

**ANALISIS SENTIMEN PUBLIK TERHADAP KONTROVERSI GUS ELHAM YAHYA
DIMEDIA SOSIAL TWITTER/X MENGGUNAKAN
PENDEKATAN LEXICON-BASED DAN NAIVE BAYES**

Reza Atreji¹, Zaehol Fatah², Syarif Aminil Khoiri³
Universitas Ibrahimy
rezaatreji19@gmail.com

ABSTRAK

Kontroversi video interaksi Gus Elham Yahya dengan anak kecil saat kegiatan dakwah memicu perdebatan sengit di media sosial Twitter/X. Publik terbagi antara kecaman atas perlindungan anak serta narasi tabayyun. Volume data yang besar membuat analisis manual sulit dilakukan. Studi ini bertujuan memetakan sentimen publik dan mengukur akurasi penggabungan metode berbasis leksikon serta algoritma Naive Bayes. Penelitian kuantitatif ini menggunakan sampel sebanyak 1.509 tweet berbahasa Indonesia yang dikumpulkan melalui teknik purposive sampling. Analisis data menerapkan kerangka Knowledge Discovery in Database yang meliputi tahapan pra-pemrosesan, ekstraksi fitur TF-IDF, pelabelan leksikon, dan klasifikasi Naive Bayes. Hasil analisis leksikon menunjukkan distribusi 9,15% positif, 27,24% netral, dan 63,62% negatif. Klasifikasi Naive Bayes mencapai tingkat akurasi tinggi sebesar 98,88% dengan performa terbaik pada kelas negatif. Dominasi sentimen negatif konsisten mencerminkan kritik tajam publik terhadap tindakan tokoh publik yang dianggap melanggar norma perlindungan anak. Temuan ini menegaskan bahwa perilaku figur publik berdampak signifikan terhadap kepercayaan masyarakat di ruang digital. Implikasi praktisnya, figur publik wajib menjaga etika interaksi serta memiliki strategi komunikasi krisis yang efektif guna memulihkan reputasi.

Kata Kunci: Analisis Sentimen, Twitter/X, Lexicon Based, Naive Bayes, Gus Elham Yahya

ABSTRACT

The controversy over the video of Gus Elham Yahya's interaction with small children during a da'wah activity that prompted fierce debate on Twitter/X social media. The public is divided between criticism of child protection and tabayyun narrative. The large volume of data makes manual analysis difficult. This study aims to map public sentiment and measure the accuracy of combining lexicon-based methods and the Naive Bayes algorithm. This quantitative research uses a sample of 1,509 Indonesian-language tweets collected through a purposive sampling technique. Data analysis applies the Knowledge Discovery framework in the Database which covers pre-processing stages, TF-IDF feature extraction, lexicon labeling, and Naive Bayes classification. The results of the lexicon analysis show a distribution of 9.15% positive, 27.24% neutral, and 63.62% negative. Naive Bayes classification achieves a high level of accuracy of 98.88% with the best performance on the negative class. The dominance of negative sentiment consistently reflects the public's sharp criticism of the actions of public figures that violate child protection norms. This finding confirms that the behavior of public figures has a significant impact on public trust in the digital space. The practical implication is that public figures must maintain interaction ethics and have an effective crisis communication strategy to restore reputation.

Keywords Sentiment Analysis, Twitter/X, Lexicon Based, Naive Bayes, Gus Elham Yahya

Pendahuluan

Di era kemajuan teknologi saat ini, berbagai sisi kehidupan masyarakat mengalami peningkatan pesat, termasuk hadirnya internet yang berperan sebagai sarana pertukaran informasi secara global [1]. Kemudahan mengakses informasi secara cepat dan instan sangat terasa di era digital ini. Hampir setiap orang di Indonesia terhubung dengan internet lewat berbagai perangkat, mulai dari ponsel, PC, hingga perangkat lainnya [2].

