

ANALISIS SENTIMEN PADA PENILAIAN APLIKASI WONDR BY BNI MENGUNAKAN SUPPORT VECTOR MACHINE (SVM) DAN NAÏVE BAYES CLASSIFIER (NBC)

Nadia Salsabila Harli^{1*}, Ria Faulina¹

¹Program Studi Statistika, Fakultas Sains dan Teknologi, Universitas Terbuka

*Penulis korespondensi: nadiaharli05@gmail.com

ABSTRAK

Perkembangan teknologi digital yang pesat sangat membantu berbagai aspek kehidupan, termasuk bidang keuangan. Industri perbankan di Indonesia terus berinovasi dalam meningkatkan digitalisasi layanan, salah satunya dengan menciptakan aplikasi *mobile banking*. PT Bank Negara Indonesia (BNI) mempunyai aplikasi *mobile banking* kedua bernama *wondr by BNI*. Setelah dilakukannya pembaruan aplikasi *wondr by BNI* ini memiliki banyak kemajuan yang lebih baik ditandai dengan meningkatnya rating yang awal mula dirilisnya 3,4 menjadi 4,8 sehingga mendorong munculnya berbagai pendapat atau persepsi pengguna terhadap fitur aplikasi *wondr by BNI* setelah diperbarui. Penelitian ini bertujuan untuk membandingkan klasifikasi sentimen pengguna aplikasi *wondr by BNI* menggunakan Support Vector Machine (SVM) dan Naïve Bayes Classifier (NBC). Data yang digunakan adalah komentar pengguna di Google Play Store yang di-scrape dalam periode 22 Mei 2025 sampai 13 November 2025. Dengan teknik text mining, analisis sentimen dilakukan menggunakan pembobotan kata TF-IDF dilanjutkan metode Support Vector Machine dan Naïve Bayes Classifier. Hasil penelitian menunjukkan bahwa akurasi menggunakan NBC lebih tinggi yaitu 78% dibandingkan menggunakan SVM yaitu sebesar 73%. Evaluasi performa menunjukkan bahwa kedua model memiliki kemampuan yang baik dalam mengklasifikasikan sentimen positif dan negatif, tetapi kesulitan dalam mengenali sentimen netral ditunjukkan dengan F1-score yang rendah sebesar 0,47 pada SVM dan 0,43 pada NBC. Hasil penelitian ini dapat membantu BNI dalam meningkatkan kualitas aplikasi dan meningkatkan kepuasan nasabah.

Kata kunci: analisis sentimen, *Naïve Bayes Classifier*, *Support Vector Machine*, *wondr by BNI*

1 PENDAHULUAN

Teknologi digital saat ini sangat berkembang pesat dari masa ke masa. Perkembangan teknologi sangat membantu dan memudahkan manusia dalam melakukan semua aktivitas dan pekerjaan sehari-hari. Banyak aktivitas masyarakat sekarang bisa dilakukan secara online dan jarak jauh, apalagi di era revolusi industri 4.0 dan ditandai dengan kemajuan teknologi seperti *Artificial Intelligence (AI)*, *Machine Learning*, dan *Internet of Things (IoT)*. Perkembangan Teknologi ini sangat mengubah lanskap bisnis maupun kehidupan sehari-hari secara signifikan, termasuk dalam sektor jasa keuangan. Menanggapi hal tersebut, industri perbankan di Indonesia terus berinovasi dalam meningkatkan digitalisasi layanan, salah satunya dengan menciptakan aplikasi *mobile banking*. Untuk memudahkan nasabah dalam bertransaksi secara *online*, seperti pengecekan saldo, transfer uang, investasi maupun pembayaran tagihan dan pembelian (Zafira dkk., 2025).

Dalam beberapa tahun terakhir, berbagai aplikasi *mobile banking* menunjukkan peningkatan signifikan dalam jumlah penggunaannya. Tiga Aplikasi teratas dalam pengguna *mobile banking di google play store* adalah *Brimo* dari BRI, *livin* dari mandiri dan *wondr* dari BNI. Dari tiga aplikasi teratas tersebut aplikasi *wondr by BNI* merupakan aplikasi dengan rating

paling bagus diantaranya dua lainnya. *Wondr by BNI*, yang merupakan layanan perbankan digital pengganti *BNI mobile banking* dan resmi diperkenalkan pada Juli 2024. Setelah dilakukannya pembaruan aplikasi *wondr by bni* ini memiliki banyak kemajuan yang lebih baik ditandai dengan meningkatnya rating yang awal mula dirilisnya 3,4 menjadi 4,8 sehingga mendorong munculnya berbagai pendapat atau persepsi pengguna terhadap fitur aplikasi *wondr by bni* setelah diperbarui (Essa, 2024).

