

ANALISIS SENTIMEN PENGGUNA X MENGENAI OPINI MASYARAKAT TENTANG DINASTI POLITIK MENGGUNAKAN ALGORITMA NAÏVE BAYES

Sevi Andini¹, Rudi Kurniawan², Saeful Anwar³

^{1,2,3}STMIK IKMI Cirebon; Jl. Perjuangan No. 10B, Karyamulya, Kec. Kesambi, Kota Cirebon, Jawa Barat; (0231) 490480

Received: 2 Maret 2025
Accepted: 27 Maret 2025
Published: 14 April 2025

Keywords:

Naïve Bayes; Teknik
Sampling; Dinasti Politik; X;
Sentimen.

Correspondent Email:

selviandini69@gmail.com

Abstrak. Fenomena dinasti politik di Indonesia sering memicu perdebatan, terutama terkait dampaknya terhadap demokrasi dan representasi rakyat. Media sosial, seperti X, menjadi wadah utama bagi masyarakat untuk mengekspresikan opini mereka mengenai isu ini. Namun, analisis sentimen terhadap opini publik di media sosial menghadapi tantangan, terutama karena data teks yang tidak terstruktur dan mengandung noise. Penelitian ini bertujuan untuk mengidentifikasi kata-kata yang sering muncul dalam sentimen terkait dinasti politik, mengevaluasi akurasi algoritma Naïve Bayes dalam analisis sentimen, serta menentukan teknik *sampling* yang paling optimal untuk meningkatkan akurasi model. Metode yang digunakan adalah *Knowledge Discovery in Database* (KDD) dengan algoritma Naïve Bayes sebagai model klasifikasi sentimen. Hasil penelitian menunjukkan bahwa kata-kata seperti “politik,” “dinasti,” “banten,” “rakyat,” dan “keluarga” sering muncul dalam sentimen terkait dinasti politik. Teknik linear sampling memberikan akurasi tertinggi sebesar 77.78%, dengan *precision* 76.19% dan *recall* 84.21%. Penelitian ini menunjukkan bahwa distribusi data yang seimbang dan teknik *sampling* yang tepat berpengaruh signifikan terhadap performa model.

Abstract. The phenomenon of political dynasties in Indonesia often sparks debate, especially regarding its impact on democracy and popular representation. Social media, such as X, has become the main forum for people to express their opinions on this issue. However, sentiment analysis of public opinion on social media faces challenges, especially because text data is unstructured and contains noise. This research aims to identify words that frequently appear in sentiment related to political dynasties, evaluate the accuracy of the Naïve Bayes algorithm in sentiment analysis, and determine the most optimal sampling technique to increase model accuracy. The method used is *Knowledge Discovery in Database* (KDD) with the Naïve Bayes algorithm as a sentiment classification model. The research results show that words such as “politics,” “dynasty,” “banten,” “people,” and “family” often appear in sentiments related to political dynasties. The linear sampling technique provides the highest accuracy of 77.78%, with precision of 76.19% and recall of 84.21%. This research shows that balanced data distribution and appropriate sampling techniques have a significant effect on model performance.

1. PENDAHULUAN

Kemajuan teknologi informasi dan komunikasi telah memberikan dampak signifikan dalam berbagai aspek kehidupan, termasuk dalam dinamika politik. Media sosial telah menjadi ruang publik digital yang memungkinkan masyarakat menyuarakan pendapat mereka tentang berbagai isu politik, salah satunya adalah fenomena dinasti politik. Platform seperti media sosial X telah menjadi alat utama dalam menyampaikan opini publik, membentuk persepsi, serta mempengaruhi diskursus politik secara luas [1].

