



Analisis Sentimen Wisata Situ Patenggang Menggunakan Metode Multinomial Naive Bayes

Mochamad Akbar Oktavio¹, Ulya Anisatur Rosyidah², Nur Qodariyah Fitriyah³

^{1, 2, 3} Fakultas Teknik, Universitas Muhammadiyah Jember

e-mail : viomav98@gmail.com, ulyaanisatur@gmail.com

Article Info

Article history:

Received Juni 27, 2026

Revised Juli 03, 2026

Accepted Juli 14, 2026

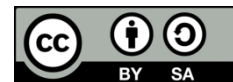
Keywords:

Sentiment Analysis, Google Maps, Multinomial Naive Bayes, Situ Patenggang, SMOTE, TF-IDF

ABSTRACT

The development of digital media has enabled tourists to express their opinions through online review platforms such as Google Maps. The data used in this study consisted of 603 Google Maps reviews collected between 2020 and 2024. The research stages included text preprocessing, sentiment labeling using a lexicon-based approach, TF-IDF weighting, data balancing using the Synthetic Minority Over-sampling Technique (SMOTE), and classification using the Multinomial Naive Bayes method. The model was evaluated using an 80:20 training and testing data split ratio. The analysis results indicated that the majority of reviews toward Situ Patenggang Tourism were positive. Before the application of SMOTE, the model achieved an accuracy of 85.12%, precision of 83.49%, recall of 100%, and F1-score of 91.00%. After applying SMOTE, the model's performance improved, achieving an accuracy of 88.43%, while precision, recall, and F1-score each reached 92.31%. This improvement demonstrates that the Multinomial Naive Bayes method is capable of effectively classifying tourism review sentiments, particularly when supported by the SMOTE data balancing technique. Therefore, the Multinomial Naive Bayes method can be considered an effective approach for text-based sentiment analysis of tourism reviews, producing more accurate and balanced classification results.

This is an open access article under the [CC BY-SA](https://creativecommons.org/licenses/by-sa/4.0/) license.



Article Info

Article history:

Received Juni 27, 2026

Revised Juli 03, 2026

Accepted Juli 14, 2026

Kata kunci:

Analisis Sentimen, Google Maps, Multinomial Naive Bayes, Situ Patenggang, SMOTE, TF-IDF

ABSTRACT

Perkembangan media digital memungkinkan wisatawan menyampaikan opini melalui platform ulasan daring seperti Google Maps. Data yang digunakan dalam penelitian ini berupa 603 ulasan Google Maps periode 2020–2024. Tahapan penelitian meliputi preprocessing teks, pelabelan sentimen menggunakan metode lexicon-based, pembobotan TF-IDF, penyeimbangan data dengan SMOTE, serta klasifikasi menggunakan Multinomial Naive Bayes. Pengujian dilakukan menggunakan rasio data latih dan data uji sebesar 80:20. Hasil analisis menunjukkan bahwa mayoritas ulasan terhadap Wisata Situ Patenggang memiliki sentimen positif. Sebelum penerapan SMOTE, model memperoleh nilai accuracy sebesar 85,12%, precision 83,49%, recall 100%, dan F1-score 91,00%. Setelah penerapan SMOTE, performa model meningkat dengan accuracy sebesar 88,43%, sedangkan nilai precision, recall, dan F1-score masing-masing mencapai 92,31%. Peningkatan kinerja tersebut menunjukkan bahwa metode Multinomial Naive Bayes mampu mengklasifikasikan sentimen ulasan wisata dengan baik, terutama ketika didukung oleh teknik penyeimbangan data SMOTE. Oleh karena itu, metode Multinomial Naive Bayes dapat digunakan sebagai pendekatan yang efektif dalam analisis sentimen ulasan wisata berbasis teks untuk menghasilkan klasifikasi yang lebih akurat dan seimbang.



This is an open access article under the [CC BY-SA](#) license.



