

Optimasi Analisis Sentimen Banjir Sumatra Menggunakan Naive Bayes dan SVC Berbasis Seleksi Fitur

Sri Karnila¹, Indra Prasetyo², Aurea Ivana Chandra³, Aghitsna Salsabila⁴, ⁵Astria Hijriani

^{1,2,3,4}Fakultas Ilmu Komputer, Program Studi Sains Data, Institut Informatika dan Bisnis Darmajaya, Bandar Lampung, Indonesia

⁵Fakultas Matematika dan Ilmu Pengetahuan Alam, Program Studi Ilmu Komputer, Bandar Lampung, Indonesia

Email: ^{1,*}srikarnila_dj@darmajaya.ac.id, ²indra.2411080002@mail.darmajaya.ac.id, ³aureaich06@gmail.com, ⁴aghitsna.2411080001@mail.darmajaya.ac.id, ⁵astria.hijriani@fmipa.unila.ac.id

^{*}) Email Penulis Utama

Abstrak—Bencana alam banjir bandang yang melanda berbagai wilayah di Sumatra telah menimbulkan dampak fisik dan ekonomi yang signifikan, sekaligus memicu gelombang opini dari masyarakat luas di ruang digital. Dalam konteks mitigasi dan tanggap darurat, media sosial *twitter* beralih fungsi menjadi platform utama bagi warga untuk menyalurkan informasi terkini, keluhan, simpati, serta kritik terhadap efektivitas penanganan bencana. Mengingat besarnya volume data opini publik tersebut, diperlukan pendekatan komputasional untuk mengevaluasi respons masyarakat secara terukur guna memberikan umpan balik kepada pemerintah. Penelitian ini bertujuan untuk mengidentifikasi dan mengklasifikasikan sentimen masyarakat terhadap isu bencana banjir Sumatra sesuai dengan alur penemuan pengetahuan berbasis data teks. Metode yang digunakan bertumpu pada kerangka kerja sistematis yang mencakup pengumpulan 2.000 data komentar dari *twitter*, tahapan prapemrosesan teks (meliputi *cleansing*, *case folding*, *tokenizing*, *filter stopwords*, dan *stemming*), pelabelan leksikon, serta pembobotan fitur menggunakan *Term Frequency-Inverse Document Frequency* (TF-IDF) berbasis *N-Grams*. Data kemudian diklasifikasikan ke dalam sentimen positif dan negatif menggunakan algoritma *Multinomial Naive Bayes* dan *Support Vector Classifier* (SVC). Hasil pengujian komparatif mengungkapkan bahwa algoritma SVC dengan optimasi *class weight* berkinerja unggul dengan akurasi 87,02%, sementara algoritma *Naive Bayes* memperoleh akurasi 78,30%. Berdasarkan 447 dataset uji yang dievaluasi, model berhasil mengelompokkan ulasan mayoritas ke dalam sentimen negatif yang menyoroti keterlambatan distribusi bantuan, serta sentimen positif yang merepresentasikan apresiasi solidaritas. Kesimpulan dari penelitian ini menegaskan bahwa algoritma klasifikasi SVC, terbukti lebih akurat memproses data teks kebencanaan yang tidak seimbang. Hasil analisis data teks ini dapat dijadikan bahan evaluasi bagi pemangku kepentingan dalam pengambilan keputusan penanganan bencana alam banjir.

Kata Kunci: Analisis Sentimen, Banjir Sumatra, Naive Bayes, Support Vector Classifier, Twitter

Abstract—The flash flood natural disaster that hit various regions in Sumatra has caused significant physical and economic impacts, while also triggering a wave of public opinion in the digital space. In the context of mitigation and emergency response, Twitter has become the main platform for citizens to channel the latest information, complaints, sympathy, and criticism regarding the effectiveness of disaster management. Given the large volume of public opinion data, a computational approach is needed to evaluate community responses measurably to provide feedback to the government. This study aims to identify and classify public sentiment on the issue of the Sumatran flood disaster according to the text-based knowledge discovery workflow. The method used relies on a systematic framework which includes collecting 2,000 comment data from Twitter, text preprocessing stages (including *cleansing*, *case folding*, *tokenizing*, *stopword filtering*, and *stemming*), lexicon labeling, and feature weighting using *N-Grams* based *Term Frequency-Inverse Document Frequency* (TF-IDF). The data is then classified into positive and negative sentiments using *Multinomial Naive Bayes* and *Support Vector Classifier* (SVC) algorithms. The comparative test results reveal that the SVC algorithm with *class weight* optimization performs superiorly with an accuracy of 87.02%, while the *Naive Bayes* algorithm obtains an accuracy of 78.30%. Based on the 447 evaluated test datasets, the model successfully grouped the majority reviews into negative sentiment highlighting delays in aid distribution, and positive sentiment representing appreciation of solidarity. The conclusion of this study confirms that classification algorithms, especially SVC, are proven reliable in processing imbalanced disaster texts, which can be used as evaluation material for stakeholders.

Keywords: Sentiment Analysis, Naive Bayes, Sumatra Flood, Support Vector Classifier, Twitter

1. PENDAHULUAN

Bencana alam berupa banjir yang menimpa beberapa daerah di Sumatra tidak hanya meninggalkan kerusakan infrastruktur dan kerugian ekonomi yang mendalam bagi masyarakat terdampak, tetapi juga memicu gelombang opini publik yang sangat masif. Di era disrupsi digital saat ini, media sosial seperti *Twitter* menjadi saluran komunikasi instan dan alami bagi masyarakat luas untuk menyuarakan ragam keluhan, simpati, maupun kritik tajam terhadap proses penanganan bencana yang dilakukan oleh instansi berwenang. Memahami sentimen

masyarakat secara aktual sangat penting bagi pemerintah dan lembaga terkait untuk mengevaluasi efektivitas respons darurat, pemerataan distribusi logistik di lapangan, dan manajemen krisis komunikasi. Tanpa sebuah analisis komputasional yang komprehensif, volume data komentar yang jumlahnya mencapai ribuan per hari di media sosial akan sangat sulit untuk diukur, dipetakan, dan dimanfaatkan sebagai parameter evaluasi kebijakan.

Penelitian mengenai analisis sentimen di platform media sosial telah banyak dilakukan untuk mengukur opini publik pada berbagai isu kritis. Beberapa studi terdahulu menerapkan teknik *text mining* menggunakan algoritma Machine Learning untuk mengekstraksi dan mengklasifikasikan opini ke dalam kelas-kelas polaritas [1]. Penelitian terkait analisis data Twitter menunjukkan bahwa penerapan klasifikasi teks berbasis metode komputasional pada dataset media sosial memiliki kemampuan yang efektif dalam mengidentifikasi, memetakan, dan mendeteksi kecenderungan opini publik secara otomatis dengan tingkat akurasi yang memadai, sehingga dapat digunakan sebagai instrumen analisis untuk memahami persepsi, respons, serta dinamika komunikasi masyarakat terhadap berbagai isu yang berkembang di ruang digital [2], [3], [4]. Berbagai studi kasus, seperti analisis opini masyarakat terhadap penggunaan kendaraan listrik, kebijakan pemerintah terkait RUU TNI, isu penundaan pemilu, aktivitas marketplace di Indonesia, hingga kampanye boikot produk tertentu, menunjukkan bahwa algoritma klasifikasi memiliki kapabilitas yang baik dalam memahami pola bahasa publik, mengidentifikasi kecenderungan sentimen masyarakat, serta merepresentasikan dinamika komunikasi digital yang berkembang pada media sosial [5], [6], [7], [8]. Penerapan analisis sentimen media sosial juga telah dimanfaatkan pada berbagai sektor layanan digital dan kebijakan publik, seperti aplikasi e-health SATUSEHAT, program akademik Kampus Merdeka, serta sektor perbankan, guna menganalisis respons masyarakat, mengevaluasi kualitas layanan, dan mendukung proses pengambilan keputusan secara lebih efektif berbasis data [9], [10], [11].

