

ANALISIS SENTIMEN ULASAN PENGGUNA APLIKASI JOBSTREET MENGUNAKAN METODE ALGORITMA NAÏVE BAYES DAN SUPPORT VECTOR MACHINE

Muhammad Fauzan Abdillah¹, Martanto², Yudhistira Arie Wijaya³, Ade Rizki Rinaldi⁴

Program Studi Teknik Informatika¹
Program Studi Manajemen Informatika²
Program Studi Sistem Informasi³
Program Studi Rekayasa Perangkat Lunak⁴

STMIK IKMI Cirebon
<https://ikmi.ac.id/page/18/?lang=de>
muhammadfauzanabdillah71@gmail.com

(*) Corresponding Author : muhammadfauzanabdillah71@gmail.com
Published : 16 Juli 2026

Abstract—This study aims to analyze user sentiment toward the JobStreet application available on the Google Play Store using text classification methods. Sentiment analysis is conducted to understand user perceptions of the application's service quality and to identify common issues experienced by users. The dataset used in this study consists of 3,000 user reviews collected through a web scraping technique. The research process begins with data collection, followed by data preprocessing, which includes lowercase transformation, cleaning and normalization, stopword removal, tokenizing, and stemming. The data is then labeled based on user ratings into three sentiment categories: positive, negative, and neutral. Next, term weighting is performed using the Term Frequency-Inverse Document Frequency (TF-IDF) method to convert textual data into numerical representations. To address the issue of imbalanced data, the Synthetic Minority Over-sampling Technique (SMOTE) is applied, resulting in a balanced dataset. The data is then divided into training and testing sets with an 80:20 ratio. The classification models used in this study are Naïve Bayes and Support Vector Machine (SVM). The results show that both methods are capable of classifying sentiment effectively. The Naïve Bayes model achieved an accuracy of 82.67%, while the SVM model achieved a higher accuracy of 84.19%. Based on evaluation metrics such as precision, recall, and F1-score, SVM demonstrates more stable performance compared to Naïve Bayes, particularly in handling high-dimensional text data. Furthermore, the analysis results indicate that most user reviews are positive, suggesting that the JobStreet application is generally perceived as helpful and effective in supporting job searching activities. However, negative reviews still exist, mainly related to technical issues such as application errors, login difficulties, and performance instability. These findings suggest that although the application provides significant benefits, improvements are still needed to enhance user satisfaction. The scientific contribution of this study lies in providing a comparative analysis of Naïve Bayes and SVM methods in Indonesian text sentiment analysis, particularly in the domain of job search applications. In addition, this study highlights the importance of preprocessing and handling imbalanced data in improving classification performance. The results are expected to serve as a reference for future research and provide valuable insights for application developers to improve service quality.

Keywords: Sentiment Analysis, JobStreet, Naïve Bayes, Support Vector Machine, TF-IDF, SMOTE.

Abstrak—Penelitian ini bertujuan untuk menganalisis sentimen ulasan pengguna aplikasi JobStreet yang tersedia di Google Play Store dengan menggunakan metode klasifikasi teks. Analisis sentimen dilakukan untuk mengetahui persepsi pengguna terhadap kualitas layanan aplikasi serta mengidentifikasi permasalahan yang sering dikeluhkan oleh pengguna. Data yang digunakan dalam penelitian ini merupakan data primer berupa 3000 ulasan pengguna yang diperoleh melalui teknik *web scraping*. Tahapan penelitian dimulai dari proses pengumpulan data, kemudian dilanjutkan dengan preprocessing data yang meliputi *lowercase*, *cleaning* dan normalisasi, *stopword removal*, *tokenizing*, serta *stemming*. Setelah itu, dilakukan pelabelan data berdasarkan rating menjadi tiga kategori sentimen, yaitu positif, negatif, dan netral. Selanjutnya, dilakukan pembobotan kata menggunakan metode *Term Frequency-Inverse Document Frequency* (TF-IDF) untuk mengubah data teks menjadi representasi numerik. Untuk mengatasi ketidakseimbangan data, digunakan metode *Synthetic Minority Over-sampling Technique* (SMOTE) sehingga