Corresponding Author:

Mochamad Akbar Oktavio
Universitas Muhammadiyah Jember
Email: viomav98@gmail.com

PENDAHULUAN

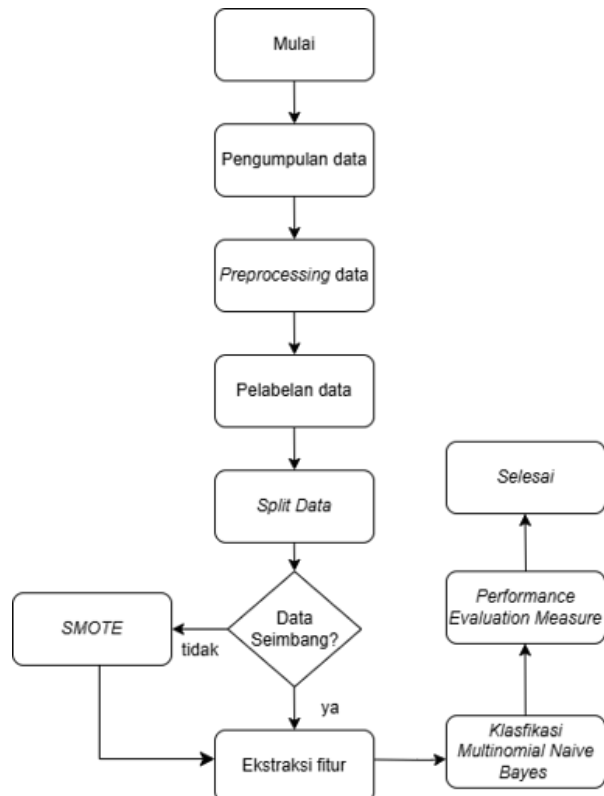
Situ Patenggang merupakan objek wisata danau yang dikelilingi kebun teh hijau di Kecamatan Rancabali, Kabupaten Bandung, yang memadukan keindahan alam, nuansa romantis, dan cerita legenda (<https://cozzy.id>). Peningkatan jumlah kunjungan wisatawan ke destinasi ini mendorong pentingnya analisis sentimen masyarakat guna memahami pengalaman mereka secara objektif. Hasil sentimen tersebut mencerminkan kepuasan pengunjung yang dapat memengaruhi keputusan calon wisatawan lain, di mana opini positif maupun negatif ini kini banyak dibagikan melalui platform ulasan digital seperti Google Maps (Dewi dkk., 2020). Untuk mengkategorikan opini tersebut, salah satu metode yang paling populer digunakan adalah *Multinomial Naïve Bayes* (Rifki dkk., 2022). Metode ini bekerja memprediksi kelas data berdasarkan probabilitas ulasan guna memberikan gambaran jelas mengenai pengalaman pengunjung. Alasan utama penerapan metode *Multinomial Naïve Bayes* dalam penelitian ini adalah karena algoritma tersebut memiliki performa yang baik pada data teks, cepat, efisien, serta mudah diimplementasikan (Abbas dkk., 2019). Efektivitas algoritma ini juga telah dibuktikan oleh berbagai riset terdahulu, meliputi analisis publik Sirkuit Mandalika dengan akurasi 78% (Mujahidin dkk., 2022), ulasan vila di Bali dengan hasil tertinggi mencapai 91% (Ginantra dkk., 2022), ulasan layanan GO-JEK berakurasi 68% (Indarwati & Februariyanti, 2023), hingga sektor pariwisata di Kabupaten Rembang sebesar 76% (Amaniputra dkk., 2023). Kendati demikian, analisis sentimen kerap dihadapkan pada kendala ketidakseimbangan data (*imbalanced data*) yang memicu bias model terhadap kelas mayoritas (Liu, 2012). Untuk mengatasinya, penelitian ini mengintegrasikan metode *Synthetic Minority Over-sampling Technique* (SMOTE) guna menyeimbangkan distribusi antar-kelas melalui pembuatan data sintetik pada kelas minoritas. Penelitian ini dirancang untuk mengatasi kesenjangan (*research gap*) studi terdahulu yang umumnya terbatas pada klasifikasi saja tanpa mengaitkannya dengan pembuatan rekomendasi berbasis kata kunci yang sering muncul (Singgalen, 2022). Melalui pendekatan ini, hasil analisis diharapkan memberikan kontribusi nyata bagi pengelola dalam meningkatkan layanan fasilitas serta mendukung pengembangan pariwisata di Kabupaten Bandung, khususnya kawasan Situ Patenggang

