

# IMPLEMENTASI HYBRID *LEXICON-BASED* DAN SVM UNTUK KLASIFIKASI ANALISIS SENTIMEN TERHADAP PELATIHAN BBPSDMP KOMINFO MAKASSAR

Nur Alam<sup>1)</sup>, Muhammad Faisal<sup>\*2)</sup>, Rizki Yusliana Bakti<sup>3)</sup>, Muhammad Syafaat<sup>4)</sup>,  
Andi Makbul Syamsuri<sup>5)</sup>, Muhyiddin AM Hayat<sup>6)</sup>, Andi Lukman Anas<sup>7)</sup>

1. Informatika, Universitas Muhammadiyah Makassar  
105841103621@student.unismuh.ac.id
2. Informatika, Universitas Muhammadiyah Makassar  
muhfaisal@unismuh.ac.id
3. Informatika, Universitas Muhammadiyah Makassar  
rizkiyusliana@unismuh.ac.id
4. Teknik Pengairan, Universitas Muhammadiyah Makassar  
syafaat\_skuba@unismuh.ac.id
5. Teknik Pengairan, Universitas Muhammadiyah Makassar  
amakbulsyamsuri@unismuh.ac.id
6. Informatika, Universitas Muhammadiyah Makassar  
muhyiddin@unismuh.ac.id
7. Informatika, Universitas Muhammadiyah Makassar  
lukmananas@unismuh.ac.id

## **Abstract**

*The evaluation of government training programs is often hindered by manual analysis of unstructured qualitative feedback, making the process inefficient and subjective. This study aims to implement and evaluate a sentiment classification model using a hybrid Lexicon-Based and Support Vector Machine approach to analyze participants' perceptions of the Vocational School Graduate Academy training organized by BBPSDMP Kominfo Makassar, as well as to compare the performance of a standard SVM model with a model optimized using Particle Swarm Optimization. This quantitative research employs 2,313 unstructured review data, which undergo text preprocessing, initial lexicon-based labeling, and TF-IDF feature extraction before being classified using an SVM with an RBF kernel. The results show that the SVM model optimized with PSO consistently outperforms the standard model across all four evaluation aspects, with the most significant accuracy improvement observed in the instructor category from 84.71% to 89.02% and in the assessor category reaching 91.46%. PSO optimization has proven effective in enhancing the model's ability to identify negative sentiments, which represent the minority class. The hybrid approach with PSO optimization is capable of producing a more accurate and balanced classification system, with practical implications as an objective automated evaluation tool.*

**Kata Kunci :** Analisis Sentimen, Hybrid Method, Lexicon-Based, Particle Swarm Optimization, Support Vector Machine

## A. PENDAHULUAN

Perkembangan teknologi informasi dan komunikasi mendorong pemerintah untuk menyelenggarakan program strategis yang bertujuan meningkatkan kualitas sumber daya manusia, sejalan dengan agenda transformasi digital nasional. Salah satu inisiatif utamanya adalah program *Vocational School Graduate Academy* (VSGA) yang dirancang untuk membekali angkatan kerja dengan keterampilan digital relevan [1]. Efektivitas program semacam ini sangat bergantung pada mekanisme evaluasi yang akurat. Namun, proses evaluasi yang ada saat ini di BBPSDMP Kominfo Makassar menghadapi tantangan signifikan dalam mengolah umpan balik peserta yang bersifat kualitatif dan tidak terstruktur. Analisis manual terhadap data kritik dan saran yang naratif tidak hanya memakan waktu, tetapi juga rentan terhadap bias interpretasi, sehingga mengurangi objektivitas hasil evaluasi [2].

Data umpan balik yang tidak terstruktur, dengan ragam bahasa informal dan variasi ekspresi, memerlukan pendekatan komputasional untuk ekstraksi makna yang objektif dan terukur. Variasi bahasa dan kurangnya struktur formal dapat memperumit proses pemaknaan informasi, yang pada akhirnya berisiko menghasilkan evaluasi yang kurang efisien dan subjektif [3]. Untuk mengatasi permasalahan ini, teknik analisis sentimen menawarkan solusi yang mampu mengotomatisasi proses klasifikasi opini ke dalam polaritas positif, negatif, atau netral [4]. Pendekatan ini memungkinkan pengolahan data dalam volume besar secara cepat dan konsisten, sehingga dapat mengubah data kualitatif naratif menjadi wawasan kuantitatif yang dapat ditindaklanjuti.