distribusi data menjadi seimbang. Data kemudian dibagi menjadi data latih dan data uji dengan perbandingan 80:20. Model klasifikasi yang digunakan dalam penelitian ini adalah *Naïve Bayes* dan *Support Vector Machine* (SVM). Hasil penelitian menunjukkan bahwa kedua metode mampu mengklasifikasikan data sentimen dengan baik. Model *Naïve Bayes* memperoleh nilai akurasi sebesar 82,67%, sedangkan model SVM memperoleh nilai akurasi yang lebih tinggi yaitu sebesar 84,19%. Berdasarkan hasil evaluasi menggunakan metrik *precision*, *recall*, dan *F1-score*, SVM menunjukkan performa yang lebih stabil dibandingkan *Naïve Bayes*, terutama dalam menangani data teks dengan dimensi tinggi. Selain itu, hasil analisis menunjukkan bahwa mayoritas ulasan pengguna memiliki sentimen positif, yang mengindikasikan bahwa aplikasi *JobStreet* dinilai cukup baik dan membantu dalam proses pencarian kerja. Namun demikian, masih terdapat ulasan negatif yang berkaitan dengan kendala teknis seperti error, kesulitan *login*, dan performa aplikasi yang kurang optimal. Temuan ini menunjukkan bahwa meskipun aplikasi telah memberikan manfaat yang baik, masih diperlukan peningkatan kualitas layanan untuk meningkatkan kepuasan pengguna. Kontribusi ilmiah dari penelitian ini adalah memberikan perbandingan performa antara metode *Naïve Bayes* dan SVM dalam analisis sentimen teks berbahasa Indonesia, khususnya pada domain aplikasi pencarian kerja. Selain itu, penelitian ini juga menunjukkan pentingnya tahap *preprocessing* dan penanganan data *imbalance* dalam meningkatkan akurasi model klasifikasi. Hasil penelitian ini diharapkan dapat menjadi referensi bagi penelitian selanjutnya serta memberikan masukan bagi pengembang aplikasi dalam meningkatkan kualitas layanan.

Kata Kunci: Analisis Sentimen, *JobStreet*, *Naïve Bayes*, *Support Vector Machine*, TF-IDF, SMOTE

INTRODUCTION

Perkembangan pesat di bidang informatika telah membawa perubahan signifikan dalam berbagai aspek kehidupan manusia. Kemajuan teknologi informasi dan komunikasi memungkinkan akses data yang lebih cepat, akurat, dan efisien, sehingga mendukung aktivitas di berbagai sektor seperti bisnis, pendidikan, kesehatan, dan pemerintahan. Dalam dunia bisnis, teknologi informatika dimanfaatkan untuk meningkatkan efisiensi operasional dan pengambilan keputusan berbasis data [1]. Di bidang pendidikan, pemanfaatan teknologi digital telah mendorong munculnya sistem pembelajaran daring yang fleksibel dan mudah diakses [2]. Selain itu, perkembangan informatika juga memengaruhi cara masyarakat berinteraksi dan menyampaikan pendapat melalui *platform* digital seperti aplikasi dan media sosial, yang menjadi sumber data penting dalam berbagai penelitian [3].

Meskipun manfaatnya besar, kemajuan informatika juga menimbulkan tantangan signifikan, terutama terkait kelolaan data digital berukuran besar yang dihasilkan setiap hari dalam bentuk teks tidak terstruktur seperti ulasan pengguna, komentar, dan opini di *platform* digital [4]. Volume tinggi, variasi bahasa, dan sifat subjektif data tersebut menjadikan analisis manual tidak praktis, sehingga diperlukan metode otomatis untuk memahami data tersebut secara efektif [2]. Kenaikan jumlah pengguna teknologi di Indonesia yang tersebar luas hingga daerah

terpencil turut meningkatkan kompleksitas data opini yang dihasilkan [2]. Oleh karena itu, analisis sentimen menjadi solusi relevan untuk mengidentifikasi dan mengklasifikasikan opini pengguna ke dalam kategori seperti positif dan negatif, sehingga mendukung evaluasi produk serta pengambilan keputusan berbasis data di bidang informatika [4].

