

OPTIMALISASI MODEL KLASIFIKASI MENGGUNAKAN SUPPORT VECTOR MACHINE DENGAN SMOTE PADA DATA SENTIMEN ULASAN APLIKASI BY.U

Wandi¹, Martanto², Fatihanurasi Dikananda³, Ratna Dewi Lestiyorini⁴.

Program Studi Teknik Informatika¹³
Program Studi Manajemen Informatika²
Program Studi Sistem Informasi⁴

STMIK IKMI Cirebon
<https://ikmi.ac.id/page/18/?lang=de>
wandiw720@gmail.com

(*) Corresponding Author : wandiw720@gmail.com
Published : 30 Juni 2026

Abstract—This study aims to analyze user sentiment toward the by.U application using the Support Vector Machine (SVM) algorithm and to optimize model performance through the Synthetic Minority Oversampling Technique (SMOTE). The dataset consists of user reviews that have undergone preprocessing and labeling into three sentiment classes: negative, neutral, and positive. The main issue addressed in this study is the imbalance in data distribution, which can affect the performance of classification models. To overcome this issue, SMOTE was applied to balance the number of samples in each class. The dataset was then split into training and testing sets, followed by model training and evaluation using the SVM algorithm. Model performance was evaluated using a confusion matrix, classification report, and ROC Curve with Area Under the Curve (AUC) metrics. The results show that the SVM model without SMOTE achieved an accuracy of 82.39%, while the SVM model with SMOTE achieved an accuracy of 80.55%. Although there is a slight decrease in accuracy, the application of SMOTE improves the balance of model performance, particularly for minority classes. The AUC values range from 0.95 to 0.96, indicating that the model has excellent classification capability. In conclusion, SMOTE is effective in handling imbalanced datasets and improving overall classification performance, although it does not always increase accuracy. This study is expected to serve as a reference for sentiment analysis using machine learning techniques on imbalanced data.

Keywords: Sentiment Analysis, SVM, SMOTE, Classification, By.U

Abstrak—Penelitian ini bertujuan untuk menganalisis sentimen ulasan pengguna aplikasi by.U menggunakan algoritma Support Vector Machine (SVM) serta mengoptimalkan performa model dengan metode Synthetic Minority Oversampling Technique (SMOTE). Data yang digunakan merupakan ulasan pengguna yang telah melalui tahap preprocessing dan pelabelan menjadi tiga kelas sentimen, yaitu negatif, netral, dan positif. Permasalahan utama dalam penelitian ini adalah ketidakseimbangan distribusi data yang dapat mempengaruhi kinerja model klasifikasi. Untuk mengatasi hal tersebut, dilakukan penerapan SMOTE guna menyeimbangkan jumlah data pada setiap kelas. Selanjutnya, data dibagi menjadi data latih dan data uji dengan proporsi tertentu, kemudian dilakukan proses pelatihan dan pengujian menggunakan algoritma SVM. Evaluasi model dilakukan menggunakan confusion matrix, classification report, serta ROC Curve dan nilai Area Under Curve (AUC). Hasil penelitian menunjukkan bahwa model SVM tanpa SMOTE memperoleh akurasi sebesar 82,39%, sedangkan model SVM dengan SMOTE menghasilkan akurasi sebesar 80,55%. Meskipun terjadi sedikit penurunan akurasi, penerapan SMOTE mampu meningkatkan keseimbangan performa model, khususnya pada kelas minoritas. Nilai AUC yang diperoleh berada pada rentang 0,95–0,96 yang menunjukkan bahwa model memiliki kemampuan klasifikasi yang sangat baik. Berdasarkan hasil tersebut, dapat disimpulkan bahwa metode SMOTE efektif dalam menangani permasalahan data tidak seimbang dan meningkatkan kualitas klasifikasi secara keseluruhan, meskipun tidak selalu meningkatkan nilai akurasi. Penelitian ini diharapkan dapat menjadi referensi dalam pengembangan analisis sentimen berbasis machine learning pada data yang tidak seimbang.

