

Analisis Perbandingan Model Machine Learning untuk Klasifikasi Sentimen Isu Kesehatan Mental

Yuwan Jumaryadi

Program Studi Sistem Informasi, Universitas Mercu Buana, Indonesia

yuwan.jumaryadi@mercubuana.ac.id

Abstrak: Perkembangan media sosial telah menjadikan platform digital sebagai ruang utama bagi masyarakat untuk berbagi pandangan mengenai isu kesehatan, termasuk kesehatan mental. Penelitian ini bertujuan untuk menganalisis sentimen masyarakat terhadap isu kesehatan mental di platform X (sebelumnya Twitter). Data dikumpulkan sebanyak 5.457 tweet pada periode April hingga Mei 2025. Proses pelabelan dilakukan secara otomatis menggunakan pendekatan berbasis leksikon, yang menghasilkan 3.415 data positif (69,61%) dan 1.491 data negatif (30,39%). Setelah melalui tahapan *preprocessing*—meliputi *cleaning*, *case folding*, normalisasi, *tokenizing*, *stopword removal*, dan *stemming*—data diklasifikasikan menggunakan empat algoritma *machine learning*: *Support Vector Machine (SVM)*, *Random Forest*, *Naïve Bayes*, dan *Logistic Regression*. Berdasarkan hasil evaluasi, algoritma SVM menunjukkan performa klasifikasi terbaik dengan akurasi sebesar 81,57%, diikuti oleh *Logistic Regression* (81,36%), *Random Forest* (80,04%), dan *Naïve Bayes* (70,44%). Hasil penelitian ini menunjukkan bahwa kombinasi pelabelan berbasis leksikon dan algoritma *machine learning* efektif dalam memetakan persepsi publik terhadap isu kesehatan mental, yang dapat menjadi acuan bagi pemangku kepentingan dalam merumuskan kebijakan serta strategi komunikasi kesehatan yang lebih responsif dan tepat sasaran.

Kata Kunci: Analisis Sentimen; Kesehatan Mental; Pembelajaran Mesin; Support Vector Machine; TF-IDF.

Abstract: Berikut adalah versi bahasa Inggris (*Abstract*) dari abstrak tersebut:
ABSTRACT

The rapid development of digital technology has transformed social media into a primary space for the public to share opinions on health issues, including mental health. This study aims to analyze public sentiment toward mental health issues on platform X (formerly Twitter). Data consisting of 5,457 tweets was collected from April to May 2025. Automated data labeling was performed using a lexicon-based approach, resulting in 3,415 positive data (69.61%) and 1,491 negative data (30.39%). After undergoing the preprocessing stages—including cleaning, case folding, normalization, tokenizing, stopwords removal, and stemming—the data was classified using four machine learning algorithms: Support Vector Machine (SVM), Random Forest, Naïve Bayes, and Logistic Regression. Based on the evaluation results, the SVM algorithm demonstrated the best classification performance with an accuracy of 81.57%, followed by Logistic Regression (81.36%), Random Forest (80.04%), and Naïve Bayes (70.44%). The findings indicate that combining lexicon-based labeling with machine learning algorithms is highly effective in mapping public perceptions

of mental health issues, which can serve as a reference for stakeholders in formulating more responsive and targeted health policies and communication strategies.

Keywords: Sentiment Analysis; Mental Health; Machine Learning; Support Vector Machine; TF-IDF.

1. PENDAHULUAN

Di tengah pesatnya perkembangan teknologi digital, media sosial telah berkembang menjadi salah satu sarana utama bagi masyarakat untuk bertukar informasi, berbagi pengalaman, serta menyampaikan pendapat mengenai berbagai topik, termasuk isu-isu Kesehatan [1]. Berbagai platform, seperti Twitter, Facebook, dan Instagram, memberikan ruang bagi pengguna untuk mengemukakan pandangan mereka terkait kebijakan kesehatan, penyakit, maupun kualitas layanan kesehatan yang diterima [2]. Informasi yang tersebar melalui media sosial juga kerap dijadikan rujukan oleh masyarakat dalam menentukan keputusan yang berkaitan dengan Kesehatan [3]. Oleh sebab itu, analisis terhadap sentimen yang berkembang di media sosial menjadi penting untuk memperoleh gambaran mengenai persepsi publik terhadap berbagai isu Kesehatan [4].

Analisis sentimen merupakan salah satu teknik dalam Natural Language Processing (NLP) yang bertujuan untuk mengidentifikasi, mengekstraksi, dan mengelompokkan opini atau emosi yang terkandung dalam suatu teks [5], [6]. Melalui teknik ini, respons masyarakat terhadap suatu isu kesehatan dapat diklasifikasikan ke dalam kategori positif, da negatif. Informasi yang dihasilkan dari analisis sentimen dapat dimanfaatkan oleh pemerintah, institusi kesehatan, maupun penyedia layanan medis sebagai dasar dalam merumuskan kebijakan serta menyusun strategi komunikasi yang lebih efektif [7]. Berbagai algoritma machine learning telah banyak digunakan dalam analisis sentimen, di antaranya Support Vector Machine (SVM), K-Nearest Neighbor (K-NN), Naïve Bayes, dan Random Forest.

