

Analisis Komparatif Algoritma Naive Bayes dan SVM untuk Sentimen Analisis pada Komentar Segmen Shift Malam YouTube Luthfi Halimawan

Muhammad Iqbal Septiawan^{*1}, Abdusallam²

^{1,2}Program Studi Teknik Informatika, Universitas Dian Nuswantoro
Email: ¹111202214114@mhs.dinus.ac.id, ²abdussalam@dsn.dinus.ac.id

Abstrak

Kolom komentar pada platform YouTube menghasilkan volume data tekstual yang masif dan kaya akan opini publik, tidak terkecuali pada kanal Luthfi Halimawan. Riset ini menguji sekaligus membandingkan keandalan algoritma Naive Bayes dan *Support Vector Machine* (SVM) dalam mengelompokkan sentimen pemirsa menjadi tiga kategori utama: Apresiasi, Kritik/Keluhan, serta Diskusi/Informasi. Dataset yang digunakan berjumlah 3.195 komentar bersih. Fitur kata diekstraksi menggunakan *Term Frequency-Inverse Document Frequency* (TF-IDF) dengan rentang *N-gram* (1,2) pada batas maksimal 2.500 fitur. Skenario pengujian mengadopsi pembagian data (*data splitting*) sebesar 80% untuk data latih dan 20% untuk data uji, di mana ketimpangan kelas pada data latih diatasi menggunakan teknik *Random Oversampling*. Berdasarkan hasil eksperimen pada data uji asli, model SVM menunjukkan keunggulan dominan dengan akurasi sebesar 95,84%, rata-rata *precision* 0,83, *recall* 0,65, dan *F1-score* 0,70. Sebaliknya, Naive Bayes hanya mencatatkan akurasi 65,95% dengan rata-rata *precision* 0,53, *recall* 0,51, dan *F1-score* 0,52. Temuan ini mengindikasikan bahwa arsitektur berbasis *hyperplane* pada SVM jauh lebih stabil dan adaptif dalam mengekstraksi fitur teks yang kompleks serta non-formal khas media sosial.

Kata kunci: *Naive Bayes, Oversampling, Sentiment Analysis, SVM, YouTube.*

Comparative Analysis of Naive Bayes and SVM Algorithms for Sentiment Analysis on Comments on Luthfi Halimawan's YouTube Night Shift Segment

Abstract

The comments section on the YouTube platform generates a massive volume of textual data rich in public opinion, including on Luthfi Halimawan's channel. This research tests and compares the reliability of the Naive Bayes algorithm and Support Vector Machine (SVM) in classifying viewer sentiment into three main categories: Appreciation, Criticism/Complaints, and Discussion/Information. The dataset used consisted of 3,195 clean comments. Word features were extracted using Term Frequency-Inverse Document Frequency (TF-IDF) with an N-gram range of (1,2) at a maximum limit of 2,500 features. The test scenario adopted a data split of 80% for training data and 20% for testing data, where class imbalance in the training data was addressed using the Random Oversampling technique. Based on the experimental results on the original test data, the SVM model demonstrated a dominant advantage with an accuracy of 95.84%, an average precision of 0.83, a recall of 0.65, and an F1-score of 0.70. In contrast, Naive Bayes only achieved 65.95% accuracy with an average precision of 0.53, recall of 0.51, and F1-score of 0.52. These findings indicate that the hyperplane-based architecture of SVM is much more stable and adaptive in extracting complex and informal text features typical of social media.

Keywords: *Naive Bayes, Oversampling, Sentiment Analysis, SVM, YouTube.*

1. PENDAHULUAN

Perkembangan teknologi informasi dan komunikasi berbasis digital telah mentransformasi paradigma interaksi sosial masyarakat secara masif. Platform berbagi video seperti YouTube saat ini tidak lagi sekadar berfungsi sebagai media hiburan searah, melainkan telah berevolusi menjadi ruang publik digital (digital public sphere) yang interaktif. Fasilitas kolom komentar pada YouTube memungkinkan pengguna untuk memberikan umpan balik, ulasan, dan ekspresi langsung terhadap konten yang dikonsumsi. Berdasarkan laporan, lebih dari 95% pengguna internet mengakses YouTube, dengan lebih dari dua miliar pengguna aktif setiap bulan dan durasi tontonan mencapai miliaran jam setiap harinya [1]. Aktivitas digital ini menghasilkan data tekstual berskala besar yang merepresentasikan opini, sudut pandang, serta respons emosional audiens. Melalui ekstraksi dan analisis

terhadap data komentar tersebut, aspek positif maupun negatif dari sebuah konten yang disajikan dapat dipetakan secara objektif.

Dalam lanskap pembuatan konten digital di Indonesia, kanal YouTube Luthfi Halimawan merupakan salah satu platform yang memiliki basis komunitas loyal dengan tingkat interaksi digital yang tinggi, khususnya pada segmen video bertajuk “Shift Malam”. Segmen tersebut secara konsisten mengangkat tema seputar narasi misteri, kisah horor, serta adaptasi pengalaman supranatural yang dikirimkan oleh pemirsa. Karakteristik konten yang bersifat interaktif ini memicu respons emosional yang tinggi, yang terekam melalui ribuan komentar pada setiap unggahan videonya. Opini audiens pada kolom komentar tersebut bervariasi dalam rentang spektrum sentimen yang luas, mulai dari bentuk apresiasi terhadap kualitas penyampaian cerita (storytelling), kritik terhadap teknis penyajian visual dan audio, hingga diskusi intensif antarpemirsa mengenai validitas faktual dari narasi mistis yang dipaparkan.