Dalam upaya memahami persepsi pengguna aplikasi, analisis sentimen adalah pendekatan yang paling relevan. Analisis sentimen merupakan proses dalam pemahaman, penarikan, dan pengolahan data dalam bentuk teks untuk mendapatkan dan mengidentifikasi pola teks yang mencerminkan sentiment atau opini tertentu (Hidayat dkk., 2021). Sentimen bisa berupa positif, negatif maupun netral. Metode ini dimanfaatkan untuk menelusuri arah atau kecenderungan emosi dalam sebuah teks. Pendekatan tersebut berperan penting dalam menilai pandangan masyarakat terhadap suatu produk, layanan, maupun isu tertentu. Selain kuantitas data yang cukup besar, kompleksitas bahasa yang digunakan merupakan tantangan dalam analisis sentimen yang dapat menyulitkan penentuan sentimen yang sebenarnya dan berisiko menimbulkan kesalahan dalam analisis (Nurwanda, Suarna, & Prihartono, 2024). Maka dari itu, tahapan *preprocessing* data menjadi hal yang sangat penting dalam memastikan keakuratan analisis. Tahapannya *preprocessing* meliputi proses *case folding*, *tokenizing*, *stopword removal*, dan *stemming*. Banyak cara klasifikasi yang digunakan dalam analisis sentimen, diantaranya *Decision Tree*, *K-Nearest Neighbor*, *Naïve bayes classifier*, *Support Vector Machine* (SVM), dan *Random forest*. Algoritma *Support vector machine* dan *Naive bayes* adalah suatu metode klasifikasi untuk mengolah data berupa teks dengan tingkat akurasi yang baik (Puji Astuti dkk., 2022).

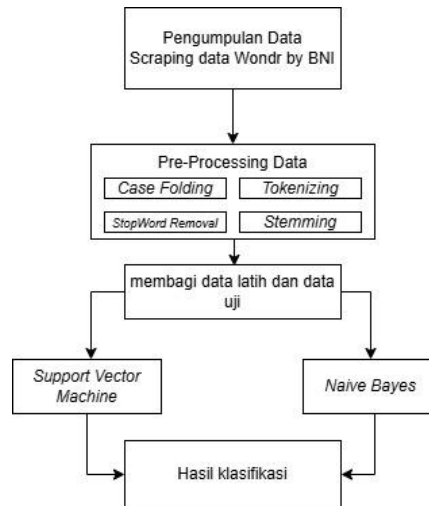
Penelitian ini akan menggunakan algoritma *Support Vector Machine* yang dikolaborasikan dengan *Naïve bayes* sebagai perbandingannya. SVM merupakan metode klasifikasi jenis terpadu (*supervised*) yang memprediksi kelas dari model atau pola berdasarkan hasil proses pelatihan. SVM merupakan algoritma yang bekerja menggunakan pemetaan nonlinear untuk mengubah data pelatihan asli ke data ruang lebih tinggi. Dalam hal ini dimensi baru akan mencari hyperplane untuk memisahkan secara linear (Handayanto & Herlawati, 2020). Algoritma *Naive Bayes* berdasarkan teorema Bayes bahwa semua kegiatan memberikan sebuah kontribusi yang sama penting atau saling bebas pada pemilihan kelas tertentu (Dedi Darwis dkk., 2020).

2 METODE

2.1 Data

Data yang digunakan dalam penelitian ini diperoleh dari aplikasi *Wondr by BNI* yang diambil melalui platform *Google Play Store*. Periode pengambilan data dilakukan pada tanggal 22 Mei 2025 sampai 13 November 2025. Pengumpulan data dilakukan dengan *Web scraping* di Python. *Web scraping* adalah sebuah teknik yang digunakan untuk mengekstraksi data dalam skala besar dari situs web, sehingga dapat mengumpulkan informasi secara sistematis dan tersimpan dalam format *file csv* yang dibutuhkan dalam analisis. Pengambilan data dilakukan menggunakan pustaka *google play scraper* di Python, dimulai dari proses inisialisasi *library*, penentuan parameter *scarping* seperti jumlah *revies*, bahasa, dan negara, hingga menyimpan hasilnya dalam bentuk data tabel.

Berikut tahapan penelitian yang ada pada Gambar 1.