Penerapan analisis sentimen pada bidang kesehatan memiliki peranan penting dalam memahami persepsi, sikap, dan respons masyarakat terhadap layanan maupun kebijakan Kesehatan [8]. Sentimen yang disampaikan masyarakat mengenai rumah sakit, klinik, ataupun layanan kesehatan berbasis daring dapat menjadi indikator kualitas pelayanan yang diberikan [9]. Analisis terhadap ulasan dan umpan balik pasien memungkinkan peneliti maupun penyedia layanan kesehatan untuk mengidentifikasi aspek-aspek pelayanan yang masih memerlukan perbaikan. Selain itu, analisis sentimen juga dapat dimanfaatkan sebagai alat untuk memantau tingkat kepercayaan masyarakat terhadap institusi kesehatan dan layanan medis [10]. Dominasi sentimen negatif dapat mengindikasikan adanya ketidakpuasan ataupun penurunan kepercayaan masyarakat yang memerlukan tindak lanjut dari pihak terkait [11]. Sebaliknya, dominasi sentimen positif dapat menjadi bahan evaluasi untuk mengetahui faktor-faktor yang berkontribusi terhadap meningkatnya kepuasan masyarakat terhadap layanan kesehatan [4].

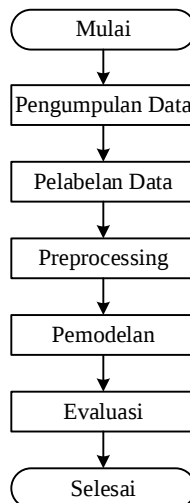
Salah satu pendekatan yang banyak digunakan dalam analisis sentimen adalah pendekatan berbasis leksikon (lexicon-based approach), yaitu metode yang menentukan polaritas sentimen berdasarkan kecocokan kata dalam teks dengan kamus kata yang telah diberi label positif maupun negatif [12], [13], [14]. Di sisi lain, berbagai penelitian juga telah memanfaatkan algoritma machine learning klasik, seperti Naïve Bayes, Support Vector Machine (SVM), dan Random Forest, untuk mengklasifikasikan sentimen masyarakat terhadap beragam isu kesehatan. Data yang digunakan umumnya berasal dari media sosial, seperti Twitter, Facebook, dan Instagram, untuk mengkaji persepsi masyarakat terhadap berbagai topik kesehatan, termasuk pandemi COVID-19, program vaksinasi, serta penyakit menular lainnya.

Berbagai penelitian terdahulu telah menerapkan algoritma Support Vector Machine (SVM), Naïve Bayes, maupun Decision Tree dalam analisis sentimen di bidang kesehatan.

Salah satu contohnya adalah penelitian Müller dan Salathé yang menganalisis sentimen masyarakat terhadap vaksin COVID-19 melalui data Twitter [15]. Meskipun demikian, metode-metode tersebut masih menghadapi kendala dalam menangkap konteks serta nuansa bahasa informal yang lazim digunakan oleh pengguna media sosial. Oleh karena itu, penelitian ini memanfaatkan pendekatan berbasis leksikon menggunakan VADER untuk melakukan pelabelan sentimen secara otomatis sebelum data diproses lebih lanjut menggunakan algoritma klasifikasi machine learning. Pendekatan ini diharapkan mampu menghasilkan pelabelan sentimen yang lebih konsisten sekaligus mendukung proses klasifikasi sentimen terhadap isu kesehatan secara lebih efektif.

2. METODE PENELITIAN

Gambar 1 menunjukkan tahapan penelitian yang dilakukan, dengan penjelasan sebagai berikut.



Gambar 1. Tahapan Penelitian

2.1. Pengumpulan Data

Pengumpulan data pada penelitian ini dilakukan melalui platform X (sebelumnya Twitter) menggunakan kata kunci "Kesehatan mental" untuk memperoleh tweet yang relevan dengan topik penelitian. Data dikumpulkan selama periode 4 April 2025 hingga 28 Mei 2025 dengan tujuan memperoleh persepsi dan opini masyarakat mengenai isu kesehatan mental. Seluruh data yang digunakan merupakan tweet yang tersedia secara publik dan diperoleh melalui teknik *web crawling* sesuai dengan ketentuan yang berlaku.

2.2. Pelabelan Data

Pelabelan sentimen pada penelitian ini dilakukan secara otomatis menggunakan metode VADER (*Valence Aware Dictionary and sEntiment Reasoner*), yaitu metode analisis sentimen berbasis leksikon yang dirancang untuk mengidentifikasi polaritas sentimen pada teks informal, khususnya unggahan di media sosial. VADER menghasilkan skor sentimen berupa nilai *compound* dengan rentang -1 yang menunjukkan sentimen sangat negatif hingga +1 yang menunjukkan sentimen sangat positif. Berdasarkan skor tersebut, data diberi label positif apabila nilai *compound* $\geq 0,05$ dan negatif apabila nilai *compound* $\leq -0,05$. Sementara itu, data yang memiliki nilai *compound* di antara -0,05 hingga 0,05 tidak disertakan dalam proses analisis karena tidak termasuk ke dalam kategori sentimen positif maupun negatif. Pendekatan ini memungkinkan proses pelabelan dilakukan secara otomatis, konsisten, dan efisien tanpa memerlukan intervensi manual.

