

ANALISIS SENTIMEN BERBASIS ASPEK PADA ULASAN APLIKASI MIDI KRIING MENGGUNAKAN SUPPORT VECTOR MACHINE (SVM)

Rohmat Ubaidillah Fahmi, Amalia Anjani Arifiyanti, Tri Luhur Indayanti Sugata

Sistem Informasi, Universitas Pembangunan Nasional Veteran Jawa Timur

Jl.Raya Rungkut Madya, Gunung Anyar, Surabaya

21082010099@student.upnjatim.ac.id

ABSTRAK

Midi Kriing merupakan aplikasi belanja online yang dikembangkan oleh PT Midi Utama Indonesia Tbk, yang menawarkan berbagai produk untuk memenuhi kebutuhan konsumen. Untuk meningkatkan kualitas layanan dan pengalaman pengguna, analisis sentimen terhadap aplikasi Midi Kriing sangat diperlukan. Penelitian ini bertujuan untuk menganalisis sentimen berdasarkan aspek-aspek tertentu dari ulasan pengguna aplikasi Midi Kriing dengan menggunakan algoritma Support Vector Machine (SVM), dengan fokus pada tiga aspek utama: produk, layanan, dan fungsionalitas aplikasi. Berbagai skenario SVM diuji, skenario tersebut meliputi pembagian data dengan rasio 80:20 dan 75:25 untuk aspek produk, layanan, dan fungsionalitas aplikasi. Algoritma SVM yang digunakan melibatkan pemilihan kernel linear, polynomial, dan radial basis function (RBF), dengan pengujian nilai parameter C pada nilai 0.1, 1, dan 10. Evaluasi dilakukan menggunakan Classification Report dan Confusion Matrix untuk mengukur kinerja model dalam mengklasifikasikan sentimen, dengan metrik utama seperti precision, recall, F1-score, dan akurasi. Hasil penelitian menunjukkan bahwa model terbaik untuk masing-masing aspek diperoleh dengan skenario pembagian data 80:20 untuk produk dan layanan menggunakan kernel linear, serta pembagian data 75:25 untuk fungsionalitas aplikasi dengan kernel RBF, menghasilkan akurasi tertinggi sebesar 87% pada model produk dan fungsionalitas aplikasi dan 85% pada model layanan.

Kata kunci : *Midi Kriing, SVM, Analisis Sentimen Berbasis Aspek, Fleiss Kappa, Confussion Matrix.*

1. PENDAHULUAN

Pertumbuhan ekonomi Indonesia telah menjadi prioritas pemerintah dalam beberapa dekade terakhir. Salah satu faktor penting dalam transformasi ekonomi ini adalah kemajuan teknologi informasi dan komunikasi yang memicu perubahan besar dalam struktur ekonomi. Revolusi digital tidak hanya mempengaruhi cara kita berkomunikasi dan berinteraksi, tetapi juga mendorong inovasi dan pertumbuhan dalam sektor bisnis yang lebih berfokus pada teknologi. Digitalisasi membuka peluang baru dalam ekonomi, terutama di sektor-sektor tradisional seperti retail, yang kini beralih dari model konvensional menuju platform digital [1].

Salah satu aplikasi yang memfasilitasi perubahan ini adalah Midi Kriing, yang dikembangkan oleh PT Midi Utama Indonesia Tbk untuk mempermudah pengalaman berbelanja online. Aplikasi ini menawarkan berbagai fitur seperti pencarian produk yang mudah, pengaturan pengiriman fleksibel, dan sistem pembayaran yang aman. Midi Kriing memungkinkan pelanggan untuk berbelanja secara praktis melalui ponsel mereka, tanpa perlu datang ke toko fisik, yang meningkatkan efisiensi dan kepuasan pelanggan. Dengan lebih dari 1.781.613 transaksi, aplikasi ini menjadi salah satu bukti kesuksesan PT Midi Utama Indonesia Tbk dalam menghadapi perkembangan teknologi [2].

Untuk mempertahankan tingkat kepuasan pelanggan, perusahaan memanfaatkan ulasan aplikasi sebagai umpan balik langsung yang dapat memberikan wawasan tentang kelebihan dan kekurangan layanan. Ulasan ini menjadi sumber informasi penting bagi

perusahaan dalam meningkatkan pengalaman berbelanja. Ulasan tertulis di Google Play Store memberikan detail yang lebih jelas dibandingkan dengan rating numerik, yang sering kali tidak mengungkapkan aspek-aspek spesifik dari pengalaman pengguna. Oleh karena itu, analisis sentimen berbasis aspek dapat membantu perusahaan untuk lebih memahami perhatian pelanggan terhadap fitur tertentu dari aplikasi [3].

Analisis sentimen berbasis aspek (Aspect-Based Sentiment Analysis) memberikan pendekatan yang lebih mendalam dalam mengevaluasi pengalaman pengguna. Sementara rating numerik memberikan gambaran umum, ulasan berbasis teks memungkinkan perusahaan untuk mengidentifikasi aspek layanan yang paling dihargai atau yang perlu perbaikan. Dengan menggunakan ABSA, PT Midi Utama Indonesia Tbk dapat menganalisis berbagai aspek seperti harga produk, layanan, dan fungsionalitas aplikasi, serta menyesuaikan layanannya sesuai dengan umpan balik pengguna [4]. Metode ini juga memberikan cara untuk mengetahui sentimen positif atau negatif terkait setiap aspek yang dibahas dalam ulasan [5].

Support Vector Machine (SVM) merupakan algoritma machine learning yang sangat efektif untuk klasifikasi sentimen berbasis aspek. Dalam penelitian yang dilakukan oleh Jessica Widyadhana Iskandar dan Yessica Nataliani pada tahun 2021 mengenai Perbandingan Naïve Bayes, SVM, dan k-NN untuk Analisis Sentimen Gadget Berbasis Aspek, SVM memiliki rata-rata accuracy SVM sebesar 96.43% dan menjadi paling tinggi dibandingkan accuracy

algoritma lainnya [6]. Dengan menggunakan SVM, perusahaan dapat memperoleh klasifikasi sentimen yang lebih tepat, membantu mereka dalam mengambil keputusan strategis untuk meningkatkan layanan. Penelitian ini bertujuan untuk menghasilkan model yang dapat mengidentifikasi sentimen pelanggan terhadap aspek-aspek aplikasi, yang nantinya akan membantu PT Midi Utama Indonesia Tbk dalam meningkatkan kepuasan pelanggan dan kualitas layanan Midi Kriing.

