qris', 'jukir', 'eparkir', 'medan', 'bidparkir', 'kayak', 'sulit', 'terap', 'tertib', 'cc']	-3	Negative

4.4. Pembobotan Kata TF-IDF

Hasil dari preprocessing ini berupa kata dasar dari setiap sentimen yang ada yang nantinya akan di proses kembali menggunakan metode TF-IDF untuk mendapatkan bobot untuk setiap kata yang ada.

Tabel 3. Hasil TF-IDF

Term	Rank
jakarta	98.281029
min	94.111450
iya	77.614135
nya	63.948575
moga	62.139727

4.5. Klasifikasi Multinomial Naïve Bayes dan Super Vector Machine

Tahap pemodelan pada penelitian ini menggunakan dua model yaitu metode *Multinomial Naïve Bayes* dan *Support Vector Machine* untuk klasifikasi.

4.5.1. Super Vector Machine

Dalam tahap ini, dataset dibagi menjadi data latih dan data uji dengan variasi rasio percobaan seperti 90% data latih dan 10% data uji, 80% data latih dan 20% data uji, 70% data uji dan 30% data latih, kemudian 60% data uji dan 40% data latih hingga 50% data latih dan 50% data uji. Metode Support Vector Machine (SVM) diterapkan dengan kernel linear dan radial basis function (RBF). Untuk mencari parameter optimal pada kedua kernel, dilakukan eksperimen menggunakan grid search cross validation. Rentang nilai cost yang dieksplorasi adalah 0.1, 1, 10, dan 100, sedangkan rentang nilai gamma adalah 0.001, 0.01, 0.1, dan 1. Rincian mengenai kernel dan parameter dievaluasi terdapat dalam tabel berikut.

Tabel 4. Grid Search SVM

Parameter	Konfigurasi
Karnel	[Linear, RBF]
C	[0.1, 1, 10, 100]
gamma	[0.001, 0.01, 0.1, 1]

Metode Support Vector Machine (SVM) dengan kernel linear dan rbf dievaluasi menggunakan grid search cross validation untuk menemukan parameter optimal. Hasil grid search menunjukkan bahwa parameter terbaik untuk kernel rbf adalah C=100 dan gamma=0.1, dengan nilai akurasi sebesar 0.7877, presisi sebesar 0.7910, dan recall sebesar 0.7828. Kernel linear memberikan kinerja yang konsisten dengan nilai akurasi, presisi, dan recall yang tinggi. Proses pemilihan parameter terbaik berfokus pada kriteria kinerja keseluruhan yang diukur melalui akurasi, presisi, dan recall.

Tabel 5. Akurasi, Presisi, Recall Kernel RBF

Data train : data test	Akurasi	Presisi	Recall
90%:10%	82%	82%	82%
80%:20%	81%	81%	81%
70%:30%	81%	81%	81%
60%:40%	78%	78%	78%
50%:50%	77%	77%	77%

Kinerja terbaik pada kernel rbf adalah pada rasio data set 90:10 dengan nilai akurasi adalah 82%. Hasil *confusion matrix* dari metode *support vector machine* pada rasio 90%:10% didapatkan bahwa prediksi benar pada sentimen positif atau *true positive* sebanyak 96, dan prediksi benar pada sentimen negatif atau *true negative* sebanyak 81.

Tabel 6. Hasil Confussion Matrix SVM

Hasil Aktual	Nilai prediksi	
	0	1
0	81	22
1	16	96

Dari *confusion matrix* tersebut, juga dilakukan perhitungan akurasi, presisi dan recall.

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN} \times 100$$

$$= \frac{81 + 96}{81 + 96 + 22 + 16} = \frac{96}{115} = 82.93\%$$

$$Precision \text{ kelas negative} = \frac{TN}{TN + FN} = \frac{96}{96 + 16} = 85.71\% \quad (4.2)$$

$$Precision \text{ kelas positive} = \frac{TP}{TP + FP} = \frac{81}{81 + 22} = 78.57\% \quad (4.3)$$

$$Recall \text{ kelas negative} = \frac{TN}{TN + FP} = \frac{96}{96 + 22} = 81.82\% \quad (4.4)$$

$$Recall \text{ kelas positive} = \frac{TP}{TP + FN} = \frac{81}{81 + 16} = 83.67\% \quad (4.5)$$

Setelah mendapatkan rasio dengan hasil rata-rata akurasi, presisi, dan recall terbesar, kemudian dilakukan cross validation untuk mendapatkan akurasi yang maksimal. Pada penelitian ini digunakan 10-folds cross validation dengan hasil akurasi, presisi, dan recall tertinggi pada n ke-5 sebesar 86%, 86%, dan 86%.

