

Komparasi Kinerja Support Vector Machine dan Naive Bayes Teroptimasi pada Sentimen Ulasan Blu by BCA

Fathony Mursyid¹, Marcellino Andelta Pinem², Abraham Imanuel Sinaga³, Haichal Ryan Saputra⁴, Eldika Rubiana⁵, Fuad Nur Hassan^{6*}

^{1,2,3,4,5,6}Informatika, Universitas Bina Sarana Informatika, Depok, Indonesia

Email: ¹gusfathon.fm@gmail.com, ²marcellpinem@gmail.com, ³abrhm.imnl@gmail.com, ⁴haikalrian04@gmail.com, ⁵eldikarubiana@gmail.com

Email Penulis Korespondensi: ⁶fuad.fnu@bsi.ac.id

Abstrak – Pertumbuhan pesat aplikasi perbankan digital Blu by BCA Digital menghasilkan volume ulasan pengguna yang besar dan tidak terstruktur. Penelitian ini bertujuan untuk menentukan model klasifikasi sentimen biner yang optimal dengan membandingkan kinerja algoritma Support Vector Machine (SVM) dan Naive Bayes (NB). Kedua algoritma dioptimasi menggunakan Particle Swarm Optimization (PSO) untuk memperoleh hiperparameter terbaik. Optimasi ini dilakukan guna meningkatkan performa klasifikasi pada data teks berbahasa Indonesia yang bersifat informal. Evaluasi model dilakukan menggunakan metrik Akurasi dan F1-Score. Hasil pengujian menunjukkan bahwa model PSO-SVM memiliki kinerja yang lebih unggul dengan akurasi 86,01% dan F1-Score (Negatif) 76,66%. Sebagai perbandingan, model PSO-NB menghasilkan akurasi 85,03% dan F1-Score (Negatif) 74,70%. Berdasarkan hasil tersebut, PSO-SVM direkomendasikan sebagai model yang paling andal untuk sistem pemantauan sentimen otomatis, khususnya dalam mendeteksi keluhan pelanggan pada industri *Fintech*.

Kata Kunci: Analisis Sentimen, Support Vector Machine, Naive Bayes, Particle Swarm Optimization, Fintech

Abstract – The rapid growth of the digital banking application Blu by BCA Digital has generated a large volume of unstructured user reviews. This study aims to determine the optimal binary sentiment classification model by comparing the performance of Support Vector Machine (SVM) and Naive Bayes (NB) algorithms. Both algorithms were optimized using Particle Swarm Optimization (PSO) to obtain the best hyperparameters. This optimization aims to improve classification performance on informal Indonesian text data. Model evaluation was conducted using Accuracy and F1-Score metrics. The results indicate that the PSO-SVM model achieved superior performance with an accuracy of 86.01% and an F1-Score (Negative) of 76.66%. In comparison, the PSO-NB model recorded an accuracy of 85.03% and an F1-Score (Negative) of 74.70%. These findings indicate that PSO-SVM is the most reliable model for implementation in automated sentiment monitoring systems, particularly for detecting customer complaints in the Fintech industry.

Keywords: Sentiment Analysis, Support Vector Machine, Naive Bayes, Particle Swarm Optimization, Fintech

1. PENDAHULUAN

Transformasi digital di Indonesia telah mengubah lanskap ekonomi secara fundamental, khususnya pada sektor perbankan. Adopsi teknologi *Financial Technology* (Fintech) kini menjadi kebutuhan utama masyarakat modern yang menuntut layanan keuangan yang cepat, efisien, dan mudah diakses. Kondisi ini mendorong munculnya berbagai bank digital (*neobank*) yang bersaing ketat, salah satunya adalah Blu by BCA Digital yang berhasil mengakuisisi jutaan pengguna. Seiring pertumbuhan ini, volume *User Generated Content* (UGC) berupa ulasan pada Google Play Store meningkat signifikan. Data ulasan ini memuat wawasan strategis, mulai dari keluhan teknis hingga kepuasan pelanggan, yang vital bagi pemantauan reputasi bisnis [1], [2].

Namun, pemanfaatan data ulasan dalam skala besar menghadirkan tantangan kompleksitas data. Ulasan pengguna sering kali tidak terstruktur dan ditulis dalam bahasa Indonesia informal yang penuh dengan singkatan tidak baku, istilah *slang*, serta kesalahan tipografi [3]. Kondisi ini menyebabkan analisis manual menjadi tidak efisien dan subjektif, sehingga diperlukan pendekatan *Sentiment Analysis* berbasis *Machine Learning* untuk mengekstraksi informasi secara otomatis.

Dalam literatur klasifikasi teks, *Support Vector Machine* (SVM) dan *Naive Bayes* (NB) adalah algoritma yang dominan digunakan. SVM unggul dalam menangani data berdimensi tinggi melalui prinsip *Structural Risk Minimization* (SRM) [4], [5], sedangkan Naive Bayes menawarkan efisiensi komputasi tinggi dengan pendekatan probabilistik [6], [7]. Kendati populer, kinerja kedua algoritma ini sangat sensitif terhadap konfigurasi *hyperparameter*. Penggunaan parameter *default* sering kali menyebabkan *underfitting* atau *overfitting*. Untuk mengatasi hal ini, optimasi metaheuristik seperti *Particle Swarm Optimization* (PSO) terbukti efektif dalam menemukan kombinasi *hyperparameter* optimal secara otomatis [8], [9], [10].