Sejumlah penelitian terdahulu telah membahas analisis sentimen menggunakan berbagai metode *machine learning*. Misalnya, [1] menggunakan *Naïve Bayes*, menunjukkan bahwa *Naïve Bayes* relatif lebih stabil pada dataset kecil. Penelitian lain menunjukkan bahwa *Naïve Bayes* efektif dalam mengklasifikasikan sentimen ulasan teks jika didukung *preprocessing* seperti *tokenizing*, *stopword removal*, dan *stemming*, serta pembobotan TF-IDF untuk peningkatan kualitas representasi fitur dan performa model [2]. Namun, beberapa studi menghadapi keterbatasan seperti dominasi bahasa non-Indonesia pada dataset, ukuran dataset yang terbatas, serta belum sepenuhnya menangkap karakteristik bahasa informal Indonesia. Fenomena ini membuka peluang penelitian lebih lanjut dengan fokus pada data berbahasa Indonesia dan konteks objek yang lebih spesifik agar analisisnya lebih relevan dan akurat [1].

Tujuan penelitian ini adalah menganalisis sentimen ulasan pengguna aplikasi *JobStreet* di *Google Play Store* dengan menggunakan metode *Naïve Bayes* dan SVM, untuk mengklasifikasikan opini ke dalam sentimen positif dan negatif. Penelitian ini juga bertujuan menguji efektivitas *Naïve Bayes* dan SVM dalam mengolah data teks berbahasa Indonesia pada ulasan aplikasi yang

bersifat tidak terstruktur. Signifikansi penelitian terletak pada upaya mengisi kesenjangan penelitian sebelumnya yang sering fokus pada dataset kecil atau bahasa non-Indonesia, serta memberikan panduan praktis bagi pengembang aplikasi dan peneliti di bidang informatika tentang penerapan teknik *machine learning* untuk data tidak terstruktur berbahasa Indonesia [5].

Metodologi ringkas yang dirancang meliputi pengumpulan data melalui *web scraping* ulasan *Google Play Store*, *preprocessing* (*case folding, cleansing, tokenizing, stopword removal, stemming*), pembagian data *training/testing*, konstruksi model *Naïve Bayes*, dan evaluasi menggunakan *confusion matrix* serta metrik akurasi, *precision*, *recall*, dan *F1-score*. Diharapkan dengan pendekatan ini, klasifikasi sentimen dapat mencapai tingkat akurasi yang layak untuk data teks berbahasa Indonesia dan berkontribusi pada peningkatan evaluasi layanan aplikasi berdasarkan opini pengguna [6].

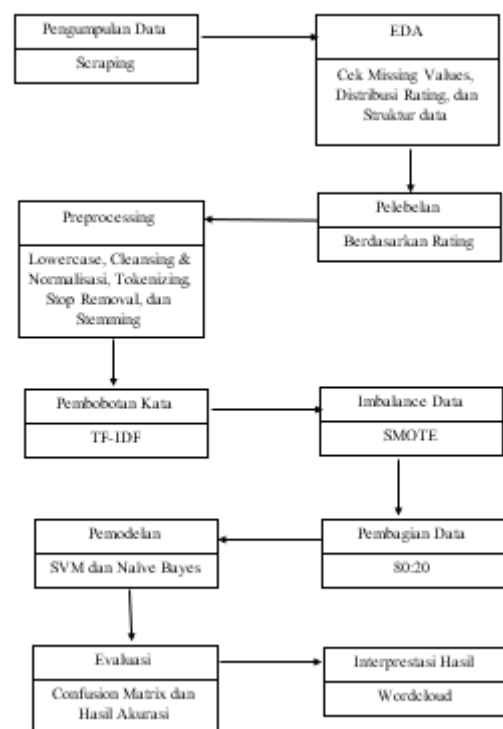
Penelitian-penelitian Indonesia terkait analisis sentimen ulasan aplikasi di *Google Play Store* menunjukkan variasi objek aplikasinya (misalnya Grab, Mola, WeTV, Blu BCA, MyPertamina, TikTok Shop, KoinWorks, dan lain-lain) dan konsisten menggunakan NBC dengan TF-IDF atau kombinasi *lexicon-based* untuk peningkatan kinerja. Hasilnya sering menunjukkan akurasi di kisaran 60–90% tergantung dataset dan metrik evaluasi. Studi-studi ini menekankan pentingnya *preprocessing* (*case folding, cleansing, tokenization, stopword removal, stemming*) dan pembobotan TF-IDF untuk meningkatkan representasi fitur serta performa NB dalam bahasa Indonesia (Allorerung [5] [6] [4] [7] [8]).