2. TINJAUAN PUSTAKA

2.1. Analisis Sentimen

Analisis sentimen adalah metode untuk memahami opini atau sikap seseorang pada suatu topik, produk, atau layanan melalui teks, yang sering dikategorikan sebagai positif, negatif, atau netral [7]. Teknik ini penting dalam pemrosesan bahasa alami (NLP) dan banyak digunakan dalam bisnis, politik, dan penelitian sosial untuk mendapatkan wawasan dari ulasan online atau media sosial. Dalam bisnis, analisis ini membantu perusahaan menilai kepuasan pelanggan dan membuat keputusan strategis untuk meningkatkan layanan atau produk. Metode yang digunakan dalam analisis sentimen meliputi pembelajaran mesin (seperti SVM, Naive Bayes, dan Random Forest) dan teknik lebih lanjut seperti Aspect-Based Sentiment Analysis (ABSA) untuk analisis yang lebih mendalam.

2.2. Data Preprocessing

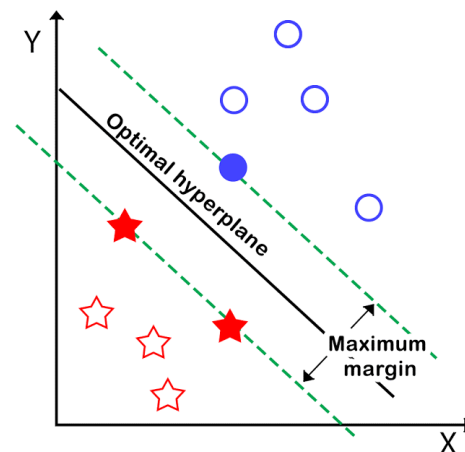
Proses data preprocessing adalah langkah awal untuk menyiapkan data agar siap dianalisis dalam data mining, NLP, dan machine learning [8]. Tahapan ini meliputi teknik seperti case folding (mengubah teks menjadi huruf kecil), data cleaning (menghapus elemen mengganggu seperti tautan dan simbol), normalisasi teks, stopwords removal (menghilangkan kata umum yang tidak penting), dan stemming (mengubah kata berimbuhan ke bentuk dasarnya). Teknik-teknik ini memastikan teks siap untuk dianalisis dan diekstraksi menjadi informasi bermakna.

2.3. Fleiss Kappa

Fleiss' Kappa adalah metode untuk mengukur tingkat kesepakatan antar penilai/annotator dalam memberikan label pada data, yang dapat melibatkan lebih dari dua penilai, sehingga lebih fleksibel dibandingkan Cohen's Kappa. Metode ini menghitung kesepakatan yang terjadi melebihi kemungkinan kebetulan, memberikan gambaran yang lebih jelas tentang konsistensi penilaian antar penilai. Fleiss' Kappa umumnya digunakan untuk mengukur reliabilitas dalam klasifikasi yang melibatkan beberapa penilai, misalnya dalam penelitian yang melibatkan analisis data kualitatif atau penandaan kategori. Dengan demikian, ini menjadi alat penting dalam menilai seberapa konsisten hasil klasifikasi yang diberikan oleh berbagai annotator atau penilai [9].

2.4. Support Vector Machine (SVM)

Support Vector Machine (SVM) adalah teknik yang digunakan untuk prediksi, baik dalam klasifikasi maupun regresi. Dasar dari prinsip SVM adalah sebagai linear classifier, namun kemudian dikembangkan untuk menangani masalah non-linear. SVM dapat diterapkan pada masalah non-linear dengan memanfaatkan konsep kernel yang membawa data ke ruang berdimensi tinggi. Di ruang berdimensi tinggi ini, SVM berusaha menemukan hyperplane yang memaksimalkan margin antara dua kelas data [10].



Gambar 1. Cara Kerja SVM

Secara sederhana, konsep SVM berfokus pada pencarian hyperplane terbaik yang dapat memisahkan dua kelas dalam ruang input. Seperti yang terlihat pada Gambar 1, hyperplane yang memisahkan kedua kelas akan dicari dengan mengukur margin dan memilih nilai maksimum margin tersebut. Margin ini merupakan jarak antara hyperplane dan pola data yang paling dekat dari masing-masing kelas.

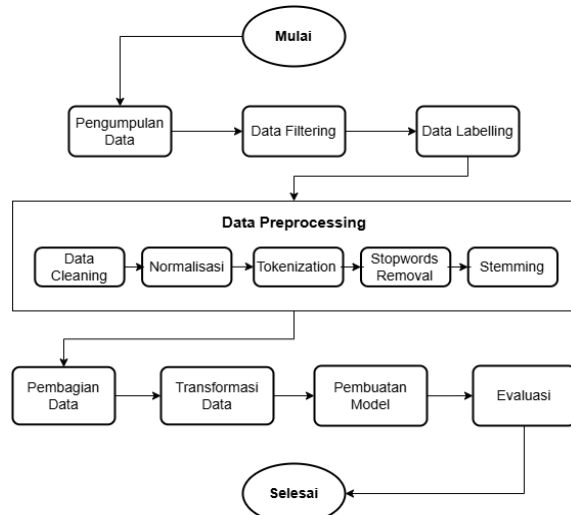
Untuk mengimplementasikan SVM dalam Python, library scikit-learn adalah salah satu pilihan yang paling populer. Library ini menyediakan antarmuka yang sederhana untuk membangun dan melatih model SVM, serta berbagai alat untuk pra pemrosesan data dan evaluasi model. Scikit-learn sering digunakan dalam berbagai literatur, terutama dalam tugas klasifikasi teks, seperti analisis sentimen atau deteksi sarkasme, berkat kemudahannya implementasinya dan fungsi-fungsi yang disediakan.

