

Analisis Sentimen Masyarakat Terhadap Penanganan Banjir Bandang di Pulau Sumatra Berdasarkan Komentar Youtube Menggunakan Metode *Support Vector Machine (SVM)*

Andi Saputra¹, Irna Nurdiyani², Ulfah Siti Nurhidayah³, Siti Maesaroh⁴
Teknik Informatika, Fakultas Teknik, Universitas Mayasari Bakti, Tasikmalaya Jawa Barat, Indonesia^{1, 2, 3, 4}

*Email Korespodensi: andiisaputra06@gmail.com

Sejarah Artikel:

Diterima 09-01-2026
Disetujui 19-01-2026
Diterbitkan 21-01-2026

ABSTRACT

Flash floods that struck Sumatra Island at the end of 2025 caused significant impacts on local communities and triggered diverse public responses on social media, particularly YouTube. This study aims to analyze public sentiment toward the handling of flash flood disasters in Sumatra Island using the Support Vector Machine (SVM) algorithm. The research data consist of 3,458 Indonesian-language comments extracted from 25 YouTube videos related to the Sumatra flash floods during the period January–March 2024. The research stages include data collection through web scraping, text preprocessing consisting of case folding, tokenization, stopword removal, and stemming, followed by feature weighting using the Term Frequency–Inverse Document Frequency (TF-IDF) method. The weighted data are then classified using an SVM algorithm with a linear kernel into three sentiment categories: positive, negative, and neutral. Model performance evaluation is conducted using a confusion matrix, producing accuracy, precision, recall, and F1-score metrics. The results show that the SVM algorithm is able to classify sentiment with an accuracy of 85.42%. The sentiment distribution indicates a dominance of negative sentiment (2,200 comments), followed by neutral sentiment (1,700 comments), and positive sentiment (600 comments), reflecting public dissatisfaction with the speed and effectiveness of disaster response. This study is expected to provide input for the government to improve the quality of disaster management systems and to serve as a reference for future research in social media–based sentiment analysis.

Keywords: Sentiment Analysis; Flash Floods in Sumatra; YouTube; Support Vector Machine; TF-IDF

ABSTRAK

Banjir bandang yang melanda Pulau Sumatra pada akhir tahun 2025 menimbulkan dampak signifikan terhadap masyarakat dan memicu beragam respons publik di media sosial, khususnya YouTube. Penelitian ini bertujuan menganalisis sentimen masyarakat terhadap penanganan banjir bandang di Pulau Sumatra menggunakan algoritma *Support Vector Machine (SVM)*. Data penelitian terdiri dari 3.458 komentar berbahasa Indonesia yang diekstraksi dari 25 video YouTube terkait banjir bandang Sumatra periode Januari–Maret 2024. Tahapan penelitian mencakup pengumpulan data melalui *web scraping*, pra-pemrosesan teks meliputi *case folding*, tokenisasi, *stopword removal*, dan *stemming*, dilanjutkan dengan pemberian bobot fitur menggunakan metode *Term Frequency-Inverse Document Frequency (TF-IDF)*. Data yang telah diberi bobot selanjutnya diklasifikasikan menggunakan algoritma SVM dengan *kernel linear* ke dalam tiga kategori sentimen: positif, negatif, dan netral. Evaluasi kinerja model dilakukan menggunakan *confusion matrix* yang menghasilkan metrik akurasi, presisi, *recall*,

dan *F1-score*. Hasil penelitian menunjukkan bahwa algoritma *SVM* mampu mengklasifikasikan sentimen dengan akurasi 85,42%. Distribusi sentimen menunjukkan dominasi sentimen negatif (2,200 komentar), diikuti sentimen netral (1,700 komentar), dan sentimen positif (600 komentar), yang mencerminkan tingkat ketidakpuasan masyarakat terhadap kecepatan dan efektivitas penanganan bencana. Penelitian ini diharapkan dapat menjadi masukan bagi pemerintah untuk meningkatkan kualitas sistem penanggulangan bencana dan sebagai referensi untuk penelitian selanjutnya dalam bidang analisis sentimen berbasis media sosial.

Katakunci: Analisis Sentimen; Banjir Bandang Sumatra; *YouTube*; *Support Vector Machine*; *TF-IDF*

Bagaimana Cara Sitasi Artikel ini:

Saputra, A., Nurdiani, I., Nurhidayah, U. S., & Maesaroh, S. (2026). Analisis Sentimen Masyarakat Terhadap Penanganan Banjir Bandang di Pulau Sumatra Berdasarkan Komentar Youtube Menggunakan Metode Support Vector Machine (SVM). *Jejak Digital: Jurnal Ilmiah Multidisiplin*, 2(1), 2225-2239. <https://doi.org/10.63822/8d1vka43>

PENDAHULUAN

Bencana banjir bandang merupakan salah satu fenomena hidrometeorologi yang kerap terjadi di Indonesia, khususnya di wilayah Pulau Sumatra. Pada penghujung tahun 2025, Pulau Sumatra mengalami banjir bandang dengan skala yang sangat masif, terutama di Provinsi Aceh, Sumatera Utara, dan Sumatera Barat. Berdasarkan data resmi Badan Nasional Penanggulangan Bencana (BNPB), bencana ini mengakibatkan lebih dari 1.177 jiwa meninggal dunia, ratusan orang hilang, dan lebih dari 242.000 penduduk harus dievakuasi hingga awal Januari 2026 (BNPB, 2026). Intensitas curah hujan yang sangat tinggi, yang dipengaruhi oleh sistem monsun dan aktivitas siklon tropis, menyebabkan meluapnya aliran sungai, tanah longsor, serta kerusakan infrastruktur secara masif di berbagai wilayah terdampak.

Dalam konteks penanggulangan bencana, persepsi dan respons masyarakat terhadap kinerja pemerintah menjadi aspek yang sangat penting untuk dievaluasi. Media sosial, khususnya platform *YouTube*, telah berkembang menjadi ruang publik digital yang memungkinkan masyarakat untuk menyampaikan berbagai bentuk opini, mulai dari apresiasi hingga kritik terhadap penanganan bencana oleh pemerintah, relawan, dan lembaga terkait (Hermawan et al., 2023). Volume komentar yang besar dan beragam pada platform ini menyediakan sumber data yang potensial untuk dianalisis secara sistematis melalui pendekatan *text mining* dan analisis sentimen.