Meskipun optimasi model klasifikasi teks telah banyak dikaji, masih terdapat beberapa kesenjangan penelitian (*research gap*) yang signifikan. Pertama, sebagian besar penelitian membandingkan SVM dan Naive Bayes tanpa perlakuan optimasi *hyperparameter* yang setara, sehingga perbandingan kinerja yang dihasilkan cenderung bias [4], [5], [7]. Kedua, studi yang menerapkan optimasi PSO umumnya hanya berfokus pada peningkatan satu algoritma saja, baik

PSO-SVM maupun PSO-NB, tanpa menyediakan pembandingan yang dioptimasi secara identik [10], [11]. Ketiga, penelitian komparatif secara langsung antara PSO-SVM dan PSO-NB pada dataset ulasan *fintech* berbahasa Indonesia yang memiliki karakteristik bahasa informal dan tingkat *noise* tinggi masih sangat terbatas.

Untuk mengisi kesenjangan tersebut, kebaruan (*novelty*) penelitian ini terletak pada evaluasi komparatif yang adil (*fair comparison*) di mana kedua algoritma (SVM dan NB) dioptimasi menggunakan metode yang sama (PSO) pada dataset yang sama. Penelitian ini bertujuan untuk: (1) mengembangkan *pipeline preprocessing* teks yang *robust* untuk bahasa informal; (2) menerapkan PSO untuk optimasi *hyperparameter* SVM dan NB; serta (3) merekomendasikan model terbaik berdasarkan metrik akurasi, presisi, *recall*, dan *F1-Score*. Kontribusi ini diharapkan dapat memberikan panduan empiris yang lebih akurat bagi pengembangan sistem analisis sentimen di industri perbankan digital.

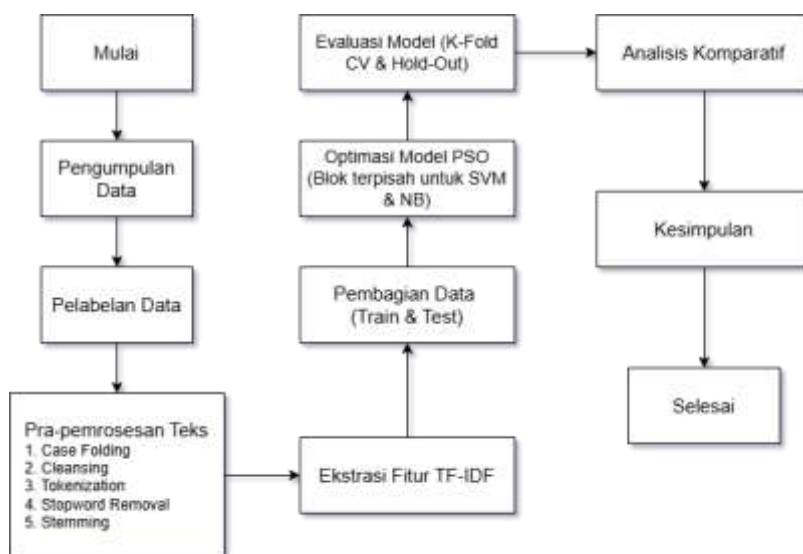
2. METODOLOGI PENELITIAN

Penelitian ini menerapkan dan membandingkan dua algoritma klasifikasi populer yang dioptimasi menggunakan metode metaheuristik. Berikut adalah penjelasan teoretis singkat dari metode-metode yang digunakan:

- Support Vector Machine (SVM):** Support Vector Machine adalah algoritma supervised learning yang bertujuan menemukan hyperplane (bidang pemisah) terbaik yang dapat memisahkan dua kelas data dengan margin maksimum. Dalam konteks teks, SVM sangat efektif untuk data berdimensi tinggi [1], [4]. Kinerjanya sangat bergantung pada pemilihan hyperparameter, terutama parameter C (Cost) yang mengontrol penalti terhadap kesalahan klasifikasi, dan γ yang menentukan pengaruh satu data latih pada *kernel non-linear* seperti RBF.
- Naive Bayes (NB):** Naive Bayes adalah algoritma klasifikasi probabilistik yang didasarkan pada Teorema Bayes. Disebut "naive" karena mengasumsikan bahwa setiap fitur bersifat independen satu sama lain [6], [7]. Untuk klasifikasi teks, varian Multinomial Naive Bayes sering digunakan dengan hyperparameter Alpha (α) sebagai parameter smoothing untuk mengatasi masalah probabilitas nol pada kata yang tidak muncul di data latih.
- Particle Swarm Optimization (PSO):** PSO adalah algoritma optimasi metaheuristik yang terinspirasi dari perilaku sosial kawanan [8], [10]. PSO bekerja dengan populasi "partikel", di mana setiap partikel merepresentasikan satu set solusi (kombinasi hyperparameter). Partikel-partikel ini bergerak di dalam ruang pencarian, belajar dari pengalaman terbaiknya sendiri (*personal best*) dan pengalaman terbaik kawanan (*global best*), untuk menemukan solusi optimal yang memaksimalkan *fitness function* [9], [11].

Terbaiknya sendiri (*personal best*) dan pengalaman terbaik dari kawanan (*global best*), untuk secara kolaboratif menemukan solusi optimal yang memaksimalkan *fitness function* (dalam penelitian ini, *F1-Score*).

Rancangan penelitian ini disusun mengikuti alur kerja yang terstruktur untuk memastikan validitas hasil dan kemudahan replikasi. Setiap tahapan, mulai dari akuisisi data hingga analisis akhir, dijelaskan secara detail dalam sub-bagian berikut, dengan alur keseluruhan yang dapat divisualisasikan melalui sebuah diagram.



Gambar 1. Diagram Alur Penelitian

Gambar 1 mengilustrasikan kerangka kerja sistematis yang diterapkan dalam penelitian ini. Tahapan penelitian dirancang secara berurutan, dimulai dari pengumpulan data hingga analisis komparatif, yang secara rinci dibahas pada Sub-bab 2.1 hingga 2.4.

