

IMPLEMENTASI ANALISIS SENTIMEN MASYARAKAT MENGENAI KENAIKAN HARGA BBM DENGAN METODE NAIVE BAYES

Ardiansyah¹⁾, Nur'aini²⁾

^{1,2} Program Studi Informatika, Fakultas Ilmu Komputer, Universitas Amikom Yogyakarta, Jl. Ring Road
Utara, Condong Catur, Sleman, Yogyakarta

email: *1ardiansyah.05@students.amikom.ac.id, 2nuraini@amikom.ac.id

Abstract

With the increase in fuel prices society reaps many opinions of the pros and cons. This means the public's statement on the government's policy must be following the situation on the ground. Regardless of the cause of the increase in fuel prices must adjust to world oil prices. Twitter is one of the favorite platforms that people use to interact, especially in expressing opinions, an aspirations and discussing certain topics. Sentiment analysis is research related to a problem that creates a perception of the problem. Aims to classify texts in categorizing positive or negative. This research uses the naive bayes algorithm with a dataset in this study of 4239 data originating from Twitter. Through the processing stage, it is reduced to 2312 datasets which will be labeled manually resulting in 1014 positive, 920 neutral and 384 negative. The research produced a predictive dataset that totaled 1210 positives, 813 neutrals and 294 negatives. To see the performance of the data distribution parameters using 2 schemes and schemes namely, 463 Training Data, and 1855 Test Data while the second scheme, 811 Training Data, 1507 Test Data. The results of the evaluation stage for the accuracy of data sharing are the same for the Manual Dataset 92% while the Prediction Dataset is 72%. For manual dataset measurements, the precision is 88%, 95% recall and 89% F1-score, while the prediction dataset produces 54% precision, 85% recall, and 53% F1-score.

Keywords: Sentiment Analysis, Naïve Bayes, Fuel Prices, Classification

Abstrak

Kenaikan harga BBM di tengah masyarakat banyak menuai pendapat pro dan kontra. Hal tersebut membuat statement masyarakat atas kebijakan pemerintah tersebut harus sesuai dengan keadaan di lapangan. Terlepas dari penyebab kenaikan harga BBM harus menyesuaikan harga minyak dunia. *twitter* merupakan salah satu platform favorit yang digunakan masyarakat dalam berinteraksi terutama dalam menyampaikan pendapat, aspirasi dan berdiskusi mengenai pembahasan tertentu. Analisis sentimen merupakan penelitian yang berkaitan dengan suatu masalah menimbulkan persepsi pada masalahnya. Bertujuan mengklasifikasi teks dalam mengkategorikan positif atau negatif. Penelitian ini menggunakan algoritma *naive bayes* dengan *dataset* pada penelitian ini sebanyak 4239 data yang berasal dari *twitter*. Melalui proses tahap *processing* berkurang menjadi 2312 *dataset* yang akan di labelkan manual menghasilkan 1014 positif, 920 Netral, dan 384 Negatif. Dalam penelitian menghasilkan *dataset* prediksi yang berjumlah 1210 positif 813 Netral dan 294 Negatif. Untuk melihat kinerja parameter pembagian data menggunakan 2 skema dan skema yakni, 463 Data Latih, dan 1855 Data Uji sedangkan skema kedua, 811 Data Latih, 1507 Data Uji. Hasil tahapan evaluasi akurasi pembagian data menghasilkan sama untuk *dataset* manual 92% sedangkan *dataset* prediksi 72%. untuk pengukuran *dataset* manual *presisi* 88%, *recall* 95% dan 89% *f1-score* sedangkan untuk *dataset* prediksi menghasilkan 54% *presisi*, 85% *recall*, dan 53% *f1-score*.

Kata kunci: Analisis Sentimen, Naïve Bayes, Harga BBM, Klasifikasi

1. PENDAHULUAN

Kenaikan harga BBM banyak menuai pendapat yang berbagai macam dari kalangan masyarakat. Oleh sebab itu, Bahan Bakar Minyak (BBM) menjadi komoditas utama dalam aktivitas berkendara. Aktivitas kebanyakan dari masyarakat lebih menggunakan kendaraan pribadi. Dalam penelitian mengangkat masalah setelah kebijakan pemerintah menaikkan harga bahan bakar minyak pada bulan September 2022. ada beberapa faktor melatarbelakangi

pemerintah mengenai kebijakan tersebut. Dimulai dari konflik antar negara, pengurangan subsidi untuk anggaran bahan bakar minyak dan sebagainya (Hrp & Aslami, 2022). Platform media sosial untuk masyarakat memungkinkan dalam mengutarakan pendapat secara publik yakni twitter. Alasan penelitian ini menggunakan platform media sosial berupa twitter yakni data berupa tweets bisa diakses publik dan sebagian masyarakat memberikan opini tanpa ada batasan dalam hal berpendapat.