Dalam beberapa tahun terakhir, analisis sentimen telah berkembang menjadi metode yang semakin populer dalam mengkaji opini publik yang terdapat di media sosial. Salah satu teknik yang sering digunakan dalam penelitian analisis sentimen adalah metode Naïve Bayes, yang terbukti efektif dalam memproses data dalam jumlah besar serta mengklasifikasikan sentimen ke dalam kategori positif, negatif, dan netral [2]. Penelitian terdahulu telah menunjukkan bahwa Naïve Bayes dapat diterapkan pada berbagai isu, seperti sentimen terhadap aplikasi pendidikan [3], opini publik tentang vaksin COVID-19 [4], serta kebijakan pemerintah seperti Pembatasan Kegiatan Masyarakat (PPKM) [5]. Namun, kajian terkait penggunaan metode ini dalam menganalisis persepsi publik terhadap dinasti politik masih terbatas.

Teori yang berkaitan dengan penelitian ini mencakup teori analisis sentimen dalam pemrosesan bahasa alami (*Natural Language Processing/NLP*) serta teori tentang opini publik dan politik digital. Teori analisis sentimen berfokus pada bagaimana data teks yang berasal dari media sosial dapat diolah dan dikategorikan untuk memahami kecenderungan opini masyarakat. Sementara itu, teori opini publik menjelaskan bagaimana persepsi masyarakat terhadap isu politik terbentuk dan berkembang seiring dengan dinamika sosial yang terjadi [6].

Meskipun terdapat banyak penelitian yang menerapkan metode Naïve Bayes dalam analisis sentimen, penerapannya dalam konteks dinasti politik masih minim, terutama dalam kajian akademik yang berfokus pada media sosial di Indonesia. Penelitian sebelumnya lebih banyak menyoroti aspek-aspek lain seperti pendidikan dan kebijakan kesehatan, sementara

diskursus politik di media sosial, terutama terkait dinasti politik, masih jarang dikaji secara mendalam. Padahal, isu ini sangat relevan mengingat semakin meningkatnya peran media sosial dalam membentuk opini publik serta pengaruhnya terhadap keputusan politik masyarakat.

Kesenjangan lain yang ditemukan adalah kurangnya penelitian yang membandingkan teknik sampling dalam analisis sentimen berbasis Naïve Bayes untuk meningkatkan akurasi klasifikasi. Dengan demikian, penelitian ini bertujuan untuk mengisi celah tersebut dengan mengeksplorasi teknik sampling yang paling efektif untuk meningkatkan akurasi model prediktif dalam analisis sentimen terkait dinasti politik di media sosial.

Penelitian ini bertujuan untuk mengidentifikasi dan menganalisis sentimen publik terhadap fenomena dinasti politik di Indonesia melalui media sosial X dengan menggunakan metode Naïve Bayes. Melalui penelitian ini, diharapkan dapat diketahui distribusi sentimen publik (positif, negatif, dan netral) terhadap dinasti politik serta kata-kata atau frasa yang paling sering muncul dalam tweet yang membahas isu tersebut. Selain itu, penelitian ini bertujuan untuk mengevaluasi akurasi metode Naïve Bayes dalam klasifikasi sentimen serta membandingkan efektivitas teknik sampling yang berbeda dalam meningkatkan akurasi model. Hasil penelitian ini diharapkan dapat memberikan wawasan yang lebih mendalam mengenai persepsi publik terhadap dinasti politik, serta berkontribusi pada pengembangan metode analisis sentimen yang lebih akurat dan efisien dalam studi politik digital. Temuan penelitian ini juga diharapkan dapat memberikan manfaat bagi pemerintah, akademisi, serta praktisi politik dalam memahami opini publik dan merancang strategi komunikasi politik yang lebih tepat sasaran.

2. TINJAUAN PUSTAKA

2.1. Algoritma Naïve Bayes

Algoritma Naïve Bayes adalah salah satu algoritma klasifikasi dalam machine learning yang didasarkan pada perhitungan probabilitas dan statistik yang diperkenalkan oleh Thomas Bayes. Algoritma ini berfungsi untuk memperkirakan probabilitas suatu kejadian di

masa depan berdasarkan data historis. Pendekatan Naïve Bayes mengasumsikan bahwa setiap atribut dalam data bersifat independen satu sama lain [7].