2.5. Confusion Matrix

Confusion Matrix adalah alat yang digunakan untuk mengevaluasi kinerja model klasifikasi dengan membandingkan hasil prediksi model dengan kelas yang sebenarnya. Matriks ini terdiri dari empat elemen utama: True Positive (TP), True Negative (TN), False Positive (FP), dan False Negative (FN). Dari elemen-elemen ini, beberapa metrik dapat dihitung, seperti Akurasi (Accuracy), yang mengukur proporsi prediksi yang benar, Precision, yang menunjukkan seberapa banyak prediksi positif yang sebenarnya positif,

Recall, yang menilai sejauh mana sampel positif berhasil teridentifikasi dengan benar, dan F1-Score, yang menggabungkan Precision dan Recall untuk memberikan gambaran umum tentang kinerja model [11].

3. METODE PENELITIAN



Gambar 2. Metodologi Penelitian

3.1. Analisis Kebutuhan

Untuk mendukung penelitian ini, diperlukan adanya data ulasan yang nantinya akan diolah, serta spesifikasi hardware maupun software yang memadai untuk mengolah data ulasan.

a. Kebutuhan Data

Dalam penelitian ini, data yang dibutuhkan yaitu data ulasan pengguna aplikasi yang diberikan oleh pengguna aplikasi Midi Kriing di Platform Google Play Store dan App Store. Data ulasan yang diambil adalah data ulasan yang berbahasa Indonesia yang diambil dalam kurun waktu bulan September 2021 sampai Oktober 2024. Data ulasan tersebut menjadi bahan utama yang digunakan analisis sentimen berbasis aspek pada penelitian ini.

b. Kebutuhan Software dan Hardware

Dalam penelitian ini, peneliti menggunakan beberapa software yang meliputi web browser, Google Colab, Google Spreadsheet, Google AppSheet, Figma, Microsoft Excel, Visual Studio Code, dan Google Docs. Adapun untuk kebutuhan hardware, peneliti menggunakan laptop Asus TUF GAMING F15 dengan RAM 8GB + 16GB dan penyimpanan internal SSHD 477 GB.

3.2. Pengumpulan Data

Proses pengumpulan data ulasan dilakukan dengan teknik scraping menggunakan bahasa pemrograman Python pada platform Google Colab yang bersumber dari Google Play Store dan App Store. Peneliti menginstall library app-store-scraper dan google-play-scraper untuk mengambil data dari kedua platform tersebut. Dua kode scraping digunakan untuk

mengumpulkan ulasan, satu untuk Google Play Store dan satu untuk App Store, yang menghasilkan file CSV terpisah. Kedua file CSV tersebut digabungkan menjadi satu file "Ulasan_MIDI_All.csv" yang berisi empat kolom: Rating, Review Text, Date, dan Platform. Hasil pengumpulan data terdiri dari 7158 ulasan, dengan 6598 data dari Google Play Store dan 560 data dari App Store.

3.3. Data Filtering

Pada tahap penyaringan data, ulasan Midi Kriing yang diperoleh melalui scraping diseleksi agar hanya data yang relevan yang digunakan dalam analisis. Proses dimulai dengan case folding untuk mengubah seluruh teks menjadi huruf kecil, menghapus baris yang hanya berisi emotikon yang tidak memberikan informasi penting, serta menghapus data duplikat untuk menjaga konsistensi dataset. Setelah melalui tahapan ini, jumlah ulasan yang tersisa adalah 5724, yang kemudian digunakan untuk tahap pelabelan manual oleh tiga pelabel. Data yang telah difilter ini siap digunakan dalam model analisis sentimen atau topik, memastikan dataset final yang bersih dan relevan.

3.4. Data Labelling

Setelah proses data filtering yang menghasilkan 5724 data ulasan, tahap selanjutnya adalah data labelling, di mana setiap ulasan dilabeli berdasarkan aspek dan sentimen yang terkandung. Tiga aspek yang ditentukan adalah produk, layanan, dan fungsionalitas aplikasi. Proses pelabelan dilakukan secara manual oleh tiga pelabel untuk meningkatkan akurasi dan mengurangi bias. Pedoman pelabelan memastikan konsistensi, dengan aspek produk mencakup kualitas, kelengkapan, dan harga produk, aspek layanan mencakup pengalaman pelanggan mengenai pengiriman atau pelayanan, dan aspek fungsional aplikasi mencakup kinerja, desain, dan fitur aplikasi.

3.5. Data Preprocessing

Tahap Data Preprocessing dimulai setelah pengumpulan data dan bertujuan untuk memastikan bahwa data yang digunakan sudah dalam format yang sesuai dan dapat diolah dengan baik. Proses ini dilakukan dalam beberapa langkah sebagai berikut :

a. Data Cleaning

Pada tahap ini, data yang terkumpul disaring dari elemen-elemen yang tidak relevan seperti emoticon, tanda baca, angka, tanda kurung, tanda petik, dan karakter khusus yang dapat mengganggu proses analisis.

b. Normalisasi

Proses normalisasi bertujuan untuk menyamakan variasi kata menjadi bentuk yang baku dan standar. Hal ini dilakukan dengan mengganti singkatan dengan bentuk formal, menghilangkan duplikasi karakter, dan menggunakan kamus bahasa baku yang sudah disusun.

c. Tokenization

Tokenization bertujuan untuk memecah teks menjadi unit-unit kecil yang disebut token. Setiap token akan menjadi elemen yang dapat dianalisis lebih lanjut dalam model analisis sentimen.

d. Stopwords Removal

Proses ini menghapus kata-kata yang sering muncul dalam teks namun tidak memberikan informasi signifikan dalam analisis sentimen, seperti kata hubung dan kata pengisi.

e. Stemming

Stemming mengubah kata turunan menjadi bentuk dasar. Dengan demikian, kata-kata yang memiliki makna serupa akan disederhanakan, yang memungkinkan model untuk menganalisis teks secara lebih efisien.