Analisis sentimen atau *opinion mining* merupakan penelitian yang berkaitan dengan suatu masalah menimbulkan persepsi pada masalahnya. Analisis sentimen bertugas mengklasifikasi teks dalam mengkategorikan positif atau negatif (Lina et al., 2022). Klasifikasi tersendiri adalah suatu hasil dari penggambaran beberapa kelas dari dari penting (Riyanto, 2019).

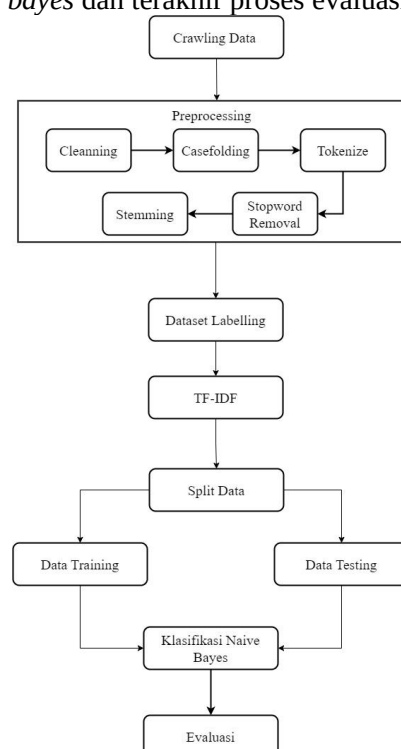
naive bayes adalah algoritma yang mengandalkan teori probabilitas dengan bertujuan menelusuri seberapa besar peluang dari pengelompokan klasifikasi berdasarkan melihat nilai dari penskoran pada data latih (Purwati, 2022). Algoritma tersebut beranggapan terdapat atribut objeknya bersifat independen. Hal ini probabilitas menghasilkan prediksi akhir dari hasil jumlah dari frekuensi berdasarkan tabel keputusan (Watratan et al., 2020). Kelebihan penggunaan algoritma *naive bayes* memiliki tingkat *error rate* lebih rendah dibandingkan algoritma klasifikasi lainnya (Marcellino et al., 2023).

Pada penelitian ini menggunakan 3 klasifikasi untuk menganalisis sentimen yang terdapat dalam *dataset* mengenai kenaikan harga BBM. Sehingga pada perhitungan menggunakan *confusion matrix* 3 kelas untuk mengidentifikasi hasil akhir akurasi, *presisi*, *recall* dan *f1-score*.

Penelitian ini bertujuan melihat analisis sentimen dari pendapat masyarakat yang membahas bbm di waktu yang bersamaan dengan pelabelan tertentu dengan melihat presentase positif, netral dan negatif. Penelitian ini untuk mengetahui penerapan nilai evaluasi analisis sentimen menggunakan metode *naive bayes classifier*.

2. METODE PENELITIAN

Pada tahapan penelitian ini menggunakan beberapa tahapan. dimulai dari pengambilan data kemudian dilanjutkan pada tahapan *Preprocessing*, Labelling Data, TF-IDF, proses klasifikasi menggunakan algoritma *naive bayes* dan terakhir proses evaluasi.



Gambar 1. Alur Penelitian

1. Crawling Data

Pemrosesan Crawling menggunakan software RapidMiner dengan bertujuan untuk mempermudah pengambilan data tanpa perlu menggunakan akses API langsung dari situs Twitter langsung dan menunggu persetujuan untuk mendapatkan API key dan secret key. Topik pembahasan pada opini masyarakat terhadap kenaikan harga BBM dengan pengambilan data dengan kata kunci “BBM”, “Bahan Bakar Minyak”.