2.3. Preprocessing

Preprocessing merupakan proses mengubah data teks yang masih tidak terstruktur menjadi data yang terstruktur sehingga siap untuk dianalisis [16]. Tahapan *preprocessing* yang dilakukan pada penelitian ini meliputi sebagai berikut.

a. Cleaning

Cleaning merupakan proses pembersihan data dengan menghapus karakter yang tidak diperlukan, seperti tanda baca, angka, karakter non-huruf, *mention*, alamat email, tautan (*URL*), serta simbol-simbol lain yang tidak memiliki makna dalam analisis. Tahapan ini bertujuan menghasilkan data yang lebih bersih dan berkualitas.

b. Case Folding

Case folding merupakan proses mengubah seluruh huruf dalam teks menjadi huruf kecil (*lowercase*) sehingga tidak terdapat perbedaan antara huruf kapital dan huruf kecil selama proses analisis.

c. Normalisasi

Normalisasi bertujuan memperbaiki kata-kata yang mengandung huruf berulang, penulisan tidak baku, singkatan, maupun kesalahan pengetikan (*typo*) sehingga kata-kata tersebut sesuai dengan bentuk bakunya. Tahapan ini dilakukan untuk meningkatkan kemampuan sistem dalam mengenali makna kata dan memperbaiki kualitas representasi teks yang akan dianalisis.

d. Tokenizing

Tokenizing merupakan proses memecah kalimat menjadi kata-kata tunggal (*token*) sehingga setiap kata dapat diproses secara individual pada tahapan berikutnya.

e. Stopword Removal

Stopword removal merupakan proses menghapus kata-kata umum yang tidak memberikan informasi penting terhadap makna kalimat, seperti *yang*, *dan*, *dari*, serta kata-kata lainnya yang terdapat pada daftar *stopword*. Pada penelitian ini, proses tersebut menggunakan daftar *stopword* yang tersedia pada pustaka Sastrawi.

f. Stemming

Stemming merupakan proses mengubah kata yang masih mengandung imbuhan menjadi kata dasar dengan menghilangkan awalan, akhiran, sisipan, maupun konfiks. Proses stemming dilakukan menggunakan pustaka Sastrawi.

2.4. Pemodelan

Dataset pada penelitian ini dibagi menjadi data latih (*training set*) dan data uji (*testing set*) dengan rasio 80:20. Penelitian ini menggunakan algoritma Naïve Bayes Classifier dan Support Vector Machine (SVM) sebagai model klasifikasi sentimen.

Pada tahap ini, model terlebih dahulu dilatih menggunakan data latih untuk membentuk pola pembelajaran (*learning model*). Selanjutnya, model yang telah terbentuk digunakan untuk mengklasifikasikan data uji ke dalam kelas sentimen positif, dan negatif. Hasil pemodelan kemudian dievaluasi menggunakan confusion matrix dan classification report yang menampilkan nilai precision, recall, F1-score, support, serta accuracy.

2.5. Evaluasi

Evaluasi model dilakukan menggunakan Confusion Matrix, yaitu metode yang digunakan untuk mengukur kinerja model klasifikasi dengan membandingkan hasil prediksi model terhadap label sebenarnya. Berdasarkan confusion matrix, diperoleh nilai True Positive (TP), True Negative (TN), False Positive (FP), dan False Negative (FN) yang selanjutnya digunakan untuk menghitung metrik evaluasi berupa akurasi, *precision*, recall, dan F1-score. Representasi confusion matrix yang digunakan dalam penelitian ini ditunjukkan pada Tabel 1.

Tabel 1. Confusion Matrix

		<i>Actual Values</i>	
		Positif	Negatif
Prediksi	Positif	TP	FN ₁
	Negatif	FP ₁	TN

Nilai akurasi digunakan untuk mengukur tingkat ketepatan model dalam mengklasifikasikan seluruh data. Akurasi dihitung berdasarkan perbandingan antara jumlah prediksi yang benar terhadap jumlah seluruh data uji sebagaimana ditunjukkan pada Persamaan (1).

$$\text{Akurasi} = \frac{TP + TN}{TP + TN + FP + FN} \quad (1)$$

Selanjutnya, presisi digunakan untuk mengukur tingkat ketepatan model dalam memprediksi kelas positif. Nilai presisi dihitung berdasarkan perbandingan antara jumlah data yang diprediksi positif secara benar dengan seluruh data yang diprediksi sebagai positif, sebagaimana ditunjukkan pada Persamaan (2).

$$\text{Precision} = \frac{TP}{TP + FP} \quad (2)$$

Kemampuan model dalam mengenali seluruh data yang benar-benar termasuk ke dalam kelas positif diukur menggunakan recall. Nilai recall dihitung berdasarkan perbandingan antara jumlah data positif yang berhasil diprediksi dengan seluruh data yang sebenarnya berlabel positif sebagaimana ditunjukkan pada Persamaan (3).

$$\text{Recall} = \frac{TP}{TP + FN} \quad (3)$$

Untuk memperoleh ukuran performa yang mempertimbangkan keseimbangan antara presisi dan *recall*, digunakan F1-score. Nilai F1-score dihitung sebagai rata-rata harmonik antara presisi dan *recall* sebagaimana ditunjukkan pada Persamaan (4).

$$\text{F1-Score} = 2 \times \frac{\text{Presisi} \times \text{Recall}}{\text{Presisi} + \text{Recall}} \quad (4)$$

3. HASIL DAN PEMBAHASAN

Pengumpulan Data

Penelitian ini memanfaatkan dataset yang diperoleh dari platform X (sebelumnya Twitter) dengan menggunakan kata kunci "kesehatan mental". Proses pengumpulan data dilakukan melalui teknik web crawling pada periode tahun 2025, sehingga diperoleh sebanyak 5457 tweet yang berkaitan dengan isu kesehatan mental. Dataset tersebut kemudian digunakan sebagai bahan analisis dalam penelitian ini.