3.6. Pembagian Data

Setelah tahap Data Preprocessing, data dibagi menjadi dua bagian: 90% untuk pelatihan dan 10% untuk validasi. Data pelatihan digunakan untuk melatih model Support Vector Machine (SVM) dalam mengenali pola sentimen, sementara data validasi digunakan untuk menguji kemampuan model dalam menggeneralisasi pada data yang belum pernah dipelajari sebelumnya. Untuk memastikan distribusi sentimen yang seimbang, diterapkan teknik stratifikasi, sehingga setiap label sentimen (positif, negatif, tidak ada) terdistribusi secara proporsional di kedua dataset.

3.7. Transformasi Data

Setelah tahap pembagian data langkah berikutnya adalah transformasi data menggunakan teknik TF-IDF (Term Frequency-Inverse Document Frequency) untuk mengubah teks menjadi format numerik yang dapat dipahami oleh algoritma Support Vector Machine (SVM). Teknik ini memberikan bobot pada kata-kata berdasarkan seberapa sering mereka muncul dalam dokumen dan korpus, sehingga kata-kata yang lebih penting akan memiliki bobot yang lebih tinggi. Implementasi TF-IDF dilakukan dengan menggunakan `TfidfVectorizer` dari `Sklearn` pada data pelatihan, untuk memastikan konsistensi dan mencegah data leakage. Hasil pembobotan ini akan digunakan dalam pemodelan untuk meningkatkan akurasi dalam mendeteksi pola sentimen oleh SVM.

3.8. Pembuatan Model

Setelah tahap transformasi data, langkah berikutnya adalah pembuatan model analisis sentimen menggunakan algoritma Support Vector Machine (SVM), yang efektif untuk klasifikasi data berdimensi tinggi, seperti teks. Pada tahap ini, model SVM dilatih untuk mengklasifikasikan sentimen berdasarkan fitur numerik yang dihasilkan dari teknik TF-IDF. Data pelatihan yang telah diproses digunakan untuk melatih model agar dapat mengenali pola-pola sentimen seperti positif, negatif, atau netral. Proses pelatihan dilakukan dengan membagi data menjadi beberapa

skenario, seperti pembagian 80:20 dan 75:25 antara data pelatihan dan data pengujian, untuk menguji kemampuan model dalam menggeneralisasi data yang belum pernah dilihat sebelumnya. Selain itu, eksperimen dilakukan dengan menggunakan tiga jenis kernel SVM, yaitu kernel linear, kernel RBF, dan kernel polynomial, serta pengaturan nilai parameter C yang bervariasi, yaitu 0.1, 1, dan 10, untuk melihat dampaknya terhadap kinerja model. Melalui berbagai skenario dan konfigurasi ini, diharapkan model SVM dapat menghasilkan klasifikasi sentimen yang akurat dan optimal, serta memberikan gambaran tentang bagaimana model menangani variasi dalam data dan seberapa baik ia memprediksi sentimen berdasarkan pelatihan yang telah dilakukan.

3.9. Evaluasi

Setelah model Support Vector Machine (SVM) dilatih, tahap berikutnya adalah evaluasi untuk mengukur kinerjanya dalam mengklasifikasikan sentimen pada data yang belum pernah dilihat. Evaluasi dilakukan dengan menggunakan Classification Report dan Confusion Matrix.

Classification Report memberikan metrik seperti akurasi, precision, recall, dan F1-score untuk menilai kinerja model dalam mengklasifikasikan sentimen. Precision mengukur keakuratan prediksi yang relevan, recall mengukur seberapa banyak prediksi relevan yang seharusnya dilakukan, dan F1-score memberikan gambaran seimbang antara precision dan recall.

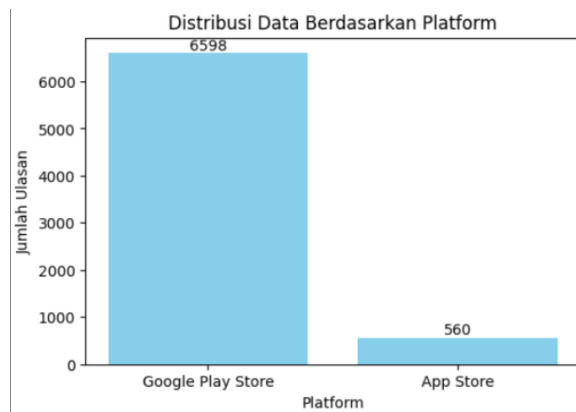
Confusion Matrix memberikan visualisasi yang lebih rinci tentang prediksi yang benar dan salah, menunjukkan empat kategori: True Positive (TP), True Negative (TN), False Positive (FP), dan False Negative (FN). Dengan kedua alat ini, kita bisa menilai performa model SVM pada data pengujian dan memilih model terbaik berdasarkan hasil evaluasi tersebut.

4. HASIL DAN PEMBAHASAN

Bab ini berisi hasil dan pembahasan dari penelitian yang dilakukan terkait analisis sentimen berbasis aspek pada ulasan aplikasi Midi Kriing dengan menggunakan algoritma support vector machine.

4.1. Pengumpulan Data

Pada tahap pengumpulan data ulasan, peneliti menggunakan teknik scraping dengan bahasa pemrograman Python pada platform Google Colab untuk mengumpulkan ulasan dari dua sumber utama, yaitu Google Play Store dan App Store. Proses pengumpulan dimulai dengan menginstal dua library yang diperlukan, yaitu `app-store-scraper` untuk App Store dan `google-play-scraper` untuk Google Play Store. Hasil pengumpulan data terdiri dari 7158 ulasan, dengan 6598 data dari Google Play Store dan 560 data dari App Store.