Gambar 1. Tahapan penelitian

2.2 Preprocessing Data

Preprocessing data adalah langkah penting sebelum melakukan analisis sentimen. Tahap *preprocessing* ini bertujuan untuk mengonversi data mentah menjadi data dengan format yang lebih rapi, terorganisir, dan siap digunakan dalam proses klasifikasi. Adapun beberapa tahapan dalam *preprocessing* data adalah sebagai berikut:

1. *Case Folding*
Seluruh kata pada review diubah menjadi huruf kecil (*lowercase*) dengan tujuan menghindari perbedaan arti antara huruf kecil dan huruf besar
2. *Stopword Removal*
Menghapus kata-kata umum yang tidak memiliki makna penting dalam analisis sentiment, seperti kata “dan”, “yang”, “di”, sehingga meningkatkan kualitas data yang dianalisis
3. *Tokenizing*
Kalimat review akan dipecah menjadi unit-unit kata secara individual untuk memudahkan proses analisis lebih lanjut
4. *Stemming*
Kata akan dikembalikan ke bentuk dasarnya dengan menghilangkan imbuhan pada kata, dengan tujuan menyamakan kata yang memiliki makna serupa namun bentuk yang berbeda.

2.3 Pembobotan Kata

Salah satu tahap paling penting dalam analisis sentimen adalah pemberian bobot pada setiap kata yang terdapat pada teks. Pembobotan kata dilakukan dengan menggunakan *Term-Frequency-Inverse Document Frequency* yaitu teknik untuk mengubah teks menjadi vektor numerik dengan memberi bobot pada setiap kata pada teks berdasarkan frekuensi kemunculannya dalam teks dan seberapa unik kata tersebut pada kumpulan seluruh teks. (Ratnaswari dkk., 2025). Perhitungan manual TF-IDF dapat dilakukan menggunakan persamaan berikut.

$$TF(t, d) = \frac{\text{Jumlah kemunculan kata } t \text{ dalam dokumen } d}{\text{Jumlah total kata dalam dokumen } d}$$

$$IDF(t) = \log\left(\frac{N}{df(t)}\right)$$

Keterangan:

N = Jumlah Total Dokumen

$df(t)$ = Jumlah Dokumen yang mengandung kata t

$$TF - IDF(t, d) = TF(t, d) \times IDF(t)$$

2.4 Text Mining

2.4.1 Support Vector Machine

Algoritma klasifikasi *Support Vector Machine* (SVM) yang diperkenalkan oleh Vladimir Vapnik merupakan metode yang mampu mengidentifikasi kelas tertentu dengan memanfaatkan pola yang dipelajari melalui pendekatan machine learning (supervised learning). Secara umum, SVM bekerja dengan mencari hyperplane paling optimal yang berfungsi memisahkan dua kategori data dalam ruang input. Dengan menghitung nilai margin dan mencari titik maksimumnya, maka akan mendapatkan hyperplane optimal (Pratama dkk., 2023).

Untuk data yang tidak bisa dipisahkan secara linear, SVM akan berupaya memperbesar jarak atau margin antar kelompok data. Untuk tetap melakukan pemisahan kelas, metode ini menggunakan fungsi kernel. Pendekatan tersebut membuat SVM mampu bekerja dengan baik pada data yang berdimensi tinggi maupun data yang pola hubungannya tidak linier, sehingga algoritma ini banyak diterapkan dalam analisis sentiment (Fitriyah dkk., 2020).

Persamaan dalam prediksi kelas secara umum:

$$y(x) = \text{sign}(w \cdot x + b)$$

Keterangan:

w = Vektor Bobot

x = Vektor Fitur

b = Bias

2.4.2 Naïve Bayes

Naïve bayes merupakan metode klasifikasi yang tergolong sederhana namun mampu menghasilkan akurasi yang baik, terutama ketika diterapkan pada dataset berukuran besar. Secara umum, pendekatan ini beroperasi dengan asumsi bahwa setiap atribut saling bebas (independent), sehingga kemunculan suatu atribut tidak dipengaruhi oleh atribut lain, sesuai dengan prinsip dalam teorema bayes (Nadira dkk., 2023).

Berikut rumus dasar untuk menghitung probabilitas suatu kelas C berdasarkan fitur(kata):

$$P(C|w_1, w_2, \dots, w_n) \propto P(C) \cdot \prod_{i=1}^n P(w_i|C)$$

2.5 Evaluasi performa model

Untuk mengetahui seberapa baik model klasifikasi dapat memprediksi sentimen pada data uji, dilakukan evaluasi model menggunakan tabel *confusion matrix*. Tabel *confusion matrix* menunjukkan berapa jumlah prediksi benar baik pada kelas positif, negative maupun netral, serta menunjukkan berapa jumlah prediksi salahnya. Melalui confusion matrix, performa model klasifikasi dievaluasi menggunakan beberapa metrik yaitu *Accuracy*, *Precision*, *Recall*, dan *F1-score* untuk memastikan bahwa model klasifikasi dapat diandalkan dalam pengambilan keputusan. Adapun rumus perhitungan metrik:

$$\text{Accuracy} = \frac{\sum TP}{\text{Total}} = \frac{TP_{Pos} + TP_{Neg} + TP_{Net}}{\text{Total}}$$

$$Precision_{kelas} = \frac{TP_{kelas}}{TP_{kelas} + FP_{kelas}}$$

$$Recall_{kelas} = \frac{TP_{kelas}}{TP_{kelas} + FN_{kelas}}$$

$$F1_{kelas} = 2 \times \frac{Precision_{kelas} \times Recall_{kelas}}{Precision_{kelas} + Recall_{kelas}}$$