Pelabelan

Pelabelan data pada penelitian ini dilakukan secara otomatis menggunakan VADER (Valence Aware Dictionary and sEntiment Reasoner), yaitu metode analisis sentimen berbasis leksikon yang menentukan polaritas sentimen berdasarkan skor compound pada setiap teks [17]. Nilai compound berada pada rentang -1 yang menunjukkan sentimen sangat negatif hingga +1 yang menunjukkan sentimen sangat positif. Pada penelitian ini, data diklasifikasikan ke dalam dua kategori, yaitu positif untuk skor compound $\geq 0,05$ dan negatif untuk skor compound $\leq -0,05$. Sementara itu, data yang memiliki skor compound pada rentang -0,05 hingga 0,05 tidak digunakan dalam proses klasifikasi. Hasil pelabelan

menghasilkan 4.906 data yang terdiri atas 3.415 data berlabel positif (69,61%) dan 1.491 data berlabel negatif (30,39%). Distribusi hasil pelabelan ditunjukkan pada Tabel 2.

Tabel 2. Pelabelan

Label	Total	Persentase (%)
Positif	3415	69,61%
Negatif	1491	30,3%
Total	4906	100%

. Preprocessing

Data yang berhasil ditarik (*crawled*) merupakan data mentah yang belum dapat diproses untuk analisis sentimen. Oleh karena itu, diperlukan tahapan *preprocessing* untuk mendapatkan data terstruktur yang dapat digunakan dalam pembuatan model *machine learning*. Penelitian ini berhasil mendapatkan data yang siap digunakan dalam proses pengembangan model *machine learning* untuk analisis sentimen. Tabel 3 menunjukkan dataset yang diperoleh dari hasil *crawling*.

Tabel 3. Dataset

full_text
@flash_poem Gen Z ini unik.. mereka sadar pentingnya healing batasan diri dan kesehatan mental. Bukan berarti malas tp tahu kapan harus berhenti sebelum hancur. Sementara Milenial dibentuk dr zaman survive mode Tapi Gen Z? Mereka bilang: Kalau burnout siapa yang tanggung jawab?
....
....
□ Kurang belanja impor Lebih bijak atur uang keuangan prioritaskan kebutuhan Kalau punya utang dolar siapin strategi biar ga makin berat Investasi di aset yang lebih stabil https://t.co/ofX5eWssdO

Data di atas telah melalui tahap *crawling*, yang bertujuan untuk mengekstraksi data dari platform media sosial yang digunakan dalam penelitian ini, khususnya Twitter. Kata kunci yang digunakan dalam pengumpulan data di Twitter adalah "Kesehatan Mental".

Cleansing, setelah menyelesaikan tahap *case folding*, proses dilanjutkan ke tahap *cleansing* (pembersihan) untuk menghapus elemen cuitan khusus (seperti sebutan *username* dengan awalan @, tagar #, URL, emoji, dan karakter khusus lainnya). Tabel 4 menunjukkan hasil penerapan tahap *cleansing* pada tahapan *preprocessing* terhadap data yang diperoleh.

Tabel 4. Cleaning

full_text
Gen ini unik mereka sadar pentingnya healing batasan diri dan kesehatan mental Bukan berarti malas tp tahu kapan harus berhenti sebelum hancur Sementara Milenial dibentuk dr zaman survive mode Tapi Gen Mereka bilang Kalau burnout siapa yang tanggung jawab
....
....
Kurang belanja impor Lebih bijak atur uang keuangan prioritaskan kebutuhan Kalau punya utang dolar siapin strategi biar ga makin berat Investasi di aset yang lebih stabil

Case Folding, tahap ini mengubah teks pada dataset hasil *crawling* yang awalnya memiliki campuran huruf kapital menjadi huruf kecil seluruhnya (*lowercase*). Tabel 5 menunjukkan hasil penerapan tahap *Case Folding* pada *preprocessing* terhadap data yang diperoleh.

Tabel 5. Case Folding

full_text

gen ini unik mereka sadar pentingnya healing batasan diri dan kesehatan mental bukan berarti malas tp tahu kapan harus berhenti sebelum hancur sementara milenial dibentuk dr zaman survive mode tapi gen mereka bilang kalau burnout siapa yang tanggung jawab

....

....

kurang belanja impor lebih bijak atur uang keuangan prioritaskan kebutuhan kalau punya utang dolar siapin strategi biar ga makin berat investasi di aset yang lebih stabil

Tabel 5 memperlihatkan hasil dataset yang telah melalui tahap *case folding* sebelumnya. Atribut *full_text* yang merepresentasikan data cuitan dari setiap akun pengguna kini menghasilkan teks berhuruf kecil.

Normalisasi merupakan proses mengubah kata tidak baku, singkatan, penggunaan huruf berulang, serta kesalahan pengetikan (*typo*) menjadi bentuk kata baku. Tahapan ini bertujuan meningkatkan kualitas data dengan menyeragamkan penulisan kata sehingga setiap kata dapat dikenali dan diproses secara lebih akurat pada tahapan analisis berikutnya. Hasil penerapan proses normalisasi pada data penelitian ditunjukkan pada Tabel 6.

