

OPTIMASI SVM DAN *RANDOM FOREST* DENGAN PSO PADA *ASPECT-BASED SENTIMENT ANALYSIS* ULASAN APLIKASI QPON

Dea Ananda Refiza Rahma ¹⁾, Abdul Rezha Efrat Najaf ²⁾, Reisa Permatasari ³⁾
^{1,2,3)} Program Studi Sistem Informasi, Fakultas Ilmu Komputer,
Universitas Pembangunan Nasional “Veteran” Jawa Timur
Email: deanandarefiza@gmail.com ¹⁾

Dikirim: 13 Juni 2026; Disetujui: 20 Juni 2026; Dipublikasikan: 31 Juli 2026

ABSTRAK

Ulasan pengguna aplikasi QPon memuat berbagai opini terhadap aspek layanan, seperti fungsionalitas, transaksi dan pembayaran, serta customer service. Analisis sentimen umum belum mampu menangkap opini secara spesifik pada setiap aspek, sehingga penelitian ini bertujuan mengimplementasikan algoritma machine learning yang dioptimasi menggunakan *Particle Swarm Optimization* (PSO) untuk *Aspect-Based Sentiment Analysis* (ABSA) ulasan aplikasi QPon. Data diperoleh melalui *scraping* ulasan *Google Play Store*, kemudian melalui tahap *filtering*, pelabelan manual, *preprocessing*, dan pembobotan fitur menggunakan TF-IDF dengan maksimum 5.000 fitur. Model yang digunakan adalah SVM dan *Random Forest* dengan enam skenario, yaitu *baseline*, PSO, dan SMOTE-PSO. Evaluasi dilakukan menggunakan *accuracy*, *precision*, *recall*, dan F1-score. Hasil penelitian menunjukkan bahwa pengaruh PSO bersifat kondisional. Pada aspek fungsionalitas, RF + SMOTE + PSO memperoleh F1-score terbaik sebesar 72,46%, meningkat 11,66 poin persentase dibanding RF *baseline*. Pada aspek transaksi dan pembayaran, SVM + PSO memperoleh F1-score 62,98%, meningkat 0,36 poin persentase dibanding SVM *baseline*. Sementara itu, pada aspek *customer service*, SVM + PSO menurunkan F1-score sebesar 4,95 poin persentase dibanding SVM *baseline*. Kontribusi ilmiah penelitian ini adalah memberikan bukti empiris bahwa efektivitas PSO dalam ABSA tidak bersifat universal, melainkan bergantung pada karakteristik data setiap aspek. Dengan demikian, pemilihan model terbaik perlu dilakukan secara spesifik pada setiap aspek.

Kata Kunci : *Aspect-Based Sentiment Analysis*, SVM, *Random Forest*, PSO, QPon.

ABSTRACT

User reviews of the QPon application contain various opinions on service aspects, such as functionality, transactions and payments, and customer service. General sentiment analysis is unable to capture user opinions specifically for each aspect; therefore, this study aims to implement machine learning algorithms optimized using *Particle Swarm Optimization* (PSO) for *Aspect-Based Sentiment Analysis* (ABSA) of QPon application reviews. The data were obtained through *scraping* reviews from the *Google Play Store*, followed by *filtering*, manual labeling, *preprocessing*, and feature weighting using TF-IDF with a maximum of 5,000 features. The models used were SVM and *Random Forest* with six scenarios, namely *baseline*, PSO, and SMOTE-PSO. Model evaluation was conducted using *accuracy*, *precision*, *recall*, and F1-score. The results show that the effect of PSO is conditional. In the functionality aspect, RF + SMOTE + PSO achieved the best F1-score of 72.46%, increasing by 11.66 percentage points compared to RF *baseline*. In the transactions and payments aspect, SVM + PSO achieved an F1-score of 62.98%, increasing by 0.36 percentage points compared to SVM *baseline*. Meanwhile, in the customer service aspect, SVM + PSO decreased the F1-score by 4.95 percentage points compared to SVM *baseline*. The scientific contribution of this study is to provide empirical evidence that the effectiveness of PSO in ABSA is not universal, but depends on the data characteristics of each aspect. Thus, the best model selection should be conducted specifically for each aspect.

Keywords : *Aspect-Based Sentiment Analysis*, SVM, *Random Forest*, PSO, QPon.

1. PENDAHULUAN

Perkembangan aplikasi *mobile* dalam ekosistem layanan digital mendorong pengguna untuk semakin aktif menyampaikan pengalaman, penilaian, dan keluhan melalui ulasan pada platform distribusi aplikasi. Ulasan pengguna merupakan salah satu bentuk umpan balik digital yang dapat merepresentasikan persepsi pengguna terhadap kualitas, kinerja, serta kepuasan terhadap suatu aplikasi. Informasi yang terkandung dalam ulasan dapat dimanfaatkan oleh pengembang untuk mengidentifikasi kelebihan layanan, menemukan permasalahan teknis, serta menentukan prioritas perbaikan aplikasi. Namun, jumlah ulasan yang besar, bentuk data yang tidak terstruktur, serta variasi bahasa pengguna membuat proses analisis secara manual menjadi tidak efisien dan berpotensi menghasilkan penilaian yang subjektif [1].

Salah satu pendekatan yang dapat digunakan untuk mengolah opini pengguna secara otomatis adalah analisis sentimen. Analisis sentimen merupakan bagian dari *Natural Language Processing* (NLP) yang berfokus pada proses identifikasi opini, emosi, sikap, atau penilaian seseorang terhadap suatu entitas melalui data teks. Melalui pendekatan ini, ulasan pengguna dapat diklasifikasikan ke dalam kategori sentimen positif, negatif, atau netral [2]. Meskipun demikian, analisis sentimen umum masih memiliki keterbatasan karena hanya memberikan gambaran polaritas secara keseluruhan dan belum mampu menunjukkan aspek layanan tertentu yang menjadi sumber kepuasan atau ketidakpuasan pengguna.

Aspect-Based Sentiment Analysis (ABSA) hadir sebagai pendekatan yang mampu mengatasi keterbatasan analisis sentimen umum dengan menguraikan ulasan ke dalam aspek-aspek tertentu, seperti pembayaran, antarmuka aplikasi, atau layanan pelanggan, kemudian menghubungkannya dengan polaritas sentimen. Berbeda dengan analisis sentimen konvensional yang hanya menilai kecenderungan opini secara keseluruhan, ABSA dapat memberikan informasi yang lebih rinci mengenai aspek yang memperoleh respons positif maupun negatif. Misalnya, pengguna dapat memberikan apresiasi terhadap promo makanan, tetapi pada saat yang sama menyampaikan ketidakpuasan terhadap proses redeem voucher. Informasi

semacam ini penting bagi pengambilan keputusan bisnis karena mampu menunjukkan bagian layanan yang perlu dipertahankan maupun diperbaiki. Pendekatan ABSA juga telah banyak diterapkan secara efektif pada berbagai bidang industri [3].