3 HASIL DAN PEMBAHASAN

Data yang digunakan dalam penelitian ini diperoleh dari 1.200 ulasan aplikasi Wondr by BNI yang diambil melalui platform Google Play Store. Untuk contoh data yang telah di scraping dapat dilihat pada Gambar 2.

reviewId	userName	userImage	content	score	thumbsUpCount	reviewCreatedVersion	at	replyContent	repliedAt	appVersion	sortOrder	appId
e96af9e9-36a4-42af-b6ea-74a5863f86d7	Andhika Widyia Kurniawan	lh.googleusercontent.com/a-/ALV-U...	Aplikasi yang sangat membantu memudahkan nasabah...	3	7	1.3.2	2025-05-22 05:29:14	Hai Kak Andhika Widyia Kurniawan, terima kasih ...	2025-05-22 05:36:27	1.3.2	most_relevant	id.bni.wondr
f5aa54d7-b082-4347-8876-f6b90abdb057	Ahmad Rico Erlangga	lh.googleusercontent.com/a-/ALV-U...	fiturnya bagus...tapi kenapa pas saya bikin te...	3	1	1.3.2	2025-05-24 03:02:49	Halo Kak Ahmad Rico Erlangga, maaf atas kendal...	2025-05-24 03:47:19	1.3.2	most_relevant	id.bni.wondr
b9c2aa30-99b7-4b3a-b45f-309d3b2f4a60	Anton Kurniawan	lh.googleusercontent.com/a-/ALV-U...	agak sedikit berbelit karena harus hafal paswo...	4	4	1.3.2	2025-05-26 08:21:29	Hai Kak Anton Kurniawan, mohon maaf atas ketid...	2025-05-26 08:39:31	1.3.2	most_relevant	id.bni.wondr
3176a862-70f0-4ef7-9134-e3ab4ef9e7d5	Dimas Fatriansyah	lh.googleusercontent.com/a-/ALV-U...	aplikasi nya bagus tapi untuk pengiriman kartu ...	3	1	1.3.2	2025-05-27 15:54:59	Hai Kak Dimas Fatriansyah, maaf atas ketidakny...	2025-05-28 01:16:43	1.3.2	most_relevant	id.bni.wondr
dcdabc30-65e8-4823-baf3-69676f34cd49	Ervita Vita	lh.googleusercontent.com/a/ACg8oc...	kenapa saldo saya ga bisa di pake buat transak...	3	1	1.3.2	2025-06-01 11:46:14	Hai Kak Ervita Vita, maaf sekali ya untuk kend...	2025-06-01 11:47:03	1.3.2	most_relevant	id.bni.wondr

Gambar 2. Hasil scrapping data

Tahapan awal dari proses analisis adalah pelabelan data secara manual. Pada bagian ini, setiap data diberi label berdasarkan makna atau isi teks yang terdapat di dalamnya. Ulasan kemudian dikelompokkan ke dalam tiga kategori sentimen, yaitu positif, negative, dan netral. Sentimen positif mencakup komentar yang mengekspresikan apresiasi, kepuasan, atau penilaian yang menguntungkan terhadap aplikasi *Wondr by BNI*. Sebaliknya, sentimen negative memuat kritik, keluhan, atau kekecewaan pengguna terkait layanan aplikasi. Adapun sentimen netral berisi pernyataan yang bersifat informatif atau deskriptif tanpa menunjukkan kecenderungan emosional positif maupun negatif secara jelas. Untuk contoh Hasil dari proses pelabelan disajikan pada Tabel 1.

Tabel 1. Contoh hasil labeling

No	Content	Sentimen
1	Saran untuk tarik tunai tanpa kartu mending kode OTP nya dikirim ke WHATSAPP aja deh biar lebih mempermudah dan simple.	Netral
2	Aplikasi nya sangat berguna sehingga ketika kartu bermasalah tidak repot2 pergi ke bank ... tinggal buka apk dan di perbaiki permasalahan nya selesai	Positif

- 3 Ini BNI makin lama makin kacau, narik uang tidak keluar, TPI saldo berkurang. tolong di perbaiki dong. belum lagi ni wonder susah amat di buka banyak errornya. Negatif

Tahap *preprocessing* adalah tahap yang penting dalam *text mining*. Data awal yang digunakan dalam text tidak selalu bersih dari noise, sehingga diperlukan penanganan yang sesuai. Oleh karena itu, tahap *text processing* perlu diselesaikan dengan baik agar analisis dapat berjalan dengan baik dan lancar (Agni dkk., 2024).