2. Preprocessing

Pada tahapan ini melalui beberapa tahapan dari *Cleanning*, *Casefolding*, *Tokenize*, *Stopword Removal*, dan *Stemming*. Tahapan *Preprocessing* ini lebih membersihkan data yang tidak perlu seperti kata duplikat, kata ambigu dan mengubah kalimat menjadi kata dasar.

a. *Cleanning*

Tabel 1. Hasil Tahapan *Cleanning*

Text
Ada yang turun tapi bukan harga BBM atau Sembako ini hukuman buat koruptor Narasi Daily Hadir Membantu Rakyat Ditpolairud Polda Sumatera Utara Sumut telah mengungkap penyalahgunaan BBM di Perairan Laut Belawan Sumut Senin BersatuJagaNegeri JKMUKHERJEE Mau tampil keren dgn style anime kamu follow distroartonline SMS BBM BA F

Tahap pembersihan pertama untuk *dataset* seperti menghilangkan simbol, kata RT, kalimat duplikat, kata mengandung URL, dan simbol emoji.

b. *Casefolding*

Tabel 2. Hasil Tahapan *Casefolding*

Text
ada yang turun tapi bukan harga bbm atau sembako ini hukuman buat koruptor narasi daily hadir membantu rakyat ditpolairud polda sumatera utara sumut telah mengungkap penyalahgunaan bbm di perairan laut belawan sumut senin bersatujaganegeri jkmukherjee mau tampil keren dgn style anime kamu follow distroartonline sms bbm ba f

Pada tahap ini beberapa kalimat yang terdapat huruf kapital dirubah menjadi huruf kecil atau disebut dengan *lower uppercase*.

c. *Tokenize*Tabel 3. Hasil Tahapan *Tokenize*

Text
['ada', 'yang', 'turun', 'tapi', 'bukan', 'harga', 'bbm', 'atau', 'sembako', 'ini', 'hukuman', 'buat', 'koruptor', 'narasi', 'daily', '']
['hadir', 'membantu', 'rakyat', 'ditpolairud', 'polda', 'sumatera', 'utara', 'sumut', 'telah', 'mengungkap', 'penyalahgunaan', 'bbm', 'di', 'perairan', 'laut', 'belawan', 'sumut', 'senin', 'bersatujaganege', '']
['jkmukherjee', 'mau', 'tampil', 'keren', 'dgn', 'style', 'anime', 'kamu', 'follow', 'distroartonline', 'sms', 'bbm', 'ba', 'f', '']

Tokenize atau *Tokenizing* adalah proses pemisahan teks menjadi kata, symbol, maupun karakter sehingga menjadi kata yang dapat dianalisa (Syarifuddin, 2020). Dengan mengubah dari semula berupa *dataset* utuh berupa kalimat menjadi beberapa kata.

d. *Stopword Removal*Tabel 4. Hasil Tahapan *Stopword Removal*

Text
['turun', 'harga', 'bbm', 'sembako', 'hukuman', 'koruptor', 'narasi', 'daily', '']
['hadir', 'membantu', 'rakyat', 'ditpolairud', 'polda', 'sumatera', 'utara', 'sumut', 'mengungkap', 'penyalahgunaan', 'bbm', 'perairan', 'laut', 'belawan', 'sumut', 'senin', 'bersatujaganege', '']
['jkmukherjee', 'tampil', 'keren', 'dgn', 'style', 'anime', 'follow', 'distroartonline', 'sms', 'bbm', 'ba', 'f', '']

Stopword removal adalah pemrosesan kata menghilangkan kata imbuhan yang tidak perlu dan memilih kata yang bisa mewakili kalimat tersebut. Pada pemrosesan ini kata imbuhan seperti “di”, “ber”, “ke”, dan sebagainya. Dengan tujuan menjadi kata dasar dan sebelumnya sudah dipisah menjadi kata dasar yang tidak menjadi kalimat utuh (Paramita & Ibrahim, 2023).

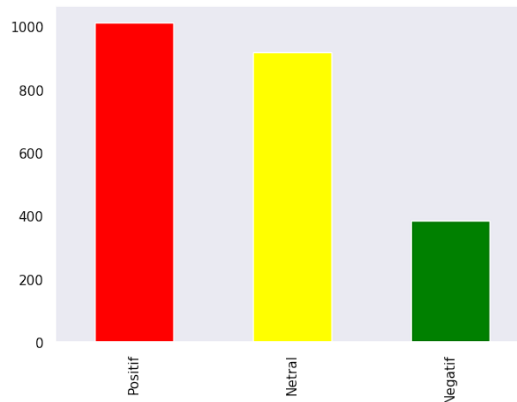
e. *Stemming*Tabel 5. Hasil Tahapan *Stemming*