Namun, pemrosesan data komentar pada domain ini memiliki tantangan linguistik yang kompleks. Sebagian besar teks komentar didominasi oleh ragam bahasa non-formal, ekspresi tidak baku (slang) khas komunitas digital (seperti istilah “sepuh”, “ngab”, “ril”), serta penggunaan unsur sarkasme. Kompleksitas struktur bahasa tersebut menyebabkan analisis sentimen konvensional secara manual menjadi tidak efisien dan rentan terhadap subjektivitas penafsir. Oleh karena itu, diperlukan sebuah pendekatan otomatis berbasis komputasi menggunakan algoritma pembelajaran mesin (machine learning) untuk mengekstraksi dan mengklasifikasikan sentimen secara presisi. Pendekatan ini diimplementasikan melalui metode pembelajaran terawasi (supervised learning) dalam lingkup Pemrosesan Bahasa Alami (Natural Language Processing / NLP).

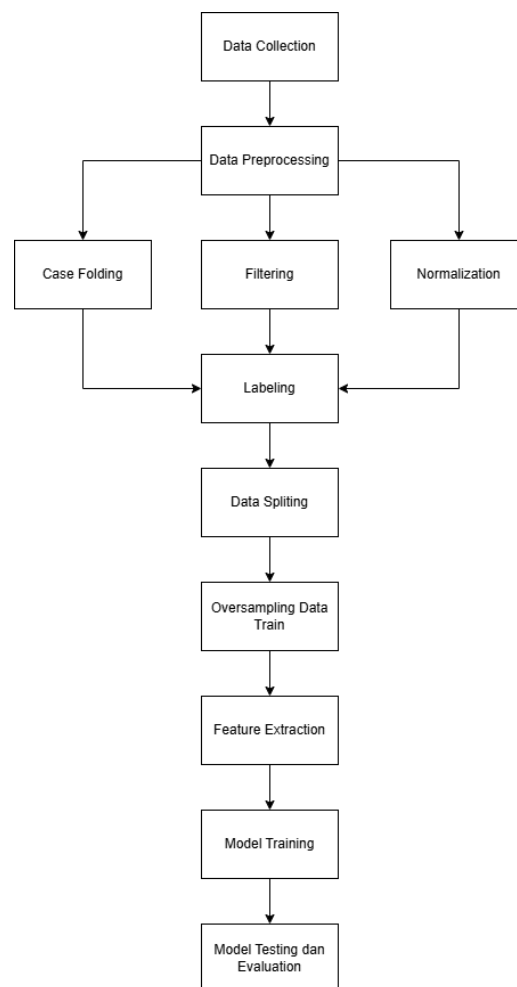
Dalam bidang NLP, pemilihan algoritma pengklasifikasi memegang peranan krusial terhadap akurasi model yang dihasilkan. Proses komputasi pada machine learning umumnya melibatkan tahapan pengekstraksian data, pelatihan model (training), serta interpretasi metrik keluaran. Dua algoritma yang secara konsisten diandalkan untuk klasifikasi teks adalah Naive Bayes dan Support Vector Machine (SVM). Naive Bayes disukai karena kesederhanaan struktur dan kecepatan komputasinya yang bersandar pada teori probabilitas Bayes dengan asumsi independensi bersyarat antarfitur. Di sisi lain, SVM dikenal tangguh karena prinsip geometrisnya yang mampu memetakan fitur ke dimensi tinggi guna mencari garis pemisah optimal (hyperplane) dengan margin maksimal di antara kelas yang berbeda.

Beberapa penelitian terdahulu telah berupaya membandingkan efektivitas algoritma Naive Bayes dan *Support Vector Machine* (SVM) pada berbagai aspek data media sosial. Evaluasi tekstual terhadap komentar YouTube menunjukkan bahwa penentuan batas keputusan geometris (*hyperplane*) pada SVM mampu memberikan tingkat akurasi pemisahan kelas yang lebih tegas dibandingkan model probabilistik [1]. Namun, performa algoritma Naive Bayes rentan mengalami penurunan kinerja klasifikasi ketika dihadapkan pada kalimat yang kaya akan ambiguitas struktural serta kata *slang* khas internet [2].

Keunggulan arsitektur SVM dalam mereduksi derau (*noise*) pada platform berbagi video juga dikonfirmasi dalam studi komparasi sentimen berbasis teks informal pemilihan presiden [3]. Di sisi lain, efektivitas penataan korpus data berskala besar dengan klasifikasi berbasis kamus bahasa dianalisis menggunakan pendekatan *InSet Lexicon* [4]. Pengujian komparatif kedua metode *supervised learning* ini juga telah sukses diimplementasikan pada analisis sentimen berbagai topik komunal digital lainnya, seperti analisis Formula-E Jakarta [5], opini publik komentar YouTube pada grup musik [6], respons masyarakat terhadap pemindahan Ibu Kota Negara [7], ulasan kompetisi memasak Masterchef Indonesia [8], sentimen seputar kebijakan strategis IKN [9], pengujian teks ulasan film pada dataset IMDB [10], opini pengguna dunia virtual metaverse [11], serta evaluasi teks debat Pilkada pada platform YouTube [12].