Dalam konteks penelitian ini, QPon dipilih sebagai objek penelitian karena merupakan aplikasi voucher digital yang memuat ulasan pengguna dengan beragam aspek layanan, seperti fungsionalitas, transaksi dan pembayaran, serta *customer service*. Berdasarkan data ulasan yang digunakan dalam penelitian ini, pengguna tidak hanya menyampaikan apresiasi terhadap variasi promo dan kemudahan penukaran *voucher*, tetapi juga menyampaikan keluhan terkait proses checkout, pembayaran, dan respons layanan pelanggan. Kondisi tersebut menunjukkan bahwa analisis sentimen umum belum cukup untuk menggambarkan persepsi pengguna secara menyeluruh, sehingga diperlukan pendekatan ABSA agar opini pengguna dapat dianalisis secara lebih spesifik berdasarkan aspek layanan.

Berdasarkan penelitian Kurniawan et al. [4] teknik analisis sentimen ini telah banyak diterapkan pada ulasan layanan *fintech* dan aplikasi *Point of Sale* (POS) dengan akurasi mencapai 80% menggunakan algoritma dasar *Support Vector Machine* (SVM) dan *Random Forest*. Namun, penelitian tersebut masih berfokus pada analisis sentimen secara umum, sehingga hasil klasifikasi belum menunjukkan sentimen pengguna secara spesifik pada setiap aspek layanan. Keterbatasan ini menjadi penting karena satu ulasan pengguna dapat memuat lebih dari satu opini, misalnya pengguna memberikan penilaian positif terhadap fitur aplikasi, tetapi menyampaikan keluhan terhadap transaksi atau layanan pelanggan. Oleh karena itu, pendekatan analisis sentimen konvensional belum sepenuhnya mampu menangkap kompleksitas opini pengguna pada ulasan aplikasi yang bersifat multi-aspek.

Untuk meningkatkan kinerja model dalam proses klasifikasi, *Particle Swarm Optimization* (PSO) dapat digunakan sebagai metode optimasi *swarm intelligence* yang terinspirasi dari perilaku kawanan burung dalam mencari solusi terbaik secara efisien pada ruang pencarian global [5]. Menurut Anggita et al. [6] dan Angkoso et al. [7] pada penelitiannya

membuktikan bahwa PSO telah berhasil mengoptimalkan hyperparameter sehingga meningkatkan akurasi hingga 11-18% dibanding dengan konfigurasi biasa. Meskipun demikian, penelitian tersebut umumnya masih menilai peningkatan performa model berdasarkan hasil klasifikasi secara keseluruhan, bukan berdasarkan karakteristik data pada setiap aspek. Dengan demikian, belum dijelaskan secara rinci apakah PSO selalu efektif pada seluruh aspek data atau hanya memberikan peningkatan pada kondisi tertentu, misalnya ketika jumlah data, distribusi kelas, dan pola ulasan memiliki karakteristik tertentu.

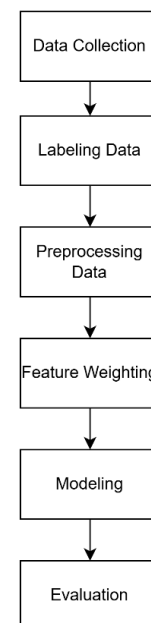
Berdasarkan keterbatasan tersebut, research gap dalam penelitian ini terletak pada belum banyaknya penelitian yang mengevaluasi efektivitas PSO secara spesifik dalam konteks *Aspect-Based Sentiment Analysis* (ABSA), terutama pada ulasan aplikasi voucher digital seperti QPon. Kebaruan penelitian ini tidak hanya terletak pada penggunaan PSO sebagai metode optimasi *hyperparameter*, tetapi pada strategi evaluasi berbasis aspek dengan membandingkan enam skenario model, yaitu SVM *baseline*, SVM + PSO, SVM + SMOTE + PSO, *Random Forest baseline*, *Random Forest* + PSO, dan *Random Forest* + SMOTE + PSO. Evaluasi dilakukan pada tiga aspek layanan QPon, yaitu fungsionalitas, transaksi dan pembayaran, serta *customer service*. Dengan pendekatan tersebut, penelitian ini tidak hanya mengidentifikasi model dengan performa terbaik, tetapi juga menganalisis kondisi ketika PSO mampu meningkatkan performa model dan ketika PSO tidak memberikan peningkatan yang signifikan.

Penelitian ini bertujuan untuk mengimplementasikan algoritma *machine learning* yang dioptimasi menggunakan PSO pada analisis sentimen berbasis aspek terhadap ulasan aplikasi QPon. Evaluasi model dilakukan menggunakan metrik *accuracy*, *precision*, *recall*, dan F1-score untuk mengetahui performa setiap skenario pada masing-masing aspek. Selain itu, penelitian ini juga membandingkan model *baseline*, model dengan optimasi PSO, serta model dengan kombinasi SMOTE dan PSO untuk melihat pengaruh optimasi hyperparameter dan penanganan ketidakseimbangan data terhadap kinerja

klasifikasi. Kontribusi ilmiah penelitian ini adalah memberikan bukti empiris bahwa efektivitas PSO dalam ABSA bersifat kondisional dan bergantung pada karakteristik data setiap aspek. Secara praktis, hasil penelitian ini diharapkan dapat membantu pengembang aplikasi QPon dalam memahami aspek layanan yang memperoleh respons positif maupun negatif, sehingga dapat menjadi dasar dalam menentukan prioritas pengembangan layanan.

2. METODE

Penelitian ini bertujuan untuk mengimplementasikan algoritma *machine learning* yang dioptimasi menggunakan PSO pada ABSA terhadap ulasan aplikasi QPon. Data yang digunakan dalam penelitian ini bersumber dari ulasan pengguna aplikasi QPon yang dipublikasikan pada platform *Google Play Store*. Proses penelitian terdiri dari enam tahapan, yaitu *data collection*, *labeling data*, *preprocessing data*, *feature weighting* menggunakan TF-IDF, *modeling* menggunakan SVM dan *Random Forest*, dan *evaluation*, pengujian skenario yang meliputi model *baseline*, PSO, dan SMOTE-PSO, serta *evaluation* menggunakan metrik *accuracy*, *precision*, *recall*, dan F1-score, sebagaimana ditunjukkan pada Gambar 1.



Gambar 1. Tahapan Penelitian

2.1 Data Collection

Penelitian ini menggunakan data ulasan pengguna aplikasi QPon yang diperoleh dari *Google Play Store*. Proses pengambilan data

dilakukan melalui teknik *web scraping* untuk mengekstraksi data dari sumber berbasis web dan menyimpannya dalam format yang terstruktur. Data yang telah dikumpulkan kemudian diseleksi dan dievaluasi berdasarkan kebutuhan penelitian untuk memastikan bahwa hanya ulasan yang relevan yang digunakan dalam proses analisis. Selanjutnya, data ulasan yang telah diproses disimpan dalam format CSV guna memudahkan tahapan analisis berikutnya [4].