Text
['turun', 'harga', 'bbm', 'sembako', 'hukum', 'koruptor', 'narasi', 'daily', '']
['hadir', 'bantu', 'rakyat', 'ditpolairud', 'polda', 'sumatera', 'utara', 'sumut', 'ungkap', 'penyalahgunaan', 'bbm', 'air', 'laut', 'belawan', 'sumut', 'senin', 'bersatujaganege', '']
['jkmukherjee', 'tampil', 'keren', 'dgn', 'style', 'anime', 'follow', 'distroartonline', 'sms', 'bbm', 'ba', 'f', '']

Stemming merupakan tahapan *Preprocessing* dengan mengembalikan yang sebelumnya kata dasar lebih dirubah ke dalam makna dasar (Yuyun et al., 2021). Dengan cara mengubah kata yang masih terdapat imbuhan menjadi kata dasar. Untuk penggunaan library pada proses yakni *sastrawi* berbahasa Indonesia yang berbasis algoritma *Nazief* dan *Adriani* (Pravina et al., 2019).

3. Dataset Labelling

Dari Hasil proses pada teks tersebut mendapatkan *dataset* sebanyak 2312 data. Dengan beragam sentimen dari kalangan masyarakat membuat beragam opini.



Gambar 2. Diagram Hasil Klasifikasi Labelling Data Manual

Dari pelabelan tersebut menghasilkan dominasi Sentimen Netral terlalu tinggi dibandingkan kategori Sentimen Negatif. Hal ini akan berpengaruh dalam penelitian yang berfokus dalam keterkaitan dengan kata bahan bakar minyak.

4. Term Frequency - Invers Document Frequency

Term Frequency merupakan metode pembobotan kata dengan berdasarkan dominasi kata (frekuensi) dalam sebuah satu kalimat. *Inverse Document Frequency* adalah kemunculan banyaknya jumlah kata dalam sebuah term yang muncul (Eko et al., 2019). TF-IDF adalah suatu proses ekstraksi dengan memberikan nilai pada tiap kata pada data yang sudah diklasifikasikan (Pravina et al., 2019).

$$W_{tf,d} = \begin{cases} 1 + \log_{10} tf_{t,d}, & \text{if } tf_{t,d} > 0 \\ 0, & \text{otherwise} \end{cases} \quad (1)$$

$$idf_{t,d} = \log_{10} \left(\frac{N}{df_t} \right) \quad (2)$$

$$W_{t,d} = W_{tf,d} \times idf_t \quad (3)$$

Keterangan

$W_{tf,d}$ = Bobot kata dalam setiap dokumen

$tf_{t,d}$ = jumlah frekuensi muncul kata t dalam dokumen d

N = jumlah seluruh dokumen

df = jumlah dokumen yang mengandung term

$idf_{t,d}$ = bobot *inverse* dalam df

$W_{t,d}$ = bobot TF-IDF

5. Algoritma Naïve Bayes

naive bayes merupakan algoritma klasifikasi menggunakan probabilitas (Heliyanti Susana, 2022). Berdasarkan nilai probabilitas dengan menjumlahkan frekuensi serta kombinasi nilai dari sebuah dokumen. *naive bayes* memiliki tingkat akurasi serta penerapan bisa ke dalam database kualitatif serta kuantitatif serta diterapkan *dataset* berjumlah banyak maupun sedikit.

Dibandingkan metode selain algoritma *naive bayes*, metode ini dapat meminimalkan kesalahan sehingga tingkat nilai akurasi lebih baik (Liberti et al., 2023). Pada proses

pengklasifikasi tersebut dengan sebelumnya pelabelan secara manual agar hasil pemrosesan tersebut berjalan maksimal sesuai hasil akurasi.

Metode rumus naïve bayes:

$$P(H|X) = \frac{P(X|H) \times P(H)}{P(X)}$$

(Ridwan, 2020)

Keterangan:

H = Hipotesis hasil dari Y lebih spesifik

X = Data yang belum diketahui

$P(H|X)$ = Probabilitas hipotesis H berdasarkan kondisi X (probabilitas posterior)

$P(H)$ = Probabilitas hipotesis H (probabilitas prior)

$P(X|H)$ = Probabilitas X berdasarkan kondisi tersebut (probabilitas prior)

$P(X)$ = Probabilitas X

6. Evaluasi

Pada tahap evaluasi memperlihatkan hasil akhir tersebut sesuai dengan tujuan penelitian di tetapkan. Dan akhir penelitian serta memberikan informasi serta acuan pemutakhiran data. Secara spesifik memperlihatkan beberapa penilaian yakni akurasi, *presisi*, *recall* dan *f1-score* dengan berdasarkan perhitungan *confusion matrix*.