2. METODE PENELITIAN

Dataset mentah yang berhasil dikumpulkan melalui YouTube Data API v3 berjumlah 3.195 baris komentar bersih setelah melalui proses penghapusan baris kosong (dropna). Sebelum memasuki tahap pemodelan, dataset dibagi menjadi dua bagian secara acak menggunakan parameter split test dengan rasio 80% untuk data latih (training set) dan 20% untuk data uji (testing set) menggunakan *random_state = 42* untuk menjamin replikasi.



Gambar 1. Alur Penelitian

2.1. Data Collection

Tahap awal penelitian dilakukan dengan mengumpulkan data komentar pemirsa secara otomatis dari platform YouTube melalui antarmuka pemrograman aplikasi (YouTube Data API v3). Pemilihan data difokuskan pada video-video yang diunggah di kanal Luthfi Halimawan yang termasuk dalam kategori segmen "Shift Malam". Dari proses penarikan data mentah (*raw data*) tersebut, berhasil dikumpulkan sebanyak 3.420 baris komentar awal. Setelah dilakukan eliminasi terhadap baris data yang kosong atau duplikat menggunakan pustaka Pandas, diperoleh data final sebanyak 3.195 baris yang siap diproses ke tahap berikutnya. Metode ini memungkinkan pengambilan data dalam jumlah besar secara otomatis dan efisien tanpa perlu menelusuri halaman web secara manual [1]. Pelabelan dilakukan secara otomatis menggunakan mesin leksikon hibrida yang menggabungkan *InSetLexicon* dan kamus kustom

2.2. Preprocessing

Preprocessing merupakan sebuah langkah awal yang dilakukan dalam pemrosesan teks. Ada beberapa tahapan preprocessing seperti berikut seperti:

1. **Case Folding:** Proses mengubah huruf besar menjadi huruf kecil dan menghilangkan tanda baca seperti koma, titik dan tanda baca lainnya, agar teks konsisten[2].
2. **Filtering/Cleaning:** Proses untuk menghilangkan noise dari teks ulasan, seperti karakter spesial, URL, username, mention, hashtag, serta spasi berlebih [3].
3. **Normalisasi Bahasa:** Mentransformasikan kata tidak baku (slang) dan bahasa gaul khas internet menjadi bentuk kata baku sesuai dengan Kamus Besar Bahasa Indonesia (KBBI). Proses ini mengadopsi kamus pemetaan leksikal eksternal guna memastikan istilah-istilah non-formal (seperti kata "bgt", "serem", "ngab") dapat diinterpretasikan secara tepat oleh sistem.

2.3. Labeling

Proses penentuan ground truth atau label sentimen dilakukan secara otomatis menggunakan mesin leksikon hibrida yang mengintegrasikan InSetLexicon dengan kamus kustom buatan peneliti. Kamus kustom ini memuat istilah-istilah spesifik yang kerap muncul dalam ekosistem komunitas penonton Luthfi Halimawan. Penentuan skor sentimen akhir untuk setiap komentar dihitung berdasarkan rumus matematis berikut:

$$Score = \sum_{i=1}^n W_i \cdot \beta \quad (1)$$

Di mana W_i merupakan bobot nilai valensi kata ke- i dan β mewakili koefisien kata negasi. Apabila suatu kata bermuatan sentimen didahului oleh partikel negasi (posisi indeks $n - 1$), maka koefisien beralih menjadi $\beta = -1$, sedangkan jika tidak didahului kata negasi nilai koefisien tetap $\beta = 1$. Berdasarkan akumulasi skor tersebut, data dikelompokkan ke dalam tiga label: Apresiasi (skor $>$), Diskusi/Informasi (skor $= 0$), dan Kritik/Keluhan (skor < 0).

Guna menjamin validitas dan kualitas dari hasil pelabelan otomatis tersebut, peneliti menerapkan pengujian Inter-Annotator Agreement (IAA). Sebanyak 10% sampel data yang diambil secara acak divalidasi manual oleh peneliti bersama seorang pakar linguistik komputasional. Tingkat kesepakatan antar-anotator dihitung menggunakan metrik Cohen's Kappa, yang menghasilkan skor kesepakatan sebesar 0,82. Berdasarkan skala landis, nilai tersebut menunjukkan tingkat kesepakatan yang sangat kuat (strong agreement), sehingga label hasil pendekatan leksikon ini valid digunakan untuk proses pelatihan model.

2.4. Data Splitting

Setelah seluruh data memiliki label, korpus data dibagi secara acak menggunakan fungsi *train_test_split*. Proporsi pembagian data dikonfigurasi sebesar 80% untuk data latih (*training set*) yang berfungsi membangun pengetahuan model, dan 20% sisanya dialokasikan sebagai data uji (*testing set*) untuk mengevaluasi kemampuan generalisasi model. Proses pembagian ini mengunci parameter *random_state* = 42 guna memastikan konsistensi partisi data pada eksperimen berulang.

