

PENINGKATAN KINERJA ALGORITMA NAIVE BAYES DAN SUPPORT VECTOR MACHINE DENGAN SAMPLING HIBRIDA UNTUK MENANGANI KETIDAKSEIMBANGAN SENTIMEN ULASAN STEAM

Dimas Eka Syahputra¹, Rudi Kurniawan², Bani Nurhakim³, Aris Pratama Putra⁴.

Program Studi Teknik Informatika¹²
Program Studi Manajemen Informatika³⁴

STMIK IKMI Cirebon
<https://ikmi.ac.id/page/18/?lang=de>
dimasekasyahputra345@gmail.com

(*) Corresponding Author : dimasekasyahputra345@gmail.com
Published : 30 Juni 2026

Abstract—The imbalance of classes in user review data often causes sentiment analysis models to ignore minority classes, resulting in less fair and less reliable decisions. This study aims to analyze and improve the performance of Naive Bayes and Support Vector Machine (SVM) algorithms in classifying the imbalance of Steam review sentiment through the application of data balancing techniques. The secondary dataset of Steam Reviews from Kaggle was processed through text cleaning, normalization, tokenization, stop word removal, and stemming, then represented using Term Frequency-Inverse Document Frequency (TF-IDF). Experiments were conducted in three scenarios: baseline without sampling, single resampling with Synthetic Minority Over-sampling Technique (SMOTE), and hybrid sampling with SMOTE-Edited Nearest Neighbour (SMOTE-ENN). Model performance was evaluated using accuracy, precision, recall, and F1-score metrics to assess the balance of performance in each class. The results show that class imbalance reduces the recall and F1-score of the negative class in both models. The application of SMOTE on Naive Bayes was able to maintain an accuracy of around 79 percent while increasing the recall of the negative class to 0.84 and the F1-score to 0.80, resulting in the best trade-off among the metrics. In contrast, the use of SMOTE-ENN tends to excessively alter the data distribution, resulting in a decrease in accuracy and F1-score, even though the recall of the minority class increases. Based on these findings, the combination of Naive Bayes and SMOTE is the most effective approach to handle the imbalance of Steam review sentiment data and can serve as a practical reference for the development of sentiment analysis systems on similar digital platforms. The novelty of this research lies in the systematic evaluation of the influence of various resampling scenarios on the performance of text models, particularly on the English-language dataset of digital game reviews from the globally popular Steam platform.

Keywords : sentiment analysis, Naive Bayes, Support Vector Machine (SVM), SMOTE, SMOTE-ENN, data imbalance, Steam reviews.

Abstrak—Ketidakseimbangan kelas pada data ulasan pengguna sering menyebabkan model analisis sentimen mengabaikan kelas minoritas sehingga keputusan yang dihasilkan menjadi kurang adil dan kurang dapat diandalkan. Penelitian ini bertujuan menganalisis serta meningkatkan kinerja algoritma *Naive Bayes* dan *Support Vector Machine (SVM)* dalam mengklasifikasikan ketidakseimbangan sentimen ulasan *Steam* melalui penerapan teknik penyeimbangan data. Dataset sekunder *Steam Reviews* dari *Kaggle* diproses melalui tahapan pembersihan teks, normalisasi, tokenisasi, penghapusan *stopword*, dan *stemming*, kemudian direpresentasikan menggunakan *Term Frequency-Inverse Document Frequency (TF-IDF)*. Eksperimen dilakukan pada tiga skenario, yaitu *baseline* tanpa *sampling*, *resampling* tunggal dengan *Synthetic Minority Over-sampling Technique (SMOTE)*, dan *sampling* hibrida dengan *SMOTE-Edited Nearest Neighbour (SMOTE-ENN)*. Kinerja model dievaluasi menggunakan metrik akurasi, *precision*, *recall*, dan *F1-score* untuk menilai keseimbangan performa pada tiap kelas. Hasil penelitian menunjukkan bahwa ketidakseimbangan kelas menurunkan *recall* dan *F1-score* kelas negatif pada kedua model. Penerapan *SMOTE* pada *Naive Bayes* mampu mempertahankan akurasi sekitar 79 persen sekaligus meningkatkan *recall* kelas negatif menjadi 0,84 dan *F1-score* menjadi 0,80 sehingga menghasilkan *trade-off* terbaik antar metrik. Sebaliknya, penggunaan *SMOTE-ENN* cenderung mengubah distribusi data secara berlebihan sehingga menurunkan akurasi dan *F1-score* meskipun *recall* kelas minoritas meningkat. Berdasarkan