Tabel 6. *Confusion matrix*

Manual	Sistem		
	Positif	Netral	Negatif
Positif	(TP)	(FNt)	(FN)
Netral	(FP)	(TNt)	(TN)
Negatif	(FP)	(FNt)	(TN)

(1) (Dhina Nur Fitriana & Yuliant Sibaroni, 2020)

$$Akurasi = \frac{TP+TNt+TN}{TP+TN+TNt+FNt+FP+FN} \times 100\% \quad (2)$$

$$Presisi = \frac{TP}{TP+FP} \quad (3)$$

$$Recall = \frac{TP}{TP+FNt+FN} \times 100\% \quad (4)$$

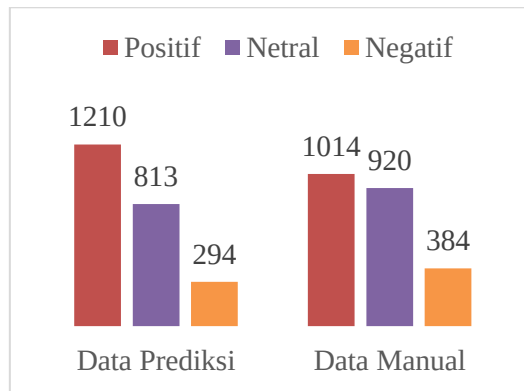
$$f1 - score = 2 \times \frac{Presisi \times Recall}{Presisi+Recall} \quad (5)$$

3. HASIL DAN PEMBAHASAN

Penelitian ini menggunakan *dataset* yang berasal dari sosial media *Twitter* kemudian dilanjutkan proses crawling dengan menggunakan *software RapidMiner* dan kemudian di analisis menggunakan bahasa pemrograman *python* dengan *text editor Google Colab*. dari pengumpulan data tersebut menghasilkan *dataset* berjumlah 4239 data. Tetapi setelah melalui tahapan preprocessing menghasilkan *dataset* berjumlah 2319 data.

3.1 KLASIFIKASI DATA PREDIKSI

Sebelum mencari hasil prediksi nilai evaluasi. Pada penelitian ini memprediksi tabel klasifikasi terhadap data yang sudah dilabelkan secara manual. Pada tahapan ini tidak membagi data tersebut melainkan dengan skenario pembobotan nilai TF-IDF dengan *Dataset* Manual secara Keseluruhan. Dari hasil tersebut *dataset* prediksi pada kolom klasifikasi prediksi menghasilkan data yang tidak seimbang. Disebabkan beberapa nilai pembobotan TF-IDF tidak optimal.



Gambar 3. Grafik Perbandingan Data prediksi dengan Data Manual

3.2 PEMBAGIAN DATA

Kemudian untuk mencari nilai pada tahapan evaluasi dengan menggunakan dua skenario pembagian data untuk mencari nilai akurasi jika penggunaan data latih semakin banyak maka nilai tersebut semakin meningkat.

Tabel 7. Skenario Pembagian Data

	Data Latih	Data Uji
Skenario 1	463	1855
Skenario 2	811	1507

3.3 HASIL EVALUASI

Dari pembagian data tersebut menghasilkan nilai akurasi yang sama antara skenario 1 dengan skenario 2. Dengan penerapan *Dataset* Manual lebih tinggi dibandingkan *Dataset* prediksi hasil dari klasifikasi Naïve Bayes. Hal ini dikarenakan dominasi kategori klasifikasi terlalu tinggi sehingga nilai akurasi dengan melakukan pembagian data tersebut menghasilkan nilai akurasi yang sama.