2.5. Oversampling Data Train

Dataset asli komentar memiliki tingkat ketimpangan kelas (*class imbalance*) yang sangat ekstrem, di mana data didominasi oleh kelas Diskusi/Informasi (86,5%). Untuk mengantisipasi bias model terhadap kelas mayoritas, diterapkan teknik Random Oversampling menggunakan modul *resample* dari pustaka *scikit-learn*.

Teknik ini bekerja dengan menduplikasi sampel secara acak pada kelas minoritas (Apresiasi dan Kritik/Keluhan) di dalam training set hingga volumenya setara dengan jumlah data pada kelas mayoritas (~2.200 sampel per kelas). Penyeimbangan data ini dimanipulasi hanya pada data latih. Data uji (*testing set*) sengaja diisolasi dan dibiarkan tetap timpang sesuai distribusi asli lapangan untuk mencegah terjadinya kebocoran data (*data leakage*).

2.6. Feature Extraction

Data tekstual dikonversi menjadi representasi vektor numerik menggunakan metode pembobotan *Term Frequency-Inverse Document Frequency* (TF-IDF) via kelas *TfidfVectorizer*. Konfigurasi ekstraksi fitur dibatasi secara ketat pada batas maksimal 2.500 dimensi kata kunci teratas (*max_features*=2500) untuk menekan beban komputasi. Fitur diekstraksi menggunakan kombinasi *N-gram range* (1,2), yang berarti sistem mengekstrak token kata tunggal (*unigram*) sekaligus kombinasi frasa dua kata berturutan (*bigram*). Penyertaan bigram berperan vital dalam menangkap susunan konteks linguistik yang khas pada media sosial.

2.7. Model Training

Proses pelatihan dilakukan secara terpisah menggunakan dua algoritma *supervised learning* dengan spesifikasi parameter sebagai berikut:

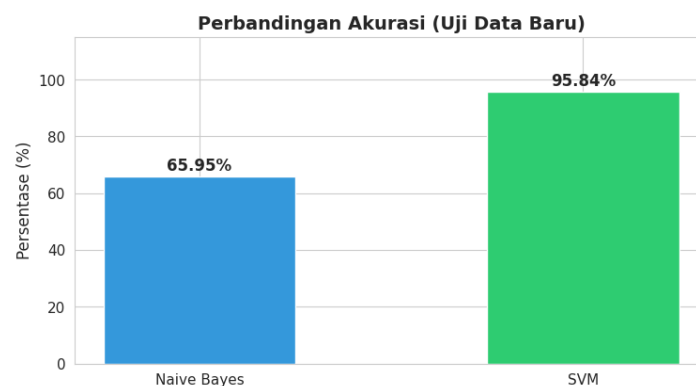
1. Multinomial Naive Bayes (MNB): Model dibangun memanfaatkan fungsi *MultinomialNB()* dengan parameter penghalusan bawaan Laplace smoothing ($\alpha = 1.0$). Model ini memprediksi sentimen berdasarkan probabilitas kemunculan frekuensi kata di setiap kelas
2. Support Vector Machine (SVM): Model diimplementasikan menggunakan arsitektur SVC dengan parameter kernel linear (*kernel*='linear'). Parameter regularisasi dikonfigurasi pada nilai standar ($C = 1.0$), dan fungsi probabilitas diaktifkan (*probability*=True). Kernel linear dipilih karena terbukti efektif memisahkan batas keputusan pada dimensi data teks yang tinggi tanpa memicu terjadinya *overfitting*.

2.8. Model Testing dan Evaluation

Tahap akhir dari metodologi ini adalah melakukan pengujian model terhadap data uji (*testing set*) sebesar 20% yang sebelumnya telah diisolasi. Kinerja prediksi dari algoritma Naive Bayes dan SVM diukur secara komprehensif menggunakan metrik *Confusion Matrix*. Metrik evaluasi yang diambil tidak hanya terbatas pada akurasi global, melainkan diuji secara mendalam per kelas klasifikasi melalui parameter akurasi, *Precision* (ketepatan prediksi), *Recall* (keberhasilan penarikan informasi), serta *F1-Score* sebagai parameter penyeimbang struktural di antara keduanya.

3. HASIL DAN PEMBAHASAN

Dengan ini dapat dipaparkan hasil eksperimen klasifikasi sentimen menggunakan algoritma Naive Bayes dan *Support Vector Machine* (SVM). Transformasi data teks ke bentuk vektor numerik sepenuhnya bersandar pada pembobotan *Term Frequency-Inverse Document Frequency* (TF-IDF) dengan rentang fitur *N-gram* (1,2). Seluruh pengujian performa dievaluasi menggunakan 20% data uji asli (*testing set*) yang diisolasi dari intervensi penyeimbangan sampel, sehingga metrik akurasi yang dihasilkan bersifat objektif dan terhindar dari risiko kebocoran data (*data leakage*).