temuan tersebut, dapat disimpulkan bahwa kombinasi *Naive Bayes* dan *SMOTE* merupakan pendekatan yang paling efektif untuk menangani ketidakseimbangan data sentimen ulasan *Steam* serta dapat dijadikan rujukan praktis bagi pengembangan sistem analisis sentimen pada *platform* digital serupa. *Novelty* penelitian ini terletak pada fokus evaluasi sistematis pengaruh berbagai skenario *resampling* terhadap performa model teks, khususnya pada dataset sentimen ulasan *game* digital berbahasa Inggris dari *platform Steam* yang populer secara global.

Kata Kunci : analisis sentimen, *Naive Bayes*, *Support Vector Machine*, *SMOTE*, *SMOTE-ENN*, ketidakseimbangan data, ulasan *Steam*.

INTRODUCTION

Perkembangan teknologi informasi yang dipengaruhi oleh *machine learning*, *big data*, dan analisis otomatis mengalami kemajuan pesat sehingga berdampak pada teknologi, bisnis, dan pendidikan melalui peningkatan kemampuan pengambilan keputusan dan prediksi (M. Mujahid, et al., 2024). Distribusi data yang tidak seimbang serta pengabaian kelas minoritas dalam mengumpulkan sebuah data sentimen ulasan dan umpan balik, merupakan tantangan kritis bagi model klasifikasi seperti *Naive Bayes* dan *Support Vector Machine (SVM)* (M. Mujahid, et al., 2024). Studi terbaru menunjukkan bahwa strategi *resampling (SMOTE)* dan kerangka kerja *sampling* hibrida lainnya secara signifikan meningkatkan *recall* klasifikasi dan pengenalan kelas minoritas dalam ketidakseimbangan data (M. Mujahid, et al., 2024). Selain itu, peningkatan kinerja yang lebih baik untuk data sentimen berbasis teks juga ditunjukkan dalam studi perbandingan yang menggabungkan *sampling* dengan rekayasa fitur atau adaptasi model spesifik (M. Mujahid, et al., 2024). Temuan kontemporer ini memberikan landasan empiris yang kuat untuk mengeksplorasi *sampling* hibrida guna meningkatkan kinerja *Naive Bayes* dan *Support Vector Machine (SVM)* dalam klasifikasi ketidakseimbangan sentimen pada ulasan *Steam* (M. Mujahid, et al., 2024).

Klasifikasi seperti *Naive Bayes* atau *Support Vector Machine (SVM)* pada ketidakseimbangan dataset sentimen ulasan dikarenakan contoh kelas minoritas (misalnya, ulasan negatif atau netral) seringkali kurang terwakili. Hal ini menyebabkan hasil pembelajaran yang bias akibat dominasi kelas mayoritas (Sun et al., 2023). Kelangkaan ulasan untuk kategori tertentu dan bias positif pengguna semakin memperparah ketidakseimbangan dalam analisis sentimen seperti mengumpulkan ulasan *platform*, menghambat wawasan yang berarti (Suhaeni & Yong, 2024). Selain itu, di antara klasifikasi metode *sampling* sendiri tidak memiliki keseimbangan. Studi empiris menemukan bahwa tidak ada pendekatan

resampling tunggal yang menjamin kinerja superior di semua algoritma dan metrik, terutama untuk sentimen berbasis teks (Sun et al., 2023). Studi terdahulu telah mengidentifikasi kesenjangan dalam strategi *sampling* hibrida atau spesifik algoritma dan evaluasi sistematisnya dalam pengaturan bahasa alami (Chen, Yang, et al., 2024). Singkatnya, mengatasi ketidakseimbangan dalam mengumpulkan sentimen hasil ulasan melalui *sampling* hibrida yang disesuaikan merupakan masalah yang sangat relevan dan layak diselesaikan untuk meningkatkan kinerja klasifikasi.

