

Analisis Sentimen di Media Sosial X tentang IKN dengan Naïve Bayes

Gilang Mario Conroy Paridy Man¹, Alfry Aristo Jansen Sinlae^{*2}, Emerensiana Ngaga³

^{1,2,3} Program Studi Ilmu Komputer, Fakultas Teknik, Universitas Katolik Widya Mandira, Kupang-Indonesia

¹gilangmcpm@gmail.com, ^{2*}alfry.aj@unwira.ac.id, ³emerensianangaga@unwira.ac.id

Abstrak

Pembangunan Ibu Kota Nusantara (IKN), khususnya proyek Istana Garuda, memicu beragam opini masyarakat di media sosial, terutama di platform X (sebelumnya Twitter). Penelitian ini bertujuan untuk menganalisis sentimen publik terhadap Istana Garuda menggunakan pendekatan *Naïve Bayes*. Data dikumpulkan melalui teknik *crawling* menggunakan *TweetHarvest* pada periode April hingga Oktober 2024, menghasilkan 585 *tweet* relevan. Proses *preprocessing* mencakup pembersihan data, *case folding*, tokenisasi, penghapusan *stopword*, dan *stemming* menggunakan Python. Sentimen dilabeli secara otomatis menggunakan kamus *Lexicon InSet* dan diklasifikasikan dengan algoritma *Naïve Bayes Multinomial*, dengan pembagian data 80% latih dan 20% uji. Hasil klasifikasi menunjukkan sebaran sentimen positif sebesar 45,6%, negatif 29,1%, dan netral 25,3%. Evaluasi model menghasilkan akurasi 61%, *precision* 71%, *recall* 61%, dan *F1-score* 54%. Selain itu, analisis topik mengidentifikasi tiga isu dominan, yaitu dukungan terhadap pembangunan, kritik terhadap anggaran, dan diskusi netral terkait progres konstruksi. Simpulan dari penelitian ini menunjukkan bahwa algoritma *Naïve Bayes* cukup efektif dalam mengklasifikasi sentimen publik secara moderat. Mayoritas opini masyarakat bersifat positif, mencerminkan dukungan terhadap proyek IKN. Hasil ini dapat digunakan sebagai acuan bagi pemangku kebijakan untuk merumuskan strategi komunikasi publik yang lebih tepat sasaran. Rekomendasi lanjutan mencakup perluasan dataset dan eksplorasi algoritma lain untuk peningkatan akurasi.

Kata kunci: X, naïve bayes, istana garuda, IKN, analisis sentimen

1. Pendahuluan

Indonesia saat ini tengah melaksanakan proyek ambisius pemindahan ibu kota negara dari Jakarta ke Nusantara. Inisiatif ini tidak hanya merupakan langkah strategis dalam pengembangan wilayah, tetapi juga memicu beragam reaksi dari masyarakat (Yusuf et al., 2023). Respons tersebut, baik yang berupa dukungan maupun penolakan, terwujud dalam berbagai bentuk diskusi dan opini di media sosial. Salah satu fokus perhatian publik adalah Istana Garuda, yang berfungsi sebagai simbol kekuasaan pemerintahan di Ibu Kota Nusantara (IKN) (Al-Deen et al., 2024). Media sosial, terutama X (sebelumnya Twitter), telah menjadi platform utama bagi masyarakat untuk menyampaikan pendapat, berinteraksi, dan membentuk opini publik (Abbas et al., 2019; Appel et al., 2020). Sentimen yang terungkap di media sosial dapat mencerminkan persepsi masyarakat secara luas terhadap isu-isu penting, termasuk pembangunan IKN dan Istana Garuda (Appel et al., 2020; Arisanty et al., 2020). Menurut penelitian (Dwivedi et al., 2021), media sosial berfungsi sebagai alat penting untuk memahami dinamika opini publik, khususnya dalam konteks perubahan sosial dan politik.

Analisis sentimen di media sosial sangat penting untuk memahami pandangan publik secara lebih objektif. Dengan jutaan pengguna aktif, media sosial X memungkinkan penyampaian opini dalam bentuk

teks singkat yang dapat mencerminkan pandangan positif, negatif, atau netral terhadap proyek pembangunan tersebut. Pemahaman terhadap sentimen masyarakat ini menjadi krusial bagi pemerintah dan pemangku kepentingan untuk mengukur persepsi publik, menilai keberhasilan sosialisasi proyek, serta memitigasi risiko komunikasi yang mungkin timbul akibat kesalahpahaman atau penolakan masyarakat. Sejalan dengan ini, penelitian (Surya & Subbulakshmi, 2019; Zulfiker et al., 2022) menunjukkan bahwa analisis sentimen dapat membantu pengambil keputusan dalam merancang strategi komunikasi yang lebih efektif dan responsif terhadap opini publik (Meme et al., 2022).

Urgensi penelitian ini terletak pada perlunya pemahaman mendalam mengenai opini publik yang berkembang secara *real-time* di media sosial X terhadap proyek strategis nasional seperti IKN. Mengingat Istana Garuda merupakan simbol utama dari otoritas pemerintahan baru, persepsi masyarakat terhadapnya akan sangat memengaruhi legitimasi sosial dan keberlanjutan proyek ini. Oleh karena itu, analisis sentimen tidak hanya penting secara teknis dalam ranah pemrosesan bahasa alami, tetapi juga memiliki dampak strategis dalam mendukung perumusan kebijakan dan strategi komunikasi pemerintah.

Dalam konteks ini, penelitian ini bertujuan untuk menganalisis sentimen masyarakat di media

sosial X terhadap pembangunan Istana Garuda IKN dengan menggunakan algoritma klasifikasi *Naïve Bayes*. Metode ini telah terbukti efektif dalam berbagai studi analisis sentimen, seperti yang diungkapkan (Aprinastya et al., 2024; Samsir et al., 2022; Zulfiker et al., 2022) yang mengonfirmasi bahwa *Naïve Bayes* dapat memberikan akurasi tinggi dalam klasifikasi teks. Beberapa penelitian terkini juga mendukung penggunaan metode ini untuk analisis sentimen di media sosial. Misalnya, penelitian (Sundara et al., 2020) menunjukkan bahwa *Naïve Bayes* dapat mencapai akurasi 86% dalam menganalisis isu radikalisme dengan menggunakan data set 550 cuitan. Penelitian (Suhardiman & Purwaningtias, 2020) menemukan bahwa analisis sentimen terhadap virus *corona* menggunakan Twitter API menghasilkan 93,4% presisi untuk sentimen positif.