Pada tahap Pengumpulan dan Pelabelan Data (Sub-bab 2.1), data mentah diakuisisi dan divalidasi. Proses dilanjutkan dengan Pra-pemrosesan Teks (Sub-bab 2.2) untuk mereduksi *noise* dan menstandarisasi fitur input menggunakan TF-IDF. Inti dari penelitian ini terletak pada tahap Optimasi Hyperparameter (Sub-bab 2.3), di mana algoritma PSO digunakan untuk mencari kombinasi parameter terbaik bagi SVM dan Naive Bayes.

Sesuai catatan pada desain alur, penentuan ruang pencarian (*search space*) pada PSO tidak dilakukan secara acak, melainkan dibatasi berdasarkan rekomendasi literatur terdahulu untuk menjaga efisiensi komputasi. Pada model SVM, parameter regularisasi (C) ditetapkan pada rentang $[0,1-100]$ dan parameter kernel Gamma (γ) pada $[0,001-1]$ dengan mengacu pada studi Janaah dan Nugroho serta tinjauan teoritis oleh Cervantes [12], [13]. Sementara itu, bobot inersia PSO diatur menurun secara linier dari $(0.9 \rightarrow 0.4)$ sesuai praktik umum pada algoritma metaheuristik [8]. Tahap akhir penelitian adalah evaluasi (Sub-bab 2.4) menggunakan skema *K-Fold Cross-Validation* untuk memastikan validitas model.

2.1 Pengumpulan dan Pelabelan Data

Objek utama penelitian adalah data ulasan pengguna aplikasi Blu by BCA Digital yang bersumber dari platform publik Google Play Store. Proses pengumpulan data menghasilkan total **24,361 ulasan** unik. Mengingat tujuan penelitian adalah klasifikasi sentimen biner, data ini kemudian dilabeli secara otomatis berdasarkan sistem rating bintang yang diberikan oleh pengguna. Ulasan dengan rating 4 hingga 5 bintang dikategorikan sebagai sentimen **Positif**, menghasilkan **16,969 data**. Sementara itu, ulasan dengan rating 1 hingga 3 bintang dikategorikan sebagai sentimen **Negatif**, menghasilkan **7,392 data**. Terdapat ketidakseimbangan kelas (*class imbalance*) dalam dataset ini, di mana kelas Positif lebih dominan. Oleh karena itu, penggunaan metrik F1-Score menjadi sangat relevan. Dataset ini kemudian dibagi menjadi dua bagian: **90% data latih** yang digunakan untuk training dan validasi model, dan **10% data uji** yang digunakan untuk pengujian final guna mengukur kemampuan generalisasi model.

2.2 Pra-pemrosesan Teks dan Ekstraksi Fitur

Tahap ini krusial untuk mengubah data teks mentah yang tidak terstruktur menjadi format yang bersih dan terstruktur yang dapat diolah oleh algoritma *machine learning*. Proses yang dilakukan meliputi:

- Case Folding:** Mengubah seluruh teks menjadi huruf kecil untuk memastikan konsistensi.
- Cleansing:** Menghapus semua karakter yang tidak relevan seperti angka, tanda baca, simbol, dan URL.
- Tokenization:** Memecah setiap ulasan menjadi unit-unit kata individual (token).
- Stopword Removal:** Menghapus kata-kata umum dalam Bahasa Indonesia yang tidak memiliki makna sentimen (misalnya: "yang", "di", "dan", "adalah") menggunakan daftar stopwords dari library Sastrawi.
- Stemming:** Mengubah setiap kata ke bentuk dasarnya (kata dasar) untuk mengurangi variasi kata dengan makna yang sama (misalnya: "membantu", "bantuan" menjadi "bantu"). Proses ini menggunakan algoritma Nazief & Adriani yang diimplementasikan dalam library Sastrawi.

Setelah teks dibersihkan, langkah selanjutnya adalah mengubahnya menjadi representasi vektor numerik menggunakan metode pembobotan **Term Frequency-Inverse Document Frequency (TF-IDF)**. TF-IDF mengukur pentingnya sebuah kata dalam sebuah dokumen relatif terhadap keseluruhan korpus. Bobot TF-IDF dihitung menggunakan rumus berikut:

$$w_{i,j} = tf_{i,j} \times \log\left(\frac{N}{df_i}\right)$$

di mana $w_{i,j}$ adalah bobot TF-IDF kata i dalam dokumen j , $tf_{i,j}$ adalah frekuensi kata i dalam dokumen j , N adalah jumlah total dokumen dalam korpus, dan df_i adalah jumlah dokumen yang mengandung kata i .

2.3 Pemodelan dan Optimasi Hyperparameter

Penelitian ini membandingkan dua model yang dioptimasi menggunakan PSO dengan parameter 15 partikel dan 20 iterasi, menggunakan F1-Score sebagai fitness function

- Support Vector Machine (SVM):** Algoritma klasifikasi yang bertujuan mencari *hyperplane* pemisah terbaik yang memaksimalkan margin antara dua kelas data. *Hyperparameter* yang dioptimasi untuk SVM dengan kernel RBF adalah:
 - C (Parameter Regularisasi):** Mengontrol *trade-off* antara memaksimalkan margin dan meminimalkan kesalahan klasifikasi. Rentang pencarian: $[0.01, 100.0]$. Pemilihan batas rentang ini didasarkan pada rekomendasi survei komprehensif yang menyatakan bahwa rentang tersebut ideal untuk menghindari *overfitting* pada klasifikasi data berdimensi tinggi [13].
 - Gamma (Koefisien Kernel):** Mendefinisikan seberapa besar pengaruh satu contoh training. Rentang pencarian: $[0.0001, 1.0]$. Rentang ini dipilih merujuk pada rekomendasi yang menyatakan bahwa nilai Gamma rendah dibutuhkan untuk menangkap pola global pada data teks berdimensi tinggi, sedangkan batas atas 1.0 dibatasi untuk menjaga generalisasi model agar tetap optimal [14].