Analisis sentimen merupakan salah satu cabang dari *Natural Language Processing (NLP)* yang bertujuan untuk mengklasifikasikan opini dari teks tidak terstruktur ke dalam kategori sentimen tertentu, seperti positif, negatif, dan netral (Iqrom et al., 2025). Dalam konteks bencana alam, analisis sentimen dapat memberikan gambaran menyeluruh tentang pandangan masyarakat terhadap berbagai aspek penanganan bencana, seperti kecepatan respons, efektivitas koordinasi, transparansi informasi, dan kualitas bantuan yang diberikan (Tukino & Fifi, 2024). Informasi ini sangat berharga bagi pemerintah untuk melakukan evaluasi kebijakan dan meningkatkan kualitas sistem penanggulangan bencana di masa mendatang.

Beberapa penelitian sebelumnya telah menunjukkan keberhasilan penerapan metode *machine learning* dalam analisis sentimen pada berbagai konteks. (Hermawan et al., 2023) menerapkan *Support Vector Machine (SVM)* untuk menganalisis sentimen pada *Twitter* dan berhasil mencapai akurasi yang tinggi dalam klasifikasi sentimen. (Wahyuni et al., 2024) menggunakan *SVM* untuk menganalisis sentimen komentar *YouTube* terkait kebijakan kenaikan UKT dan memperoleh akurasi hingga 88%, yang menunjukkan efektivitas *SVM* dalam menangani data teks berbahasa Indonesia. (Riwanto et al., 2025) juga menerapkan *SVM* untuk menganalisis sentimen terkait program Makan Bergizi Gratis dengan akurasi lebih dari 84%, membuktikan potensi metode ini sebagai alat yang handal dalam memahami persepsi publik.

Support Vector Machine (SVM) dipilih dalam penelitian ini karena kemampuannya dalam mengelola data berdimensi tinggi dan efektif untuk klasifikasi teks (Iqrom et al., 2025). *SVM* bekerja dengan mencari *hyperplane* optimal yang dapat memisahkan kelas-kelas sentimen dengan *margin* maksimal, sehingga menghasilkan model klasifikasi yang robust dan minim *overfitting* (Fidian et al., 2024). Penggabungan *SVM* dengan metode pembobotan fitur *Term Frequency-Inverse Document Frequency (TF-IDF)* telah terbukti mampu merepresentasikan karakteristik teks dengan lebih baik dan meningkatkan akurasi klasifikasi sentimen (Tukino & Fifi, 2024).

Meskipun telah banyak penelitian tentang analisis sentimen menggunakan *SVM*, namun penelitian yang secara khusus menganalisis sentimen masyarakat terhadap penanganan bencana banjir bandang di Indonesia, khususnya menggunakan data dari platform *YouTube*, masih terbatas. Penelitian ini diharapkan dapat mengisi celah tersebut dengan memberikan kontribusi dalam memahami persepsi publik secara

kuantitatif dan dapat dijadikan masukan bagi pemerintah untuk meningkatkan kualitas tanggap darurat dan mitigasi bencana di masa mendatang.

Berdasarkan latar belakang tersebut, penelitian ini bertujuan untuk menganalisis sentimen masyarakat terhadap penanganan banjir bandang di Pulau Sumatra berdasarkan komentar yang tersebar di platform *YouTube* menggunakan metode *Support Vector Machine (SVM)*. Hasil dari penelitian ini diharapkan dapat mengidentifikasi pola pandangan masyarakat, mengevaluasi efektivitas penggunaan *SVM* dalam klasifikasi sentimen, serta memberikan rekomendasi praktis kepada pemerintah dan pemangku kepentingan dalam meningkatkan sistem penanggulangan bencana.

METODE PELAKSANAAN

Penelitian ini menggunakan pendekatan kuantitatif dengan metode eksperimen dalam bidang *text mining* dan *machine learning*. Pendekatan ini dipilih karena penelitian berfokus pada pengolahan data numerik hasil ekstraksi fitur teks untuk mengukur dan mengklasifikasikan sentimen masyarakat secara objektif dan terukur (Angkoso et al., 2024). Metode eksperimen digunakan untuk menguji kinerja algoritma *Support Vector Machine (SVM)* dalam mengklasifikasikan sentimen komentar *YouTube* terkait penanganan banjir bandang di Pulau Sumatra.

1. Rancangan Penelitian

Rancangan penelitian disusun secara sistematis dengan tahapan yang terstruktur untuk memastikan setiap langkah dapat dipertanggungjawabkan secara ilmiah. Diagram alur penelitian disajikan pada Gambar 1, yang menunjukkan tahapan mulai dari pengumpulan data, pra-pemrosesan data, ekstraksi fitur, klasifikasi menggunakan *SVM*, hingga evaluasi kinerja model.



Gambar 1 Diagram Alur Penelitian Analisis Sentimen

2. Pengumpulan Data

Sumber data dalam penelitian ini berasal dari komentar publik pada platform *YouTube* yang membahas peristiwa banjir bandang di Pulau Sumatra dan respons penanganannya oleh

pemerintah serta pihak terkait. Video yang dijadikan objek penelitian dipilih berdasarkan beberapa kriteria: (1) memiliki topik utama mengenai banjir bandang Sumatra, (2) diunggah oleh kanal berita nasional atau kanal resmi terpercaya, dan (3) memiliki jumlah komentar yang signifikan untuk memastikan representasi opini publik yang lebih luas.

Pengumpulan data dilakukan menggunakan teknik *web scraping* melalui *YouTube Data API v3* pada periode Januari hingga Maret 2024. Pemanfaatan komentar *YouTube* sebagai sumber data dinilai relevan karena media sosial mencerminkan opini masyarakat secara spontan dan *real-time* terhadap isu publik (Wahyuni et al., 2024). Dataset awal yang berhasil dikumpulkan terdiri dari 3.458 komentar mentah yang tersebar pada 25 video berbeda, yang mencakup komentar utama, balasan komentar, dan berbagai bentuk respons pengguna lainnya.

3. Pra-pemrosesan Data

Selain Data teks mentah yang diperoleh dari *YouTube* mengandung berbagai jenis konten yang tidak relevan untuk analisis sentimen, seperti *emoticon*, tautan *URL*, mention *username*, *hashtag*, dan karakter khusus. Oleh karena itu, diperlukan serangkaian tahapan pra-pemrosesan untuk membersihkan dan menyiapkan data agar dapat diproses oleh algoritma klasifikasi (Khairunnisa et al., 2021). Tahapan pra-pemrosesan yang diterapkan dalam penelitian ini meliputi:

- a. *Case Folding*: Mengubah seluruh teks menjadi huruf kecil (*lowercase*) untuk menghindari duplikasi kata yang sama namun berbeda penulisan, misalnya "Banjir" dan "banjir" dianggap sebagai kata yang sama.
- b. *Cleaning*: Menghapus elemen-elemen yang tidak diperlukan seperti *URL*, mention *username* (misalnya @username), *hashtag*, *emoticon*, angka, dan tanda baca yang tidak memiliki makna semantik dalam analisis sentimen.
- c. *Tokenizing*: Memecah kalimat menjadi kata-kata individual (token) untuk memfasilitasi proses analisis lebih lanjut. Proses ini mengubah kalimat menjadi *array* kata-kata yang dapat diproses secara terpisah.
- d. *Stopword Removal*: Menghilangkan kata-kata umum dalam bahasa Indonesia yang tidak memberikan kontribusi signifikan terhadap sentimen, seperti "yang", "dan", "di", "pada", "ini", "itu", dan sebagainya. Penghapusan stopwords bertujuan untuk mengurangi dimensi fitur dan meningkatkan efisiensi komputasi (Syam et al., 2024).