Tujuan utama penelitian ini adalah untuk memetakan, mengoptimasi, serta membandingkan opini masyarakat terhadap musibah banjir di Sumatra berdasarkan data percakapan pada media sosial Twitter. Penelitian ini diharapkan mampu memberikan gambaran yang komprehensif mengenai persepsi publik terhadap penanganan bencana, efektivitas distribusi bantuan, serta respons pemerintah selama kondisi darurat berlangsung. Adapun manfaat dari penelitian ini adalah menyediakan wawasan strategis berbasis data yang terukur dan objektif bagi pemerintah maupun badan penanggulangan bencana sebagai bahan evaluasi dalam meningkatkan kualitas respons tanggap darurat, memperbaiki strategi komunikasi publik, serta mendukung pengambilan keputusan yang lebih cepat, adaptif, dan tepat sasaran.

Penggunaan algoritma *Naive Bayes* dalam penelitian ini didasarkan pada kemampuannya dalam melakukan klasifikasi teks secara efisien melalui pendekatan probabilitas independen bersyarat yang terbukti cepat, ringan, dan efektif untuk pengolahan data teks berskala besar. Metode ini memiliki kompleksitas komputasi yang relatif rendah sehingga sangat sesuai diterapkan pada data media sosial yang memiliki volume besar, tidak terstruktur, dan terus bertambah secara dinamis. Selain itu, *Naive Bayes* juga dikenal mampu memberikan performa klasifikasi yang baik pada kasus analisis sentimen karena efektif dalam mengidentifikasi pola kemunculan kata pada dokumen teks meskipun menggunakan jumlah data pelatihan yang relatif terbatas [12], [13]. *Naive Bayes* sendiri merupakan sebuah metode klasifikasi yang berakar pada teorema Bayes, metode pengklasifikasian dengan menggunakan metode probabilitas dan statistik yang dikemukakan oleh ilmuwan Inggris Thomas Bayes, yaitu memprediksi peluang di masa depan berdasarkan pengalaman di masa sebelumnya sehingga dikenal sebagai Teorema Bayes. Ciri utama dari *Naive Bayes Classifier* ini adalah asumsi yang sangat kuat akan independensi dari masing-masing kondisi atau kejadian, adapun keuntungan penggunaan metode ini hanya membutuhkan jumlah data pelatihan yang kecil untuk menentukan estimasi parameter yang diperlukan dalam proses pengklasifikasian. Karena yang diasumsikan sebagai *variable independent*, maka hanya varians dari suatu *variable* dalam sebuah kelas yang dibutuhkan untuk menentukan klasifikasi, bukan keseluruhan dari matriks kovarians [14]. Di sisi lain, *Support Vector Classifier* (SVC) yang merupakan implementasi dari konsep dasar *Support Vector Machine* (SVM) untuk kasus klasifikasi, secara empiris terbukti tangguh dalam memetakan ruang dimensi fitur tinggi menggunakan fungsi *hyperplane*, menjadikannya kompetitor yang solid dalam domain linguistik komputasional. Selain itu, kombinasi dan perbandingan kedua algoritma ini sering membuktikan bahwa keduanya memegang peranan vital dalam sentimen analitik [15], [16].

Selain itu, karakteristik data teks pada media sosial yang tidak terstruktur, mengandung *noise*, serta memiliki variasi bahasa yang tinggi menjadi tantangan tersendiri dalam proses analisis sentimen. Kondisi ini menuntut penggunaan metode klasifikasi yang tidak hanya mampu mengolah data dalam skala besar, tetapi juga memiliki ketahanan terhadap ketidakseimbangan kelas dan kompleksitas pola linguistik yang beragam. Beberapa penelitian terbaru menunjukkan bahwa performa model klasifikasi sangat dipengaruhi oleh teknik prapemrosesan dan metode yang digunakan, terutama pada data media sosial yang bersifat dinamis, sehingga pemilihan metode yang tepat menjadi faktor penting dalam menghasilkan analisis sentimen yang akurat dan konsisten [17].

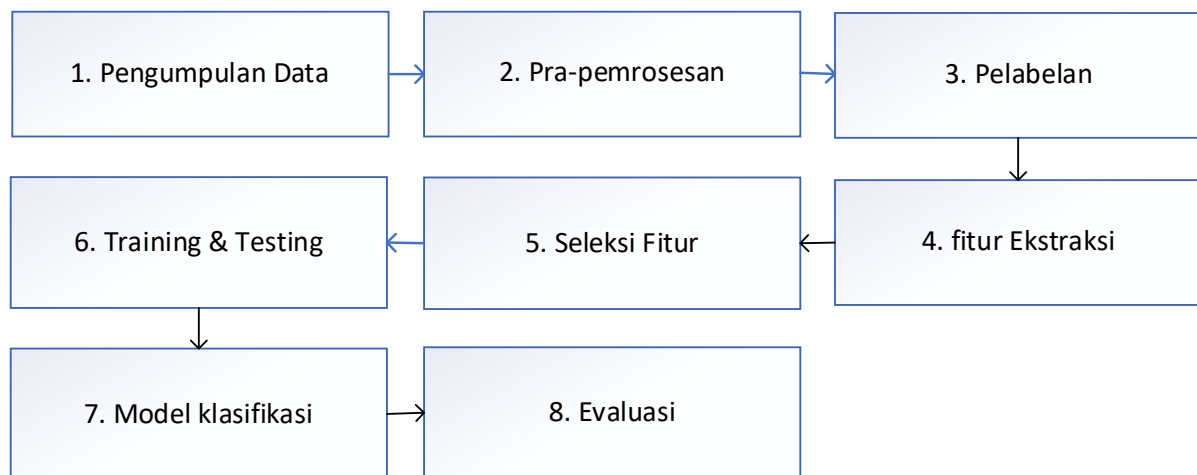
Berdasarkan urgensi tersebut, penelitian ini bertujuan untuk memetakan dan membandingkan pendapat masyarakat terhadap musibah banjir di Sumatra melalui data percakapan pada media sosial Twitter. Secara spesifik, penelitian ini mengimplementasikan tahapan *text mining* secara menyeluruh sesuai alur *flowchart* penelitian, mulai dari proses pra-pemrosesan data, transformasi teks, proses klasifikasi, hingga tahap evaluasi model menggunakan algoritma *Naive Bayes* dan *Support Vector Classifier* (SVC). Penggunaan istilah SVC dalam

penelitian ini dimaksudkan untuk merepresentasikan modul klasifikasi berbasis skema kelas diskrit yang dalam berbagai kajian literatur terdahulu sering digeneralisasikan sebagai Support Vector Machine (SVM). Kontribusi luaran penelitian ini dirancang sebagai instrumen wawasan teknis berbasis sains data yang dapat dimanfaatkan untuk mendukung pengembangan dashboard monitoring mitigasi dan evaluasi tanggap darurat bencana di Indonesia secara lebih terukur, adaptif, dan berbasis data.

2. METODE PENELITIAN

2.1 Tahapan Penelitian

Penelitian ini dilakukan secara bertahap yaitu pengumpulan data mentah (*Start*), dimana fase ini merupakan tahapan awal sebagai fase *Data Collection*, dilanjutkan ke tahapan *Preprocessing*, *Lexicon Labeling*, *Feature Extraction*, *Feature Selection*, *Train & Testing*, *Classification*, hingga *Evaluation* untuk mengukur hasil akurasi. Adapun tahapan penelitian secara rinci tampilan pada Gambar 1.



