



Analisis Sentimen Masyarakat Di Media Sosial X Terhadap Kemenkes Dengan Naive Bayes dan SVM

Freddy Andrew Ryandi¹, Dian Pratiwi², Syandra Sari³

^{1,2}Informatika, Universitas Trisakti, Jakarta, Indonesia

³Sistem Informasi, Universitas Trisakti, Jakarta, Indonesia

Email Penulis Korespondensi: dian.pratiwi@trisakti.ac.id

Abstrak– Penelitian ini menganalisis sentimen masyarakat di platform media sosial X terhadap Kementerian Kesehatan Indonesia selama pandemi COVID-19 dengan menggunakan algoritma Naive Bayes dan Support Vector Machine (SVM). Data dari unggahan yang menyebut akun resmi Kemenkes (@KemenkesRI) dipreproses dan diberi label sentimen menggunakan VADER. Klasifikasi sentimen dilakukan dengan pembobotan TF-IDF, dan kedua algoritma dievaluasi. Hasil menunjukkan SVM mencapai akurasi yang sedikit lebih tinggi (79%) dibandingkan Naive Bayes (77%), menandakan efektivitasnya dalam menangani struktur bahasa yang kompleks, meskipun memerlukan sumber daya komputasi lebih tinggi. Penelitian ini menyoroti potensi SVM dalam analisis sentimen terhadap kebijakan kesehatan.

Kata Kunci: SVM, Naive Bayes, Analisis Sentimen, Sosial Media X, COVID-19

Abstract– This study examines public sentiment on social media platform X regarding Indonesia's Ministry of Health during the COVID-19 pandemic, using Naive Bayes and Support Vector Machine (SVM) algorithms. Posts mentioning the Ministry's official account (@KemenkesRI) were preprocessed and labeled using the VADER tool. Sentiment classification was performed with TF-IDF word weighting, and both algorithms were evaluated. Results show SVM achieved slightly higher accuracy (79%) than Naive Bayes (77%), indicating its effectiveness in handling complex language structures, though it requires more computational resources. This research underscores the utility of SVM for analyzing public sentiment on health policies..

Keywords: SVM, Naive Bayes, Sentiment Analysis, Social Media X, COVID-19

I. PENDAHULUAN

Pandemi COVID-19 yang terjadi di tahun 2020 adalah sebuah fenomena global yang memiliki dampak yang merambah ke banyak aspek kehidupan. Dampak yang luas ini diantara lain adalah pada aspek kesehatan masyarakat, ekonomi, dan standar budaya. Disaat yang krusial seperti ini Kementrian kesehatan RI memiliki peran penting dalam mengkoordinasikan penyebaran informasi yang relevan dan merancang kebijakan respons muncul sebagai pusat kesejahteraan masyarakat.

Seiring berkembangnya teknologi internet dan mulai ramai digunakannya media sosial, masyarakat Indonesia mulai sering menggunakan media sosial sebagai medium utama percakapan, penyebaran gagasan, dan juga ekspresi sentimen. Media sosial adalah . Salah satu media sosial yang kerap kali digunakan masyarakat Indonesia adalah X hal ini dapat disimpulkan karena Indonesia merupakan salah satu dari lima pengguna media sosial terbesar di dunia [1]. X adalah media sosial yang sering digunakan untuk membuat posting yang berisi gagasan seorang individu [4] tersebut dan juga untuk mengekspresikan sentimen individu tersebut terhadap suatu hal. Sebagai contoh, adalah sentimen terhadap Kementerian Kesehatan Republik Indonesia.

Data posting dari aplikasi X dapat menjadi basis yang baik untuk melakukan sebuah penelitian analisis sentimen [2], dengan Indonesia menjadi nomor 4 pengguna X terbanyak di dunia. Ada banyak data yang dapat diambil dari posting masyarakat tentang sesuatu, seperti contohnya

sentimen masyarakat terhadap Kementerian Kesehatan Republik Indonesia selama masa pandemi COVID-19.

Analisis sentimen adalah proses mengevaluasi sebuah opini, perasaan, atau sikap yang terdapat dalam sebuah badan teks, seperti ulasan, komentar, atau posting. Tujuannya adalah untuk menentukan apakah suatu badan teks menyuarakan sentimen negatif atau positif, terhadap sesuatu tertentu yang dibahas pada badan teks tersebut [3] [1].