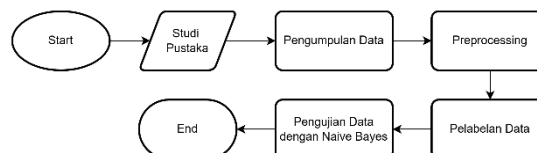
Lebih lanjut, (Aripiyanto et al., 2022) meneliti sentimen terhadap Ibu Kota Nusantara dan menemukan bahwa dengan 5112 cuitan, presisi mencapai 82% dan akurasi 79%. Sementara itu, (Arya & Zufria, 2024) melaporkan akurasi 84,99% dalam analisis sentimen terhadap kampanye di X. (Soffi Arfan et al., 2024) juga menunjukkan bahwa algoritma *Naïve Bayes* mampu mencapai akurasi 86,62% dalam analisis *cyber bullying*. Penelitian-penelitian ini menunjukkan betapa pentingnya teknik analisis sentimen dalam memahami reaksi masyarakat terhadap isu-isu terkini.

Dengan pendekatan ini, diharapkan dapat diperoleh wawasan yang mendalam mengenai bagaimana masyarakat merespons proyek pembangunan ini. Rumusan masalah dalam penelitian ini mencakup dua pertanyaan utama: pertama, bagaimana sentimen masyarakat terhadap Istana Garuda yang terungkap di media sosial X? Kedua, seberapa akurat algoritma *Naïve Bayes* dalam mengklasifikasikan sentimen masyarakat terhadap pembangunan Istana Garuda.

Untuk membatasi ruang lingkup penelitian, terdapat beberapa batasan yang ditetapkan. Pertama, analisis sentimen akan difokuskan pada sentimen publik terhadap Istana Garuda. Kedua, penggunaan bahasa gaul dalam cuitan tidak akan mempengaruhi hasil penelitian. Ketiga, data yang akan dianalisis diambil dari cuitan di media sosial X selama periode April 2024 hingga Oktober 2024. Dengan tujuan tersebut, penelitian ini diharapkan dapat memberikan gambaran komprehensif mengenai respons masyarakat terhadap proyek pembangunan Istana Garuda dan mengetahui tingkat akurasi *Naïve Bayes* dalam menganalisis persepsi masyarakat, yang pada gilirannya dapat membantu pengambilan keputusan dan pengembangan strategi komunikasi pemerintah terkait proyek IKN. Penelitian ini sejalan dengan studi-studi sebelumnya yang menunjukkan pentingnya analisis sentimen dalam pengembangan kebijakan publik dan komunikasi efektif (Samsir et al., 2022; Zulfikar et al., 2023; Zulfiker et al., 2022).

2. Metode

Metode yang digunakan dalam penelitian ini merupakan kegiatan yang merencanakan penelitian dengan sistematis dan ilmiah yang ditunjukkan dalam Gambar 1. Dalam konteks ini, penelitian bertujuan untuk menganalisis sentimen masyarakat terhadap pembangunan Istana Garuda di Ibu Kota Nusantara (IKN) menggunakan metode *Naïve Bayes*. Media sosial X (sebelumnya Twitter) menjadi sumber data utama, mengingat platform ini menghasilkan berbagai opini masyarakat yang dapat dianalisis.



Gambar 1. Alur tahapan penelitian

Proses analisis sentimen dilakukan dengan langkah-langkah berikut:

- Studi Pustaka**
Studi pustaka digunakan untuk mempelajari klasifikasi *Naïve Bayes*, aplikasi media sosial yang digunakan dalam penelitian ini, serta konteks Ibu Kota Nusantara berdasarkan berbagai jurnal. Penelitian-penelitian sebelumnya seperti yang dilakukan oleh (Aripiyanto et al., 2022) dan (Arya & Zufria, 2024) menunjukkan efektivitas metode *Naïve Bayes* dalam analisis sentimen di media sosial.
- Pengumpulan Data**
Pengumpulan data atau *crawling* data dilakukan menggunakan *tools TweetHarvest* dengan bahasa pemrograman *Python*. Data yang dikumpulkan mencakup periode dari April 2024 hingga Oktober 2024, yang terdiri dari empat *file* yang kemudian digabungkan menjadi satu data set.
- Preprocessing Data**
Proses *preprocessing* dilakukan dalam beberapa tahap, yaitu *cleaning*, *case folding*, *tokenization*, *filtering/stopwords removal*, dan *stemming*. Proses *cleaning* menghilangkan *link* URL, HTML, *emoji*, penomoran, dan simbol. *Case folding* dilakukan untuk mengubah huruf kapital menjadi huruf kecil. *Tokenization* membagi kalimat sentimen menjadi pecahan kata-kata. *Filtering/stopwords removal* menghilangkan kata-kata yang tidak memiliki makna, dan langkah terakhir adalah *stemming* untuk menghilangkan imbuhan dalam suatu kata. Semua proses ini dilakukan menggunakan bahasa pemrograman *Python*.
- Labeling Data**
Pelabelan data merupakan proses untuk membagi data sentimen menjadi kategori positif, negatif, dan netral. Proses ini menggunakan kamus *Lexicon InSet* (Indonesia *Sentiment Lexicon*) dan dilakukan secara otomatis dengan bahasa pemrograman *Python*.

e. Pengujian Data Menggunakan *Naïve Bayes*

Setelah pelabelan, dilakukan klasifikasi menggunakan *Naïve Bayes* terhadap data yang telah diberi label. Data set dibagi menjadi 80% untuk data latih dan 20% untuk data uji. Pengujian ini bertujuan untuk mengetahui nilai akurasi, presisi, *recall*, dan *F1-Score* dari sentimen yang telah dianalisis. Pengujian ini juga dilakukan menggunakan bahasa pemrograman *Python*.

Dengan pendekatan ini, diharapkan penelitian dapat memberikan wawasan yang mendalam mengenai sentimen masyarakat terhadap pembangunan Istana Garuda, serta mengukur akurasi dari metode *Naïve Bayes* dalam klasifikasi sentimen di media sosial. Penelitian ini sejalan dengan penelitian sebelumnya yang menunjukkan pentingnya analisis sentimen dalam pengembangan kebijakan publik dan komunikasi efektif.

3. Hasil dan Pembahasan

Berdasarkan tahapan-tahapan dalam metode, maka langkah pertama yang dilakukan dalam menyelesaikan penelitian ini adalah studi pustaka. Studi dilakukan untuk mendapatkan referensi-referensi dan literatur yang sesuai. Hasil studi pustaka menunjukkan bahwa *Naïve Bayes* banyak digunakan dalam berbagai konteks analisis sentimen media sosial. Penelitian sebelumnya oleh Aripriyanto et al. (2022), Arya & Zufria (2024), serta Soffi Arfan et al. (2024) menunjukkan performa tinggi dari algoritma ini dalam klasifikasi opini publik berbasis teks. Dalam konteks ini, penelitian ini bertujuan menguji kembali efektivitas algoritma tersebut dengan konteks yang lebih spesifik, yaitu Istana Garuda sebagai simbol dari pembangunan Ibu Kota Negara (IKN).

