

Analisis Sentimen Publik Terhadap Implementasi Kurikulum Merdeka: Pendekatan Algoritma Random Forest dan Support Vector Machine

1st Senna Yoga Abira
Program Studi Sains Data
Universitas Telkom
Surabaya, Indonesia
sennayga@gmail.com

2nd Rifdatun Ni'mah
Program Studi Sains Data
Universitas Telkom
Surabaya, Indonesia
rifdatun@telkomuniversity.ac.id

3rd Regita Putri Permata
Program Studi Sains Data
Universitas Telkom
Surabaya, Indonesia
regitapermata@telkomuniversity.ac.id

Abstrak—Krisis pendidikan di Indonesia mendorong pemerintah menerapkan Kurikulum Merdeka sebagai strategi peningkatan mutu pembelajaran. Namun, respons publik terhadap implementasinya masih menuai pro dan kontra, terutama di media sosial. Penelitian ini menganalisis sentimen publik terhadap Kurikulum Merdeka menggunakan 10.066 data dari platform X yang dikumpulkan selama satu tahun (Oktober 2023–Oktober 2024). Data dianalisis menggunakan algoritma *Support Vector Machine* (SVM) dan *Random Forest* (RF), melalui tahapan pelabelan manual, praproses teks, ekstraksi fitur dengan TF-IDF, serta penambahan data sintetis ke kelas minoritas menggunakan metode SMOTE. Model dievaluasi dalam tiga tahap *hyperparameter tuning*. Hasil menunjukkan bahwa model SVM memberikan performa terbaik dengan akurasi 72,42% dan *F1-score* makro 69,67%, dibandingkan RF yang mencapai akurasi 68,49% dan *F1-score* makro 66,25%. Sentimen netral dan negatif lebih mendominasi opini publik, sementara sentimen positif relatif rendah. Penambahan data sintetis terbukti meningkatkan kemampuan model dalam mengenali kelas minoritas. Penelitian ini memberikan gambaran empiris mengenai persepsi publik terhadap kebijakan pendidikan, sekaligus menunjukkan potensi analisis sentimen berbasis media sosial sebagai alat evaluasi kebijakan secara *real-time*.

Kata kunci—analisis sentimen, kurikulum merdeka, media sosial, *random forest*, *support vector machine*, X

I. PENDAHULUAN

Isi Krisis pendidikan yang berkepanjangan di Indonesia telah menjadi sorotan berbagai studi, terutama terkait rendahnya kemampuan dasar siswa dalam matematika, literasi, dan pemahaman bacaan [1][2]. Sebagai respons terhadap kondisi tersebut, pemerintah meluncurkan Kurikulum Merdeka pada tahun 2022 sebagai bagian dari strategi pemulihan pembelajaran pascapandemi dan reformasi pendidikan. Kurikulum ini dirancang untuk memberikan fleksibilitas bagi guru dan sekolah dalam menyesuaikan pembelajaran dengan kebutuhan lokal, karakteristik siswa, dan konteks sekolah [3]. Salah satu fokus utamanya adalah menyederhanakan materi, mendorong kreativitas guru, dan meningkatkan keterlibatan siswa melalui pembelajaran berbasis proyek [4].

Meskipun Kurikulum Merdeka memiliki tujuan yang progresif, implementasinya di lapangan menimbulkan respons publik yang beragam. Sebagian masyarakat mendukung pendekatan yang lebih fleksibel dan kontekstual, namun sebagian lainnya menunjukkan keraguan terhadap efektivitas dan kesiapan sistem pendidikan dalam menerapkan perubahan ini. Oleh karena itu, penting untuk memahami bagaimana persepsi publik berkembang, khususnya melalui media sosial yang menjadi sarana ekspresi terbuka masyarakat.

Salah satu pendekatan yang relevan untuk mengevaluasi persepsi publik adalah analisis sentimen. Analisis ini memungkinkan identifikasi opini, sikap, dan perasaan publik terhadap topik tertentu dengan memanfaatkan data teks tidak terstruktur dalam skala besar [5]. Platform X (sebelumnya Twitter) menjadi salah satu sumber utama karena karakteristiknya yang memungkinkan diskusi publik terbuka. Platform X memiliki lebih dari 590 juta pengguna aktif global pada Oktober 2024 [6].

Dalam penelitian ini, data dikumpulkan dari platform X menggunakan *Tweet Harvest*, alat pengambil data berbasis *command-line* yang memanfaatkan *Playwright*. Proses dilanjutkan dengan pelabelan manual ke dalam tiga kelas sentimen (positif, negatif, netral) dan dianalisis menggunakan pendekatan *text mining* serta algoritma *machine learning*, yaitu *Random Forest* (RF) dan *Support Vector Machine* (SVM). Kedua algoritma dipilih karena efektivitasnya dalam klasifikasi teks dan kemampuan menangani data berdimensi tinggi [7][8].

Penelitian ini berkontribusi dalam memberikan gambaran empiris mengenai persepsi publik terhadap Kurikulum Merdeka, serta mengevaluasi performa dua algoritma klasifikasi dalam konteks data sosial media berbahasa Indonesia. Hasil penelitian diharapkan dapat menjadi masukan bagi pemerintah dan pemangku kebijakan dalam mengevaluasi implementasi kebijakan pendidikan secara lebih kontekstual dan berbasis data.

II. KAJIAN TEORI

A. Penelitian Terdahulu

Beberapa studi sebelumnya telah mengkaji analisis sentimen di media sosial, termasuk yang berfokus pada Kurikulum Merdeka serta penerapan berbagai algoritma klasifikasi. Studi pada [9] menganalisis sentimen terhadap implementasi Kurikulum Merdeka di Indonesia menggunakan algoritma *K-Nearest Neighbor* (K-NN) dan teknik *Forward Selection* (FS). Hasil penelitian tersebut menunjukkan bahwa sentimen negatif lebih dominan dibandingkan sentimen positif, dengan peningkatan akurasi dari 73,64% menjadi 76,82% setelah seleksi fitur.

Studi pada [10] menggunakan algoritma *Multinomial Naïve Bayes* untuk mengklasifikasikan sentimen publik terhadap Kurikulum Merdeka berdasarkan data tweet. Meskipun dataset yang digunakan terbatas, model tersebut mencapai akurasi sebesar 96% dengan presisi 100% dan recall 95%, menekankan pentingnya tahap preprocessing dalam meningkatkan performa klasifikasi.