Tabel 7. Hasil dari 10-folds cross validation SVM

n	Akurasi	Presisi	Recall
1	81%	81%	81%
2	80%	80%	79%
3	85%	86%	85%
4	80%	80%	80%
5	86%	86%	86%
6	84%	84%	84%
7	82%	82%	82%
8	85%	86%	85%
9	84%	84%	83%
10	80%	80%	80%

4.5.2. Multinomial Naïve Bayes

Pada langkah ini, dataset dipisahkan menjadi dua bagian, yakni data latih dan data uji, dengan variasi rasio percobaan yang berbeda, termasuk 90% data latih dan 10% data uji, 80% data latih dan 20% data uji, 70% data latih dan 30%, 60% data latih dan 40%, serta 50% data latih dan 50% data uji. Dalam penelitian ini, metode Multinomial Naive Bayes dieksplorasi menggunakan *grid search* untuk menemukan parameter optimal. Grid parameter yang dijelajahi mencakup nilai *alpha* sebesar 0.1, 0.5, 1.0, dan 2.0 dengan *Fit_prior True* dan *False*.

Tabel 8. Grid Search SVM

Parameter	Konfigurasi
Multinomialnb_alpha	[0.1, 0.5, 1.0, 2.0]
Fit_prior	[True, False]

Dalam proses analisis menggunakan metode klasifikasi *Multinomial Naive Bayes* untuk melakukan klasifikasi sentimen, langkah awal yang krusial adalah pembagian data *train* dan data *test* dengan berbagai variasi rasio, hasil akurasi dari pembagian data 90%:10% mendapatkan akurasi sebesar 80%, kemudian untuk pembagian data 80%:20% menghasilkan akurasi sebesar 79%, kemudian skala pembagian data 70%:30% mendapatkan akurasi 78%, pada pembagian data 60%:40% didapatkan hasil akurasi 77%, dan pada pembagian data 50%:50% menghasilkan akurasi sebesar 76%.

Pada metode klasifikasi Multinomial Naive Bayes, digunakan parameter-parameter tertentu, seperti *alpha* dengan nilai [0.1, 0.5, 1.0, 2.0], serta *fit_prior* dengan nilai [True, False]. Proses penentuan parameter terbaik dilakukan melalui metode *grid search cross validation* untuk mengidentifikasi kombinasi parameter yang optimal. Hasil *grid search cross validation* untuk setiap kombinasi parameter dapat ditemukan pada table dibawah ini.

Parameter terbaik dipilih berdasarkan kombinasi yang memberikan kinerja paling optimal. Dilakukan beberapa pembagian rasio dataset dengan menggunakan nilai *alpha* 2.0 dimana *fit_prior* kondisi true.

Tabel 9. Akurasi, Presisi, Recall alpha 2.0

Data train : data test	Akurasi	Presisi	Recall
90%:10%	80%	80%	80%
80%:20%	80%	80%	80%
70%:30%	78%	79%	78%
60%:40%	77%	77%	77%
50%:50%	76%	76%	76%

Kinerja terbaik pada *alpha* 2.0 adlah pada rasio data set 90:10 dengan nilai akurasi adalah 80%. Hasil *confusion matrix* dari metode *multinomial naïve bayes* pada rasio 90% : 10% didapatkan bahwa prediksi benar pada sentimen positif atau *true positive* sebanyak 77, dan prediksi benar pada sentimen negatif atau *true negative* sebanyak 95.

Tabel 10. Hasil Confussion Matrix MNB

Hasil Aktual	Nilai prediksi	
	0	1
0	77	26
1	17	95

Dari *confusion matrix* tersebut, juga dilakukan perhitungan akurasi, presisi dan recall.

$$Accuracy = \frac{TP+TN}{TP+TN+FP+FN} \times 100 = \frac{77+95}{77+95+26+17} = 80.27\%$$

$$Precision \text{ kelas negative} = \frac{TN}{TN+FN} = \frac{95}{95+17} = 84.78\%$$

$$Precision \text{ kelas positive} = \frac{TP}{TP+FP} = \frac{77}{77+26} = 74.77\%$$

$$Recall \text{ kelas negative} = \frac{TN}{TN+FP} = \frac{95}{95+26} = 78.81\%$$

$$Recall \text{ kelas positive} = \frac{TP}{TP+FN} = \frac{77}{77+17} = 81.95\%$$

Setelah mendapatkan rasio dengan hasil rata-rata akurasi, presisi, dan recall terbesar, kemudian dilakukan *cross validation* untuk mendapatkan akurasi yang maksimal. Pada penelitian ini digunakan 10-folds *cross validation* dengan hasil akurasi, presisi, dan recall tertinggi pada n ke-3 sebesar 82%, 83%, dan 82%.