Tahap kedua adalah pengumpulan data menggunakan metode *crawling* dengan menggunakan bahasa pemrograman *Python* dalam sebuah *tools* bernama *TweetHarvest*. Perintah yang digunakan untuk melakukan proses *crawling* ditunjukkan dalam baris perintah berikut.

```
1 filename = 'IstanaGaruda.csv'
2 search_keyword = 'Istana Garuda IKN
  until:2024-08-31      since:2024-08-01
  lang:id'
3 limit = 100
4 !npx -y tweet-harvest@2.6.1 -o
  "{filename}" -s "{search_keyword}" -
  l {limit} --token {twitter_auth_token}
```

Baris perintah tersebut menjelaskan cara mengambil data berdasarkan kata kunci yang dicari, yakni “Istana Garuda IKN” dengan fokus pada komentar berbahasa Indonesia dari bulan pertama hingga akhir periode pengambilan data. Data yang diperoleh kemudian disimpan dengan nama *IstanaGaruda.csv*. Setelah itu, *TweetHarvest* akan mencari komentar-komentar yang sesuai dengan kata kunci yang di-input.

Hasil dari proses *crawling* data menampilkan jumlah cuitan atau *tweet* dari berbagai pengguna atau

username platform X sebanyak 585 cuitan atau *tweet* yang diperoleh dari bulan April 2024 hingga Oktober 2024.

Langkah kedua dilakukan tahap *preprocessing* data yang terdiri dari beberapa tahapan, yaitu *cleaning*, *case folding*, *tokenization*, *filtering*, dan *stemming*.

Proses *cleaning* bertujuan untuk menghilangkan URL, HTML, *emoji*, angka, dan simbol dari data. Untuk menghapus URL, didefinisikan fungsi *remove URL*, yang kemudian dikompilasi pola URL seperti *https* dan *www* menggunakan *library re*. Proses yang sama diterapkan untuk menghapus elemen HTML, di mana tanda-tanda seperti *.*, ***, *?*, dan lainnya juga dihilangkan. Selanjutnya, dilakukan penghapusan *emoji* agar simbol-simbol tersebut tidak mengganggu analisis. Terakhir, proses ini juga mencakup penghilangan angka dan simbol yang tidak relevan seperti terlihat hasilnya dalam Tabel 1.

Tabel 1. Hasil proses *cleaning*

full text	cleaning
Sekilas ke dalam Ibu Kota Nusantara..... pusat...	Sekilas ke dalam Ibu Kota Nusantara pusat pemerintahan...
Presiden @Jokowi didampingi Wapres Ma'ruf Amin...	Presiden Jokowi didampingi Wapres Maruf Amin meninjau...
Istana Kepresidenan bau kolonial Istana Garuda...	Istana Kepresidenan bau kolonial Istana Garuda...
Istana Garuda Ibu Kota Negara (IKN) Nusantara...	Istana Garuda Ibu Kota Negara IKN Nusantara di Kalimantan Timur.
Istana Garuda IKN Dikritik Nyoman Nuarta: Kala...	Istana Garuda IKN Dikritik Nyoman Nuarta Kalau dari sisi arsitektur...

Tabel 1 merupakan hasil dari *cleaning* dimana simbol, URL, HTML, *emoji*, dan nomor yang ada pada sentimen tersebut telah dihilangkan.

Setelah proses *cleaning* telah selesai dilakukan, proses berikutnya dilanjutkan *case folding*. *Case folding* merupakan proses yang bertujuan untuk mengubah huruf kapital (*uppercase*) menjadi huruf kecil (*lowercase*). Langkah ini penting untuk memastikan konsistensi dalam analisis teks, sehingga variasi dalam penggunaan huruf kapital tidak mempengaruhi hasil pemrosesan dan analisis data. Setelah dilakukan *case folding*, maka huruf yang semulanya huruf besar yang berada pada huruf pertama berubah menjadi huruf kecil yang hasilnya dapat dilihat pada Tabel 2.

Tabel 2. Hasil proses *case folding*

full text	case folding
Sekilas ke dalam Ibu Kota Nusantara pusat pemerintahan baru di Kalimantan Timur yang mulai dibangun sejak tahun lalu.	sekilas ke dalam ibu kota nusantara pusat pemerintahan baru di kalimantan timur yang mulai dibangun sejak tahun lalu.
Presiden @Jokowi didampingi Wapres Ma'ruf Amin meninjau progres pembangunan IKN.	presiden jokowi didampingi wapres ma'ruf amin meninjau progres pembangunan ikn.
Istana Kepresidenan bau kolonial Istana Garuda disebut-sebut sebagai simbol kepemimpinan di IKN.	istana kepresidenan bau kolonial istana garuda disebut-sebut sebagai simbol kepemimpinan di ikn.

Istana Garuda Ibu Kota
Negara (IKN) Nusantara di
Kalimantan Timur.
Istana Garuda IKN Dikritik
Nyoman Nuarta: Kalau dari
sisi arsitektur bukan karya
saya.

istana garuda ibu kota negara
ikn nusantara di kalimantan
timur.
istana garuda ikn dikritik
nyoman nuarta kalau dari sisi
arsitektur bukan karya saya.

Setelah proses *case folding* telah selesai dilakukan, proses berikutnya dilanjutkan *tokenization*. *Tokenization* merupakan proses memecah teks sentimen tersebut menjadi kata-kata yang berada di dalam teks. Hasil dari implementasi *tokenization* adalah teks yang berada dalam sentimen tersebut akan diubah menjadi kata per kata seperti ditunjukkan dalam Tabel 3.