Gambar 3. Grafik Distribusi Data Ulasan Midi Kriing

4.2. Data Filtering

Setelah data ulasan Midi Kriing terkumpul, tahap berikutnya adalah data filtering untuk menghapus data yang tidak relevan sebelum dilakukan proses pelabelan. Pada tahap ini, data awal yang terkumpul berjumlah 7158 ulasan. Langkah pertama dalam data filtering adalah case folding, di mana semua teks ulasan diubah menjadi huruf kecil (lowercase) untuk menghindari perbedaan antara huruf kapital dan huruf kecil. Dengan demikian, kata yang sama seperti "MIDI" dan "midi" dianggap setara. Selain itu, dilakukan penghapusan baris yang hanya berisi emoticon, karena emoticon tidak memberikan informasi yang berguna dalam analisis sentimen atau topik.

Selanjutnya, dilakukan penghapusan data duplikat untuk memastikan setiap ulasan dihitung hanya sekali dalam analisis. Proses ini memeriksa dan menghapus ulasan yang identik, baik dari segi konten maupun format lainnya. Setelah melalui semua tahapan filtering, jumlah data yang tersisa adalah 5724 ulasan. Data yang telah difilter ini kemudian akan dilanjutkan ke tahap pelabelan manual oleh tiga pelabel untuk proses analisis sentimen selanjutnya.

4.3. Data Labelling

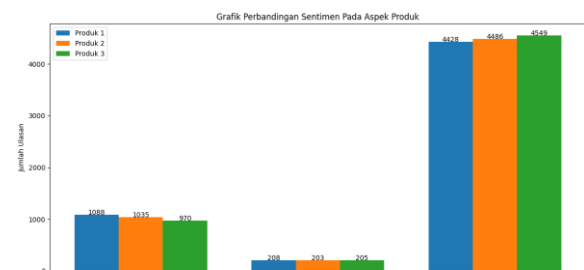
Setelah melalui data filtering, menghasilkan dataset ulasan sebanyak 5724 data yang akan dilabeli oleh 3 orang pelabel secara manual. Dalam proses pelabelan data ulasan, diberikan pedoman mengenai jenis ulasan yang termasuk dalam aspek tertentu, dengan tujuan untuk menjaga konsistensi dalam pelabelan. Berikut adalah pedoman pelabelan untuk data ulasan Midi Kring :

- Aspek Produk = Ulasan yang membahas mengenai kualitas, kelengkapan, atau harga produk yang dijual.
- Aspek Layanan = Ulasan yang membahas mengenai pengalaman pelanggan mengenai layanan pengiriman, service dari pelayan toko, atau layanan lainnya.
- Aspek Fungsional Aplikasi = Ulasan yang membahas mengenai kinerja aplikasi, desain, fitur, atau masalah teknis.

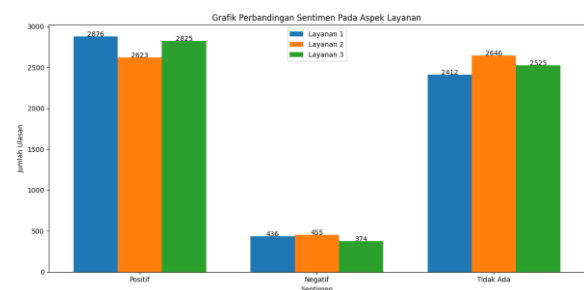
Ada tiga label sentimen pada masing-masing aspek, yaitu positif, negatif, dan tidak ada. Berikut adalah penjelasan mengenai sentimen dalam pelabelan.

- Positif : Label "Positif" digunakan ketika aspek yang bersangkutan memiliki nilai yang baik dalam ulasan.
- Negatif : Label "Negatif" digunakan ketika aspek yang bersangkutan memiliki nilai yang buruk dalam ulasan.
- Tidak Ada : Label "Tidak Ada" digunakan apabila aspek yang bersangkutan tidak disebutkan atau tidak ada dalam ulasan.

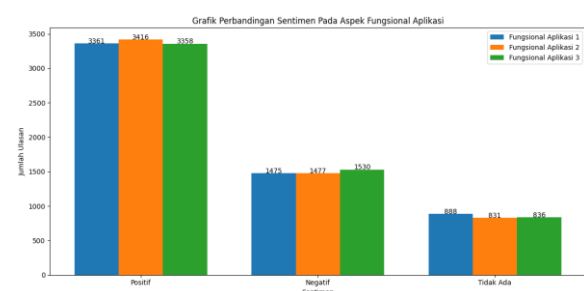
Pedoman ini digunakan untuk memastikan bahwa setiap ulasan diberi label yang sesuai dengan topik yang dibahas, sehingga pelabelan tetap konsisten. Berikut merupakan hasil ulasan yang telah dilabeli dalam Google Spreadsheet.



Gambar 4. Hasil Pelabelan Aspek Produk



Gambar 5. Hasil Pelabelan Aspek Layanan



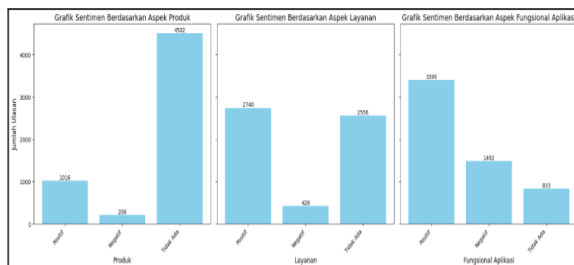
Gambar 6. Hasil Pelabelan Aspek Fungsional Aplikasi

Setelah pelabelan data ulasan selesai, peneliti menghitung nilai kesepakatan antar pelabel menggunakan metode Fleiss Kappa. Berikut merupakan hasilnya.

Tabel 1. Hasil Fleiss Kappa

Aspek	Nilai Kappa	Interpretasi
Produk	0.79	Kesepakatan Substansial
Layanan	0.8	Kesepakatan Substansial
Fungsional Aplikasi	0.88	Kesepakatan Hampir Sempurna

Hasil menampilkan pada aspek produk menunjukkan nilai 0.79 (kesepakatan substansial), pada aspek layanan menunjukkan nilai 0.8 (kesepakatan substansial) dan pada aspek fungsional aplikasi menunjukkan nilai 0.88 (kesepakatan hampir sempurna). Dengan ketiga nilai fleiss kappa pada masing-masing aspek, membuktikan bahwa tingkat kesepakatan pelabel tinggi dan pelabel cukup konsisten dalam melabeli data ulasan. Tahapan selanjutnya adalah menentukan sentimen final pada dataset ulasan. Penentuan sentimen final menggunakan modus, dan berikut merupakan diagram sentimen pada dataset ulasan.