METODE PENELITIAN

Tahapan Penelitian

Tahap penelitian ini dimulai dari pengumpulan data ulasan wisatawan mengenai Situ Patenggang yang diambil melalui proses web scraping dari situs *Google Maps*. data yang diperoleh kemudian di proses melalui tahap *preprocessing*, seperti pembersihan teks, konversi huruf, tokenisasi, penghapusan kata umum (stopword), dan stemming. Setelah data dibersihkan, dilakukan pelabelan secara manual ke dalam kategori sentimen positif, negatif, dan netral. Tahap berikutnya adalah ekstraksi fitur menggunakan metode Tf-id untuk mengubah data teks menjadi bentuk numerik dan memvisualisasikan data menggunakan

wordcloud untuk memudahkan wisatawan awam dalam mengerti poin utama yang sering dipuji maupun dikeluhkan. Terakhir proses klasifikasi dilakukan menggunakan metode *Multinomial Naive Bayes* untuk membangun model yang mampu mengklasifikasikan komentar berdasarkan sentimen yang terkandung di dalamnya.



Gambar 1. Rangkaian Tahapan Penelitian

Dataset

Dataset pada penelitian ini terdiri atas 603 ulasan yang diperoleh dari platform Google Maps. Pengambilan data dilakukan dengan bantuan bahasa pemrograman *python*, yang digunakan untuk menulis skrip dalam mengakses endpoint dari *SerpApi*. Proses ini bertujuan untuk mendapatkan data mentah yang nantinya akan diolah dan dianalisis untuk keperluan analisis selanjutnya.

Preprocessing Data

Pada tahap preprocessing, seluruh data dibersihkan agar siap diolah secara optimal, sekaligus untuk meningkatkan akurasi dan kinerja model (Harsemadi dkk., 2023). Langkah pertama berupa proses pembersihan teks, di mana berbagai komponen yang tidak relevan seperti tautan, angka, simbol, dan karakter lain yang tidak diperlukan dalam analisis. Setelah itu, dilakukan penyeragaman huruf bentuk huruf melalui case folding sehingga seluruh teks dikonversi menjadi huruf kecil. Kemudian teks dinormalisasi agar kata-kata yang tidak baku atau slang disesuaikan dengan padanan yang sesuai KBBI. Teks yang sudah bersih kemudian diproses melalui tokenisasi untuk memisahkan kalimat menjadi deretan kata. Berikutnya dilakukan stopwords removal menggunakan Sastrawi untuk menghilangkan kata-kata yang tidak memiliki nilai informatif terhadap analisis. Tahap terakhir adalah stemming menggunakan library Sastrawi, untuk mengembalikan setiap kata ke bentuk dasarnya, sehingga model dapat membaca makna kata dengan lebih konsisten.



Pelabelan

Pada penelitian ini, proses pelabelan dilakukan secara otomatis menggunakan metode *Lexicon-Based*, yaitu pendekatan analisis sentimen yang memanfaatkan kamus kata atau frasa dengan bobot polaritas manual (Manullang dkk., 2023). Proses pelabelan dilakukan dengan menjumlahkan skor setiap kata dalam kalimat—di mana kata positif bernilai 1, netral bernilai 0, dan negatif bernilai -1—untuk menentukan kategori sentimen akhir berupa positif, negatif, atau netral. Dalam penelitian ini memfokuskan analisis pada dua kategori sentimen, yaitu positif dan negatif. Berdasarkan proses pelabelan tersebut, diperoleh 453 ulasan positif dan 149 ulasan negatif. Data yang sudah terlabeli tersebut selanjutnya digunakan sebagai dasar dalam proses pelatihan dan pengujian model klasifikasi sentimen. Hasil dari proses pelabelan ditampilkan pada tabel 1.

Tabel 1. Hasil Pelabelan Otomatis

No	Hasil preprocessing	Skor lexicon	Sentimen label
1	bagus nyaman untuk dingin otak hati hehe	2	Positif
2	bayar tiket masuk lumayan mahal tidak ingin foto pinggir kenal tarif titip hlm kenal biaya kecewa trlalu biaya pdahal numpang doang	-1	Negatif
3	suasa sangat syahdu nyaman udara segar dingin	1	Positif
4	pungli bayar aneh parkir main tembak harga motor tanpa karcis resmi mohon benah	-4	Negatif
5	tiket wahana lain tinggal kasih liat tiket perahu bayar rp an tidak layak cok pungli bayar bayar mulu an kocak	-4	Negatif
...	...	...	...
598	memukau	1	Positif
599	tempat bagus	1	Positif
600	asik untuk cari inspirasi	1	Positif
601	alam bagus	1	Positif
602	mantap tampak sisi luar indo	1	Positif