Studi sebelumnya telah mengkaji ketidakseimbangan kelas dan metode *sampling* dalam *machine learning* secara mendalam dengan membandingkan 16 metode *sampling* dan menunjukkan ketidakstabilan di antara berbagai model pembelajaran, menyarankan bahwa tidak ada pendekatan *resampling* yang cocok untuk semua kasus sehingga menghasilkan wawasan yang relevan untuk meningkatkan *Naive Bayes* dan *Support Vector Machine (SVM)* dalam tugas analisis sentimen (Sun et al., 2023). Penelitian serupa lainnya, meninjau aliran ketidakseimbangan data dan menyoroti pergeseran konsep serta celah evaluasi dalam konteks analisis sentimen *streaming* (Aguiar et al., 2024). Peneliti membuat hipotesis mengenai kemajuan terbaru dalam pembelajaran ketidakseimbangan data dan mengidentifikasi kebutuhan akan metode hibrida yang menggabungkan strategi tingkat data dan algoritma (Chen, Yang, et al., 2024a). Variasi *oversampling* yang diuji peneliti saat meninjau ketidakseimbangan dataset medis menunjukkan bahwa rekayasa fitur berinteraksi dengan pilihan *resampling* dan menekankan kekhawatiran tentang ketahanan dan reproduibilitas [6]. Berdasarkan beberapa studi ini melaporkan keuntungan penting dari *resampling* tetapi meninggalkan celah dalam desain *sampling* hibrida yang ditargetkan untuk data sentimen teks (misalnya ulasan *Steam*), penyesuaian spesifik klasifikasi, dan pengujian *benchmarking* sistematis. Celah-celah ini mendorong eksperimen penelitian yang menggabungkan *Naive Bayes* dan *Support Vector Machine (SVM)* secara efektif (Salmi, Atif, et al., 2024).

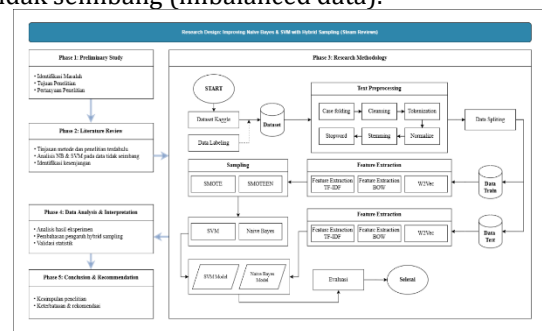
Berdasarkan hasil penelitian sebelumnya yang menyoroti ketidakstabilan persisten dari metode *resampling* tunggal di berbagai model pembelajaran, sehingga mendorong pengembangan kerangka kerja hibrida dan paralel (Sun et al., 2023; Zheng et al., 2024). Penelitian dilakukan dengan maksud untuk mengetahui peningkatan kinerja *Naive Bayes* dan *Support Vector Machine (SVM)* pada ketidakseimbangan data sentimen ulasan *Steam* dengan mengembangkan strategi *sampling* hibrida yang menggabungkan *oversampling* dan *undersampling* dengan penyesuaian yang sadar klasifikasi (Chen, Yang, et al., 2024; M. Mujahid, et al., 2024). Pendekatan hibrida berbasis GAN dan yang sadar fitur telah menunjukkan potensi dalam menghasilkan sampel minoritas yang realistis sambil mengurangi tumpang tindih dan *noise* (M. Mujahid, et al., 2024; Zhu et al., 2022). Fokus pada prapemrosesan teks spesifik, fitur *TF-IDF/embedding* kata, dan kalibrasi algoritma spesifik untuk *Naive Bayes* dan *Support Vector Machine (SVM)*, proyek ini bertujuan untuk meningkatkan *recall* minoritas, *F1* seimbang, dan generalisasi yang tangguh dalam pengaturan ulasan *streaming* (Chen, Yang, et al., 2024). Hasilnya akan mengisi celah metodologis dalam pengambilan *sampling* hibrida untuk penambangan sentimen teks dan menyediakan alat praktis untuk analitik *platform* dan moderasi, serta alat sumber terbuka secara global (Chen, Yang, et al., 2024).

