

ANALISIS SENTIMEN PENGGUNA APLIKASI MYMRTJ DI GOOGLE PLAY STORE MENGGUNAKAN METODE NAIVE BAYES DAN SUPPORT VECTOR MACHINE

Septi Cesarani Helmi^{1*}, Siti Hadijah Hasanah²

^{1,2}Program Studi Statistika, Fakultas Sains dan Teknologi, Universitas Terbuka, Indonesia

*Penulis korespondensi: septi.rani45@gmail.com

ABSTRAK

Mass Rapid Transit (MRT) merupakan salah satu moda transportasi umum yang banyak digunakan oleh masyarakat Jakarta. Agar dapat memaksimalkan pelayanan jasa pada MRT pemerintah membuat aplikasi dengan nama “MyMRTJ” sebagai tempat untuk mengakses dan mendapatkan informasi terkini untuk para pengguna. Di dalam aplikasi tersebut yang tersedia di *Google Play Store*, terdapat berbagai macam ulasan dari pengguna tentang pengalaman mereka dalam menggunakan MRT di Jakarta. Dalam hal ini peneliti ingin mengetahui algoritma klasifikasi mana yang terbaik dan akurat dalam menganalisis sentimen penggunaan aplikasi MyMRTJ dengan cara membandingkan metode klasifikasi *Naive Bayes* dan metode *Support Vector Machine* (SVM). Data yang ingin di analisis adalah data ulasan pengguna pada aplikasi MyMRTJ dari tahun 2023 sampai 2025 dengan melalui tahap berupa pembersihan data, *tokenisasi*, pelabelan sentimen, dan proses pelatihan serta pengujian model. Setelah melalui beberapa proses tahapan didapatkan hasil analisis berupa nilai akurasi 91.67%, nilai *precision* 94.81% dan nilai *recall* 90.12% pada metode SVM serta nilai akurasi 77.78%, nilai *precision* 94.55% dan nilai *recall* 64.20% pada metode *Naive Bayes*. Dari data yang didapat menunjukkan semua hasil pada nilai akurasi, nilai *precision*, dan nilai *recall* SVM lebih besar dari *Naive Bayes*.

Kata kunci: Analisis sentimen, Mymrtj, Naive Bayes, dan SVM.

1 PENDAHULUAN

MRT merupakan transportasi umum terbaru yang terdapat di Ibukota Jakarta sebagai salah satu bentuk perhatian pemerintah untuk mengurangi tingkat kemacetan. Banyaknya antusias dari masyarakat dengan kemunculan MRT ini membuat pemerintah lebih meningkatkan kembali dalam sistem pelayanan. Salah satu cara untuk dapat memberikan pelayanan terbaik dengan mengembangkan sebuah aplikasi MyMRTJ yang berisikan mengenai informasi penting terkait pelayanan MRT Jakarta. Dalam aplikasi tersebut disediakan pula tempat untuk para pengguna memberikan komentar atau ulasan mengenai pengalaman mereka dalam menggunakan MRT Jakarta. Tanggapan, saran, maupun kritik dari pengguna MRT harus dianalisis dan diolah dengan benar supaya dapat menghasilkan informasi yang bermanfaat untuk PT MRT Jakarta (Iryana, Indriati, dan Adikara, 2021).

Meningkatnya jumlah pengguna MRT Jakarta dari tahun ke tahun membuat ulasan yang terdapat di aplikasi menjadi semakin banyak. Dalam aplikasi MyMRTJ terdapat berbagai macam ulasan yang diberikan oleh para pengguna agar sistem pelayanan dapat diubah menjadi lebih baik lagi. Ulasan tersebut dapat dikelompokkan berdasarkan sentimen positif maupun sentimen negatif yang nantinya akan diproses menjadi sebuah data untuk melihat apakah pengguna sudah merasakan manfaat yang terdapat di aplikasi tersebut atau belum. Jika proses pengelompokan dilakukan secara manual tentu akan memerlukan waktu dan tenaga dalam

jumlah besar. Dalam hal ini, untuk mengidentifikasi dan mengklasifikasikan opini pengguna secara otomatis ke dalam kategori tertentu diperlukan suatu proses analisis sentimen (Putri, Indriati, dan Wihandika, 2020).

Pada penelitian sebelumnya yang dilakukan oleh Putri, Indriati, dan Wihandika (2020) mengenai analisis sentimen terhadap ulasan pengguna MRT Jakarta dengan metode *Neighbor-Weighted K-Nearest Neighbor* (NWKNN) dengan seleksi fitur *Information Gain* terdapat kelemahan yaitu berupa kinerja menjadi lambat ketika jumlah data meningkat karena harus menghitung jarak ke setiap data pelatihan.