Sementara itu, studi pada [11] membandingkan performa *Support Vector Machine* (SVM) dan *Random Forest* (RF) dalam menganalisis sentimen terhadap program vaksinasi menggunakan data dari platform Twitter. Studi tersebut menerapkan teknik SMOTE untuk mengatasi ketidakseimbangan kelas serta melakukan *hyperparameter tuning* guna meningkatkan akurasi model. Hasil evaluasi menunjukkan bahwa SVM memberikan performa lebih unggul dengan *F1-score* sebesar 0,52 dibandingkan RF. Temuan ini menggarisbawahi pentingnya penggunaan SMOTE dan tuning dalam meningkatkan performa model pada data yang tidak seimbang, sebagaimana juga umum terjadi dalam klasifikasi sentimen publik.

B. Kurikulum Merdeka

Kurikulum Merdeka diluncurkan pada tahun 2022 oleh Kementerian Pendidikan, Kebudayaan, Riset, dan Teknologi (Kemendikbudristek) sebagai pengganti Kurikulum 2013 (K13) [12]. Kebijakan ini bertujuan untuk membekali peserta didik dengan kompetensi yang relevan dalam menghadapi perubahan zaman, serta memberikan fleksibilitas bagi satuan pendidikan dalam menyusun pembelajaran yang sesuai dengan karakteristik peserta didik [13].

Hingga 27 Maret 2024, Kurikulum Merdeka telah diimplementasikan di lebih dari 300 ribu satuan pendidikan di seluruh Indonesia [14]. Implementasi ini menuntut kesiapan dan pemahaman mendalam dari para guru, khususnya dalam merancang pembelajaran yang berpusat pada siswa, menggunakan pendekatan tematik integratif, serta mengembangkan modul pembelajaran yang sesuai [15]. Meskipun demikian, proses penerapannya masih menghadapi berbagai tantangan yang memerlukan evaluasi dan perbaikan berkelanjutan.

C. Media Sosial X

Media sosial merupakan platform digital berbasis daring yang memungkinkan pengguna untuk berinteraksi, berbagi informasi, dan membangun jaringan sosial baik secara individu maupun kelompok secara efektif dan efisien. Jumlah

pengguna aktif media sosial secara global tercatat sebanyak 5,22 miliar per Oktober 2024, mencerminkan proporsi besar terhadap total populasi dunia pada periode tersebut [16].

Salah satu platform media sosial yang digunakan dalam penelitian ini adalah X (sebelumnya dikenal sebagai Twitter). Pada Oktober 2022, platform ini diakuisisi oleh Elon Musk dengan nilai transaksi sebesar 44 miliar dolar [17]. Perubahan nama dan logo dari Twitter menjadi X diumumkan pada 23 Juli 2023, disertai pengalihan domain situs dari “twitter.com” ke “x.com” [17]. Platform X digunakan secara luas untuk berbagi informasi, menyampaikan opini, dan berdiskusi mengenai isu-isu publik. Pada Oktober 2024, jumlah pengguna aktif bulanan platform X mencapai 590 juta, menempatkannya di posisi ke-12 sebagai media sosial dengan pengguna aktif terbanyak secara global, setelah Kuaishou [16].

D. Analisis Sentimen

Analisis sentimen, atau dikenal juga sebagai opinion mining, merupakan metode komputasi yang digunakan untuk mengidentifikasi pendapat, sikap, dan emosi seseorang terhadap suatu entitas melalui ulasan berbasis teks [5]. Teknik ini mengandalkan pendekatan *Natural Language Processing* (NLP) untuk mengekstraksi, mengubah, serta menafsirkan opini, yang kemudian diklasifikasikan ke dalam kategori sentimen positif, negatif, atau netral [18]. Tujuan utama dari analisis sentimen adalah untuk memahami pandangan publik terhadap suatu isu, layanan, produk, atau entitas tertentu dengan lebih sistematis [19]. Metode ini banyak diterapkan pada berbagai sumber data berbasis teks seperti ulasan produk, komentar media sosial, artikel berita, dan forum diskusi daring. Beberapa pendekatan umum yang digunakan dalam analisis sentimen mencakup algoritma berbasis machine learning, seperti *Naive Bayes*, *Support Vector Machine* (SVM), dan *Random Forest*, serta pendekatan berbasis leksikon.

E. SMOTE dan TF-IDF

Synthetic Minority Oversampling Technique (SMOTE) merupakan metode yang digunakan untuk mengatasi ketidakseimbangan kelas (*imbalanced data*) dengan cara menghasilkan data sintesis pada kelas minoritas sehingga distribusi kelas menjadi seimbang [20]. Teknik ini hanya diterapkan pada data latih dan bekerja dengan membuat sampel sintesis berdasarkan kedekatan jarak menggunakan algoritma *K-Nearest Neighbor* (KNN) [20].

Sementara itu, *Term Frequency-Inverse Document Frequency* (TF-IDF) adalah metode pembobotan kata dalam representasi teks umum digunakan pada teks mining dan analisis sentimen. Nilai bobot dihitung berdasarkan frekuensi kemunculan suatu kata dalam sebuah dokumen (TF) dan seberapa jarang kata tersebut muncul di keseluruhan korpus (IDF) [21]. Nilai akhir bobot TF-IDF untuk kata t pada dokumen d dirumuskan sebagai:

$$TF(t, d) = \frac{n_{i,j}}{\sum_k n_{k,j}} \quad (1)$$

$$IDF(t) = \log \left(\frac{N}{df_t} \right) \quad (2)$$

$$w_{t,d} = TF(t, d) \times IDF(t) \quad (3)$$

Keterangan:

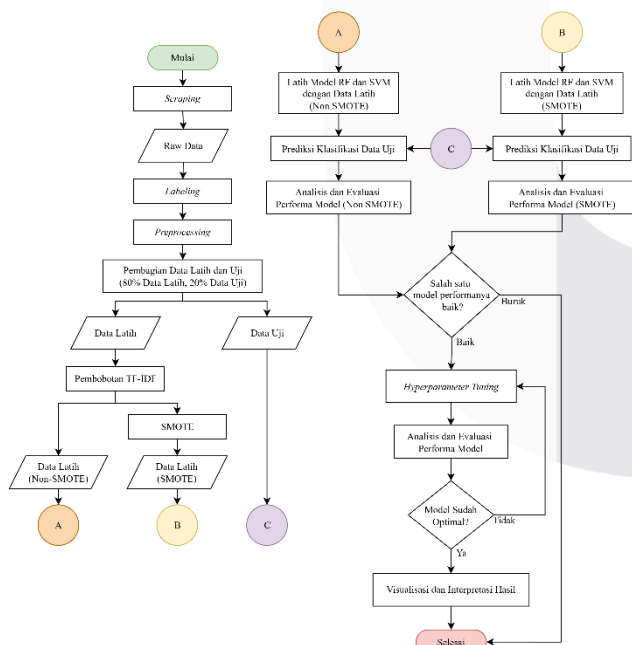
- $n_{i,j}$: Frekuensi kata i dalam dokumen j
 $\sum_k n_{k,j}$: Total kata dalam dokumen j
 N : Jumlah total dokumen
 df_t : Jumlah dokumen mengandung kata t
 $w_{t,d}$: Bobot akhir kata t dalam dokumen d

F. Random Forest dan Support Vector Machine

Algoritma *Random Forest* (RF) merupakan metode *ensemble learning* yang membangun banyak pohon keputusan (*decision tree*) dengan pendekatan *bootstrap aggregating* (bagging). Setiap pohon dilatih secara acak dan hasil prediksi akhir diperoleh melalui teknik *majority voting* untuk klasifikasi atau *average* untuk regresi [7]. RF dikenal efektif dalam menangani data berdimensi tinggi, tahan terhadap *outlier*, serta mampu mengestimasi variabel penting dalam model [22].