Naive Bayes adalah metode yang dapat digunakan dalam melakukan analisis sentimen. Metode ini didasarkan pada teorema Bayes [2] dan mengasumsikan dengan sederhana atau "naif" bahwa setiap kata dalam sebuah badan teks berdiri sendiri dan tidak berhubungan dengan kata yang lain, walaupun pada kebanyakan kasus hubungan antar kata dalam suatu badan teks bisa sangat kompleks, dan sering kali suatu kata memiliki hubungan dengan kata yang lain sehingga tidak dapat dianalisis kata tersebut secara independen. Naive Bayes dalam sentimen analisis menggunakan probabilitas kemunculan kata-kata dalam kategori negatif atau positif untuk menghitung probabilitas suatu teks apakah masuk ke dalam kategori positif atau negative [6].

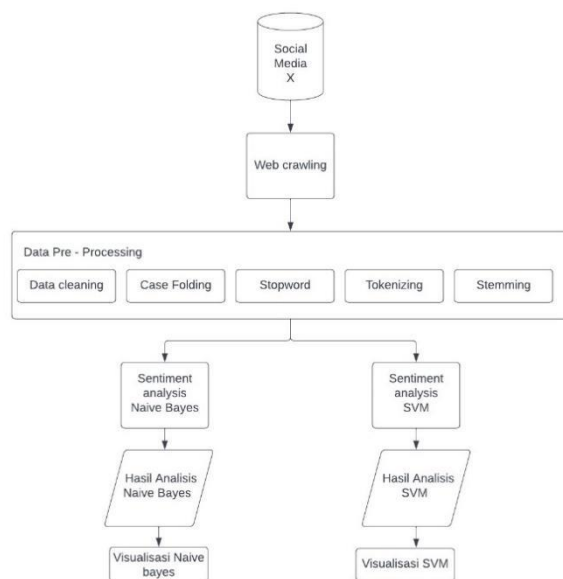
Support Vector Machine (SVM) juga adalah salah satu metode yang dapat digunakan untuk melakukan analisis sentimen. SVM bekerja dengan cara mencari garis atau batas terbaik yang memisahkan sentimen positif dan negatif dalam sebuah badan teks [5] [9]. SVM dapat bekerja dengan baik pada struktur kalimat yang kompleks, dimana kata pada teks tersebut memiliki hubungan. Walau begitu, metode SVM memakan lebih banyak waktu untuk belajar dari data.



Penelitian ini akan membandingkan keakuratan dari metode Naive Bayes dan SVM dalam melakukan analisis sentimen, dengan data yang digunakan adalah hasil scraping dari posting masyarakat Indonesia terhadap akun Kementerian Kesehatan Republik Indonesia, yaitu @KemenkesRI.

Berbagai penelitian telah dilakukan untuk analisis sentimen menggunakan algoritma yang berbeda. Dalam penelitian oleh Samsir, algoritma Naive Bayes digunakan untuk klasifikasi opini masyarakat terhadap pembelajaran daring dengan akurasi 97,15% [6]. Merinda Lestandy, membandingkan algoritma RNN (TF-IDF) dan Naive Bayes (TF-IDF) pada 5000 posting, dengan hasil akurasi tertinggi oleh RNN sebesar 97,77% dan Naive Bayes sebesar 80% [2]. Elisabet Sinta Romaito, menganalisis sentimen pilkada menggunakan SVM dan Naive Bayes, dengan akurasi masing-masing 80,7% dan 81,7% [4]. Bobby Pranata, menggunakan VADER untuk labeling dan SVM untuk klasifikasi, menghasilkan akurasi 60% [10]. Sabar Sautomo membandingkan SVM dan Naive Bayes pada data Twitter tentang mudik selama pandemi, dengan SVM unggul pada akurasi 87%, sedangkan Naive Bayes sebesar 82% [5].

II. METODE PENELITIAN



Gambar 1. Flowchart Penelitian

A. Web Crawling

Crawling adalah proses pengunduhan dan pengumpulan data dari website [4]. *Crawling* menjadi tahap awal penelitian ini. Pada penelitian ini, *crawling* akan dijalankan pada media sosial X. Data yang diambil adalah data *posting* yang terdapat "@KemenkesRI". Kata-kata seperti "Malaysia," "Malay," dan "nak" dibuang dikarenakan pada awal *crawling* terdapat banyak posting dari warga negara Malaysia.