2.2. Analisis Sentimen

Analisis sentimen adalah bidang penelitian yang terus berkembang yang mencakup berbagai bidang seperti data mining, pemrosesan bahasa alami *Natural Language Processing* (NLP), dan pembelajaran mesin yang berfokus pada mengekstrak sentimen dalam sebuah kalimat berdasarkan isi kalimat tersebut. Analisis sentimen, juga dikenal sebagai analisis opini, adalah proses memahami, mengekstrak, dan mengolah data teks secara otomatis untuk mendapatkan informasi yang terkandung dalam suatu kalimat. Tujuan dari analisis sentimen adalah untuk mengidentifikasi kecenderungan pendapat atau opini seseorang terhadap suatu masalah atau objek, apakah itu memiliki kecenderungan positif, negatif, atau netral [8].

2.3. Dinasti Politik

Dinasti politik merupakan fenomena yang masih berlangsung di Indonesia dan menjadi isu penting yang perlu dibenahi agar etika dalam berpolitik benar-benar mencerminkan nilai-nilai demokrasi. Karakteristik dinasti politik terlihat dari adanya penguasaan jabatan-jabatan strategis dalam pemerintahan oleh sekelompok individu yang berasal dari keluarga atau memiliki hubungan kekerabatan. Fenomena politik dinasti yang sangat erat kaitannya dengan faktor kekerabatan ini berdampak pada dinamika partai politik di Indonesia [9].

2.4. Media Sosial X

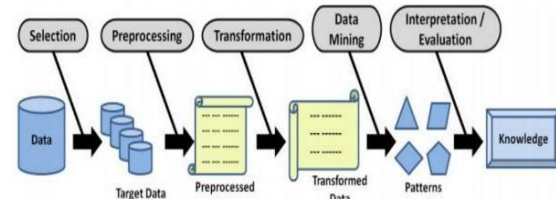
Media sosial x atau yang sebelumnya dikenal twitter, salah satu media sosial yang banyak digunakan oleh masyarakat untuk berkomunikasi dan memperoleh informasi, memungkinkan pengguna untuk menyampaikan berbagai macam opini dan komentar tentang berbagai masalah. Selain itu, opini dan komentar yang disampaikan pengguna melalui tweets mereka dapat digunakan untuk menganalisis sentiment [10].

2.5. Knowledge Discovery in Database

Knowledge Discovery in Databases (KDD) adalah proses kompleks yang bertujuan untuk mencari dan mengidentifikasi pola dalam data. Pola yang ditemukan harus memiliki karakteristik valid, baru, bermanfaat, dan dapat dipahami, sehingga dapat memberikan wawasan yang bermakna [11].

3. METODE PENELITIAN

Tahapan metode penelitian data dilakukan dengan mengadopsi metode *Knowledge Discovery in Databases* (KDD). Proses KDD diterapkan untuk memastikan bahwa setiap langkah yang dilakukan berlangsung secara terstruktur dan efisien. Tahapan-tahapan dalam metode KDD meliputi pemilihan data (*Data Selection*), pra proses data (*Pre-processing*), transformasi data (*Transformation*), penambangan data (*Data Mining*), evaluasi hasil (*Evaluation*), dan akhirnya pengolahan pengetahuan (*Knowledge*). Penelitian ini menggunakan KDD dan Algoritma Naive Bayes untuk mengklasifikasikan opini publik tentang dinasti politik. Tahapan metode penelitian ditampilkan pada Gambar 1.



Gambar 1. Tahapan Metode Penelitian

3.1. Data Selection

Proses ini bertujuan untuk mengumpulkan data tweet terkait dinasti politik dari platform media sosial X, menggunakan teknik *crawling*.

3.2. Pre-processing

Pre-processing merupakan tahap penting dalam analisis data yang bertujuan untuk mempersiapkan data mentah agar siap untuk dianalisis lebih lanjut. Pada tahap ini, data yang telah dikumpulkan akan dibersihkan, diubah, dan dipersiapkan sedemikian rupa agar kualitasnya lebih baik dan cocok untuk analisis. Adapun tahapan dari *pre-processing* sebagai berikut:

1. *Cleansing*: merupakan tahapan membersihkan data dari elemen yang tidak relevan atau tidak

berdampak pada proses klasifikasi pada data teks.