Penelitian ini mengusulkan implementasi model klasifikasi sentimen dengan pendekatan *hybrid* yang

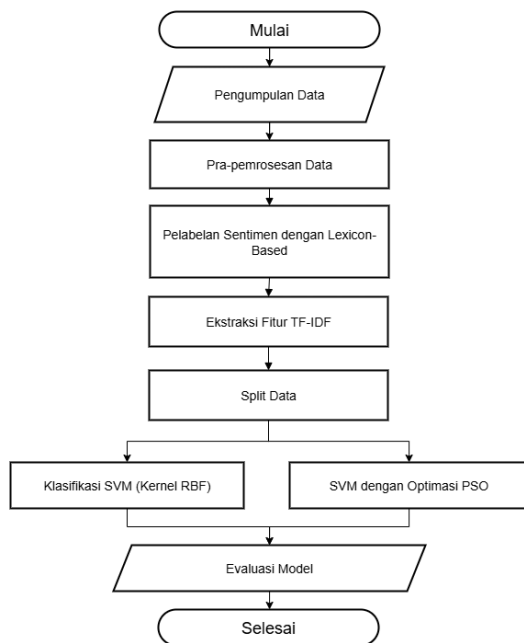
mengintegrasikan metode *Lexicon-Based* dan *Support Vector Machine* (SVM). Metode *Lexicon-Based* dimanfaatkan untuk pelabelan awal yang efisien, sementara SVM digunakan karena kemampuannya yang terbukti andal dalam menangani data teks berdimensi tinggi [5]. Lebih lanjut, untuk meningkatkan performa klasifikasi, terutama pada kelas sentimen minoritas, penelitian ini juga menerapkan *Particle Swarm Optimization* (PSO) untuk optimasi hyperparameter pada model SVM.

Oleh karena itu, penelitian ini bertujuan untuk: (1) mengimplementasikan model *hybrid Lexicon-SVM* untuk klasifikasi sentimen terhadap umpan balik peserta pelatihan VSGA, dan (2) menganalisis perbedaan performa antara model SVM dengan kernel RBF standar dan model SVM yang telah dioptimasi menggunakan PSO. Kontribusi utama penelitian ini terletak pada penerapan metode *hybrid* yang dioptimasi dalam domain evaluasi program pelatihan pemerintah yang spesifik, di mana studi sebelumnya lebih banyak berfokus pada data media sosial. Hasilnya diharapkan tidak hanya mengisi celah penelitian tersebut, tetapi juga menyediakan model evaluasi otomatis yang dapat memberikan rekomendasi berbasis data untuk peningkatan kualitas program secara berkelanjutan.

## B. METODE PENELITIAN

Penelitian ini menerapkan pendekatan kuantitatif dengan desain penelitian yang berfokus pada implementasi dan evaluasi komparatif model klasifikasi sentimen. Objek penelitian adalah data umpan balik kualitatif berupa kritik dan saran dari peserta pelatihan *Vocational School Graduate Academy* yang diselenggarakan oleh BBPSDMP Kominfo Makassar pada tahun 2024. Alur penelitian dirancang secara sistematis yang mencakup tahapan

pengumpulan data, pra-pemrosesan, pelabelan, ekstraksi fitur, pemodelan, hingga evaluasi, sebagaimana diilustrasikan pada Gambar 1.



**Gambar 1.** Rancangan Penelitian

### 1. Pengumpulan Data

Data penelitian dikumpulkan melalui formulir digital Monitoring dan Evaluasi yang diisi oleh peserta pada akhir sesi pelatihan. Dataset yang berhasil dihimpun berjumlah 2.313 ulasan tidak terstruktur yang mencakup empat aspek evaluasi, yaitu fasilitas, makanan, pengajar, dan penguji. Karakteristik data yang bersifat naratif dan mengandung ragam bahasa informal menjadi tantangan utama yang diatasi dalam penelitian ini [2]. Kompleksitas data semacam ini, yang seringkali memuat singkatan dan ekspresi subjektif, memerlukan pendekatan pemrosesan otomatis yang cermat untuk mengekstraksi makna secara efektif [6]. Pengumpulan data ini merupakan bagian integral dari evaluasi program untuk memastikan relevansi pelatihan dengan kebutuhan industri [1].