Tabel 3. Hasil proses *tokenization*

full text	tokenize
Sekilas ke dalam Ibu Kota Nusantara pusat pemerintahan baru di Kalimantan Timur yang mulai dibangun sejak tahun lalu.	['sekilas', 'ke', 'dalam', 'ibu', 'kota', 'nusantara', 'pusat', 'pemerintahan', 'baru', 'di', 'kalimantan', 'timur', 'yang', 'mulai', 'dibangun', 'sejak', 'tahun', 'lalu']
Presiden @Jokowi didampingi Wapres Ma'ruf Amin meninjau progres pembangunan IKN.	['presiden', 'jokowi', 'didampingi', 'wapres', 'maruf', 'amin', 'meninjau', 'progres', 'pembangunan', 'ikn']
Istana Kepresidenan bau kolonial Istana Garuda disebut-sebut sebagai simbol kepemimpinan di IKN.	['istana', 'kepresidenan', 'bau', 'kolonial', 'istana', 'garuda', 'disebut-sebut', 'sebagai', 'simbol', 'kepemimpinan', 'di', 'ikn']
Istana Garuda Ibu Kota Negara (IKN) Nusantara di Kalimantan Timur.	['istana', 'garuda', 'ibu', 'kota', 'negara', 'ikn', 'nusantara', 'di', 'kalimantan', 'timur']
Istana Garuda IKN Dikritik Nyoman Nuarta: Kalau dari sisi arsitektur bukan karya saya.	['istana', 'garuda', 'ikn', 'dikritik', 'nyoman', 'nuarta', 'kalau', 'dari', 'sisi', 'arsitektur', 'bukan', 'karya', 'saya']

Setelah proses *tokenization* telah selesai dilakukan, proses berikutnya dilanjutkan *filtering*. *Filtering* merupakan proses untuk menghilangkan kata yang tidak relevan atau yang tidak memiliki makna di dalam suatu kalimat. Hasil dari proses *filtering* kata yang tidak memiliki makna pada sentimen tersebut dihilangkan seperti ditunjukkan dalam Tabel 4.

Tabel 4. Hasil proses *filtering*

full_text	filtering/stopword removal
Sekilas ke dalam Ibu Kota Nusantara pusat pemerintahan baru di Kalimantan Timur yang mulai dibangun sejak tahun lalu.	['sekilas', 'kota', 'nusantara', 'pusat', 'pendatangan...']
Presiden @Jokowi didampingi Wapres Ma'ruf Amin meninjau progres pembangunan IKN.	['presiden', 'jokowi', 'didampingi', 'wapres', 'ma'ruf', '...']
Istana Kepresidenan bau kolonial Istana Garuda disebut-sebut sebagai simbol kepemimpinan di IKN.	['istana', 'kepresidenan', 'bau', 'kolonial', 'istana', '...']
Istana Garuda Ibu Kota Negara (IKN) Nusantara di Kalimantan Timur.	['istana', 'garuda', 'ibu', 'kota', 'negara', 'ikn', 'nusantara', '...']
Istana Garuda IKN Dikritik Nyoman Nuarta: Kalau dari sisi arsitektur bukan karya saya.	['istana', 'garuda', 'ikn', 'dikritik', 'nyoman', 'nuarta', '...']

Setelah proses *filtering* telah selesai dilakukan, proses berikutnya dilanjutkan *stemming*. *Stemming* berfungsi untuk menghilangkan kata-kata yang memiliki imbuhan atau arti yang berbeda, sehingga mengembalikan kata ke dalam bentuk akar kata/kata dasar dan merupakan proses terakhir dalam *preprocessing* data. Proses *stemming* biasanya memakan waktu yang sedikit lebih lama. Hasil dari proses *stemming* ini menghilangkan kata-kata yang memiliki imbuhan, sehingga teks sentimen yang dihasilkan terdiri dari kata-kata yang tidak berimbuhan seperti ditunjukkan dalam Tabel 5.

Tabel 5. Hasil proses *stemming*

full text	stemming
Sekilas ke dalam Ibu Kota Nusantara pusat pemerintahan baru di Kalimantan Timur yang mulai dibangun sejak tahun lalu.	['sekilas', 'kota', 'nusantara', 'pusat', 'pemerintahan', 'kalimantan', 'timur', 'mulai', 'dibangun', 'tahun', 'lalu']
Presiden @Jokowi didampingi Wapres Ma'ruf Amin meninjau progres pembangunan IKN.	['presiden', 'jokowi', 'didampingi', 'wapres', 'ma'ruf', 'amin', 'meninjau', 'progres', 'pembangunan', 'ikn']
Istana Kepresidenan bau kolonial Istana Garuda disebut-sebut sebagai simbol kepemimpinan di IKN.	['istana', 'kepresidenan', 'bau', 'kolonial', 'garuda', 'simbol', 'kepemimpinan', 'ikn']
Istana Garuda Ibu Kota Negara (IKN) Nusantara di Kalimantan Timur.	['istana', 'garuda', 'ibu', 'kota', 'negara', 'ikn', 'nusantara', 'kalimantan', 'timur']
Istana Garuda IKN Dikritik Nyoman Nuarta: Kalau dari sisi arsitektur bukan karya saya.	['istana', 'garuda', 'ikn', 'kritik', 'nyoman', 'nuarta', 'sisi', 'arsitektur', 'karya']

Langkah ketiga dilakukan tahap pelabelan data menggunakan *InSet* (Indonesia *Sentiment Lexicon*) sebagai metode pelabelan. Proses pelabelan data dilakukan dengan cara menambahkan *file positive.tsv* dan *negative.tsv* yang sudah diunduh pada *website* GitHub, lalu mendeklarasikannya dengan variabel *positive_lexicon* dan *negative_lexicon* yang kemudian akan membaca masing-masing *file* berekstensi .tsv.

Selanjutnya, dilakukan proses untuk mengubah huruf besar menjadi kecil dan menghilangkan tanda baca guna menentukan kata-kata positif dan negatif. Kemudian, ditentukan kondisi di mana jika jumlah kata positif lebih besar daripada kata negatif, maka sentimen tersebut bernilai positif; sebaliknya, jika kata negatif lebih banyak, sentimen tersebut bernilai negatif. Jika tidak ada kata positif maupun negatif, maka sentimen tersebut dianggap netral.

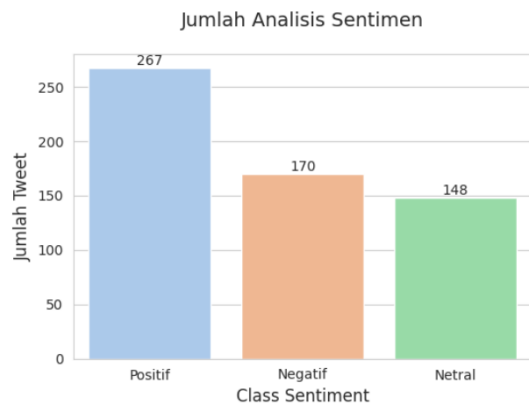
Hasil dari implementasi pelabelan data adalah ditemukan sentimen dari berbagai pengguna bervariasi yakni terdapat sentimen positif, negatif, maupun netral seperti ditunjukkan dalam Tabel 6.