Selain itu beberapa penelitian lain juga telah membandingkan performa algoritma klasifikasi pada analisis sentimen, khususnya *Naive Bayes* dan *Support Vector Machine* pada berbagai data ulasan di media sosial yang menunjukkan bahwa kedua metode tersebut memiliki performa yang kompetitif untuk klasifikasi teks (Kurniawan dkk (2023); Iwandini dkk (2023); Azzahara dkk (2023)).

Dalam kesempatan kali ini, penulis ingin mengetahui cara yang paling efisien untuk mendapatkan hasil analisis sentimen dalam mengklasifikasikan ulasan pengguna aplikasi MyMRTJ di Google Play Store dengan dua metode. Di antaranya *Naive Bayes Classification* (NBC) dan *Support Vector Machine* (SVM). Menurut Siregar, Ladayya, Albaqi, dan Wardana (2023) Metode *Naive Bayes* merupakan teknik klasifikasi berbasis probabilitas sederhana yang mengasumsikan bahwa setiap fitur bersifat independen satu sama lain sementara itu, *Support Vector Machine* (SVM) merupakan algoritma *supervised learning* yang bekerja dengan mencari fungsi *hyperplane* atau batas pemisah terbaik untuk memisahkan data ke dalam kelas tertentu. Keduanya dipilih karena memiliki tingkat keakuratan yang lebih tinggi dalam mengklasifikasi teks serta telah terbukti lebih akurat dalam menganalisis data ulasan pengguna secara sistematis. Tujuan dalam penelitian ini untuk mengetahui metode yang paling cocok dalam menghasilkan kinerja yang lebih optimal pada analisis sentimen terhadap ulasan pengguna aplikasi tersebut dan untuk mengetahui apakah pengguna sudah merasakan keunggulan yang ada di dalam aplikasi MyMRTJ.

2 METODE

2.1 Sumber Data

Dalam penelitian ini penulis menggunakan data ulasan pada periode Januari 2023 hingga Oktober 2025 dengan jumlah ulasan sebanyak 478 di aplikasi MyMRTJ. Data tersebut dikumpulkan dengan menggunakan teknik *web scraping*, yang menghasilkan informasi berupa nama pengguna, *review*, waktu ulasan (tanggal), serta *rating* atau bintang yang diberikan pada aplikasi tersebut. Berikut untuk hasil *scrapping* datanya.

No.	Nama	Tanggal	Review	Bintang	
0	1	Andriyan	12/10/2025	aplikasi payah parah baru download uda daftar ...	1
1	2	Arief Rahman Hanafi	14/09/2025	opsi pembayaran untuk bank tidak lengkap... ma...	1
2	3	berryosis	09/09/2025	kalo bikin aplikasi yg niat kalo hujan knp ha...	1
3	4	Ixsan Maulana	26/08/2025	ke balik milih setasiun buat masuk, udah nyamp...	1
4	5	Teuku Raziq Fadjri Fahari	11/08/2025	Gak bisa bayar pakai Qris.	1

Gambar 1. Hasil Scrapping Data

2.2 Metode Analisis Sentimen

Analisis sentimen merupakan studi yang memiliki fungsi menganalisis pendapat, sentimen, penilaian, perilaku dan emosi dalam masyarakat melalui suatu objek seperti, pelayanan publik (Siregar, Ladayya, Albaqi, dan Wardana, 2023). Dalam analisis sentimen, suatu opini atau pendapat yang terdapat dalam teks akan diekstraksi dan diubah menjadi matriks yang berisi bobot kata. Jumlah ekstraksi sentimen positif dan negatif tersebut kemudian dihitung untuk mengetahui kecenderungan opini masyarakat secara keseluruhan terhadap objek yang diteliti. Pada tabel 1 telah diketahui terdapat sebanyak 334 data tentang ulasan positif (sentimen positif) dan 144 data tentang ulasan negatif (sentimen negatif) di ulasan aplikasi MyMRTJ.

Tabel 1. Jumlah Ekstraksi Sentimen Positif dan Negatif

Label	Jumlah Data
Positif	334
Negatif	144

2.3 Naive Bayes

Naive Bayes merupakan sebuah metode dimana dapat memprediksi barisan kemungkinan secara sederhana dengan didasarkan diterapkannya aturan bayes menggunakan anggapan bebas serta kuat (Permana, Widiastuti, dan Saepudin, 2023). Tingkat keakuratan *Naive Bayes* merupakan penggolongan dengan tingkat keakuratan paling bagus dibanding bentuk klasifikasi yang lain. Dalam proses klasifikasi menggunakan algoritma *Naive Bayes* menggunakan persamaan berikut:

$$P(x|y) = \frac{P(y|x)P(x)}{P(y)}$$

Dimana :

y : Data (label)