Kata Kunci: Analisis Sentimen, SVM, SMOTE, Klasifikasi, by.U

INTRODUCTION

Perkembangan teknologi informasi dan komunikasi yang pesat telah mendorong meningkatnya penggunaan aplikasi berbasis mobile, termasuk aplikasi layanan telekomunikasi seperti by.U. Aplikasi ini memungkinkan pengguna untuk mengelola layanan secara mandiri, sehingga kualitas layanan sangat bergantung pada pengalaman pengguna yang tercermin dalam ulasan di platform digital seperti Google Play Store. Ulasan pengguna tersebut menjadi sumber data penting untuk mengevaluasi tingkat kepuasan dan kualitas layanan secara berkelanjutan [1].

Analisis sentimen merupakan salah satu pendekatan dalam text mining yang digunakan untuk mengidentifikasi opini pengguna menjadi kategori positif, negatif, atau netral. Dalam beberapa tahun terakhir, metode machine learning seperti Support Vector Machine (SVM) banyak digunakan dalam klasifikasi teks karena kemampuannya dalam menangani data berdimensi tinggi dan menghasilkan performa yang baik [2]. Namun, dalam implementasinya, data ulasan pengguna seringkali memiliki distribusi kelas yang tidak seimbang (imbalanced data), di mana jumlah data pada satu kelas jauh lebih dominan dibandingkan kelas lainnya.

Permasalahan ketidakseimbangan data ini dapat menyebabkan model klasifikasi menjadi bias terhadap kelas mayoritas dan menurunkan performa dalam mendeteksi kelas minoritas [3]. Oleh karena itu, diperlukan teknik penanganan imbalance seperti Synthetic Minority Oversampling Technique (SMOTE), yang mampu menghasilkan data sintetis pada kelas minoritas untuk meningkatkan keseimbangan distribusi data [4]. Penerapan SMOTE telah terbukti dapat meningkatkan akurasi dan performa model klasifikasi pada berbagai studi machine learning.

Penelitian ini bertujuan untuk mengoptimalkan model klasifikasi sentimen pada ulasan pengguna aplikasi by.U dengan menggunakan algoritma Support Vector Machine (SVM) yang dikombinasikan dengan teknik Synthetic Minority Oversampling Technique (SMOTE). Tujuan utama dari penelitian ini adalah meningkatkan performa model klasifikasi, khususnya dalam menangani permasalahan ketidakseimbangan data yang dapat mempengaruhi hasil prediksi. Selain itu, penelitian ini juga bertujuan untuk menghasilkan model yang lebih akurat dan

robust dalam mengidentifikasi sentimen positif dan negatif dari ulasan pengguna. Signifikansi penelitian ini terletak pada kontribusinya dalam mengisi kesenjangan penelitian sebelumnya yang umumnya belum mengintegrasikan metode SVM dengan teknik penanganan data imbalance secara optimal pada domain aplikasi by.U. Secara praktis, hasil penelitian ini diharapkan dapat memberikan insight yang lebih akurat bagi pengembang aplikasi dalam meningkatkan kualitas layanan serta pengalaman pengguna berbasis data ulasan digital [5].

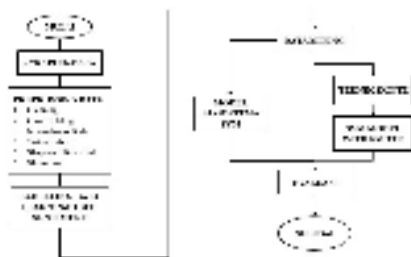
Dalam penelitian ini, pendekatan yang digunakan adalah metode klasifikasi berbasis machine learning untuk menganalisis sentimen ulasan pengguna aplikasi by.U. Algoritma yang diterapkan adalah Support Vector Machine (SVM) karena memiliki kemampuan yang baik dalam menangani data teks berdimensi tinggi serta menghasilkan performa klasifikasi yang optimal. Untuk mengatasi permasalahan ketidakseimbangan data yang sering terjadi pada dataset ulasan pengguna, penelitian ini mengintegrasikan teknik Synthetic Minority Oversampling Technique (SMOTE) guna meningkatkan representasi kelas minoritas. Proses penelitian meliputi tahapan pengumpulan data, preprocessing teks, pembobotan fitur menggunakan TF-IDF, serta proses pelatihan dan pengujian model klasifikasi. Evaluasi kinerja model dilakukan menggunakan metrik seperti akurasi, precision, recall, dan F1-score untuk mengukur efektivitas model yang diusulkan [6].