Tabel 6. Normalisasi

full_text

gen ini unik mereka sadar pentingnya healing batasan diri dan kesehatan mental bukan berarti malas tapi tahu kapan harus berhenti sebelum hancur sementara milenial dibentuk dari zaman bertahan mode tapi gen mereka bilang kalau burnout siapa yang tanggung jawab

....

....

kurang belanja impor lebih bijak atur uang keuangan prioritaskan kebutuhan kalau punya utang dolar siapin strategi biar tidak makin berat investasi di aset yang lebih stabil

Tokenizing merupakan proses memecah teks hasil normalisasi menjadi unit-unit kata (*token*), sehingga setiap kata dapat diproses secara individual pada tahapan berikutnya. Proses ini menghasilkan representasi teks dalam bentuk kumpulan token yang siap digunakan pada tahap *stopword removal* dan *stemming*. Hasil penerapan proses *tokenizing* pada data penelitian ditunjukkan pada Tabel 7.

Tabel 7. Tokenizing

full_text

['gen', 'ini', 'unik', 'mereka', 'sadar', 'pentingnya', 'healing', 'batasan', 'diri', 'dan', 'kesehatan', 'mental', 'bukan', 'berarti', 'malas', 'tapi', 'tahu', 'kapan', 'harus', 'berhenti', 'sebelum', 'hancur', 'sementara', 'milenial', 'dibentuk', 'dari', 'zaman', 'bertahan', 'mode', 'tapi', 'gen', 'mereka', 'bilang', 'kalau', 'burnout', 'siapa', 'yang', 'tanggung', 'jawab']

....

....

['kurang', 'belanja', 'impor', 'lebih', 'bijak', 'atur', 'uang', 'keuangan', 'prioritaskan', 'kebutuhan', 'kalau', 'punya', 'utang', 'dolar', 'siapin', 'strategi', 'biar', 'tidak', 'makin', 'berat', 'investasi', 'di', 'aset', 'yang', 'lebih', 'stabil']

Stopword Removal, setelah tahap *Tokenizing* selesai, dilanjutkan dengan tahap *Stopword Removal* yang berfungsi untuk menyaring dan menghapus kata-kata yang tidak

penting di dalam dokumen. Tabel 8 menunjukkan hasil penerapan tahap *Stopword Removal* pada *Preprocessing* terhadap data yang telah diperoleh.

Tabel 8. Stopword Removal

full_text
['gen', 'unik', 'sadar', 'healing', 'batasan', 'kesehatan', 'mental', 'malas', 'berhenti', 'hancur', 'milenial', 'dibentuk', 'zaman', 'bertahan', 'mode', 'gen', 'bilang', 'burnout', 'tanggung']
....
....
['belanja', 'impor', 'bijak', 'atur', 'uang', 'keuangan', 'prioritaskan', 'kebutuhan', 'utang', 'dolar', 'siapin', 'strategi', 'biar', 'berat', 'investasi', 'aset', 'stabil']

Stemming merupakan tahapan akhir dalam proses *preprocessing*, yaitu menghapus akhiran, sisipan, dan awalan dari kata-kata berimbuhan (turunan) dengan menggunakan pustaka (*library*) Sastrawi. *Library* Sastrawi adalah *library* Python yang berfungsi membantu mereduksi kata-kata berimbuhan bahasa Indonesia menjadi bentuk kata dasarnya. Tabel 9 menunjukkan hasil penerapan tahap *Stemming* pada *preprocessing* terhadap data yang diperoleh.

Tabel 9. Stemming

full_text
['gen', 'unik', 'sadar', 'healing', 'batas', 'sehat', 'mental', 'malas', 'henti', 'hancur', 'milenial', 'bentuk', 'zaman', 'tahan', 'mode', 'gen', 'bilang', 'burnout', 'tanggung']
....
....
['belanja', 'impor', 'bijak', 'atur', 'uang', 'uang', 'prioritas', 'butuh', 'utang', 'dolar', 'siapin', 'strategi', 'biar', 'berat', 'investasi', 'aset', 'stabil']

Tabel 9 menunjukkan hasil dari proses *stemming*, yang mencari kata dasar pada setiap dokumen cuitan yang telah di-*crawl* sebelumnya. Kata-kata dasar ini diambil dari keseluruhan data yang terdapat pada atribut Teks.

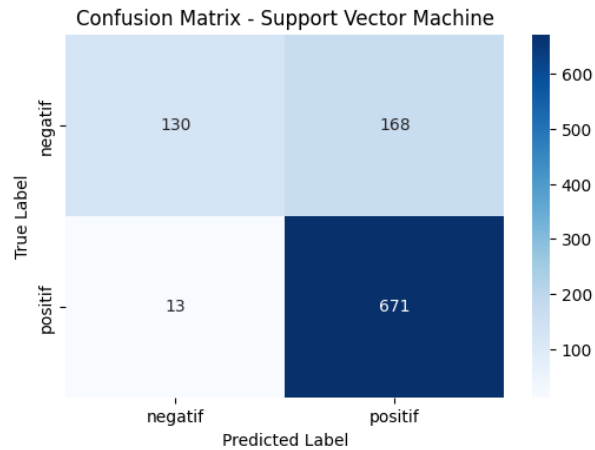
Pemodelan

Pemodelan merupakan tahapan membangun model *machine learning* untuk mengklasifikasikan sentimen teks yang berkaitan dengan isu kesehatan mental pada platform X ke dalam dua kelas, yaitu positif dan negatif. Sebelum proses pemodelan, dataset dibagi menjadi data latih (*training set*) dan data uji (*testing set*) dengan rasio 80:20. Selanjutnya, proses klasifikasi dilakukan menggunakan empat algoritma *machine learning*, yaitu Support Vector Machine (SVM), Random Forest, Naïve Bayes, dan Logistic Regression. Model yang dihasilkan kemudian dievaluasi untuk membandingkan kinerja masing-masing algoritma dalam mengklasifikasikan sentimen.