x : Hipotesa y adalah label

$P(x|y)$: Peluang x pada y

$P(y|x)$: Peluang y pada x

$P(x)$: Peluang x

2.4 SVM

SVM merupakan algoritma populer untuk klasifikasi, dengan kemampuan utama mengidentifikasi *hyperplane* yang memisahkan kelas berbeda secara optimal dengan tingkat akurasi dan kualitas yang tinggi (Kurniawan, Ferdian, dan Hidayati, 2025). Adapun untuk proses klasifikasi algoritma *Support Vector Machine*, menggunakan persamaan berikut:

$$F(xd) = \sum_{i=1}^{ns} a_i y_i \rightarrow_{xi} \rightarrow_{xd} + b$$

Dimana :

ns : Jumlah support vector

ai = Nilai bobot setiap titik data

yi = Kelas data

$\rightarrow_{xi}$ = Variabel support vector

$\rightarrow_{xd}$ = Data yang akan diklasifikasikan

b = Nilai error atau bias

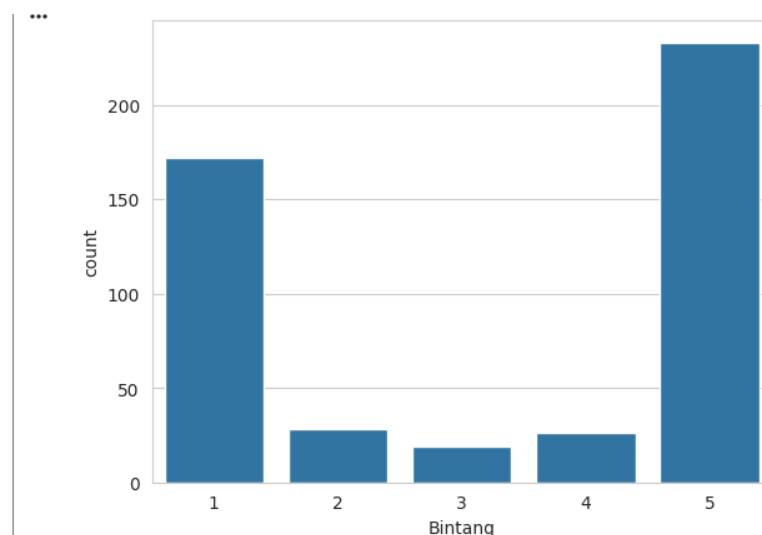
2.5 Rangkaian Pengolahan Data Menggunakan SVM dan Naive Bayes

Dalam penelitian ini, dilakukan proses *web scraping* pada aplikasi MyMRTJ di platform *Google Play Store*. Proses pengolahan data dilakukan dengan membandingkan kedua metode dari kinerja algoritma *Naive Bayes* dan SVM. Kedua metode ini dipilih karena memiliki keunggulan yang sama dalam mengolah data yaitu terbukti secara efektif dalam menganalisis teks khususnya dalam mengolah data ulasan para pengguna secara sistematis. Setelah data ulasan terkumpul kemudian dilakukan beberapa tahapan pengolahan berupa proses pengumpulan data ulasan, *tokenisasi*, pelabelan data berdasarkan kategori sentimen serta pemerataan dan analisis data dalam setiap kelas sentimen. Data yang sudah diolah tersebut akan digunakan untuk menghasilkan klasifikasi ke dalam metode algoritma dan SVM. Hasil yang didapat dari kedua klasifikasi akan dilihat lebih lanjut melalui analisis performa dan penyusunan *confusion matrix* dalam mengidentifikasi sentimen pengguna aplikasi MyMRTJ.

3. HASIL DAN PEMBAHASAN

3.1 Pengumpulan Data

Untuk proses pengumpulan data ulasan MyMRTJ dari aplikasi *Google Play Store* dilakukan cara dengan menggunakan program *Instant Data Scraper*. Program tersebut berfungsi sebagai alat untuk melihat seluruh komentar pengguna. Setelah data sudah tersedia dapat dilanjutkan dengan mengolah dan memproses lebih lanjut menggunakan *Google Colab* sehingga setiap tahapan pengolahan dapat dilakukan secara sistematis dan terstruktur.



Gambar 2. Jumlah data dalam bentuk grafik

Grafik yang ada pada gambar 2 menunjukkan jumlah data ulasan pengguna aplikasi MyMRTJ periode Januari 2023 hingga Oktober 2025 berdasarkan bintang. Pada bintang “1” jumlah ulasan mencapai kisaran sekitar 150-200 ulasan. Pada bintang “2” dan “3” serta “4” memiliki jumlah ulasan yang hampir sama sekitar kurang dari 50 ulasan sementara untuk

bintang “5” jumlah ulasan mencapai angka 200. Hal ini dapat disimpulkan bahwa pengguna sangat menyukai transportasi umum MRT.