Hasil dari penelitian ini diharapkan dapat memberikan kontribusi yang signifikan dalam pengembangan metode analisis sentimen berbasis machine learning, khususnya dalam penanganan data yang tidak seimbang. Dengan berhasilnya penerapan algoritma Support Vector Machine yang dioptimalkan menggunakan teknik SMOTE, penelitian ini dapat memperkaya pemahaman mengenai efektivitas kombinasi metode tersebut dalam meningkatkan performa klasifikasi teks. Selain itu, temuan penelitian ini diharapkan dapat menjadi referensi bagi peneliti selanjutnya dalam mengembangkan model klasifikasi yang lebih akurat dan robust pada berbagai domain data teks. Bagi praktisi, khususnya pengembang aplikasi digital seperti by.U, hasil penelitian ini dapat dimanfaatkan untuk memperoleh insight yang lebih mendalam terkait persepsi dan kepuasan pengguna berdasarkan ulasan yang diberikan. Informasi tersebut dapat digunakan sebagai dasar dalam pengambilan keputusan strategis untuk meningkatkan kualitas layanan dan pengalaman pengguna. Di sisi lain, penelitian ini juga berpotensi

mendorong pemanfaatan teknik penanganan data imbalance secara lebih luas dalam berbagai studi machine learning di bidang informatika. Dengan demikian, kontribusi penelitian ini tidak hanya bersifat teoritis, tetapi juga memiliki dampak praktis dalam pengembangan teknologi berbasis data dan peningkatan kualitas layanan digital secara berkelanjutan.

MATERIALS AND METHODS

Penelitian ini menggunakan pendekatan kuantitatif dengan desain eksperimental, yang dirancang untuk mengevaluasi secara sistematis kinerja model klasifikasi teks melalui manipulasi variabel bebas dan pengukuran terkontrol terhadap variabel terikat. Pendekatan kuantitatif dipilih karena mampu memberikan hasil yang objektif, terukur, dan dapat diuji secara statistik, sehingga memungkinkan peneliti menilai secara empiris pengaruh metode praproses, teknik penyeimbangan data, serta algoritma klasifikasi terhadap performa analisis sentimen [7]. Lalu daripada itu desain penelitian di gambarkan pada visual berikut ini :



Gambar 1 Alur Penelitian

Dalam desain eksperimental ini, variabel bebas yang diuji mencakup jenis representasi fitur (misalnya TF-IDF), metode klasifikasi (SVM), dan teknik penyeimbangan data (SMOTE). Sementara itu, variabel terikatnya adalah metrik kinerja model klasifikasi, yang meliputi akurasi, recall, presisi, dan F1-score. Melalui manipulasi terencana terhadap variabel bebas tersebut, penelitian ini bertujuan untuk mengidentifikasi hubungan kausal antara metode yang digunakan dan hasil klasifikasi sentimen terhadap dataset ulasan aplikasi By.U.

Desain eksperimental ini mendukung proses evaluasi yang reproducible dan comparable, dengan cara menstandarkan pembagian dataset (train-validation-test split), tahapan praproses, serta prosedur evaluasi model [8]. Dengan standarisasi tersebut, hasil penelitian dapat dibandingkan secara langsung dengan studi lain dan diulang kembali untuk verifikasi. Selain itu, penerapan uji signifikansi statistik serta analisis kesalahan (error analysis) dilakukan untuk memastikan bahwa perbedaan

performa antar model bukan disebabkan oleh faktor acak, melainkan oleh efek nyata dari perlakuan eksperimental [9].