Evaluasi

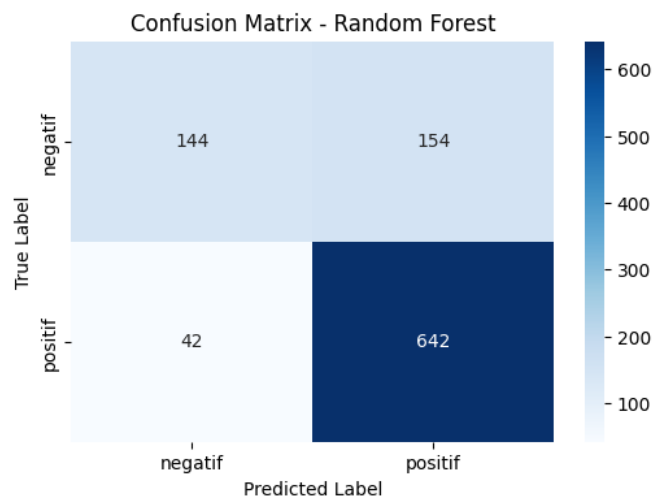
Gambar 2 menunjukkan hasil *confusion matrix* pada algoritma SVM. Berdasarkan hasil pengujian, model *Support Vector Machine* (SVM) menunjukkan performa terbaik dengan akurasi sebesar 82%, yang merupakan nilai tertinggi dibandingkan model lainnya. Model ini memperoleh nilai presisi sebesar 91% untuk kelas negatif dan 80% untuk kelas positif, yang menunjukkan kemampuan yang baik dalam menghasilkan prediksi yang tepat. Dari sisi *recall*, SVM memiliki nilai 44% untuk kelas negatif dan 98% untuk kelas positif, sehingga sangat efektif dalam mengidentifikasi tweet positif, namun masih kurang optimal dalam mendeteksi seluruh tweet negatif. Hal tersebut juga tercermin pada nilai *F1-Score*,

yaitu 59% untuk kelas negatif dan 88% untuk kelas positif. Secara keseluruhan, hasil ini menunjukkan bahwa SVM merupakan model dengan performa terbaik dalam penelitian ini, meskipun masih diperlukan peningkatan pada kemampuan mendeteksi sentimen negatif.



Gambar 2. Confusion Matrix SVM

Gambar 2 menunjukkan hasil *confusion matrix* pada algoritma Random Forest. Berdasarkan hasil pengujian, model *Random Forest* memperoleh akurasi sebesar 80%. Model ini menunjukkan performa yang sangat baik dalam mengidentifikasi tweet positif dengan nilai presisi 94%, *recall* 99%, dan *F1-Score* 97%, yang merupakan hasil terbaik untuk kelas positif dibandingkan model lainnya. Sementara itu, untuk kelas negatif, model memperoleh nilai presisi 81%, *recall* 30%, dan *F1-Score* 44%, yang menunjukkan bahwa meskipun prediksi negatif yang dihasilkan cukup akurat, model masih kurang optimal dalam mendeteksi seluruh tweet yang benar-benar memiliki sentimen negatif. Secara keseluruhan, *Random Forest* sangat andal dalam mengklasifikasikan sentimen positif, namun masih memerlukan peningkatan dalam mengenali sentimen negatif.



Gambar 3. Confusion Matrix Random Forest

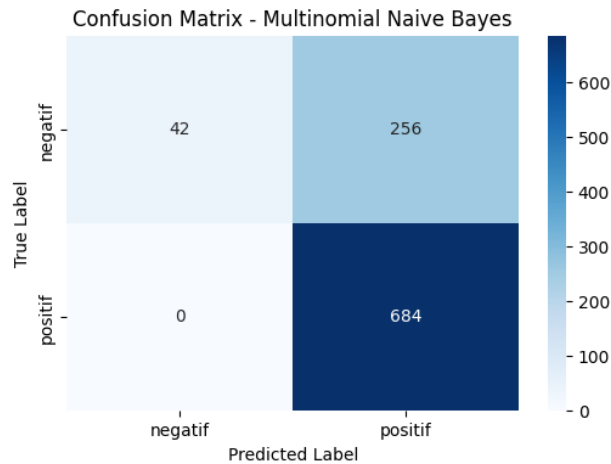
Gambar 3 menunjukkan hasil *confusion matrix* pada algoritma *Naïve Bayes*. Berdasarkan hasil pengujian, model *Multinomial Naïve Bayes* memperoleh akurasi sebesar 74%, yang merupakan nilai terendah dibandingkan model lainnya. Model ini memiliki presisi 100% untuk kelas negatif dan 73% untuk kelas positif, menunjukkan bahwa setiap