Tabel 2. Jumlah ulasan pengguna MRT berdasarkan bintang

Bintang	Jumlah Ulasan
1	172
2	28
3	19
4	26
5	233

Dari Tabel 2. Dapat dilihat bahwa jumlah ulasan tertinggi terdapat pada bintang “5” dengan jumlah sebanyak 233 ulasan, kemudian ulasan tertinggi kedua adalah bintang “1” dengan jumlah mencapai 172 ulasan. Sementara pada bintang “2”, “3” dan “4” jumlah ulasan hanya selisih sedikit yaitu 28, 19 dan 26 ulasan. Dari tabel tersebut dapat dilihat bahwa pengguna lebih spesifik dalam memberikan bintang tertinggi maupun bintang terendah daripada memilih untuk bintang netral.

3.2 Tokenisasi

Merupakan sebuah proses untuk memecah sebuah kalimat yang terdapat dalam ulasan menjadi kata-kata individu. Proses ini bertujuan untuk menghilangkan semua karakter selain huruf (Putri, Indriati, dan Wihandika, 2020).

***	Bintang	Review	cleaned_text	label	text_len	punct	tokens
0	1	aplikasi payah parah baru download uda daftar ...	aplikasi payah parah baru download uda daftar ...	0	63	0.0	[aplikasi, payah, parah, baru, download, uda, ...]
1	1	opsi pembayaran untuk bank tidak lengkap... ma...	opsi pembayaran untuk bank tidak lengkap ma...	0	56	0.0	[opsi, pembayaran, untuk, bank, tidak, lengkap...]
2	1	kalo bikin aplikasi yg niat kalo hujan knp ha...	kalo bikin aplikasi yg niat kalo hujan knp ha...	0	152	0.0	[kalo, bikin, aplikasi, yg, niat, kalo, hujan...]
3	1	ke balik milih setasiun buat masuk, udah nyamp...	ke balik milih setasiun buat masuk udah nyamp...	0	185	0.0	[ke, balik, milih, setasiun, buat, masuk, udah...]
4	1	Gak bisa bayar pakai Qris.	gak bisa bayar pakai qris	0	21	0.0	[gak, bisa, bayar, pakai, qris]

Gambar 3. Hasil *Tokenisasi*

Hasil *tokenisasi* pada Gambar 3 menunjukkan hasil berupa setiap kalimat dalam *review* dipecah menjadi unit kata (token). Salah satu contoh kalimat pada nomor 4 “Gak bisa bayar pakai Qris” menjadi *gak, bisa, bayar, pakai, qris*. Proses ini dilakukan untuk mempermudah dalam tahapan analisis sentimen dan supaya dapat memungkinkan algoritma pemodelan untuk mengolah fitur bahasa menjadi lebih terstruktur.

3.3 Pelabelan Data

Pelabelan data dilakukan dengan cara mengumpulkan data yang telah tersedia yang selanjutnya akan dilakukan proses analisis sentimen.

	Review	Bintang	label
0	aplikasi payah parah baru download uda daftar ...	1	Negatif
1	opsi pembayaran untuk bank tidak lengkap... ma...	1	Negatif
2	kalau bikin aplikasi yg niat kalo hujan knp ha...	1	Negatif
3	ke balik milih setasiun buat masuk, udah nyamp...	1	Negatif
4	Gak bisa bayar pakai Qris.	1	Negatif

Gambar 4. Pelabelan data sentimen

Pada Gambar 4 ditampilkan hasil pengelompokan jenis sentimen. *Review* pada baris 0 hingga 4 seluruhnya memiliki *rating* bintang 1 yang menunjukkan bahwa ulasan tersebut tergolong ke dalam sentimen negatif.

3.4 Pemerataan dan Analisis Kata dalam Kelas Sentimen

Kelas sentimen dibagi menjadi dua kategori, yaitu kelas positif dan kelas negatif pada ulasan pengguna aplikasi MyMRTJ. Dengan masing-masing total ulasan positif dan negatif adalah 334 dan 144. Pola kemunculan kata pada kelas sentimen positif dan negatif dalam penelitian ini sejalan dengan studi sebelumnya pada data kepuasan pengguna MRT yang dilakukan menggunakan pendekatan NLP dan klasifikasi sentimen (Pribadi dan Ernawati 2024).

Jumlah data sentimen positif dari ulasan pengguna aplikasi MyMRTJ adalah sebanyak 334. Selanjutnya visualisasi distribusi kata dilakukan menggunakan *Word Cloud*.