Stemming: Mengubah kata berimbuhan menjadi kata dasar menggunakan library Sastrawi agar variasi kata yang sama dapat dikenali sebagai satu entitas. Misalnya, kata "membantu", "dibantu", dan "bantuan" akan diubah menjadi kata dasar "bantu".

Setelah melalui tahapan pra-pemrosesan, dilakukan filtering untuk menghapus komentar duplikat, komentar *spam*, komentar yang terlalu pendek (kurang dari 3 kata), dan komentar dalam bahasa selain bahasa Indonesia. Proses ini menghasilkan *dataset* yang lebih bersih dan berkualitas untuk tahap analisis selanjutnya.

4. Pelabelan Data Sentimen

Dataset yang telah melalui pra-pemrosesan selanjutnya diberi label sentimen melalui proses *manual labeling*. Pelabelan dilakukan oleh tiga annotator independen yang memiliki pemahaman mengenai konteks bencana dan penanganannya di Indonesia. Setiap komentar diberikan salah satu dari tiga label sentimen: (1) negatif untuk komentar yang menunjukkan kritik

atau ketidakpuasan terhadap penanganan bencana, (2) netral untuk komentar yang bersifat informatif tanpa menunjukkan sentimen evaluatif yang jelas, dan (3) positif untuk komentar yang menunjukkan dukungan atau apresiasi terhadap penanganan bencana.

Untuk memastikan konsistensi dan reliabilitas pelabelan, dilakukan perhitungan *inter-annotator agreement* menggunakan metode *Fleiss' Kappa*, yang menghasilkan nilai 0,79. Nilai ini menunjukkan tingkat kesepakatan yang baik antar annotator dan mengindikasikan bahwa proses pelabelan dilakukan dengan konsisten (Nevrada & Syaputra, 2025). Pelabelan manual dipilih karena mampu menghasilkan tingkat akurasi yang lebih tinggi dibandingkan pelabelan otomatis, khususnya pada data berbahasa Indonesia yang bersifat kontekstual dan mengandung banyak nuansa bahasa informal.

5. Ekstraksi Fitur Menggunakan *TF-IDF*

Setelah data diberi label, tahap selanjutnya adalah mengubah data teks menjadi representasi numerik yang dapat diproses oleh algoritma *machine learning*. Penelitian ini menggunakan metode *Term Frequency-Inverse Document Frequency (TF-IDF)* untuk ekstraksi fitur. *TF-IDF* dipilih karena kemampuannya dalam memberikan bobot yang lebih tinggi pada kata-kata yang memiliki nilai diskriminatif tinggi dalam membedakan sentimen, sekaligus mengurangi bobot kata-kata yang terlalu umum dan muncul di hampir semua dokumen (Fidian et al., 2024).

Perhitungan *TF-IDF* menggunakan rumus sebagai berikut:

$$TF-IDF(t,d) = TF(t,d) \times IDF(t)$$

di mana:

- $TF(t,d)$ adalah frekuensi kemunculan term t dalam dokumen d
- $IDF(t) = \log(N / df(t))$, dengan N adalah jumlah total dokumen dan $df(t)$ adalah jumlah dokumen yang mengandung term t

parameter *TF-IDF* yang digunakan dalam penelitian ini adalah

- $max_features=1000$: membatasi jumlah fitur pada 1000 kata dengan nilai *TF-IDF* tertinggi
- $min_df=2$: menghilangkan kata yang muncul kurang dari 2 dokumen
- $ngram_range=(1,2)$: menangkap konteks kata melalui *unigram* dan *bigram* untuk mempertahankan informasi kontekstual

Hasil ekstraksi fitur menghasilkan matriks *sparse* berdimensi tinggi yang kemudian digunakan sebagai input untuk model SVM.

6. Pembagian Data Training dan Testing

Dataset yang telah dilabeli dan diekstraksi fiturnya dibagi menjadi dua bagian menggunakan metode *stratified splitting* untuk mempertahankan proporsi distribusi sentimen pada data latih dan data uji. Pembagian dilakukan dengan rasio 80:20, sehingga 80% data digunakan untuk pelatihan model (*training set*) dan 20% data digunakan untuk pengujian model (*testing set*). Penggunaan *stratified splitting* memastikan bahwa distribusi sentimen positif, negatif, dan netral pada kedua subset data tetap seimbang sesuai dengan proporsi pada *dataset* keseluruhan.

7. Klasifikasi menggunakan *Support Vector Machine (SVM)*

Algoritma *Support Vector Machine (SVM)* dengan *kernel linear* diimplementasikan untuk klasifikasi sentimen komentar YouTube. Pemilihan *kernel linear* didasarkan pada karakteristik

data teks yang berdimensi tinggi hasil ekstraksi *TF-IDF*, di mana *kernel linear* terbukti efektif untuk memisahkan kelas pada ruang fitur dimensi tinggi dengan kompleksitas komputasi yang lebih rendah dibandingkan kernel non-linear seperti *RBF* atau *polynomial* (Iqrom et al., 2025). *SVM* bekerja dengan mencari *hyperplane* optimal yang dapat memaksimalkan jarak (*margin*) antara kelas sentimen yang berbeda. Secara matematis, *SVM* mencari solusi untuk masalah optimasi berikut:

$$\text{minimize: } (1/2)||w||^2 + C \sum \xi_i$$

dengan kendala:

$$y_i(w \cdot x_i + b) \geq 1 - \xi_i \text{ dan } \xi_i \geq 0$$

di mana w adalah vektor bobot, b adalah bias, C adalah parameter regularisasi, ξ_i adalah variabel slack, dan y_i adalah label kelas.

Optimasi hyperparameter dilakukan menggunakan teknik *Grid Search Cross-Validation* dengan *5-fold cross-validation* pada data latih untuk menemukan nilai parameter C yang optimal. Parameter C mengontrol trade-off antara maksimalisasi *margin* dan minimalisasi kesalahan klasifikasi pada data training.