Untuk mempermudah tahap preprocessing, data mentah yang telah didapatkan akan diseleksi terlebih dahulu. Tujuannya untuk memfokuskan data hanya pada informasi yang relevan, yaitu content (isi ulasan). Selanjutnya melakukan text processing dengan melakukan *Case Folding*. Berikut contoh data setelah dilakukan *Case Folding* pada Tabel 2.

Tabel 2. Hasil *case folding*

Content	Cleaned_Text
Saran untuk tarik tunai tanpa kartu mending kode OTP nya dikirim ke WHATSAPP aja deh biar lebih mempermudah dan simple.	saran untuk tarik tunai tanpa kartu mending kode otp nya dikirim ke whatsapp aja deh biar lebih mempermudah dan simple
aplikasi nya sangat berguna sehingga ketika kartu bermasalah tidak repot2 pergi ke bank ... tinggal buka apk dan di perbaiki permasalahan nya selesai	aplikasi nya sangat berguna sehingga ketika kartu bermasalah tidak repot2 pergi ke bank tinggal buka apk dan di perbaiki permasalahan nya selesai
ini BNI makin lama makin kacau, narik uang tidak keluar, TPI saldo berkurang. tolong di perbaiki dong. belum lagi ni wonder susah amat di buka banyak errornya.	ini bni makin lama makin kacau narik uang tidak keluar tpi saldo berkurang tolong di perbaiki dong belum lagi ni wonder susah amat di buka banyak errornya

Setelah dilakukan *Case Folding*, Langkah selanjutnya adalah *Stopword Removal*, yaitu proses penghapusan kata-kata umum yang tidak memiliki makna yang signifikan dalam analisis, sehingga penghapusannya dapat meningkatkan kinerja model klasifikasi. Berikut contoh data setelah dilakukan proses *Stopword Removal* pada Tabel 3.

Tabel 3. Hasil *stopword removal*

cleaned_text	no_stopword_text
saran untuk tarik tunai tanpa kartu mending kode otp nya dikirim ke whatsapp aja deh biar lebih mempermudah dan simple	['saran', 'tarik', 'tunai', 'kartu', 'mending', 'kode', 'otp', 'nya', 'dikirim', 'whatsapp', 'aja', 'deh', 'biar', 'mempermudah', 'simple']
aplikasi nya sangat berguna sehingga ketika kartu bermasalah tidak repot2 pergi ke bank tinggal buka apk dan di perbaiki permasalahan nya selesai	['aplikasi', 'nya', 'berguna', 'kartu', 'bermasalah', 'repot2', 'pergi', 'bank', 'tinggal', 'buka', 'apk', 'perbaiki', 'permasalahan', 'nya', 'selesai']
ini bni makin lama makin kacau narik uang tidak keluar tpi saldo berkurang tolong di perbaiki dong belum lagi ni wonder susah amat di buka banyak errornya	['bni', 'kacau', 'narik', 'uang', 'tpi', 'saldo', 'berkurang', 'tolong', 'perbaiki', 'ni', 'wonder', 'susah', 'buka', 'errornya']

Setelah tahap stopword removal dilakukan, Langkah berikutnya adalah Tokenizing. Berikut contoh data setelah dilakukan proses Tokenisasi pada Tabel 4

Tabel 4. Hasil proses tokenisasi

no_stopword_text	tokenized_text
['saran', 'tarik', 'tunai', 'kartu', 'mending', 'kode', 'otp', 'nya', 'dikirim', 'whatsapp', 'aja', 'deh', 'biar', 'mempermudah', 'simple']	['saran', 'tarik', 'tunai', 'kartu', 'mending', 'kode', 'otp', 'nya', ' kirim', ' whatsapp', ' aja', ' deh', ' biar', ' mudah', ' simple']
['aplikasi', 'nya', 'berguna', 'kartu', 'bermasalah', 'repot2', 'pergi', 'bank', 'tinggal', 'buka', 'apk', 'perbaiki', 'permasalahan', 'nya', 'selesai']	['aplikasi', 'nya', 'guna', 'kartu', 'masalah', 'repot2', 'pergi', 'bank', 'tinggal', 'buka', 'apk', 'baik', 'masalah', 'nya', 'selesai']
['bni', 'kacau', 'narik', 'uang', 'tpi', 'saldo', 'berkurang', 'tolong', 'perbaiki', 'ni', 'wonder', 'susah', 'buka', 'errornya']	['bni', 'kacau', 'narik', 'uang', 'tpi', 'saldo', 'kurang', 'tolong', 'baik', 'ni', 'wonder', 'susah', 'buka', 'errornya']

Setelah tahap tokenisasi selesai, proses selanjutnya adalah *Stemming* menggunakan pustaka Sastrawi, yang secara khusus dirancang untuk Bahasa Indonesia. Proses Tahapan ini bertujuan untuk menyederhanakan variasi kata yang memiliki makna serupa, sehingga meningkatkan efisiensi dalam analisis serta akurasi dalam proses klasifikasi. Berikut contoh data setelah dilakukan proses Stemming pada Tabel 5.