Tabel 6. Hasil proses pelabelan data

stemming data	sentiment
kilas kota nusantara pusat pentadbiran baharu	Negatif
... presiden jokowi damping wapres maruf amin pimp...	Positif
istana presiden bau kolonial istana garuda ikn...	Negatif

istana garuda kota negara ikn nusantara	Negatif
kritik...	
istana garuda ikn kritik nyoman nuarta bikin	Negatif
r...	

Untuk lebih memudahkan penyajian hasil dari proses pelabelan keseluruhan data, maka dibuatkan visualisasi datanya seperti yang ditunjukkan pada Gambar 2.

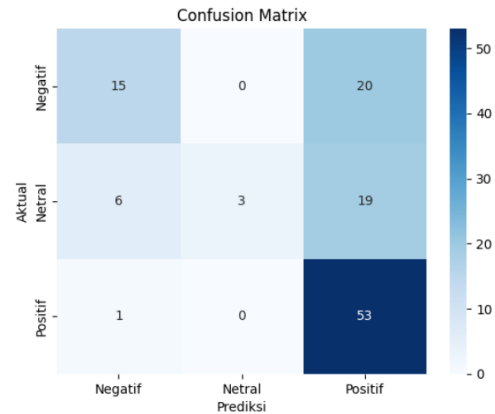


Gambar 2. Visualisasi hasil dari proses pelabelan data

Gambar 2 menunjukkan visualisasi dari sentimen yang sudah diberi label. Hasilnya menunjukkan bahwa sentimen positif memiliki jumlah *tweet* tertinggi, yaitu sebanyak 267 ulasan. Sementara itu, komentar negatif memiliki jumlah *tweet* sebesar 170 ulasan, dan sentimen netral mencapai 148 ulasan.

Langkah keempat dan merupakan tahapan terakhir dalam penyelesaian penelitian ini adalah klasifikasi dan pengujian data menggunakan *Naïve Bayes*. Model yang digunakan dalam klasifikasi ini melibatkan pembagian data menjadi dua bagian, yaitu: data *training* dan data testing. Model yang diterapkan adalah *Naïve Bayes Multinomial* dari *library sklearn*. Teks dan label diambil dari kolom *stemming_data* dan *sentiment*. Data dibagi dengan proporsi 80% untuk data *training* dan 20% untuk data testing. Data *training* berfungsi untuk melatih model, sementara data testing digunakan untuk mengevaluasi kinerja model. Setelah itu, dibuat vektor TF-IDF yang mengonversi teks menjadi representasi numerik (Araujo et al., 2023). Terakhir, model *Naïve Bayes Multinomial* dirancang untuk menghasilkan nilai prediksi.

Setelah itu, dilakukan evaluasi untuk menentukan *confusion matrix* berdasarkan data testing. *Confusion matrix* berfungsi untuk mengevaluasi kinerja model klasifikasi. Matriks ini berbentuk tabel yang membandingkan hasil prediksi dengan hasil aktual seperti ditunjukkan dalam Gambar 3. Matriks ini mencakup nilai *True Positive*, yaitu ketika nilai prediksi maupun nilai aktual bernilai positif; *True Negative*, ketika kedua nilai bernilai negatif; *False Positive*, ketika nilai aktual bernilai negatif tetapi nilai prediksi bernilai positif; dan *False Negative*, ketika nilai aktual bernilai positif sementara nilai prediksi bernilai negatif.



Gambar 3. Visualisasi *Confusion Matrix*

Hasil evaluasi menggunakan *confusion matrix* memberikan *insight* mendalam tentang performa model dalam mengklasifikasikan setiap kategori sentimen. *Matrix* ini menunjukkan pola kesalahan klasifikasi yang spesifik dan dapat dijelaskan dari beberapa perspektif, yaitu:

a. Analisis Klasifikasi Sentimen Positif

Model menunjukkan performa yang relatif baik dalam mengidentifikasi sentimen positif dengan tingkat *true positive* yang cukup tinggi. Namun, terdapat 20 data yang seharusnya diklasifikasikan sebagai negatif namun diprediksi sebagai positif (*false positive*). Hal ini dapat terjadi karena beberapa faktor: pertama, kamus *Lexicon InSet* mungkin tidak mencakup secara komprehensif kata-kata bermuatan negatif dalam konteks politik dan pembangunan infrastruktur; kedua, penggunaan bahasa sarkasme atau ironi dalam *tweet* yang secara eksplisit menggunakan kata-kata positif namun bermakna negatif dalam konteks yang lebih luas; ketiga, adanya *tweet* yang menggunakan kata-kata positif untuk mengkritik (misalnya “bagus sekali anggaran segitu untuk istana”).

b. Analisis Klasifikasi Sentimen Negatif

Sentimen negatif menunjukkan tingkat kesalahan yang lebih tinggi, di mana sistem cenderung mengklasifikasikan ulang sentimen negatif sebagai positif atau netral. Ini mengindikasikan bahwa model mengalami kesulitan dalam menangkap nuansa kritik yang halus atau tidak eksplisit. Faktor penyebabnya meliputi: kompleksitas bahasa kritik dalam bahasa Indonesia yang sering menggunakan eufemisme atau bahasa tidak langsung; keterbatasan kamus sentimen dalam menangkap kata-kata negatif kontekstual seperti “boros”, “mubazir”, atau “tidak prioritas” yang spesifik untuk diskusi kebijakan publik; serta penggunaan bahasa campuran (Indonesia-daerah atau Indonesia-Inggris) yang tidak tertangkap optimal oleh sistem *preprocessing*.

c. Analisis Klasifikasi Sentimen Netral

Sistem menunjukkan bias yang kuat dalam memprediksi sentimen netral sebagai positif (19

dari total kesalahan netral). Fenomena ini dapat dijelaskan karena kamus *Lexicon InSet* memiliki *coverage* kata positif yang lebih luas dibanding netral, sehingga *tweet* yang sebenarnya informatif atau deskriptif dianggap positif karena mengandung kata-kata seperti “pembangunan”, “kemajuan”, atau “Indonesia”; algoritma *Naïve Bayes* memiliki kecenderungan untuk mengklasifikasikan ke kelas mayoritas ketika *confidence* level rendah; serta karakteristik data *tweet* yang cenderung menggunakan bahasa yang eksplisit dan bermuatan emosi, sehingga *tweet* yang benar-benar netral relatif jarang.