### 2. Pra-pemrosesan Data

Tahap pra-pemrosesan bertujuan untuk mengubah data mentah menjadi format yang bersih dan terstruktur agar siap diolah oleh model komputasional. Proses ini sangat krusial untuk meningkatkan performa klasifikasi dan akurasi pelabelan sentimen [7]. Tahapan yang dilakukan secara berurutan meliputi: (1) *Data cleaning* untuk menghapus karakter non-alfabetik, angka, dan simbol; (2) *Case folding* untuk mengonversi seluruh teks menjadi huruf kecil; (3) Normalisasi untuk mengubah kata tidak baku ke bentuk standar; (4) Tokenisasi untuk memecah kalimat menjadi unit kata; (5) *Stopword removal* untuk menghapus kata umum yang tidak memiliki bobot sentimen; dan (6) *Stemming* untuk mereduksi kata berimbuhan ke bentuk dasarnya. Untuk proses stemming Bahasa Indonesia, digunakan stemmer Sastrawi guna memastikan akurasi pengembalian kata ke bentuk dasarnya [8]. Data yang telah terstandarisasi ini akan menentukan efektivitas proses analisis selanjutnya sebelum pemodelan dengan SVM [9].

### 3. Pelabelan Sentimen

Setelah pra-pemrosesan, dilakukan pelabelan sentimen awal secara otomatis menggunakan pendekatan *Lexicon-Based*. Penelitian ini memanfaatkan *InSet Lexicon*, sebuah kamus sentimen yang dirancang khusus untuk Bahasa Indonesia [10]. Proses ini bekerja dengan menghitung skor sentimen untuk setiap ulasan berdasarkan jumlah kata positif dan negatif yang ditemukan. Pendekatan ini memungkinkan pelabelan yang cepat tanpa memerlukan data latih manual yang besar [11].

### 4. Ekstraksi Fitur TF-IDF

Data teks yang telah bersih selanjutnya dikonversi menjadi representasi numerik melalui proses vektorisasi menggunakan metode *Term Frequency-Inverse Document Frequency*. Metode TF-IDF

dipilih karena kemampuannya dalam memberikan bobot yang merefleksikan tingkat kepentingan sebuah kata dalam dokumen relatif terhadap keseluruhan korpus, yang sangat efektif untuk input model SVM [12].

## 5. Pemodelan dan Klasifikasi Sentimen

Penelitian ini membangun dan membandingkan dua model klasifikasi. Pertama, model dasar SVM dengan fungsi RBF dibangun sebagai tolok ukur. Kernel RBF dipilih karena fleksibilitasnya dalam menangani hubungan data yang non-linear [13]. Kedua, model SVM dioptimasi menggunakan algoritma PSO. PSO diterapkan untuk mencari kombinasi *hyperparameter*  $C$  dan  $\gamma$  yang optimal secara sistematis, dengan tujuan memaksimalkan akurasi klasifikasi [14]. Pendekatan ini bertujuan untuk mengatasi kelemahan penentuan parameter secara manual dan menemukan konfigurasi model yang lebih baik [15]. Untuk kedua skenario pemodelan, dataset dibagi menggunakan metode *hold-out* dengan proporsi 80% untuk data latih dan 20% untuk data uji.

## 6. Teknik Analisis Data

Analisis data dilakukan untuk mengevaluasi dan membandingkan kinerja kedua model. Performa model diukur menggunakan empat metrik evaluasi standar yang dihitung dari *Confusion Matrix*. Penggunaan metrik ini telah terbukti efektif dalam memberikan gambaran komprehensif mengenai efektivitas sebuah model klasifikasi [16]. Analisis komparatif dilakukan dengan membandingkan nilai *weighted average* dari setiap metrik antara model SVM standar dan model SVM yang dioptimasi PSO. Pendekatan ini memungkinkan penilaian yang adil terhadap dampak optimasi, terutama pada dataset dengan potensi kelas yang tidak seimbang [17]. Hasil analisis ini digunakan untuk menjawab rumusan masalah penelitian

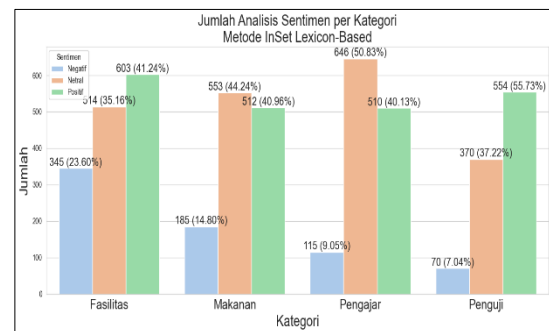
mengenai model mana yang menunjukkan performa lebih unggul [18].