2. *Case Folding*: merupakan tahapan pengolahan teks yang bertujuan untuk menyamakan bentuk huruf dalam sebuah teks menjadi satu format yang konsisten, biasanya dengan mengubah semua huruf menjadi huruf kecil (*lowercase*).
3. *Tokenize*: merupakan tahapan memecah teks menjadi unit-unit yang lebih kecil, biasanya berupa kata atau frasa yang disebut token.
4. *Stemming*: merupakan tahapan untuk mengubah kata-kata modern atau singkatan menjadi bentuk dasar.
5. *Stopword*: merupakan tahapan untuk menghapus kata-kata yang tidak bermakna atau kata penghubung yang tidak memengaruhi konteks kalimat.

3.3. Transformation

Tahap ini dilakukan dengan menerapkan metode *Term Frequency-Inverse Document Frequency* (TF-IDF). Teknik ini memberikan nilai bobot kepada setiap kata berdasarkan frekuensi kemunculannya dalam dokumen, dengan tujuan untuk menilai pentingnya kata tersebut dalam konteks keseluruhan.

3.4. Data Mining

Data Mining adalah proses mengidentifikasi pola dalam data yang akan ditambang dengan menggunakan teknik tertentu, dalam penelitian ini digunakan metode Naive Bayes [12]. Naive Bayes merupakan sebuah model pengklasifikasi dengan bias tinggi dan varians rendah, serta dapat membangun model yang baik bahkan dengan set data yang kecil.

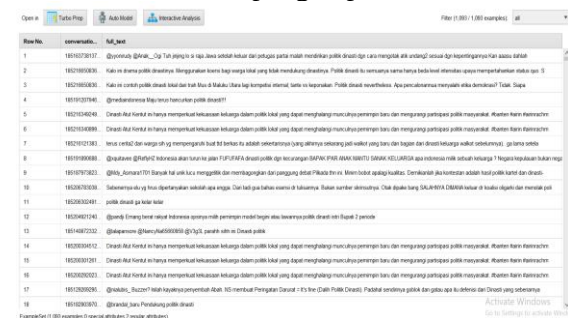
3.5. Evaluation

Penelitian ini juga melakukan perhitungan manual terhadap hasil nilai *Confusion Matrix*. Dari *Confusion Matrix* ini, dapat menghitung berbagai matrik evaluasi seperti akurasi, presisi, dan *recall*.

4. HASIL DAN PEMBAHASAN

4.1. Data Selection

Pengumpulan data dalam penelitian ini berasal dari sumber data media sosial X (Twitter). Data dikumpulkan menggunakan *tweet-harvest* dengan bahasa pemrograman *Python* di *Google Colaboratory*, dengan kata kunci “Politik Dinasti.” Pengumpulan data dilakukan pada bulan Oktober hingga November 2024, menghasilkan maksimal 1.093 tweet. Berikut hasil pengumpulan dataset.



Gambar 2. Hasil Pengumpulan Dataset

4.2. Pre-processing

Proses *pre-processing* ini melalui beberapa langkah, yaitu *cleansing*, *case folding*, *tokenizing*, *stopword*, dan *stemming*. Berikut ini adalah tahapan pre-processing yang dilakukan:

1. *Cleansing*: Tahap ini menghapus seperti tanda baca, angka, URL, mention, teks yang memiliki duplikasi, dan karakter khusus lainnya dalam analisis teks.

Tabel 1. Hasil Data *Cleansing*

Sebelum	Sesudah
@hasyimmah Kalo mau bangun dinasti politik secara konstitusional ya harus punya partai dong.. Masih nebeng kok mau bangun dinasti politik.	Kalo mau bangun dinasti politik secara konstitusional ya harus punya partai dong Masih nebeng kok mau bangun dinasti politik

2. *Case Folding*: Tahap ini untuk mengubah semua huruf kapital menjadi huruf kecil (*lowercase*) pada data teks agar seragam.