Tabel 5. Hasil proses *stemming*

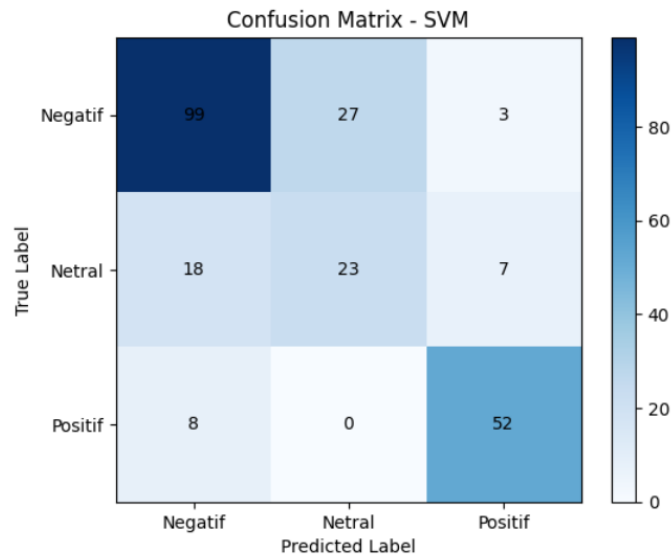
tokenized_text	stemmed_text
['saran', 'tarik', 'tunai', 'kartu', 'mending', 'kode', 'otp', 'nya', ' kirim', ' whatsapp', ' aja', ' deh', ' biar', ' mudah', ' simple']	saran tarik tunai kartu mending kode otp nya kirim whatsapp aja deh biar mudah simple
['aplikasi', 'nya', 'guna', 'kartu', 'masalah', 'repot2', 'pergi', 'bank', 'tinggal', 'buka', 'apk', 'baik', 'masalah', 'nya', 'selesai']	aplikasi nya guna kartu masalah repot2 pergi bank tinggal buka apk baik masalah nya selesai
['bni', 'kacau', 'narik', 'uang', 'tpi', 'saldo', 'kurang', 'tolong', 'baik', 'ni', 'wonder', 'susah', 'buka', 'errornya']	bni kacau narik uang tpi saldo kurang tolong baik ni wonder susah buka errornya

Setelah seluruh tahapan *preprocessing* dilakukan, maka langkah berikutnya adalah melakukan pembagian dataset menjadi dua bagian, yaitu 80% sebagai data latih dan 20% sebagai data uji. Pembagian ini dilakukan menggunakan fungsi `train_test_split` dari pustaka `scikit-learn` dengan menetapkan parameter `random_state=0` agar hasil pemisahan data tetap sama setiap kali kode dijalankan. Pada tahap ini, fitur yang digunakan berasal dari hasil *preprocessing* data (text), sedangkan variabel targetnya adalah kategori sentimen (sentiment). Pembagian ini bertujuan memastikan bahwa model dilatih pada Sebagian data dan diuji secara objektif menggunakan data yang belum pernah digunakan selama proses pelatihan.

Langkah berikutnya setelah data dibagi adalah mengubah teks ke dalam bentuk numerik dengan menerapkan dua metode representasi fitur, yaitu TF-IDF (*Term Frequency-Inverse Document Frequency*) dan *Count Vectorizer*. Kedua pendekatan ini berfungsi untuk mengonversi data teks menjadi matriks fitur sehingga dapat diproses oleh algoritma klasifikasi.

Proses transformasi tersebut menghasilkan representasi vector dari setiap dokumen dalam bentuk array numerik.