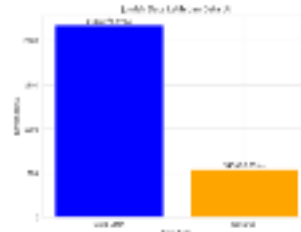
Menjamin validitas eksperimental merupakan aspek penting dalam penelitian ini. Validitas internal dan eksternal dipertahankan melalui dokumentasi metodologis yang rinci, mencakup komposisi dataset, prosedur pembagian data, tahapan praproses teks, parameterisasi model SVM (termasuk pemilihan kernel), dan penerapan SMOTE [10], [11]. Kontrol terhadap random seed diterapkan untuk menjaga konsistensi hasil dan mengurangi bias acak dalam pelatihan model.

Selanjutnya, validitas konstruk dan reliabilitas hasil diperkuat dengan pengulangan percobaan beberapa kali (multiple runs) serta pelaporan hasil dalam bentuk rata-rata dan varians, bukan sekadar satu kali hasil eksperimen. Praktik pelaporan terbuka, termasuk penyediaan kode sumber dan dataset penelitian, mengikuti rekomendasi dari studi replikasi independen dalam bidang NLP yang menekankan transparansi dan traceability [12], [13].

RESULTS AND DISCUSSION

Pada tahap ini, dilakukan proses pemodelan menggunakan algoritma Support Vector Machine (SVM) untuk mengklasifikasikan sentimen ulasan pengguna berdasarkan data yang telah melalui tahap praproses dan pelabelan. Sebelum proses pelatihan model, dataset dibagi menjadi data latih (training data) dan data uji (testing data) untuk mengevaluasi kinerja model secara objektif.

Data dibagi menjadi data latih (79,97%) dan data uji (20,03%), masing-masing sebanyak 2.176 dan 545 entri. Pembagian ini bertujuan agar model dapat mempelajari pola dari data latih dan kemudian diuji kemampuannya dalam mengklasifikasikan data baru yang belum pernah dilihat sebelumnya. Dapat dilihat pada gambar berikut:

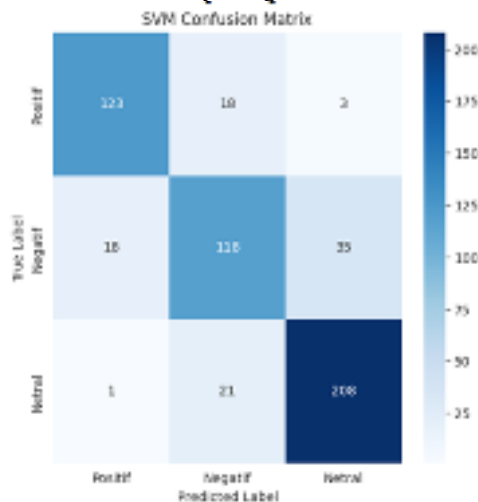


Gambar 2 Grafik pembagian data latih dan data uji

Representasi teks pada penelitian ini dilakukan menggunakan metode Term Frequency-Inverse Document Frequency (TF-IDF) untuk mengubah data teks menjadi bentuk vektor numerik yang dapat diproses oleh algoritma klasifikasi. Selanjutnya, model Support Vector

Machine (SVM) dengan kernel linear dilatih menggunakan data latih, kemudian dievaluasi menggunakan data uji untuk mengukur kinerja model dalam mengklasifikasikan sentimen.

Untuk memperoleh gambaran yang lebih rinci mengenai kinerja model dalam melakukan klasifikasi, digunakan confusion matrix sebagai alat evaluasi. Confusion matrix mampu menunjukkan perbandingan antara hasil prediksi model dengan label aktual pada setiap kelas sentimen. Hasil evaluasi tersebut kemudian disajikan dalam bentuk grafik confusion matrix pada gambar berikut.