Twitter, sebagai salah satu media sosial paling populer, berperan sebagai wadah komunikasi masyarakat. Platform ini memungkinkan siapa saja di seluruh dunia terhubung dengan keluarga, teman, dan kerabat lewat komputer atau ponsel. Fitur utamanya adalah membuat *tweet* (pesan status) berisi opini pengguna tentang berbagai topik dengan batas 140 karakter. Hal ini membuat *Twitter/X* menjadi platform penyedia kumpulan data opini masyarakat dunia [3]. Di tahun 2022, *Twitter/X* mempunyai 220 juta pengguna aktif di dunia dan Indonesia tercatat sebagai negara pengguna *Twitter/X* ke 5 terbanyak di dunia dengan jumlah pengguna aplikasi ini mencapai 18,45 juta. Selain itu, terdapat 500 juta tweet yang di unggah setiap harinya. Berisi opini dan sentimen dari pengguna terhadap suatu topik atau kejadian yang menarik perhatian, tidak heran jumlah konten atau video yang di buat status ini terus meningkat secara drastis selama beberapa tahun kedepan [4]. Twitter menghasilkan interaksi dan data dalam jumlah besar, menjadikannya lahan subur untuk analisis sentimen. Tujuannya adalah untuk menangkap persepsi publik secara komprehensif dan terkini. Platform media sosial ini memungkinkan pengguna untuk menyuarakan berbagai pandangan—baik dukungan, keberatan, atau kekhawatiran tentang berbagai isu penting. Oleh karena itu, analisis sentimen menghadirkan alat yang

ampuh untuk memproses banyaknya teks tidak terstruktur ini dan mengkategorikan opini publik ke dalam kategori positif, negatif, atau netral [5].

Salah satu metode yang sering digunakan adalah pendekatan berbasis *lexicon Based*, di mana kumpulan kata atau istilah dengan makna sentimen yang diketahui digunakan untuk memberikan skor positif atau negatif pada setiap fitur kata. Pendekatan ini bekerja dengan mengandalkan daftar kata atau frasa yang terkait dengan sentimen tertentu. Metode ini sangat efektif untuk klasifikasi sentimen secara real-time, terutama karena fleksibilitasnya dalam beradaptasi dengan dinamika dan perubahan konten media sosial [6]. Dalam penelitian ini, algoritma *Naïve Bayes* dipilih sebagai metode klasifikasi. Algoritma ini beroperasi berdasarkan prinsip probabilitas dan dikenal karena efisiensi dan presisinya yang unggul, terutama untuk data teks yang kompleks dan beragam. Keunggulan lainnya adalah *Naïve Bayes* mudah diimplementasikan dan tetap andal bahkan saat menggunakan dataset besar. Karakteristik ini membuatnya sangat cocok untuk analisis sentimen di platform media sosial seperti Twitter. *Naïve Bayes* telah terbukti sebagai pendekatan yang sangat sukses ketika diterapkan pada tugas analisis sentimen. Inti dari metode ini adalah membangun model probabilitas yang menghitung kemungkinan sebuah kata atau frasa muncul di setiap kategori sentimen, seperti positif, negatif, atau netral. Agar berfungsi, model ini membutuhkan data pelatihan yang terdiri dari teks dengan label sentimen. Dari data ini, *Naïve Bayes* akan belajar mengenali pola dan asosiasi antara kata-kata dan kecenderungan sentimen tertentu. Setelah proses pembelajaran selesai, model yang terlatih siap digunakan untuk mengklasifikasikan teks baru ke dalam kelas sentimen yang paling relevan [7].

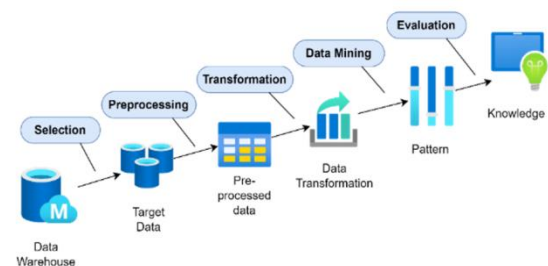
Sebuah video yang menunjukkan Gus Elham Yahya berinteraksi dengan seorang balita selama dakwahnya baru-baru ini telah memicu gelombang kritik yang signifikan. Sebagian besar orang percaya bahwa tindakan tersebut telah melampaui batas dan berpotensi melanggar norma perlindungan anak, yang pada akhirnya memicu rasa cemas kolektif di kalangan orang tua dan profesional hukum. Kompleksitas masalah ini diperparah oleh munculnya dua kubu yang berlawanan. Di satu sisi, ada kelompok yang mengatakan itu tidak pantas dilakukan seorang pendawah (sentimen negatif), menekankan perlindungan anak dan menjunjung tinggi otoritas lembaga keagamaan. Sementara itu, di sisi lain, ada kelompok yang mendukung (sentimen positif), menekankan pentingnya *tabayyun* (pemeriksaan fakta), dan menganggap insiden tersebut sebagai *kehilafan* manusiawi. Perbincangan di media sosial tentang insiden ini sangat ramai. Akibatnya, terkumpul data teks dalam jumlah yang sangat besar dan tidak beraturan (tidak terstruktur). Data sebanyak itu tidak mungkin dianalisis hanya dengan membaca satu per satu secara manual. Oleh karena itu, penelitian ini menjadi sangat penting. Tujuannya adalah untuk mengetahui sejauh mana perilaku seorang tokoh publik seperti Gus Elham Yahya dapat mempengaruhi kepercayaan masyarakat.