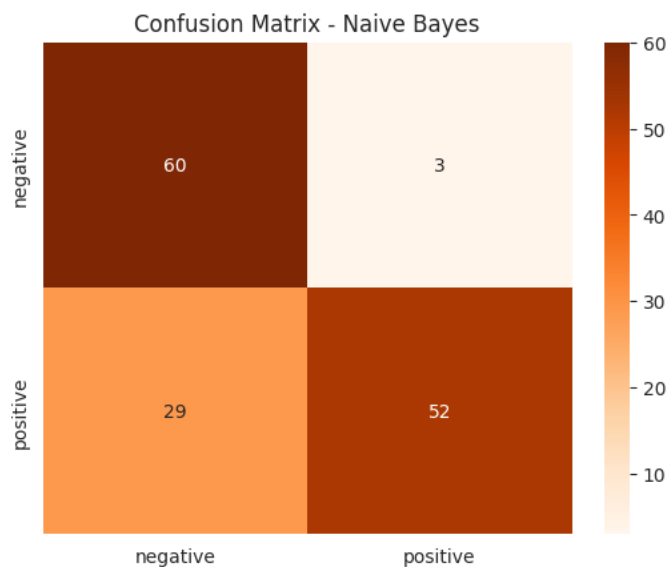
Gambar 5. Positive Reviews Word Cloud

Gambar 5 menunjukkan informasi sentimen positif berupa kata-kata pengguna mengenai ulasan aplikasi MyMRTJ. Beberapa kata positif yang muncul dengan frekuensi tinggi di antaranya "OK", "bagus", "mantap", dan "keren". Berikut adalah contoh kata-kata yang mengandung sentimen positif :

1. "ok, tidak perlu antri di stasiun, bisa prepare selama perjalanan, dan pembayaran banyak cara, qr, gopay, dana. Memudahkan"

Accuracy			0,78	144
macro avg	0,81	0,80	0,78	144
weighted avg	0,83	0,78	0,78	144

Berdasarkan evaluasi model *Naive Bayes* pada data ulasan aplikasi MyMRTJ diperoleh hasil yang positif. Pada tabel 3 diperlihatkan angka “0” yang menunjukkan nilai negatif sementara angka “1” menunjukkan nilai positif. Pada nilai *precision*, nilai positif menunjukkan angka yang lebih besar daripada dengan nilai negatif (95% > 67%). Sementara pada nilai *recall* dan *F1-score* nilai negatif menunjukkan angka lebih besar dibandingkan dengan nilai positif. Pada nilai negatif diperoleh angka 95% untuk *recall* dan 79% untuk *F1-score* sedangkan nilai positif diperoleh angka 64% untuk *recall* dan 76% untuk *F1-score*. Untuk nilai *macro average* dan *weighted average* masing-masing mencapai angka sebesar 78% yang menunjukkan bahwa performa model stabil pada kedua kelas.



Gambar 7. *Confusion Matriks Naive Bayes.*

Pada gambar 7 klasifikasi menunjukkan kinerja dengan hasil yang sangat baik. Pada kelas negatif model dapat mengidentifikasi 60 sampel dan hanya melakukan 3 kesalahan sebagai positif. Sementara pada kelas positif klasifikasi *Naive Bayes* hanya mampu mengidentifikasi 52 sampel dengan jumlah sampel yang keliru sebagai negatif sebanyak 29. Dalam hasil *confusion matriks* ini menunjukkan bahwa *Naive Bayes* memiliki kemampuan yang kuat dalam membedakan ulasan positif dan negatif. Hasil pada tabel dan *confusion matriks Naive Bayes* menunjukkan bahwa algoritma *Naive Bayes* mampu memberikan hasil yang konsisten serta dapat digunakan untuk menganalisis sentimen secara umum dengan nilai akurasi mencapai 78%. Dalam hal ini menunjukkan bahwa model *Naive Bayes* terbukti efektif dalam mengklasifikasi ulasan aplikasi MyMRTJ.