Gambar 1. Tahapan Penelitian

Pada Gambar 1 diatas menjelaskan setiap tahapan dalam penelitian ini, supaya penelitian lebih terarah dan sistematis sehingga mendapatkan hasil analisis sentimen yang sesuai dengan kebutuhan dalam penyelesaian masalah untuk penanganan banjir. Setiap tahapan di deskripsikan sebagai berikut :

2.2 Tahap Pengumpulan Data

Tahap awal dalam penelitian ini adalah proses pengumpulan data (*data crawling*) yang bertujuan untuk memperoleh data komentar masyarakat dari media sosial *Twitter* terkait peristiwa banjir di Sumatra. Pengumpulan data dilakukan menggunakan metode *crawling* dengan memanfaatkan kata kunci tertentu yang relevan dengan topik penelitian, seperti “banjir sumatera”, “bantuan”, “lambat”, “warga”, dan beberapa istilah lain yang sering digunakan masyarakat dalam membahas kondisi bencana. Proses *crawling* dilakukan pada rentang waktu eskalasi bencana agar data yang diperoleh merepresentasikan opini publik secara aktual selama terjadinya banjir. Penggunaan media sosial *Twitter* dipilih karena platform tersebut memungkinkan masyarakat menyampaikan informasi, kritik, keluhan, maupun dukungan secara cepat dan *real-time*, sehingga dapat menjadi sumber data yang relevan dalam analisis sentimen kebencanaan.

Melalui proses *crawling* tersebut, berhasil diperoleh sebanyak 2.000 *raw data* komentar *Twitter*. Data yang dikumpulkan terdiri atas beberapa atribut penting, seperti source id, waktu scraping, username pengguna *Twitter*, isi komentar, serta waktu pengguna memposting komentar. Selanjutnya, data hasil *crawling* melalui tahap *preprocessing* awal berupa filtrasi duplikasi (*remove duplicates*) untuk menghapus komentar yang memiliki isi sama atau terulang akibat proses pengambilan data.

2.3 Prapemrosesan data

Pada tahap prapemrosesan data komentar terlebih dahulu dibersihkan menggunakan teknik text cleaning untuk mengurangi *noise* yang dapat memengaruhi kualitas analisis. Proses ini dilakukan dengan menghapus berbagai elemen yang tidak memiliki kontribusi semantik terhadap isi teks, seperti tautan (URL), mention pengguna (@username), hashtag (#tagar), angka, serta tanda baca. Selain itu, seluruh karakter pada teks diubah menjadi huruf kecil (*lowercase*) melalui proses case folding agar sistem dapat mengenali kata dengan format yang

seragam. Tahapan ini penting karena data yang diperoleh dari media sosial umumnya bersifat tidak terstruktur dan mengandung banyak variasi penulisan yang dapat menurunkan performa model analisis teks [18], [19].

Setelah proses pembersihan selesai, teks yang telah bersih disimpan pada atribut *cleaned_comment* untuk digunakan pada tahapan analisis berikutnya. Pembersihan data bertujuan untuk mempertahankan kata-kata yang benar-benar merepresentasikan isi opini atau sentimen pengguna, sehingga informasi yang diproses menjadi lebih relevan. Penghapusan elemen-elemen seperti URL, mention, dan tanda baca juga membantu mengurangi redundansi data serta mencegah model mempelajari pola yang tidak berkaitan dengan konteks sentimen. Dengan demikian, representasi teks yang dihasilkan menjadi lebih konsisten dan mudah dipahami oleh sistem pemrosesan bahasa alami (*Natural Language Processing*).

Tahapan preprocessing ini merupakan langkah penting sebelum dilakukan proses lanjutan, seperti tokenisasi, penghapusan *stopword*, *stemming*, maupun pelabelan sentimen. Data yang telah dibersihkan akan menghasilkan intisari teks dengan kualitas yang lebih baik sehingga dapat meningkatkan akurasi dalam proses klasifikasi maupun analisis sentimen. Selain itu, normalisasi teks melalui preprocessing juga membantu mengurangi variasi kata yang tidak diperlukan, sehingga proses ekstraksi fitur menjadi lebih optimal dan efisien dalam membangun model analisis data teks.

2.4. Pelabelan Data

Tahap pelabelan data dilakukan menggunakan pendekatan *lexicon-based labeling*, yaitu proses pemberian label sentimen berdasarkan keberadaan kata-kata tertentu yang telah dikelompokkan ke dalam kamus sentimen positif dan negatif. Pada penelitian ini, kamus leksikon yang digunakan tidak mengadopsi pustaka (*library*) eksternal secara mentah, melainkan disusun dan dikembangkan secara manual oleh peneliti dengan penyesuaian khusus terhadap karakteristik dan konteks linguistik bencana banjir di Sumatra. Pendekatan kustomisasi kamus secara manual ini merujuk pada penelitian [20] yang menyatakan bahwa penyusunan kata leksikon secara mandiri oleh peneliti efektif untuk mendongkrak kepekaan model terhadap istilah spesifik atau bahasa informal yang berkembang dinamis di media sosial Twitter.

Dalam implementasinya, kamus kata positif memuat daftar kosakata bernuansa dukungan, harapan, apresiasi, atau solidaritas seperti "*bagus*", "*hebat*", "*sigap*", "*amanah*", "*membantu*", "*bantu*", "*doa*", dan "*peduli*". Sementara itu, kamus kata negatif memuat kosakata bernuansa keluhan, kritik, kerugian, atau kekecewaan seperti "*parah*", "*gagal*", "*lambat*", "*korupsi*", "*mengecewakan*", "*korup*", dan "*buruk*". Pemilihan pendekatan leksikon ini dinilai efektif untuk mengidentifikasi kecenderungan opini awal pada data komentar Twitter yang ringkas dan ekspresif tanpa memerlukan anotasi manual dalam jumlah besar di awal fase.

Proses kalkulasi pelabelan dilakukan dengan menghitung frekuensi kemunculan kata-kata positif dan negatif pada setiap entitas komentar yang telah melewati tahapan prapemrosesan teks secara terarah. Aturan keputusan logis untuk menentukan polaritas sentimen akhir ditetapkan berdasarkan bobot skor berikut:

Jika akumulasi kemunculan skor positif lebih besar dibandingkan kata negatif, maka akan menghasilkan sentimen positif

$$\text{Skor Positif} > \text{Skor Negatif}$$

Jika akumulasi skor negatif lebih besar dibandingkan skor kata positif, maka akan menghasilkan sentimen negatif.

$$\text{Skor Negatif} > \text{Skor Positif}$$

Jika sebuah komentar tidak mengandung kata positif atau negatif sama sekali, maka akan menghasilkan sentiment netral.

$$\text{Skor} = 0$$

Proses validasi terhadap hasil pelabelan leksikon manual ini diuji secara komputasional melalui pengujian performa pemodelan *Machine Learning*. Pertama, dataset biner hasil pelabelan leksikon (601 data) dipartisi menggunakan fungsi *train_test_split* dengan proporsi proporsional, yaitu 80% alokasi untuk data latih (*training data*) dan 20% alokasi untuk data uji (*testing data* yang berjumlah 121 sampel data biner). Kedua, algoritma komputer (Multinomial Naive Bayes dan SVC) dilatih menggunakan porsi data 80% tersebut untuk menguji apakah mesin dapat mengidentifikasi hubungan matematis yang konsisten antara bobot kata TF-IDF dengan label target yang dirancang peneliti. Ketiga, kemampuan prediktif model dievaluasi secara komprehensif menggunakan sisa data uji 20% melalui instrumen *Confusion Matrix* untuk mengukur metrik performa ilmiah berupa nilai akurasi dan *F1-Score*.

2.5. Ekstraksi Fitur

Untuk meningkatkan kualitas klasifikasi sentimen sekaligus mengurangi bias akibat distribusi data yang tidak seimbang, penelitian ini menerapkan pendekatan *Binary Classification* dengan mengeliminasi kelas Netral pada

tahap filtering data. Penghapusan kelas Netral dilakukan karena komentar dengan polaritas netral umumnya tidak mengandung kecenderungan opini yang kuat, sehingga berpotensi menurunkan kemampuan model dalam membedakan sentimen positif dan negatif secara lebih jelas. Dengan hanya memfokuskan klasifikasi pada dua kelas utama, yaitu Positif dan Negatif, model diharapkan mampu mempelajari pola subjektivitas opini secara lebih optimal. Pendekatan ini juga membantu mengurangi kompleksitas klasifikasi dan meningkatkan sensitivitas model terhadap ekspresi sentimen yang bersifat eksplisit.