b. **Naive Bayes (NB):** Algoritma klasifikasi probabilistik yang didasarkan pada Teorema Bayes dengan asumsi independensi antar fitur. *Hyperparameter* yang dioptimasi untuk Multinomial Naive Bayes adalah:

1. **Alpha (Smoothing Parameter):** Parameter untuk *Laplace/Lidstone smoothing* yang menangani masalah kata yang tidak ada di data latih. Rentang pencarian: [0.001, 1.0]. Pengaturan ini merujuk pada penelitian ulasan aplikasi Bukalapak [15], yang membuktikan bahwa penggunaan nilai alpha desimal (mendekati 0) sangat krusial untuk mengatasi sparsity data akibat banyaknya kata unik atau bahasa gaul (slang) yang tidak muncul di data latih, dibandingkan menggunakan nilai default 1.0.

2.4 Skenario Pengujian dan Metrik Evaluasi

Kinerja model dievaluasi menggunakan dua skenario pengujian utama:

- a. **K-Fold Cross-Validation (K=10):** Digunakan pada 90% data latih untuk mengukur stabilitas dan performa rata-rata model. Data latih dibagi menjadi 10 bagian (fold), di mana model dilatih pada 9 fold dan diuji pada 1 fold yang tersisa, proses ini diulang sebanyak 10 kali.
- b. **Hold-out Test:** Pengujian final yang dilakukan pada 10% data uji yang sama sekali belum pernah dilihat oleh model selama proses training atau optimasi. Ini memberikan estimasi yang tidak bias tentang kinerja model di dunia nyata.

Metrik evaluasi yang digunakan adalah Akurasi dan F1-Score. F1-Score dipilih sebagai metrik utama karena kemampuannya menangani dataset tidak seimbang [4], [16]:

$$F1 - Score = 2 \times \frac{Presisi \times Recall}{Presisi + Recall}$$

3. HASIL DAN PEMBAHASAN

Pada bagian ini, disajikan hasil kuantitatif dari seluruh skenario pengujian, diikuti dengan pembahasan mendalam yang menginterpretasikan temuan, membandingkannya dengan penelitian terkait, dan mendiskusikan implikasi praktisnya.

3.1 Hasil Eksperimen

Proses eksperimen menghasilkan serangkaian data kuantitatif yang menjadi dasar analisis komparatif, mencakup *hyperparameter* optimal yang ditemukan, hasil uji stabilitas, dan kinerja final model pada data uji.

3.1.1 Hasil Optimasi Hyperparameter

Setelah menjalankan proses optimasi PSO dengan 15 partikel dan 20 iterasi, kombinasi *hyperparameter* terbaik yang memaksimalkan *F1-Score* pada data validasi internal berhasil ditemukan untuk kedua model. Hasilnya disajikan pada Tabel 1. Nilai-nilai ini, yang berbeda dari nilai *default* pada umumnya, menunjukkan bahwa proses optimasi berhasil menyesuaikan model dengan karakteristik spesifik dari dataset ulasan Blu by BCA.

Tabel 1. Hyperparameter Optimal Hasil

Model	Hyperparameter	Nilai Optimal
PSO-SVM	C	2.4057
	Gamma	0.0518
PSO-NB	Alpha	0.4942

3.1.2 Hasil Uji Kinerja Stabilitas Model (K-Fold Cross-Validation)

Stabilitas model dievaluasi menggunakan 10-Fold Cross-Validation pada data latih. Tabel 2 merangkum rata-rata F1-Score dan standar deviasi (σ) yang dihasilkan. Standar deviasi yang lebih rendah mengindikasikan kinerja model yang lebih konsisten dan stabil di berbagai subset data.

Tabel 2. Hasil Rata-Rata 10-Fold Cross-Validation

Model	F1-Score Rata-Rata	Standar Deviasi (σ)
PSO-SVM	0.8445	0.0076
PSO-NB	0.8458	0.0086

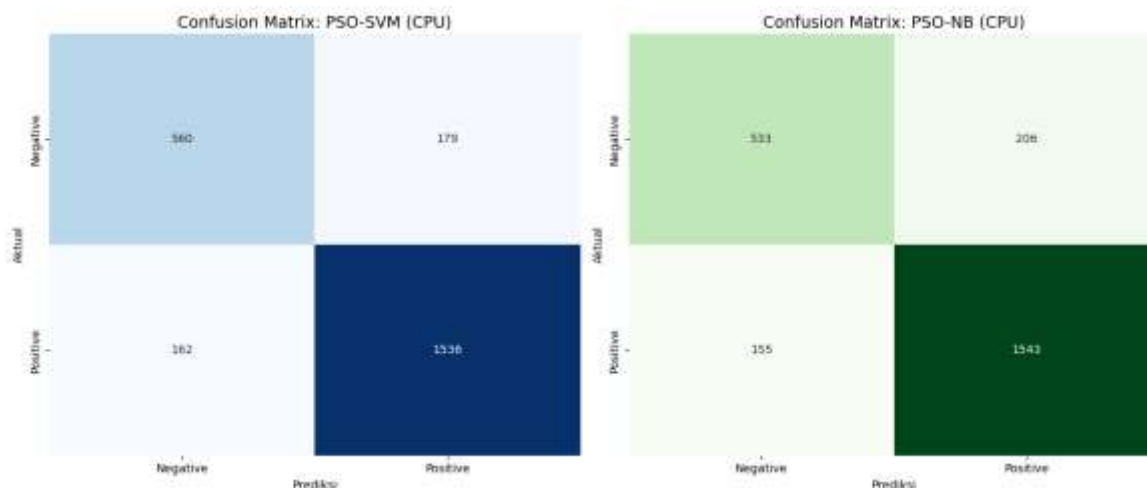
Hasil pada Tabel 2 menunjukkan bahwa kedua model memiliki performa rata-rata yang sangat kompetitif. Model PSO-NB memiliki F1-Score rata-rata sedikit lebih tinggi, namun juga memiliki standar deviasi yang lebih besar, yang mengindikasikan bahwa kinerja model PSO-SVM sedikit lebih stabil di berbagai segmen data latih.