Temuan yang berhasil dicapai dalam penelitian ini akan secara signifikan meningkatkan pemahaman tentang penanganan ketidakseimbangan data dalam klasifikasi sentimen berbasis teks di bidang Informatika. Pengembangan kerangka kerja *sampling* hibrida yang secara khusus dioptimalkan untuk *Naive Bayes* dan *Support Vector Machine (SVM)* dapat berfungsi sebagai acuan metodologis untuk meningkatkan kinerja klasifikasi pada ketidakseimbangan dataset secara alami, seperti ulasan yang dihasilkan pengguna (Chen, Yang, et al., 2024; M. Mujahid, et al., 2024). Kemajuan ini akan memperluas teori pembelajaran ketidakseimbangan data dengan mengintegrasikan strategi *sampling* yang sadar fitur dan spesifik algoritma (Sun et al., 2023). Pendekatan yang diusulkan menyediakan model yang dapat direproduksi dan terbuka aksesnya yang menjembatani kesenjangan antara studi *sampling* generik dan aplikasi teks dunia nyata (Zheng et al., 2024; Zhu et al., 2022). Di bidang analitik *game*, pemasaran digital, dan *e-commerce*, peningkatan deteksi sentimen dapat memperbaiki pemantauan pengalaman

pengguna dan proses pengambilan keputusan. Pada akhirnya, penelitian ini berkontribusi pada Informatika dengan memperkuat metode berbasis bukti untuk analisis sentimen yang seimbang, dapat dijelaskan, dan skalabel dalam lingkungan data dinamis.

MATERIALS AND METHODS

Desain penelitian yang berfokus pada peningkatan performa algoritma klasifikasi teks, yaitu *Naive Bayes* dan *Support Vector Machine (SVM)*, dengan pendekatan *hybrid sampling*. Penelitian ini disusun secara sistematis dalam beberapa fase, mulai dari studi pendahuluan hingga penarikan kesimpulan. Setiap tahapan dirancang untuk memastikan bahwa proses pengolahan data, pemodelan, dan evaluasi dilakukan secara terstruktur dan menghasilkan model yang optimal, khususnya dalam menangani permasalahan data tidak seimbang (*imbalanced data*).



Gambar 1. Alur Penelitian

Berdasarkan gambar tersebut, penelitian dimulai dari Phase 1: Preliminary Study, yang mencakup identifikasi masalah, tujuan penelitian, serta perumusan pertanyaan penelitian. Selanjutnya pada Phase 2: Literature Review, dilakukan kajian terhadap penelitian terdahulu, termasuk analisis metode *Naive Bayes* dan *SVM* serta identifikasi kesenjangan penelitian, khususnya terkait masalah ketidakseimbangan data.

Pada Phase 3: Research Methodology, proses dimulai dari pengambilan dataset (misalnya dari Kaggle) dan pelabelan data. Data kemudian melalui tahap *text preprocessing* yang meliputi *case folding*, *cleansing*, *tokenization*, *stopword removal*, *stemming*, dan *normalisasi*. Setelah itu, data dibagi menjadi data latih dan data uji. Tahap berikutnya adalah *feature extraction*, menggunakan metode seperti *TF-IDF*, *Bag of Words (bow)*, dan *Word2Vec* untuk mengubah teks menjadi representasi numerik.