Setelah proses filtering selesai dilakukan, data teks ditransformasikan ke dalam bentuk representasi numerik menggunakan metode *Term Frequency-Inverse Document Frequency* (TF-IDF). Metode TF-IDF digunakan karena mampu memberikan bobot lebih tinggi pada kata-kata yang sering muncul dalam suatu dokumen tetapi jarang muncul pada keseluruhan korpus, sehingga kata yang memiliki nilai informasi tinggi dapat lebih diperhatikan oleh model klasifikasi [21], [22]. Pada penelitian ini, jumlah fitur dibatasi sebanyak 1000 fitur menggunakan parameter *max_features=1000* untuk mengontrol dimensi data dan mengurangi kompleksitas komputasi. Selain itu, ekstraksi fitur juga mempertimbangkan penggunaan pendekatan *N-Grams*, khususnya kombinasi unigram dan bigram (1,2), agar model mampu memahami hubungan antar kata dalam konteks kalimat. Pendekatan ini penting untuk menangkap pola linguistik tertentu, seperti frasa negasi “tidak bagus”, “kurang cepat”, atau “tidak lambat”, yang sering kali memiliki makna berbeda apabila kata dianalisis secara terpisah.

Tahapan berikutnya adalah pembagian dataset menjadi data latih (*Data Train*) dan data uji (*Data Test*) dengan komposisi 80:20 menggunakan metode *train_test_split*. Sebanyak 80% data digunakan untuk melatih model agar dapat mempelajari pola sentimen, sedangkan 20% sisanya digunakan untuk mengukur kemampuan generalisasi model terhadap data yang belum pernah dilihat sebelumnya. Pada penelitian ini, proses pelatihan dilakukan menggunakan algoritma *Support Vector Machine* (SVM) dengan kernel linear. Pemilihan kernel linear didasarkan pada karakteristik data teks berdimensi tinggi yang umumnya lebih efektif dipisahkan menggunakan *hyperplane* linear. Algoritma SVM bekerja dengan mencari batas pemisah optimal (*optimal hyperplane*) yang mampu memaksimalkan jarak antar kelas sehingga proses klasifikasi menjadi lebih akurat. Setelah model selesai dilatih menggunakan data latih, model kemudian diuji menggunakan data uji untuk mengevaluasi performa klasifikasi sentimen berdasarkan pola fitur TF-IDF yang telah dihasilkan.

2.6. Seleksi Fitur

Setelah proses ekstraksi fitur selesai dilakukan, tahap selanjutnya adalah seleksi fitur guna menyaring dan memilih representasi kata yang paling relevan dari kamus leksikon sentimen sebelum diproses oleh algoritma. Kamus sentimen positif memuat kata-kata bernuansa dukungan seperti "bantu", "doa", dan "peduli", sementara kamus negatif memuat kata-kata bernuansa keluhan seperti "lambat", "korup", dan "parah". Proses ini secara default menghasilkan tiga kelas polaritas sentimen, yakni Positif, Negatif, dan Netral.

Namun, untuk mengatasi potensi ketimpangan data dan memfokuskan mesin untuk mempelajari opini subjektif murni, dilakukan tahapan seleksi fitur (*filtering*) melalui penghapusan seluruh entitas teks yang terdeteksi sebagai kelas Netral. Transformasi skema dari klasifikasi multikelas menjadi *Binary Classification* (Positif dan Negatif) ini sangat efektif untuk menyederhanakan varians data latih serta mengeliminasi *noise* dari cuitan yang hanya bersifat informatif. Setelah kelas target biner didapatkan, teks komputasional ditransformasikan ke dalam matriks numerik bobot menggunakan teknik *Term Frequency-Inverse Document Frequency* (TF-IDF). Parameter pembobotan dikalibrasi menggunakan metode *N-Grams* berukuran (1, 2) atau *bigram* untuk mendongkrak kepekaan algoritma terhadap konteks linguistik yang terdiri atas dua kata berturutan (seperti "tidak lambat" atau "sangat buruk").

2.7. Proses latih dan Uji

Matriks fitur TF-IDF yang telah terbentuk selanjutnya dipartisi untuk keperluan pengujian validitas model. Dataset dibagi ke dalam dua proporsi, yaitu 80% alokasi untuk Data Latih (*Training Data*) dan 20% alokasi untuk Data Uji (*Testing Data*) menggunakan fungsi *train_test_split* dari pustaka Scikit-Learn dengan parameter *random state* yang diatur tetap guna memastikan proses komputasi yang berulang.

Pelatihan model dieksekusi secara simultan untuk membandingkan kinerja dua algoritma *Machine Learning*. Algoritma pertama adalah *Multinomial Naive Bayes* yang dilatih menggunakan parameter *alpha smoothing = 0,1* untuk mencegah probabilitas bernilai nol pada kelas fitur yang tidak terlihat. Algoritma kedua adalah *Support Vector Classifier* (SVC) yang dilatih menggunakan *linear kernel* serta parameter penalti *class_weight = 'balanced'* untuk memberikan pembobotan adaptif pada margin *hyperplane* guna mengatasi sisa ketidakseimbangan frekuensi data latih antar kelas. Proses ini diakhiri dengan fase pengujian model menggunakan data uji 20%, di mana kemampuan prediktif algoritma dievaluasi secara komprehensif menggunakan instrumen *Confusion Matrix* yang menghasilkan metrik Akurasi, Presisi, *Recall*, dan *F1-Score*.

2.8. Model Klasifikasi

Tahap model klasifikasi dilakukan untuk mengelompokkan data teks ke dalam kelas sentimen yang telah ditentukan, yaitu positif, negatif, dan netral. Proses klasifikasi merupakan inti dari analisis sentimen karena pada

tahap ini sistem mulai mempelajari pola hubungan antara fitur teks hasil ekstraksi TF-IDF dengan label sentimen yang telah diperoleh dari proses pelabelan sebelumnya. Data komentar yang telah melalui preprocessing dan transformasi fitur akan direpresentasikan dalam bentuk vektor numerik sehingga dapat diproses oleh algoritma machine learning. Pada penelitian ini digunakan dua algoritma klasifikasi, yaitu *Naive Bayes* dan *Support Vector Classifier* (SVC), yang dipilih karena keduanya merupakan metode yang paling banyak digunakan dalam penelitian analisis sentimen berbasis teks, khususnya pada data media sosial seperti Twitter, serta memiliki kemampuan yang baik dalam menangani data berdimensi tinggi.

Naive Bayes merupakan metode klasifikasi berbasis probabilitas yang mengacu pada Teorema Bayes dengan asumsi bahwa setiap fitur bersifat independen terhadap fitur lainnya. Dalam konteks analisis sentimen, fitur yang dimaksud berupa kata atau term hasil ekstraksi teks. Penelitian ini menggunakan varian *Multinomial Naive Bayes* karena algoritma tersebut dirancang khusus untuk menangani data diskrit berbasis frekuensi kata dan sangat sesuai digunakan pada representasi TF-IDF. Prinsip kerja algoritma ini adalah menghitung probabilitas suatu dokumen termasuk ke dalam kelas tertentu berdasarkan distribusi kemunculan kata pada data pelatihan. Setiap kata pada dokumen akan memberikan kontribusi probabilitas terhadap kelas sentimen sehingga model dapat menentukan kategori dengan probabilitas tertinggi sebagai hasil prediksi. Selain memiliki proses pelatihan yang relatif cepat, *Multinomial Naive Bayes* juga dikenal efisien dalam penggunaan sumber daya komputasi dan mampu memberikan performa yang stabil pada dataset teks berskala besar [23][24]. Pada penelitian ini dilakukan penyesuaian parameter $\alpha = 0,1$ sebagai teknik smoothing untuk menghindari probabilitas nol pada kata-kata yang jarang muncul di data latih. Pengaturan parameter ini bertujuan agar model tetap dapat melakukan prediksi secara optimal meskipun terdapat variasi kata yang tidak muncul pada seluruh kelas sentimen.