Tabel 2. Hasil *Case Folding*

Sebelum	Sesudah
Kalo mau bangun dinasti	kalo mau bangun dinasti

politik secara konstitusional ya harus punya partai dong Masih nebeng kok mau bangun dinasti politik	politik secara konstitusional ya harus punya partai dong masih nebeng kok mau bangun dinasti politik
---	---

3. *Tokenize*: Pada tahap *tokenize*, merupakan tahapan memecah teks menjadi unit-unit yang lebih kecil, biasanya berupa kata atau frasa yang disebut token.

Tabel 3. Hasil *Tokenize*

Sebelum	Sesudah
kalo mau bangun dinasti politik secara konstitusional ya harus punya partai dong masih nebeng kok mau bangun dinasti politik	“kalo” “mau” “bangun” “dinasti” “politik” “secara” “konstitusional” “ya” “harus” “punya” “partai” “dong” “masih” “nebeng” “kok” “mau” “bangun” “dinasti” “politik”

4. *Stemming*: Pada tahap ini berfungsi untuk mengubah kata-kata dalam data teks ke bentuk dasarnya. Proses ini mengurangi pengulangan, meningkatkan konsistensi data, dan memudahkan algoritma untuk mengidentifikasi pola sentimen dengan tepat.

Tabel 4. Hasil *Stemming*

Sebelum	Sesudah
“kalo” “mau” “bangun” “dinasti” “politik” “secara” “konstitusional” “ya” “harus” “punya” “partai” “dong” “masih” “nebeng” “kok”	“kalo” “mau” “bangun” “dinasti” “politik” “secara” “konstitusi” “ya” “harus” “punya” “partai” “dong”

“mau” “bangun” “dinasti” “politik”	“masih” “nebeng” “kok” “mau” “bangun” “dinasti” “politik”
--	--

5. *Stopword*: Tahap ini digunakan untuk menghapus kata-kata umum yang tidak memiliki nilai informatif dalam analisis teks. Tujuan penghapusan stopwords ini adalah untuk meningkatkan kompleksitas data dan memberi fokus pada kata-kata yang relevan untuk analisis sentimen.

Tabel 5. Hasil *Stopword*

Sebelum	Sesudah
“kalo” “mau” “bangun” “dinasti” “politik” “secara” “konstitusi” “ya” “harus” “punya” “partai” “dong” “masih” “nebeng” “kok” “mau” “bangun” “dinasti” “politik”	“kalo” “mau” “bangun” “dinasti” “politik” “secara” “konstitusi” “harus” “punya” “partai” “masih” “mau” “bangun” “dinasti” “politik”

4.3. Transformation

Pada tahap ini, teks tweet yang telah melalui proses *pre-processing* data akan diubah menjadi representasi data numerik. Untuk mencapai tahap ini, metode *Term Frequency-Inverse Document Frequency* (TF-IDF) digunakan. Tujuan pembobotan ini adalah untuk menentukan nilai untuk setiap kata berdasarkan seberapa sering kata muncul dalam dokumen atau seberapa umum kata digunakan.

Word	Attribute Name	Total Occurrences ↓	Document Occurrences
politik	politik	1075	963
dinasti	dinasti	1057	961
banten	banten	176	162
rakyat	rakyat	107	94
keluarga	keluarga	92	83
indonesia	indonesia	81	70
dukung	dukung	71	62
negara	negara	71	60
anak	anak	58	53
kuasa	kuasa	58	52

Gambar 3. Hasil Tahap Data Transformation TF-IDF

Pada gambar 3 mencakup kolom untuk *Word* (kata), *Attribute Name* (nama atribut yang mewakili kata tersebut), *Total Occurrences* (jumlah total kemunculan kata dalam seluruh kumpulan tweet), dan *Document Occurrences* (jumlah dokumen atau tweet unik yang mengandung kata tersebut). Seperti, kata "politik" muncul sebanyak 1,075 kali dalam keseluruhan dataset, dan ditemukan di 963 tweet yang berbeda. Kata "dinasti" muncul 1,057 kali dalam total data dan ada dalam 961 dokumen. Frekuensi tinggi dari kata-kata seperti "politik," "dinasti," "banten," "rakyat," dan "keluarga" menunjukkan bahwa kata-kata ini adalah topik utama atau sering dibahas dalam tweet terkait, yang relevan untuk analisis sentimen terhadap isu politik dinasti.