2.2 Labeling Data

Proses pelabelan pada penelitian ini dilakukan secara manual oleh dua anotator untuk memastikan reliabilitas label aspek dan sentimen. Setiap ulasan pengguna diberi label berdasarkan dua kategori, yaitu aspek ulasan dan polaritas sentimen. Label aspek terdiri dari fungsionalitas, transaksi dan pembayaran, serta *customer service*, sedangkan label sentimen terdiri dari positif, negatif, dan netral. Sebelum proses pelabelan dilakukan, kedua anotator menggunakan pedoman pelabelan yang sama untuk meminimalkan perbedaan interpretasi. Selanjutnya, setiap ulasan diberi label secara independen oleh masing-masing anotator.

Setelah proses pelabelan selesai, tingkat kesepakatan antara kedua anotator diukur menggunakan koefisien *Cohen's Kappa*. Pengukuran ini digunakan untuk mengetahui konsistensi hasil pelabelan manual dengan memperhitungkan kemungkinan kesamaan penilaian antar-anotator yang terjadi secara acak. Ulasan yang memiliki perbedaan label kemudian dievaluasi kembali melalui diskusi untuk memperoleh label akhir [8]. Dengan demikian, dataset berlabel akhir yang digunakan dalam penelitian ini diperoleh melalui kombinasi anotasi independen, pengukuran kesepakatan, dan validasi label.

2.3 Preprocessing Data

Preprocessing data dilakukan untuk mempersiapkan data teks sebelum masuk ke tahap pembobotan fitur dan pemodelan. Data ulasan mentah yang diperoleh dari *Google Play Store* umumnya masih tidak terstruktur dan mengandung *noise*, seperti tanda baca, karakter khusus, angka, kata tidak baku, kata berimbuhan, serta elemen teks yang tidak relevan. Dengan demikian, tahap *preprocessing* dilakukan untuk mengubah data teks mentah menjadi data yang lebih bersih, terstandarisasi,

dan siap digunakan dalam proses analisis lanjutan. [9].

Pada penelitian ini, tahapan *preprocessing* meliputi *case folding*, *cleaning*, *normalization*, *tokenizing*, *stopword removal*, dan *stemming*. *Case folding* dilakukan sebagai tahap awal *preprocessing* untuk menyeragamkan penulisan teks dengan mengubah seluruh huruf menjadi bentuk *lowercase* [10]. *Cleaning* digunakan untuk menghapus karakter yang tidak relevan, simbol, angka, serta elemen teks yang tidak diperlukan [11]. *Normalization* dilakukan untuk mengubah kata slang atau kata tidak baku menjadi bentuk baku dalam bahasa Indonesia [12]. *Tokenizing* berfungsi untuk memecah teks menjadi bagian-bagian kecil untuk merepresentasikan suatu makna sehingga dapat dianalisis secara individual pada tahap pemodelan [10]. Tahap *stopword removal* bertujuan menghapus kata-kata yang sering muncul namun tidak memiliki pengaruh besar terhadap analisis sentimen, sedangkan *stemming* digunakan untuk menyederhanakan kata ke bentuk dasar. [10].

2.4 Feature Weighting

Pada tahap pembobotan fitur, penelitian ini menerapkan metode *Term Frequency-Inverse Document Frequency* (TF-IDF) untuk merepresentasikan data teks hasil *preprocessing* ke dalam bentuk numerik. TF-IDF memberikan bobot pada setiap kata berdasarkan tingkat kemunculannya dalam dokumen serta relevansinya terhadap keseluruhan korpus [13]. Semakin sering suatu kata muncul dalam dokumen tertentu dan semakin jarang kata tersebut ditemukan pada dokumen lain, maka semakin besar bobot yang diberikan karena dianggap memiliki nilai pembeda yang lebih kuat. Ekstraksi fitur TF-IDF pada penelitian ini dibatasi hingga maksimum 5.000 fitur menggunakan parameter *max_features=5000*. Pemilihan batas tersebut mengacu pada penelitian analisis sentimen terdahulu yang menunjukkan bahwa konfigurasi 5.000 fitur mampu memberikan performa klasifikasi yang stabil dengan menyaring fitur kurang relevan tanpa mengurangi informasi penting pada data teks [14]. Hasil pembobotan fitur menggunakan TF-IDF selanjutnya digunakan sebagai input pada model SVM dan *Random Forest* dalam tahap klasifikasi. Berikut untuk persamaan TF-IDF.

$$TF = \frac{tf}{\max(f)} \quad (1)$$

$$IDF = \log \frac{D}{df} \quad (2)$$

$$w = TF \times IDF \quad (3)$$

2.5 Modeling

Tahap pemodelan dilakukan dengan menerapkan algoritma SVM dan *Random Forest* menggunakan fitur hasil pembobotan TF-IDF sebagai input. Pada penelitian ini, terdapat enam skenario pengujian yang dirancang untuk membandingkan performa model ini ditunjukkan pada Tabel 1.

Tabel 1. Skenario Model

Skenario	Model
1	SVM baseline
2	SVM + PSO
3	SVM+ SMOTE + PSO
4	RF baseline
5	RF + PSO
6	RF+ SMOTE + PSO

Data dibagi menggunakan skema 80:20, dengan 80% data sebagai data latih dan 20% data sebagai data uji. Pembagian data dilakukan secara terpisah pada aspek fungsionalitas, transaksi dan pembayaran, serta customer service agar proses pemodelan dapat mengevaluasi karakteristik data pada setiap aspek secara spesifik.

Tahap pemodelan dilakukan dengan menerapkan SVM dan *Random Forest* menggunakan fitur hasil pembobotan TF-IDF sebagai input. Penelitian ini menggunakan enam skenario pengujian, yaitu SVM *baseline*, SVM + PSO, SVM + SMOTE + PSO, *Random Forest baseline*, *Random Forest* + PSO, dan *Random Forest* + SMOTE + PSO. Skenario *baseline* digunakan untuk mengetahui performa awal model tanpa optimasi, skenario PSO digunakan untuk mengetahui pengaruh optimasi *hyperparameter*, sedangkan skenario SMOTE + PSO digunakan untuk mengetahui pengaruh kombinasi penanganan ketidakseimbangan data dan optimasi *hyperparameter* terhadap performa klasifikasi.

Pada skenario yang melibatkan PSO, proses optimasi dilakukan pada data latih menggunakan *Stratified K-Fold Cross*

Validation dengan 5 *fold*, *shuffle=True*, dan *random_state=42*. Metrik yang digunakan sebagai fungsi objektif adalah F1-score *macro* karena metrik ini mampu mengevaluasi performa model secara lebih seimbang pada data dengan distribusi kelas yang tidak merata. Karena algoritma PSO bekerja dengan mencari nilai minimum, nilai F1-score *macro* diubah menjadi nilai negatif pada fungsi objektif. Setelah parameter terbaik diperoleh, model final dilatih menggunakan data latih dan dievaluasi pada data uji.