Sementara itu, *Support Vector Classifier* (SVC) merupakan implementasi dari metode *Support Vector Machine* (SVM) yang bekerja dengan mencari hyperplane optimal untuk memisahkan data antar kelas pada ruang fitur berdimensi tinggi. Pada data teks, jumlah fitur yang dihasilkan dari TF-IDF umumnya sangat besar sehingga SVC menjadi salah satu metode yang efektif karena mampu melakukan pemisahan kelas secara optimal meskipun dimensi data tinggi dan kompleks. Algoritma ini bekerja dengan menentukan *support vector*, yaitu titik data yang berada paling dekat dengan batas pemisah antar kelas. *Hyperplane* terbaik dipilih berdasarkan margin terbesar antara dua kelas sehingga model memiliki kemampuan generalisasi yang lebih baik terhadap data baru. Penelitian ini menggunakan linear kernel karena karakteristik data teks umumnya lebih efektif diproses menggunakan pemisahan linear dibandingkan kernel non-linear. Selain itu, penggunaan linear kernel juga mampu mengurangi kompleksitas komputasi dan mempercepat proses pelatihan model tanpa mengurangi performa klasifikasi secara signifikan [25]. Untuk mengatasi ketidakseimbangan distribusi data sentimen, model SVC juga menerapkan parameter *class weight* = balanced sehingga bobot tiap kelas dapat disesuaikan secara otomatis berdasarkan proporsi jumlah data pada masing-masing kelas.

Beberapa penelitian sebelumnya menunjukkan bahwa perbandingan antara *Naive Bayes* dan SVC sering digunakan untuk mengetahui algoritma yang paling optimal dalam analisis sentimen berbasis media sosial. *Naive Bayes* memiliki keunggulan dari sisi efisiensi, kesederhanaan model, dan kecepatan komputasi, sedangkan SVC unggul dalam kemampuan pemisahan data kompleks dan akurasi klasifikasi pada data berdimensi tinggi. Oleh karena itu, penggunaan kedua algoritma secara komparatif dapat memberikan gambaran yang lebih objektif terkait performa model dalam memahami pola sentimen masyarakat pada data Twitter. Dalam penelitian ini, kedua model dilatih menggunakan data latih sebesar 80% dan diuji menggunakan data uji sebesar 20%. Hasil prediksi kemudian dievaluasi menggunakan beberapa metrik performa seperti akurasi, dan *F1-Score* untuk mengetahui kemampuan masing-masing algoritma dalam mengklasifikasikan sentimen secara tepat dan konsisten.

3. HASIL DAN PEMBAHASAN

3.1 Implementasi *Binary Classification* dan Distribusi Sentimen

Berdasarkan hasil pra-pemrosesan dan proses pelabelan awal menggunakan pendekatan leksikon, diperoleh distribusi data sentimen yang menunjukkan dominasi sangat tinggi pada kelas Netral dibandingkan kelas Positif dan Negatif. Kondisi ini mengindikasikan bahwa sebagian besar komentar hanya berisi informasi umum, pernyataan singkat, atau teks yang tidak menunjukkan kecenderungan opini secara eksplisit. Dalam analisis sentimen, dominasi kelas tertentu dapat menyebabkan terjadinya ketimpangan data (*data imbalance*), yaitu kondisi ketika jumlah data pada satu kelas jauh lebih besar dibandingkan kelas lainnya. Ketidakseimbangan tersebut berpotensi memengaruhi proses pembelajaran model machine learning karena algoritma cenderung lebih fokus mempelajari pola dari kelas mayoritas dan mengabaikan karakteristik kelas minoritas. Akibatnya, sensitivitas model dalam mendeteksi sentimen positif maupun negatif dapat menurun dan menghasilkan performa klasifikasi yang bias.

Untuk mengatasi permasalahan tersebut, penelitian ini menerapkan pendekatan *Binary Classification* dengan mengeliminasi seluruh data pada kelas Netral. Penghapusan kelas Netral dilakukan karena komentar netral umumnya tidak mengandung ekspresi emosi, dukungan, maupun kritik yang kuat sehingga kurang relevan dalam

3.2 Hasil ekstraksi fitur

Pada tahap ekstraksi fitur, seluruh korpus teks yang telah melalui proses preprocessing dan filtering ditransformasikan ke dalam bentuk representasi numerik menggunakan metode *Term Frequency-Inverse Document Frequency* (TF-IDF). Transformasi ini bertujuan agar data teks yang awalnya berupa *unstructured text* dapat diproses oleh algoritma machine learning dalam bentuk matriks fitur numerik. Metode TF-IDF bekerja dengan menghitung tingkat kepentingan suatu kata berdasarkan frekuensi kemunculannya pada sebuah dokumen dibandingkan dengan frekuensi kemunculan kata tersebut pada seluruh dokumen dalam korpus. Dengan pendekatan ini, kata-kata yang sering muncul pada satu komentar namun jarang muncul pada keseluruhan data akan memperoleh bobot lebih tinggi karena dianggap memiliki nilai informasi yang lebih kuat dalam merepresentasikan sentimen tertentu.

Pada penelitian ini, proses ekstraksi fitur tidak hanya menggunakan unigram, tetapi juga menerapkan pendekatan *N-Grams* dengan kombinasi rentang (1,2) yang mencakup unigram dan bigram. Penggunaan bigram memungkinkan model menangkap hubungan kontekstual antar kata yang sering kali memiliki makna berbeda apabila dipisahkan. Sebagai contoh, frasa seperti “bantuan lambat”, “pemerintah gagal”, “tidak sigap”, atau “saling bantu” memiliki makna sentimen yang lebih jelas dibandingkan analisis kata tunggal secara terpisah. Pendekatan ini menjadi penting karena data media sosial umumnya mengandung pola bahasa informal, ekspresi spontan, dan frasa pendek yang maknanya sangat dipengaruhi oleh kombinasi antar kata. Dengan penggunaan *N-Grams*, sistem mampu merekam ribuan kombinasi *vocabulary* penting yang merepresentasikan opini masyarakat secara lebih mendalam dan kontekstual.

Selain meningkatkan pemahaman konteks kalimat, metode TF-IDF juga membantu meminimalkan pengaruh kata-kata umum yang sering muncul pada seluruh dokumen namun tidak memiliki kontribusi signifikan terhadap proses klasifikasi sentimen. Kata-kata dengan frekuensi tinggi tetapi kurang informatif akan memperoleh penalti bobot yang lebih rendah, sedangkan kata-kata yang menjadi pembeda utama antar kelas sentimen akan memperoleh bobot lebih tinggi. Sebagai contoh, kata atau frasa seperti “gagal”, “korupsi”, “lambat”, “kecewa”, “sigap”, “peduli”, atau “solidaritas” menjadi fitur penting dalam membedakan sentimen negatif dan positif pada data Twitter yang dianalisis. Hasil transformasi TF-IDF kemudian menghasilkan matriks fitur berdimensi tinggi yang berisi representasi numerik dari setiap komentar berdasarkan bobot kata dan frasa yang dimiliki.

Matriks fitur inilah yang menjadi dasar utama dalam proses pembelajaran model klasifikasi pada tahap berikutnya. Kualitas representasi fitur sangat menentukan kemampuan algoritma machine learning dalam mengenali pola sentimen secara akurat. Semakin baik metode ekstraksi fitur dalam menangkap karakteristik bahasa dan konteks opini masyarakat, maka semakin tinggi pula kemampuan model dalam membedakan sentimen positif dan negatif. Oleh karena itu, kombinasi TF-IDF dan *N-Grams* pada penelitian ini berperan penting dalam meningkatkan kualitas representasi data teks sehingga algoritma klasifikasi dapat melakukan prediksi sentimen secara lebih optimal dan konsisten.