4.4. Data Mining

Pada tahap ini, akan ditentukan data latih dan data uji serta teknik sampling yang menghasilkan nilai akurasi terbaik. Untuk melakukan proses ini, operator split data membagi data menjadi tiga pengujian, yaitu perbandingan 70:30, 80:20, dan 90:10. Perbandingan ini secara otomatis sesuai dengan perbandingan dari dataset yang sudah bersih. Pada tahap ini, algoritma Naive Bayes akan digunakan untuk memproses data.

4.5. Evaluation

Pada tahap ini, evaluasi klasifikasi sentimen dilakukan menggunakan algoritma Naive Bayes. Confusion matrix digunakan untuk menganalisis kinerja model dan menghitung nilai akurasi, presisi, dan recall untuk menunjukkan efektivitasnya. Untuk mendapatkan pemahaman yang lebih baik tentang kinerja model ini, setiap pembagian

data dengan rasio 70:30, 80:20, dan 90:10 dievaluasi.

Tabel 6. Hasil Evaluasi

Rasio	Akurasi	Presisi	Recall
70:30	72.22%	72.88%	75.44%
80:20	69.44%	71.05%	71.05%
90:10	77.78%	76.19%	84.21%

Hasil evaluasi ini diperoleh dari sampling type linear sampling, yang menunjukkan nilai akurasi terbaik dibandingkan *sampling type* lainnya. Dapat dilihat bahwa pada rasio data 70:30, model mencapai akurasi 72.22%, dengan presisi 72.88%, dan *recall* 75.44%. Dari rasio 80:20, model menunjukkan nilai akurasi 69.44%, presisi 71.05%, dan *recall* 75.05%. Sementara itu rasio 90:10, model mencapai performa terbaiknya dengan nilai akurasi 77.78%, presisi 76.19%, dan *recall* 84.21%.

5. KESIMPULAN

- Secara keseluruhan kata-kata seperti "politik", "dinasti", "banten", "rakyat", dan "keluarga" memiliki frekuensi kemunculan yang tinggi. Hal ini menunjukkan bahwa kata-kata tersebut sering dibicarakan pada isu terkait dinasti politik. Penelitian ini dapat dengan mudah mengidentifikasi pola sentimen publik, tetapi tidak dapat melihat aspek emosional yang lebih kompleks seperti sarkasme atau konteks budaya.
- Pengukuran kinerja model Naive Bayes menunjukkan hasil evaluasi dengan akurasi sebesar 77.78%, *precision* sebesar 76.19%, dan *recall* sebesar 84.21%. Nilai-nilai ini menunjukkan performa model yang cukup baik dalam mengklasifikasikan sentimen publik terkait dinasti politik, tetapi asumsi independensi antar fitur dapat mempengaruhi akurasi klasifikasi.
- Teknik sampling linear terbukti memberikan hasil terbaik dalam meningkatkan akurasi model Naive Bayes pada penelitian ini. Dengan akurasi yang mencapai 77.78%. Teknik ini dapat dipertimbangkan sebagai pendekatan yang efektif dalam analisis sentimen untuk kasus

yang serupa, khususnya dalam konteks media sosial.

UCAPAN TERIMA KASIH

Penulis mengucapkan terima kasih kepada pihak-pihak terkait yang telah memberi dukungan terhadap penelitian ini.