Sementara itu, *Support Vector Machine* (SVM) adalah algoritma *supervised learning* yang bekerja dengan membentuk sebuah *hyperplane* untuk memisahkan kelas dalam ruang fitur. Tujuannya adalah menemukan *hyperplane* terbaik dengan margin maksimum antara kelas yang berbeda. Titik-titik terdekat dari setiap kelas yang digunakan untuk menentukan margin dikenal sebagai *support vectors* [5][7].

III. METODE



GAMBAR 1
(Flowchart Alur Penelitian)

A. Dataset dan Sumber Data

Data dalam penelitian ini diperoleh dari platform media sosial X dengan menggunakan proses *data scraping*. Kata kunci "Kurikulum Merdeka" digunakan untuk memastikan

relevansi data dengan fokus penelitian. Selain itu, rentang waktu pengambilan data ditetapkan pada periode 20 Oktober 2023 hingga 20 Oktober 2024, dengan pertimbangan kesesuaian terhadap masa jabatan pemerintahan dan menteri terkait.

Untuk pengambilan data, digunakan *Tweet Harvest*, sebuah alat yang dipilih karena kemampuannya dalam melakukan scraping hingga 500 entri per jam serta kemudahan dalam pengoperasian. Setelah data berhasil dikumpulkan, proses pelabelan dilakukan secara manual oleh peneliti ke dalam tiga kategori sentimen: positif, netral, dan negatif.

B. Data Preprocessing

Setelah proses pengumpulan dan pelabelan data selesai, tahap selanjutnya adalah *data preprocessing*, yang bertujuan untuk membersihkan serta mempersiapkan data agar siap untuk dianalisis. Tahapan-tahapan *preprocessing* yang dilakukan dalam penelitian ini meliputi:

1. *Cleaning*
Menghapus elemen tidak relevan seperti URL, username, dan simbol khusus yang tidak berkontribusi pada analisis sentimen.
2. *Case Folding*
Mengonversi seluruh teks menjadi huruf kecil untuk menyamakan representasi dan menghindari duplikasi makna.
3. *Tokenizing*
Memecah kalimat menjadi token kata berdasarkan spasi atau tanda baca untuk mempermudah analisis kata per kata.
4. *Filtering*
Menghapus kata-kata umum (*stopwords*) yang tidak bermakna secara semantik dalam analisis sentimen karena tidak memberikan kontribusi informasi yang signifikan.
5. *Stemming* atau *Lemmatization*
Mengubah kata ke bentuk dasarnya (*stemming*) dengan menghilangkan imbuhan, agar variasi kata yang berasal dari akar yang sama dapat disatukan dalam analisis.
6. *Normalization*
Menstandarkan penulisan kata-kata yang memiliki arti sama namun ditulis dengan ejaan yang bervariasi, termasuk dalam bentuk slang atau singkatan.

C. Pembobotan TF-IDF

Term Frequency-Inverse Document Frequency (TF-IDF) merupakan teknik pembobotan kata yang digunakan untuk merepresentasikan teks dalam bentuk vektor numerik berdasarkan frekuensi kemunculan kata pada sebuah dokumen dan seluruh korpus. Teknik ini bertujuan untuk menekankan kata-kata penting yang bersifat spesifik terhadap suatu dokumen, serta mengurangi bobot kata-kata yang terlalu umum.

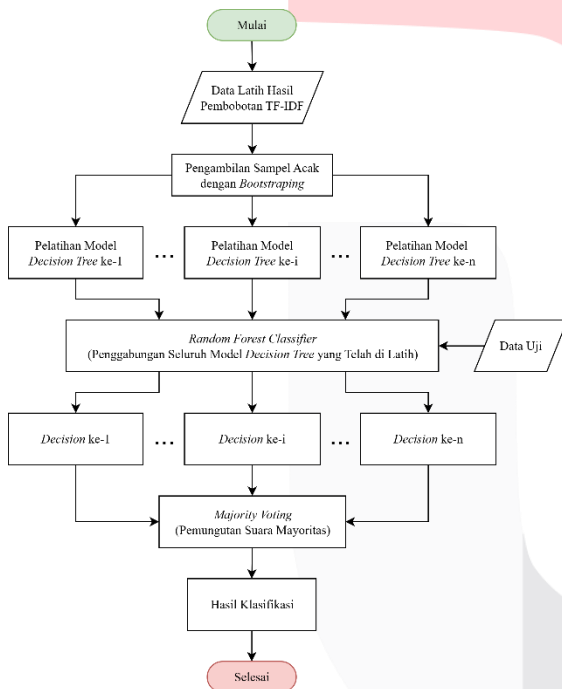
Setelah data melalui tahapan preprocessing, setiap kata dalam data latih diberi bobot menggunakan metode TF-IDF. Proses ini dilakukan dengan memanfaatkan pustaka *TfidfVectorizer* dari *scikit-learn* dalam bahasa pemrograman Python.

D. Data Balancing

Penyeimbangan data dilakukan untuk mengatasi ketimpangan jumlah sampel antar kelas, terutama saat kelas netral jauh lebih dominan dibanding positif atau negatif. Penelitian ini menggunakan *Synthetic Minority Over-sampling Technique* (SMOTE) yang menambahkan sampel sintetis pada kelas minoritas berdasarkan interpolasi data yang ada, tanpa menyamakan jumlah kelas secara absolut.

SMOTE efektif untuk data berdimensi tinggi seperti TF-IDF dan mengurangi risiko *overfitting* dibanding metode *oversampling* biasa. Meskipun algoritma seperti *Random Forest* cukup tangguh terhadap ketidakseimbangan, penyeimbangan tetap diperlukan untuk meningkatkan sensitivitas model terhadap kelas minoritas. SMOTE diterapkan setelah TF-IDF dan sebelum pelatihan model RF dan SVM.