Pada skenario yang melibatkan SMOTE, proses *oversampling* hanya diterapkan pada data latih untuk menghindari kebocoran data terhadap data uji. Nilai *k_neighbors* pada SMOTE ditentukan secara otomatis menggunakan rumus $\min(5, \text{jumlah data kelas minoritas} - 1)$. Penyesuaian ini dilakukan agar proses SMOTE tetap dapat diterapkan pada aspek dengan jumlah data kelas minoritas yang relatif kecil. Data uji tidak melalui proses SMOTE, sehingga hasil evaluasi tetap merepresentasikan distribusi data asli. Detail konfigurasi PSO ditampilkan pada Tabel 2.

Tabel 2. Konfigurasi PSO

Komponen	Konfigurasi
Jumlah partikel PSO	10
Jumlah iterasi PSO	5
Bobot inersia PSO	0,5
Learning factor individu	0,5
Learning factor sosial	0,5

Detail ruang pencarian *hyperparameter* ditampilkan pada Tabel 3.

Tabel 3. Ruang Pencarian *Hyperparameter*

Model	Parameter	Ruang Pencarian	Keterangan
SVM + PSO	<i>C</i> , <i>gamma</i>	<i>C</i> = 0,1–100; <i>gamma</i> = 0,0001–1	Kernel RBF
SVM + SMOT E + PSO	<i>C</i> , <i>gamma</i>	<i>C</i> = 0,1–100; <i>gamma</i> = 0,0001–1	SMOTE pada data latih
Rando m Forest + PSO	<i>n_estimato rs</i> , <i>max_depth</i>	<i>n_estimato rs</i> = 50–300; <i>max_depth</i> = 5–50	random_state=42, n_jobs=-1
Rando m Forest	<i>n_estimato rs</i> , <i>max_depth</i>	<i>n_estimato rs</i> = 50–300;	SMOTE pada data latih

+	max_depth
SMOT	$= 5-50$
E +	
PSO	

2.6 Evaluation

Evaluasi model dalam penelitian ini menggunakan metode *5-fold cross-validation*. Dataset dibagi menjadi lima subset, kemudian pada setiap iterasi empat subset digunakan sebagai data latih dan satu subset digunakan sebagai data uji. Proses ini dilakukan secara bergantian hingga seluruh subset pernah berperan sebagai data uji, sehingga hasil evaluasi model menjadi lebih objektif dan stabil. Nilai performa akhir diperoleh dari rata-rata hasil evaluasi pada seluruh *fold*.

Performa setiap skenario pengujian dievaluasi menggunakan metrik *accuracy*, *precision*, *recall*, dan F1-score. Metrik tersebut digunakan untuk membandingkan kinerja model *baseline*, model yang dioptimasi menggunakan PSO, serta model dengan kombinasi SMOTE dan PSO pada masing-masing aspek ulasan. Berikut untuk persamaan tiap metrik.

$$Precision = \frac{TP}{(TP + FP)} \quad (4)$$

$$Recall = \frac{TP}{(TP + FN)} \quad (5)$$

$$F-1 \text{ Score} = \frac{(2 \times Recall \times Precision)}{(Recall + Precision)} \quad (6)$$

3. HASIL DAN PEMBAHASAN

Pada bagian ini memaparkan hasil dari langkah-langkah pada metode penelitian yang mana diproses menggunakan bahasa pemrograman *Python* dan *Google Sheets*.

3.1. Data Collection

Proses pengumpulan data dilakukan dengan melakukan scraping terhadap ulasan pengguna berbahasa Indonesia pada aplikasi QPon di Google Play Store. Hasil proses scraping memperoleh sebanyak 4.101 data ulasan mentah yang kemudian disimpan dalam format CSV untuk diproses lebih lanjut. Sebelum memasuki tahap pelabelan, dilakukan proses filtering data untuk meningkatkan kualitas dataset dengan menghapus ulasan duplikat, konten yang tidak

relevan, serta ulasan yang tidak informatif. Setelah proses filtering dilakukan, jumlah data berkurang menjadi 3.460 ulasan yang selanjutnya digunakan sebagai data akhir untuk proses pelabelan manual.

3.2. Labeling Data

Hasil pelabelan kemudian dievaluasi menggunakan *Cohen's Kappa* untuk mengukur tingkat kesepakatan antar-anotator. Nilai *Cohen's Kappa* yang diperoleh sebesar 0,9793 untuk pelabelan sentimen dan 0,9727 untuk pelabelan aspek, yang menunjukkan tingkat kesepakatan *almost perfect*. Distribusi label akhir yang digunakan dalam penelitian ini dirangkum pada Tabel 4.

Tabel 4. Distribusi Label

Aspek	Negatif	Positif	Netral
Fungsional	1.407	666	74
Transaksi & Pembayaran	508	441	16
Customer Service	238	110	0

3.3. Preprocessing Data

Tahap preprocessing dilakukan terhadap 3.460 ulasan yang telah diberi label untuk menghasilkan data teks yang lebih bersih dan terstandarisasi sebelum proses pembobotan fitur. Tahapan ini meliputi *case folding*, *cleaning*, *normalization*, *tokenizing*, *stopword removal* dan *stemming*. Setelah proses *stemming*, terdapat tujuh ulasan yang dihapus karena teks akhir hanya menyisakan satu kata yang tidak informatif, seperti “tidak”, “terima”, dan “pahala”. Kata-kata tersebut tidak lagi merepresentasikan konteks ulasan yang jelas untuk proses klasifikasi sentimen. Dengan demikian, dataset akhir yang digunakan pada tahap pembobotan fitur dan pemodelan berjumlah 3.453 ulasan. Secara keseluruhan, preprocessing membantu mengurangi noise dan menghasilkan fitur teks yang lebih konsisten untuk proses pembobotan menggunakan TF-IDF. Berikut untuk hasil dari *preprocessing* dapat dilihat pada Tabel 3.

Tabel 5. Hasil Preprocessing

Sebelum	Sesudah
Masuk pake otp gagal terus. Plus ada limited otp nya lagi. Aplikasi butut. 🤖🤖	masuk otp gagal plus limited otp aplikasi butut

3.4. Feature Weighting

Setelah tahap *preprocessing*, sebanyak 3.453 ulasan akhir ditransformasikan ke dalam bentuk fitur numerik menggunakan metode

Term Frequency-Inverse Document Frequency (TF-IDF). Pada penelitian ini, proses ekstraksi fitur TF-IDF dibatasi hingga maksimum 5.000 fitur untuk mempertahankan term yang paling relevan sekaligus mengurangi kompleksitas fitur. Proses ini mengubah data ulasan berbentuk teks menjadi representasi numerik berbobot sehingga dapat diproses oleh algoritma *machine learning*. Hasil fitur TF-IDF selanjutnya digunakan sebagai input pada model SVM dan *Random Forest* dalam tahap klasifikasi.