3.1.3 Hasil Uji Kinerja Final Model (Hold-out Test)

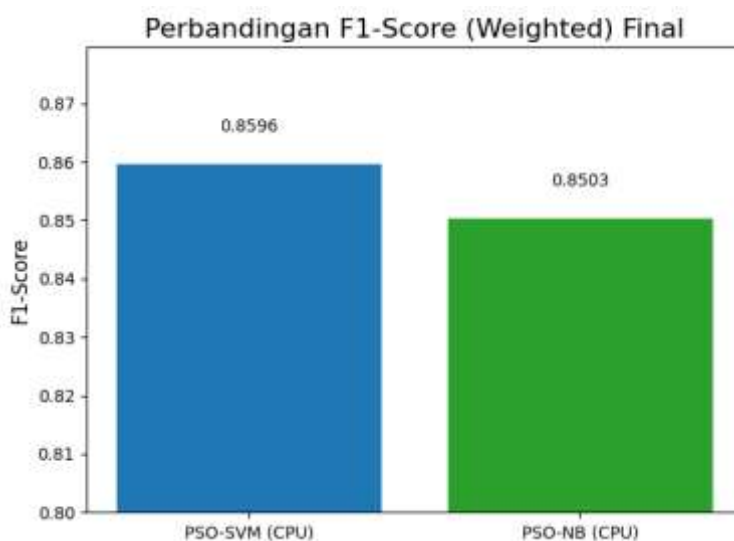
Kinerja final dan kemampuan generalisasi model diukur menggunakan 10% data uji yang terpisah. Tabel 3 menyajikan hasil perbandingan metrik Akurasi dan berbagai F1-Score. Untuk visualisasi yang lebih jelas, Gambar 2 menyajikan *confusion matrix* dari kedua model, dan Gambar 3 menampilkan perbandingan F1-Score secara grafis.

Tabel 3. Metrik Komparatif Final pada Data Uji

Metrik	PSO-SVM	PSO-NB
Akuras	0.8601	0.8519
F1-Score (Negatif)	0.7666	0.7470
F1-Score (Positif)	0.9001	0.8953
F1-Score (Weighted)	0.8596	0.8503



Gambar 2. Confusion Matrix Untuk PSO-SVM Dan PSO-NB



Gambar 3. Grafik Batang Perbandingan F1-Score Final

3.2 Pembahasan

Hasil eksperimen secara konsisten menunjukkan bahwa model PSO-SVM memiliki keunggulan kinerja dibandingkan PSO-NB, meskipun dengan selisih yang tidak terlalu besar. Namun, narasi yang lebih dalam terletak pada interpretasi mengapa keunggulan ini terjadi, analisis stabilitas, dan *trade-off* fundamental antara kinerja dan efisiensi komputasi, yang memiliki implikasi praktis yang signifikan.

3.2.1 Analisis Mendalam Keunggulan PSO-SVM dibandingkan PSO-NB

Berdasarkan hasil komparasi metrik pada Tabel 3 serta visualisasi perbandingan pada Gambar 3, algoritma PSO-SVM menunjukkan superioritas kinerja dibandingkan PSO-NB dengan capaian Akurasi 86,01% dan F1-Score (Negatif) 76,66%. Keunggulan performa ini bukan merupakan kejadian acak, melainkan hasil dari kemampuan intrinsik SVM dalam menangani karakteristik data ulasan berdimensi tinggi yang gagal diakomodasi secara optimal oleh Naive Bayes. Penyebab mendasar keunggulan ini dapat dijelaskan melalui analisis parameter optimal (Tabel 1), pola kesalahan klasifikasi (Gambar 2), dan stabilitas validasi (Tabel 2).

Faktor utama keunggulan SVM terletak pada fleksibilitas decision boundary yang dihasilkan melalui "*Kernel Trick*". Merujuk pada Tabel 1, optimasi PSO berhasil menemukan nilai parameter Gamma (λ) yang rendah, yaitu 0.0518, untuk model SVM. Dalam teori Support Vector, nilai Gamma yang rendah memperluas radius pengaruh data latih, sehingga batas pemisah antar kelas menjadi lebih halus (*smoother*) dan adaptif [13]. Hal ini sangat krusial untuk menangani data ulasan aplikasi mobile yang sering kali mengandung noise atau kata-kata tidak baku. Sebaliknya, Naive Bayes bekerja dengan asumsi independensi fitur yang kaku, di mana setiap kata dianggap tidak memiliki korelasi dengan kata lainnya dalam menentukan sentimen [6], [7].

Kelemahan asumsi independensi pada Naive Bayes terbukti secara empiris pada Gambar 2. Matriks kebingungan menunjukkan bahwa PSO-NB menghasilkan jumlah *False Positive* yang signifikan lebih tinggi (206 ulasan) dibandingkan PSO-SVM (179 ulasan). Tingginya *False Positive* pada NB mengindikasikan kegagalan model dalam menangkap konteks kalimat yang kompleks (misalnya ulasan sarkas atau ulasan yang mengandung kata positif namun bermakna keluhan), karena NB hanya mengkalkulasi probabilitas kemunculan kata secara terisolasi. Sementara itu, PSO-SVM dengan parameter Regularisasi (C) sebesar 2.4057 (Tabel 1) mampu memetakan hubungan non-linear antar fitur secara lebih efektif, meminimalkan kesalahan pelabelan ulasan negatif menjadi positif.