Pada penelitian ini, digunakan *tweet-harvest* untuk *crawling* data dari X.

B. Data Cleaning

Data Cleaning adalah proses pada analisis sentimen yang bertujuan untuk mengurangi noise pada dataset, seperti URL, "@" untuk mention, dan beberapa karakter lainnya yang tidak berguna untuk analisis sentiment [9]. Dalam Data Cleaning juga dilakukan penghapusan baris duplikat. Berikut adalah proses Data Cleaning.

Sebelum Data Cleaning	Setelah Data Cleaning
Vaksinasi Covid ini sebenarnya Bagaimana ya?Apa Hrs Vaksinasi sesuai Asal KTP?Jika demikian ,seandainya yg bersangkutan bekerja diluar daerah asal harus kembali ke daerah asal KTP untuk Vaksinasi,Kok Ribet?Serius nanya @jokowi ,@KemenkesRI	Vaksinasi Covid ini sebenarnya Bagaimana ya?Apa Hrs Vaksinasi sesuai Asal KTPJika demikian seandainya yg bersangkutan bekerja diluar daerah asal harus kembali ke daerah asal KTP untuk VaksinasiKok RibetSerius nanya

C. Case Folding

Selanjutnya, dilakukan *Case Folding*. *Case folding* adalah langkah dalam pengolahan data yang bertujuan untuk mengubah kata dalam sebuah kalimat menjadi huruf kecil, karena tidak semua kata dalam kalimat konsisten dalam penggunaan huruf kapital [9]. Maka dari itu, semua kata diubah menjadi *lowercase*

Sebelum Case Folding	Setelah Case Folding
Semua merasa tenang dng sudah di vaksin Tapi lupa dengan cara penanggulangan pertama dng penyemprotan disinfektan besarn dng mobilitas tinggi NASIONAL Polri TNI ad	semua merasa tenang dng sudah di vaksin tapi lupa dengan cara penanggulangan pertama dng penyemprotan disinfektan besarn dng mobilitas tinggi nasional polri tni ad

D. Stopword Removal

Selanjutnya adalah menghapus *stopwords* yang terdapat pada teks. *Stopwords* removal berfungsi untuk menghapus kata-kata pada teks yang dianggap tidak penting [9].

Sebelum Stopwords Removal	Setelah Stopwords Removal
kapan sih kita ini yang masyarakat umum divaksin nakes udah pejabat udah tni polri udah	kapan sih ini masyarakat umum divaksin nakes udah pejabat udah tni polri udah lansia tokoh agama

- Df 1 : 289 Negative dan 239 Positive
- Df 2 : 240 Negative dan 288 Positive
- Df 3 : 235 Negative dan 293 Positive
- Df 4 : 224 Negative dan 304 Positive
- Df 5 : 221 Negative dan 307 Positive
- Df 6 : 249 Negative dan 279 Positive
- Df 7 : 261 Negative dan 267 Positive
- Df 8 : 245 Negative dan 283 Positive

E.Tokenizing

Proses selanjutnya adalah dilakukan *tokenizing*. *Tokenizing* adalah proses di mana sebuah kalimat dipecah per kata yang disebut token [2] [5].

Sebelum Tokenizing	Setelah Tokenizing
kapan sih kita ini yang	['kapan', 'sih', 'ini',
masyarakat umum	'masyarakat', 'umum',
divaksin nakes udah	'divaksin', 'nakes', 'udah',
pejabat udah tnipolri udah	'pejabat', 'udah', 'tnipolri',
lansia tokoh agama	'udah', 'lansia', 'tokoh',
masyarakat udah yang	'agama', 'masyarakat',
masyarakat umum	'udah', 'masyarakat',
pedagang pekerja lepas	'umum', 'pedagang',
pekerja kontrak non asn	'pekerja', 'lepas', 'pekerja',
kapan	'kontrak', 'non', 'asn',
	'kapan']

F. Stemming

Selanjutnya, dilakukan *stemming*. *Stemming* adalah proses analisis sentimen yang bertujuan untuk mengubah kata berimbuhan menjadi kata dasarnya [1] [2]. Pada penelitian ini, untuk proses stemming digunakan *library Sastrawi*,