3.5. Modeling

Tahap *modeling* dilakukan setelah seluruh data ulasan direpresentasikan ke dalam bentuk fitur numerik menggunakan TF-IDF. Pada tahap ini, pemodelan dilakukan secara terpisah pada masing-masing aspek karena setiap aspek memiliki jumlah data, distribusi kelas, dan pola ulasan yang berbeda. Penelitian ini menerapkan enam skenario pengujian, yaitu SVM *baseline*, SVM + PSO, SVM + SMOTE + PSO, *Random Forest baseline*, *Random Forest* + PSO, dan *Random Forest* + SMOTE + PSO. Namun, tabel hasil optimasi parameter hanya menampilkan skenario yang melibatkan PSO, yaitu skenario 2, 3, 5, dan 6.

Tabel 6 menunjukkan bahwa setiap skenario menghasilkan kombinasi *hyperparameter* terbaik yang berbeda pada masing-masing aspek. Pada skenario SVM + PSO, parameter terbaik untuk aspek fungsionalitas adalah $C = 89,2730$ dan $\gamma = 0,0130$ dengan F1-score sebesar 65,64%. Pada aspek transaksi dan pembayaran, parameter terbaik yang diperoleh adalah $C = 3,1191$ dan $\gamma = 0,6577$ dengan F1-score sebesar 62,98%. Sementara itu, pada aspek *customer service*, parameter terbaik adalah $C = 87,4788$ dan $\gamma = 0,1125$ dengan F1-score sebesar 88,25%. Perbedaan nilai C dan γ menunjukkan bahwa setiap aspek membutuhkan tingkat kompleksitas model yang berbeda dalam membentuk batas keputusan klasifikasi.

Tabel 6. Hasil Optimasi Parameter

Skenario	Aspek	Parameter Terbaik	F1-score (%)
2	1	$C = 89,2730$ $\gamma = 0,0130$	65,64
	2	$C = 3,1191$ $\gamma = 0,6577$	62,98

Skenario	Aspek	Parameter Terbaik	F1-score (%)
3	3	$C = 87,4788$ $\gamma = 0,1125$	88,25
	1	$C = 66,3979$ $\gamma = 0,0119$	67,55
	2	$C = 46,4138$ $i = 0,0028$	62,61
	3	$C = 0,5614$ $\gamma = 0,3266$	93,20
5	1	$n_estimators = 183$ $max_depth = 49$	60,84
	2	$n_estimators = 210$ $max_depth = 24$	61,58
	3	$n_estimators = 147$ $max_depth = 50$	91,11
6	1	$n_estimators = 300$ $max_depth = 22$	72,46
	2	$n_estimators = 216$ $max_depth = 21$	61,93
	3	$n_estimators = 242$ $max_depth = 43$	91,38

Keterangan:

Skenario 2 = SVM + PSO

Skenario 3 = SVM + SMOTE + PSO

Skenario 5 = *Random Forest* + PSO

Skenario 6 = *Random Forest* + SMOTE + PSO

Aspek 1 = Fungsionalitas

Aspek 2 = Transaksi & Pembayaran

Aspek 3 = *Customer Service*

Pada skenario *Random Forest* + PSO, parameter terbaik untuk aspek fungsionalitas adalah $i = 183$ dan $max_depth = 49$ dengan F1-score sebesar 60,84%. Pada aspek transaksi dan pembayaran, parameter terbaik adalah $n_estimators = 210$ dan $max_depth = 24$ dengan F1-score sebesar 61,58%. Sementara itu, pada aspek *customer service*, parameter terbaik adalah $n_estimators = 147$ dan $max_depth = 50$ dengan F1-score sebesar 91,11%. Hasil tersebut menunjukkan bahwa optimasi PSO pada *Random Forest* belum selalu menghasilkan performa yang lebih tinggi dibandingkan kombinasi dengan SMOTE, terutama pada aspek yang memiliki distribusi kelas tidak seimbang.

Pada skenario *Random Forest* + SMOTE + PSO, aspek fungsionalitas memperoleh parameter terbaik $n_estimators = 300$ dan $max_depth = 22$ dengan F1-score sebesar 91,38%. Aspek transaksi dan pembayaran

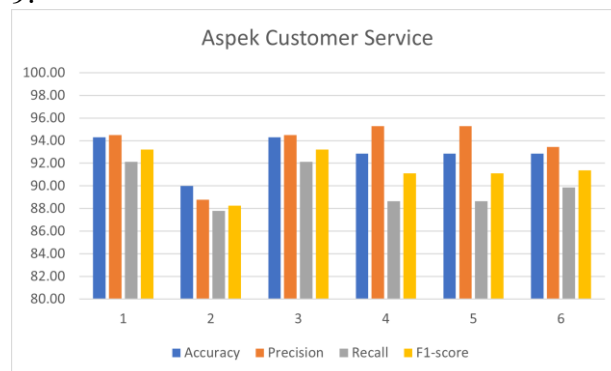
memperoleh $n_estimators = 216$ dan $max_depth = 21$ dengan F1-score sebesar 61,93%. Sementara itu, aspek *customer service* memperoleh $n_estimators = 242$ dan $max_depth = 43$ dengan F1-score sebesar 72,46%. Berdasarkan hasil tersebut, skenario *Random Forest + SMOTE + PSO* memberikan performa yang sangat baik pada aspek fungsionalitas, tetapi tidak memberikan hasil terbaik pada aspek *customer service*. Hal ini menunjukkan bahwa penanganan ketidakseimbangan data dan optimasi *hyperparameter* tidak selalu memberikan dampak yang sama pada seluruh aspek.

Secara umum, hasil *modeling* menunjukkan bahwa PSO tidak menghasilkan satu kombinasi parameter yang berlaku secara seragam untuk seluruh aspek. Setiap aspek memiliki karakteristik data yang berbeda, sehingga parameter terbaik dan nilai F1-score yang dihasilkan juga berbeda. Pada aspek fungsionalitas, performa terbaik diperoleh oleh *Random Forest + SMOTE + PSO* dengan F1-score sebesar 91,38%. Pada aspek transaksi dan pembayaran, performa terbaik diperoleh oleh SVM + PSO dengan F1-score sebesar 62,98%. Sementara itu, pada aspek *customer service*, performa terbaik diperoleh oleh SVM + SMOTE + PSO dengan F1-score sebesar 93,22%. Temuan ini memperkuat bahwa efektivitas PSO dan SMOTE bersifat kondisional, bergantung pada jumlah data, distribusi kelas, serta kompleksitas pola ulasan pada masing-masing aspek.