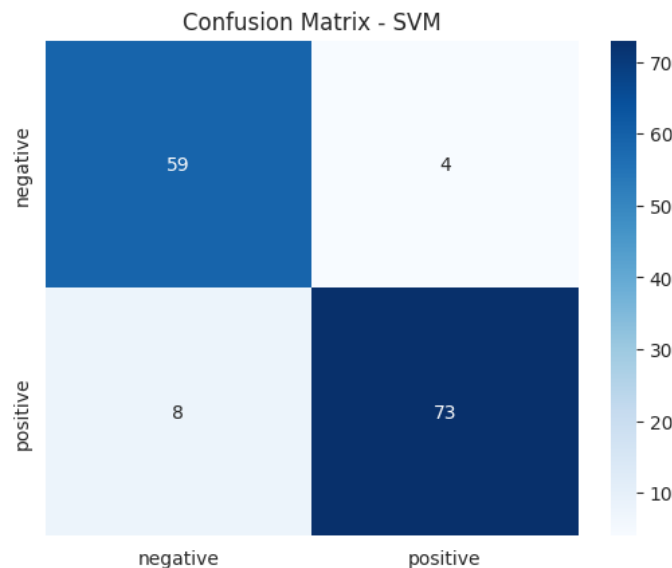
3.6 Hasil Klasifikasi SVM

Tabel 4. Hasil Klasifikasi SVM

	precision	Recall	f1-score	support
0	0,88	0,95	0,91	63
1	0,95	0,90	0,92	81

Accuracy			0,92	144
macro avg	0,91	0,92	0,92	144
weighted avg	0,92	0,92	0,92	144

Berdasarkan evaluasi model SVM pada data ulasan aplikasi MyMRTJ diperoleh hasil yang positif. Pada tabel 4 diperlihatkan angka “0” yang menunjukkan nilai negatif, sementara angka “1” menunjukkan nilai positif. Untuk nilai *precision* dan *F1-score* menghasilkan nilai positif yang lebih besar daripada dengan nilai negatif. Pada nilai positif diperoleh angka 95% untuk *precision* dan 92% untuk *F1-score* sedangkan untuk nilai negatif diperoleh angka 88% untuk *precision* dan 91% untuk *F1-score*. Sementara pada *recall* nilai negatif menunjukkan angka lebih besar dibandingkan dengan nilai positif (95% > 90%). Untuk nilai *macro average* dan *weighted average* masing-masing mencapai angka sebesar 92% yang menunjukkan bahwa performa model stabil pada kedua kelas.



Gambar 8. *Confusion Matriks SVM*

Pada Gambar 8 klasifikasi menunjukkan kinerja dengan hasil yang sangat baik. Untuk kelas negatif model dapat mengidentifikasi 59 sampel dan hanya melakukan 4 sampel kesalahan sebagai positif. Sementara pada kelas positif klasifikasi SVM berhasil mengidentifikasi 73 sampel dengan jumlah sampel yang keliru sebagai negatif sebanyak 8 sampel. Dalam hasil *confusion matriks* ini menunjukkan bahwa SVM memiliki kemampuan yang kuat dalam membedakan ulasan positif dan negatif. Hasil pada tabel dan *confusion matriks* SVM menunjukkan bahwa algoritma SVM mampu memberikan hasil yang konsisten serta dapat digunakan untuk menganalisis sentimen secara umum dengan nilai akurasi mencapai 92%. Dalam hal ini menunjukkan bahwa model SVM terbukti efektif dalam mengklasifikasi ulasan aplikasi MyMRTJ.

3.7 Analisis Performa Algoritma SVM dan Naive Bayes

Penggunaan variasi rasio pembagian data *training-testing* dalam uji model klasifikasi sentimen juga telah dibahas pada penelitian sebelumnya yang hasilnya menunjukkan bahwa komposisi data latih yang lebih besar cenderung menghasilkan performa model yang lebih stabil dan akurat (Prasetyo, Utami dan Yaqin 2024).

Data set dalam penelitian ini merupakan data ulasan pada aplikasi MyMRTJ yang diambil dari *Google play store* periode Januari 2023 hingga Oktober 2025. Data ulasan yang diperoleh sebanyak 478 ulasan. Untuk mengukur kinerja model secara objektif, data tersebut dibagi ke dalam beberapa skema *training-testing*, dengan variasi rasio 90:10 dan 80:20. Berikut untuk hasil nilai akurasi pada model algoritma Naive Bayes dan SVM.

Tabel 5. Nilai Akurasi Data

Class	Nilai akurasi berbagai rasio data	
	90:10	80:20
Naive Bayes	94,00%	91,00%
SVM	92,00%	89,00%

Berdasarkan hasil pengujian yang disajikan pada Tabel 5, terlihat bahwa variasi komposisi data *training-testing* memberikan pengaruh terhadap performa kedua algoritma. Pada algoritma *Naive Bayes*, akurasi tertinggi diperoleh pada komposisi 90:10 dengan nilai 94,00%, sedangkan akurasi terendah muncul pada komposisi 80:20 dengan nilai 91,00%. Sementara itu pada algoritma SVM, akurasi tertinggi sebesar 92,00% pada komposisi 90:10, serta akurasi terendah sebesar 89,00% pada komposisi 80:20. Pola ini mengindikasikan bahwa kedua algoritma cenderung mencapai performa optimal ketika proporsi data latih lebih besar dibandingkan data uji. Dengan demikian, penggunaan rasio 90:10 dapat dianggap sebagai komposisi data yang paling efektif dalam menghasilkan model dengan akurasi terbaik pada penelitian ini.

3.8 Perbandingan Hasil Model Naive Bayes dan SVM

Berikut menampilkan perbandingan kinerja dari kedua model, yaitu Naive Bayes, dan SVM.