Apabila tujuan penelitian ini berhasil dicapai, maka hasilnya diharapkan dapat memberikan kontribusi signifikan dalam meningkatkan pemahaman terhadap analisis sentimen pada data teks berbahasa Indonesia, khususnya pada ulasan pengguna aplikasi. Penelitian ini dapat menjadi referensi dalam penerapan metode *machine learning*, terutama *Naïve Bayes*, untuk mengolah data tidak terstruktur secara efektif dan efisien. Selain itu, temuan dari penelitian ini juga dapat dimanfaatkan oleh pengembang aplikasi sebagai bahan evaluasi untuk meningkatkan kualitas layanan berdasarkan opini pengguna. Bagi praktisi di bidang informatika, hasil penelitian ini dapat menjadi dasar dalam mengembangkan sistem analisis sentimen yang lebih canggih dan adaptif terhadap berbagai jenis data teks. Di sisi

lain, bagi peneliti, penelitian ini dapat membuka peluang untuk pengembangan lebih lanjut, seperti penggunaan metode lain yang lebih kompleks atau pengolahan data dengan skala yang lebih besar. Penelitian ini juga berpotensi mendorong pemanfaatan data ulasan pengguna sebagai sumber informasi strategis dalam pengambilan keputusan berbasis data. Dengan demikian, penelitian ini tidak hanya memberikan kontribusi secara akademis, tetapi juga memiliki dampak praktis terhadap perkembangan teknologi dan pemanfaatannya di berbagai sektor.

MATERIALS AND METHODS

Alur penelitian dalam penelitian ini disusun secara sistematis untuk memastikan setiap tahapan dilakukan secara terstruktur dan menghasilkan output yang sesuai dengan tujuan penelitian. Proses penelitian dimulai dari pengumpulan data hingga tahap interpretasi hasil dan penarikan kesimpulan.



Gambar 1 Alur Penelitian

Tahap pertama adalah pengumpulan data, yaitu dengan melakukan *scraping* terhadap ulasan pengguna aplikasi *JobStreet* pada *Google Play Store*. Data yang dikumpulkan meliputi teks ulasan, rating, dan tanggal ulasan. Data yang diperoleh kemudian disimpan dalam format terstruktur seperti CSV atau *Excel* untuk memudahkan proses pengolahan selanjutnya.

Tahap kedua adalah *Exploratory Data Analysis* (EDA) dan *Data Cleansing*. Pada fase ini, EDA dilakukan terlebih dahulu sebagai langkah pemeriksaan untuk memvalidasi kualitas data yang telah dikumpulkan. Jika hasil eksplorasi menunjukkan bahwa data sudah lengkap, konsisten, dan tidak memiliki masalah seperti *missing value* atau ulasan duplikat, maka tahap *cleansing* tidak perlu dilakukan secara ekstensif. Namun, jika ditemukan adanya ketidaksesuaian atau data yang kosong, maka tindakan *cleansing* menjadi wajib untuk memastikan data benar-benar bersih. Dengan kata lain, *cleansing* adalah langkah kondisional yang sangat bergantung pada temuan saat proses EDA guna menjaga integritas data sebelum dianalisis lebih lanjut.

Tahap ketiga dalam penelitian ini adalah pelabelan data berdasarkan rating. Pada tahap ini, data ulasan diklasifikasikan ke dalam tiga kategori sentimen, yaitu positif, netral, dan negatif berdasarkan nilai rating yang diberikan oleh pengguna. Rating tinggi (4–5) dikategorikan sebagai sentimen positif, rating sedang (3) sebagai sentimen netral, dan rating rendah (1–2) sebagai sentimen negatif. Pendekatan ini umum digunakan dalam penelitian analisis sentimen karena rating dianggap merepresentasikan opini pengguna terhadap suatu produk atau layanan. Beberapa penelitian menyatakan bahwa terdapat hubungan yang kuat antara nilai rating dan sentimen teks, sehingga rating dapat digunakan sebagai dasar pelabelan data dalam proses klasifikasi [9]. Namun demikian, terdapat kemungkinan ketidaksesuaian antara isi teks ulasan dan nilai rating yang diberikan pengguna, sehingga perlu diperhatikan dalam proses analisis lebih lanjut [9].

Tahap keempat adalah preprocessing data, yang bertujuan untuk membersihkan dan menyiapkan data teks agar siap digunakan dalam proses analisis. Tahapan *preprocessing* meliputi *case folding*, *tokenizing*, *stopword removal*, dan *stemming*. Proses ini dilakukan untuk menghilangkan *noise* dan menyederhanakan bentuk teks sehingga lebih mudah diproses oleh algoritma *machine learning*.