Gambar 7. Distribusi Sentimen Pada Masing-Masing Aspek

4.4. Data Preprocessing

Tahap Data Preprocessing merupakan langkah krusial untuk mempersiapkan data sebelum dilakukan analisis lebih lanjut. Tujuan utamanya adalah mengubah format data agar siap digunakan dalam pelatihan model analisis sentimen berbasis aspek dengan algoritma Support Vector Machine (SVM). Proses ini mencakup beberapa tahapan, seperti data cleaning, case folding, normalisasi, tokenisasi, penghapusan stopwords, dan stemming.

4.5. Data Cleaning

Proses penghapusan karakter atau kata yang tidak relevan, seperti tanda baca atau simbol yang mengganggu analisis teks.

Tabel 2. Data Cleaning

Review Text	Cleaning Text
penipuan !!! transaksi dibatalkan tapi uang tidak kembali,	penipuan transaksi dibatalkan tapi uang tidak kembali

4.6. Normalisasi

Langkah untuk menyamakan bentuk kata atau frasa agar konsisten, seperti mengganti kata-kata yang tidak baku menjadi bentuk baku.

Tabel 3. Normalisasi

Cleaning Text	Normalisasi Text
penipuan transaksi dibatalkan tapi uang tidak kembali	penipuan transaksi dibatalkan tapi uang tidak kembali

4.7. Tokenization

Proses memecah teks menjadi unit-unit terkecil berupa kata atau token yang dapat dianalisis lebih lanjut.

Tabel 4. Tokenization

Normalisasi Text	Tokenization Text
penipuan transaksi dibatalkan tapi uang tidak kembali	['penipuan', 'transaksi', 'dibatalkan', 'tapi', 'uang', 'tidak', 'kembali']

4.8. Stopwords Removal

Penghapusan kata-kata umum yang tidak memberikan makna penting dalam analisis, seperti "dan", "atau", "yang".

Tabel 5. Stopwords Removal

Tokenization Text	Stopwords Removal Text
['penipuan', 'transaksi', 'dibatalkan', 'tapi', 'uang', 'tidak', 'kembali']	['penipuan', 'transaksi', 'dibatalkan', 'uang']

4.9. Stemming

Proses pengembalian kata ke bentuk dasarnya, menghapus imbuhan yang tidak relevan untuk analisis.

Tabel 6. Stemming

Stopwords Removal Text	Stemming Text
['penipuan', 'transaksi', 'dibatalkan', 'uang']	['tipu', 'transaksi', 'batal', 'uang']

4.10. Pembagian Data

Pada tahap pembagian data, dataset hasil preprocessing akan dipisahkan menjadi dua bagian, yaitu data untuk model dan data untuk validasi. Sebanyak 90% dari data digunakan untuk melatih model, sementara 10% sisanya digunakan untuk proses validasi. Pembagian ini dilakukan untuk memastikan model dapat diuji dengan data yang belum pernah dilihat sebelumnya. Untuk menjaga keseimbangan distribusi sentimen di seluruh aspek, digunakan teknik stratifikasi dalam pembagian data. Dengan demikian, setiap label sentimen akan terwakili secara proporsional dalam kedua dataset tersebut. Hasil pembagian data menggunakan teknik stratifikasi sebagai berikut :

- Dataset Model = 5149 Data Ulasan
- Dataset Validasi = 573 Data Ulasan

4.11. Transformasi Data

Pada tahap ini, proses TF-IDF diterapkan pada tiga model, yaitu produk, layanan, dan fungsi aplikasi, dengan pembagian data 80:20 dan 75:25. Dalam setiap skenario, 80% atau 75% data digunakan untuk pelatihan, dan sisanya untuk pengujian. Proses ini

mengubah teks ulasan menjadi representasi numerik, di mana setiap kata diberi bobot berdasarkan pentingnya dalam dokumen.

Setiap model diproses secara terpisah, menghasilkan matriks numerik yang mempermudah analisis. Penerapan TF-IDF ini membantu model fokus pada kata-kata yang paling relevan dalam setiap kategori ulasan, meningkatkan akurasi analisis dan prediksi.

4.12. Pembuatan Model

Pada tahap Pembuatan Model, peneliti membangun model pengklasifikasi sentimen untuk setiap aspek menggunakan algoritma Support Vector Machine (SVM). Proses ini melibatkan serangkaian eksperimen dengan berbagai skenario untuk memastikan bahwa model yang dihasilkan memiliki performa yang optimal. Secara keseluruhan, sebanyak 36 skenario diuji untuk setiap model aspek yang digunakan dalam penelitian ini.

Dalam eksperimen ini, pembagian data dilakukan dengan dua proporsi yang berbeda: 80:20 dan 75:25, yang berarti 80% atau 75% dari data digunakan sebagai data pelatihan (training data), sementara sisanya digunakan sebagai data pengujian (testing data). Pembagian ini bertujuan untuk mengevaluasi konsistensi model dalam berbagai proporsi data yang digunakan.

Selanjutnya, algoritma SVM diterapkan dengan tiga jenis kernel yang berbeda, yaitu linear, polynomial, dan RBF (Radial Basis Function). Masing-masing kernel ini memiliki karakteristik yang berbeda dalam menangani masalah klasifikasi, di mana kernel linear digunakan untuk data yang dapat dipisahkan secara linier, kernel polynomial digunakan untuk data dengan hubungan yang lebih kompleks, dan kernel RBF digunakan untuk menangani data non-linear dengan memetakan data ke ruang dimensi tinggi.