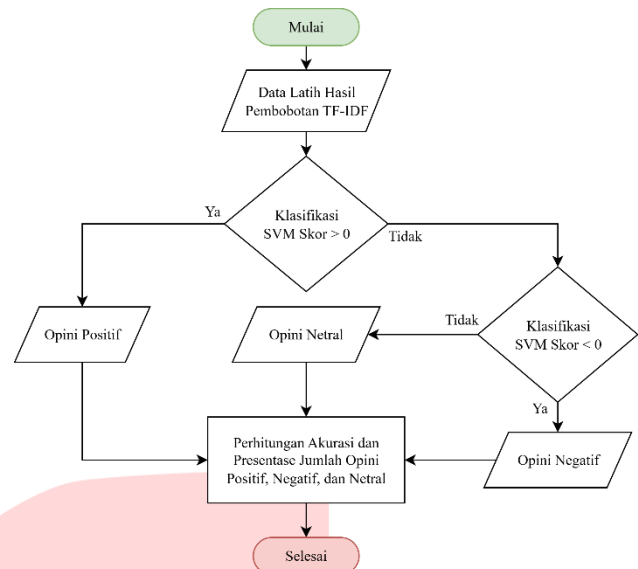
E. Implementasi Algoritma RF dan SVM



GAMBAR 2
(Flowchart Proses Klasifikasi dengan Algoritma RF)

Algoritma klasifikasi pertama yang digunakan adalah *Random Forest* (RF), yaitu algoritma *ensemble learning* yang membangun sejumlah *decision tree* dari subset data latih dan menghasilkan prediksi akhir melalui mekanisme *majority voting*. Setiap subset mencakup data dari kelas positif, negatif, dan netral.

RF efektif untuk data berdimensi tinggi seperti TF-IDF dan mampu mengurangi *overfitting* dengan menggabungkan banyak pohon keputusan. Selain memberikan prediksi yang stabil, RF juga menawarkan *feature importance*, yang berguna dalam mengidentifikasi fitur paling relevan, sebagaimana ditunjukkan dalam studi oleh Susanto et al. [10] dan Darmawan [9].



GAMBAR 3
(Flowchart Proses Klasifikasi dengan Algoritma SVM)

Support Vector Machine (SVM) adalah algoritma klasifikasi yang bekerja dengan membangun *hyperplane* untuk memisahkan kelas, seperti sentimen positif, negatif, dan netral. Tujuan utamanya adalah memaksimalkan margin antar kelas. SVM efektif digunakan dalam klasifikasi, deteksi outlier, dan regresi, terutama saat menangani data berdimensi tinggi seperti hasil ekstraksi TF-IDF. Penelitian ini menguji beberapa jenis kernel, yaitu linear, polynomial, RBF, dan sigmoid, untuk memperoleh performa terbaik.

SVM dikenal unggul dalam memisahkan kelas secara optimal, bahkan dalam kondisi data kompleks. Studi sebelumnya [10][9] menunjukkan bahwa SVM sering menghasilkan akurasi lebih tinggi dibandingkan algoritma seperti K-NN, terutama ketika distribusi data cukup seimbang. Kemampuan kernel SVM dalam mengenali pola non-linear menjadikannya salah satu metode yang efektif untuk analisis sentimen.

F. Metrik Evaluasi Performa Model

Evaluasi performa model dilakukan untuk mengukur kinerja algoritma yang digunakan dalam klasifikasi sentimen. Pada penelitian ini, metode evaluasi yang digunakan adalah *confusion matrix*, dengan empat metrik utama sebagai berikut:

1. *Accuracy*
Mengukur seberapa sering model memprediksi dengan benar dari seluruh data yang diuji.
2. *F1-Score*
Merupakan rata-rata harmonis dari *precision* dan *recall*, untuk menilai keseimbangan antara ketepatan dan kelengkapan prediksi positif.
3. *Recall*
Menilai seberapa baik model mengenali seluruh data yang termasuk dalam suatu kelas, khususnya data positif.
4. *Precision*

Menunjukkan seberapa akurat prediksi positif model dengan menghitung proporsi yang benar-benar relevan.

G. Visualisasi dan Interpretasi Hasil

Hasil dari analisis sentimen serta evaluasi kinerja algoritma *Random Forest* (RF) dan *Support Vector Machine* (SVM) disajikan dalam bentuk visualisasi grafik dan tabel untuk memudahkan interpretasi. Beberapa teknik visualisasi yang digunakan antara lain:

1. Word Cloud

Menampilkan kata-kata paling sering muncul, dengan ukuran kata mencerminkan frekuensinya. Visualisasi ini membantu mengidentifikasi tema dominan dalam opini publik.

2. Time Series Plot

Menggambarkan perubahan distribusi sentimen dari waktu ke waktu untuk menganalisis dinamika opini publik selama periode tertentu.

3. Pie Chart dan Bar Chart

Menyajikan proporsi kelas sentimen secara visual untuk memudahkan pemahaman sebaran data secara keseluruhan.

4. Tabel Perbandingan

Membandingkan performa RF dan SVM berdasarkan metrik evaluasi sebagai dasar pemilihan algoritma klasifikasi terbaik.

IV. HASIL DAN PEMBAHASAN

A. Persiapan Data

Penelitian ini menggunakan 10.066 data tweet dari platform X, yang mencakup 15 atribut berisi konten dan metadata. Atribut penting meliputi `full_text` (isi tweet), `created_at`, `favorite_count`, `retweet_count`, `reply_count`, serta `conversation_id_str` sebagai ID percakapan. Data tambahan seperti `username`, `lang`, `location`, dan `image_url` (jika ada) turut disertakan. Tabel 1 menampilkan cuplikan data mentah hasil scraping yang merepresentasikan struktur dataset yang dianalisis.

TABEL 1
(Sampel Data dari Platform X Hasil *Scraping*)

conv_id_str	created_at	fav_count	...	username
1.85E+18	Sun Oct 20 17:47	0	...	AhmadGa16937372
1.85E+18	Sun Oct 20 17:32	0	...	renxijun
1.85E+18	Sun Oct 20 17:05	0	...	koko_busa
1.85E+18	Sun Oct 20 17:02	2	...	AshbornDeSoul
1.85E+18	Sun Oct 20 16:57	0	...	insomneyac
...	...	...	...	...
1.72E+18	Fri Oct 20 01:07	2	...	tempodotco

Seluruh data diberi label secara manual ke dalam kategori positif, negatif, atau netral berdasarkan isi kolom `full_text`. Label disimpan dalam kolom label dan digunakan sebagai acuan dalam pelatihan dan evaluasi model. Setelah pelabelan, data diproses melalui tahap *preprocessing* untuk

membersihkan teks sebelum digunakan oleh algoritma. Contoh data hasil pelabelan ditampilkan pada Tabel 2.