## C. HASIL DAN PEMBAHASAN

Bagian ini menyajikan hasil analisis data secara bertahap, mulai dari distribusi sentimen awal hingga perbandingan kinerja model klasifikasi. Setiap temuan diinterpretasikan untuk menjawab rumusan masalah dan menunjukkan kontribusi penelitian.

### 1. Distribusi Sentimen

Analisis awal dilakukan dengan pelabelan sentimen berbasis leksikon menggunakan InSet Lexicon untuk memetakan persepsi umum peserta terhadap empat aspek utama pelatihan. Distribusi sentimen yang dihasilkan disajikan pada Gambar 2.



**Gambar 2.** Distribusi Sentimen

Hasil visualisasi menunjukkan perbedaan persepsi yang signifikan antar kategori. Aspek Ruangan dan Fasilitas mencatatkan jumlah sentimen negatif tertinggi (345 ulasan; 23,60%), mengindikasikan bahwa ini adalah area yang paling banyak dikeluhkan dan memerlukan perhatian utama. Sebaliknya, aspek Penguji menunjukkan profil sentimen paling positif, dengan 554 ulasan (55,73%) bernada positif dan hanya 70 ulasan (7,04%) yang negatif. Temuan ini menggarisbawahi bahwa profesionalisme dan kompetensi penguji merupakan kekuatan utama dari program pelatihan ini.

Aspek Makanan menunjukkan sentimen yang lebih terpolarisasi, dengan jumlah ulasan netral yang tinggi (44,24%) namun juga diiringi sentimen negatif yang cukup signifikan (14,80%), lebih tinggi dibandingkan aspek pengajar dan penguji. Sementara itu, aspek Pengajar didominasi oleh sentimen netral (50,83%) dan positif (40,13%), yang menyiratkan bahwa kualitas pengajaran dinilai baik, meskipun banyak umpan balik yang bersifat deskriptif daripada ekspresif. Distribusi yang tidak seimbang ini, terutama pada jumlah data negatif yang sedikit untuk kategori Pengajar dan Penguji, menjadi tantangan utama bagi model klasifikasi.

## 2. Kinerja Model Klasifikasi SVM Kernel RBF

Model SVM dengan kernel RBF standar diimplementasikan sebagai tolok ukur untuk mengevaluasi kemampuannya dalam mengklasifikasikan sentimen pada setiap kategori. Kinerja model dievaluasi menggunakan metrik akurasi, presisi, recall, dan F1-score, yang dirangkum pada Tabel 1.

**Tabel 1.** Hasil Kinerja Model SVM Kernel RBF

| Kategori  | Akurasi | Presisi | Recall | F1-Score |
|-----------|---------|---------|--------|----------|
| Fasilitas | 85,67 % | 0.86    | 0.86   | 0.86     |
| Makanan   | 84,40 % | 0.85    | 0.84   | 0.84     |
| Pengajar  | 84,71 % | 0.85    | 0.85   | 0.84     |
| Penguji   | 87,44 % | 0.87    | 0.87   | 0.87     |

Secara umum, model SVM standar menunjukkan kinerja yang baik dengan akurasi di atas 84% untuk semua kategori. Namun, analisis lebih mendalam pada *classification report* mengungkapkan kelemahan signifikan, yaitu nilai recall

yang sangat rendah untuk kelas negatif pada kategori Pengajar (0.30) dan Penguji (0.46). Hal ini mengindikasikan bahwa model kesulitan mengidentifikasi sentimen negatif yang merupakan kelas minoritas. Kegagalan ini dapat diatribusikan pada ketidakseimbangan data dan kemungkinan ekspresi kritik yang lebih halus, sehingga model cenderung memprediksi kelas mayoritas yaitu kelas positif dan netral. Temuan ini sejalan dengan tantangan yang diidentifikasi oleh [16] terkait penanganan dataset yang tidak seimbang.

## 3. Kinerja Model SVM dengan Optimasi PSO

Untuk mengatasi keterbatasan model standar, diterapkan optimasi hyperparameter menggunakan PSO. PSO bertugas mencari kombinasi parameter C dan gamma yang optimal untuk memaksimalkan akurasi klasifikasi. Hasil pencarian parameter optimal disajikan pada Tabel 2.