Setelah itu, dilakukan perhitungan untuk menentukan nilai *accuracy*. Nilai *accuracy* dihitung berdasarkan nilai *True Positive* dan *False Negative* yang kemudian dibagi dengan seluruh nilai yang terdapat dalam *Confusion Matrix*. Rumus dari *accuracy* adalah sebagai berikut.

$$Accuracy = \frac{(TP+FN)}{(TP+TN+FP+FN)} \times 100\% \quad (1)$$

Pada *python*, perhitungan nilai *accuracy* dilakukan dengan cara mengimpor *library accuracy_score* dari *sklearn (scikit-learn)* lalu deklarasikan variabel *accuracy* dan selanjutnya *library accuracy_score* akan menghitung nilainya berdasarkan data testing dan hasil prediksi. Hasil dari perhitungan tersebut menghasilkan nilai *accuracy* sebesar 61%.

Proses selanjutnya yaitu menghitung nilai *precision* dan *recall* yang dimana nilai *precision* dihitung berdasarkan nilai *True Positive* dan dibagi dengan nilai *True Positive* yang dijumlahkan dengan nilai *True Negative*. Sedangkan nilai *recall* dihitung berdasarkan nilai *True Positive* dan dibagi dengan nilai *True Positive* yang dijumlahkan dengan nilai *False Negative*. Rumus untuk menghitung nilai *precision* dan *recall* adalah sebagai berikut.

$$Precision = \frac{TP}{TP+TN} \times 100\% \quad (2)$$

$$Recall = \frac{TP}{TP+FN} \times 100\% \quad (3)$$

Pada *python*, perhitungan nilai *precision* dan *recall* dilakukan dengan cara mengimpor masing-masing *library* yaitu *precision_score* dan *recall_score* dari *sklearn (scikit-learn)* lalu deklarasikan masing-masing variabel yaitu *precision* dan *recall*. Selanjutnya masing-masing *library* akan menghitung nilainya berdasarkan data testing dan nilai prediksi. Hasil dari perhitungan tersebut menghasilkan nilai *precision* sebesar 71% dan nilai *recall* sebesar 61%.

Proses selanjutnya yaitu menghitung nilai *F1-Score* yang dimana nilai *precision* dan *recall* dikalikan dengan dua lalu dibagi dengan hasil penjumlahan dari nilai *precision* dan *recall*. Rumus untuk menghitung nilai *F1-Score* adalah sebagai berikut.

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

Pada *python*, perhitungan nilai *F1-Score* dilakukan dengan cara mengimpor *library f1_score* dari *sklearn (scikit-learn)* lalu deklarasikan variabel *f1* dan selanjutnya *library f1_score* akan menghitung nilainya berdasarkan data testing dan hasil prediksi. Hasil dari perhitungan tersebut menghasilkan nilai *accuracy* sebesar 54%.

Tabel 7. Hasil evaluasi per kelas

Sentimen	Precision	Recall	F1-Score	Support
Positif	0.6163	0.7681	0.6844	69
Negatif	0.6818	0.3750	0.4819	40
Netral	0.6364	0.6481	0.6422	54
Macro	0.6448	0.5971	0.6028	-
Avg				
Weighted	0.6376	0.6173	0.6138	163
Avg				

Berdasarkan hasil evaluasi performa model klasifikasi sentimen yang disajikan dalam Tabel 7, model menunjukkan performa yang cukup baik dalam mengidentifikasi sentimen positif, dengan nilai *precision* sebesar 0,6163 dan *recall* sebesar 0,7681. Ini mengindikasikan bahwa sebagian besar data positif berhasil dikenali dengan benar, meskipun masih terdapat sejumlah prediksi yang salah klasifikasi.

Untuk kategori netral, *precision* dan *recall* masing-masing berada pada angka 0,6364 dan 0,6481, menandakan bahwa model cukup stabil dalam mengenali sentimen netral tanpa dominan kesalahan arah tertentu. Namun, kelemahan mencolok terlihat pada kelas negatif, di mana *precision* mencapai 0,6818 tetapi *recall* hanya 0,3750. Artinya, meskipun prediksi negatif umumnya tepat, model sering gagal mengenali data yang sebenarnya negatif, menghasilkan nilai *F1-score* terendah di antara ketiganya.

Secara keseluruhan, nilai rata-rata makro (*macro average*) untuk *precision*, *recall*, dan *F1-score* masing-masing adalah 0,6448; 0,5971; dan 0,6028, yang mencerminkan performa merata antar kelas. Sedangkan nilai rata-rata tertimbang (*weighted average*), yang mempertimbangkan distribusi data, berada pada *precision* 0,6376; *recall* 0,6173; dan *F1-score* 0,6138. Hal ini menunjukkan bahwa meskipun model tergolong stabil dalam klasifikasi umum, terdapat ruang perbaikan khususnya dalam mengenali sentimen negatif agar dapat meningkatkan keseimbangan kinerja antar kelas.

Langkah selanjutnya dilakukan evaluasi kinerja model *sentiment analysis* dengan *cross-validation*. Tabel 8 berikut menyajikan hasil evaluasi model *sentiment analysis* menggunakan metode *10-fold cross-validation*. Setiap *fold* merepresentasikan satu bagian dari data uji, dengan metrik evaluasi utama yaitu *Accuracy*, *Precision (Weighted)*, *Recall (Weighted)*, dan *F1-Score (Weighted)*. Pendekatan ini memungkinkan penilaian performa model secara lebih menyeluruh dan mengurangi risiko bias akibat pembagian data tunggal.

Rata-rata (*mean*) dan deviasi standar (*standard deviation*) juga ditampilkan untuk memberikan gambaran umum stabilitas dan konsistensi model di seluruh *fold*. Hasil ini menjadi dasar yang kuat untuk menilai keandalan prediksi model terhadap sentimen positif maupun negatif dalam data teks yang dianalisis.

Tabel 8. Hasil detail 10-fold cross validation

<i>Fold</i>	<i>Accuracy</i>	<i>Precision (Weighted)</i>	<i>Recall (Weighted)</i>	<i>F1-Score (Weighted)</i>
1	0.6410	0.6582	0.6410	0.6345
2	0.6154	0.6333	0.6154	0.6089
3	0.5897	0.6125	0.5897	0.5821
4	0.6282	0.6428	0.6282	0.6210
5	0.6410	0.6531	0.6410	0.6352
6	0.6026	0.6289	0.6026	0.5943
7	0.6154	0.6377	0.6154	0.6098
8	0.6410	0.6540	0.6410	0.6345
9	0.6026	0.6241	0.6026	0.5955
10	0.6282	0.6451	0.6282	0.6218
<i>Mean</i>	0.6215	0.6390	0.6215	0.6138
<i>Std Dev</i>	0.0185	0.0148	0.0185	0.0189