Penelitian ini bertujuan untuk analisis sentimen pada opini publik di *Twitter/X* terkait kontroversi Gus Elham Yahya. Algoritma *Naïve Bayes* akan digunakan untuk mengklasifikasikan *tweet* ke dalam tiga kategori sentimen, positif, negatif, dan netral. Dengan bantuan pelabelan *Lexicon Based* hasilnya diharapkan dapat memetakan sejauh mana perilaku tokoh publik memengaruhi kepercayaan publik.

Materi dan Metode

Jenis penelitian yang digunakan Adalah kuantitatif dengan metode analisis deskriptif. Hal ini dikarenakan penelitian berfokus pada pengolahan data numerik probabilitas dan frekuensi untuk menggambarkan fenomena opini masyarakat. Selain itu, penelitian ini juga bersifat *Eksperimental Laboratorium* karena melibatkan pengujian algoritma pada dataset tertentu.

Dalam penelitian ini, teknik pengumpulan data yang diterapkan adalah penemuan pengetahuan dalam basis data yaitu *Knowledge Discovery in Database* (KDD). Langkah-langkah yang termasuk dalam metode KDD meliputi, pemilihan data, pembersihan data, transformasi data, eksekusi penambangan, dan evaluasi model yang dihasilkan [8]. Seperti gambar berikut.



Gambar : 1 Langkah *Knowledge Discovery in Database* (KDD).

1. Data Selection

Ini adalah tahap Pada fase pemilihan data KDD, tiga atribut dipilih dari 12 atribut yang tersedia dalam dataset tweet: teks tweet, label sentimen, dan waktu posting. Proses ini mengekstrak data yang dibutuhkan untuk analisis sentimen dan membuang atribut yang tidak relevan [9].

2. Preprocessing

Dalam pengolahan data, tahap preprocecing menempati posisi paling krusial demi mengoptimalkan luaran dari analisis sentimen. Proses ini bertugas membersihkan data dari berbagai cacat seperti duplikasi dan nilai yang kosong . hasil akhirnya berupa data

yang tertera rapi dan terstruktur, sehingga informasi yang dikasih menjadi akurat serta proses analisis data berjalan lebih lancar [10].

3. Transformation

Pada fase transformasi, representasi atau bentuk data diubah agar lebih relevan dan menyederhanakan proses analisis. Salah satu teknik yang digunakan adalah pengkodean variabel kategorikal, yang mengubah data kategorikal menjadi format yang dapat dipahami oleh algoritma. Transformasi ini pada akhirnya bertujuan untuk memperkuat kemampuan algoritma penambangan data dalam menemukan pola atau wawasan yang bermakna dari data yang telah melalui tahap persiapan [11].

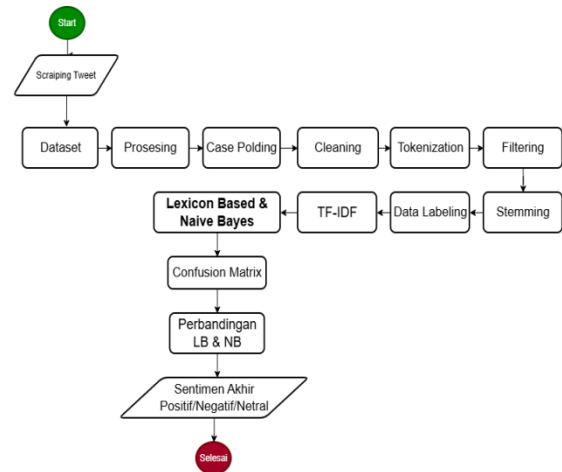
4. Data Mining

Yang dimaksud dengan data mining adalah serangkaian proses untuk mengelola data dalam skala besar sedemikian rupa sehingga data tersebut dapat menghasilkan informasi yang akurat. Informasi yang diperoleh kemudian dapat dimanfaatkan untuk mempermudah penyelesaian masalah maupun proses pengambilan keputusan [10].