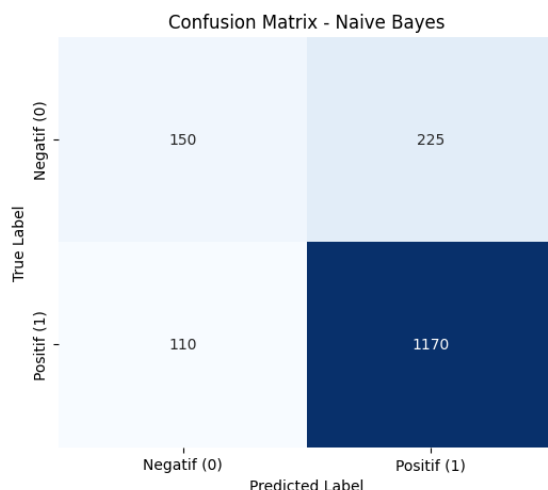
Untuk mengatasi ketidakseimbangan kelas, diterapkan teknik *sampling* seperti *SMOTE* dan *SMOTEENN*. Data yang telah diproses kemudian digunakan untuk melatih model *SVM* dan *Naive*

Bayes. Kedua model ini menghasilkan model prediksi yang selanjutnya dievaluasi untuk mengetahui performanya.

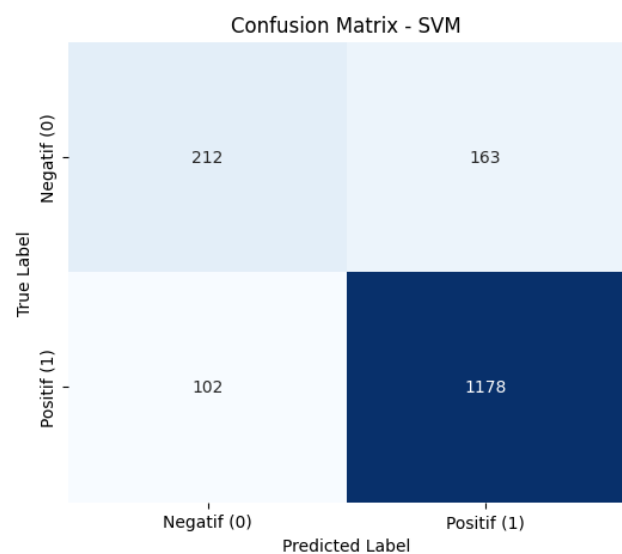
Pada Phase 4: Data Analysis & Interpretation, dilakukan analisis hasil eksperimen, pembahasan pengaruh metode hybrid sampling, serta validasi statistik untuk memastikan keandalan hasil. Terakhir, pada Phase 5: Conclusion & Recommendation, disusun kesimpulan penelitian beserta keterbatasan dan rekomendasi untuk penelitian selanjutnya. Alur ini menunjukkan pendekatan komprehensif dalam membangun model klasifikasi teks yang lebih akurat dan robust.

RESULTS AND DISCUSSION

Skenario pertama ini berfungsi sebagai tolak ukur (*benchmark*) dasar untuk memahami kinerja algoritma *Naive Bayes* dan *Support Vector Machine (SVM)* pada kondisi data asli tanpa intervensi. Pada pengujian ini, model dilatih menggunakan dataset yang memiliki ketidakseimbangan kelas moderat. Hasil evaluasi menunjukkan bahwa *Support Vector Machine (SVM)* mencapai akurasi tertinggi sebesar 84%, sedikit mengungguli *Naive Bayes* yang berada di angka 80%. Namun, *Recall* untuk kelas minoritas (*Negatif*) pada *Naive Bayes* relatif rendah (0.40), yang mengindikasikan bahwa model probabilistik ini kesulitan mengenali sentimen negatif tanpa bantuan penyeimbangan data. Temuan ini menegaskan perlunya teknik sampling untuk memperbaiki sensitivitas model.