Yuwan Jumaryadi: *Penulis Korespondensi



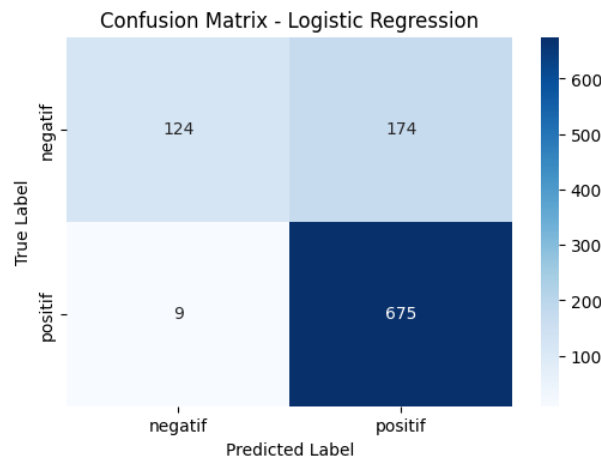
Copyright © 2026, Yuwan Jumaryadi.

prediksi negatif yang dihasilkan selalu benar. Namun, nilai *recall* untuk kelas negatif hanya 14%, sedangkan *recall* kelas positif mencapai 100%, yang mengindikasikan bahwa model sangat baik dalam mengenali tweet positif tetapi kurang mampu mendeteksi sebagian besar tweet negatif. Hal tersebut juga tercermin pada nilai *F1-Score*, yaitu 25% untuk kelas negatif dan 84% untuk kelas positif. Secara keseluruhan, model *Multinomial Naïve Bayes* cenderung mendominasi prediksi ke kelas positif sehingga masih kurang optimal dalam mengidentifikasi sentimen negatif.



Gambar 4. Confusion Matrix Naïve Bayes

Gambar 4 menunjukkan hasil *confusion matrix* pada algoritma Logistic Regression. Berdasarkan hasil pengujian, model *Logistic Regression* memperoleh akurasi sebesar 81%, yang menunjukkan bahwa sebagian besar prediksi yang dihasilkan telah sesuai dengan label sebenarnya. Model ini memiliki nilai presisi sebesar 93% untuk kelas negatif dan 80% untuk kelas positif, sehingga mampu menghasilkan prediksi yang cukup akurat. Dari sisi *recall*, model memperoleh nilai 42% untuk kelas negatif dan 99% untuk kelas positif, yang menunjukkan bahwa model sangat baik dalam mengidentifikasi tweet positif, namun masih kurang optimal dalam mendeteksi seluruh tweet negatif. Hal ini juga tercermin pada nilai *F1-Score*, yaitu 58% untuk kelas negatif dan 88% untuk kelas positif. Secara keseluruhan, *Logistic Regression* memiliki performa yang baik, terutama dalam mengklasifikasikan sentimen positif, tetapi masih memerlukan peningkatan dalam mengenali sentimen negatif.



Gambar 5. Confusion Matrix Logistic Regression

4. KESIMPULAN

Berdasarkan hasil penelitian yang telah dilakukan, dapat disimpulkan bahwa pendekatan berbasis leksikon terbukti efektif dalam melakukan pelabelan sentimen secara otomatis pada data tekstual dari platform X. Hasil analisis menunjukkan bahwa opini masyarakat terhadap isu kesehatan mental didominasi oleh sentimen positif, yang mengindikasikan tingginya tingkat kesadaran dan dukungan publik terhadap isu tersebut. Berdasarkan evaluasi menggunakan metrik accuracy, precision, recall, dan F1-score, algoritma Support Vector Machine (SVM) memberikan performa terbaik dengan akurasi sebesar 82%, diikuti oleh Logistic Regression (81%), Random Forest (80%), dan Multinomial Naïve Bayes (74%). Selain memiliki akurasi tertinggi, SVM juga menunjukkan keseimbangan performa yang lebih baik dibandingkan model lainnya, meskipun seluruh model masih memiliki keterbatasan dalam mendeteksi sentimen negatif yang ditunjukkan oleh nilai recall kelas negatif yang relatif rendah. Dengan demikian, kombinasi pelabelan berbasis leksikon dan algoritma SVM dapat menjadi pendekatan yang andal untuk memetakan opini publik mengenai isu kesehatan mental serta memberikan informasi yang bermanfaat bagi instansi kesehatan dan para pemangku kebijakan dalam merancang strategi komunikasi dan layanan kesehatan mental yang lebih tepat sasaran.

Pada penelitian selanjutnya, pengembangan dapat dilakukan dengan memanfaatkan model bahasa berbasis transformer seperti IndoBERT atau IndoRoBERTa untuk meningkatkan kemampuan model dalam memahami konteks bahasa Indonesia yang lebih kompleks. Selain itu, analisis sentimen dapat diperluas menjadi Aspect-Based Sentiment Analysis (ABSA) untuk mengidentifikasi sentimen terhadap aspek-aspek tertentu dari isu kesehatan mental, serta dikombinasikan dengan emotion classification untuk mengenali emosi spesifik, seperti bahagia, sedih, marah, takut, atau cemas, yang terkandung dalam setiap unggahan.