5. Evaluation

Hasil dari fase data mining, dalam bentuk model data, dapat disajikan dalam format yang lebih mudah dipahami oleh para pemangku kepentingan. Tujuan dari fase ini adalah untuk memverifikasi apakah kebijakan atau informasi yang dihasilkan bertentangan dengan fakta atau informasi lain yang telah diketahui sebelumnya [12].

Adapun perencanaan Sistem dapat dilihat pada Gambar 2:



Gambar : 2 Perancangan Sistem

Alur ini merupakan proses komprehensif untuk menganalisis sentimen (opini) dari tweet yang dikumpulkan dari media sosial Twiter/X. Tujuannya adalah untuk mengkategorikan setiap tweet ke dalam salah satu dari tiga kategori: Positif, Negatif, atau Netral.

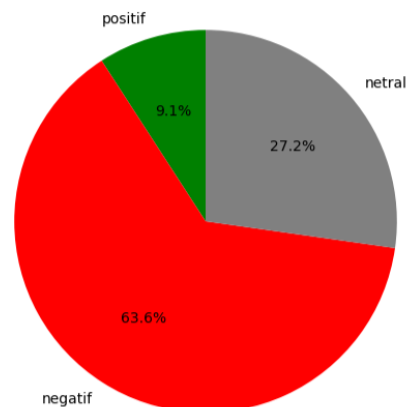
Hasil dan Pembahasan

Hasil

Adapun hasil dari penelitian ini perencanaan sistem ini dapat dari gambar dibawah ini:

a. Lexicond Based

Distribusi Sentimen Akhir (Lexicon)



Gambar : 3 Hasil Dari Pelabelan Lexicond Based

Ini adalah hasil distribusi sentimen dari data tweet yang dianalisis. Data disajikan dalam

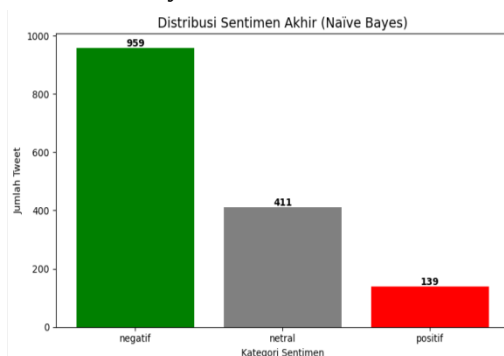
bentuk tabel frekuensi yang menunjukkan persentase setiap kategori sentimen.

Tabel : 1 Distribusi LExicond

Kategori	Jumlah Data	Presentase Sentimen
Negatif	960	63, 618290
Netral	411	27, 236581
Positif	138	9, 145129

Tabel tersebut menampilkan hasil akhir klasifikasi sentimen dari data tweet yang telah dianalisis. Data disajikan dalam bentuk tabel yang memuat jumlah data (frekuensi) dan persentase sentimen untuk masing-masing kategori.

b. Naïve Bayes

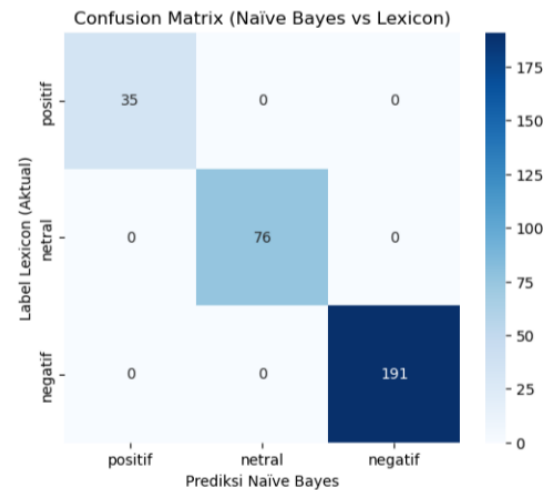


Gambar : 4 Hasil Dari Algoritma Naive Bayes

Ini adalah hasil distribusi sentimen akhir dari proses klasifikasi yang dilakukan menggunakan metode Naïve Bayes. Data disajikan dalam bentuk tabel frekuensi yang menunjukkan jumlah tweet untuk setiap kategori sentimen.

c. Kofusion Matrix dari Naive Bayes dan Lexicond Based

Confusion Matrix ini menunjukkan bahwa model Naïve Bayes yang digunakan memiliki performa (100% akurasi) dalam mengklasifikasikan 302 data tweet ke dalam tiga kategori sentimen (Positif, Netral, Negatif).