TABEL 2
(Sampel Data Hasil *Manual Labeling*)

No	Teks	Label
1.	@CakKhum @nadiemmakarim Kelas X?? Loh model begini kok bisa lulus SD?? Zaman dia SD kan belum kurikulum merdeka juga..	netral
2.	stuju kurikulum merdeka ngabisin duit	negatif
3.	@bangonae @txtdrjkt Merdeka. Alhamdulillah semenjak kurikulum merdeka anak sekolah gak punya beban berat untuk ngafalin.	positif
4.	@alwaysanehhh nah rata tuh tapi kurikulum merdeka ga ada ipa ips	netral
5.	@TirtoID Mending kurikulum merdeka dihapuskan saja.	negatif
6.	@CakKhum @nadiemmakarim pak @nadiemmakarim halooooo bagus bgt pak hasil kurikulum merdeka	positif

B. Data Preprocessing

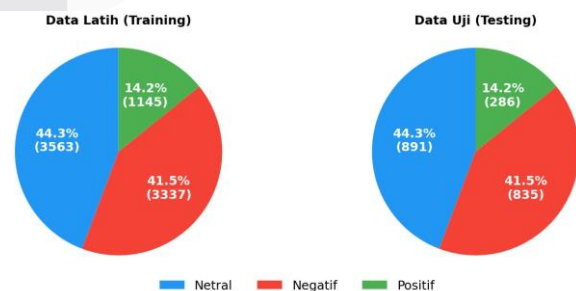
Proses *preprocessing* mencakup ekstraksi waktu, penghapusan duplikasi, normalisasi dan pembersihan teks, penghapusan *stopword*, serta *tokenization* dan *stemming* untuk menyederhanakan kata. Tahap akhir berupa *label encoding* dilakukan agar label sentimen dapat dikenali oleh model. Hasil dari proses *preprocessing* ditampilkan pada Tabel 3 berikut ini.

TABEL 3
(Sampel Data Hasil *Preprocessing*)

No	Teks
1.	kelas x model lulus sd zaman sd kurikulum merdeka
2.	sepatutnya kurikulum merdeka habis uang
3.	merdeka alhamdulillah semenjak kurikulum merdeka anak sekolah beban berat hafal
4.	kurikulum merdeka ipa ips
5.	kurikulum merdeka hapus
6.	bagus kurikulum merdeka

C. Pembagian Data Latih dan Uji

Pada penelitian ini, sebanyak 10.057 data yang telah melalui tahap preprocessing dibagi menjadi data latih dan data uji dengan rasio 80:20. Data latih terdiri dari 3.337 data negatif, 3.563 netral, dan 1.145 positif, sedangkan data uji terdiri atas 835 data negatif, 891 netral, dan 286 positif.



GAMBAR 4
(Visualisasi Proporsi Pembagian Data Latih dan Uji)

D. Ekstraksi Fitur dengan TF-IDF

Ekstraksi fitur dilakukan menggunakan metode TF-IDF untuk mengubah data teks menjadi bentuk numerik berbobot. Proses ini membantu mengenali kata-kata penting dalam dokumen. Beberapa penyesuaian juga diterapkan untuk menyaring kata yang kurang relevan, menstandarkan teks, dan menjaga konsistensi panjang vektor fitur.

TABEL 4
(Informasi Transformasi TF-IDF Sebelum Pengaturan Fitur)

Jumlah fitur TF-IDF	8.741
Jumlah dokumen training	8.045
Jumlah dokumen testing	2.012
Sparsity data training	99,89%
Sparsity data testing	99,90%

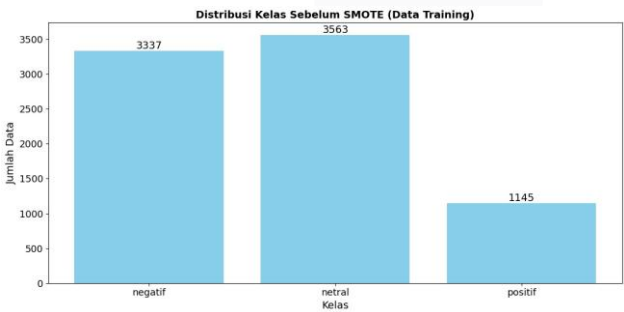
TABEL 5
(Informasi Transformasi TF-IDF Setelah Pengaturan Fitur)

Jumlah fitur TF-IDF	2.384
Jumlah dokumen training	8.045
Jumlah dokumen testing	2.012
Sparsity data training	99,73%
Sparsity data testing	99,74%

Tabel 5 menunjukkan bahwa proses penyaringan berhasil menurunkan jumlah fitur secara signifikan, serta mengurangi *sparsity* pada data pelatihan dan pengujian. Hal ini menghasilkan representasi fitur yang lebih efisien dan mendukung kinerja klasifikasi yang lebih baik.

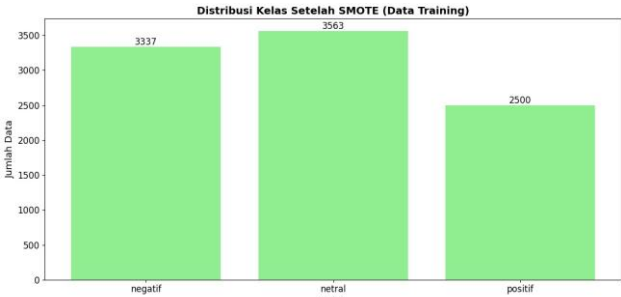
E. Penerapan SMOTE

Data latih memiliki distribusi kelas yang tidak seimbang, terutama pada kelas positif. Untuk mengatasi hal ini, digunakan metode SMOTE guna menyeimbangkan data. Model SVM dan RF kemudian dilatih pada data asli dan data yang telah diseimbangkan untuk dibandingkan performanya.



GAMBAR 5
(Visualisasi Distribusi Kelas Sebelum Penerapan SMOTE)

Gambar 5 memperlihatkan ketimpangan kelas, di mana data positif jauh lebih sedikit. Ketidakseimbangan ini dapat menyebabkan model bias. Untuk mengatasinya, diterapkan SMOTE guna menambah data sintesis pada kelas positif agar distribusi lebih seimbang.