Split Data

Pada Tahap ini, Split Validation diterapkan untuk membagi dataset menjadi data latih dan data uji. Teknik ini dapat menggunakan pendekatan validasi silang bergulir, yaitu melatih model dengan sebagian kecil data awal lalu mengujinya secara bertahap menggunakan data urutan berikutnya (Amaniputra dkk., 2023). Pembagian data yang digunakan adalah rasio 80 : 20.

Tabel 2. Hasil Split Data

No	Ulasan	Label	Tipe Data
1	bagus nyaman untuk dingin otak hati hehe	Positif	Latih
2	bayar tiket masuk lumayan mahal tidak ingin foto pinggir kenal tarif titip hlm kenal biaya kecewa trlalu biaya pdahal numpang doang	Negatif	Latih
3	suasa sangat syahdu nyaman udara segar dingin	Positif	Latih
4	pungli bayar aneh parkir main tembak harga motor tanpa karcis resmi mohon benah	Negatif	Latih
5	tiket wahana lain tinggal kasih liat tiket perahu bayar rp an tidak layak cok pungli bayar bayar mulu an kocak	Negatif	Latih
...	...	...	...
598	memukau	Positif	Uji
599	tempat bagus	Positif	Uji
600	asik untuk cari inspirasi	Positif	Uji



No	Ulasan	Label	Tipe Data
601	alam bagus	Positif	Uji
602	mantap tampak sisi luar indo	Positif	Uji

Proses *split data* menghasilkan 480 data latih (361 positif, 120 negatif) dan 121 data uji (91 positif, 30 negatif). Karena distribusi kelas pada data latih tidak seimbang dan didominasi kelas positif, metode *Synthetic Minority Over-sampling Technique* (SMOTE) diterapkan sebelum pelatihan model untuk menyeimbangkan jumlah data.

Synthetic Minority Over-Sampling (SMOTE)

Pada tahap ini diterapkan metode SMOTE untuk menyeimbangkan data yang tidak seimbang. Metode ini dilakukan dengan mengambil data minoritas dan membuat data sintetik baru di titik-titik yang terletak di antara sampel tersebut serta sampel terdekat lainnya dalam kelas minoritas (Nurhopipah & Magnolia, 2023).

Tabel 3. Hasil SMOTE

Sebelum SMOTE		Setelah SMOTE	
Positif	Negatif	Positif	Negatif
362	119	362	362

Berdasarkan Tabel 3, sebelum dilakukan proses SMOTE jumlah data pada kelas positif sebanyak 362 data dan kelas negatif sebanyak 119 data. Ketidakseimbangan ini berpotensi menyebabkan model lebih cenderung mempelajari pola dari kelas mayoritas. Oleh karena itu, diterapkan metode SMOTE untuk menghasilkan data sintetis pada kelas minoritas sehingga jumlah data pada kedua kelas menjadi seimbang, yaitu masing-masing 362 data. Dengan kondisi data yang lebih seimbang, model diharapkan dapat melakukan klasifikasi secara lebih objektif dan menghasilkan performa yang lebih baik.

Ekstraksi Fitur

Pada Tahap ekstraksi fitur, hasil data teks dirubah menjadi bentuk numerik menggunakan tf-idf dengan tujuan memberikan bobot pada setiap kata. Kata yang sering muncul dalam sebuah dokumen namun jarang ditemukan di dokumen lain akan memiliki bobot yang lebih tinggi, sedangkan kata yang tersebar di banyak dokumen akan memiliki bobot yang lebih rendah (Septian dkk., 2019). Selanjutnya, hasil pembobotan TF-IDF divisualisasikan menggunakan WordCloud untuk menampilkan kata-kata yang paling dominan berdasarkan bobotnya sehingga memudahkan interpretasi karakteristik setiap sentimen.