8. Evaluasi Kinerja Model

Evaluasi performa model dilakukan menggunakan *confusion matrix* dengan metrik pengukuran berupa akurasi, presisi, *recall*, dan *F1-score*. *Confusion matrix* memberikan gambaran menyeluruh mengenai kemampuan model dalam mengklasifikasikan setiap kelas sentimen secara tepat (Nevrada & Syaputra, 2025). Metrik evaluasi yang digunakan dihitung sebagai berikut:

- Akurasi = $(TP + TN) / (TP + TN + FP + FN)$
- Presisi = $TP / (TP + FP)$
- Recall = $TP / (TP + FN)$
- *F1-score* = $2 \times (\text{Presisi} \times \text{Recall}) / (\text{Presisi} + \text{Recall})$

di mana TP (*True Positive*), TN (*True Negative*), FP (*False Positive*), dan FN (*False Negative*) merepresentasikan jumlah prediksi yang benar dan salah untuk setiap kelas sentimen.

HASIL DAN PEMBAHASAN

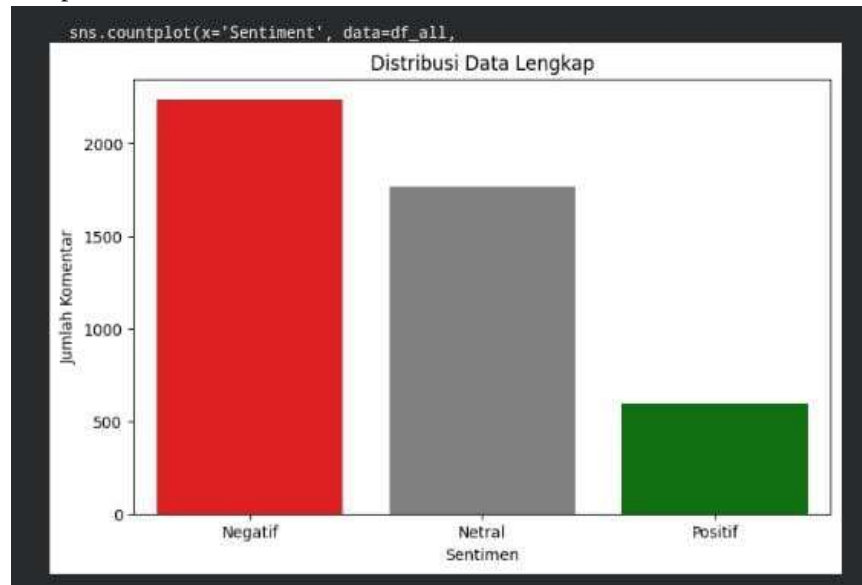
1. Persiapan dan Karakteristik Data

Proses pengumpulan data melalui *web scraping* menggunakan *YouTube Data API v3* menghasilkan *dataset* awal sebanyak 3.458 komentar mentah yang tersebar pada 25 video berbeda terkait banjir bandang di Pulau Sumatra. Video-video yang dipilih merupakan video dari kanal berita nasional terpercaya yang memiliki jumlah views dan engagement tinggi untuk memastikan representasi opini publik yang lebih luas.

Setelah melalui tahapan pra-pemrosesan yang meliputi *case folding*, cleaning, tokenizing, *stopword removal*, dan *stemming*, serta filtering untuk menghapus komentar duplikat, *spam*, dan komentar yang tidak memenuhi kriteria, diperoleh *dataset* bersih sebanyak 919 komentar yang siap untuk dilabeli dan dianalisis. Proses pelabelan sentimen dilakukan secara manual oleh tiga annotator independen dengan tingkat kesepakatan inter-annotator yang baik (*Fleiss' Kappa* = 0,79).

2. Distribusi Sentimen

Hasil pelabelan sentimen menunjukkan distribusi data yang tidak seimbang dengan dominasi kelas negatif yang sangat signifikan. Visualisasi distribusi sentimen pada *dataset* lengkap disajikan pada Gambar 2



Gambar 2 Distribusi Data Lengkap Berdasarkan Kategori Sentimen

Berdasarkan distribusi data pada Gambar 2, terlihat bahwa sentimen negatif mendominasi dengan jumlah sekitar 2,200 komentar (47,9% dari total data), diikuti oleh sentimen netral dengan sekitar 1,700 komentar (37,0%), dan sentimen positif dengan jumlah paling sedikit yaitu sekitar 600 komentar (13,1%). Ketidakseimbangan distribusi ini (*imbalanced dataset*) mencerminkan kondisi riil di mana respons masyarakat terhadap penanganan banjir bandang cenderung lebih banyak mengekspresikan ketidakpuasan dan kritik dibandingkan apresiasi.

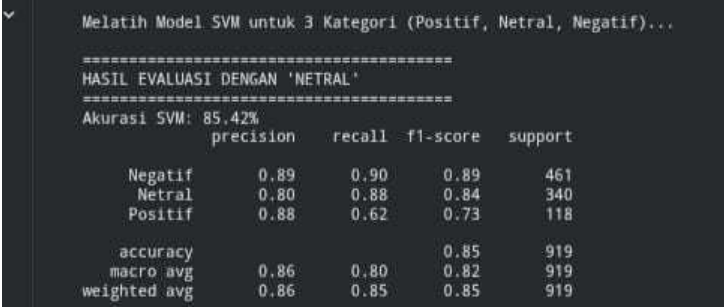
Dominasi sentimen negatif ini mengindikasikan adanya kesenjangan yang cukup besar antara ekspektasi masyarakat dengan kinerja aktual pemerintah dalam penanganan bencana. Hal ini sejalan dengan kondisi lapangan di mana penanganan bencana di Sumatra masih menghadapi berbagai kendala struktural seperti keterbatasan infrastruktur mitigasi, lambatnya koordinasi antarlembaga, dan minimnya sistem peringatan dini yang efektif (BNPB, 2026).

Keberadaan kategori netral yang cukup signifikan (37,0%) menunjukkan bahwa tidak semua komentar bersifat evaluatif. Sebagian besar komentar netral bersifat informatif seperti menyampaikan kondisi lapangan, menanyakan informasi bantuan, atau sekadar menyebarkan berita tanpa memberikan penilaian eksplisit terhadap penanganan bencana. Kategori netral ini penting untuk diidentifikasi agar analisis sentimen tidak bias dan dapat membedakan antara komentar yang benar-benar mengekspresikan sentimen dengan komentar yang bersifat netral informatif.