GAMBAR 6
(Visualisasi Distribusi Kelas Setelah Penerapan SMOTE)

Setelah SMOTE diterapkan, jumlah data kelas positif meningkat menjadi 2.500, sementara kelas lainnya tetap. Penambahan ini membantu menyeimbangkan distribusi data. Kombinasi dengan *class_weight* saat pelatihan juga dilakukan agar model memperhatikan semua kelas secara adil.

F. Evaluasi dan Perbandingan Model

Subbab ini menampilkan hasil evaluasi model *Random Forest* dan SVM, baik pada data asli (*imbalanced*) maupun data yang telah diseimbangkan dengan SMOTE. Perbandingan dilakukan menggunakan metrik akurasi, *precision*, *recall*, dan *F1-score*.

TABEL 6
(Perbandingan Performa Model Sebelum dan Sesudah SMOTE)

Model	Training (10-Fold CV)				Testing
	Accuracy	Precision	Recall	F1-Macro	F1-Macro
RF (tanpa SMOTE)	68.69%	71.68%	62.51 %	64.95 %	63.24%
RF + SMOTE	74.69%	75.57%	75.67 %	75.57 %	65.23%
SVM (tanpa SMOTE)	71.61%	75.33%	65.84 %	68.51 %	66.27%
SVM + SMOTE	74.46%	77.16%	74.37 %	75.35 %	68.77%

Tabel 6 menunjukkan bahwa penerapan SMOTE memberikan peningkatan performa pada kedua model, dengan SVM+SMOTE mencatatkan hasil terbaik di sebagian besar metrik, baik pada data latih maupun uji. Hal ini menunjukkan bahwa SVM mampu menangani ketidakseimbangan kelas secara lebih efektif. Untuk menyoroti kemampuan model dalam mengenali kelas minoritas, Tabel 7 membandingkan nilai *recall* pada kelas positif dari masing-masing model.

TABEL 7
(Perbandingan Nilai *Recall* pada Kelas Positif)

Model	Recall (Kelas Positif)
RF (tanpa SMOTE)	0.42
RF + SMOTE	0.56
SVM (tanpa SMOTE)	0.43
SVM + SMOTE	0.51

Penerapan SMOTE terbukti meningkatkan sensitivitas model terhadap kelas positif. Meskipun RF+SMOTE mencatat *recall* tertinggi pada kelas positif, SVM+SMOTE tetap unggul secara keseluruhan karena menghasilkan metrik yang lebih seimbang, terutama pada *F1-score macro* di data uji.

Selanjutnya, dilakukan *hyperparameter tuning* dalam tiga tahap untuk masing-masing model guna mengoptimalkan performa. Tabel 8 merangkum hasil akhir dari enam model, mencakup nilai validasi silang, akurasi, *F1-score* pada data uji, serta selisih performa antara validasi dan pengujian.

TABEL 8
(Ringkasan Performa Model Tuning)

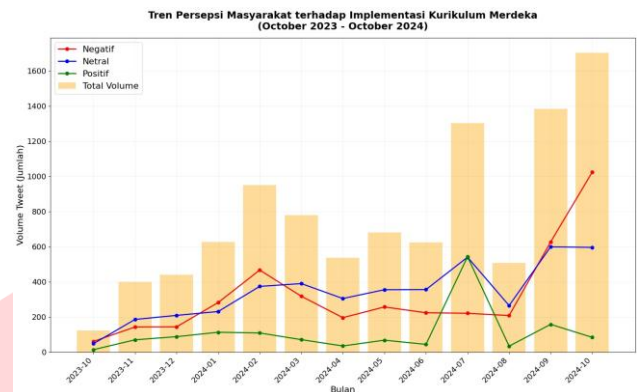
Model	Test Accuracy	Test F1-Macro	F1 Gap (CV-Test)
RF Tuning-1	0.6890	0.6600	0.0613 (6.13%)
SVM Tuning-1	0.7177	0.6907	0.0956 (9.56%)
RF Tuning-2	0.6859	0.6618	0.0640 (6.40%)
SVM Tuning-2	0.7226	0.6947	0.0913 (9.13%)
RF Tuning-3	0.6849	0.6625	0.0635 (6.35%)
SVM Tuning-3	0.7422	0.6967	0.0895 (8.95%)

Tabel 8 menunjukkan bahwa model SVM mencatatkan performa paling unggul, baik pada validasi silang maupun pengujian. SVM Tuning-3 menjadi konfigurasi terbaik dengan akurasi 74,22% dan *F1-score macro* sebesar 69,67% pada data uji. Meskipun terdapat selisih *F1-score* sebesar 8,95%, model ini tetap menunjukkan kemampuan generalisasi yang stabil. Penggunaan SMOTE juga terbukti efektif dalam meningkatkan sensitivitas terhadap kelas minoritas tanpa mengurangi akurasi secara keseluruhan.

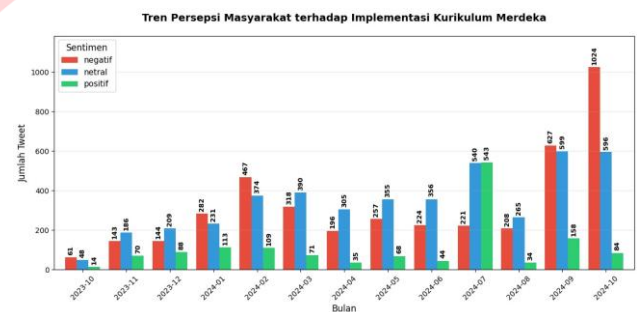
Secara keseluruhan, hasil evaluasi menunjukkan bahwa SVM Tuning-3 dengan SMOTE merupakan model dengan performa terbaik. Model ini mampu menghasilkan klasifikasi yang paling seimbang, menunjukkan sensitivitas tinggi terhadap kelas minoritas, serta memiliki metrik makro yang lebih baik dibandingkan model lainnya. Oleh karena itu, model ini dipilih sebagai model akhir dalam penelitian.

G. Visualisasi Tren Sentimen Publik

Subbab ini menyajikan hasil visualisasi sentimen terhadap opini publik mengenai Kurikulum Merdeka menggunakan data dari media sosial X. Tujuannya adalah mengidentifikasi kecenderungan sentimen positif, negatif, atau netral terhadap implementasi kebijakan tersebut.