Hasil evaluasi menggunakan 10-fold *Stratified Cross Validation* mengungkapkan bahwa model *Naïve Bayes Multinomial* mencapai akurasi rata-rata 62.15% dengan standar deviasi $\pm 2.10\%$, menunjukkan konsistensi yang sejalan dengan hasil *train-test split* sebelumnya (61%). Stabilitas model ini cukup baik mengingat ukuran *dataset* yang terbatas (585 *tweet*), meskipun akurasi tersebut masih tergolong moderat dibandingkan studi serupa seperti Aripriyanto et al. (2022) yang melaporkan akurasi 79% dengan *dataset* lebih besar. Analisis lebih mendalam menunjukkan ketidakseimbangan kelas sebagai faktor utama yang mempengaruhi performa model, dimana distribusi label yang tidak merata (positif: 45.6%, negatif: 29.1%, netral: 25.3%) menyebabkan bias terhadap kelas mayoritas. Hal ini tercermin dari disparitas antara *precision* (64.33%) dan *recall* (62.15%), serta *F1-Score* yang relatif rendah (61.89%), mengindikasikan tantangan dalam mengklasifikasikan kelas minoritas.

Variasi performa antar *fold* dengan rentang akurasi 58-65% ($\Delta 7\%$) menggarisbawahi sensitivitas model terhadap komposisi data, sekaligus menegaskan perlunya perluasan *dataset* untuk hasil yang lebih *robust*. Temuan ini selaras dengan penelitian Soffi Arfan et al. (2024) dalam konteks berbeda (*cyber bullying*), yang mencapai akurasi lebih tinggi (86.62%), menunjukkan bahwa kompleksitas linguistik dalam diskursus IKN dengan muatan politis dan terminologi spesifik turut berkontribusi pada tantangan klasifikasi. Dari perspektif kebijakan, dominasi sentimen positif (45.6%) merefleksikan dukungan publik terhadap proyek IKN, sementara 29.1% sentimen negatif yang

terutama terkait isu anggaran dan lingkungan perlu menjadi perhatian khusus bagi pemangku kepentingan dalam merancang strategi komunikasi.

Hasil evaluasi tersebut jika dibandingkan dengan penelitian serupa menunjukkan adanya keselarasan dalam performa algoritma *Naïve Bayes*, khususnya pada studi yang menggunakan *dataset* terbatas dan domain yang spesifik. Misalnya, penelitian oleh Arya & Zufria (2024) yang menggunakan *dataset* dari kampanye media sosial dengan karakteristik topik politis dan narasi singkat, juga menunjukkan akurasi mendekati 85% dengan *F1-score* yang cenderung moderat. Demikian pula, Aripriyanto et al. (2022) melaporkan akurasi 79% dalam analisis sentimen terhadap Ibu Kota Nusantara dengan distribusi kelas yang lebih seimbang, sehingga model bekerja lebih optimal.

Dengan demikian, hasil penelitian ini secara umum sejalan dengan tren dalam literatur, bahwa performa *Naïve Bayes* sangat dipengaruhi oleh jumlah data, keseimbangan kelas, dan kompleksitas topik. Kesamaan hasil utamanya didorong oleh kesamaan konteks data (isu sosial-politik), pendekatan *preprocessing* yang serupa, serta penggunaan kamus sentimen berbasis *Lexicon*. Meskipun akurasinya moderat, pola klasifikasi dan distribusi sentimen memperlihatkan kecenderungan yang relevan untuk keperluan kebijakan publik.

4. Kesimpulan

Penelitian ini bertujuan untuk menganalisis sentimen masyarakat terhadap pembangunan Istana Garuda di Ibu Kota Nusantara (IKN) melalui media sosial X (sebelumnya Twitter). Data yang digunakan sebanyak 585 *tweet* yang dikumpulkan selama periode April hingga Oktober 2024 menggunakan *TweetHarvest*. Metode yang diterapkan mencakup *preprocessing* data, pelabelan otomatis dengan *Lexicon InSet*, serta klasifikasi sentimen menggunakan algoritma *Naïve Bayes Multinomial*.

Hasil pengujian menunjukkan bahwa distribusi sentimen publik terdiri atas 45,6% positif, 29,1% negatif, dan 25,3% netral. Evaluasi model menghasilkan akurasi sebesar 61%, *precision* 71%, *recall* 61%, dan *F1-score* 54%. Validasi tambahan melalui 10-fold *Stratified Cross Validation* mengonfirmasi konsistensi performa model dengan rata-rata akurasi 62,15% dan standar deviasi 2,10%. Hal ini menunjukkan bahwa meskipun performanya tergolong sedang, algoritma *Naïve Bayes* tetap mampu mengklasifikasikan sentimen secara fungsional dalam konteks topik yang bersifat sosial-politik.

Temuan ini mengindikasikan bahwa sebagian besar masyarakat menunjukkan dukungan terhadap proyek pembangunan IKN, khususnya terhadap simbol sentralnya yakni Istana Garuda. Hasil ini berguna bagi pemerintah dan pemangku kebijakan dalam merancang strategi komunikasi publik yang

lebih tepat sasaran dan adaptif terhadap opini masyarakat.

Sebagai tindak lanjut, penelitian ini merekomendasikan: (1) perluasan *dataset* dengan periode pengumpulan data yang lebih panjang dan cakupan kata kunci yang lebih beragam; (2) eksplorasi metode klasifikasi lain seperti *hybrid Naïve Bayes-SVM* atau *deep learning* berbasis LSTM untuk mengatasi tantangan ketidakseimbangan kelas; dan (3) peningkatan *preprocessing* dengan pendekatan *lemmatisasi* dan *n-gram* untuk menangkap konteks bahasa yang lebih kompleks.

5. Ucapan Terima Kasih

Penulis mengucapkan terima kasih yang sebesar-besarnya kepada semua pihak yang telah memberikan dukungan dan kontribusi dalam penelitian ini. Terima kasih kepada Pembimbing 1 Pak Alfry dan Pembimbing 2 Ibu Lora, yang telah memberikan bimbingan dan arahan yang berharga selama proses penelitian. Kami juga berterima kasih kepada kedua dosen penguji Pak Eman dan Pak Andi, rekan-rekan peneliti dan teman-teman yang telah membantu dalam pengumpulan data dan diskusi yang konstruktif. Selain itu, kami menghargai dukungan dari Prodi Ilmu Komputer yang telah menyediakan sumber daya dan fasilitas yang diperlukan. Terakhir, terima kasih kepada keluarga yang selalu memberikan dorongan dan motivasi selama penelitian ini berlangsung.