**Tabel 2.** Parameter Optimal Hasil Pencarian PSO

| Kategori  | Parameter C | Parameter Gamma |
|-----------|-------------|-----------------|
| Fasilitas | 696.3381    | 0.0125          |
| Makanan   | 915.1470    | 0.0570          |
| Pengajar  | 309.7231    | 0.2062          |
| Penguji   | 563.4800    | 0.4049          |

Dengan menggunakan parameter optimal tersebut, model SVM dilatih ulang dan dievaluasi. Hasil kinerjanya menunjukkan peningkatan yang konsisten di seluruh kategori, sebagaimana dirangkum pada Tabel 3.

**Tabel 3.** Hasil Kinerja Model SVM dengan Optimasi PSO

| Kategori  | Akurasi | Presisi | Recall | F1-Score |
|-----------|---------|---------|--------|----------|
| Fasilitas | 85,67 % | 0.86    | 0.86   | 0.86     |
| Makanan   | 84,40 % | 0.85    | 0.84   | 0.84     |
| Pengajar  | 84,71 % | 0.85    | 0.85   | 0.84     |
| Penguji   | 87,44 % | 0.87    | 0.87   | 0.87     |

|           |            |          |          |      |
|-----------|------------|----------|----------|------|
| Fasilitas | 88,05<br>% | 0.8<br>8 | 0.8<br>8 | 0.88 |
| Makanan   | 88,00<br>% | 0.8<br>8 | 0.8<br>8 | 0.88 |
| Pengajar  | 89,02<br>% | 0.8<br>9 | 0.8<br>9 | 0.89 |
| Penguji   | 91,46<br>% | 0.9<br>1 | 0.9<br>1 | 0.91 |

Model yang telah dioptimasi berhasil mencapai akurasi yang lebih tinggi, dengan peningkatan paling menonjol pada kategori Pengajar dan Penguji. Peningkatan ini mengindikasikan bahwa penentuan *hyperparameter* yang tepat secara fundamental memperbaiki kemampuan generalisasi model.

#### 4. Analisis Komparatif dan Implikasi Hasil

Perbandingan langsung antara model SVM standar dan model yang dioptimasi PSO menunjukkan keunggulan signifikan dari pendekatan optimasi. Tabel 4 menyajikan perbandingan kinerja kedua model berdasarkan metrik F1-Score yang menyeimbangkan presisi dan recall.

**Tabel 4.** Perbandingan Kinerja Model

| Kategori  | Mode<br>l | Akurasi | F1-<br>Score<br>Weigh<br>ted |
|-----------|-----------|---------|------------------------------|
| Fasilitas | SVM       | 85,67%  | 0.86                         |
|           | SVM       | 88,05%  | 0.88                         |
|           | +<br>PSO  |         |                              |
| Makanan   | SVM       | 84,40%  | 0.84                         |
|           | SVM       | 88,00%  | 0.88                         |
|           | +<br>PSO  |         |                              |
| Pengajar  | SVM       | 84,71%  | 0.84                         |
|           | SVM       | 89,02%  | 0.89                         |
|           | +<br>PSO  |         |                              |
| Penguji   | SVM       | 87,87%  | 0.87                         |

|     |        |      |
|-----|--------|------|
| SVM | 91,46% | 0.91 |
| +   |        |      |
| PSO |        |      |

Penerapan PSO secara konsisten meningkatkan performa di seluruh metrik dan kategori. Peningkatan paling substansial terjadi pada kategori Pengajar (kenaikan akurasi 4,31 poin) dan Penguji (kenaikan 4,02 poin). Hal ini sangat penting karena peningkatan tersebut diatribusikan pada kemampuan PSO dalam menemukan hyperplane yang lebih sensitif terhadap kelas minoritas. Nilai recall untuk kelas negatif pada kategori Pengajar meningkat drastis dari 0.30 menjadi 0.70 setelah optimasi. Ini membuktikan bahwa optimasi PSO berhasil mengatasi kelemahan utama model standar.