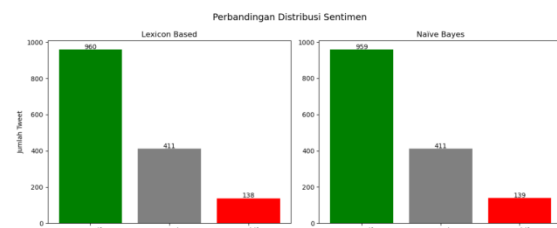


Gambar : 5 Confucion Matrix

Tabel : 2 Perhitungan Performa

Keterangan	Precision	recall	F1-score	support
Negatif	1.00	1.00	1.00	191
Netral	1.00	1.00	1.00	76
Positif	1.00	1.00	1.00	35

d. Perbandingan Lexicond Based dan Naive Bayes



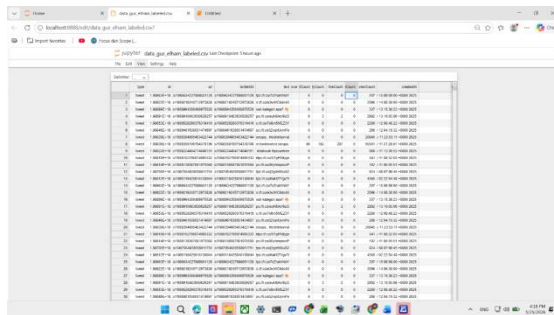
Gambar : 6 Hasil Perbandingan Lexicond dan Naive Bayes

Kedua metode (berbasis leksikon dan Naive Bayes) menunjukkan kinerja yang hampir identik dalam klasifikasi sentimen. Perbedaannya hanya terletak pada satu titik data (0,07% dari total), yang menunjukkan bahwa kedua metode tersebut berkinerja sama baiknya pada dataset ini.

Pembahasan Dataset

Penelitian ini menggunakan data tweet berbahasa Indonesia yang diambil dari platform Twitter/X dengan kata kunci

"kontroversi Gus Elham Yahya" dan "Gus Elham mencium anak kecil". Proses scraping menggunakan Apify Actor menghasilkan 1.509 tweet setelah pembersihan data awal (menghapus duplikat dan nilai kosong). Data kemudian melalui tahap pra-pemrosesan yang meliputi case folding, pembersihan (menghapus penyebutan, URL, tanda baca), tokenisasi, penghapusan stopword, dan stemming menggunakan Sastrawi. Hasil pra-pemrosesan akhir adalah teks dalam bentuk stemmed yang siap untuk dianalisis.



Gambar : 7 Hasil Scriping dari twiterr/X

Ini adalah hasil dari scrapping dari platform Twitter/X berupa data mentah yang belum di proses. Data yang dikumpulkan terdiri dari berbagai jenis konten. tweet teks biasa, tweet dengan gambar, tweet dengan video, dan tweet kutipan, tweet yang mengutip tweet lain. Beberapa tweet juga berisi jejak pendapat dan tautan berita. Keragaman format ini menunjukkan bahwa diskusi publik tentang kontroversi Gus Elham Yahya meluas melampaui opini langsung hingga mencakup konten multimedia dan berita. Semua data kemudian menjalani pra-pemrosesan (pembersihan, normalisasi, dan stemming) sebelum dianalisis sentimennya.

Selection

Tahap seleksi merupakan langkah awal dalam metodologi Knowledge Discovery in Database (KDD), yang bertujuan untuk mengurutkan dan memilih atribut (kolom) yang relevan dari seluruh dataset mentah, sambil mengabaikan atribut yang tidak diperlukan untuk penelitian. Berdasarkan hasil scraping menggunakan Apify Actor.