5. REFERENCES

- [1] R. R. Septa and B. Kusumasari, "Opini Publik Terkait Tren Isu Kesehatan: Analisis Konten pada Twitter dan Portal Berita di Yogyakarta," *PESIRAH J. Adm. Publik*, vol. 2, no. 2, pp. 34–43, 2023, doi: 10.47753/pjap.v2i2.34.
- [2] M. D. Al Fahreza, A. Luthfiarta, M. Rafid, and M. Indrawan, "Analisis Sentimen: Pengaruh Jam Kerja Terhadap Kesehatan Mental Generasi Z," *J. Appl. Comput. Sci. Technol.*, vol. 5, no. 1, pp. 16–25, 2024.
- [3] Y. Jumaryadi, R. Fajriah, U. Salamah, B. Priambodo, and A. Lystha, "Machine Learning Approaches to Sentiment Analysis of Mental Health Discussions on Platform X," *PIKSEL Penelit. Ilmu Komput. Sist. Embed. Log.*, vol. 13, no. 2, pp. 235–246, 2025, doi: 10.33558/piksel.v13i2.11350.
- [4] I. G. B. A. Budaya, L. Putu Safitri Pratiwi, and D. Panji Agustino, "Klasifikasi Sentimen untuk Analisis Kepuasan Pelayanan Puskesmas Berbasis Arsitektur LSTM," *Smart Comp Jurnalnya Orang Pint. Komput.*, vol. 12, no. 4, 2023, doi: 10.30591/smartcomp.v12i4.5361.
- [5] E. S. Aritonang and Y. Jumaryadi, "Analisis Sentimen Ulasan Pengguna QRIS pada Aplikasi GoPay : Studi Komparatif Algoritma Support Vector Machine dan Decision Tree Berbasis TF – IDF," *TIN Terap. Inform. Nusantara*, vol. 6, no. 12, pp. 2323–2330, 2026, doi: 10.47065/tin.v6i12.9582.
- [6] Ferdi and Vina Ayumi, "Analisa Sentimen Mengenai Kenaikan Harga Bbm Menggunakan Metode Naïve Bayes Dan Support Vector Machine," *JSAI (Journal Sci. Appl. Informatics)*, vol. 6, no. 1, pp. 1–10, 2023, doi: 10.36085/jsai.v6i1.4628.
- [7] N. Zelina and A. Afiyati, "Analisis Sentimen Ulasan Pengguna Aplikasi M-Banking Menggunakan Algoritma Support Vector Machine dan Decision Tree," *J. Linguist. Komputasional*, vol. 7, no. 1, pp. 31–37, 2024, doi: 10.26418/jlk.v7i1.169.

- [8] I. A. Kashim, "Public Health Surveillance for COVID-19 Using Twitter Sentiment Analysis," *Texila Int. J. Public Heal.*, vol. 12, no. 3, 2024, doi: 10.21522/TIJPH.2013.12.03.Art034.
- [9] A. I. Pandesenda, R. R. Yana, E. A. Sukma, A. Yahya, P. Widharto, and A. N. Hidayanto, "Sentiment Analysis of Service Quality of Online Healthcare Platform Using Fast Large-Margin," in *Proceedings - 2nd International Conference on Informatics, Multimedia, Cyber, and Information System, ICIMCIS 2020*, Jakarta: IEEE, 2020, pp. 121–125. doi: 10.1109/ICIMCIS51567.2020.9354295.
- [10] O. Oyeboade *et al.*, "Health, psychosocial, and social issues emanating from the COVID-19 pandemic based on social media comments: Text mining and thematic analysis approach," *JMIR Med. Informatics*, vol. 9, no. 4, pp. 1–26, 2021, doi: 10.2196/22734.
- [11] A. Ruangkanjanases and T. Hariguna, "Exploring the synergy of guided numeric and text analysis in e-commerce: a comprehensive investigation into univariate and multivariate distributions," *PeerJ Comput. Sci.*, vol. 10, pp. 1–23, 2024, doi: 10.7717/peerj-cs.2288.
- [12] Y. Jumaryadi, R. Meiyanti, R. Fajriah, A. N. Mahsyar, and P. S. Anggraeni, "Implementasi Algoritma Random Forest untuk Analisis Sentimen Ulasan Pengguna Aplikasi Merdeka Mengajar," *Bull. Comput. Sci. Res.*, vol. 5, no. 4, pp. 813–820, 2025, doi: 10.47065/bulletincsr.v5i4.530.
- [13] B. K. Subrata, Y. Jumaryadi, F. P. Sulistyio, and S. Rahayu, "Analisis Sentimen Masyarakat terhadap Profesionalisme Generasi Z di Dunia Kerja Menggunakan Support Vector Machine (SVM)," *J. Artif. Intell. Technol. Inf.*, vol. 4, no. 2, pp. 314–329, 2026.
- [14] K. Novianto, H. Herlawati, and A. Hidayat, "Klasifikasi Sentimen iPhone Bekas di Tokopedia menggunakan Naïve Bayes dan Support Vector Machine," *J. Ilm. FIFO*, vol. 17, no. 2, pp. 212–224, 2025, doi: 10.22441/fifo.2025.v17i2.010.
- [15] M. Müller and M. Salathé, "Addressing machine learning concept drift reveals declining vaccine sentiment during the COVID-19 pandemic," *arXiv Prepr.*, pp. 1–12, 2020, [Online]. Available: <http://arxiv.org/abs/2012.02197>
- [16] O. A. Malik, N. Ismail, B. R. Hussein, and U. Yahya, "Automated real-time identification of medicinal plants species in natural environment using deep learning models—a case study from Borneo Region," *Plants*, vol. 11, no. 15, p. 1952, 2022.
- [17] D. W. Seno and A. Wibowo, "Analisis Sentimen Data Twitter Tentang Pasangan Capres-Cawapres Pemilu 2019 Dengan Metode Lexicon Based Dan Support Vector Machine," *J. Ilm. FIFO*, vol. 11, no. 2, p. 144, 2019, doi: 10.22441/fifo.2019.v11i2.004.