Selain itu, nilai C yang digunakan dalam eksperimen ini diuji dengan tiga nilai yang berbeda, yaitu 0.1, 1, dan 10. Parameter C ini berfungsi untuk mengontrol trade-off antara memaksimalkan margin dan meminimalkan kesalahan klasifikasi. Nilai C yang lebih besar cenderung membuat model lebih sensitif terhadap kesalahan pelatihan, sedangkan nilai C yang lebih kecil memungkinkan model untuk lebih generalis dalam menghadapi data uji.

Skenario yang sama juga diterapkan pada data yang telah mengalami proses over-sampling menggunakan SMOTE (Synthetic Minority Over-sampling Technique) untuk mengatasi masalah ketidakseimbangan kelas. SMOTE digunakan untuk menghasilkan data sintesis dari kelas minoritas, sehingga membantu model dalam mengatasi bias terhadap kelas mayoritas dan meningkatkan akurasi prediksi pada kelas yang lebih jarang.

Dengan pendekatan ini, diharapkan dapat diperoleh model SVM yang tidak hanya akurat dalam melakukan klasifikasi, tetapi juga mampu memberikan

hasil yang lebih seimbang meskipun menghadapi ketidakseimbangan kelas dalam data.

4.13. Evaluasi

Pada tahap evaluasi, model diuji menggunakan Classification Report dan Confusion Matrix untuk menilai kinerja model dalam mengklasifikasikan sentimen. Classification Report menghitung precision, recall, F1-score, dan accuracy menggunakan rumus berikut :

- Precision : $\frac{TP}{TP+FP}$
- Recall : $\frac{TP}{TP+FN}$
- F1 Score : $\frac{2 * Precision * recall}{Precision + recall}$
- Accuracy : $\frac{TP+TN}{TP+TN+FP+FN}$

Dalam penelitian ini, 36 skenario diuji untuk masing-masing model aspek. Model terbaik dipilih berdasarkan nilai precision, recall, F1-score, dan accuracy tertinggi. Model terbaik untuk masing-masing aspek adalah sebagai berikut :

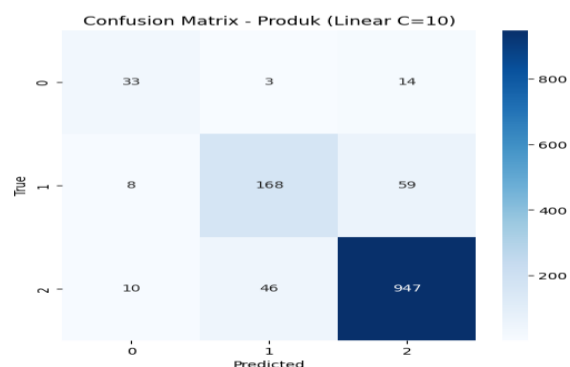
Tabel 7. Skenario Terbaik

Model	Precision	Recall	F1 Score	Accuracy
Produk	0.75	0.74	0.74	0.87
Layanan	0.84	0.83	0.83	0.85
Fungsional Aplikasi	0.82	0.8	0.81	0.87

- Produk: Pembagian data 80:20, kernel linear, C = 10
- Layanan: Pembagian data 80:20, kernel linear, C = 1
- Fungsional Aplikasi: Pembagian data 75:25, kernel RBF, C = 10

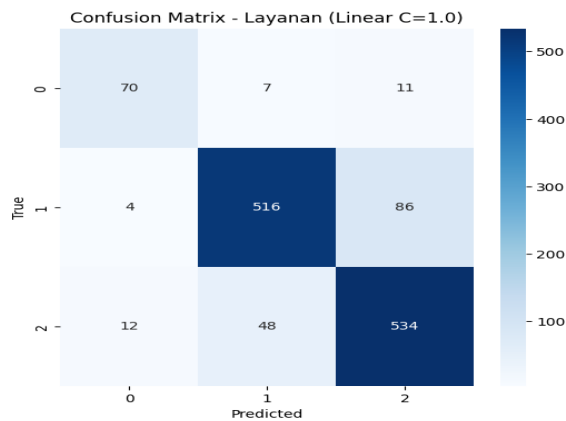
Confusion Matrix memberikan visualisasi lebih rinci tentang prediksi yang benar dan salah, yaitu True Positive (TP), True Negative (TN), False Positive (FP), dan False Negative (FN). Berikut merupakan contoh confusion matrix pada 3 model pada aspek produk, layanan dan fungsional aplikasi. Kelas sentimen di encoding sebagai berikut :

- Sentimen “Tidak Ada” bernilai 0
- Sentimen “Positif” bernilai 1
- Sentimen “Negatif” bernilai 2



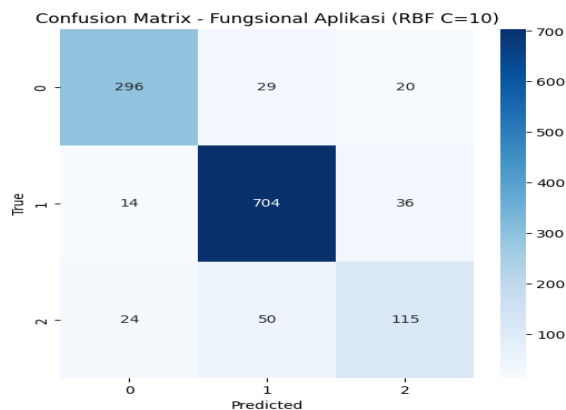
Gambar 8. Model Aspek Produk

Pada gambar 8 menunjukkan output confusion matrix pada model produk dengan skenario terbaik yaitu, pembagian data 80:20, kernel linear, nilai $C = 10$. Model klasifikasi ini menunjukkan kinerja yang baik dengan tingkat akurasi tinggi, terutama pada kelas 2, meskipun masih terdapat kesalahan prediksi pada kelas 0 dan 1.



Gambar 9. Model Aspek Layanan

Pada gambar 9 menunjukkan output confusion matrix pada model layanan dengan skenario terbaik yaitu, pembagian data 80:20, kernel linear, nilai $C = 1$. Model klasifikasi untuk layanan menunjukkan akurasi yang baik pada kelas 1, namun masih terdapat kesalahan prediksi pada kelas 0 dan 2 yang perlu diperbaiki.