Kemudian, seluruh hasil preprocessing ini akan disimpan ke dalam sebuah file CSV baru untuk selanjutnya ke proses labeling. File CSV ini bernama "hasil_TextPreProcessing.csv"

Sebelum Stemming	Setelah Stemming
['aku', 'keluarga', 'ku', 'kapan', 'dapat', 'vaksin', 'covid', 'gratisnva']	aku keluarga ku kapan dapat vaksin covid gratis

G.Text Labelling

Proses *Text Labelling* merupakan proses dalam Analisis Sentimen yang bertujuan untuk memberi label positif atau negatif pada setiap sentimen sesuai dengan sifat dari sentimen tersebut.

Pada penelitian ini, untuk proses *labelling* digunakan VADER (*Valence Aware Dictionary and Sentiment Reasoner*). VADER bekerja dengan cara memberikan nilai pada sebuah teks, dan berdasarkan nilai tersebut ditentukan sentimen tersebut negatif atau positif [12].

Labelling dilakukan dengan membagi file sebelumnya yaitu `hasil_TextPreProcessing.csv` menjadi 8 file csv baru dan menyimpannya ke variable dengan nama `df1` sampai dengan `df8` untuk dijalankan proses *labelling* secara parallel, berikut adalah hasilnya.

H. Penerapan Naïve Bayes

Untuk *Naïve Bayes* data dipersiapkan dengan membagi dataset yang telah ada menjadi data *training* dan *testing*, di mana 80% digunakan untuk *training* dan 20% untuk *testing*. Proses berikutnya adalah transformasi teks menjadi representasi numerik menggunakan *TfidfVectorizer*, yang membantu dalam menangkap fitur penting dari teks dengan memprioritaskan kata-kata yang lebih informatif.

Setelah itu, dilakukan *feature selction* dengan Chi-Square, yang menyaring fitur untuk hanya menyisakan 1000 fitur teratas yang paling relevan. Model Naive Bayes, yaitu MultinomialNB, diinisialisasi dengan parameter alpha yang disetel untuk memperhalus model. lalu pelatihan model dilakukan dengan *cross validation* dengan *K-Fold*, di mana data dibagi menjadi lima *fold*. Di setiap iterasi, model dilatih pada *training fold* dan diuji pada *testing fold*, dan hasil prediksi serta label disimpan untuk evaluasi selanjutnya.

Setelah *cross validation* selesai, *confusion matrix* dibuat untuk menunjukan keakuratan model, yang divisualisasikan dalam bentuk heatmap menggunakan *seaborn*, memberikan gambaran jelas tentang jumlah prediksi yang benar dan salah untuk masing-masing kelas (negatif dan positif). Kemudian, *classification report* dihasilkan untuk mengetahui nilai *precision*, *recall*, dan *F1-score*. Terakhir, akurasi keseluruhan model setelah *cross-validation* dihitung dan ditampilkan.

I.Penerapan SVM

Untuk penerapan SVM Hampir sama dengan *Naïve Bayes*, Data dibagi menjadi data *training* dan *testing* menggunakan fungsi `train_test_split`, dengan 80% data digunakan untuk *training* dan 20% untuk *testing* sama seperti pada *naïve bayes*. Setelah itu, teks diubah menjadi representasi numerik menggunakan `TfidfVectorizer`

Selanjutnya, *feature selection* dilakukan dengan *Chi-Square* untuk menyaring dan hanya menyisakan 1000 fitur paling relevan. Lalu untuk *parameter grid SVM* mencakup variasi kernel, nilai parameter C, dan gamma. *Grid search* dilakukan dengan menggunakan *GridSearchCV*, yang melakukan *cross-validation* untuk menemukan kombinasi parameter terbaik berdasarkan performa model.