Selain itu, PSO-SVM menunjukkan stabilitas yang lebih baik dalam generalisasi data. Walaupun rata-rata performa validasi kedua model pada Tabel 2 terlihat kompetitif, PSO-SVM mampu mempertahankan performa yang konsisten saat diterapkan pada data uji (*unseen data*). Hal ini tercermin pada Tabel 3, di mana selisih F1-Score kelas Negatif antara SVM dan NB mencapai hampir 2% (76.66% vs 74.70%). Kemampuan SVM untuk menyeimbangkan margin kesalahan melalui optimasi nilai C menjadikannya metode yang lebih robust dibandingkan Naive Bayes yang murni bergantung pada probabilitas statistik frekuensi [4], [14]. Oleh karena itu, SVM terbukti lebih andal untuk dijadikan model inti dalam analisis sentimen ulasan perbankan digital Blu by BCA.

3.2.2 Analisis Kesalahan (Error Analysis) dan Sampel Misklasifikasi

Berdasarkan hasil evaluasi, meskipun PSO-SVM menunjukkan performa yang lebih unggul dibandingkan PSO-NB, analisis kesalahan (*error analysis*) tetap dilakukan untuk memahami karakteristik kegagalan prediksi pada kedua model. Tabel 4 berikut menyajikan distribusi kesalahan prediksi (*misclassification*) dari total 2.437 data uji.

Tabel 4. Perbandingan Distribusi Kesalahan Prediksi

Model	Total Error	Error Rate	False Positive (FP)	False Negative (FN)
PSO-SVM	341	13.99%	179	162
PSO-NB	361	14.81%	206	155

Terlihat bahwa PSO-SVM memiliki tingkat kesalahan yang lebih rendah, terutama dalam menekan angka *False Positive* (179 kasus) dibandingkan PSO-NB (206 kasus). Untuk memahami penyebab kesalahan, Tabel 5 menampilkan sampel ulasan yang mengalami *misclassification* pada PSO-SVM.

Tabel 5. Sampel Misklasifikasi PSO-SVM (False Positive & False Negative)

Jenis Error	Teks Ulasan (<i>Cleaned</i>)	Label Aktual	Prediksi PSO-SVM
False Positive	"transfer g gratis"	Negatif	Positif
False Positive	"aja kyk induk aplikasi sulit msh bnyak bank digital yg mudah nyesel downloadn"	Negatif	Positif
False Negative	"akun nomer rekening salah gimana nih g tuk transaksi"	Positif	Negatif
False Negative	"tarik tunai password d reset ya ngulang masuk blu"	Positif	Negatif

Berdasarkan sampel misclassification yang ditampilkan pada Tabel 5, penyebab kesalahan klasifikasi dapat dikelompokkan ke dalam dua faktor utama, yaitu kegagalan dalam mendeteksi negasi dan slang, serta ambiguitas label dan kompleksitas struktur kalimat.

Kesalahan *False Positive*, di mana model memprediksi sentimen positif pada ulasan yang berlabel negatif, umumnya disebabkan oleh kegagalan model dalam menangkap konteks negasi yang ditulis menggunakan singkatan tidak baku. Sebagai contoh, pada ulasan “transfer g gratis”, model PSO-SVM cenderung memberikan bobot positif yang tinggi pada kata “gratis” dan gagal mengasosiasikan singkatan “g” (tidak/gak) sebagai penanda negasi. Temuan ini menunjukkan bahwa meskipun SVM mampu menangani pola non-linear, representasi fitur berbasis *Bag-of-Words* atau TF-IDF standar masih memiliki keterbatasan dalam mempertahankan informasi urutan kata, yang berperan penting dalam mendeteksi struktur negasi pada teks informal [15].

Sementara itu, pada kasus *False Negative* (label aktual positif namun diprediksi negatif), ditemukan indikasi adanya ambiguitas label atau label noise pada data. Ulasan seperti “tarik tunai password d reset ya ngulang masuk blu” secara linguistik mengandung ekspresi keluhan yang merefleksikan sentimen negatif, namun pada dataset diberi label positif. Dalam konteks ini, prediksi negatif yang dihasilkan PSO-SVM justru lebih selaras dengan makna semantik kalimat dibandingkan label aktualnya. Fenomena ini mengonfirmasi temuan Janaah dan Nugroho [12] yang menyatakan bahwa tantangan utama dalam klasifikasi sentimen tidak hanya berasal dari keterbatasan algoritma, tetapi juga dari inkonsistensi pelabelan yang bersifat subjektif.

3.2.3 Evaluasi Kebutuhan Data Tambahan

Berdasarkan hasil analisis kesalahan klasifikasi yang telah dibahas sebelumnya, kebutuhan akan penambahan data dapat dievaluasi dari aspek kuantitas serta kualitas dan keseimbangan data. Dari sisi kuantitas, dataset yang digunakan dalam penelitian ini terdiri dari 24.361 data ulasan, yang secara umum sudah memadai untuk proses pembelajaran model. Oleh karena itu, penambahan data dalam jumlah besar tanpa disertai peningkatan kualitas data diperkirakan tidak akan memberikan penurunan error rate yang signifikan.

Sebaliknya, permasalahan utama yang teridentifikasi terletak pada kualitas dan distribusi data. Dominasi kesalahan *False Positive* mengindikasikan adanya bias model terhadap kelas mayoritas, serta ditemukan pula noise pada label, seperti ulasan bernada negatif yang diberi rating tinggi. Kondisi ini menunjukkan bahwa peningkatan kinerja model tidak bergantung pada penambahan data semata, melainkan pada penerapan strategi penanganan data yang lebih tepat. Beberapa pendekatan yang relevan meliputi pembersihan label secara manual pada kasus bias rating serta penerapan teknik oversampling yang lebih canggih untuk memperkaya variasi kosakata slang negatif, seperti penggunaan singkatan “g”, “gk”, dan “gbs”.