Gambar 3. Confusion Matrik

Baris pada confusion matrix menunjukkan kelas aktual, sedangkan kolom menunjukkan kelas prediksi. Hasil tersebut menunjukkan bahwa:

Pada kelas negatif, dari total data yang ada, sebanyak 123 data berhasil diklasifikasikan dengan benar, sedangkan 18 data salah diklasifikasikan sebagai netral dan 3 data sebagai positif.

Pada kelas netral, model mampu mengklasifikasikan dengan benar sebanyak 118 data, namun terdapat 18 data salah diprediksi sebagai negatif dan 35 data sebagai positif.

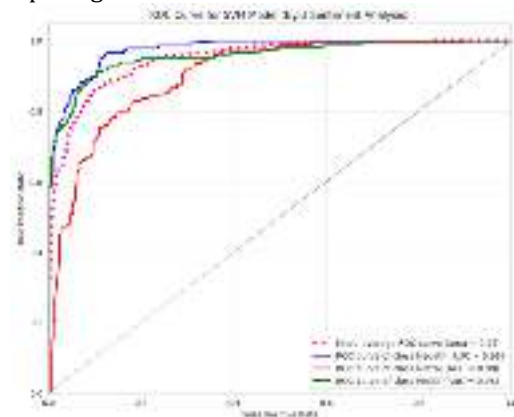
Pada kelas positif, performa model tergolong sangat baik dengan 208 data terklasifikasi dengan benar, sementara hanya 1 data salah diprediksi sebagai negatif dan 21 data sebagai netral.

Secara keseluruhan, nilai pada diagonal utama (123, 118, dan 208) menunjukkan jumlah prediksi yang benar untuk masing-masing kelas. Hal ini mengindikasikan bahwa model SVM memiliki kemampuan klasifikasi yang cukup baik, terutama pada kelas positif yang memiliki tingkat akurasi tertinggi.

Namun demikian, masih terdapat kesalahan klasifikasi yang cukup signifikan pada kelas netral, khususnya yang cenderung diprediksi sebagai kelas positif. Hal ini menunjukkan bahwa model mengalami kesulitan dalam membedakan sentimen netral dengan sentimen positif, yang kemungkinan disebabkan oleh kemiripan karakteristik kata atau distribusi data yang tidak seimbang.

Selain menggunakan confusion matrix, kinerja model juga dievaluasi menggunakan kurva Receiver Operating Characteristic (ROC). Kurva ROC digunakan untuk mengukur kemampuan model dalam membedakan antar kelas dengan membandingkan nilai True Positive Rate (TPR) dan False Positive Rate (FPR) pada berbagai ambang batas (threshold).

Melalui kurva ROC, dapat diketahui seberapa baik model dalam mengklasifikasikan data secara keseluruhan, yang ditunjukkan melalui nilai Area Under Curve (AUC). Semakin mendekati nilai 1, maka semakin baik kemampuan model dalam membedakan antar kelas sentimen. Hasil evaluasi tersebut kemudian disajikan dalam bentuk grafik ROC pada gambar berikut.



Gambar 4 Grafik ROC Curver for SVM

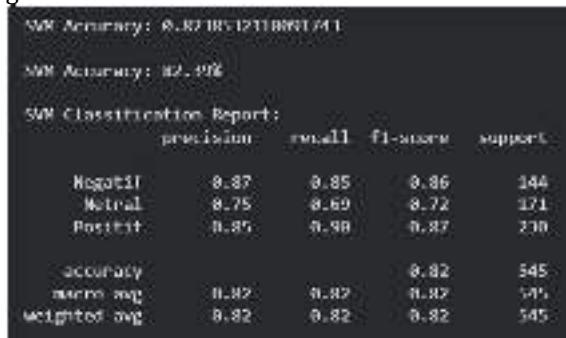
Berdasarkan grafik ROC yang ditampilkan, diperoleh nilai Area Under Curve (AUC) untuk setiap kelas, yaitu kelas negatif sebesar 0,98, kelas netral sebesar 0,90, dan kelas positif sebesar 0,96. Selain itu, nilai micro-average AUC sebesar 0,95 menunjukkan performa model secara keseluruhan.