Gambar 7. Perbandingan Model Naive Bayes dan SVM

Pada Gambar 7 menunjukkan perbandingan akurasi antara model *Naive Bayes* dan SVM pada proses klasifikasi sentimen. Berdasarkan visualisasi, model SVM memperoleh akurasi yang lebih tinggi, yaitu mendekati 0,90, sedangkan model *Naive Bayes* menunjukkan akurasi sekitar 0,80. Perbedaan ini mengindikasikan bahwa SVM memiliki kemampuan generalisasi yang lebih baik dibandingkan *Naive Bayes* dalam memproses data ulasan aplikasi

MyMRTJ. Oleh karena itu, SVM dapat dianggap sebagai model yang lebih optimal untuk digunakan pada penelitian ini. Selain itu, pada tabel 6 dan 7 akan disajikan pula hasil perbandingan *training testing* dengan kedua model algoritma.

Tabel 6. Hasil Perbandingan *Training-Testing* Metode SVM

Tabel Training Testing SVM		
Aspek Evaluasi	Training-Testing 90:10	Training-Testing 80:20
Akurasi	92%	89%
Kelas Negatif (0)	Precision 0.90 ; Recall 0.90 ; F1 0.90	Precision 0.81 ; Recall 0.95 ; F1 0.87
Kelas Positif (1)	Precision 0.93 ; Recall 0.93 ; F1 0.93	Precision 0.96 ; Recall 0.84 ; F1 0.90
Macro avg	0.91	0.88
Weighted avg	0.92	0.89
Jumlah data uji	48 Sampel	96 Sampel

Pada Tabel 6 skenario *training-testing* 90:10, mencapai akurasi 92% dengan performa yang seimbang pada kedua kelas. Hal ini ditunjukkan oleh nilai *F1-score* 0,90 pada kelas negatif dan 0,93 pada kelas positif. Pada skenario 80:20, akurasi model sedikit menurun menjadi 89%, namun metrik evaluasi tetap stabil dengan *F1-score* 0,87 pada kelas negatif dan 0,90 pada kelas positif. Perbedaan jumlah data uji dari 48 sampel menjadi 96 sampel juga menunjukkan bahwa model tetap menghasilkan performa yang konsisten meskipun diuji pada data yang lebih besar.

Tabel 7. Hasil Perbandingan *Training-Testing* Metode Naive Bayes

Tabel Training Testing Naive Bayes		
Aspek Evaluasi	Training-Testing 90:10	Training-Testing 80:20
Akurasi	94%	91%
Kelas Negatif (0)	Precision 0.95 ; Recall 0.90 ; F1 0.92	Precision 0.86 ; Recall 0.93 ; F1 0.89
Kelas Positif (1)	Precision 0.93 ; Recall 0.96 ; F1 0.95	Precision 0.94 ; Recall 0.89 ; F1 0.92
Macro avg	0.94	0.90
Weighted avg	0.94	0.91
Jumlah data uji	48 Sampel	96 Sampel

Pada Tabel 7 skenario *training-testing* 90:10, mencapai akurasi 94% dengan performa yang seimbang pada kedua kelas. Hal ini ditunjukkan oleh nilai *F1-score* 0,92 pada kelas negatif dan 0,95 pada kelas positif. Pada skenario 80:20, akurasi model sedikit menurun menjadi 91%, namun metrik evaluasi tetap stabil dengan *F1-score* 0,89 pada kelas negatif dan 0,92 pada kelas positif. Perbedaan jumlah data uji dari 48 sampel menjadi 96 sampel juga menunjukkan bahwa model tetap menghasilkan performa yang konsisten meskipun diuji pada data yang lebih besar.

4 KESIMPULAN

Berdasarkan hasil penelitian, dapat disimpulkan bahwa kedua algoritma, yaitu *Naive Bayes* dan SVM, mampu mengolah data ulasan aplikasi MyMRTJ secara efektif. Namun, keduanya menunjukkan keunggulan pada aspek yang berbeda. Pada proses klasifikasi utama, SVM menghasilkan akurasi tertinggi sebesar 92%, sementara *Naive Bayes* menghasilkan akurasi sebesar 78%. Hal ini menunjukkan bahwa SVM lebih unggul dalam memprediksi sentimen pada keseluruhan dataset. Sebaliknya, pada pengujian menggunakan skenario *training-testing*, *Naive Bayes* memberikan performa lebih stabil, dengan akurasi 94% pada perbandingan 90:10 dan 91% pada perbandingan 80:20. Sementara itu, SVM memperoleh akurasi 92% pada perbandingan 90:10 dan menurun menjadi 89% pada perbandingan 80:20. Perbedaan ini menunjukkan bahwa SVM lebih kuat ketika dilatih pada keseluruhan data, sedangkan *Naive Bayes* cenderung lebih konsisten saat diuji dengan variasi proporsi data latih dan data uji. Selain itu, hasil klasifikasi menunjukkan bahwa mayoritas ulasan berada pada kelas positif. Temuan ini mengindikasikan bahwa para pengguna secara umum memberikan kesan yang baik terhadap aplikasi MyMRTJ. Dominasi sentimen positif tersebut semakin menguatkan bahwa aplikasi mampu memberikan pengalaman penggunaan yang memuaskan, baik dari segi fungsi maupun kemudahan akses. Dengan demikian, kinerja model yang efektif serta kecenderungan ulasan positif dari pengguna memberikan gambaran bahwa aplikasi MyMRTJ telah diterima dengan baik oleh masyarakat.