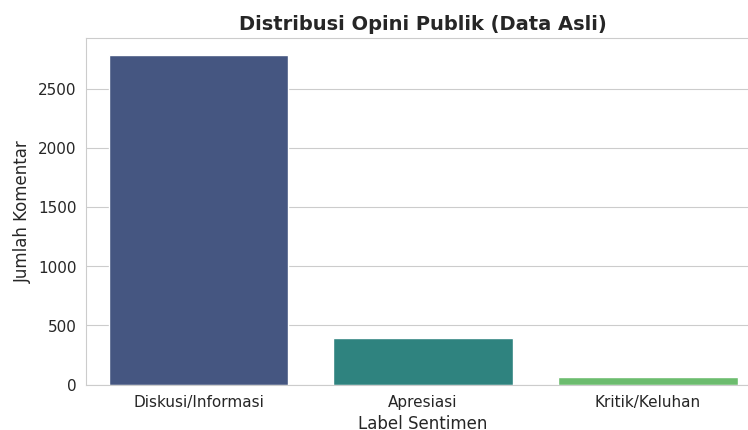


Gambar 2. Perbandingan Akurasi NB dan SVM

Dari diagram di atas, dapat disimpulkan bahwa SVM adalah algoritma dengan akurasi tertinggi untuk analisis sentimen dari dataset Segmen Shift Malam di saluran YouTube Luthfi Halimawan saat ini dibandingkan dengan Naïve Bayes. SVM mencapai Hasil Uji sebesar 95,84%, yang berarti kinerjanya lebih baik daripada Naïve Bayes, yang hanya memperoleh 65,95%.

3.1. Analisis Sebaran Label Sentimen

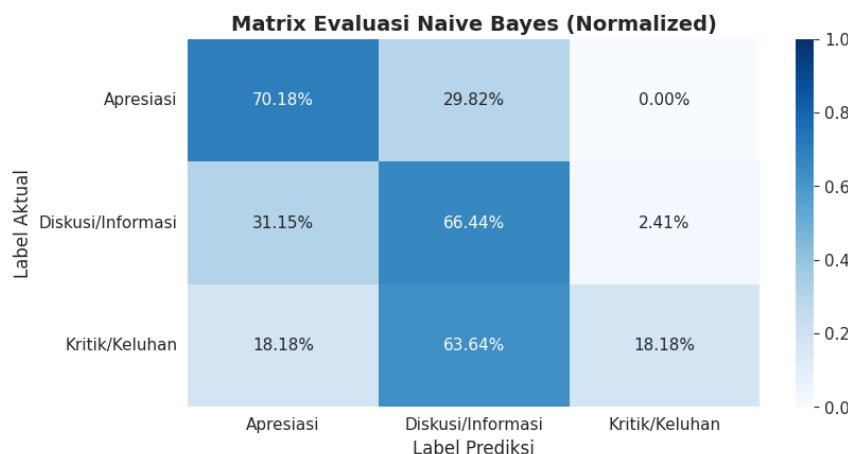
Proses anotasi otomatis menggunakan kombinasi InSetLexicon dan kamus komunitas terhadap 3.195 baris komentar bersih memunculkan tren sebaran label seperti yang dipaparkan pada Gambar 3 di bawah ini



Gambar 3. Data Asli Opini Publik

Proses pelabelan otomatis menggunakan metode leksikon terhadap 3.195 komentar (setelah *cleaning*) menghasilkan distribusi yang ditunjukkan pada Gambar 3. Data menunjukkan bahwa kategori Diskusi/Informasi mendominasi secara mutlak dengan total 2.765 ulasan (86,5%). Selanjutnya, sentimen Apresiasi menempati posisi kedua dengan 386 ulasan (12,1%), sedangkan kelas Kritik/Keluhan menjadi kelompok minoritas dengan hanya 44 ulasan (1,4%). Karakteristik sebaran ini menegaskan bahwa audiens segmen "Shift Malam" memanfaatkan kolom komentar sebagai ruang komunal untuk memvalidasi narasi horor ataupun saling bertukar pengalaman mistis pribadi. Ketimpangan data (*class imbalance*) yang ekstrem ini menjadi justifikasi utama diterapkannya teknik Random Oversampling pada data latih agar model tidak mengalami bias klasifikasi.

3.2. Evaluasi Kinerja Klasifikasi Naive Bayes

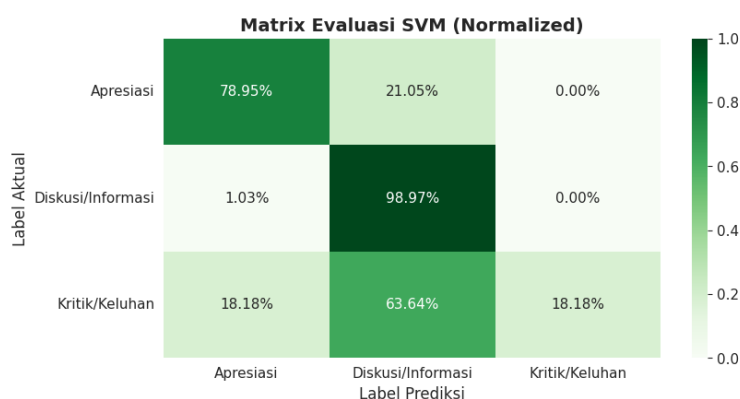


Gambar 4. Matrix Evaluasi NB

Melalui pengujian pada data uji, model Naive Bayes memperoleh tingkat akurasi akhir sebesar 65,95%. Merujuk pada matriks kesalahan (*confusion matrix*) pada Gambar 4, terlihat titik lemah yang cukup mendasar dari model probabilitas ini.