Nilai AUC yang mendekati 1 menunjukkan bahwa model memiliki kemampuan klasifikasi yang sangat baik. Kelas negatif memiliki nilai AUC tertinggi, yang mengindikasikan bahwa model sangat efektif dalam mengidentifikasi ulasan negatif. Sementara itu, kelas netral memiliki nilai AUC paling rendah dibandingkan kelas lainnya, yang menunjukkan bahwa model masih mengalami kesulitan dalam membedakan sentimen netral dengan sentimen lainnya.

Secara keseluruhan, grafik ROC menunjukkan bahwa model SVM memiliki performa yang sangat

baik dalam klasifikasi sentimen, meskipun masih terdapat potensi peningkatan, khususnya dalam meningkatkan akurasi pada kelas netral.

Untuk mengukur kinerja model dalam mengklasifikasikan sentimen, dilakukan evaluasi menggunakan metrik *accuracy*, *precision*, *recall*, dan *f1-score*. Evaluasi ini bertujuan untuk mengetahui sejauh mana model Support Vector Machine (SVM) mampu mengklasifikasikan data uji secara tepat pada masing-masing kelas sentimen. Bisa dilihat pada gambar berikut :



SVM Classification Report:				
	precision	recall	f1-score	support
Negatif	0.87	0.85	0.86	144
Netral	0.75	0.69	0.72	121
Positif	0.85	0.98	0.92	210
accuracy			0.82	545
macro avg	0.82	0.82	0.82	545
weighted avg	0.82	0.82	0.82	545

Gambar 5 Hasil Accuracy Model SVM

Berdasarkan hasil evaluasi yang ditunjukkan pada gambar di atas, diperoleh nilai akurasi model SVM sebesar 82,39%, yang menunjukkan bahwa model mampu mengklasifikasikan sebagian besar data uji dengan baik dari total 545 data.

CONCLUSION

Berdasarkan hasil penelitian yang telah dilakukan mengenai analisis sentimen ulasan aplikasi by.U menggunakan algoritma Support Vector Machine (SVM) serta optimasi dengan metode SMOTE (Synthetic Minority Oversampling Technique), dapat disimpulkan beberapa hal sebagai berikut:

Model SVM tanpa penanganan ketidakseimbangan data mampu menghasilkan akurasi sebesar 82,39%, dengan performa yang cukup baik terutama pada kelas positif dan negatif. Namun, masih terdapat kelemahan dalam mengklasifikasikan kelas netral yang memiliki nilai recall lebih rendah.

Penerapan metode SMOTE berhasil menyeimbangkan distribusi data antar kelas, di mana setiap kelas memiliki jumlah data yang sama (907 data). Hal ini bertujuan untuk mengurangi bias model terhadap kelas mayoritas.

Model SVM dengan SMOTE menghasilkan akurasi sebesar 80,55%, sedikit lebih rendah dibandingkan SVM tanpa SMOTE. Meskipun

demikian, performa klasifikasi menjadi lebih seimbang antar kelas, khususnya pada kelas netral yang mengalami peningkatan recall.

Berdasarkan evaluasi menggunakan confusion matrix, model SVM dengan SMOTE menunjukkan distribusi prediksi yang lebih merata dibandingkan model tanpa SMOTE, sehingga mampu meningkatkan kemampuan generalisasi model terhadap seluruh kelas.

Hasil evaluasi menggunakan ROC Curve menunjukkan bahwa kedua model memiliki performa yang sangat baik, dengan nilai AUC mencapai 0,95–0,96, yang mengindikasikan bahwa model memiliki kemampuan tinggi dalam membedakan setiap kelas sentimen.