[illegible]

Gambar: Jumlah data yang didapatkan dari scripting

Tahap ini ada memilih data yang sekiranya layak untuk di proses selanjutnya atau data yang relvan untuk di proses ketahap selanjutnya.

Propresing

Setelah melakukan scraping menggunakan Apify Actor, diperoleh 1.509 tweet sebagai data mentah. Data tersebut kemudian menjalani serangkaian langkah pra-pemrosesan, termasuk konversi huruf besar/kecil, pembersihan (menghapus penyebutan, URL, tanda baca, dan angka), tokenisasi, penghapusan stopword, dan stemming menggunakan pustaka Sastrawi. Hasil dari setiap langkah disimpan dalam kolom terpisah di dalam dataframe untuk memudahkan analisis.



Gambar: Tahap Prossing

a. Case Folding

adalah proses mengganti huruf ke huruf kecil semua. Langkah ini bertujuan untuk menyeragamkan kata sehingga kata yang sama dengan variasi huruf kapital di anggap identik.

Tabel: hasil dari tahap prosesing

Sebelum Stemming	Setelah Stemming
Kontroversi Gus Rlham Yahva mencium anak kecil	kontroversi gus elham

	yahya <i>cium</i> anak kecil
Ini bukanlh mendidik orang kejalan yang benar	Ini <i>bukan didik</i> orang <i>jalan</i> yang benar

b. Cleaning

Bertujuan untuk meghilangkan elemen elemen yang tidak di gunakan atau tidak dibutuhkan. Seperti, URL, emoji, simbol simbol, tanda tanya dan sebagainya inti dari clening ini adah pembersihan komentar yang tidak di butuhkan.

c. Tokenisasi

Tahap ini adalah memisah kata perkata atau memecah kata satu satu seperti "gus elham mencium anak kecil" akan menjadi seperti ('gus', 'elham', 'yahya', 'mencium', 'anak', 'kecil').

d. Stopword

Proses menghilangkan kata kata umum yang tidak memiliki arti atau makna seperti (dan, yang, kok, di, dan lain lain).

e. Stemming

Adalah proses mengubah kata menjadi kata dasar, tujuanya untuk menyatukan kata yang memiliki makna dasar seperti "mencium menjadi cium, menulis menjadi nulis, memakan menjadi makan".

Transformation

Transformasi merupakan langkah ketiga dalam metodologi KDD setelah seleksi dan pra-pemrosesan. Tujuan utama transformasi adalah untuk mengubah data yang telah dibersihkan dan distandarisasi menjadi format yang lebih sesuai untuk penambangan data. Dalam penelitian analisis sentimen, data teks non-numerik harus diubah menjadi representasi numerik agar dapat diproses oleh algoritma klasifikasi seperti Naïve Bayes. Transformasi dalam penelitian ini mencakup dua sub-tahap utama: pengkodean variabel target (label sentimen) dan vektorisasi teks menggunakan TF-IDF.

Tabel: Hasil dari TF-IDF mengubah teks ke angka

Jumlah Data yang didapatkan	1509,144
------------------------------------	----------

Data Latih	1207,144
Data Uji	302,144

ini adalah hasil dari transformation atau TF-IDF yang mana tahap ini ada mengubah teks ke angka dan nanti akan di baca atau di hitung oleh agoritma yang kita gunakan.

Data Mining

Dalam penelitian analisis sentimen, data mining bertujuan untuk mengklasifikasikan opini publik ke dalam kelas sentimen (positif, netral, negatif) berdasarkan teks tweet. Penelitian ini menerapkan dua pendekatan utama, Berbasis Leksikon (sebagai pembentukan data pelatihan) dan Naïve Bayes (sebagai model klasifikasi terawasi). Seluruh proses data mining diimplementasikan menggunakan bahasa pemrograman Python dengan bantuan pustaka scikit-learn, pandas, nltk, dan Sastrawi.

Evaluasi

Tahap evaluasi adalah dima kita akan melihat hasil kualitas dari keandalan model atau pola yang dihasilkan dari proses antra dua metode yang kita gunakan seperti pendekatan Lexicon Based dan Naive Bayes, Sejauh mana ketepatan dalam menghitung sentimen.