Temuan ini mengonfirmasi bahwa optimasi parameter merupakan langkah krusial untuk mencapai performa maksimal, sejalan dengan penelitian oleh [14] dan [15]. Secara akademik, penelitian ini berkontribusi dengan menunjukkan bahwa pendekatan hybrid Lexicon-SVM yang diperkuat dengan optimasi PSO mampu menciptakan model klasifikasi yang tidak hanya akurat tetapi juga seimbang dan andal. Implikasinya, model ini dapat berfungsi sebagai alat evaluasi otomatis yang objektif, memungkinkan institusi untuk mengolah umpan balik secara efisien dan membuat keputusan berbasis data untuk peningkatan mutu program secara berkelanjutan.

#### D. KESIMPULAN DAN SARAN

Penelitian ini berhasil menunjukkan bahwa implementasi model *hybrid Lexicon-Based* dan SVM yang dioptimasi menggunakan *Particle Swarm Optimization* secara signifikan lebih unggul dibandingkan model SVM standar untuk klasifikasi sentimen pada umpan balik peserta pelatihan. Model yang

dioptimasi mampu meningkatkan akurasi klasifikasi di seluruh kategori, dengan peningkatan tertinggi pada aspek Pengajar (89,02%) dan Penguji (91,46%). Temuan utama menjawab rumusan masalah dengan membuktikan bahwa optimasi PSO efektif mengatasi kelemahan model standar dalam mengidentifikasi sentimen negatif yang merupakan kelas minoritas, yang tercermin dari peningkatan nilai recall secara substansial. Secara praktis, penelitian ini menghasilkan sebuah model evaluasi otomatis yang objektif dan andal, yang dapat dimanfaatkan oleh BBPSDMP Kominfo Makassar untuk mendukung pengambilan keputusan berbasis data. Secara teoritis, hasil ini memperkuat relevansi penerapan metode optimasi metaheuristik untuk meningkatkan ketahanan dan akurasi model klasifikasi pada data teks tidak terstruktur yang tidak seimbang.

Meskipun demikian, penelitian ini memiliki keterbatasan, antara lain fokus data yang hanya berasal dari satu institusi dan periode waktu tertentu, serta penggunaan kamus sentimen yang bersifat umum. Oleh karena itu, penelitian selanjutnya dapat diarahkan pada beberapa pengembangan. Pertama, menerapkan model ini pada dataset dari program pelatihan lain untuk menguji generalisasinya. Kedua, mengembangkan kamus sentimen yang spesifik untuk domain evaluasi pendidikan guna meningkatkan akurasi pelabelan. Ketiga, melakukan analisis komparatif dengan algoritma deep learning seperti BERT atau LSTM untuk mengeksplorasi potensi peningkatan performa lebih lanjut. Terakhir, mengimplementasikan *Aspect-Based Sentiment Analysis* (ABSA) untuk memperoleh wawasan yang lebih granular mengenai aspek spesifik yang dikeluhkan atau dipuji oleh peserta.

## E. REFERENSI

- [1] Kementerian Komunikasi dan Informatika Republik Indonesia. (2021). Peluncuran program literasi digital nasional. Badan Pengembangan SDM Kominfo. <https://bpsdm.komdigi.go.id/berita-peluncuran-program-literasi-digital-nasional-44-777>
- [2] Syahputra, H. (2021). Sentiment analysis of community opinion on online store in Indonesia on Twitter using support vector machine algorithm (SVM). *Journal of Physics: Conference Series*, 1819(1), 012030. <https://doi.org/10.1088/1742-6596/1819/1/012030>
- [3] Nurmallasari, D., Hidayatul, D., Authors, Q., Chairani, N., & Yuliantoro, H. R. (2024). Discovering user sentiment patterns in libraries with a hybrid machine learning and lexicon-based approach. *International Journal* (Vol. 13, pp. 311–317).
- [4] Savira, R., Solichin, A., & Syafrullah, M. (2023). Analisis sentimen pada Twitter terhadap kenaikan BBM 2022 dengan Lexicon dan support vector machine. *Prosiding Seminar Nasional Mahasiswa Fakultas Teknologi Informasi (SENAFTI)*, 2(1), 211–218. <https://senafti.budiluhur.ac.id/index.php/senafti/article/view/564>
- [5] Pratama, M. R., Ramadhan, Y. R., & Komara, M. A. (2023). Analisis sentimen BRI mo dan BCA mobile