3.2.4 Peran Optimasi PSO dan Stabilitas Model

Hasil optimasi menunjukkan PSO berhasil menemukan hyperparameter yang jauh dari nilai default, mengonfirmasi hipotesis bahwa optimasi adalah langkah krusial untuk memaksimalkan potensi model [8], [9], [10]. Temuan ini sejalan dengan penelitian yang melaporkan peningkatan signifikan setelah menerapkan PSO untuk optimasi hyperparameter.

Analisis stabilitas menunjukkan meskipun PSO-NB memiliki rata-rata F1-Score sedikit lebih tinggi, PSO-SVM menunjukkan standar deviasi lebih rendah (0.0076 vs 0.0086). Ini mengindikasikan kinerja PSO-SVM lebih stabil dan konsisten di berbagai subset data, yang merupakan atribut penting untuk implementasi sistem jangka panjang [17].

3.2.5 Perbandingan dengan Penelitian Terdahulu

Akurasi PSO-SVM (86.01%) menunjukkan kinerja yang kompetitif jika dibandingkan beberapa penelitian sejenis pada analisis sentimen aplikasi *Fintech* [1], [3]. Temuan ini sejalan dengan hasil penelitian terdahulu yang melaporkan bahwa SVM dengan optimasi PSO dapat meningkatkan performa klasifikasi teks [4], [16].

Lebih lanjut, studi-studi sebelumnya juga menunjukkan bahwa penerapan PSO untuk optimasi hyperparameter berkontribusi positif terhadap peningkatan kinerja model machine learning [9], [10], [11]. Hal ini menegaskan bahwa proses tuning hyperparameter merupakan faktor penting dalam pengembangan sistem analisis sentimen berbasis pembelajaran mesin.

3.2.6 Implikasi Praktis: Trade-off Kinerja vs. Efisiensi

Meskipun PSO-SVM unggul dalam kinerja, perbedaan signifikan terletak pada efisiensi komputasi. SVM membutuhkan sumber daya komputasi lebih besar dibanding Naive Bayes [6], [7]. Hal ini menyajikan trade-off klasik:

- Skenario Prioritas Akurasi (PSO-SVM):** Untuk analisis mendalam periodik (laporan bulanan/triwulanan) di mana akurasi maksimal adalah prioritas dan waktu komputasi bukan kendala utama [1], [4].

- b. **Skenario Prioritas Efisiensi (PSO-NB):** Untuk sistem pemantauan real-time seperti dashboard yang mengklasifikasikan ulasan baru secara instan, di mana kecepatan PSO-NB menjadikannya pilihan lebih praktis meskipun dengan sedikit penurunan akurasi [6], [11].

3.2.7 Limitasi Penelitian

Beberapa batasan penelitian ini meliputi: (1) dataset hanya dari Google Play Store, sehingga sentimen dari platform lain tidak terwakili; (2) fokus pada klasifikasi sentimen biner tanpa analisis *aspect-based* atau *topic modeling* yang lebih mendalam; (3) perbandingan terbatas pada dua algoritma klasik, sementara model *deep learning* seperti IndoBERT mungkin memberikan kinerja lebih tinggi meskipun dengan biaya komputasi lebih besar [18].

4. KESIMPULAN

Penelitian ini telah berhasil menjawab tantangan klasifikasi sentimen pada ulasan pengguna aplikasi Blu by BCA Digital yang memiliki karakteristik volume besar, tidak terstruktur, dan dominan menggunakan ragam bahasa informal. Melalui pendekatan komparatif yang sistematis, penelitian ini membuktikan bahwa integrasi teknik optimasi *Particle Swarm Optimization* (PSO) memegang peranan vital dalam meningkatkan performa algoritma *Support Vector Machine* (SVM) dan *Naive Bayes* (NB). Berdasarkan hasil evaluasi komprehensif, model PSO-SVM secara konsisten mengungguli PSO-NB di seluruh skenario pengujian dengan mencatatkan akurasi tertinggi sebesar 86,01% dan F1-Score untuk kelas negatif sebesar 76,66%. Keunggulan signifikan ini mengindikasikan bahwa kemampuan SVM dengan kernel RBF dalam memetakan data ke ruang dimensi tinggi untuk membentuk *hyperplane* pemisah yang optimal jauh lebih andal dalam menangani kompleksitas data ulasan dibandingkan pendekatan probabilistik yang mengasumsikan independensi fitur pada *Naive Bayes*.

Sebagai implikasi praktis, PSO-SVM direkomendasikan sebagai metode utama dalam pengembangan sistem pemantauan sentimen otomatis untuk mitigasi risiko reputasi di industri perbankan digital. Meskipun kinerja model sudah optimal, penelitian ini masih memiliki ruang untuk penyempurnaan di masa depan (*works for future*). Keterbatasan utama yang perlu diatasi adalah penanganan aspek linguistik yang kompleks, seperti deteksi negasi bertingkat dan majas ironi yang kerap memicu kesalahan klasifikasi. Oleh karena itu, penelitian selanjutnya sangat disarankan untuk menerapkan teknik penyeimbangan data (*oversampling*) seperti SMOTE atau mengeksplorasi arsitektur *Deep Learning* guna meningkatkan sensitivitas terhadap kelas minoritas serta menangkap konteks semantik secara lebih mendalam.

UCAPAN TERIMAKASIH

Terima kasih disampaikan kepada pihak-pihak yang telah mendukung terlaksananya penelitian ini.