Kesimpulan

Studi ini berhasil menganalisis sentimen publik terhadap kontroversi Gus Elham Yahya menggunakan algoritma Naïve Bayes berbasis leksikon. Hasilnya menunjukkan bahwa 63,6% opini publik bersifat negatif, yang mengindikasikan bahwa perilaku tokoh publik yang dianggap melanggar norma perlindungan anak dapat secara signifikan mengurangi kepercayaan publik. Metode yang digunakan mencapai akurasi 98%, membuktikan efektivitas Naïve Bayes untuk analisis sentimen di Twitter/X.

Ucapan Terimakasih

Penulis mengucapkan terima kasih kepada Rumah Jurnal Ilmiah Universitas Muhammadiyah Muara Bungo atas dukungan

dan kesempatan yang diberikan sehingga kegiatan Seminar Nasional dan *Call for Paper* ini penulis bisa berjalan dengan baik. Penulis juga menyampaikan apresiasi kepada semua yang terlibat dalam penelitian ini sehingga sampai artikel ini diterbitkan

Daftar Pustaka

- [1] Undap, M. G., Rantung, V. P., & Rompas, P. T. D. (2021). *Analisis Sentimen Situs Pembajak Artikel Penelitian Menggunakan Metode Lexicon-Based*.
- [2] Ismail, A. R., Bagus, R., & Hakim, F. (2023). *Implementasi Lexicon Based Untuk Analisis Sentimen Dalam Mengetahui Trend Wisata Pantai Di DI Yogyakarta Berdasarkan Data Twitter*. 1(1), 37–46.
- [3] Fikri, M. I., Sabrila, T. S., & Azhar, Y. (2020). *Perbandingan Metode Naïve Bayes dan Support Vector Machine pada Analisis Sentimen Twitter*. 10, 71–76.
- [4] Metode, I., Bayes, N., Wikarsa, L., Angdresey, A., Kapantow, J. D., Informatika, P. T., & Teknik, F. (2022). *APPROACH UNTUK MENGLASIFIKASI SENTIMEN NETIZEN PADA TWEET BERBAHASA INDONESIA*. 18(1), 15–24.
- [5] Dimas, M., Chasis, R., Azahari, D., & Ibnu, M. (2025). *Analisis Sentimen Terhadap Kontroversi Pembangunan IKN di Media Sosial Twitter Menggunakan Metode Naïve Bayes*. 6(2), 97–109. <https://doi.org/10.47065/bit.v5i2.1993>
- [6] Lexicon-based, M. P., Fatah, Z., & Ningsih, R. A. (2025). *Analisis Sentimen Komentar YouTube terhadap Tragedi Demo 25 Agustus*. 4, 217–224.
- [7] Enjeli, M., & Wijaya, A. (2024). *Analisis Sentimen Penggunaan Aplikasi Kinemaster Menggunakan Metode Naive Bayes*. 2, 89–98.
- [8] Sholeh, M., Aeni, K., Informatika, P. S., Informatika, P. S., Sains, F., Peradaban, U., & Bouldin, D. (2023). *PERBANDINGAN EVALUASI METODE DAVIES BOULDIN , ELBOW DAN SILHOUETTE PADA MODEL CLUSTERING DENGAN*. 8(1).
- [9] Bina, U., Jl, D., & Yani, J. A. (2025). *Analisis Clustering Data Gizi Pada Puskesmas 1 Ulu Menggunakan Metode Algoritma K-Means*. 5(2), 209–217. <https://doi.org/10.35957/algorithm.v5i2.10707>.
- [10] Putry, N. M., Sari, B. N., & Kom, M. (2022). *KOMPARASI ALGORITMA KNN DAN NAÏVE BAYES UNTUK KLASIFIKASI DIAGNOSIS PENYAKIT DIABETES MELITUS*. 10(1).
- [11] Nasional, J., Informasi, S., Alia, Q., Sagirani, T., & Lemantara, J. (2023). *Perbandingan Algoritma Naïve Bayes dan K-Nearest Neighbor (KNN) untuk Mengetahui Keakuratan Diagnosa Penyakit Diabetes*. 03, 247–254.
- [12] Pebdika, A., Herdiana, R., Solihudin, D., Pinter, P. I., Bayes, N., Penentuan, U., & Bantuan, P. (2023). *KLASIFIKASI MENGGUNAKAN METODE NAIVE BAYES UNTUK MENENTUKAN CALON PENERIMA PIP*. 7(1), 452–458.