3.6. Evaluation

Evaluasi model dilakukan menggunakan metrik *accuracy*, *precision*, *recall*, dan F1-score untuk menilai kinerja setiap skenario klasifikasi pada masing-masing aspek ulasan. Penggunaan beberapa metrik diperlukan karena dataset memiliki distribusi kelas yang tidak seimbang, sehingga *accuracy* saja belum cukup untuk menggambarkan kemampuan model secara menyeluruh. Nilai F1-score menjadi metrik penting karena mampu menunjukkan keseimbangan antara *precision* dan *recall*, terutama dalam mengukur kemampuan model terhadap kelas minoritas. Perbandingan hasil evaluasi pada aspek fungsionalitas, transaksi dan pembayaran, serta *customer service*

ditampilkan pada Gambar serta Tabel 7, 8, dan 9.



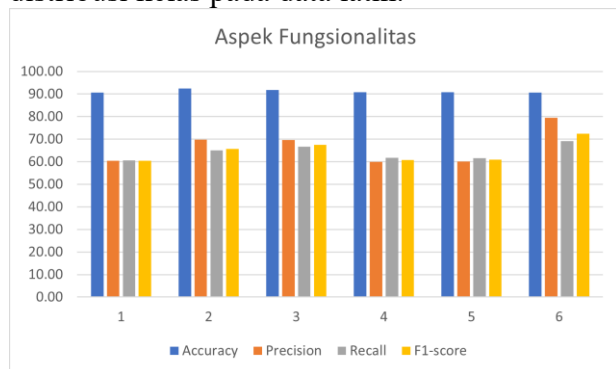
Gambar 2. Grafik Komparasi Evaluasi Aspek Customer Service

Tabel 7. Evaluasi Aspek Customer Service

Skenario	Accuracy (%)	Precision (%)	Recall (%)	F1-score (%)
1: SVM Baseline	94,29	94,50	92,14	93,20
2: SVM + PSO	90,00	88,78	87,78	88,25
3: SVM + SMOTE + PSO	94,29	94,50	92,14	93,20
4: RF Baseline	92,86	95,28	88,64	91,11
5: RF + PSO	92,86	95,28	88,65	91,11
6: RF + SMOTE + PSO	92,86	93,45	89,87	91,38

Berdasarkan Tabel 7, aspek *customer service* menunjukkan bahwa PSO tidak selalu meningkatkan performa model. Berdasarkan Tabel 7, SVM baseline dan SVM + SMOTE + PSO memperoleh F1-score tertinggi yang sama sebesar 93,20%, sedangkan SVM + PSO mengalami penurunan menjadi 88,25% atau turun 4,95 poin persentase. Penurunan ini dapat terjadi karena distribusi data pada aspek *customer service* masih tidak seimbang, yaitu didominasi oleh kelas negatif sebanyak 238 ulasan, sedangkan kelas positif berjumlah 110 ulasan dan kelas netral tidak ditemukan. Kondisi tersebut membuat model lebih banyak mempelajari pola kelas mayoritas dan tidak memperoleh representasi kelas netral, sehingga kemampuan generalisasi model terhadap variasi sentimen menjadi terbatas. Selain itu, performa SVM baseline yang sudah tinggi menunjukkan bahwa pola ulasan *customer service* relatif sudah dapat dikenali dengan baik tanpa optimasi tambahan. Oleh karena itu, penerapan PSO

tanpa penanganan *imbalance* data justru dapat menurunkan performa, sedangkan kombinasi SMOTE + PSO mampu menyamai *baseline* karena SMOTE membantu menyeimbangkan distribusi kelas pada data latih.



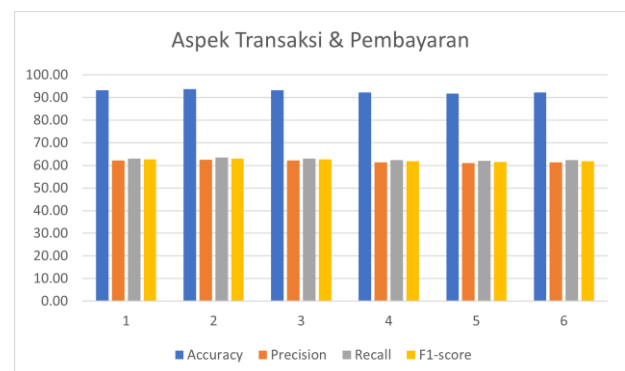
Gambar 3. Grafik Komparasi Evaluasi Aspek Fungsionalitas

Tabel 8. Evaluasi Aspek Fungsionalitas

Skenario	Accuracy (%)	Precision (%)	Recall (%)	F1-score (%)
1: SVM Baseline	90,65	60,49	60,66	60,44
2: SVM + PSO	92,52	69,79	65,04	65,64
3: SVM + SMOTE + PSO	91,82	69,64	66,65	67,55
4: RF Baseline	90,89	59,94	61,70	60,80
5: RF + PSO	90,89	60,16	61,57	60,84
6: RF + SMOTE + PSO	90,65	79,53	69,21	72,46

Pada aspek fungsionalitas, sebagaimana ditampilkan pada Tabel 8 yang menunjukkan peningkatan performa paling besar pada skenario *Random Forest* + SMOTE + PSO. Skenario ini memperoleh F1-score tertinggi sebesar 72,46%, meningkat 11,66 poin persentase dibandingkan *Random Forest baseline* yang memperoleh F1-score sebesar 60,80%. Peningkatan ini menunjukkan bahwa kombinasi SMOTE dan PSO mampu membantu *Random Forest* dalam mengenali pola sentimen secara lebih seimbang pada aspek fungsionalitas. Hal ini dapat terjadi karena aspek fungsionalitas memiliki jumlah data paling besar dibandingkan aspek lainnya, yaitu 1.407 ulasan negatif, 666 ulasan positif, dan 74 ulasan netral. Meskipun kelas negatif tetap mendominasi, jumlah data yang lebih besar memberikan lebih banyak variasi pola teks

untuk dipelajari oleh model. SMOTE membantu menyeimbangkan distribusi kelas pada data latih, terutama terhadap kelas minoritas, sedangkan PSO membantu mencari kombinasi parameter *n_estimators* dan *max_depth* yang lebih sesuai dengan karakteristik data. Meskipun *accuracy* pada skenario ini sedikit lebih rendah dibandingkan SVM + PSO, nilai *precision*, *recall*, dan F1-score yang lebih tinggi menunjukkan bahwa model lebih seimbang dalam mengenali kelas sentimen, bukan hanya dominan pada kelas mayoritas.