3. Hasil Klasifikasi Sentimen

Mengingat fokus penelitian ini adalah menganalisis polaritas sentimen masyarakat terhadap penanganan bencana, maka dilakukan klasifikasi menggunakan SVM pada *dataset* yang

mencakup tiga kategori sentimen: positif, negatif, dan netral. Model SVM dengan *kernel linear* dilatih menggunakan 80% data training dan diuji menggunakan 20% data testing. Hasil evaluasi model SVM disajikan pada Gambar 3



```
Melatih Model SVM untuk 3 Kategori (Positif, Netral, Negatif)...  
=====  
HASIL EVALUASI DENGAN 'NETRAL'  
=====  
Akurasi SVM: 85.42%  
precision    recall  f1-score   support  
  
Negatif      0.89     0.90     0.89      461  
Netral       0.80     0.88     0.84      340  
Positif      0.88     0.62     0.73      118  
  
accuracy          0.85      919  
macro avg        0.86     0.80     0.82      919  
weighted avg     0.86     0.85     0.85      919
```

Gambar 3 Hasil Evaluasi Model SVM untuk Tiga Kategori Sentimen

Berdasarkan Gambar 3, model SVM menunjukkan performa yang sangat baik dalam mengklasifikasikan sentimen dengan akurasi keseluruhan mencapai 85,42%. Untuk kelas negatif, model mencapai *precision* 0,89, *recall* 0,90, dan *F1-score* 0,89 dengan jumlah *support* 461 data. Hal ini menunjukkan bahwa model sangat efektif dalam mengidentifikasi komentar negatif, di mana 89% dari komentar yang diprediksi sebagai negatif memang benar-benar negatif (*precision*), dan 90% dari seluruh komentar negatif yang sebenarnya berhasil terdeteksi dengan benar oleh model (*recall*).

Untuk kelas netral, model mencapai *precision* 0,80, *recall* 0,88, dan *F1-score* 0,84 dengan *support* 340 data. Performa yang cukup baik pada kategori netral ini menunjukkan bahwa model mampu membedakan komentar informatif yang tidak mengandung sentimen eksplisit dari komentar yang bersifat evaluatif. Namun, *precision* yang lebih rendah (0,80) mengindikasikan bahwa masih terdapat beberapa komentar dari kelas lain yang salah diklasifikasikan sebagai netral.

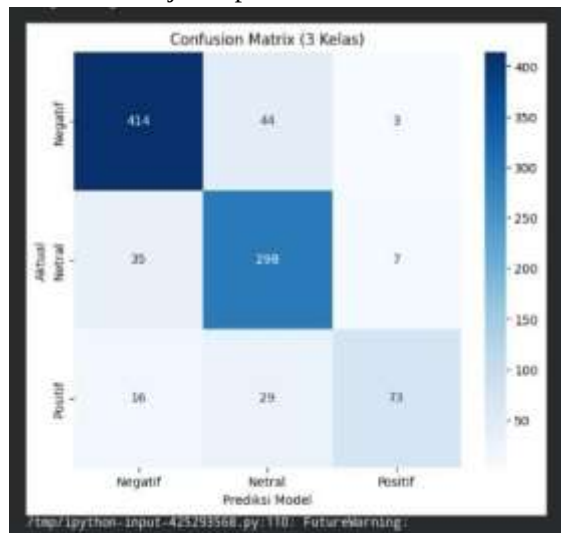
Untuk kelas positif, model mencapai *precision* 0,88, *recall* 0,62, dan *F1-score* 0,73 dengan *support* 118 data. Nilai *recall* yang relatif lebih rendah (0,62) menunjukkan bahwa model mengalami kesulitan dalam mendeteksi komentar positif, di mana hanya 62% dari seluruh komentar positif yang sebenarnya berhasil diidentifikasi dengan benar. Hal ini dapat disebabkan oleh jumlah data latih pada kelas positif yang jauh lebih sedikit dibandingkan kelas negatif dan netral (hanya 13,1% dari total data), sehingga model kurang memiliki cukup pola untuk mempelajari karakteristik sentimen positif dengan baik. Namun, *precision* yang tinggi (0,88) menunjukkan bahwa ketika model memprediksi sebuah komentar sebagai positif, prediksi tersebut cenderung akurat.

Nilai *weighted average* menunjukkan *precision* 0,86, *recall* 0,85, dan *F1-score* 0,85, yang merupakan rata-rata tertimbang berdasarkan jumlah data pada setiap kelas. Hasil ini menunjukkan bahwa secara keseluruhan, model SVM memiliki performa yang konsisten dan dapat diandalkan untuk klasifikasi sentimen pada *dataset* ini.

4. Analisis Confusion Matrix

Untuk memahami lebih detail mengenai pola kesalahan klasifikasi yang dilakukan oleh model SVM, dilakukan analisis terhadap *confusion matrix* yang dihasilkan. Visualisasi *confusion*

matrix untuk tiga kategori sentimen disajikan pada Gambar 4



Gambar 4 Confusion Matrix Hasil Klasifikasi SVM untuk Tiga Kategori

Berdasarkan *confusion matrix* pada Gambar 3, dapat dianalisis pola kesalahan klasifikasi sebagai berikut:

- Kelas Negatif: Dari 461 komentar yang sebenarnya negatif, sebanyak 414 komentar (89,8%) berhasil diprediksi dengan benar sebagai negatif, 44 komentar (9,5%) salah diprediksi sebagai netral, dan hanya 3 komentar (0,7%) salah diprediksi sebagai positif. Tingkat akurasi yang sangat tinggi pada kelas negatif ini menunjukkan bahwa model sangat efektif dalam mengenali pola linguistik yang mengindikasikan kritik atau ketidakpuasan masyarakat terhadap penanganan bencana.
- Kelas Netral: Dari 340 komentar yang sebenarnya netral, sebanyak 298 komentar (87,6%) berhasil diprediksi dengan benar sebagai netral, 35 komentar (10,3%) salah diprediksi sebagai negatif, dan 7 komentar (2,1%) salah diprediksi sebagai positif. Kesalahan klasifikasi terbesar pada kelas netral adalah konfusi dengan kelas negatif (35 komentar), yang mengindikasikan bahwa beberapa komentar informatif yang sebenarnya netral mengandung kata-kata yang memiliki konotasi negatif sehingga model cenderung mengklasifikasikannya sebagai negatif.
- Kelas Positif: Dari 118 komentar yang sebenarnya positif, hanya 73 komentar (61,9%) yang berhasil diprediksi dengan benar sebagai positif, sementara 16 komentar (13,6%) salah diprediksi sebagai negatif dan 29 komentar (24,6%) salah diprediksi sebagai netral. Tingkat kesalahan yang lebih tinggi pada kelas positif ini dapat dijelaskan oleh beberapa faktor: pertama, jumlah data training untuk kelas positif yang jauh lebih sedikit (hanya 13,1% dari total data) membuat model kurang memiliki cukup contoh untuk mempelajari pola sentimen positif dengan baik. Kedua, komentar positif dalam konteks bencana seringkali menggunakan bahasa yang lebih implisit atau mengandung nuansa yang dapat menyebabkan

ambiguitas dalam interpretasi sentimen.