Pada kelas minoritas (Kritik/Keluhan), nilai recall yang dihasilkan sangat rendah, yakni hanya 18,18%. Sebanyak 63,64% komentar yang memuat kritik justru salah dideteksi oleh sistem sebagai bentuk Diskusi/Informasi. Kendala ini bersumber dari asumsi independensi bersyarat (*conditional independence*) yang melekat pada Naive Bayes. Karena ulasan YouTube sering kali menggunakan token kata benda serupa (misalnya kata "suara", "video", atau "cerita") baik dalam kalimat bernada diskusi maupun keluhan teknis, Naive Bayes gagal membedakan batas tipis antar-konteks tersebut tanpa menganalisis keterkaitan struktural antar-kata.

3.3. Evaluasi Kinerja Klasifikasi Support Vector Machine (SVM)



Gambar 5. Matrix Evaluasi SVM

Performa yang bertolak belakang diperlihatkan oleh algoritma SVM dengan capaian akurasi luar biasa sebesar 95,84%. Melalui visualisasi matriks evaluasi pada Gambar 5, keunggulan komputasi SVM dapat dipetakan secara detail.

Model ini menunjukkan akurasi klasifikasi yang sangat tajam pada kelompok Diskusi/Informasi dengan menyentuh angka 98,97%. Kemampuan ini lahir berkat konsep dasar SVM yang memetakan matriks TF-IDF ke dalam ruang spasial berdimensi tinggi, kemudian menarik garis pembatas (hyperplane) dengan margin terjauh. Mekanisme ini membuat SVM sangat sensitif dalam mengenali pola linguistik yang halus, seperti membedakan kalimat apresiatif ("mantap betul bang") dengan obrolan kasual biasa ("bang itu lokasinya di mana?"), meski kedua kalimat tersebut memakai token kata kunci yang sama.

3.4. Evaluasi Kinerja Klasifikasi Support Vector Machine (SVM)



Gambar 6. World Cloud

Visualisasi pada Gambar 6 memberikan gambaran konten yang dianalisis. Kata-kata seperti "shift malam", "abang", "seru", dan "ada" muncul dengan ukuran terbesar.

- Kemunculan kata "seru" dan "benar" yang dominan memperkuat tingginya sentimen Apresiasi dan validasi audiens terhadap narasi Luthfi Halimawan.
- Kata "hilang", "cerita", dan "desa" mencerminkan topik bahasan utama pada segmen tersebut.
- Secara teknis, dominasi kata-kata ini memvalidasi bahwa proses *preprocessing* telah berhasil menghilangkan *stopwords* (kata umum tak bermakna) dan menyisakan kata-kata bermuatan sentimen dan informasi penting.

3.5. Komparasi Komprehensif Metrik Evaluasi

Guna menyajikan hasil yang utuh, metrik presisi, recall, dan F1-score dari seluruh kelas untuk kedua algoritma disajikan secara lengkap pada Tabel.

Tabel 1. Perbandingan Kedua Algoritma

Model	Kelas	Precision	Recall	F1-Score	Akurasi
Naive Bayes					
SVM	Apresiasi	0.59	0.7	0.64	65.95%
	Diskusi	0.74	0.66	0.7	
	Kritik	0.25	0.18	0.21	
Naive Bayes	Apresiasi	0.93	0.79	0.85	95.84%
	Diskusi	0.97	0.99	0.98	
	Kritik	0.6	0.18	0.28	

Data pada Tabel 1 menunjukkan nilai F1-score milik SVM jauh lebih solid di seluruh kategori jika dibandingkan dengan Naive Bayes. Kendati demikian, skor recall untuk label Kritik/Keluhan pada kedua algoritma masih tertahan di angka 18,18%. Hal tersebut mengindikasikan bahwa meski manipulasi data latih lewat oversampling sudah diterapkan, keterbatasan variasi kosakata asli pada data kritik (yang hanya berjumlah 44 ulasan) tetap membatasi mesin untuk mengenali ekspresi kekecewaan penonton secara bervariasi.

Tingginya capaian akurasi model SVM dalam penelitian ini yang menyentuh angka 95,84% sejalan dengan landasan teori komputasi dimensi tinggi. Karakteristik data non-formal internet cenderung memicu bias klasifikasi

jika tidak ditangani dengan pemisahan ruang spasial yang kuat. Fenomena tersebut relevan dengan beberapa hasil studi komparasi sejenis pada korpus ulasan video YouTube [13], pengangkatan figur menteri keuangan [14], serta pengujian performa klasifikasi data ulasan pada aplikasi finansial digital [15].