3.3 Hasil Seleksi Fitur

Proses seleksi fitur melalui tahapan filtering memberikan dampak yang sangat signifikan terhadap komposisi dataset penelitian. Sebelum dilakukan seleksi, distribusi sentimen pada 2.000 *raw data* yang diekstraksi sangat didominasi oleh kelas Netral. Dominasi kelas Netral ini berpotensi menjadi *noise* yang mengaburkan sensitivitas algoritma dalam mempelajari dikotomi opini positif (dukungan/apresiasi) dan opini negatif (kritik/keluhan) secara presisi. Oleh karena itu, penyederhanaan ruang dimensi masalah dari klasifikasi multikelas (Positif, Negatif, Netral) menjadi *Binary Classification* (Positif dan Negatif) dieksekusi sebagai strategi data preprocessing lanjutan. Hasil dari tahapan seleksi fitur biner ini sukses menghasilkan korpus data yang sangat bersih dan seimbang (*balanced corpus*). Dari total data mentah, proses filtrasi secara spesifik menyisakan 601 data opini murni.

3.4 Perbandingan Kinerja Algoritma Klasifikasi

Pengujian performa prediktif dilakukan menggunakan instrumen *Confusion Matrix* terhadap 20% porsi data uji (121 sampel data biner). Parameter pengujian berfokus pada kemampuan model dalam menebak kelas minoritas (Negatif dan Positif) tanpa mengorbankan tebakan secara keseluruhan.

Tabel 1. Perbandingan Algoritma *Naive Bayes* dengan SVC

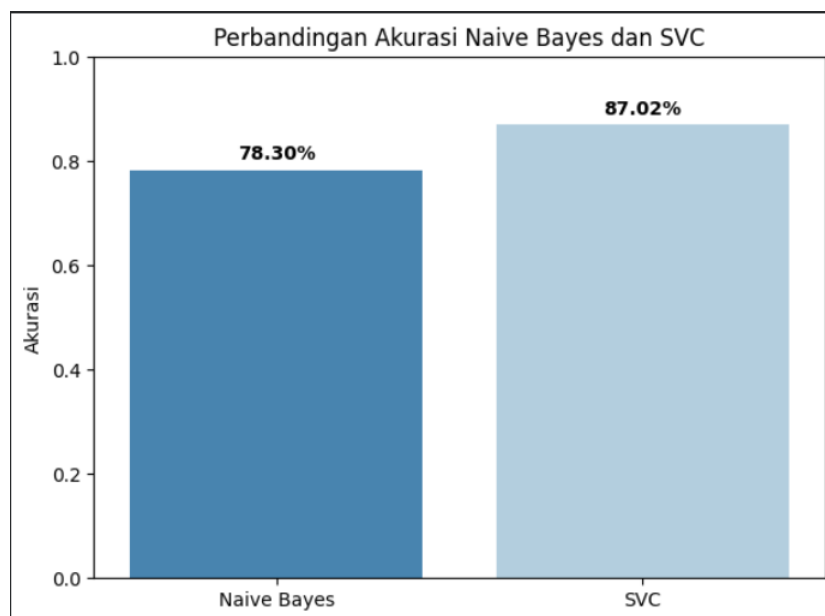
Algoritma	Akurasi	F1-Score
Naive Bayes	78.30%	0,62
Support Vector Classifier	87.02%	0,79

Tabel 1 menunjukkan hasil perbandingan performa algoritma *Multinomial Naive Bayes* dan *Support Vector Classifier* (SVC) pada proses klasifikasi sentimen data Twitter terkait banjir. Berdasarkan hasil pengujian, algoritma SVC memperoleh performa yang lebih tinggi dibandingkan *Naive Bayes* dengan akurasi sebesar 87,02% dan *F1-Score* sebesar 0,79. Sementara itu, algoritma *Multinomial Naive Bayes* menghasilkan akurasi sebesar 78,30% dengan *F1-Score* sebesar 0,62. Perbedaan nilai tersebut menunjukkan bahwa SVC memiliki kemampuan yang lebih baik dalam mengenali pola sentimen positif maupun negatif pada data teks yang telah direpresentasikan menggunakan TF-IDF dan bigram.

Tingginya performa SVC dipengaruhi oleh kemampuannya dalam membentuk hyperplane optimal pada ruang fitur berdimensi tinggi hasil transformasi TF-IDF. Pada penelitian ini, penggunaan fitur bigram membantu model memahami hubungan antar kata yang memiliki makna kontekstual, seperti frasa “bantuan lambat”, “pemerintah gagal”, atau “saling bantu”. Karakteristik tersebut dapat diproses dengan lebih baik oleh SVC karena algoritma ini bekerja dengan mencari batas pemisah terbaik antar kelas berdasarkan support vector yang berada paling dekat dengan margin klasifikasi. Selain itu, penerapan parameter *class weight = balanced* membantu model mengurangi bias terhadap jumlah data pada masing-masing kelas sehingga kemampuan klasifikasi sentimen positif dan negatif menjadi lebih seimbang.

Di sisi lain, performa *Multinomial Naive Bayes* masih berada di bawah SVC meskipun telah dilakukan optimasi parameter *alpha*. Hal ini disebabkan oleh sifat dasar *Naive Bayes* yang mengasumsikan bahwa setiap fitur bersifat independen satu sama lain. Pada data teks, khususnya penggunaan bigram, hubungan antar kata sebenarnya saling berkaitan dan membentuk konteks tertentu. Akibatnya, *Naive Bayes* mengalami keterbatasan dalam memahami keterkaitan makna antar kata sehingga performa klasifikasi menjadi kurang optimal dibandingkan SVC. Meskipun demikian, *Naive Bayes* tetap memiliki keunggulan dari sisi efisiensi komputasi dan kecepatan proses pelatihan model.

Untuk memberikan gambaran visual yang lebih jelas mengenai perbedaan performa kedua algoritma, hasil akurasi dan *F1-Score* divisualisasikan dalam bentuk grafik batang pada Gambar 3. Grafik tersebut memperlihatkan bahwa batang performa SVC berada lebih tinggi dibandingkan *Naive Bayes* pada kedua metrik evaluasi. Perbedaan tinggi batang pada grafik menunjukkan adanya selisih performa yang cukup signifikan, khususnya pada nilai *F1-Score* yang menggambarkan keseimbangan antara *precision* dan *recall* dalam proses klasifikasi sentimen.



Gambar 3. Grafik perbandingan Akurasi *Naive Bayes* dan SVC

Berdasarkan grafik batang pada Gambar 3, terlihat secara visual bahwa algoritma *Support Vector Classifier* (SVC) memberikan hasil klasifikasi yang lebih stabil dan konsisten dibandingkan *Multinomial Naive Bayes*. Tingginya nilai akurasi dan *F1-Score* pada SVC menunjukkan bahwa model mampu mengenali pola sentimen positif dan negatif dengan tingkat kesalahan yang lebih rendah. Dengan demikian, dapat disimpulkan bahwa algoritma SVC lebih efektif digunakan pada klasifikasi sentimen berbasis data Twitter dibandingkan *Multinomial Naive Bayes* dalam penelitian ini.

3.5 Implikasi

Secara fungsional, hasil implementasi machine learning pada penelitian ini menunjukkan bahwa pendekatan text mining dapat digunakan sebagai instrumen intelijen krisis yang efektif dalam menganalisis opini publik secara cepat dan terstruktur. Kemampuan sistem dalam mengolah ribuan komentar Twitter menjadi informasi sentimen yang terklasifikasi membuktikan bahwa media sosial dapat dimanfaatkan sebagai sumber data real-time untuk memantau kondisi sosial masyarakat selama terjadinya bencana. Informasi yang sebelumnya berbentuk unstructured text berhasil ditransformasikan menjadi pengetahuan yang lebih terukur melalui tahapan preprocessing, ekstraksi fitur TF-IDF, hingga klasifikasi menggunakan algoritma machine learning. Dengan demikian, analisis sentimen tidak hanya berfungsi sebagai kajian akademik, tetapi juga memiliki potensi implementasi langsung dalam mendukung pengambilan keputusan berbasis data.