- menggunakan support vector machine dan lexicon based. JUTISI: Jurnal Ilmiah Teknik Informatika dan Sistem Informasi, 12(3), 1439–1450.
- [6] Al-Khowarizmi, I. P., Sari, S., & Maulana, H. (2023). Detecting cyberbullying on social media using support vector machine: A case study on Twitter. *International Journal of Safety and Security Engineering*, 13(4), 709–714.  
<https://doi.org/10.18280/ijssse.130413>
- [7] Ulya, D. I., Kunaefi, A., Rolliawati, D., & Nugroho, B. A. (2023). Unpacking public perceptions of QRIS with Twitter data: A VADER and LDA methodology. *Jurnal ELTIKOM*, 7(2), 145–159.  
<https://doi.org/10.31961/eltikom.v7i2.742>
- [8] Ependi, U., & Ahmad, N. A. (2024). A novel hybrid classification on urban opinion using ROS-RF: A machine learning approach. *Jurnal Penelitian Pendidikan IPA*, 10(8), 5816–5824.  
<https://doi.org/10.29303/jppipa.v10i8.8042>
- [9] Kristiyanti, D. A., & Hardani, S. (2023). Sentiment analysis of public acceptance of Covid-19 vaccines types in Indonesia using naïve Bayes, support vector machine, and long short-term memory (LSTM). *Jurnal RESTI* (Rekayasa Sistem dan Teknologi Informasi), 7(3), 722–732.  
<https://doi.org/10.29207/resti.v7i3.4737>
- [10] Muttakin, F., & Andrika, N. (2025). Sentiment analysis of shoe product reviews on Indonesian e-commerce platform using lexicon based and support vector machine. *International Journal* (Vol. 6, No. 2, pp. 839–854).
- [11] Artana, I. K. A. B., Pradnyana, G. A., & Darmawiguna, I. G. M. (2023). Analisis sentimen Twitter untuk menilai kesiapan pembelajaran tatap muka terbatas dengan InSet lexicon dan Levenshtein distance. *Jurnal Pendidikan Teknologi dan Kejuruan*, 20(2), 200–209.  
<https://doi.org/10.23887/jptkundiksha.v20i2.64579>
- [12] Arifin, N., Enri, U., & Sulistiyowati, N. (2021). Penerapan algoritma support vector machine (SVM) dengan TF-IDF N-Gram untuk text classification. *STRING (Satuan Tulisan Riset dan Inovasi Teknologi)*, 6(2), 129.  
<https://doi.org/10.30998/string.v6i2.10133>
- [13] Nauli, S., Berutu, S. S., Budiati, H., & Maedjaja, F. (2025). Klasifikasi kalimat perundungan pada Twitter menggunakan algoritma support vector machine. *JIPI (Jurnal Ilmiah Penelitian dan Pembelajaran Informatika)*, 10(1), 107–122.

- <https://doi.org/10.29100/jipi.v10i1.5749>
- [14] Darmawan, R., Indra, I., & Surahmat, A. (2022). Optimalisasi support vector machine (SVM) berbasis particle swarm optimization (PSO) pada analisis sentimen terhadap official account Ruang Guru di Twitter. *Jurnal Kajian Ilmiah*, 22(2), 143–152. <https://doi.org/10.31599/jki.v22i2.1130>
- [15] Risawati, R., Ernawati, S., & Maryani, I. (2020). Optimasi parameter PSO berbasis SVM untuk analisis sentimen review jasa maskapai penerbangan berbahasa Inggris. *EVOLUSI: Jurnal Sains dan Manajemen*, 8(2), 64–71. <https://doi.org/10.31294/evolusi.v8i2.9248>
- [16] Muhammadi, R. H., Laksana, T. G., & Arifa, A. B. (2022). Combination of support vector machine and lexicon-based algorithm in Twitter sentiment analysis. *Khazanah Informatika: Jurnal Ilmu Komputer dan Informatika*, 8(1), 59–71. <https://doi.org/10.23917/khif.v8i1.15213>
- [17] Alfauzan, M. F., Sibaroni, Y., & Fitriyani, F. (2023). Sentiment classification of fuel price rise in economic aspects using lexicon and SVM method. *Sinkron*, 8(4), 2526–2536. <https://doi.org/10.33395/sinkron.v8i4.12851>
- [18] Oktaviana, N. E., Sari, Y. A., & Indriati, I. (2022). Analisis sentimen terhadap kebijakan kuliah daring selama pandemi menggunakan pendekatan lexicon based features dan support vector machine. *Jurnal Teknologi Informasi dan Ilmu Komputer*, 9(2), 357–362. <https://doi.org/10.25126/jtiik.2022925625>