GAMBAR 7
(Visualisasi Perkembangan Tren dari setiap Sentimen)



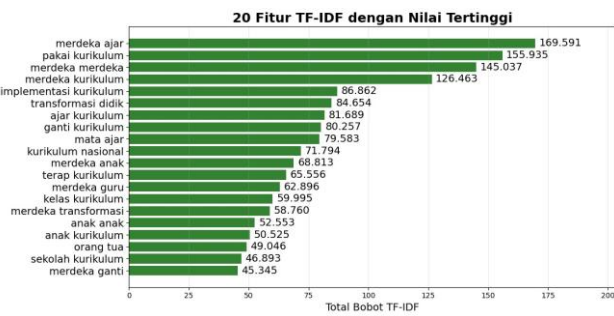
GAMBAR 8
(Visualisasi Jumlah Tren Sentimen (Bulanan))

Gambar 7 dan 8 menunjukkan fluktuasi sentimen publik terhadap Kurikulum Merdeka dari Oktober 2023 hingga Oktober 2024. Sentimen negatif melonjak pada Februari dan kembali meningkat pada September–Oktober akibat kritik terhadap implementasi kurikulum. Sebaliknya, sentimen positif sempat dominan pada Juli, dipicu unggahan apresiasi masyarakat. Pola ini mencerminkan bahwa persepsi publik bersifat dinamis dan dipengaruhi oleh pelaksanaan kebijakan.



GAMBAR 9
(Visualisasi *Word Cloud* Sebaran Kata Keseluruhan Sentimen)

Gambar 9 menunjukkan word cloud sentimen publik terhadap Kurikulum Merdeka, dengan kata dominan seperti “merdeka”, “guru”, “anak”, dan “sekolah” yang mencerminkan fokus pada pelaksana dan penerima kebijakan. Kata-kata bernuansa opini seperti “buruk”, “bingung”, “kreatif”, dan “bagus” menandakan adanya persepsi yang beragam. Visualisasi ini memperlihatkan bahwa topik Kurikulum Merdeka cukup menyita perhatian publik dan memunculkan tanggapan yang bervariasi terhadap pelaksanaannya.



GAMBAR 10
(Visualisasi 20 Fitur TF-IDF Tertinggi Seluruh Sentimen)



GAMBAR 11
(Visualisasi 20 Fitur TF-IDF Terendah Seluruh Sentimen)

Gambar 10 dan 11 menampilkan fitur dengan bobot TF-IDF tertinggi dan terendah. Fitur seperti “merdeka ajar” dan “implementasi kurikulum” mencerminkan topik spesifik Kurikulum Merdeka, sementara frasa seperti “ajar nilai” atau “semangat ajar” bersifat umum dan kurang informatif. Secara keseluruhan, fitur-fitur ini menunjukkan narasi faktual tanpa kecenderungan nilai. Frasa teknis berbobot tinggi banyak ditemukan dalam wacana netral yang menjelaskan praktik di lapangan tanpa ekspresi dukungan atau penolakan.

Tren sentimen terhadap Kurikulum Merdeka menunjukkan pola yang dinamis, dengan dominasi sentimen netral dan negatif, terutama saat awal implementasi atau muncul isu tertentu. Sentimen positif hanya meningkat pada momen kebijakan baru atau kegiatan pemerintah. Hal ini menunjukkan bahwa opini publik sangat dipengaruhi oleh konteks waktu dan situasi yang berkembang.

V. KESIMPULAN

Penelitian ini mengevaluasi sentimen publik terhadap implementasi Kurikulum Merdeka menggunakan algoritma *Random Forest* (RF) dan *Support Vector Machine* (SVM). Mayoritas sentimen yang diidentifikasi bersifat netral

(44,3%) dan negatif (41,5%), sementara sentimen positif hanya mencakup sebagian kecil. Dominasi opini netral dan negatif mencerminkan adanya kekritisian publik terhadap kebijakan tersebut, terutama terkait kesiapan implementasi dan beban guru.

Analisis temporal menunjukkan fluktuasi sentimen selama periode observasi, dengan lonjakan sentimen negatif pada awal tahun dan peningkatan positif sementara pada pertengahan tahun. Visualisasi TF-IDF dan word cloud memperkuat temuan ini melalui kata-kata dominan yang mencerminkan respons publik.

Model SVM menunjukkan performa terbaik dengan F1-macro sebesar 69,67% dan akurasi 72,42%, mengungguli RF. Penerapan SMOTE berhasil meningkatkan akurasi terhadap kelas minoritas, khususnya sentimen positif. Hasil ini mengindikasikan bahwa kombinasi teknik *preprocessing*, penyeimbangan data, dan tuning parameter memiliki pengaruh signifikan terhadap peningkatan kinerja model dalam analisis sentimen.

REFERENSI

- [1] A. C. Widyaningrum and S. Suparni, “Inovasi Pembelajaran Matematika Dengan Model Discovery Learning Pada Kurikulum Merdeka,” *Sepren*, vol. 4, no. 02, pp. 186–193, May 2023, doi: 10.36655/SEPREN.V4I02.887.
- [2] D. L. Anggraini, M. Yulianti, S. Faizah, A. Putri, and B. Pandiangan, “PERAN GURU DALAM MENGEMBANGAN KURIKULUM MERDEKA,” *Jurnal Ilmu Pendidikan dan Sosial*, vol. 1, no. 3, pp. 290–298, Dec. 2022, doi: 10.58540/JIPSI.V1I3.53.
- [3] F. Fatmawati, L. Rusdi, A. Mardhiah, P. Husna, and F. Fuady, “TAHAP-TAHAP PENYUSUNAN MODUL AJAR KURIKULUM MERDEKA TINGKAT SEKOLAH,” *COVIT (Community Service of Tambusai)*, vol. 2, no. 2, pp. 308–313, Sep. 2022, doi: 10.31004/COVIT.V2I2.10779.
- [4] T. Mulyadin, M. Khoiron, D. Ginanto, and K. A. Putra, “Workshop on Kurikulum Merdeka (Freedom Curriculum): Dismantling Theories and Practices,” *BEMAS: Jurnal Bermasyarakat*, vol. 3, no. 2, pp. 126–132, Mar. 2023, doi: 10.37373/BEMAS.V3I2.265.
- [5] E. Hokijulandy, “ANALISIS SENTIMEN MENGGUNAKAN METODE KLASIFIKASI SUPPORT VECTOR MACHINE DAN SELEKSI FITUR CHI-SQUARE (Studi Kasus: Ulasan Aplikasi Mobile JKN pada Google Play Store),” 2023, *Fakultas Matematika & IPA Universitas Padjadjaran*. Accessed: Jul. 26, 2025. [Online]. Available: <https://kandaga.unpad.ac.id/koleksi/repository/item/140110190034>
- [6] “Digital 2024 October Global Statshot Report - We Are Social Indonesia,” <https://wearesocial.com/>. Accessed: Jul. 26, 2025. [Online]. Available: <https://wearesocial.com/id/blog/2024/10/digital-2024-october-global-statshot-report/>