Hasil pengujian menunjukkan bahwa arsitektur *Support Vector Classifier* (SVC) dengan linear kernel dan parameter class weight = balanced memiliki kemampuan yang baik dalam menyaring serta mengidentifikasi pola sentimen masyarakat di tengah tingginya volume informasi terkait banjir. Penggunaan linear kernel terbukti efektif dalam memproses representasi fitur TF-IDF berdimensi tinggi, sedangkan parameter class weight membantu model tetap sensitif terhadap kedua kelas sentimen meskipun terdapat perbedaan distribusi data. Melalui kombinasi tersebut, model mampu mengenali komentar yang mengandung kritik, keluhan, maupun dukungan masyarakat dengan tingkat akurasi yang tinggi. Secara praktis, kemampuan ini sangat penting dalam konteks kebencanaan karena media sosial sering kali menjadi saluran utama masyarakat untuk menyampaikan kondisi lapangan, keterlambatan bantuan, kerusakan infrastruktur, hingga kebutuhan darurat yang belum terpenuhi.

Selain mampu mengidentifikasi sentimen umum, model SVC juga berpotensi digunakan untuk mendeteksi isu prioritas secara lebih dini berdasarkan pola percakapan masyarakat. Sebagai contoh, kemunculan frasa seperti “bantuan lambat”, “akses terputus”, “banjir parah”, atau “belum dievakuasi” dapat menjadi indikator awal adanya permasalahan distribusi logistik dan respons penanganan bencana di suatu wilayah. Informasi semacam ini memiliki nilai strategis karena dapat membantu institusi publik memahami kondisi lapangan secara lebih cepat dibandingkan mekanisme pelaporan konvensional. Dengan memanfaatkan analisis sentimen berbasis media sosial, lembaga penanggulangan bencana dapat memperoleh gambaran mengenai persepsi masyarakat sekaligus mengidentifikasi titik permasalahan yang paling banyak dikeluhkan oleh warga terdampak.

Dalam konteks implementasi nyata, pipeline penelitian ini dapat diintegrasikan ke dalam sistem social listening pada institusi publik seperti Badan Penanggulangan Bencana Daerah (BPBD), pemerintah daerah, maupun lembaga kemanusiaan lainnya. Integrasi tersebut memungkinkan proses pengumpulan data Twitter, preprocessing, klasifikasi sentimen, hingga visualisasi hasil dilakukan secara otomatis dan real-time. Melalui sistem ini, lonjakan keluhan masyarakat terkait keterlambatan bantuan, distribusi logistik, maupun kegagalan mitigasi dapat terdeteksi lebih awal sehingga tim respons bencana dapat segera melakukan prioritas tindakan berdasarkan data yang aktual. Pendekatan berbasis machine learning ini berpotensi meningkatkan efektivitas pengambilan keputusan karena proses mitigasi tidak hanya bergantung pada laporan manual, tetapi juga didukung oleh analisis data sosial masyarakat secara langsung.

Lebih lanjut, penggunaan sistem social listening berbasis text mining juga dapat membantu pemerintah dalam mengevaluasi efektivitas kebijakan penanganan bencana dari perspektif masyarakat. Sentimen publik yang terekam pada media sosial dapat dijadikan indikator untuk mengukur tingkat kepuasan, kepercayaan, maupun kritik masyarakat terhadap respons pemerintah selama masa krisis. Dengan demikian, implementasi machine learning dalam analisis sentimen tidak hanya berkontribusi pada pengembangan teknologi kecerdasan buatan, tetapi juga memiliki dampak strategis dalam mendukung sistem mitigasi bencana yang lebih responsif, adaptif, dan data-driven di era digital.

4. KESIMPULAN

Berdasarkan hasil penelitian yang telah dilakukan, implementasi tahapan text mining berhasil digunakan untuk mengklasifikasikan sentimen masyarakat terkait krisis banjir di Sumatra melalui data Twitter. Penelitian ini membandingkan performa algoritma *Multinomial Naive Bayes* dan *Support Vector Classifier* (SVC) menggunakan representasi fitur TF-IDF dan pendekatan bigram. Hasil pengujian menunjukkan bahwa algoritma SVC memberikan performa terbaik dengan tingkat akurasi sebesar 87,02% dan *F1-Score* sebesar 0,79, lebih unggul dibandingkan *Multinomial Naive Bayes* yang memperoleh akurasi sebesar 78,30% dan *F1-Score* sebesar 0,62. Tingginya performa SVC menunjukkan bahwa algoritma tersebut lebih efektif dalam menangani data teks berdimensi tinggi serta mampu memahami pola sentimen positif dan negatif secara lebih akurat melalui pemanfaatan fitur kontekstual pada data Twitter.

Secara praktis, hasil penelitian ini menunjukkan bahwa analisis sentimen berbasis machine learning memiliki potensi besar untuk diterapkan sebagai pendukung pengambilan keputusan dalam situasi kebencanaan. Model SVC dapat dimanfaatkan sebagai bagian dari sistem social listening untuk membantu institusi publik memantau

opini masyarakat, mendeteksi keluhan terkait distribusi bantuan, serta mengidentifikasi respons publik terhadap penanganan bencana secara real-time. Meskipun demikian, penelitian ini masih memiliki keterbatasan pada pendekatan pelabelan berbasis leksikon yang belum sepenuhnya mampu memahami konteks bahasa kompleks seperti sarkasme dan ironi. Oleh karena itu, penelitian selanjutnya disarankan untuk mengembangkan leksikon yang lebih dinamis atau menerapkan metode *Deep Learning* berbasis *Bidirectional* agar kemampuan pemahaman konteks kalimat dapat meningkat dan menghasilkan performa klasifikasi yang lebih optimal.

UCAPAN TERIMAKASIH

Terima kasih disampaikan kepada pihak-pihak yang telah mendukung terlaksananya penelitian ini.