Tabel 8. Hasil Akurasi Masing-Masing *Dataset*

Pembagian Data	Akurasi	
	Manual	prediksi
Skenario 1	92%	72%
Skenario 2	92%	72%

Kemudian untuk akurasi masing-masing kategori klasifikasi menghasilkan nilai positif lebih rendah karena banyaknya data bersentimen positif terlalu banyak dibandingkan kategori klasifikasi nilai negatif. *Dataset* prediksi nilai positif lebih rendah dibandingkan *Dataset* Manual dengan jarak sekitar 20%. Sedangkan nilai negatif pada *dataset* prediksi berjarak sekitar 7%. Sedangkan netral berjarak nilai akurasi pada *dataset* prediksi dengan *dataset* manual yakni 9%.

Tabel 9. Hasil Akurasi Masing-Masing Klasifikasi

Dataset	Pembagian Data	Akurasi		
		Positif	Netral	Negatif
Prediksi	Skenario 1	73,8%	85,2%	87,1%
	Skenario 2	73,5%	85,5%	87,1%
Manual	Skenario 1	93,5%	94,1%	94,1%
	Skenario 2	93,4%	94,2%	95,2%

Untuk hasil nilai *presisi* pada prediksi cukup rendah karena tidak adanya fitur tahapan keseimbangan data sehingga penerapan klasifikasi terdapat *error rate* serta beberapa dokumen terjadi kesalahan klasifikasi. Untuk *recall* nilai cukup baik antara *dataset* manual dengan *dataset* prediksi. Kemudian pada nilai *f1-score* melibatkan nilai *presisi* dan *recall* sehingga nilai pada *dataset* prediksi rendah.

Tabel 10. Hasil Perhitungan *Presisi*, *Recall* dan *F1-score*

	Manual	Prediksi
<i>Presisi</i>	88%	54%
<i>Recall</i>	95%	85%
<i>F1-score</i>	89%	53%

4. SIMPULAN

Dari hasil klasifikasi *dataset* prediksi hasil pada klasifikasi positif mengalami peningkatan cukup signifikan dibandingkan klasifikasi kategori Netral dan Negatif mengalami penurunan. Hal ini pada *dataset* penelitian ini tidak seimbang atau imbalanced data antara kategori sentimen positif, negatif dan netral. Mengenai tentang hasil klasifikasi tersebut pada tahapan preprocessing perlu diubah kembali karena pada *dataset* tersebut kata, huruf atau simbol terdapat kesalahan terutama kehilangan yang tadinya kata tersebut kata dasar hilang huruf salah satu sehingga kata dasar tidak utuh sehingga pada proses TF-IDF tidak optimal karena pencocokan kata yang sama tetapi kehilangan huruf maka pada sistem tersebut tidak dijumlahkan sehingga pembobotan nilai kata tidak optimal. Hal tersebut mengganggu pada proses tahapan selanjutnya.

Hasil dari penelitian mengenai analisis sentimen dengan klasifikasi menggunakan algoritma *naive bayes* dengan menggunakan pembagian data, pada penelitian ini tidak mempengaruhi signifikan terhadap nilai tingkat akurasi dengan *dataset* yang sama. Hasil akurasi dari *dataset* manual dengan *dataset* prediksi berjarak sekitar 20% yakni 92% akurasi *dataset* manual dan 72% akurasi *dataset* prediksi. Dan juga berpengaruh pada nilai *presisi*, *recall* dan *f1-score*. Terutama pada nilai *presisi* dan *f1-score* pada *dataset* prediksi.

Untuk penelitian kedepan, disarankan menggunakan pencarian data menggunakan variasi kata kunci mengenai topik, menyederhanakan pada tahap preprocessing serta penggunaan perbandingan algoritma klasifikasi seperti SVM, *Random Forest*, *Decision Tree*.