Tabel 11. Hasil dari 10-folds cross validation MNB

n	Akurasi	Presisi	Recall
1	78%	78%	77%
2	79%	79%	79%
3	82%	83%	82%
4	77%	78%	77%
5	78%	78%	78%
6	79%	78%	79%
7	81%	82%	82%
8	78%	79%	78%
9	78%	78%	77%
10	79%	78%	78%

4.6. Hasil

Hasil *crawling data* sejumlah 2526 data mentah lalu dilakukan *text pre-processing* sehingga menjadi 2149 data bersih. Selanjutnya, dilakukan pelabelan 10.218 *data set* dengan metode *lexicon*. Berdasarkan hasil pelabelan didapatkan bahwa jumlah komentar negative sebanyak 1136 dan komentar positive sebanyak 1013. Berdasarkan proses pengujian klasifikasi dengan metode Multinomial Naive Bayes dan Support Vector Machine pada data komentar didapatkan hasil akurasi dari kedua metode tersebut pada Tabel 4.30. Hasil dari penelitian menggunakan klasifikasi teks data komentar masyarakat apda akun dki jakarta menggunakan metode Multinomial Naive Bayes mendapatkan akurasi 80%. Sementara itu, akurasi yang didapatkan dengan menggunakan metode Support Vector Machine adalah 82%.

Tabel 12. Kinerja Multinomial NB dan Support Vector Machine

Rasio	SVM			MNB		
	Akurasi	Presisi	Recall	Akurasi	Presisi	Recall
90:10	82%	82%	82%	80%	80%	80%
80:20	81%	81%	81%	80%	80%	80%
70:30	81%	81%	81%	78%	79%	78%
60:40	78%	78%	78%	77%	77%	77%
50:50	77%	77%	77%	76%	76%	76%

5. KESIMPULAN DAN SARAN

Berdasarkan penelitian ini, dapat disimpulkan bahwa metode klasifikasi yang digunakan, yaitu Support Vector Machine (SVM) dan Multinomial Naive Bayes (MNB), memberikan hasil yang baik dalam melakukan klasifikasi sentimen terhadap komentar masyarakat pada akun DKI Jakarta di media sosial. SVM menunjukkan akurasi yang sedikit lebih tinggi dibandingkan dengan MNB, yakni sebesar 82%. Pembagian data sebagai bagian latih dan uji dengan variasi rasio juga memengaruhi signifikan terhadap kinerja kedua metode. Hasil eksperimen menunjukkan bahwa parameter optimal untuk SVM adalah $C=100$ dan $\gamma=0.1$, sedangkan untuk MNB adalah $\alpha=2.0$ dan $\text{fit_prior}=\text{True}$. Meskipun SVM sedikit lebih unggul dengan akurasi 82% dibandingkan MNB yang memiliki akurasi 80%, penelitian selanjutnya disarankan untuk mengembangkan model klasifikasi dengan mempertimbangkan penggunaan teknik-teknik baru dalam machine learning, mengeksplorasi metode klasifikasi lainnya, dan mempertimbangkan penggunaan fitur tambahan atau teknik feature engineering untuk analisis yang lebih mendalam terhadap jenis sentimen yang muncul. Dengan demikian, penelitian ini memberikan kontribusi penting dalam memahami sentimen masyarakat terhadap pelayanan publik DKI Jakarta dan memberikan dasar untuk pengembangan lebih lanjut dalam bidang analisis sentimen menggunakan metode-metode machine learning.