Gambar 10. Model Aspek Fungsional Aplikasi

Pada gambar 10 menunjukkan output confusion matrix pada model fungsionalitas aplikasi dengan skenario terbaik yaitu, pembagian data 75:25, kernel RBF, nilai $C = 10$. Model klasifikasi untuk fungsional aplikasi menunjukkan performa yang sangat baik, dengan akurasi tinggi pada kelas 1, meskipun terdapat beberapa kesalahan prediksi pada kelas 0 dan 2.

5. KESIMPULAN DAN SARAN

Penelitian ini berhasil menganalisis sentimen berbasis aspek pada ulasan aplikasi Midi Kriing menggunakan algoritma Support Vector Machine

(SVM), dengan fokus pada tiga aspek utama: produk, layanan, dan fungsionalitas aplikasi. Hasil penelitian menunjukkan bahwa untuk model produk dan fungsional aplikasi menghasilkan akurasi sebesar 87%. Untuk model layanan menghasilkan akurasi sebesar 85%. Temuan ini diharapkan dapat memberikan wawasan penting bagi PT Midi Utama Indonesia Tbk untuk meningkatkan kualitas aplikasi dengan fokus pada aspek yang kurang memuaskan, sambil mempertahankan keunggulan di area yang sudah diterima baik oleh pengguna. Sebagai saran, perusahaan disarankan untuk terus melakukan analisis sentimen secara berkala dan mempertimbangkan penggunaan algoritma lain untuk meningkatkan akurasi, serta menambah aspek analisis untuk memperoleh gambaran yang lebih komprehensif tentang pengalaman pengguna.

DAFTAR PUSTAKA

- [1] F. Abdillah, "Dampak Ekonomi Digital Terhadap Pertumbuhan Ekonomi di Indonesia," *Benefit J. Bussiness Econ. Finance*, vol. 2, no. 1, pp. 27–35, Feb. 2024, doi: 10.37985/benefit.v2i1.335.
- [2] Alfamidi, "Pencapaian 17 Tahun Alfamidi Melayani Indonesia," 2024. [Online]. Available: <https://alfamidiku.com/berita/perusahaan/pencapaian-17-tahun-alfamidi-melayani-indonesia>
- [3] M. Anggraini, R. Rahmadani, and S. Priyono, "Pengaruh Kualitas Produk Dan Ulasan Produk Terhadap Keputusan Pembelian Produk Pada Aplikasi Shopee Mahasiswa Pendidikan Ekonomi Semester V Universitas Nurul Huda," *JECO J. Econ. Educ. Eco-Technopreneurship*, vol. 2, no. 1, pp. 25–31, Jul. 2023, doi: 10.30599/jeco.v2i1.233.
- [4] G. Radiena and A. Nugroho, "ANALISIS SENTIMEN BERBASIS ASPEK PADA ULASAN APLIKASI KAI ACCESS MENGGUNAKAN METODE SUPPORT VECTOR MACHINE," *J. Pendidik. Teknol. Inf. JUKANTI*, vol. 6, no. 1, pp. 1–10, Apr. 2023, doi: 10.37792/jukanti.v6i1.836.
- [5] Syaifulloh Amien Pandega Perdana, Teguh Bharata Aji, and Ridi Ferdiana, "Aspect Category Classification dengan Pendekatan Machine Learning Menggunakan Dataset Bahasa Indonesia," *J. Nas. Tek. Elektro Dan Teknol. Inf.*, vol. 10, no. 3, pp. 229–235, Aug. 2021, doi: 10.22146/jnteti.v10i3.1819.
- [6] J. W. Iskandar and Y. Nataliani, "Perbandingan Naïve Bayes, SVM, dan k-NN untuk Analisis Sentimen Gadget Berbasis Aspek," *J. RESTI Rekayasa Sist. Dan Teknol. Inf.*, vol. 5, no. 6, pp. 1120–1126, Dec. 2021, doi: 10.29207/resti.v5i6.3588.
- [7] R. M. Turjaman and I. Budi, "ANALISIS SENTIMEN BERBASIS ASPEK MARKETING MIX TERHADAP ULASAN APLIKASI DOMPET DIGITAL (STUDI KASUS: APLIKASI LINKAJA PADA

- TWITTER),” *J. Darma Agung*, vol. 30, no. 2, p. 266, Aug. 2022, doi: 10.46930/ojsuda.v30i2.1672.
- [8] D. Rifaldi, Abdul Fadlil, and Herman, “Teknik Preprocessing Pada Text Mining Menggunakan Data Tweet ‘Mental Health,’” *Decode J. Pendidik. Teknol. Inf.*, vol. 3, no. 2, pp. 161–171, Apr. 2023, doi: 10.51454/decode.v3i2.131.
- [9] I. Febriansyah, M. Fikry, and Yusra, “Analisis Sentiment di Twitter terhadap Anies Baswedan sebagai Bakal Calon Presiden 2024 Menggunakan Metode K-Nearest Neighbor,” *G-Tech J. Teknol. Terap.*, vol. 7, no. 3, pp. 1061–1070, Jul. 2023, doi: 10.33379/gtech.v7i4.2723.
- [10] Oryza Habibie Rahman, Gunawan Abdillah, and Agus Komarudin, “Klasifikasi Ujaran Kebencian pada Media Sosial Twitter Menggunakan Support Vector Machine,” *J. RESTI Rekayasa Sist. Dan Teknol. Inf.*, vol. 5, no. 1, pp. 17–23, Feb. 2021, doi: 10.29207/resti.v5i1.2700.
- [11] A. F. Azmi and A. Voutama, “PREDIKSI CHURN NASABAH BANK MENGGUNAKAN KLASIFIKASI RANDOM FOREST DAN DECISION TREE DENGAN EVALUASI CONFUSION MATRIX,” vol. 13, no. 1, 2024.