Gambar 2 Confusion Matrix - Naive Bayes



Gambar 3 Confusion Matrix - SVM

Table 1 Hasil Evaluasi Naive Bayes (Baseline)

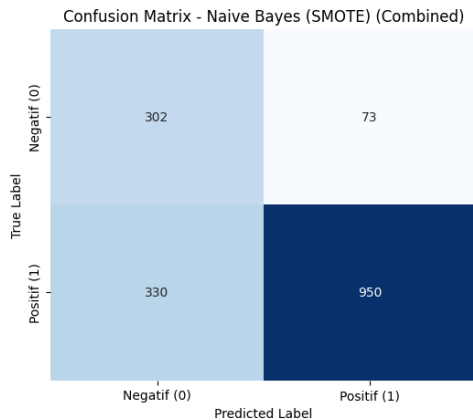
Kelas Sentimen	Precision	Recall	F1-score	Support
Negatif	0.58	0.40	0.47	375
Positif	0.84	0.91	0.87	1280
Akurasi			80%	1655
macro avg	0.71	0.66	0.67	1655
weighted avg	0.78	0.80	0.78	1655

Table 2 Hasil Evaluasi SVM (Baseline)

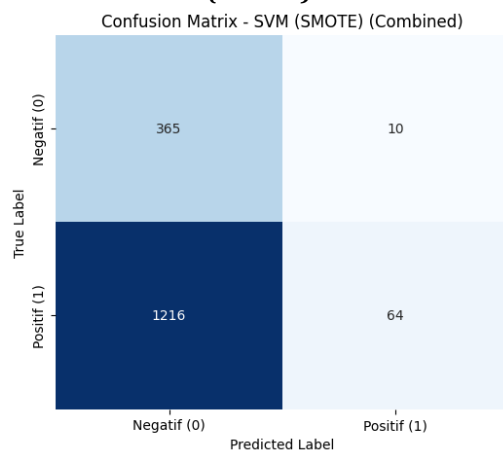
Kelas Sentimen	Precision	Recall	F1-score	Support
Negatif	0.68	0.57	0.62	375
Positif	0.88	0.92	0.90	1280
Akurasi			84%	1655
macro avg	0.78	0.74	0.76	1655
weighted avg	0.83	0.84	0.83	1655

Skenario kedua bertujuan untuk mengatasi bias model terhadap kelas mayoritas dengan menerapkan teknik *Synthetic Minority Over-sampling Technique (SMOTE)* pada data latih. Dengan menyeimbangkan jumlah sampel antar kelas, diharapkan model dapat mempelajari pola kelas minoritas dengan lebih baik. Hasil pengujian menunjukkan dampak positif yang signifikan, terutama pada algoritma Naive Bayes, di mana akurasinya meningkat menjadi 76% dan Recall kelas Negatif melonjak menjadi 0.81. Sementara itu,

SVM menunjukkan performa yang menurun dengan akurasi **26%**. Hal ini membuktikan bahwa penambahan data sintetis sangat efektif membantu Naive Bayes memperbaiki batasan keputusannya tanpa mengorbankan kinerja keseluruhan.



Gambar 4 Confusion Matrix - Naive Bayes (SMOTE)



Gambar 5 Confusion Matrix - SVM (SMOTE)

Table 3 Hasil Evaluasi - Naive Bayes (SMOTE)

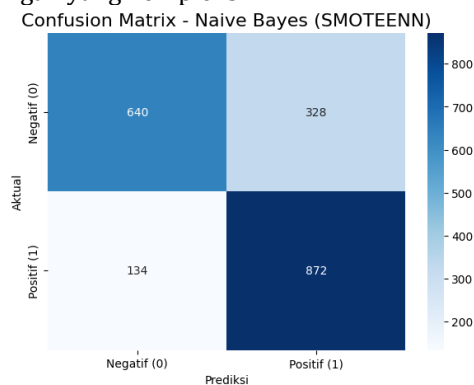
Kelas Sentimen	Precisi	Reca	F1-scor	Suppo
	on	ll	e	rt
Negatif	0.48	0.81	0.60	375
Positif	0.93	0.74	0.83	1280
Akurasi			71%	1655
macro avg	0.70	0.77	0.71	1655
weighte d avg	0.83	0.76	0.77	1655

Table 4 Hasil Evaluasi - SVM (SMOTE)