REFERENCE

- [1] M. P. Utami, R. G. Guntara, and B. M. Purwaamijaya, "Analisis Tingkat Kepuasan Pengguna Aplikasi By . U Menggunakan Metode EUCS dan IPA," vol. 4, no. 2, pp. 47–59, 2024.
- [2] E. Hokijuliandy, H. Napitupulu, and Firdaniza, "Application of SVM and chi-square feature selection for sentiment analysis of Indonesia's National Health Insurance Mobile Application," *Mathematics*, vol. 11, no. 17, p. 3765, 2023, doi: <https://doi.org/10.3390/math11173765>.
- [3] K. Aen, F. A. Tyas, and H. Rakhmawati, "OPTIMASI ALGORITMA SUPPORT VECTOR MACHINE (SVM) MENGGUNAKAN TEKNIK SMOTE UNTUK ANALISIS SENTIMEN PUBLIK TERHADAP DAMPAK PENGGUNAAN ARTIFICIAL INTELLIGENCE (AI)," vol. 9, no. 5, pp. 8287–8295, 2025.
- [4] D. Elreedy, A. F. Atiya, and F. Kamalov, "A theoretical distribution analysis of synthetic minority oversampling technique (SMOTE) for imbalanced learning," *Machine Learning*, vol. 113, pp. 4903–4923, 2023, doi: 10.1007/s10994-022-06296-4.
- [5] M. R. Gusmansyah, H. Hendrawan, and P. T. Informatika, "Peningkatan Kinerja Analisis Sentimen pada Ulasan Aplikasi Identitas Kependudukan Digital (IKD) di Indonesia Menggunakan Algoritma Support Vector Machine (SVM)," vol. 10, no. 1, pp. 185–196, 2025.
- [6] A. H. Luthfi, A. Faqih, and G. Dwilestari, "Enhancing Model Accuracy in Sentiment Analysis of the by . U Application Using Naïve Bayes and SMOTE Techniques," vol. 4, no. 2, 2025.

- [7] E. Lee, C. Lee, and S. Ahn, "Comparative Study of Multiclass Text Classification in Research Proposals Using Pretrained Language Models," *Applied Sciences*, vol. 12, no. 9, p. 4522, 2022, doi: 10.3390/app12094522.
- [8] Q. Li *et al.*, "Twenty Years of Machine-Learning-Based Text Classification: A Systematic Review," *Algorithms*, vol. 16, no. 5, p. 236, 2022, doi: 10.3390/a16050236.
- [9] C.-L. Fan, "Evaluation of Classification for Project Features with Machine Learning Algorithms," *Symmetry*, vol. 14, no. 2, p. 372, 2022, doi: 10.3390/sym14020372.
- [10] G. M. D. Minzoni and F. Nanni, "A thorough reproducibility study on sentiment classification: methodology, experimental setting, results," *Information*, vol. 14, no. 2, p. 76, 2023, doi: 10.3390/info14020076.
- [11] I. Magnusson, N. A. Smith, and J. Dodge, "Reproducibility in NLP: What have we learned from the checklist?," *Transactions of the Association for Computational Linguistics*, vol. 11, pp. 456–472, 2023, doi: 10.1162/tacl_a_00675.
- [12] A. Belz, C. Thomson, E. Reiter, and S. Mille, "Non-repeatable experiments and non-reproducible results: The reproducibility crisis in human evaluation in NLP," in *Findings of ACL 2023*, 2023, pp. 21–31. doi: 10.18653/v1/2023.findings-acl.226.
- [13] S. Martínez-Colon, J. Rodríguez-Gómez, and F. Espinosa-Rodríguez, "TextFlows: an open science NLP evaluation approach," *Language Resources & Evaluation*, vol. 59, no. 5, pp. 1125–1147, 2024, doi: 10.1007/s10579-024-09793-1.