Multinomial Naive Bayes

Multinomial Naive Bayes digunakan untuk proses klasifikasi teks dengan menghitung probabilitas suatu kelas berdasarkan distribusi kata dalam sebuah dokumen. *Naive Bayes* bekerja dengan pendekatan probabilitas kondisional dengan asumsi independensi antar fitur, sedangkan *Multinomial Naive Bayes* memodelkan data berdasarkan distribusi multinomial dengan mempertimbangkan frekuensi kemunculan setiap istilah dalam dokumen (Lesmana & Nabya, 2020).

Evaluasi

Dalam menilai efektivitas model dalam melakukan klasifikasi sentimen, evaluasi dilakukan dengan menggunakan *Performance Evaluation Measure* (PEM) merupakan metode yang digunakan untuk mengevaluasi kinerja suatu model dalam menjalankan fungsinya secara efektif dan efisien (Amaniputra dkk., 2023). Evaluasi ini bertujuan untuk mengetahui sejauh



mana hasil prediksi model sesuai dengan data aktual sehingga tingkat akurasi dan efektivitas model dapat diukur (Rahma dkk., 2021). Pada permasalahan klasifikasi, PEM umumnya direpresentasikan menggunakan *confusion matrix*, yaitu sebuah tabel yang membandingkan hasil prediksi model dengan kelas sebenarnya. *Confusion matrix* terdiri atas empat komponen, yaitu True Positive (TP), True Negative (TN), False Positive (FP), dan False Negative (FN). TP menunjukkan jumlah data positif yang diprediksi benar, TN menunjukkan jumlah data negatif yang diprediksi benar, FP merupakan data negatif yang salah diprediksi sebagai positif, sedangkan FN merupakan data positif yang salah diprediksi sebagai negatif. Berdasarkan nilai-nilai pada *confusion matrix*, diperoleh beberapa metrik evaluasi yang umum digunakan, yaitu Accuracy, Precision, Recall, dan F1-Score. Accuracy mengukur tingkat ketepatan prediksi model secara keseluruhan, Precision menunjukkan proporsi prediksi positif yang benar, Recall mengukur kemampuan model dalam mengidentifikasi seluruh data positif, sedangkan F1-Score merupakan rata-rata harmonis antara Precision dan Recall sehingga mampu menggambarkan keseimbangan kedua metrik tersebut.

$$\text{Rumus Precision (Pre)} \quad \text{Pre} = \frac{TP}{FP + TP} \quad (1)$$

$$\text{Rumus Accuracy (Acc)} \quad \text{Acc} = \frac{TN + TP}{FN + FP + TN + TP} \quad (2)$$

$$\text{Rumus Recall (Recc)} \quad \text{Recc} = \frac{TP}{FN + FP + TN + TP} \quad (3)$$

$$\text{Rumus F1-Score} \quad \text{F1-Score} = \frac{2 \times \text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}} \quad (4)$$

HASIL DAN PEMBAHASAN

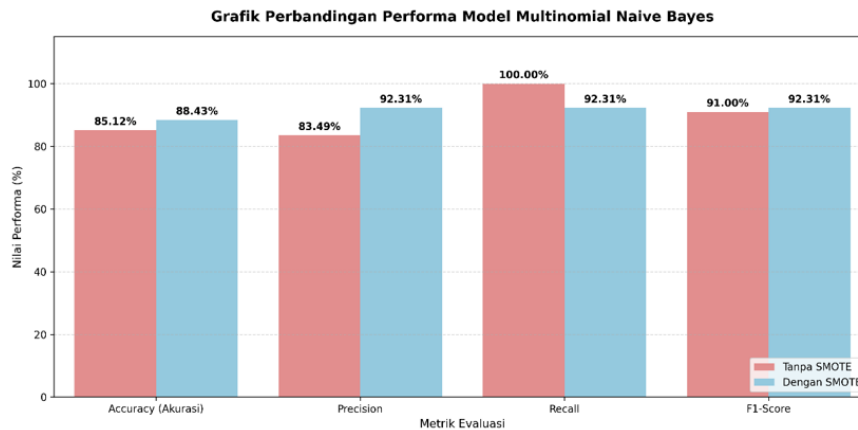
Hasil Klasifikasi

Penelitian ini dilakukan menggunakan metode Multinomial Naive Bayes dengan membandingkan hasil klasifikasi sebelum dan sesudah penerapan Synthetic Minority Over-sampling Technique (SMOTE). Perbandingan tersebut bertujuan untuk mengetahui pengaruh penyeimbangan data terhadap kemampuan model dalam mengklasifikasikan data sentimen. Hasil klasifikasi disajikan berdasarkan nilai accuracy, precision, recall, dan F1-Score.