DAFTAR PUSTAKA

- [1] D. Darwis, N. Siskawati, and Z. Abidin, "Penerapan Algoritma Naive Bayes Untuk Analisis Sentimen Review Data Twitter BMKG Nasional," *Jurnal Tekno Kompak*, vol. 15, no. 1, p. 131, Jan. 2021, doi: 10.33365/JTK.V15I1.744.
- [2] F. Nurrachmat Hidayat, "Analisis Sentimen Masyarakat Terhadap Perekrutan PPPK Pada Twitter Dengan Metode Naive Bayes ... (Fajar Nurrachmat Hidayat, Sugiyono) Analisis Sentimen Masyarakat Terhadap Perekrutan PPPK Pada Twitter Dengan Metode Naive Bayes Dan Support Vector Machine," *Jurnal Sains dan Teknologi*, vol. 5, no. 2, pp. 665–672, doi: 10.55338/saintek.v5i1.1359.
- [3] D. Darwis, E. Shintya Pratiwi, A. Ferico, and O. Pasaribu, "PENERAPAN ALGORITMA SVM UNTUK ANALISIS SENTIMEN PADA DATA TWITTER KOMISI PEMBERANTASAN KORUPSI REPUBLIK INDONESIA," 2020.
- [4] S. Berliani and S. Lestari, "Analisis Sentimen Masyarakat Terhadap Isu Pecat Sri Mulyani Pada Twitter Menggunakan Metode Naive Bayes Dan Support Vector Machine," *Jurnal Sains dan Teknologi*, vol. 5, no. 3, pp. 951–960, Apr. 2024, doi: 10.55338/saintek.v5i3.2746.
- [5] A. Agustian, T. Tukiro, and F. Nurapriani, "Analisis Sentimen, Text Mining Penerapan Analisis Sentimen Dan Naive Bayes Terhadap Opini Penggunaan Kendaraan Listrik Di Twitter," *Jurnal Tika*, vol. 7, no. 3, pp. 243–249, Dec. 2022, doi: 10.51179/TIKA.V7I3.1550.
- [6] S. Suryani, M. F. Fayyad, D. T. Savra, V. Kurniawan, and B. H. Estanto, "Sentiment Analysis of Towards Electric Cars using Naive Bayes Classifier and Support Vector Machine Algorithm," *Public Research Journal of Engineering, Data Technology and Computer Science*, vol. 1, no. 1, pp. 1–9, Jul. 2023, doi: 10.57152/PREDATECS.V1I1.814.
- [7] F. Fathoni, A. Ibrahim, F. R. Mumtaz, M. A. Zaky, M. J. Pratama, and I. A. Kurniawan, "ANALISIS SENTIMEN PUBLIC TWITTER TERHADAP KEBIJAKAN PEMERINTAH MENGGUNAKAN METODE SVM (STUDI KASUS : RUU TNI)," *JATI (Jurnal Mahasiswa Teknik Informatika)*, vol. 9, no. 4, pp. 6322–6329, May 2025, doi: 10.36040/JATI.V9I4.14036.
- [8] D. Auliya Agustina, S. Subanti, E. Zukhronah, P. Studi Statistika, and U. Sebelas Maret, "Implementasi Text Mining Pada Analisis Sentimen Pengguna Twitter Terhadap Marketplace di Indonesia Menggunakan Algoritma Support Vector Machine," *Indonesian Journal of Applied Statistics*, vol. 3, no. 2, pp. 109–122, Jan. 2021, doi: 10.13057/IJAS.V3I2.44337.
- [9] A. Tiara Susilawati, A. H. Tiara Susilawati Universitas Muhammadiyah Kalimantan Timur Nur Anjeni Lestari Universitas Muhammadiyah Kalimantan Timur Puput Alpria Nina Universitas Muhammadiyah Kalimantan Timur Jl Ir Juanda No, K. Samarinda Ulu, K. Samarinda, and K. Timur, "Analisis Sentimen Publik Pada Twitter Terhadap Boikot Produk Israel Menggunakan Metode Naive Bayes," *Nian Tana Sikka : Jurnal ilmiah Mahasiswa*, vol. 2, no. 1, pp. 26–35, Dec. 2024, doi: 10.59603/NIANTANASIKKA.V2I1.240.
- [10] A. Sentimen Terhadap Aplikasi Satu Sehat Pada Twitter Menggunakan Algoritma Naive Bayes Dan and F. Matheos Sarimole, "Analisis Sentimen Terhadap Aplikasi Satu Sehat Pada Twitter Menggunakan Algoritma Naive Bayes Dan Support Vector Machine," *Jurnal Sains dan Teknologi*, vol. 5, no. 3, pp. 783–790, Feb. 2024, doi: 10.55338/SAINTEK.V5I3.2702.
- [11] J. Homepage, R. Andini Husen, R. Astuti, L. Marlia, and L. Efrizoni, "Analisis Sentimen Opini Publik pada Twitter Terhadap Bank BSI Menggunakan Algoritma Machine Learning: Sentiment Analysis of Public Opinion on Twitter Toward," *journal.irpi.or.id*, vol. 3, pp. 211–218, 2023, doi: 10.57152/malcom.v3i2.901.

- [12] A. Perdana, A. Hermawan, and D. Avianto, "Analisis Sentimen Terhadap Isu Penundaan Pemilu di Twitter Menggunakan Naive Bayes Clasifier," *Jurnal Sisfokom (Sistem Informasi dan Komputer)*, vol. 11, no. 2, pp. 195–200, Jul. 2022, doi: 10.32736/SISFOKOM.V11I2.1412.
- [13] E. Febriyani and H. Februariyanti, "Analisis Sentimen Terhadap Program Kampus Merdeka Menggunakan Algoritma Naive Bayes Classifier Di Twitter," vol. 17, no. 1.
- [14] A. Felicia Watratan, A. B. Puspita, D. Moeis, S. Informasi, and S. Profesional Makassar, "Implementasi Algoritma Naive Bayes Untuk Memprediksi Tingkat Penyebaran Covid-19 Di Indonesia," 2020. [Online]. Available: <http://journal.isas.or.id/index.php/JACOST>
- [15] A. Tharwat, "Parameter investigation of support vector machine classifier with kernel functions," *Knowl. Inf. Syst.*, vol. 61, no. 3, pp. 1269–1302, 2019, doi: 10.1007/s10115-019-01335-4.
- [16] G. M. D. Iqbal *et al.*, "A Support Vector Machine based Logistic Regression Model Approach for Classification Problems," in *IISE Annual Conference and Expo 2023*, Institute of Industrial and Systems Engineers, IISE, 2023, doi: 10.21872/2023IISE_3358.
- [17] J. Saputra, L. Maryani, D. Wulandari, and W. Eka, "Analisis Performa Naive Bayes dan SVM terhadap Sentimen Teks Media Sosial dengan Word2Vec dan SMOTE under a Creative Commons Attribution-NonCommercial ShareAlike 4.0 International (CC BY-NC-SA 4.0)," vol. 10, no. 1, 2025, [Online]. Available: <https://journal.uin-laaluddin.ac.id/index.php/instek>
- [18] A. Kho and F. F. Tampinongkol, "Analisis Sentimen Pengguna Aplikasi STEAM dengan Algoritma Naive Bayes," *Ranah Research : Journal of Multidisciplinary Research and Development*, vol. 7, no. 6, 2025, doi: 10.38035/rj.v7i6.1803.
- [19] T. Ridwansyah, "Implementasi Text Mining Terhadap Analisis Sentimen Masyarakat Dunia Di Twitter Terhadap Kota Medan Menggunakan K-Fold Cross Validation Dan Naïve Bayes Classifier," *KLIK: Kajian Ilmiah Informatika dan Komputer*, vol. 2, no. 5, 2022, doi: 10.30865/klik.v2i5.362.
- [20] S. A. Nugraha, "PENERAPAN LEXICON BASED UNTUK ANALISIS SENTIMEN MASYARAKAT INDONESIA TERHADAP DANANTARA," 2025.
- [21] M. Salman Al Markas, S. Anraeni, and L. Budiman Ilmuwan, "Implementasi Fitur Vector Bag Of Word Dan TF IDF untuk Analisis Sentiment," *LINIER: Literatur Informatika dan Komputer*, vol. 2, no. 2, 2025, doi: 10.33096/linier.v2i2.3104.
- [22] M. F. Ramadhan, "Klasifikasi Topik dan Sentimen Judul Berita dengan Augmentasi dan TF-IDF," *RIGGS: Journal of Artificial Intelligence and Digital Business*, vol. 4, no. 2, 2025, doi: 10.31004/riggs.v4i2.1692.
- [23] A. D. Haryo and J. B. B. Darmawan, "Prosiding Nasional Rekayasa Teknologi Industri dan Informasi (ReTII) ke-20 262 Institut Teknologi Nasional Yogyakarta," pp. 262–268, 2025, [Online]. Available: <http://journal.itny.ac.id/index.php/ReTII>
- [24] M. F. Abdullah, A. Agus, and A. Miftachuddin, "Analisis Sentimen Review Aplikasi Ruang Guru Menggunakan Multinomial Naïve Bayes," 2025.
- [25] H. Hong, B. Pradhan, D. T. Bui, C. Xu, A. M. Youssef, and W. Chen, "Comparison of four kernel functions used in support vector machines for landslide susceptibility mapping: a case study at Suichuan area (China)," *Geomatics, Natural Hazards and Risk*, vol. 8, no. 2, 2017, doi: 10.1080/19475705.2016.1250112.