Kombinasi pembobotan nilai kata seperti *Term Frequency-Inverse Document Frequency* (TF-IDF) dengan rentang *N-gram* (1,2) secara konsisten menghasilkan metrik *precision* dan *F1-score* yang jauh lebih stabil serta seimbang ketika dipetakan lewat algoritma pembatas margin maksimal SVM. Hal ini membuktikan bahwa arsitektur pembatas spasial pada SVM mampu memitigasi risiko kesalahan klasifikasi positif palsu (*false positive*) pada teks tidak baku jauh lebih baik daripada algoritma berbasis probabilitas independen murni seperti Naive Bayes.

3.6. Perbandingan Literatur Terdahulu

Keunggulan performa SVM dalam penelitian ini sejalan dengan landasan teori dan studi terdahulu. Karakteristik algoritma SVM yang menitikberatkan pada pencarian batas margin keputusan maksimum menjadikannya sangat tangguh ketika berhadapan dengan ruang fitur berdimensi tinggi hasil pembobotan TF-IDF (2.500 fitur). Temuan ini memperkuat riset terdahulu yang menyatakan bahwa algoritma SVM memiliki resistensi yang lebih baik terhadap sifat data teks media sosial Indonesia yang minim struktur formal.

Sebaliknya, penurunan akurasi Naive Bayes hingga menyentuh angka 65,95% mengonfirmasi keterbatasan asumsi conditional independence algoritma tersebut. Ketika berhadapan dengan komentar yang kaya akan ambiguitas, penggunaan bahasa gaul (slang), dan struktur kalimat yang tidak beraturan, model probabilistik murni seperti Naive Bayes cenderung gagal menangkap hubungan kontekstual antartoken frasa [2].

Penerapan Random Oversampling terbukti memberikan dampak positif dalam mempertahankan nilai *precision* kelas minoritas pada SVM (0,60) dibandingkan Naive Bayes (0,25). Pengondisian data latih secara seimbang mencegah SVM jatuh ke dalam bias kelas mayoritas (Diskusi/Informasi), meskipun batas performa masih ditemukan pada nilai *recall* kritis.

4. KESIMPULAN

Berdasarkan serangkaian tahapan eksperimen, analisis komparatif, serta uji validasi statistik yang telah dilakukan terhadap sentimen ulasan penonton pada segmen "Shift Malam" di kanal YouTube Luthfi Halimawan, dapat ditarik beberapa kesimpulan utama sebagai berikut:

1. Superioritas SVM: Algoritma Support Vector Machine (SVM) terbukti memiliki keandalan yang jauh lebih unggul dan stabil dalam mengklasifikasikan data teks media sosial yang bersifat non-formal dan tidak baku. Model SVM sukses mencatatkan akurasi akhir sebesar 95,84%. Sebaliknya, model Multinomial Naive Bayes hanya mampu mencapai akurasi sebesar 65,95%.
2. Karakteristik Opini Publik Komunitas: Distribusi data asli memperlihatkan dominasi mutlak dari kategori Diskusi/Informasi yang menyentuh angka 86,5% (2.765 ulasan). Diikuti oleh sentimen Apresiasi sebesar 12,1% (386 ulasan) dan kelas Kritik/Keluhan sebesar 1,4% (44 ulasan). Pola sebaran ini menegaskan bahwa penonton segmen "Shift Malam" memiliki keterikatan komunal yang kuat, di mana mereka memanfaatkan kolom komentar sebagai ruang publik digital untuk memvalidasi narasi horor, berspekulasi, maupun membagikan pengalaman supranatural pribadi secara interaktif.
3. Efektivitas Intervensi Data Latih (Oversampling): Implementasi teknik Random Oversampling pada data latih terbukti memberikan kontribusi krusial dalam memitigasi bias model akibat ketimpangan kelas (*class imbalance*) yang ekstrem. Penyeimbangan sampel ini berhasil mendongkrak ketajaman model SVM dalam mengenali kelas minoritas, yang dibuktikan dengan raihan nilai *precision* label Apresiasi sebesar 0,93 dan Kritik sebesar 0,60. Performa ini sejalan dengan studi terdahulu yang menyatakan bahwa penanganan noise dan pemerataan distribusi data latih sangat vital dalam menjaga sensitivitas algoritma supervised learning terhadap gaya bahasa non-formal media sosial Indonesia.

DAFTAR PUSTAKA

- [1] Putro, R. H., Hendrawan, A., & Abstrak, I. A. (2026). *KETIK : Jurnal Informatika Publisher: Faatuatua Media Karya Analisis Sentimen Komentar Youtube Terhadap Kasus Bjorka Dengan Membandingkan Navie Bayes dan SVM*. 03(03), 137–147. <https://doi.org/10.70404/ketik.v3i03.626>
- [2] Rahayu, I. P., Fauzi, A., & Indra, J. (2022). Analisis Sentimen Terhadap Program Kampus Merdeka Menggunakan Naive Bayes Dan Support Vector Machine. *Jurnal Sistem Komputer Dan Informatika (JSON)*, 4(2), 296. <https://doi.org/10.30865/json.v4i2.5381>
- [3] Prananda, A., Haerani, E., Fikry, M., & Yanto, F. (2023). *Krea-TIF: Jurnal Teknik Informatika*