Tabel 4. Perbandingan hasil pengujian

	Tanpa SMOTE	Menggunakan SMOTE
Accuracy	85.12%	88.43%
Precision	83.49%	92.31%
Recall	100.00%	92.31%
F1-Score	91.00%	92.31%

Berdasarkan Tabel 3, model Multinomial Naive Bayes tanpa penerapan SMOTE memperoleh nilai accuracy sebesar 85,12%, precision sebesar 83,49%, recall sebesar 100,00%, dan F1-Score sebesar 91,00%. Setelah data latih diseimbangkan menggunakan metode SMOTE, nilai accuracy meningkat menjadi 88,43%, precision menjadi 92,31%, recall menjadi 92,31%, dan F1-Score menjadi 92,31%. Peningkatan tersebut menunjukkan bahwa penerapan SMOTE mampu meningkatkan kemampuan model dalam mengklasifikasikan data, terutama ditinjau dari nilai accuracy, precision, dan F1-Score. Meskipun nilai recall mengalami penurunan dari 100,00% menjadi 92,31%, nilai tersebut masih tergolong tinggi dan diimbangi dengan peningkatan precision, sehingga hasil klasifikasi menjadi lebih seimbang. Secara keseluruhan, penerapan SMOTE memberikan peningkatan performa klasifikasi dibandingkan model yang tidak menggunakan penyeimbangan data.



Gambar 2. Grafik Hasil Pengujian

KESIMPULAN

Dari seluruh tahapan yang telah dilaksanakan, beberapa kesimpulan dapat ditarik sebagai berikut.

1. Berdasarkan hasil analisis sentimen terhadap ulasan pengunjung Wisata Situ Patenggang yang diperoleh dari Google Maps, dapat disimpulkan bahwa sebagian besar ulasan yang diberikan pengunjung memiliki sentimen positif. Hal ini menunjukkan bahwa pengunjung secara umum merasa puas terhadap keindahan alam, suasana, serta pengalaman wisata yang diperoleh selama berkunjung ke Situ Patenggang. Proses analisis sentimen dilakukan melalui tahapan preprocessing, pelabelan menggunakan metode Lexicon Based, pembobotan TF-IDF, dan klasifikasi menggunakan algoritma Multinomial Naive Bayes.
2. Model klasifikasi sentimen menggunakan metode Multinomial Naive Bayes berhasil diterapkan pada data ulasan wisata Situ Patenggang yang diperoleh dari Google Maps. Pengujian model dilakukan dengan metode split validation menggunakan rasio data latih dan data uji sebesar 80:20. Hasil pengujian sebelum penerapan SMOTE menunjukkan nilai accuracy sebesar 85,12%, precision sebesar 83,49%, recall sebesar 100,00%, dan F1-Score sebesar 91,00%. Setelah dilakukan penyeimbangan data menggunakan metode SMOTE, performa model meningkat dengan memperoleh accuracy sebesar 88,43%, precision sebesar 92,31%, recall sebesar 92,31%, dan F1-Score sebesar 92,31%. Hasil tersebut menunjukkan bahwa penerapan SMOTE mampu meningkatkan kinerja model Multinomial Naive Bayes dan menghasilkan klasifikasi yang lebih seimbang dalam mengenali kedua kelas sentimen.

Saran

Berdasarkan pelaksanaan dan hasil dari penelitian ini, terdapat beberapa hal yang dapat dijadikan pertimbangan untuk pengembangan selanjutnya.

1. Penelitian selanjutnya dapat menggunakan jumlah data yang lebih banyak dan berasal dari berbagai platform ulasan wisata seperti X, Tripadvisor, Instagram, maupun platform media sosial lainnya agar hasil analisis sentimen menjadi lebih representatif.
2. Pelabelan data dapat dilakukan secara manual oleh beberapa anotator atau menggunakan metode hybrid antara Lexicon Based dan pelabelan manual untuk meningkatkan kualitas label sentimen.