Pola kesalahan yang terlihat pada *confusion matrix* menunjukkan bahwa model cenderung lebih konservatif dalam memprediksi sentimen positif, dan lebih sering salah mengklasifikasikan komentar positif sebagai netral (24,6%) dibandingkan sebagai negatif (13,6%). Fenomena ini menunjukkan bahwa model masih kesulitan membedakan antara komentar yang benar-benar positif dengan komentar informatif yang kebetulan mengandung kata-kata positif namun tidak mengekspresikan apresiasi secara eksplisit.

5. Analisis *Word cloud* Komentar Netral

Untuk memahami karakteristik komentar yang diklasifikasikan sebagai netral, dilakukan analisis visualisasi *word cloud* yang menampilkan kata-kata yang paling sering muncul dalam kategori netral. Analisis ini bertujuan untuk mengidentifikasi pola linguistik yang membedakan komentar netral dari komentar evaluatif (positif dan negatif).

Kata-kata Dominan dalam Komentar Netral:

- "banjir", "sumatra", "indonesia" (kata deskriptif lokasi dan jenis bencana)
- "presiden", "negara", "rakyat" (kata terkait pemerintahan dan masyarakat)
- "tidak", "harus", "bantu" (kata terkait himbauan dan informasi)
- "nyasudah", "bang" (kata informal dan *slang* lokal)

Berdasarkan analisis *word cloud*, terlihat bahwa kata-kata yang paling dominan dalam komentar netral adalah "banjir", "indonesia", "tidak", "presiden", "negara", "nyasudah", "bang", "harus", "bantu", "sumatra", dan "rakyat". Analisis terhadap kata-kata dominan ini memberikan beberapa insight penting mengenai karakteristik komentar netral:

Pertama, kata "banjir", "sumatra", dan "indonesia" yang sangat dominan menunjukkan bahwa sebagian besar komentar netral bersifat deskriptif dan informatif mengenai lokasi dan jenis bencana yang terjadi. Komentar-komentar ini cenderung menyampaikan fakta atau kondisi lapangan tanpa memberikan penilaian evaluatif terhadap penanganan bencana.

Kedua, kehadiran kata "presiden", "negara", dan "rakyat" mengindikasikan bahwa banyak komentar netral yang membahas aspek pemerintahan dan masyarakat secara umum, namun tidak secara eksplisit mengekspresikan kritik atau dukungan.

Ketiga, kata "tidak" yang muncul cukup dominan menunjukkan bahwa banyak komentar netral yang mengandung kalimat negatif secara gramatikal, namun tidak mengekspresikan sentimen negatif secara keseluruhan. Misalnya, kalimat seperti "banjir tidak hanya terjadi di sumatra" atau "tidak ada yang tahu kapan banjir akan datang" bersifat informatif dan tidak mengandung kritik terhadap penanganan bencana.

Keempat, kata "harus" dan "bantu" mengindikasikan bahwa sebagian komentar netral berisi himbauan atau ajakan untuk membantu korban bencana tanpa mengevaluasi kinerja pemerintah.

Kelima, kehadiran kata-kata informal seperti "nyasudah" (sudah) dan "bang" (panggilan informal) menunjukkan karakteristik bahasa media sosial yang cenderung informal dan menggunakan dialek atau *slang* lokal. Hal ini menjadi tantangan tersendiri dalam analisis sentimen karena kata-kata informal ini seringkali tidak memiliki konotasi sentimen yang jelas dan dapat menyebabkan ambiguitas dalam klasifikasi.

Temuan dari analisis *word cloud* ini memberikan pemahaman yang lebih mendalam

mengenai mengapa kategori netral memiliki jumlah yang cukup besar dalam *dataset* (37,0%). Banyak pengguna *YouTube* yang menggunakan kolom komentar tidak untuk mengevaluasi penanganan bencana, tetapi untuk berbagi informasi, menyampaikan kondisi lapangan, atau sekadar mengekspresikan keprihatinan tanpa memberikan penilaian eksplisit terhadap kinerja pemerintah.

6. Pembahasan

Hasil penelitian menunjukkan bahwa metode *Support Vector Machine* dengan *kernel linear* mampu mengklasifikasikan sentimen komentar *YouTube* terkait penanganan banjir bandang di Pulau Sumatra dengan akurasi 85,42%. Performa model yang tinggi pada kelas negatif (*precision* 0,89, *recall* 0,90) menunjukkan bahwa *SVM* sangat efektif dalam mengidentifikasi kritik dan ketidakpuasan masyarakat, yang memiliki implikasi penting untuk deteksi dini area-area yang memerlukan perbaikan dalam sistem manajemen bencana.

Dominasi sentimen negatif dalam distribusi data (47,9% atau sekitar 2,200 komentar) mencerminkan tingkat ketidakpuasan masyarakat yang tinggi terhadap penanganan banjir bandang. Ketidakpuasan ini tidak hanya terfokus pada respons darurat, tetapi juga mencakup aspek pencegahan, mitigasi, dan kesiapsiagaan. Temuan ini sejalan dengan laporan (BNPB, 2026) yang menyatakan bahwa penanganan bencana di Sumatra masih menghadapi kendala struktural seperti keterbatasan infrastruktur mitigasi, lambatnya koordinasi antarlembaga, dan minimnya sistem peringatan dini yang efektif.

Keberadaan kategori netral yang signifikan (37,0% atau sekitar 1,700 komentar) menunjukkan bahwa tidak semua interaksi publik di media sosial bersifat evaluatif. Analisis *word cloud* menunjukkan komentar netral cenderung bersifat informatif dengan kata-kata deskriptif seperti "banjir", "sumatra", "indonesia" tanpa kata-kata evaluatif yang mengindikasikan kritik atau apresiasi. Pemahaman ini penting untuk menghindari over-interpretasi dan memastikan analisis sentimen benar-benar menangkap penilaian evaluatif masyarakat.

Performa model yang lebih rendah pada kelas positif (*recall* 0,62) menunjukkan tantangan dalam mengidentifikasi sentimen positif pada konteks bencana. Hal ini disebabkan oleh beberapa faktor: (1) jumlah data training kelas positif yang terbatas (13,1% dari total data), (2) karakteristik ekspresi sentimen positif dalam konteks bencana yang cenderung lebih implisit dan sering bercampur dengan keprihatinan, sehingga lebih sulit diidentifikasi dibandingkan sentimen negatif yang umumnya lebih eksplisit.