Setelah mendapatkan parameter terbaik dari *grid search*, model SVM diinisialisasi dengan parameter tersebut. Proses *cross-validation* diterapkan untuk mengevaluasi model dengan membagi data menjadi lima *fold*, di mana model dilatih dan diuji pada setiap *fold*. Hasil



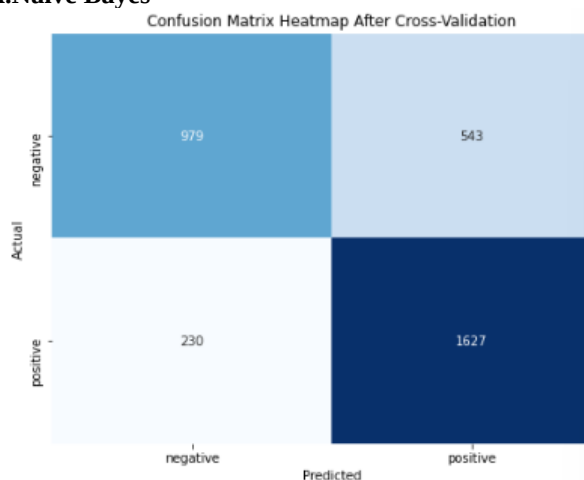
prediksi dan label sebenarnya dikumpulkan untuk analisis lebih lanjut.

Setelah *cross-validation* selesai, confusion matrix dibuat untuk menilai kinerja model, yang divisualisasikan dalam bentuk heatmap. Lalu dibuat Laporan klasifikasi untuk mendapatkan nilai *precision*, *recall*, dan *F1-score*. Terakhir, akurasi keseluruhan model dihitung dan ditampilkan.

III. HASIL DAN PEMBAHASAN

Setelah dilakukan semua proses yang dijelaskan sebelumnya maka didapatkan hasil model *naïve bayes* dan juga *svm*.

A. Naïve Bayes



Gambar 2. Confusion matrix Naïve Bayes

Classification Report:				
	precision	recall	f1-score	support
negative	0.81	0.64	0.72	1522
positive	0.75	0.88	0.81	1857
accuracy			0.77	3379
macro avg	0.78	0.76	0.76	3379
weighted avg	0.78	0.77	0.77	3379

Overall Accuracy after Cross-Validation: 0.7712340929269015

Gambar 3. Hasil Evaluasi Naïve Bayes

Dari visualisasi di atas, dapat dilihat bahwa akurasi dari model *Naïve Bayes* adalah kurang lebih sebesar 77%, atau 0,77. Hasil yang diperoleh di sini berupa *test time classification report*, yang mana berarti yang didapatkan adalah nilai *accuracy*, *precision*, *recall*, dan *f1-score*, serta *confusion matrix*.

Untuk perhitungan *accuracy* dapat dihitung dengan menggunakan rumus yang dapat dilihat di bawah ini :

$$\text{Accuracy} = \frac{TN+TP}{\text{Total data}}$$

$$\text{Accuracy} = \frac{1627+979}{3379} = 0,77$$

Dari perhitungan diatas diketahui bahwa nilai akurasi yang didapat dari model *Naïve Bayes* adalah sebesar 0,77%.

Untuk perhitungan *precision* kelas *negative*, dapat dihitung dengan rumus dibawah :

$$\text{Precision} = \frac{TN}{TN+FN}$$

$$\text{Precision} = \frac{979}{979+230} = 0,81$$

$$\text{Recall} = \frac{TN}{TN+FP}$$

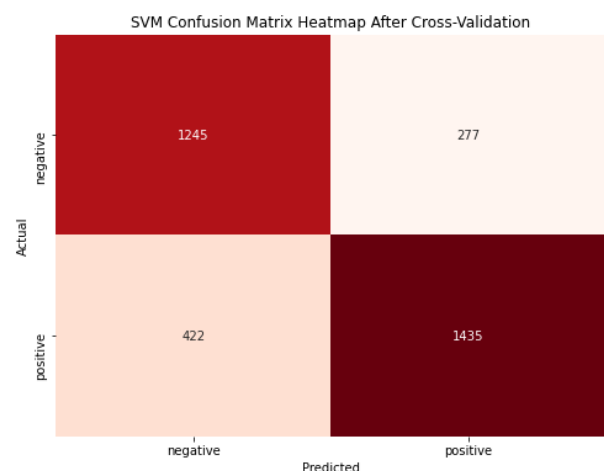
$$\text{Recall} = \frac{979}{979+543} = 0,64$$

$$\text{F1-Score} = \frac{2 \times (\text{Recall} \times \text{Precision})}{\text{Recall} + \text{Precision}}$$

$$\text{F1-Score} = \frac{2 \times (0,64 \times 0,81)}{0,64 + 0,81} = \frac{1,0368}{1,45} = 0,72$$

Dengan hasil diatas didapatkan bahwa, nilai *precision* adalah 0,81, dengan *recall* 0,64, dan *f1-score* sebesar 0,72.