- [7] M. R. Adrian, M. P. Putra, M. H. Rafialdy, and N. A. Rakhmawati, "Perbandingan Metode Klasifikasi Random Forest dan SVM Pada Analisis Sentimen PSBB," *Jurnal Informatika UPGRIS*, vol. 7, no. 1, Jun. 2021, doi: 10.26877/JIU.V7I1.7099.
- [8] K. Algoritma Machine Learning Untuk Klasifikasi Kelompok Obat, A. Wirasto, K. Nisa, U. Harapan Bangsa, J. Raden Patah No, and L. Kecamatan Kembaran Kabupaten Banyumas, "Comparison of Machine Learning Algorithms for Classification of Drug Groups," *SISFOTENIKA*, vol. 11, no. 2, pp. 196–207, Jul. 2021, doi: 10.30700/JST.V11I2.1134.
- [9] W. Darmawan, M. Kurniawan Faizal, W. Setianto, and W. Hapsoro, "Analisis Sentimen Penerapan Kurikulum Merdeka Pada Pengguna Twitter Menggunakan Metode K-Nearest Neighbor Dengan Forward Selection," *Smart Comp Jurnalnya Orang Pintar Komputer*, vol. 12, no. 1, Jan. 2023, doi: 10.30591/SMARTCOMP.V12I1.4634.
- [10] A. A. Susanto, Painem, M. Syafrullah, and R. Pradana, "ANALISIS SENTIMEN PENERAPAN KURIKULUM MERDEKA PADA TWITTER DENGAN METODE NAÏVE BAYES," *Prosiding Seminar Nasional Mahasiswa Fakultas Teknologi Informasi (SENAFTI)*, vol. 2, no. 1, pp. 227–234, 2023, Accessed: Jul. 26, 2025. [Online]. Available: <https://senafiti.budiluhur.ac.id/senafiti/article/view/576>
- [11] I. M. Sumertajaya, Y. Angraini, J. R. Harahap, and A. Fitrianto, "Sentiment Analysis on Covid-19 Vaccination in Indonesia Using Support Vector Machine and Random Forest," *JUITA: Jurnal Informatika*, vol. 10, no. 1, pp. 1–8, May 2022, doi: 10.30595/JUITA.V10I1.12394.
- [12] "Luncurkan Kurikulum Merdeka, Mendikbudristek: Ini Lebih Fleksibel!," <https://ditpsd.kemdikbud.go.id/>. Accessed: Jul. 26, 2025. [Online]. Available: <https://ditpsd.kemdikbud.go.id/artikel/detail/luncurkan-kurikulum-merdeka-mendikbudristek-ini-lebih-fleksibel>
- [13] F. Fatmawati, L. Rusdi, A. Mardhiah, P. Husna, and F. Fuady, "TAHAP-TAHAP PENYUSUNAN MODUL AJAR KURIKULUM MERDEKA TINGKAT SEKOLAH," *COVIT (Community Service of Tambusai)*, vol. 2, no. 2, pp. 308–313, Sep. 2022, doi: 10.31004/COVIT.V2I2.10779.
- [14] Kemendikbud, "Kemendikbudristek Terbitkan Payung Hukum bagi Implementasi Kurikulum Merdeka Secara Nasional," <https://www.kemdikbud.go.id/>. Accessed: Jul. 26, 2025. [Online]. Available: <https://www.kemdikbud.go.id/main/index.php/blog/2024/03/kemendikbudristek-terbitkan-payung-hukum-bagi-implementasi-kurikulum-merdeka-secara-nasional>
- [15] I. Kurikulum Merdeka Belajar di Sekolah Penggerak Restu Rahayu, R. Rosita, Y. Sri Rahayuningsih, and A. Herry Hernawan, "Implementasi Kurikulum Merdeka Belajar di Sekolah Penggerak," *Jurnal Basicedu*, vol. 6, no. 4, pp. 6313–6319, May 2022, doi: 10.31004/BASICEDU.V6I4.3237.
- [16] S. Kemp, "Digital 2024 October Global Statshot Report - We Are Social Indonesia," <https://wearesocial.com/>. Accessed: Jul. 26, 2025. [Online]. Available: <https://wearesocial.com/id/blog/2024/10/digital-2024-october-global-statshot-report/>
- [17] M. Andrivet, "Twitter's Extreme Rebrand to X: A Calculated Risk or Pure Chaos? - The Branding Journal," <https://www.thebrandingjournal.com/>. Accessed: Jul. 26, 2025. [Online]. Available: <https://www.thebrandingjournal.com/2023/08/twitter-rebrand-x/>
- [18] J. I. Matematika and H. Y. Wicaksono, "ANALISIS SENTIMEN DATA TWITTER MENGENAI PROGRAM VAKSINASI DI INDONESIA MENGGUNAKAN ALGORITMA BACKPROPAGATION," *MATHunesa: Jurnal Ilmiah Matematika*, vol. 10, no. 1, pp. 161–169, Apr. 2022, doi: 10.26740/MATHUNESA.V10N1.P161-169.
- [19] M. S. Syahlan *et al.*, "ANALISIS SENTIMEN TERHADAP TEMPAT WISATA DARI KOMENTAR PENGUNJUNG DENGAN MENGGUNAKAN METODE SUPPORT VECTOR MACHINE (SVM)," *Simtek: jurnal sistem informasi dan teknik komputer*, vol. 8, no. 2, pp. 315–319, Oct. 2023, doi: 10.51876/SIMTEK.V8I2.281.
- [20] A. A. Arifiyanti, A. A. Arifiyanti, and E. D. Wahyuni, "SMOTE: METODE PENYEIMBANG KELAS PADA KLASIFIKASI DATA MINING," *Scan : Jurnal Teknologi Informasi dan Komunikasi*, vol. 15, no. 1, pp. 34–39, Feb. 2020, doi: 10.33005/scan.v15i1.1850.
- [21] A. Kautsar and M. Syafrullah, "Implementasi Algoritme Multinomial Naïve Bayes Pada Analisis Sentimen Terhadap Isu Presiden 3 Periode," Sep. 30, 2022. Accessed: Jul. 26, 2025. [Online]. Available: <https://senafiti.budiluhur.ac.id/senafiti/article/view/77>
- [22] D. Sudrajat, A. I. Purnamasari, A. R. Dikananda, D. A. Kurnia, and A. Bahtiar, "Klasifikasi Mutu Pembelajaran Hybrid berdasarkan Algoritma C.45, Random Forest dan Naïve Bayes dengan Optimasi Bootstrap Areggating (Bagging) pada masa COVID-19," *JURIKOM (Jurnal Riset Komputer)*, vol. 9, no. 6, pp. 2227–2233, Dec. 2022, doi: 10.30865/JURIKOM.V9I6.5179.