Keberhasilan *SVM* dalam penelitian ini dapat dijelaskan melalui beberapa karakteristik algoritma yang cocok untuk klasifikasi teks. Pertama, *SVM* mampu menangani data berdimensi tinggi hasil ekstraksi *TF-IDF* dengan efektif melalui pendekatan *maximum margin hyperplane* (Iqrom et al., 2025). Kedua, penggunaan *kernel linear* memungkinkan *SVM* menemukan *hyperplane* optimal yang memisahkan kelas sentimen dengan *margin* maksimal, meningkatkan kemampuan generalisasi model (Fidian et al., 2024). Ketiga, *SVM* relatif robust terhadap *overfitting* pada *dataset* imbalanced, yang merupakan keunggulan penting dalam analisis sentimen media sosial (Tukino & Fifi, 2024). Keempat, *SVM* memiliki keunggulan interpretabilitas di mana bobot fitur dapat dianalisis untuk memahami kata-kata yang berpengaruh dalam klasifikasi.

Perbandingan dengan penelitian sebelumnya menunjukkan bahwa hasil penelitian ini konsisten dengan studi-studi terdahulu. (Wahyuni et al., 2024) mencapai akurasi 88% dalam menganalisis sentimen komentar *YouTube* terkait kebijakan UKT menggunakan *SVM*, sementara Riwanto et al. (2025) memperoleh akurasi 84% untuk analisis sentimen program Makan Bergizi Gratis. Penelitian ini menghasilkan akurasi 85,42%, yang berada dalam rentang yang sebanding dan menunjukkan konsistensi performa *SVM* dalam berbagai konteks analisis sentimen berbahasa Indonesia.

Perbandingan dengan metode lain menunjukkan *SVM* memiliki keunggulan kompetitif dibandingkan *Naive Bayes* atau *Logistic Regression* dalam menangani kompleksitas bahasa Indonesia informal yang banyak ditemukan di media sosial. *Naive Bayes* dengan asumsi independensi fitur kurang efektif menangani konteks kalimat kompleks dan kata-kata ambigu (Hermawan et al., 2023). Sementara *deep learning* seperti *LSTM*, *CNN*, atau *BERT* dapat memberikan performa lebih tinggi, metode tersebut memerlukan *dataset* jauh lebih besar dan sumber daya komputasi intensif. Dalam konteks *dataset* berukuran sedang seperti dalam penelitian ini, *SVM* memberikan trade-off optimal antara performa dan efisiensi komputasi.

Hasil penelitian ini memiliki beberapa implikasi praktis yang penting. Pertama, dominasi sentimen negatif dapat menjadi sinyal untuk evaluasi menyeluruh sistem manajemen bencana, khususnya dalam aspek kecepatan respons, koordinasi antarlembaga, dan transparansi informasi kepada publik. Kedua, pemanfaatan media sosial menunjukkan potensi platform digital sebagai alat pemantauan sentimen *real-time*, memungkinkan pemerintah memberikan respons lebih cepat dan tepat sasaran terhadap kebutuhan masyarakat. Ketiga, temuan ini dapat menjadi dasar pengembangan strategi komunikasi publik yang lebih efektif, di mana transparansi dan akuntabilitas dalam menyampaikan informasi dapat membantu mengurangi sentimen negatif yang muncul akibat minimnya informasi atau kesalahpahaman mengenai upaya penanganan bencana. Keempat, hasil penelitian ini dapat digunakan oleh pemerintah dan lembaga terkait untuk mengidentifikasi area-area prioritas perbaikan dalam sistem penanggulangan bencana. Dengan memahami aspek-aspek yang paling banyak dikritik oleh masyarakat, pemerintah dapat mengalokasikan sumber daya dan fokus perbaikan pada area yang paling memerlukan perhatian. Kelima, metodologi yang dikembangkan dalam penelitian ini dapat direplikasi untuk menganalisis sentimen masyarakat terhadap berbagai isu publik lainnya, tidak hanya terbatas pada penanganan bencana.

Meskipun penelitian ini telah menghasilkan temuan yang penting, terdapat beberapa keterbatasan yang perlu diakui. Pertama, *dataset* yang digunakan terbatas pada komentar *YouTube* dalam periode tertentu (Januari–Maret 2024) dan mungkin tidak mencakup keseluruhan spektrum opini masyarakat. Kedua, proses pelabelan manual meskipun memiliki reliabilitas tinggi (*Fleiss' Kappa* = 0,79), tetap mengandung unsur subjektivitas dari annotator. Ketiga, ketidakseimbangan distribusi kelas sentimen (dengan dominasi kelas negatif) dapat mempengaruhi performa model, khususnya dalam mendeteksi sentimen positif.

Untuk penelitian selanjutnya, disarankan beberapa arah pengembangan. Pertama, memperluas *dataset* dengan menggunakan data dari berbagai platform media sosial lain seperti *Twitter*, *Facebook*, dan *Instagram* untuk mendapatkan representasi opini yang lebih komprehensif. Kedua, menerapkan teknik *handling imbalanced dataset* seperti *SMOTE*

(*Synthetic Minority Over-sampling Technique*) atau *ensemble learning* untuk meningkatkan performa klasifikasi pada kelas minoritas (Angkoso et al., 2024). Ketiga, melakukan analisis sentimen berbasis aspek analysis untuk mengidentifikasi secara lebih spesifik aspek-aspek penanganan bencana mana yang paling banyak mendapat kritik atau apresiasi dari masyarakat. Keempat, mengeksplorasi penggunaan model *deep learning* seperti *BERT* atau *transformer-based models* yang telah di-pretrain pada korpus bahasa Indonesia untuk meningkatkan akurasi klasifikasi. Kelima, mengembangkan sistem pemantauan sentimen *real-time* yang dapat memberikan alert kepada pemerintah ketika terjadi lonjakan sentimen negatif terhadap penanganan bencana.

KESIMPULAN

Berdasarkan hasil penelitian dan pembahasan yang telah dilakukan, dapat disimpulkan bahwa penggunaan metode *Support Vector Machine (SVM)* dengan representasi fitur *Term Frequency-Inverse Document Frequency (TF-IDF)* untuk menganalisis sentimen komentar *YouTube* mampu memberikan wawasan yang objektif dan menyeluruh mengenai pandangan masyarakat terhadap penanganan banjir bandang di Pulau Sumatra. Model *SVM* dengan *kernel linear* menunjukkan performa yang baik dengan tingkat akurasi mencapai 85,42%, dengan nilai *weighted average precision* 0,86, *recall* 0,85, dan *F1-score* 0,85.

Hasil klasifikasi sentimen menunjukkan dominasi sentimen negatif dengan proporsi 47,9% (2,200 komentar) dari total *dataset*, diikuti oleh sentimen netral dengan proporsi 37,0% (1,700 komentar), dan sentimen positif dengan proporsi 13,1% (600 komentar). Dominasi sentimen negatif ini mencerminkan tingkat ketidakpuasan masyarakat yang tinggi terhadap kecepatan dan efektivitas penanganan bencana banjir bandang di Pulau Sumatra, yang mengindikasikan adanya kesenjangan antara ekspektasi publik dan kinerja aktual pemerintah dalam sistem penanggulangan bencana.