DAFTAR PUSTAKA

- Abbas, M., Memon, K. A., Jamali, A. A., Memon, S., & Ahmed, A. (2019). *Multinomial Naive Bayes Classification Model for Sentiment Analysis*. 19(3), 62–67.
- Amaniputra, M. E., Soesanto, R. P., & ... (2023). Perancangan Dashboard Sentimen Pariwisata Menggunakan Metode Complement Naïve Bayes pada Kabupaten Rembang. *EProceedings* ..., 10(3), 2452–2461. <https://openlibrarypublications.telkomuniversity.ac.id/index.php/engineering/article/view/20362%0Ahttps://openlibrarypublications.telkomuniversity.ac.id/index.php/engineering/article/view/20362/19675>
- Dewi, M. T., Herdiani, A., & Kusumo, D. S. (2020). Multi-Aspect Sentiment Analysis Komentar Wisata TripAdvisor dengan Rule-Based Classifier (Studi Kasus : Bandung Raya). *E-Proceeding of Engineering*, 5(1), 1589–1596.
- Ginantra, N. L. W. S. R., Yanti, C. P., Prasetya, G. D., Sarasvananda, I. B. G., & Wiguna, I. K. A. G. (2022). Analisis Sentimen Ulasan Villa di Ubud Menggunakan Metode Naive Bayes, Decision Tree, dan K-NN. *Jurnal Nasional Pendidikan Teknik Informatika (JANAPATI)*, 11(3), 205–215. <https://doi.org/10.23887/janapati.v11i3.49450>
- Harsemadi, I. G., Dharmendra, I. K., & Wijaya, I. M. P. P. (2023). Klasifikasi emosi pada tweet berbahasa Indonesia. *Prosiding Seminar*, 86.
- Lesmana, L., & Nabyla, F. (2020). *ANALISIS SENTIMEN PENGGUNA TWITTER PPDB MENGGUNAKAN ALGORITMA*. 1(1).
- Liu, B. (2012). *Sentiment Analysis and Opinion Mining*. May.
- Manullang, O., Prianto, C., & Harani, N. H. (2023). Analisis Sentimen Untuk Memprediksi Hasil Calon Pemilu Presiden Menggunakan Lexicon Based Dan Random Forest. *Jurnal Ilmiah Informatika*, 11(02), 159–169. <https://doi.org/10.33884/jif.v11i02.7987>
- Mujahidin, S., Hasyim, M. N., & Pratama, B. M. (2022). Implementasi Analisis Sentimen Opini Publik Mengenai Sirkuit Internasional Mandalika Pada Twitter Menggunakan Metode Multinomial Naïve Bayes Classifier. *Bianglala Informatika*, 10(2), 129–136. <https://doi.org/10.31294/bi.v10i2.13544>
- Nurhopipah, A., & Magnolia, C. (2023). Perbandingan Metode Resampling Pada Imbalanced Dataset Untuk Klasifikasi Komentar Program Mbkm. *Jurnal Publikasi Ilmu Komputer Dan Multimedia*, 2(1), 9–22. <https://doi.org/10.55606/jupikom.v2i1.862>
- Rahma, F., Farmadiansyah, A. Z., & Hidayatullah, A. F. (2021). *Deteksi Surel Spam dan Non Spam Bahasa Indonesia Menggunakan Metode Naïve Bayes*. 1–5.
- Rifki, M., Imelda, I., Informatika, T., Informasi, F. T., Luhur, U. B., Bayes, M. N., & Borobudur, C. (2022). *BOROBUDUR MENGGUNAKAN MULTINOMIAL NAÏVE BAYES ANALYSIS OF DISCOURSE SENTIMENT OF BOROBUDUR TEMPLE TICKET*. 5(2), 156–163. <https://doi.org/10.33387/jiko>
- Septian, J. A., Fachrudin, T. M., & Nugroho, A. (2019). Analisis Sentimen Pengguna Twitter Terhadap Polemik Persepakbolaan Indonesia Menggunakan Pembobotan TF-IDF dan K-Nearest Neighbor. *Journal of Intelligent System and Computation*, 1(1), 43–49. <https://doi.org/10.52985/insyst.v1i1.36>
- Singgalen, Y. A. (2022). Analisis Sentimen Wisatawan Melalui Data Ulasan Candi Borobudur di Tripadvisor Menggunakan Algoritma Naïve Bayes Classifier. *Building of Informatics, Technology and Science (BITS)*, 4(3). <https://doi.org/10.47065/bits.v4i3.2486>