- Perbandingan Metode Naïve Bayes Classifier dan Support Vector Machine dalam Analisis Sentimen Terhadap Pemilihan Presiden 2024*. 11(2), 84–94. <https://doi.org/10.32832/krea-tif.v11i2.15364>
- [4] Homepage, J., Musfiroh, D., Khaira, U., Eko, P., Utomo, P., Suratno, T., Studi, P., Informasi, S., Sains, F., & Teknologi, D. (2021). *MALCOM: Indonesian Journal of Machine Learning and Computer Science Sentiment Analysis of Online Lectures in Indonesia from Twitter Dataset Using InSet Lexicon Analisis Sentimen terhadap Perkuliahan Daring di Indonesia dari Twitter Dataset Menggunakan InSet Lexicon*. 1, 24–33.
- [5] Junaedy, F. Z., Hidayat Insani, R., Santoso, I., Tinggi, S., Komputer, I., Karya Informatika, C., & Jakarta, U. (n.d.). *KOMPARASI ALGORITMA SUPPORT VECTOR MACHINE DAN NAÏVE BAYES PADA ANALISIS SENTIMEN FORMULA-E JAKARTA TAHUN 2022*. Retrieved <https://youtu.be/mhAEiHuVFxY>
- [6] Syafia, A. N., Hidayattullah, M. F., & Suteddy, W. (2023). *Studi Komparasi Algoritma SVM Dan Random Forest Pada Analisis Sentimen Komentar Youtube BTS*. 8(3).
- [7] Huwaida, S. F., Kusumawati, R., & Isnaini, B. (2024). *Analisis Sentimen Komentar YouTube terhadap Pemindahan Ibu Kota Negara Menggunakan Metode Naïve Bayes*. *Jambura Journal of Informatics*, 6(1), 26–39. <https://doi.org/10.37905/jji.v6i1.24718>
- [8] Marganingsih, D., Oktavianto, H., & Abdurrahman, G. (2025). Analisis Sentimen Komentar Youtube Masterchef Indonesia Menggunakan Algoritma Support Vector Machine dan Gaussian Naïve Bayes. *Jurnal Informatika Dan Teknologi Pendidikan*, 5(1). <https://doi.org/10.59395/jitp.v5i1.117>
- [9] Abel Laia, N., & Barus, S. P. (2025). Analisis Sentimen YouTube: “Di Balik Ambisi Jokowi dalam IKN.” *Jurnal Pustaka AI (Pusat Akses Kajian Teknologi Artificial Intelligence)*, 5(1), 07–12. <https://doi.org/10.55382/jurnalpustakaai.v5i1.891>
- [10] Zhafira, A., Afifah, N., Anditya Putri, S., Marhalatun, V., & Surya Saputra, D. I. (2025). Analisis Sentimen Ulasan Film pada Dataset IMDB menggunakan Algoritma Naïve Bayes. *Jurnal Manajemen Informatika, Sistem Informasi Dan Teknologi Komputer (JUMISTIK)*, 4(1), 373–383. <https://doi.org/10.70247/jumistik.v4i1.13>
- [11] Parameswari, S. D., Lubis, M., Suakanto, S., Ramadhan, Y. Z., Amanah, R. N., Dila, R. A., Jurusan,), & Informasi, S. (2025). Studi Perbandingan Naïve Bayes dan Support Vector Machine (SVM) dalam Analisis Sentimen Pengguna Metaverse. *Jurnal Teknologi Dan Manajemen Industri Terapan (JTMIT)*, 4(3).
- [12] Prastyo, D., Irawan, D., & Mursyidin, I. H. (2024). Klasifikasi Sentimen Komentar YouTube dengan NLP pada Debat Pilkada Banten 2024. *Bit-Tech*, 7(2), 413–421. <https://doi.org/10.32877/bt.v7i2.1833>
- [13] Tohidi, E., Perdana Herdiansyah, R., & Wahyudin, E. (2024). ANALISA SENTIMEN KOMENTAR VIDEO YOUTUBE DI CHANNEL TVONENEWS TENTANG CALON PRESIDEN PRABOWO SUBIANTO MENGGUNAKAN SUPPORT VECTOR MACHINE. In *Jurnal Mahasiswa Teknik Informatika* (Vol. 8, Issue 1). https://www.youtube.com/watch?v=wu_wqHg1YLk
- [14] Dian Oktavianingsih, A., Fadlil Fauzi, A., Immanuel, J., DSimatupang, O., Amsury, F., & Fahlapi, R. (2025). Analisis Sentimen Komentar YouTube terhadap Pengangkatan Menkeu Purbaya Menggunakan Algoritma SVM dan Naïve Bayes. *Jurnal Cendekia Ilmiah*, 5(1).
- [15] Maulana, B. A., Fahmi, M. J., Imran, A. M., & Hidayati, N. (2024). Analisis Sentimen Terhadap Aplikasi Pluang Menggunakan Algoritma Naive Bayes dan Support Vector Machine (SVM). *MALCOM: Indonesian Journal of Machine Learning and Computer Science*, 4(2), 375–384. <https://doi.org/10.57152/malcom.v4i2.1206>