DAFTAR PUSTAKA

- [1] T. D. Mahpudin, S. Nursanti, and M. Rifai, "Penggunaan Media Sosial Instagram dalam Pembentukan Corporate Branding," *Jurnal Ilmiah Wahana Pendidikan*, vol. 9, no. 9, pp. 292-301, 2023, doi <https://doi.org/10.5281/zenodo.7951770>.
- [2] P. Y. Saputra, D. H. Subhi, and F. Z. A. Winatama, "Implementasi Sentimen Analisis Komentar Channel Video Pelayanan Pemerintah di YouTube Menggunakan Algoritma Naive Bayes" *Jurnal Teknik Informatika, Politeknik Negeri Malang, Malang, Indonesia*, vol. 5, no. 3, pp. 209, Mei 2019.
- [3] R. Yunita dan M. Kamayani, "Perbandingan Algoritma SVM Dan Naive Bayes Pada Analisis Sentimen Penghapusan Kewajiban Skripsi," *Indonesian Journal of Computer Science (IJCS)*, vol. 12, no. 5, pp. 2879- 2890, 2023. doi 10.33022/ijcs.v12i5.3415.
- [4] A. Firdaus and W. I. Firdaus, "Text Mining Dan Pola Algoritma Dalam Penyelesaian Masalah Informasi: Sebuah Ulasan," *Jurnal JUPITER*, vol. 13, no. 1, pp. 202-218, Apr. 2021.
- [5] Tan, P. N., Steinbach, M., Karpatne, A., & Kumar, V. (2018). *Introduction to Data Mining* (2nd Edition). United States: Pearson Education.
- [6] S. Hilda Kusumahadi, H. Junaedi, and J. Santoso, "Klasifikasi Helpdesk Menggunakan Metode Support Vector Machine," *J. Inform. J. Pengemb. IT*, vol. 4, no. 1, pp. 54–60, 2019, doi: 10.30591/jpit.v4i1.1125.
- [7] I. Saputra and D. A. Kristiyanti, "Machine Learning untuk Pemula," *BI-OBSSES*, 2022. [Online]. Available: <https://elibrary.bsi.ac.id/readbook/220976/machine-learning-untuk-pemula>. [Accessed: 24-Oct-2024].
- [8] A. U. Khasanah dan A. Febriyanti, "Sentiment Analysis of JNE User Perception using Naive Bayes Classifier Algorithm," *Operations and Systems Integration Perspectives (OPSIP)*, vol. 15, no. 1, pp. 124-130, Juni 2022. doi 10.31315/opsi.v15i1.7179
- [9] A. Oktavia Praneswara and N. Cahyono, "Analisis Sentimen Ulasan Aplikasi TikTok Shop Seller Center di Google Playstore Menggunakan Algoritma Naive Bayes," *Indonesian Journal of Computer Science Attribution*, vol. 12, no. 6, p. 3925.
- [10] R. Yunita dan M. Kamayani, "Perbandingan Algoritma SVM Dan Naive Bayes Pada Analisis Sentimen Penghapusan Kewajiban Skripsi," *Indonesian Journal of Computer Science (IJCS)*, vol. 12, no. 5, pp. 2879-2890, 2023. doi 10.33022/ijcs.v12i5.3415.
- [11] G. Singh, B. Kumar, L. Gaur, and A. Tyagi, "Comparison between Multinomial and Bernoulli Naive Bayes for Text Classification," 2019 Int. Conf. Autom. Comput. Technol.

- Manag. ICACTM 2019, no. May 2020, pp. 593–596, 2019, doi: 10.1109/ICACTM.2019.8776800
- [12] D. I. Purnama dan O. P. Hendarsin, "Peramalan Jumlah Penumpang Berangkat Melalui Transportasi Udara di Sulawesi Tengah Menggunakan Support Vector Regression (SVR)," *Jambura Journal of Mathematics*, vol. 2, no. 2, pp. 49-59, Jul. 2020.
- [13] R. A. Fauzan and Mufti, "Analisis Sentimen Komentar YouTube tentang Program Kampus Merdeka Berbasis Web Menggunakan Algoritma Multinomial Naïve Bayes," in *Prosiding Seminar Nasional Mahasiswa Fakultas Teknologi Informasi (SENAFTI)*, vol. 2, no. 2, pp. 864–871, 2023. Available: <https://senafiti.budiluhur.ac.id/index.php/senafiti/article/view/929>
- [14] I. Habib Kusuma and N. Cahyono, "Analisis Sentimen Masyarakat Terhadap Penggunaan E-Commerce Menggunakan Algoritma K- Nearest Neighbor," vol. 8, no. 3, 2023.
- [15] A. Z. Praghakusma and N. Charibaldi, "Komparasi Fungsi Kernel Metode Support Vector Machine untuk Analisis Sentimen Instagram dan Twitter (Studi Kasus: Komisi Pemberantasan Korupsi)," *Jurnal Sarjana Teknik Informatika*, vol. 9, no. 2, pp. 33-42, Juni 2021.
- [16] S. Suprianto, A. Indriani, dan M. Rahmawati, "Perbandingan Metode Naïve Bayes Classifier Dan Holistic Lexicon Based Dalam Analisis Sentimen Angket Mahasiswa," *JSI : Jurnal Sistem Informasi (E-Journal)*, vol. 11, no. 2, hal. 1763, Oktober 2019.