UCAPAN TERIMA KASIH

Ucapan terimakasih kepada pihak-pihak yang telah membantu dalam penelitian. Semoga hasil penelitian bermanfaat dan menjadi acuan penelitian selanjutnya.

DAFTAR PUSTAKA

- Kurniawan I, Hananto A. L, Hilabi S. S, Hananto A, Priyatna B, & Rahman A. Y. (2023). *Perbandingan Algoritma Naive Bayes Dan SVM Dalam Sentimen Analisis Marketplace Pada Twitter*. Jurnal Teknik Informatika dan Sistem Informasi Vol. 10, No. 1, Maret 2023, Hal. 731-740. ISSN 2407-4322 E-ISSN 2503-2933.
- Iryana T, M, Indriati, & Adikara P, P. (2021). *Analisis Sentimen Masyarakat Terhadap Mass Rapid Transit Jakarta Menggunakan Metode Naive Bayes Dengan Normalisasi Kata*. Jurnal Pengembangan Teknologi Informasi dan Ilmu Komputer Vol. 5, No. 6, Mei 2021, hlm. 2753-2760. e-ISSN: 2548-964X <http://j-ptiik.ub.ac.id>.
- Pribadi M, N, N, & Ernawati I. (2024). *Perbandingan Hasil Penerapan Algoritma Klasifikasi dan Natural Language Processing Terhadap Data Kepuasan Pengguna Layanan Transportasi Umum MRT Jakarta*. JURNAL INFORMATIK Edisi ke-20, Nomor 3, Desember 2024.
- Siregar D, Ladayya F, Albaqi N, Z, & Wardana B, M. (2023). *Penerapan Metode Support Vector Machines (SVM) dan Metode Naive Bayes Classifier (NBC) dalam Analisis Sentimen Publik terhadap Konsep Child-free di Media Sosial Twitter*. e-ISSN: 2620-8369.
- Kurniawan N, D, Ferdian P, R, & Hidayati N. (2025). *Analisis Sentimen Algoritma Naive Bayes, Support Vector Machine, dan Random Forest Pada Ulasan Aplikasi Ajaib*.
- Putri F, O, Indriati, & Wihandika R, C. (2020). *Analisis Sentimen pada Ulasan Pengguna MRT Jakarta Menggunakan Metode Neighbor-Weighted K-Nearest Neighbor dengan Seleksi Fitur Information Gain*. Jurnal Pengembangan Teknologi Informasi dan Ilmu Komputer Vol. 4, No. 7, Juli 2020, hlm. 2195-2203. e-ISSN: 2548-964X <http://j-ptiik.ub.ac.id>

- Permana M, A, Widiastuti S, & Saepudin S. (2023). *ANALISIS SENTIMEN PENGGUNAAN APLIKASI VIDEO CONFERENCE PADA ULASAN GOOGLE PLAY STORE MENGGUNAKAN METODE NBC (NAIVE BAYES CLASSIFIER)*. JURSISTEKNI (Jurnal Sistem Informasi dan Teknologi Informasi) Vol 5, No.1, Januari 2023: Hal 178- 191 ISSN. P: 2715-1875, E: 2715-1883.
- Azzahara S, P, Apriyanto Y, A, & Wijaya A. (2023). *ANALISIS SENTIMEN ULASAN APLIKASI DEEPL PADA GOOGLE PLAY DENGAN METODE SUPPORT VECTOR MACHINE (SVM)*.
- Iwandini I, Triyadi A, & Soepriyono G. (2023). *Analisa Sentimen Pengguna Transportasi Jakarta Terhadap Transjakarta Menggunakan Metode Naives Bayes dan K-Nearest Neighbor*. Journal of Information System Research (JOSH) Volume 4, No. 2, Januari 2023, pp 543–550 ISSN 2686-228X (media online) <https://ejurnal.seminar-id.com/index.php/josh/> DOI 10.47065/josh.v4i2.2937.
- Prasetyo Y, Utami E, Yaqin A. (2024). *Pengaruh Komposisi Split Data Terhadap Performa Akurasi Analisis Sentimen Algoritma Naïve Bayes dan SVM*. Journal of Electrical Engineering and Computer (JEECOM) Vol. 6, No. 2 (2024), DOI: 10.33650/jeeecom.v4i2 p-ISSN: 2715-0410; e-ISSN: 2715-6427