Gambar 4. Grafik Komparasi Evaluasi Aspek Transaksi & Pembayaran

Tabel 9. Evaluasi Aspek Transaksi & Pembayaran

Skenario	Accuracy (%)	Precision (%)	Recall (%)	F1-score (%)
1: SVM Baseline	93,26	62,17	63,09	62,62
2: SVM + PSO	93,78	62,50	63,47	62,98
3: SVM + SMOTE + PSO	93,26	62,24	63,03	62,61
4: RF Baseline	92,23	61,44	62,42	61,92
5: RF + PSO	91,71	61,11	62,05	61,58
6: RF + SMOTE + PSO	92,23	61,44	62,43	61,93

Sementara itu, Tabel 9 menunjukkan bahwa peningkatan performa akibat PSO relatif kecil pada aspek transaksi dan pembayaran. Skenario terbaik diperoleh oleh SVM + PSO dengan F1-score sebesar 62,98%, hanya meningkat 0,36 poin persentase dibandingkan SVM *baseline* yang memperoleh F1-score sebesar 62,62%. Hasil ini menunjukkan bahwa optimasi *hyperparameter* melalui PSO belum memberikan pengaruh yang besar pada aspek

transaksi dan pembayaran. Kondisi tersebut dapat disebabkan oleh distribusi kelas yang masih tidak seimbang, yaitu 508 ulasan negatif, 441 ulasan positif, dan hanya 16 ulasan netral. Meskipun jumlah kelas negatif dan positif relatif lebih berdekatan dibandingkan aspek lainnya, jumlah kelas netral yang sangat sedikit membuat model sulit mempelajari pola sentimen netral secara optimal. Selain itu, ulasan pada aspek ini memiliki pola yang lebih kompleks karena berkaitan dengan berbagai isu, seperti checkout, pembayaran QRIS, kegagalan transaksi, saldo, dan proses *redeem voucher*. Kompleksitas tersebut menyebabkan batas antar kelas sentimen menjadi lebih sulit dipisahkan, sehingga perubahan *hyperparameter* melalui PSO hanya memberikan peningkatan performa yang terbatas. Nilai *accuracy* yang tinggi pada seluruh skenario, yaitu di atas 91%, juga tidak selalu diikuti oleh F1-score yang tinggi. Hal ini menunjukkan bahwa model masih lebih mudah mengenali kelas dominan dibandingkan kelas minoritas, terutama kelas netral.

Secara keseluruhan, hasil evaluasi pada Tabel 7, Tabel 8, dan Tabel 9 menunjukkan bahwa tidak terdapat satu skenario model yang selalu unggul pada seluruh aspek. Pada aspek fungsionalitas, skenario terbaik adalah *Random Forest* + SMOTE + PSO dengan F1-score sebesar 72,46%. Pada aspek transaksi dan pembayaran, skenario terbaik adalah SVM + PSO dengan F1-score sebesar 62,98%. Sementara itu, pada aspek *customer service*, SVM baseline dan SVM + SMOTE + PSO memperoleh F1-score tertinggi yang sama sebesar 93,20%. Temuan ini menunjukkan bahwa efektivitas PSO dan SMOTE bersifat kondisional, bergantung pada jumlah data, distribusi kelas, serta kompleksitas pola ulasan pada masing-masing aspek.

Temuan penelitian ini sejalan dengan penelitian Anggita et al. [6] dan Angkoso et al. [7] yang menunjukkan bahwa PSO dapat meningkatkan performa model melalui proses optimasi *hyperparameter*. Namun, hasil penelitian ini juga memperluas temuan sebelumnya karena menunjukkan bahwa peningkatan performa akibat PSO tidak terjadi secara merata pada seluruh aspek. Pada aspek fungsionalitas, PSO yang dikombinasikan dengan SMOTE mampu memberikan peningkatan yang cukup signifikan. Sebaliknya,

pada aspek transaksi dan pembayaran, peningkatan yang dihasilkan sangat kecil, sedangkan pada aspek *customer service*, PSO tanpa SMOTE justru menurunkan performa model. Dengan demikian, kontribusi penelitian ini tidak hanya menunjukkan bahwa PSO dapat digunakan untuk optimasi model, tetapi juga memberikan bukti empiris bahwa efektivitas PSO dalam ABSA perlu dievaluasi berdasarkan karakteristik data pada setiap aspek.

Secara teoretis, PSO bekerja dengan mencari kombinasi *hyperparameter* terbaik dalam ruang pencarian tertentu berdasarkan fungsi objektif yang digunakan. Dalam penelitian ini, fungsi objektif yang digunakan adalah F1-score *macro*, sehingga proses optimasi diarahkan untuk memperoleh parameter yang mampu menyeimbangkan performa antar kelas sentimen. Namun, keberhasilan optimasi tidak hanya dipengaruhi oleh algoritma PSO, tetapi juga oleh karakteristik data, seperti ukuran data, distribusi kelas, dan kompleksitas pola teks. Apabila data memiliki jumlah yang cukup dan variasi pola yang dapat dipelajari, PSO dapat membantu meningkatkan performa model. Sebaliknya, pada data yang lebih kecil atau sudah memiliki performa baseline tinggi, optimasi tambahan tidak selalu memberikan peningkatan dan bahkan dapat menurunkan kemampuan generalisasi model.

Secara praktis, hasil penelitian ini dapat dimanfaatkan oleh pengembang aplikasi QPon sebagai dasar untuk menentukan prioritas peningkatan layanan. Aspek fungsionalitas perlu terus dipantau karena menjadi aspek dengan peningkatan performa model paling jelas setelah penerapan SMOTE dan PSO. Aspek transaksi dan pembayaran perlu menjadi perhatian khusus karena meskipun *accuracy* tinggi, nilai F1-score masih relatif rendah, yang menunjukkan bahwa model belum sepenuhnya optimal dalam mengenali seluruh kelas sentimen. Hal ini relevan karena aspek transaksi berkaitan langsung dengan pengalaman pengguna dalam proses checkout, pembayaran, dan penukaran voucher. Sementara itu, aspek *customer service* menunjukkan performa klasifikasi yang tinggi, tetapi tetap memerlukan pengayaan data agar model lebih stabil terhadap perubahan parameter. Dengan demikian, hasil ABSA dapat digunakan sebagai dasar

pengembangan *dashboard monitoring* opini pengguna, sistem rekomendasi prioritas perbaikan layanan, serta strategi peningkatan kualitas aplikasi QPon.

4. PENUTUP

4.1. Kesimpulan

Berdasarkan hasil penelitian, implementasi PSO pada analisis sentimen berbasis aspek ulasan aplikasi QPon menunjukkan bahwa pengaruh optimasi terhadap performa model bersifat kondisional dan bergantung pada karakteristik data setiap aspek. Pada aspek fungsionalitas, skenario *Random Forest* + SMOTE + PSO memperoleh F1-score terbaik sebesar 72,46%. Pada aspek transaksi dan pembayaran, skenario SVM + PSO memperoleh F1-score terbaik sebesar 62,98%, tetapi peningkatannya relatif kecil. Sementara itu, pada aspek *customer service*, SVM *baseline* dan SVM + SMOTE + PSO memperoleh F1-score tertinggi yang sama sebesar 93,20%, sedangkan SVM + PSO mengalami penurunan performa. Temuan ini menunjukkan bahwa PSO tidak selalu meningkatkan performa model, terutama pada data dengan karakteristik tertentu seperti distribusi kelas yang tidak seimbang atau performa *baseline* yang sudah tinggi.