DAFTAR PUSTAKA

- [1] D. A. Agustina, S. Subanti, and E. Zukhronah, "Implementasi Text Mining Pada Analisis Sentimen Pengguna Twitter Terhadap Marketplace di Indonesia Menggunakan Algoritma Support Vector Machine," *Indones. J. Appl. Stat.*, vol. 3, no. 2, pp. 109–122, 2020, doi: 10.13057/ijas.v3i2.44337.
- [2] K. C. Astuti, A. Firmansyah, and A. Riyadi, "Implementasi Text Mining Untuk Analisis Sentimen Masyarakat Terhadap Ulasan Aplikasi Digital Korlantas Polri pada Google Play Store," *REMIK Ris. dan E-Jurnal Manaj. Inform. Komput.*, vol. 8, no. 1, pp. 383–394, 2024.
- [3] A. P. Giovani, A. Ardiansyah, T. Haryanti, L. Kurniawati, and W. Gata, "Analisis Sentimen Aplikasi Ruang Guru Di Twitter Menggunakan Algoritma Klasifikasi," *J. Teknoinfo*, vol. 14, no. 2, pp. 116–124, 2020, doi: 10.33365/jti.v14i2.679.
- [4] F. Fathonah and A. Herliana, "Penerapan Text Mining Analisis Sentimen Mengenai Vaksin Covid - 19 Menggunakan Metode Naïve Bayes," *J. Sains dan Inform.*, vol. 7, no. 2, pp. 155–164, 2021, doi: 10.34128/jsi.v7i2.331.
- [5] T. Krisdiyanto and E. M. O. Nurharyanto, "Analisis Sentimen Opini Masyarakat Indonesia Terhadap Kebijakan PPKM pada Media Sosial Twitter Menggunakan Naïve Bayes Clasifiers," *J. CoreIT J. Has. Penelit. Ilmu Komput. dan Teknol. Inf.*, vol. 7, no. 1, pp. 32–37, 2021, doi: 10.24014/coreit.v7i1.12945.
- [6] E. Priansyah and T. Sutabri, "Analisis Sentimen Berbasis Naïve Bayes Pada Media Sosial Twitter Terhadap Hasil Pemilu Indonesia 2024," *IJM Indones. J. Multidiscip.*, vol. 2, no. 3, pp. 128–138, 2024, [Online]. Available: <https://journal.csspublishing/index.php/ijm>
- [7] A. Kusuma and A. Nugroho, "Analisa Sentimen Pada Twitter Terhadap Kenaikan Tarif Dasar Listrik Dengan Metode Naïve Bayes," *J. Ilm. Teknol. Inf. Asia*, vol. 15, no. 2, pp. 137–146, 2021, doi: 10.32815/jitika.v15i2.557.
- [8] A. Halim and Andri Safuwan, "Analisis Sentimen Opini Warganet Twitter Terhadap Tes Screening Genose Pendeteksi Virus Covid-19 Menggunakan Metode Naïve Bayes Berbasis Particle Swarm Optimization," *J. Inform. Teknol. dan Sains*, vol. 5, no. 1, pp. 170–178, 2023, doi: 10.51401/jinteks.v5i1.2229.
- [9] Agus Dedi, "Politik Dinasti Dalam Perspektif Demokrasi," *Moderat J. Ilm. Ilmu Pemerintah.*, vol. 8, no. 1, pp. 92–101, 2022, doi: 10.25157/moderat.v8i1.2596.
- [10] M. I. Fikri, T. S. Sabrila, and Y. Azhar, "Perbandingan Metode Naïve Bayes dan Support Vector Machine pada Analisis Sentimen Twitter," *Smatika J.*, vol. 10, no. 02, pp. 71–76, 2020.
- [11] T. I. Hermanto and M. A. Sunandar, "ANALISIS DATA SEBARAN PENYAKIT MENGGUNAKAN ALGORITMA DENSITY BASED SPATIAL CLUSTERING OF APPLICATIONS WITH NOISE," *J. Sains Komput. dan Teknol. Inf.*, vol. 3, no. 1, pp. 104–110, 2020, [Online]. Available: <http://jurnal.umt.ac.id/index.php/nyimak>
- [12] I. Print, G. A. Putri, A. Trimaysella, and A. Khoiriah, "Penerapan Klasifikasi Data Mining pada Diabetes Menggunakan Metode Naive Bayes," *J. Ilmu Komput. Teknol. Terap.*, vol. 1, no. 1, pp. 1–9, 2024.