Tahap kelima adalah pembobotan kata menggunakan metode TF-IDF (*Term Frequency – Inverse Document Frequency*). Pada tahap ini, setiap kata dalam teks akan diberikan bobot numerik berdasarkan tingkat kepentingannya dalam dokumen, sehingga data teks dapat direpresentasikan dalam bentuk vektor numerik.

Tahap keenam adalah penanganan *imbalance* data. Pada tahap ini dilakukan pengecekan distribusi jumlah data pada setiap kelas sentimen. Jika ditemukan ketidakseimbangan jumlah data antar kelas, maka dilakukan teknik penyeimbangan data seperti *SMOTE* pada kelas minoritas atau *undersampling* pada kelas mayoritas. Tujuan dari tahap ini adalah agar model tidak bias terhadap kelas tertentu dan mampu mengklasifikasikan semua kategori sentimen secara lebih akurat.

Tahap ketujuh adalah pembagian data menjadi data *training* dan data *testing*. Pembagian ini dilakukan dengan proporsi tertentu, misalnya 80% data *training* dan 20% data *testing*. Data *training* digunakan untuk melatih model, sedangkan data *testing* digunakan untuk menguji performa model.

Tahap kedelapan adalah proses klasifikasi menggunakan algoritma *Naive Bayes* dan *Support Vector Machine* (SVM). Pada tahap ini, kedua algoritma dilatih menggunakan data *training* untuk mempelajari pola sentimen dalam data, kemudian digunakan untuk memprediksi sentimen pada data *testing*.

Tahap kesembilan adalah evaluasi model menggunakan *confusion matrix*. Berdasarkan hasil klasifikasi, dihitung metrik evaluasi seperti *accuracy*, *precision*, *recall*, dan *F1-score* untuk mengetahui performa masing-masing algoritma. Perbandingan kinerja model. Pada tahap ini dilakukan analisis perbandingan antara metode *Naive Bayes* dan *Support Vector Machine* (SVM) berdasarkan hasil evaluasi yang telah diperoleh, untuk menentukan metode yang memiliki performa terbaik.

Tahap terakhir adalah interpretasi hasil dan penarikan kesimpulan. Pada tahap ini dilakukan analisis terhadap hasil klasifikasi sentimen untuk memahami kecenderungan opini pengguna terhadap aplikasi *JobStreet*, serta memberikan rekomendasi berdasarkan hasil penelitian yang telah dilakukan.

RESULTS AND DISCUSSION

Pada tahap ini dilakukan proses pemodelan menggunakan algoritma klasifikasi untuk menganalisis sentimen ulasan pengguna aplikasi *JobStreet*. Model yang digunakan dalam penelitian ini adalah *Naive Bayes* dan *Support Vector Machine* (SVM). Kedua algoritma tersebut dipilih karena memiliki performa yang baik dalam klasifikasi teks, khususnya pada analisis sentimen.

1. Hasil Pemodelan *Naive Bayes*

Berdasarkan hasil pengujian menggunakan algoritma *Naive Bayes*, diperoleh nilai akurasi

sebesar 0.8267 atau sekitar 82.67%. Model ini menunjukkan performa yang cukup baik dalam mengklasifikasikan sentimen ulasan pengguna.

Tabel 1 Hasil Naive Bayes

Kelas	Precision	Recall	F1-Score	Support
Negatif	0.89	0.9	0.89	466
Netral	0.79	0.82	0.8	450
Positif	0.79	0.77	0.78	463
Accuracy			0.83	1379

Berdasarkan tabel tersebut, model *Naive Bayes* memiliki performa terbaik pada kelas negatif dengan nilai *precision* dan *recall* yang tinggi. Hal ini menunjukkan bahwa model mampu mengenali ulasan negatif dengan baik. Namun, pada kelas netral dan positif masih terdapat beberapa kesalahan klasifikasi.

2. Hasil Pemodelan *Support Vector Machine* (SVM)

Berdasarkan hasil pengujian menggunakan algoritma *Support Vector Machine* (SVM), diperoleh nilai akurasi sebesar 0.8419 atau sekitar 84.19%. Hasil ini menunjukkan bahwa SVM memiliki performa yang lebih baik dibandingkan *Naive Bayes*.