B. SVM



Gambar 2. Confusion matrix SVM

Classification Report:				
	precision	recall	f1-score	support
negative	0.75	0.82	0.78	1522
positive	0.84	0.77	0.80	1857
accuracy			0.79	3379
macro avg	0.79	0.80	0.79	3379
weighted avg	0.80	0.79	0.79	3379

Overall Accuracy after Cross-Validation: 0.7931340633323468

Gambar 3. Hasil Evaluasi SVM

Pada visualisasi diatas dapat dilihat bahwa SVM mendapatkan hasil *accuracy* 79 % dimana *accuracy* ini lebih baik jika dibandingkan dengan *Naïve Bayes* yang



mendapatkan hasil accuracy 77% untuk perhitungan *accuracy*, *precision*, *recall*, dan *f1-score* dapat dilihat di bawah

$$\text{Accuracy} = \frac{TN+TP}{\text{Total data}}$$

$$\text{Accuracy} = \frac{1245+1435}{3379} = 0,79$$

$$\text{Precision} = \frac{TP}{TP+FP}$$

$$\text{Precision} = \frac{1435}{1435+277} = 0,84$$

$$\text{Recall} = \frac{TP}{TP+FN}$$

$$\text{Recall} = \frac{1435}{1435+422} = 0,77$$

$$\text{F1-Score} = \frac{2 \times (\text{Recall} \times \text{Precision})}{\text{Recall} + \text{Precision}}$$

$$\text{F1-Score} = \frac{2 \times (0,77 \times 0,84)}{0,77 + 0,84} = \frac{1,2936}{1,61} = 0,79$$

Dengan hasil diatas, didapatkan bahwa nilai *precision* adalah 0,84, dengan *recall* 0,77, dan *f1-score* sebesar 0,79.

IV. KESIMPULAN

Kesimpulan yang dapat diambil dari penelitian analisis sentimen terhadap kementerian Kesehatan Indonesia selama pandemi covid-19 ini adalah:

1. Dalam penelitian analisis sentimen ini, data yang digunakan berasal dari media sosial X, dan data diambil selama masa pandemi COVID-19. Data yang diambil merupakan data yang bagus untuk digunakan sebagai penelitian analisis sentimen karena X adalah media sosial yang biasa digunakan penggunaannya untuk menyuarakan pendapat tentang sesuatu. Namun, data ini masih harus dilakukan preprocessing dan labeling sebelum dapat diterapkan algoritma klasifikasi seperti Naïve Bayes dan SVM.

2. Dalam penelitian ini, digunakan dua algoritma, yaitu Naïve Bayes dan SVM. Didapatkan bahwa algoritma SVM mendapat angka akurasi yang lebih baik dibandingkan dengan Naïve Bayes, dengan SVM mendapatkan nilai akurasi sebesar 79%, sementara Naïve Bayes mendapatkan nilai akurasi sebesar 77%.

UCAPAN TERIMA KASIH

Penulis ingin mengucapkan terima kasih kepada dosen pembimbing, Ibu Dian Pratiwi, ST, MTI dan juga Ibu Syandra Sari, M.Kom, yang telah bersedia menjadi pembimbing saya dan telah membimbing saya dalam pembuatan laporan ini dengan memberikan arahan dan juga petunjuk untuk menyelesaikan penelitian ini, Penulis juga ingin berterimakasih kepada para dosen yang telah

memberikan ilmu selama masa perkuliahan hingga penulis dapat membuat penelitian ini.

Penulis juga ingin berterimakasih kepada keluarga yang telah memberikan dukungan secara emosional, memberi semangat dan waktu, serta memberikan dukungan materil yang sangat membantu penulis dalam proses penyelesaian penelitian ini.

Penulis sadar bahwa tanpa dukungan moral yang baik, akan sangat sulit bagi penulis untuk tetap semangat dalam menjalani perkuliahan dan menyelesaikan penelitian ini. Oleh karena itu, penulis juga ingin mengucapkan terimakasih kepada kekasih penulis, Thalia Teana Cahyani, yang telah memberikan dukungan moral, telah menemani penulis selama masa kuliah dari mahasiswa baru hingga menjadi mahasiswa tingkat akhir yang sedang berjuang untuk membuat penelitian ini dan telah menjadi *support system* yang baik untuk penulis sehingga penulis dapat terus semangat dan positif dalam mengerjakan penelitian ini.