Kontribusi ilmiah penelitian ini terletak pada penyajian bukti empiris bahwa efektivitas PSO dalam ABSA tidak dapat digeneralisasi pada seluruh aspek, sehingga evaluasi model perlu dilakukan secara spesifik berdasarkan karakteristik data masing-masing aspek. Secara praktis, hasil penelitian ini dapat membantu pengembang QPon dalam mengidentifikasi prioritas perbaikan layanan berdasarkan aspek yang memperoleh sentimen pengguna. Penelitian ini masih memiliki keterbatasan karena hanya menggunakan satu objek aplikasi, distribusi kelas pada beberapa aspek belum seimbang, dan representasi fitur masih berbasis TF-IDF sehingga belum sepenuhnya menangkap konteks semantik ulasan.

4.2. Saran

Penelitian selanjutnya disarankan untuk menggunakan dataset yang lebih seimbang pada setiap aspek dan kelas sentimen agar performa model lebih stabil. Selain itu, metode penanganan imbalance selain SMOTE dan metode optimasi selain PSO dapat digunakan sebagai pembandingan untuk memperoleh hasil

yang lebih komprehensif. Pengembangan model berbasis *deep learning* atau *transformer* juga dapat dipertimbangkan agar proses klasifikasi mampu menangkap konteks ulasan berbahasa Indonesia secara lebih baik.

5. DAFTAR PUSTAKA

- [1] I. Febyanti, A. S. Ddevi, S. S. M. Wara, and A. T. Damaliana, "Klasifikasi Sentimen Ulasan Pengguna Aplikasi Qpon dengan Support Vector Machine dan Logistic Regression," *Journal of Data Mining and Information Systems*, vol. 4, no. 1, Feb. 2026, Accessed: Jun. 18, 2026. [Online]. Available: <https://www.journal.yp3a.org/index.php/jdmis/article/view/4663>
- [2] A. Novantika and Sugiman, "Analisis Sentimen Ulasan Pengguna Aplikasi Video Conference Google Meet menggunakan Metode SVM dan Logistic Regression," in *PRISMA, Prosiding Seminar Nasional Matematika*, 2022, pp. 808–813. Accessed: Jun. 19, 2026. [Online]. Available: <https://journal.unnes.ac.id/sju/index.php/prisma/>
- [3] Ł. Augustyniak, T. Kajdanowicz, and P. Kazienko, "Comprehensive analysis of aspect term extraction methods using various text embeddings," *Comput. Speech Lang.*, vol. 69, p. 101217, Sep. 2021, doi: 10.1016/j.csl.2021.101217.
- [4] D. A. A. Kurniawan, E. Utami, and H. Al Fatta, "ANALISIS SENTIMEN PADA OPINI PENGGUNA APLIKASI QASIR MENGGUNAKAN SUPPORT VECTOR MACHINE DAN RANDOM FOREST," *TEKNIMEDIA: Teknologi Informasi dan Multimedia*, vol. 4, no. 1, pp. 1–8, Jun. 2023, doi: 10.46764/teknimedia.v4i1.83.
- [5] M. Saini, A. M. S. Yunus, and M. R. Djalal, *Swarm Intelligence: Craziness Particle Swarm Optimization untuk Optimasi Sistem Tenaga Listrik*. Deepublish, 2024.
- [6] S. D. Anggita and F. F. Abdulloh, "Optimasi Algoritma Support Vector Machine Berbasis PSO Dan Seleksi Fitur Information Gain Pada Analisis Sentimen," *Journal of Applied Computer*

- Science and Technology*, vol. 4, no. 1, pp. 52–57, Jul. 2023, doi: 10.52158/jacost.v4i1.524.
- [7] C. V. Angkoso, K. Asror, A. Kusumaningsih, and A. K. Nugroho, “Optimasi Algoritma Support Vector Machine Berbasis Kernel Radial Basis Function (RBF) Menggunakan Metode Particle Swarm Optimization Untuk Analisis Sentimen,” *Jurnal Teknologi Informasi dan Ilmu Komputer*, vol. 12, no. 3, pp. 705–718, Jun. 2025, doi: 10.25126/jtiik.2025129317.
- [8] I. Amal and Jayanata, “Perbandingan Pelabelan Otomatis Dan Manual Untuk Analisis Sentimen Terhadap Kenaikan Harga BBM Pertamina Pada Twitter Menggunakan Algoritma Support Vector Machine,” in *Seminar Nasional Mahasiswa Ilmu Komputer dan Aplikasinya (SENAMIKA)*, Jakarta, 2023.
- [9] I. M. Karo Karo, S. Dewi, and A. Perdana, “Implementasi Text Summarization Pada Review Aplikasi Digital Library System Menggunakan Metode Maximum Marginal Relevance,” *JEKIN - Jurnal Teknik Informatika*, vol. 4, no. 1, pp. 25–31, Apr. 2024, doi: 10.58794/jekin.v4i1.671.
- [10] C. D. Manning, P. Raghavan, and H. Schütze, *An Introduction to Information Retrieval*. Cambridge University Press, 2009.
- [11] M. Afrad, D. C. Febrianto, S. Wijayanto, and M. Y. Fathoni, “Sentiment Analysis of Visitor Reviews on Baturaden Tourist Attraction Using Machine Learning Methods,” *Edu Komputika Journal*, vol. 11, no. 1, Aug. 2024, doi: 10.15294/edukom.v11i1.10561.
- [12] G. M. Manik, “ANALISIS SENTIMEN PADA REVIEW PENGGUNA E-COMMERCE BIDANG PANGAN MENGGUNAKAN METODE SUPPORT VECTOR MACHINE : (Studi Kasus: Review Sayurbox Dan Tanihub Pada Google Play),” Universitas Pembangunan Nasional Veteran Jakarta, 2021.
- [13] D. Darwis, E. S. Pratiwi, and A. F. O. Pasaribu, “PENERAPAN ALGORITMA SVM UNTUK ANALISIS SENTIMEN PADA DATA TWITTER KOMISI PEMBERANTASAN KORUPSI REPUBLIK INDONESIA,” *Edutic - Scientific Journal of Informatics Education*, vol. 7, no. 1, Nov. 2020, doi: 10.21107/edutic.v7i1.8779.
- [14] F. I. Ramadhani, T. A. Yoga, and N. A. Verdhika, “Komparasi FastText dan TF-IDF Berbasis Random Forest pada Analisis Sentimen IKN di Youtube,” *TIN: Terapan Informatika Nusantara*, vol. 6, no. 12, pp. 2288–2301, May 2026, Accessed: Jun. 19, 2026. [Online]. Available: <https://ejurnal.seminar-id.com/index.php/tin/article/view/9749>