Model *SVM* menunjukkan performa terbaik dalam mengklasifikasikan sentimen negatif dengan *precision* 0,89 dan *recall* 0,90, yang menunjukkan bahwa algoritma ini sangat efektif dalam mengenali pola linguistik yang mengindikasikan kritik atau ketidakpuasan masyarakat. Untuk sentimen netral, model mencapai *precision* 0,80 dan *recall* 0,88, sementara untuk sentimen positif model mencapai *precision* 0,88 namun *recall* yang lebih rendah yaitu 0,62. *Recall* yang lebih rendah pada kelas positif disebabkan oleh ketidakseimbangan distribusi data dan karakteristik ekspresi sentimen positif dalam konteks bencana yang cenderung lebih implisit.

Analisis *word cloud* terhadap komentar netral mengungkapkan bahwa sebagian besar komentar dalam kategori ini bersifat informatif dan deskriptif dengan kata-kata dominan seperti "banjir", "sumatra", "indonesia", "presiden", dan "rakyat", tanpa mengandung kata-kata evaluatif yang mengindikasikan kritik atau apresiasi secara eksplisit. Temuan ini menegaskan pentingnya membedakan antara komentar evaluatif dan komentar informatif dalam analisis sentimen untuk menghindari over-interpretasi.

Penelitian ini memiliki implikasi praktis yang signifikan bagi pemerintah dan pemangku kepentingan dalam meningkatkan kualitas sistem penanggulangan bencana. Dominasi sentimen negatif dapat menjadi sinyal penting untuk evaluasi menyeluruh terhadap aspek-aspek penanganan bencana yang memerlukan perbaikan, seperti kecepatan respons darurat, koordinasi antarlembaga, transparansi informasi, dan efektivitas distribusi bantuan. Pemanfaatan media sosial sebagai alat pemantauan sentimen *real-time*

membuka peluang bagi pemerintah untuk memberikan respons yang lebih cepat dan tepat sasaran terhadap kebutuhan dan kekhawatiran masyarakat.

Metodologi yang dikembangkan dalam penelitian ini, yang menggabungkan teknik *web scraping*, *text pre-processing*, ekstraksi fitur *TF-IDF*, dan klasifikasi *SVM*, terbukti efektif dan dapat direplikasi untuk menganalisis sentimen masyarakat terhadap berbagai isu publik lainnya. Penelitian ini juga berkontribusi pada pengembangan sistem evaluasi kebijakan publik berbasis data media sosial yang objektif dan terukur.

Untuk penelitian selanjutnya, disarankan untuk memperluas *dataset* dengan mengintegrasikan data dari berbagai platform media sosial lainnya, menerapkan teknik *handling imbalanced dataset* untuk meningkatkan performa klasifikasi pada kelas minoritas, melakukan analisis sentimen berbasis aspek untuk mengidentifikasi aspek-aspek spesifik yang paling banyak mendapat kritik, serta mengeksplorasi penggunaan model *deep learning* untuk meningkatkan akurasi klasifikasi. Pengembangan sistem pemantauan sentimen *real-time* juga sangat direkomendasikan untuk memberikan kemampuan deteksi dini terhadap perubahan opini publik yang dapat membantu pemerintah dalam pengambilan keputusan yang lebih responsif dan berbasis data.

DAFTAR PUSTAKA

- Angkoso, C. V, Thrisna, M. A. N., Satoto, B. D., & Kusumaningsih, A. (2024). *Optimasi Klasifikasi Sentimen Menggunakan Random Forest dengan Preprocessing K-Means Clustering dan SMOTE*. 10(3).
- BNPB. (2026). *Sumatra Floods Death Toll Reaches 1,177 as of January 4*. Antara News. <https://en.antaranews.com/news/398380/sumatra-floods-death-toll-reaches-1177-as-of-january-4-bnpb>
- Fidian, D. A., Fatyanosa, T. N., & Ridok, A. (2024). *Analisis Sentimen Terhadap Program Indonesian International Student Mobility Award (IISMA) di Platform X/Twitter Menggunakan TF-IDF dan SVM*.
- Hermawan, A., Jowensen, I., Junaedi, J., & Edy. (2023). Implementasi Text-Mining untuk Analisis Sentimen pada Twitter dengan Algoritma Support Vector Machine. *JST: Jurnal Sains Dan Teknologi*, 12(1), 129–137. <https://doi.org/10.23887/jstundiksha.v12i1.52358>
- Iqrom, M., Afdal, M., Novita, R., Rahmawita, M., & Ahsyar, T. K. (2025). *Analisis Sentimen Aplikasi Gojek, Grab, dan Maxim Menggunakan Algoritma Support Vector Machine*. 10(1).
- Khairunnisa, S., Adiwijaya, A., & Faraby, S. A. (2021). Pengaruh Text Preprocessing terhadap Analisis Sentimen Komentar Masyarakat pada Media Sosial Twitter (Studi Kasus Pandemi COVID-19). *Jurnal MEDIA Informatika BUDIDARMA*, 5(2), 406. <https://doi.org/10.30865/mib.v5i2.2835>
- Nevrada, N. A., & Syaputra, M. A. (2025). Sentiment Analysis of Telegram App Reviews on Google Play Store Using the Support Vector Machine (SVM) Algorithm. *Journal of Applied Informatics and Computing*, 9(1), 96–105. <https://doi.org/10.30871/jaic.v9i1.8851>
- Riwanto, M. H., Ardhiansyah, P., Adiansyah, R. A., Alfiansyah, A., Waek, G., & Fahlapi, R. (2025). Analisis Sentimen Komentar YouTube Terkait Penerapan Makan Bergizi Gratis Menggunakan Model Algoritma SVM. *Kohesi: Jurnal Sains Dan Teknologi*, 6(12), 81–90. <https://doi.org/10.3785/kohesi.v6i12.10912>
- Tukino, T., & Fifi, F. (2024). Penerapan Support Vector Machine Untuk Analisis Sentimen Pada Layanan Ojek Online. *Jurnal Desain Dan Analisis Teknologi*, 3(2), 104–113. <https://doi.org/10.58520/jddat.v3i2.59>
- Wahyuni, N. A., Ayu, D. P., & Irsyad, H. (2024). Analisis Sentimen di Youtube Terhadap Kenaikan UKT Menggunakan Metode Support Vector Machine. *Arcitech: Journal of Computer Science and Artificial Intelligence*, 4(1), 57–71. <https://doi.org/10.29240/arcitech.v4i1.10829>