Tabel 2 Hasil SVM

Kelas	Precision	Recall	F1-Score	Support
Negatif	0.82	0.93	0.87	466
Netral	0.87	0.9	0.89	450
Positif	0.84	0.7	0.76	463
Accuracy			0.84	1379

Berdasarkan tabel di atas, model SVM memiliki performa yang sangat baik pada kelas negatif dan netral, ditunjukkan oleh nilai *recall* yang tinggi. Namun, pada kelas positif masih terdapat beberapa kesalahan prediksi yang menyebabkan nilai *recall* lebih rendah dibandingkan kelas lainnya.

Berdasarkan hasil pemodelan yang telah dilakukan, kedua algoritma mampu mengklasifikasikan data dengan baik. Namun, *Support Vector Machine* (SVM) memiliki performa yang lebih unggul dengan nilai akurasi yang lebih tinggi dibandingkan *Naive Bayes*. Oleh karena itu, SVM dapat dipertimbangkan sebagai model terbaik dalam penelitian ini sebelum dilakukan evaluasi lebih lanjut.

Evaluasi

Pada tahap ini dilakukan evaluasi terhadap model yang telah dibangun untuk mengetahui tingkat performa dalam

mengklasifikasikan sentimen ulasan pengguna aplikasi *JobStreet*. Evaluasi dilakukan menggunakan *confusion matrix* serta perbandingan nilai akurasi dari masing-masing model.

1. *Confusion Matrix Naive Bayes*

Tabel 3 Hasil *Confusion Matrix Naive Bayes*

Aktual \ Prediksi	Negatif	Netral	Positif
Negatif	418	17	31
Netral	22	367	61
Positif	29	79	355

Berdasarkan tabel di atas, *confusion matrix* menunjukkan perbandingan antara data aktual dan hasil prediksi model. Nilai pada diagonal utama, yaitu 418, 367, dan 355, merupakan jumlah data yang berhasil diklasifikasikan dengan benar oleh model. Sebagai contoh, terdapat 418 data dengan sentimen negatif yang berhasil diprediksi dengan benar sebagai negatif.

Sementara itu, nilai di luar diagonal menunjukkan kesalahan klasifikasi. Sebagai contoh, terdapat 22 data yang sebenarnya memiliki sentimen netral namun diprediksi sebagai negatif, serta 79 data yang sebenarnya positif namun diprediksi sebagai netral. Kesalahan ini menunjukkan bahwa model masih mengalami kesulitan dalam membedakan beberapa data yang memiliki kemiripan makna antar sentimen.

2. *Confusion Matrix Support Vector Machine* (SVM)

Tabel 4 Hasil *Confusion Matrix SVM*

Aktual \ Prediksi	Negatif	Netral	Positif
Negatif	433	13	20
Netral	3	406	41
Positif	93	48	322

Berdasarkan tabel tersebut, model SVM menunjukkan performa yang lebih baik dalam mengklasifikasikan data. Hal ini terlihat dari nilai diagonal yang lebih tinggi, seperti 433 pada kelas negatif dan 406 pada kelas netral, yang menunjukkan jumlah prediksi benar lebih banyak dibandingkan model *Naive Bayes*.

Namun, masih terdapat kesalahan klasifikasi, terutama pada kelas positif. Sebagai contoh, terdapat 93 data yang sebenarnya positif tetapi diprediksi sebagai negatif. Hal ini menunjukkan bahwa model masih mengalami kesulitan dalam mengenali sebagian data positif, meskipun secara keseluruhan performanya lebih baik.

CONCLUSION

Berdasarkan hasil penelitian analisis sentimen ulasan pengguna aplikasi *JobStreet*, maka dapat diambil beberapa kesimpulan sebagai berikut:

1. Penerapan metode *Naïve Bayes* dan *Support Vector Machine* (SVM) dalam klasifikasi sentimen ulasan pengguna berhasil dilakukan dengan baik melalui tahapan *preprocessing*, pembobotan kata menggunakan TF-IDF, penanganan data tidak seimbang menggunakan SMOTE, serta pembagian data menjadi data latih dan data uji. Kedua metode mampu mengklasifikasikan sentimen ke dalam kategori positif, negatif, dan netral.
2. Berdasarkan hasil evaluasi kinerja model, metode *Naïve Bayes* memperoleh tingkat akurasi sebesar 82.67%, sedangkan *Support Vector Machine* (SVM) memperoleh akurasi sebesar 84.19%. Selain itu, SVM juga menunjukkan performa yang lebih stabil pada metrik *precision*, *recall*, dan *F1-score* dibandingkan dengan *Naïve Bayes*.
3. Metode *Support Vector Machine* (SVM) ditentukan sebagai metode terbaik dalam penelitian ini karena memiliki nilai akurasi yang lebih tinggi serta kemampuan yang lebih baik dalam mengklasifikasikan data sentimen secara keseluruhan. Hal ini menunjukkan bahwa SVM lebih optimal digunakan dalam analisis sentimen teks berbahasa Indonesia pada ulasan aplikasi *JobStreet*.
4. Hasil analisis sentimen menunjukkan bahwa mayoritas ulasan pengguna bersifat positif, yang mengindikasikan bahwa aplikasi *JobStreet* dinilai cukup baik dan membantu dalam proses pencarian kerja. Namun, masih terdapat ulasan negatif yang berkaitan dengan kendala teknis seperti *error* dan kesulitan *login*, sehingga perlu adanya perbaikan dari pihak pengembang.

REFERENCE

- [1] Z. Annisa and B. S. S. Ulama, "Analisis Sentimen Data Ulasan Pengguna Aplikasi 'PeduliLindungi' Pada Google Play Store Menggunakan Metode *Naïve Bayes* Classifier Model Multinomial," *Jurnal Sains Dan Seni Its*, vol. 11, no. 6, 2023, doi: 10.12962/j23373520.v11i6.94064.
- [2] S. N. Salsabila, B. N. Sari, and R. Mayasari, "Klasifikasi Ulasan Pengguna Aplikasi Discord Menggunakan Metode Information Gain Dan *Naïve Bayes* Classifier," *Infotech Journal*, vol. 9, no. 2, pp. 383–392, 2023, doi: 10.31949/infotech.v9i2.6277.
- [3] D. Pratmanto, F. F. D. Imaniawan, and V. Maarif, "Analisis Sentimen Pada Ulasan Pengguna Aplikasi Identitas Kependudukan Digital Dengan Metode *Naive Bayes* Dan *K-Nearest*," *Computatio Journal of Computer Science and Information Systems*, vol. 7, no. 2, pp. 155–166, 2023, doi: 10.24912/computatio.v7i2.26322.
- [4] M. D. Hendriyanto, A. A. Ridha, and U. Enri, "Analisis Sentimen Ulasan Aplikasi Mola Pada Google Play Store Menggunakan Algoritma *Support Vector Machine*," *Intecom Journal of Information Technology and Computer Science*, vol. 5, no. 1, pp. 1–7, 2022, doi: 10.31539/intecom.v5i1.3708.
- [5] A. Rifa'i, R. Ardhani, D. Pratama, and F. Fatihanursari, "Analisis Sentimen Terhadap Layanan Aplikasi Grab Indonesia Menggunakan Metode *Naïve Bayes*," *Jati (Jurnal Mahasiswa Teknik Informatika)*, vol. 8, no. 1, pp. 303–309, 2024, doi: 10.36040/jati.v8i1.8425.
- [6] D. R. Firmansyah and E. Lestariningsih, "Analisis Sentimen Ulasan Aplikasi Smart Campus Unisbank Di Google Playstore Menggunakan Algoritma *Naive Bayes*," *Jurnal Jtik (Jurnal Teknologi Informasi Dan Komunikasi)*, vol. 8, no. 2, pp. 498–507, 2024, doi: 10.35870/jtik.v8i2.1882.
- [7] M. Umair and E. R. Susanto, "Analisis Sentimen Ulasan Pengguna Pada Aplikasi BRI Mo BRI Menggunakan Metode Klasifikasi Algoritma *Naive Bayes*," *Jurnal Media Informatika Budidarma*, vol. 8, no. 2, p. 1149, 2024, doi: 10.30865/mib.v8i2.7381.
- [8] W. Wahyudi, R. Kurniawan, and Y. A. Wijaya, "Analisis Sentimen Pengguna Terhadap Aplikasi Blu Bca Di Playstore Menggunakan Algoritma *Naïve Bayes*," *Jati (Jurnal Mahasiswa Teknik Informatika)*, vol. 8, no. 3, pp. 2511–2517, 2024, doi: 10.36040/jati.v8i3.9216.
- [9] P. Nandwani and R. Verma, "A Review on Sentiment Analysis and Emotion Detection From Text," *Soc. Netw. Anal. Min.*, vol. 11, no. 1, 2021, doi: 10.1007/s13278-021-00776-6.