Setelah melakukan ekstraksi fitur menggunakan TF-IDF, Langkah berikutnya adalah penerapan model klasifikasi untuk mengidentifikasi sentiment atau kategori pada data teks yang telah diproses. Pada penelitian ini, model *Support Vector Machine* dan *Naïve Bayes* akan diterapkan untuk menghasilkan model yang dapat diulang. Proses pelatihan dilakukan menggunakan data X_{train} dan label y_{train} yang kemudian diuji dengan data X_{test} . Selanjutnya model *Support Vector Machine* dan *Naïve bayes* di uji pada data uji dan dievaluasi menggunakan beberapa metrik yaitu akurasi, presisi, recall, dan F1-score. *Confusion Matrix* untuk model *Support Vector Machine* dapat dilihat pada Gambar 3



Gambar 3. *Confusion Matrix SVM*

$$Accuracy = \frac{99 + 23 + 52}{(99 + 27 + 3) + (18 + 23 + 7) + (8 + 0 + 52)} = \frac{174}{237} \approx 0.73$$

Hasil evaluasi model menggunakan SVM berdasarkan data uji ditunjukkan pada Tabel 6.

Tabel 6. Hasil evaluasi SVM

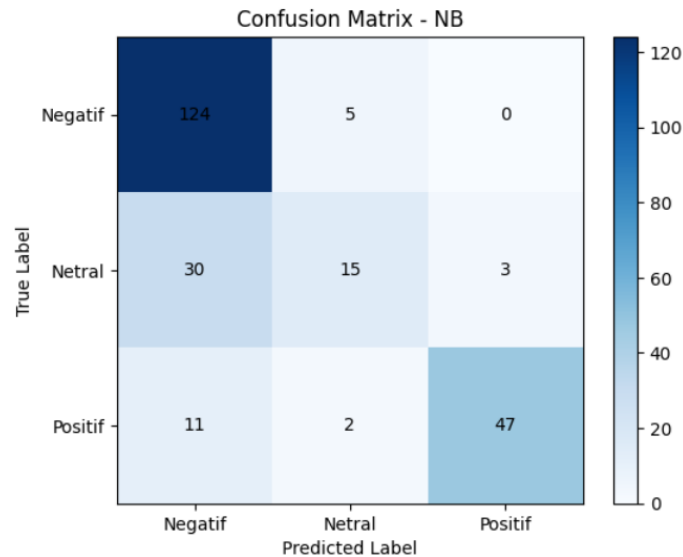
	Precision	recall	f1-score
Negatif	0,79	0,77	0,78
Netral	0,46	0,48	0,47
Positif	0,84	0,87	0,85

Berdasarkan perhitungan akurasi dari confusion matrix diperoleh nilai akurasi model Support Vector Machine sebesar 73%, yang menunjukkan bahwa model mampu mengklasifikasikan sekitar 73% data uji dengan benar. Dari hasil evaluasi berdasarkan Tabel 6, terlihat bahwa:

- **Kelas Positif** memiliki kinerja terbaik, dengan **presisi 0,84**, **recall 0,87**, dan **F1-score 0,85**. Ini berarti model sangat baik dalam mengklasifikasikan dan mengidentifikasi kelas Positif
- **Kelas Negatif** menunjukkan kinerja yang cukup seimbang, dengan **presisi 0,79**, **recall 0,77**, dan **F1-score sebesar 0,78**.

- **Kelas Netral** menunjukkan performa terendah, dengan **presisi 0,46**, **recall 0,46**, dan **f1-score sebesar 0,47**. Hal ini menunjukkan bahwa model kurang mampu mendeteksi data Netral secara selektif.

Sedangkan *Confusion Matrix* untuk model *Naïve Bayes Classifier* dapat dilihat pada Gambar 4.



Gambar 4. *Confusion Matrix* NB

$$Accuracy = \frac{124 + 15 + 47}{(124 + 5 + 0) + (30 + 15 + 3) + (11 + 2 + 47)} = \frac{186}{237} \approx 0.78$$

Hasil Evaluasi model menggunakan *Naïve Bayes* berdasarkan data uji ditunjukkan pada Tabel 7

Tabel 7. Hasil evaluasi *Naïve Bayes*

	Precision	recall	f1-score
Negatif	0,75	0,96	0,84
Netral	0,68	0,31	0,43
Positif	0,94	0,78	0,85

Berdasarkan perhitungan akurasi dari *confusion matrix* diperoleh nilai akurasi model *Naïve bayes* sebesar 78%, yang menunjukkan bahwa model mampu mengklasifikasikan sekitar 78% data uji dengan benar. Dari hasil evaluasi berdasarkan Tabel 7, terlihat bahwa:

- **Kelas Positif** memiliki kinerja terbaik, dengan **presisi 0,94**, **recall 0,78**, dan **F1-score 0,85**. Ini berarti model sangat baik dalam mengklasifikasikan dan mengidentifikasi kelas Positif
- **Kelas Negatif** menunjukkan kinerja yang cukup seimbang, dengan **presisi 0,75**, namun **recall yang sangat tinggi sebesar 0,96**, dan **F1-score sebesar 0,84**.
- **Kelas Netral** menunjukkan performa terendah, dengan **presisi 0,68**, **recall 0,31**, dan **f1-score sebesar 0,43**. Hal ini menunjukkan bahwa model kurang mampu mendeteksi data Netral secara efektif.