5. DAFTAR PUSTAKA

- Dhina Nur Fitriana, & Yuliant Sibaroni. (2020). Sentiment Analysis on KAI Twitter Post Using Multiclass Support Vector Machine (SVM). *Jurnal RESTI (Rekayasa Sistem Dan Teknologi Informasi)*, 4(5), 846–853. <https://doi.org/10.29207/resti.v4i5.2231>
- Eko, P., Utomo, P., Khaira, U., Suratno, T., & Jambi, U. (2019). Analisis Sentimen Online Review Pengguna Bukalapak Menggunakan Metode Algoritma TF-IDF. (*JUSS*) *Jurnal Sains Dan Sistem Informasi*, 2(2), 35–39.
- Heliyanti Susana. (2022). Penerapan Model Klasifikasi Metode Naive Bayes Terhadap Penggunaan Akses Internet. *Jurnal Riset Sistem Informasi Dan Teknologi Informasi (JURSISTEKNI)*, 4(1), 1–8. <https://doi.org/10.52005/jursistekni.v4i1.96>
- Hrp, G. R., & Aslami, N. (2022). Analisis Dampak Kebijakan Perubahan Publik Harga BBM terhadap Perekonomian Rakyat Indonesia. *Jurnal Ilmu Komputer, Ekonomi Dan Manajemen (JIKEM)*, 2(1), 1464–1474.
- Liberti, A., Tavares, D., Nurraharjo, E., Informatika, S. T., & Semarang, U. S. (2023). *Tweet Terkait Naiknya Kasus Omicron Menggunakan Naive Bayes Classifier*. 6(April), 1–8.
- Lina, S., Sitio, M., Studi, P., Informatika, T., Pamulang, U., & Selatan, B. T. (2022). *PADA MEDIA SOSIAL TWITTER MENGGUNAKAN PHYTON*. XVII(02), 57–61.
- Marcellino, M., Marlim, Y. N., & ... (2023). Aplikasi Pendeteksi Penyakit Hepatitis Menggunakan <https://doi.org/10.35145/joisie.v8i1.3838>

- Metode Naïve Bayes. *JOISIE (Journal Of ...*, 7(1), 155–164. <https://www.ejournal.pelitaindonesia.ac.id/ojs32/index.php/JOISIE/article/view/3298%0Ahttps://www.ejournal.pelitaindonesia.ac.id/ojs32/index.php/JOISIE/article/download/3298/1161>
- Paramita, P., & Ibrahim, A. (2023). Analisis Sentimen Terhadap Pengguna Qris (Quick Respond Code Indonesian Standart) Pada Twitter Menggunakan Metode Naïve Bayes Classifier. *JOISIE Journal Of Information System And Informatics Engineering*, 7(1), 1–6. <https://t.co/lJemg7TbKb>
- Pravina, A. M., Cholissodin, I., & Adikara, P. P. (2019). Analisis Sentimen Tentang Opini Maskapai Penerbangan pada Dokumen Twitter Menggunakan Algoritme Support Vector Machine (SVM). *Jurnal Pengembangan Teknologi Informasi Dan Ilmu Komputer*, 3(3), 2789–2797. <http://j-ptiik.ub.ac.id>
- Purwati, N. (2022). Analisa Pengaruh Iklan tanpa Label Harga pada media sosial Menggunakan Algoritma Naive Bayes. *Infomatek: Jurnal Informatika, Manajemen Dan Teknologi*, 24(1), 45–50. <https://doi.org/10.23969/INFOMATEK.V24I1.4622>
- Ridwan, A. (2020). Penerapan Algoritma Naïve Bayes Untuk Klasifikasi Penyakit Diabetes Mellitus. *Jurnal SISKOM-KB (Sistem Komputer Dan Kecerdasan Buatan)*, 4(1), 15–21. <https://doi.org/10.47970/siskom-kb.v4i1.169>
- Riyanto, U. (2019). Analisis Perbandingan Algoritma Naive Bayes Dan Support Vector Machine Dalam Mengklasifikasikan Jumlah Pembaca Artikel Online. *JIKA (Jurnal Informatika)*, 2(2), 62–72. <https://doi.org/10.31000/.v2i2.1521>
- Syarifuddin, M. (2020). Analisis Sentimen Opini Publik Mengenai Covid-19 Pada Twitter Menggunakan Metode Naïved Bayes dan KNN. *INTI Nusa Mandiri*, 15(1), 23–28. <https://doi.org/10.33480/INTI.V15I1.1347>
- Watratana, A. F., B, A. P., & Moeis, D. (2020). Implementasi Algoritma Naive Bayes Untuk Memprediksi Tingkat Penyebaran Covid-19 Di Indonesia. *Journal of Applied Computer Science and Technology*, 1(1), 7–14. <https://doi.org/10.52158/JACOST.V1I1.9>
- Yuyun, Nurul Hidayah, & Supriadi Sahibu. (2021). Algoritma Multinomial Naïve Bayes Untuk Klasifikasi Sentimen Pemerintah Terhadap Penanganan Covid-19 Menggunakan Data Twitter. *Jurnal RESTI (Rekayasa Sistem Dan Teknologi Informasi)*, 5(4), 820–826. <https://doi.org/10.29207/resti.v5i4.3146>