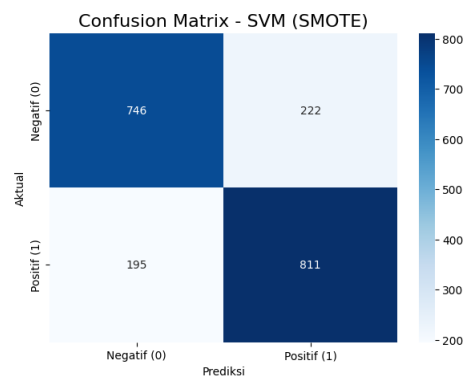
Kelas Sentimen	Precisi	Reca	F1-scor	Suppo
	on	ll	e	rt
Negatif	0.23	0.97	0.37	375
Positif	0.86	0.05	0.09	1280

Akurasi			26%	1655
macro avg	0.55	0.51	0.23	1655
weighted avg	0.72	0.26	0.16	1655

Skenario ketiga mengeksplorasi pendekatan hybrid dengan menggabungkan oversampling (SMOTE) dan undersampling (ENN) untuk membersihkan data yang dianggap noise atau tumpang tindih di perbatasan kelas. Tujuannya adalah untuk menciptakan batas pemisah kelas yang lebih tegas. Namun, hasil eksperimen menunjukkan penurunan kinerja pada kedua model dibandingkan skenario sebelumnya. Naive Bayes mengalami penurunan akurasi menjadi **76,27%**, sedangkan SVM turun menjadi **76,61%**. Penurunan ini mengindikasikan bahwa proses pembersihan data oleh ENN terlalu agresif, sehingga menghilangkan informasi variasi data yang sebenarnya penting bagi model untuk melakukan generalisasi dengan tepat pada ruang fitur gabungan yang kompleks.



Gambar 6 Confusion Matrix - Naive Bayes (SMOTE-ENN)



Gambar 7 Confusion Matrix - SVM (SMOTE-ENN)

Table 5 Hasil Evaluasi - Naive Bayes (SMOTE-ENN)

Kelas Sentimen	Precision	Recall	F1-score	Support
Negatif	0.83	0.66	0.73	968
Positif	0.73	0.87	0.79	1006
Akurasi			76,27 %	1974
macro avg	0.78	0.76	0.76	1974
weighted avg	0.78	0.77	0.76	1974

Table 6 Hasil Evaluasi untuk: SVM (SMOTE-ENN)

Kelas Sentimen	Precision	Recall	F1-score	Support
Negatif	0.83	0.70	0.76	968
Positif	0.75	0.86	0.80	1006
Akurasi			76,61 %	1974
macro avg	0.79	0.78	0.78	1974
weighted avg	0.79	0.78	0.78	1974

Berdasarkan skenario uji dan hasil yang dilakukan untuk mengukur efektivitas penerapan metode *sampling* hibrida terhadap peningkatan performa algoritma *Naive Bayes* (NB) dan *Support Vector Machine* (SVM) dalam menangani ketidakseimbangan sentimen ulasan *Steam*.

Table 7 Tabel Perbandingan

N o.	Model	Recall (0)/negatif	F1-score (Macro)	Akurasi
1	<i>Naive Bayes</i> (Baseline)	0.40	0.47	80%
2	<i>Support Vector Machine</i> (Baseline)	0.57	0.62	84%
3	<i>Naive Bayes</i> (SMOTE)	0.81	0.60	71%
4	<i>Support Vector Machine</i> (SMOTE)	0.97	0.37	26%

5	<i>Naive Bayes</i> (SMOTE-ENN)	0.95	0.75	71%
6	<i>Support Vector Machine</i> (SMOTE-ENN)	0.76	0.74	82%

CONCLUSION

Berdasarkan hasil penelitian, analisis data, serta pembahasan pada bab-bab sebelumnya, sebagai berikut:

1. Dampak ketidakseimbangan kelas terhadap performa *Naive Bayes* dan SVM. Penelitian ini menunjukkan bahwa ketidakseimbangan kelas pada ulasan *Steam* memiliki dampak signifikan terhadap kinerja model *Naive Bayes* dan SVM. Ketika data tidak diseimbangkan, kedua algoritma cenderung memberikan prediksi bias terhadap kelas mayoritas, menghasilkan nilai recall kelas minoritas yang sangat rendah dan *F1-score* yang tidak stabil. Hal ini membuktikan bahwa dominasi kelas mayoritas menghambat kemampuan model dalam mengenali pola sentimen kelas minoritas.
2. Kinerja metode resampling tunggal dibandingkan dengan algoritma *Naive Bayes* dan SVM. Hasil evaluasi menunjukkan bahwa metode resampling tunggal seperti SMOTE dan SMOTE-ENN memiliki performa yang lebih stabil dibandingkan pendekatan *hybrid*. Metode SMOTE secara khusus terbukti paling efektif dalam meningkatkan performa model, terutama pada *Naive Bayes*, dengan peningkatan recall kelas minoritas tanpa menyebabkan distorsi data yang drastis. Sebaliknya, beberapa metode undersampling menunjukkan kecenderungan menghilangkan informasi penting pada kelas mayoritas, yang berdampak pada turunnya akurasi keseluruhan.
3. Efektivitas strategi *hybrid* sampling yang dirancang secara *classifier-aware*. Eksperimen menunjukkan bahwa *hybrid sampling* berbasis SMOTE-ENN justru menghasilkan performa yang kontraproduktif pada dataset ulasan *Steam*. Meskipun metode ini mampu meningkatkan recall kelas minoritas secara signifikan, nilai akurasi dan *F1-score* model mengalami penurunan drastis. Hal ini terjadi karena ENN menghapus terlalu banyak data mayoritas yang dianggap sebagai noise, sehingga model kehilangan struktur distribusi data yang penting untuk proses klasifikasi. Dengan demikian, strategi *hybrid classifier-aware*

aware yang diuji dalam penelitian ini tidak mampu meningkatkan performa model secara seimbang dan tidak lebih unggul dibandingkan metode resampling tunggal.

REFERENCE

- [1] M. Mujahid, E. Kina, F. Rustam, E. Silva Alvarado, I. De La Torre Diez, and I. Ashraf, "Data oversampling and imbalanced datasets: An investigation of performance for machine learning and feature engineering," *J. Big Data*, vol. 11, no. Article 87, 2024, doi: 10.1186/s40537-024-00943-4.
- [2] Z. Sun, J. Zhang, X. Zhu, and D. Xu, "How far have we progressed in the sampling methods for imbalanced data classification? An empirical study," *Electronics (Basel)*, vol. 12, no. 20, 2023, doi: 10.3390/electronics12204232.
- [3] C. Suhaeni and H.-S. Yong, "Enhancing imbalanced sentiment analysis: A GPT-3-based sentence-by-sentence generation approach," *Applied Sciences*, vol. 14, no. 2, 2024, doi: 10.3390/app14020622.
- [4] W. Chen, K. Yang, Z. Yu, Y. Shi, and C. L. P. Chen, "A survey on imbalanced learning: Latest research, applications and future directions," *Artif. Intell. Rev.*, vol. 57, 2024, doi: 10.1007/s10462-024-10759-6.
- [5] G. Aguiar, A. Bifet, J. Read, J. Gama, and M. Pechenizkiy, "A survey on learning from imbalanced data streams: Taxonomy, challenges, empirical study, and reproducible experimental framework," *Mach. Learn.*, 2024, doi: 10.1007/s10994-023-06353-6.
- [6] M. Salmi, D. Atif, D. Oliva, A. Abraham, S. Ventura, and et al., "Handling imbalanced medical datasets: Review of a decade of research," *Artif. Intell. Rev.*, vol. 57, 2024, doi: 10.1007/s10462-024-10884-2.
- [7] M. Zheng *et al.*, "Exploratory parallel hybrid sampling framework for imbalanced data classification," *Eng. Appl. Artif. Intell.*, vol. 138, p. 109428, Dec. 2024, doi: 10.1016/j.engappai.2024.109428.
- [8] B. Zhu, X. Pan, S. vanden Broucke, and J. Xiao, "A GAN-based hybrid sampling method for imbalanced customer classification," *Inf. Sci. (N. Y.)*, vol. 609, pp. 1397–1411, Sep. 2022, doi: 10.1016/j.ins.2022.07.145.