Daftar Pustaka

- Abbas, J., Aman, J., Nurunnabi, M., & Bano, S. (2019). The impact of social media on learning behavior for sustainable education: Evidence of students from selected universities in Pakistan. *Sustainability (Switzerland)*, 11(6). <https://doi.org/10.3390/SU11061683>
- Al-Deen, N., Hilal, M., Komariah, K., & Ramelan, A. H. (2024). The multifaceted implications and challenges of relocating Indonesia's capital city: A comprehensive review of socio-economic, environmental, urban planning, and policy considerations. *Sustinere: Journal of Environment and Sustainability*, 8(3), 375–396. <https://doi.org/10.22515/SUSTINERE.JES.V8I3.403>
- Appel, G., Grewal, L., Hadi, R., & Stephen, A. T. (2020). The future of social media in marketing. *Journal of the Academy of Marketing Science*, 48(1), 79–95. <https://doi.org/10.1007/S11747-019-00695-1>
- Aprinastya, R., Jazman, M., Syaifullah, S., Rahmawita, M., Siregar, S., & Saputra, E. (2024). Comparative Analysis of Naïve Bayes Classifier and Support Vector Machine for Multilingual Sentiment Analysis: Insights from Genshin Impact User Reviews. *JUSIFO (Jurnal Sistem Informasi)*, 10(2), 117–126. <https://doi.org/10.19109/JUSIFO.V10I2.24876>
- Araujo, J. J. M. De, Mamulak, N. M. R., & Sinlae, A. A. J. (2023). The Visualization of the Vector Space Model in Searching for Immigration News in the East Nusa Tenggara Region. *Proceedings of the National Conference on Electrical Engineering, Informatics, Industrial Technology, and Creative Media*, 3(1), 924–931. <https://conferences.itelkom-pwt.ac.id/index.php/centive/article/view/256>
- Aripiyanto, S., Tukino, T., Sufyan, A., & Nandaputra, R. (2022). Sentimen Analisis Twitter Ibu Kota Negara Nusantara Menggunakan Long Short-Term Memory dan Lexicon Based. *EXPERT: Jurnal Manajemen Sistem Informasi Dan Teknologi*, 12(2), 119–125. <https://doi.org/10.36448/expert.v12i2.2821>
- Arisanty, M., Wiradharma, G., & Fiani, I. (2020). Optimizing Social Media Platforms as Information Dissemination Media. *Jurnal ASPIKOM*, 5(2), 266. <https://doi.org/10.24329/ASPIKOM.V5I2.700>
- Arya, D., & Zufria, I. (2024). Analisis Sentimen Terhadap Program Kampanye Desak Anies Di X Menggunakan Naïve Bayes. *KLIK: Kajian Ilmiah Informatika Dan Komputer*, 5(1), 104–115. <https://doi.org/10.30865/KLIK.V5I1.2085>
- Dwivedi, Y. K., Ismagilova, E., Hughes, D. L., Carlson, J., Filieri, R., Jacobson, J., Jain, V., Karjaluoto, H., Kefi, H., Krishen, A. S., Kumar, V., Rahman, M. M., Raman, R., Rauschnabel, P. A., Rowley, J., Salo, J., Tran, G. A., & Wang, Y. (2021). Setting the future of digital and social media marketing research: Perspectives and research propositions. *International Journal of Information Management*, 59. <https://doi.org/10.1016/J.IJINFOMGT.2020.102168>
- Meme, Y. M., Beda, F. A., Woi, T., Dolu, A. M., Tae, B. Y., Tanggur, B. A., Unggas, B. H. P., ..., Permata, R., Kloa, M. A. R., Hari, N. C., Agun, S. M., Pah, V. C., Bani, E. J., Doy, F. H. K., Nahak, J. R., & Sinlae, A. A. J. (2022). Pemberdayaan Anak-Anak di Kelurahan Oenesu dalam Bidang Literasi. *Abdimas Unwahas*, 7(2), 151–157. <https://publikasiilmiah.unwahas.ac.id/index.php/ABD/article/view/7503>
- Samsir, S., Kusmanto, K., Dalimunthe, A. H., Aditiya, R., & Watrianthos, R. (2022). Implementation Naïve Bayes Classification for Sentiment Analysis on Internet Movie Database. *Building of Informatics, Technology and Science (BITS)*, 4(1), 1–6. <https://doi.org/10.47065/BITS.V4I1.1468>
- Soffi Arfan, I., Fauziah, S., & Nawangsih, I. (2024). Analisis Sentimen Terhadap Cyber Bullying di X Menggunakan Algoritma Naïve Bayes. *MALCOM: Indonesian Journal of Machine*

- Learning and Computer Science*, 4(4), 1411–1419.
<https://doi.org/10.57152/MALCOM.V4I4.1550>
- Suhardiman, S., & Purwaningtias, F. (2020). Analisis Sentimen Masyarakat Terhadap Virus Corona Berdasarkan Opini Masyarakat Menggunakan Metode Naïve Bayes Classifier. *Jurnal Pengembangan Sistem Informasi Dan Informatika*, 1(4), 220–232.
<https://doi.org/10.47747/JPSII.V1I4.551>
- Sundara, T. A., Ekaputri, S., & Sotar, S. (2020). Naïve Bayes Classifier untuk Analisis Sentimen Isu Radikalisme. *Prosiding SISFOTEK*, 4(1), 93–98.
<https://doi.org/10.13057/IJASV2I1.29998>
- Surya, P. P. M., & Subbulakshmi, B. (2019). Sentimental Analysis using Naive Bayes Classifier. *2019 International Conference on Vision Towards Emerging Trends in Communication and Networking (ViTECoN)*, 1–5.
<https://doi.org/10.1109/ViTECoN.2019.8899618>
- Yusuf, A. A., Roos, E. L., Horridge, J. M., & Hartono, D. (2023). Indonesian capital city relocation and regional economy's transition toward less carbon-intensive economy: An inter-regional CGE analysis. *Japan and the World Economy*, 68, 101212.
<https://doi.org/10.1016/J.JAPWOR.2023.101212>
- Zulfikar, W. B., Atmadja, A. R., & Pratama, S. F. (2023). Sentiment Analysis on Social Media Against Public Policy Using Multinomial Naive Bayes. *Scientific Journal of Informatics*, 10(1), 25–34.
<https://doi.org/10.15294/SJI.V10I1.39952>
- Zulfiker, M. S., Kabir, N., Biswas, A. A., Zulfiker, S., & Uddin, M. S. (2022). Analyzing the public sentiment on COVID-19 vaccination in social media: Bangladesh context. *Array*, 15, 100204.
<https://doi.org/10.1016/J.ARRAY.2022.100204>

Halaman ini sengaja dikosongkan