Penelitian ini memberikan hasil bahwa performa dari model *Naïve bayes* lebih baik dan akurat dibandingkan model *Support Vector Machine* dalam mengklasifikasikan ulasan aplikasi

wondr by BNI yang memberikan akurasi 78% dari *naïve bayes* dan 73% dari *Support vector machine*. Dari analisis ini, dapat disimpulkan bahwa kedua model memiliki kecenderungan berpihak pada kelas mayoritas (Negatif) dan kurang mampu mengenali kelas dengan jumlah data lebih sedikit (Netral). Dengan demikian, diperlukan langkah perbaikan, misalnya melalui penyeimbangan distribusi data, agar model dapat mengidentifikasi setiap kelas secara lebih proporsional.

DAFTAR PUSTAKA

- Agni V.P.D, Rudi K, Yudhistira A.W. 2024. Akurasi Naïve Bayes untuk Analisis Sentimen Twitter Berdasarkan Split Data. *Journal of Computing Engineering, System dan Science*. Vol. 9. No.1(238-250).
- Dedi Darwis, Nery Siskawati, & Zaenal Abidin. (2020). Penerapan Algoritma Naive Bayes untuk Analisis Sentimen Review Data Twitter BMKG Nasional. *Jurnal TEKNO KOMPAK*, 15(1), 131–145.
- Essa, A. F. (2024). *Analisis Sentimen Penilaian Aplikasi Wondr By Bni Menggunakan Modified K-Nearest Neighbor (MK-NN)*. Skripsi.
- Fatma Wati, F., Eko Widodo, A., & Korespondensi, P. (2025). Analisis Sentimen Ulasan Pengguna Aplikasi Deepseek Menggunakan Algoritma Random Forest Dan Naive Bayes. *Computer and Network Technology*, 5(1), 8–15. <http://jurnal.bsi.ac.id/index.php/content8>
- Fitriyah N., Warsito B., dan Maruddani D. A. I. 2020. Analisis Sentimen Gojek Pada Media Sosial Twitter Dengan Klasifikasi Support Vector Machine (SVM). *Jurnal Gaussian*, vol. 9, no. 3, pp. 376–390. doi: 10.14710/j.gauss.9.3.376-390.
- Handayanto dan Herlawati. 2020. Prediksi Kelas Jamak dengan Deep Learning Berbasis Graphics Processing Units. *Jurnal Kajian Ilmiah (JKI)*. Vol. 20 No.1.
- Hidayat, T. F. T., Garno, G., & Ridha, A. A. (2021). Analisis Sentimen Opini Pemindahan Ibu Kota Pada Twitter Dengan Metode Support Vector Machine. *Jurnal Ilmu Komputer*, 14(1), 49. <https://doi.org/10.24843/jik.2021.v14.i01.p06>
- Nadira A, Nanang Y. S, Welly P. 2023. Analisis Sentimen pada Ulasan Aplikasi Mobbile Banking menggunakan Metode Naïve Bayes dengan Kamus Inset. *Informatic and Computational Intelligent Journal*.
- Nurwanda, Suarna, N., & Prihartono, W. 2024. Penerapan Nlp (Natural Language Processing) Dalam Analisis Sentimen Pengguna Telegram Di Playstore. Jati (*Jurnal Mahasiswa Teknik Informatika*).
- Pratama M. R, Y.R Ramadhan, dan M. A. Komara. 2023. Analisis Sentimen BRImo dan BCA Mobile Menggunakan Support Vector Machine dan Lexicon Based. *Jurnal Ilmiah Teknik Informatika dan Sistem Informasi*.
- Puji Astuti, A., Alam, S., & Jaelani, I. (2022). Komparasi Algoritma Support Vector Machine dengan Naive Bayes Untuk Analisis Sentimen Pada Aplikasi BRImo. *Jurnal Bangkit Indonesia*, 11(2), 1–6. <https://doi.org/10.52771/bangkitindonesia.v11i2.196>
- Ratnaswari, S., Wibowo, N. C., Satria, D., & Kartika, Y. (2025). Analisis Sentimen Menggunakan Metode Lexicon-Based Dan Support Vector Machine Pada Presiden Dan Wakil Presiden Indonesia. *Jurnal Informatika dan Teknik Elektro Terapan*, 13(1), 362–368.
- Zafira, Z. T., Tania, K. D., Sari, W. K., Informasi, S., & Sriwijaya, U. (2025). *Sentiment-Based Knowledge Discovery of Wondr by BNI App Reviews Using SVM , KNN , and Naive Bayes for CRM Enhancement*. 9(5), 2498–2508.