Terakhir, penulis ingin berterima kasih kepada teman-teman seperjuangan yang telah membantu penulis dalam menyusun penelitian ini dengan rasa kebersamaan yang kuat.

V. REFERENSI

- [1] Kurniawan, I., & Susanto, A. (2019). Implementasi Metode K-Means dan Naïve Bayes Classifier untuk Analisis Sentimen Pemilihan Presiden (Pilpres) 2019. *Jurnal Eksplora Informatika*, 9(1), 1-10.
- [2] Lestandy, M., Abdurrahim, A., & Syafa'ah, L. (2021). Analisis Sentimen Tweet Vaksin COVID-19 Menggunakan Recurrent Neural Network dan Naïve Bayes. *Jurnal RESTI (Rekayasa Sistem Dan Teknologi Informasi)*, 5(4), 802-808.
- [3] Ananda, F. D., & Pristyanto, Y. (2021). Analisis Sentimen Pengguna Twitter Terhadap Layanan Internet Provider Menggunakan Algoritma Support Vector Machine. *MATRIK: Jurnal Manajemen, Teknik Informatika dan Rekayasa Komputer*, 20(2), 407-416.
- [4] Romaito, E. S., Anam, M. K., & Rahmaddeni, A. N. U. PERBANDINGAN ALGORITMA SVM DAN NBC DALAM ANALISA SENTIMEN PILKADA PADA TWITTER.
- [5] Sautomo, S., Hafidz, N., Achyani, Y. E., & Gata, W. (2020). Sentiment Analysis Due to "Mudik" Prohibited of COVID-19 Through Twitter. *JITK (Jurnal Ilmu Pengetahuan dan Teknologi Komputer)*, 6(1), 7-12.
- [6] Samsir, S., Ambiyar, A., Verawardina, U., Edi, F., & Watrianthos, R. (2021). Analisis Sentimen Pembelajaran Daring Pada Twitter di Masa Pandemi COVID-19 Menggunakan Metode Naïve Bayes. *Jurnal Media Informatika Budidarma*, 5(1), 157-163.
- [7] Prabowo, W. A., & Azizah, F. (2020). Sentiment analysis for detecting cyberbullying using TF-



- IDF and SVM. Jurnal RESTI (Rekayasa Sistem dan Teknologi Informasi), 4(6), 1142-1148.
- [8] Hafidz, N., & Liliana, D. Y. (2021). Klasifikasi Sentimen pada Twitter Terhadap WHO Terkait Covid-19 Menggunakan SVM, N-Gram, PSO. Jurnal RESTI (Rekayasa Sistem Dan Teknologi Informasi), 5(2), 213-219.
- [9] Widowati, T. T. ANALISIS SENTIMEN TWITTER TERHADAP TOKOH PUBLIK DENGAN ALGORITMA NAIVE BAYES DAN SUPPORT VECTOR.
- [10] Pranata, B. (2023). Support Vector Machine untuk Sentiment Analysis Bakal Calon Presiden Republik Indonesia 2024. Indonesian Journal of Computer Science, 12(3).
- [11] Faesal, A., Muslim, A., Ruger, A. H., & Kusrini, K. (2020). Sentimen analisis pada data tweet pengguna twitter terhadap produk penjualan toko online menggunakan metode k-means. MATRIK: Jurnal Manajemen, Teknik Informatika dan Rekayasa Komputer, 19(2), 207-213.
- [12] R. Rasiban and S. . Riyadi, "Analisis Sentimen Opini Masyarakat Terhadap Stadion Jakarta Internasional Stadium (Jis) Pada Twitter Dengan Perbandingan Metode Naive Bayes Dan Support Vector Machine", *SAINTEK*, vol. 5, no. 3, pp. 1010-1017, Apr. 2024.
- [13] I. B. M. Prawira, B. Solihah, and S. Sari, "Analysis of Oil Sentiment Sentiments on Twitter using Support Vector Machine," *Intelmatiks*, vol. 3, no. 1, pp. 13-16, Jan.-Jun. 2023.