REFERENCES

- [1] M. Rahardi and V. Daarten Pandiangan, "INTERNATIONAL JOURNAL ON INFORMATICS VISUALIZATION journal homepage: www.joiv.org/index.php/joiv INTERNATIONAL JOURNAL ON INFORMATICS VISUALIZATION Sentiment Analysis of Neobank Digital Banking Using Support Vector Machine Algorithm in Indonesia." [Online]. Available: <https://play.google.com/store/apps/details?id=com.bnc.finan>
- [2] R. Alawaji and A. Aloraini, "Sentiment Analysis of Digital Banking Reviews Using Machine Learning and Large Language Models," *Electronics (Switzerland)*, vol. 14, no. 11, Jun. 2025, doi: 10.3390/electronics14112125.
- [3] M. D. Bimantara and I. Zufria, "Text Mining Sentiment Analysis on Mobile Banking Application Reviews using TF-IDF Method with Natural Language Processing Approach," *JINAV: Journal of Information and Visualization*, vol. 5, no. 1, pp. 115–123, Jul. 2024, doi: 10.35877/454ri.jinav2772.
- [4] Y. Christian, T. Wibowo, and M. Lyawati, "Sentiment Analysis by Using Naïve Bayes Classification and Support Vector Machine, Study Case Sea Bank," *Sinkron*, vol. 9, no. 1, pp. 258–275, Jan. 2024, doi: 10.33395/sinkron.v9i1.13141.
- [5] J. O. Leandro and M. I. Fianty, "INTERNATIONAL JOURNAL ON INFORMATICS VISUALIZATION journal homepage: www.joiv.org/index.php/joiv INTERNATIONAL JOURNAL ON INFORMATICS VISUALIZATION Evaluation of Sentiment Analysis Methods for Social Media Applications: A Comparison of Support Vector Machines and Naïve Bayes." [Online]. Available: www.joiv.org/index.php/joiv
- [6] I. A. H. Hidayah, R. Kusumawati, Z. Abidin, and M. Imamuddin, "Analysis of Public Sentiment Towards The TikTok Application Using The Naive Bayes Algorithm and Support Vector Machine," *Journal of Computer Networks, Architecture and High Performance Computing*, vol. 6, no. 2, pp. 881–891, Jun. 2024, doi: 10.47709/cnahpc.v6i2.3990.



- [7] M. D. Rizkiyanto, M. D. Purbolaksono, and W. Astuti, "Sentiment Analysis Classification on PLN Mobile Application Reviews using Random Forest Method and TF-IDF Feature Extraction," *INTEK: Jurnal Penelitian*, vol. 11, no. 1, pp. 37–43, Apr. 2024, doi: 10.31963/intek.v11i1.4774.
- [8] A. Hosseinalipour and R. Ghanbarzadeh, "A novel metaheuristic optimisation approach for text sentiment analysis," *International Journal of Machine Learning and Cybernetics*, vol. 14, no. 3, pp. 889–909, Mar. 2023, doi: 10.1007/s13042-022-01670-z.
- [9] E. Elgeldawi, A. Sayed, A. R. Galal, and A. M. Zaki, "Hyperparameter tuning for machine learning algorithms used for arabic sentiment analysis," *Informatics*, vol. 8, no. 4, Dec. 2021, doi: 10.3390/informatics8040079.
- [10] R. Azizah Arilya, Y. Azhar, and D. Rizki Chandranegara, "Sentiment Analysis on Work from Home Policy Using Naïve Bayes Method and Particle Swarm Optimization," *Jurnal Ilmiah Teknik Elektro Komputer dan Informatika*, vol. 7, no. 3, p. 433, Dec. 2021, doi: 10.26555/jiteki.v7i3.22080.
- [11] S. Siswanto, I. Kurniawan, and S. A. Thamrin, "Naïve Bayes Algorithm with Feature Selection Using Particle Swarm Optimization," *Jurnal Varian*, vol. 7, no. 2, pp. 127–136, Jun. 2024, doi: 10.30812/varian.v7i2.2409.
- [12] M. Janaah and A. Nugroho, "Performance of SVM Optimized with PSO as Classification Method for Sentiment Analysis UNNES's Social Media," *Jurnal Infotel*, vol. 17, no. 1, pp. 68–80, 2025, doi: 10.20895/infotel.v17i1.1266.
- [13] J. Cervantes, F. Garcia-Lamont, L. Rodríguez-Mazahua, and A. Lopez, "A comprehensive survey on support vector machine classification: Applications, challenges and trends," *Neurocomputing*, vol. 408, pp. 189–215, 2020, doi: 10.1016/j.neucom.2019.10.118.
- [14] F. Andrianto, A. Fadlil, and I. Riadi, "Linear Kernel Optimization of Support Vector Machine Algorithm on Online Marketplace Sentiment Analysis," *Komputasi: Jurnal Ilmiah Ilmu Komputer dan Matematika*, vol. 21, no. 1, pp. 68–82, 2024, doi: 10.33751/komputasi.v21i1.9266.
- [15] F. Akmal, A. D. Riyanto, and I. Darmayanti, "Optimization Naïve Bayes Algorithm in Sentiment Analysis of Bukalapak App Reviews," *Sinkron*, vol. 9, no. 1, pp. 145–151, 2024, doi: 10.33395/sinkron.v9i1.13132.
- [16] C. N. Dang, M. N. Moreno-García, and F. De La Prieta, "Hybrid Deep Learning Models for Sentiment Analysis," *Complexity*, vol. 2021, 2021, doi: 10.1155/2021/9986920.
- [17] K. Chong and N. Shah, "Comparison of Naive Bayes and SVM Classification in Grid-Search Hyperparameter Tuned and Non-Hyperparameter Tuned Healthcare Stock Market Sentiment Analysis." [Online]. Available: www.ijacsa.thesai.org
- [18] L. J. Anreaja *et al.*, "JISA (Jurnal Informatika dan Sains) Naive Bayes and Support Vector Machine Algorithm for Sentiment Analysis Opensea Mobile Application Users in Indonesia," 2022.