

Analisis Sentimen Pengguna X terhadap IKN Menggunakan Word2Vec dan IndoBERT

Luthfiana Azzalea Afifah*, Dedi Gunawan

Teknik Informatika, Komunikasi dan Informatika, Universitas Muhammadiyah Surakarta, Surakarta, Indonesia

Email: ^{1,*}l200220052@student.ums.ac.id, ²dedi.gunawan@ums.ac.id

Email Penulis Korespondensi: l200220052@student.ums.ac.id*

Submitted: 07/01/2026; Accepted: 22/01/2026; Published: 31/03/2026

Abstrak– Pembangunan Ibu Kota Nusantara (IKN) di Kalimantan Timur adalah kebijakan strategis nasional yang menimbulkan beragam respons di tengah masyarakat. Media sosial X (Twitter) salah satu ruang utama bagi publik dengan tujuan menyampaikan pandangan, baik berupa dukungan ataupun kritik, terhadap kebijakan tersebut. Fokus penelitian ini adalah untuk mengevaluasi persepsi masyarakat tentang pembangunan IKN dengan membandingkan kinerja algoritma Word2Vec yang diklasifikasikan menggunakan Support Vector Machine (SVM) dan model IndoBERT Fine-Tuning. Data penelitian diperoleh dengan proses crawling tweet berbahasa Indonesia pada periode 2022 hingga 2025 dengan memanfaatkan istilah yang relevan terhadap pembangunan IKN. Sebanyak 2.667 dataset yang telah menjalani tahap pelabelan manual dan otomatis serta preprocessing, yang terdiri dari text cleaning, case folding, normalisasi, dan tokenisasi, digunakan dalam proses pemodelan. Dataset kemudian dibagi menjadi data latih dan data uji dengan rasio 80:20. Evaluasi performa model dilakukan menggunakan confusion matrix dengan metrik accuracy, precision, recall, dan F1-score. Hasil penelitian menunjukkan bahwa model IndoBERT Fine-Tuning dengan pelabelan manual menghasilkan performa terbaik dengan nilai accuracy sebesar 0.900 dan F1-score sebesar 0.899 serta mampu mengungguli pendekatan Word2Vec–SVM. Analisis ini menegaskan bahwa model berbasis Transformer lebih efektif dalam memahami konteks bahasa Indonesia pada data media sosial. Hasil analisis sentimen selanjutnya disajikan dalam bentuk dashboard interaktif untuk memudahkan interpretasi opini publik terhadap pembangunan Ibu Kota Nusantara.

Kata Kunci: Analisis sentimen; IKN; Deep Learning; Word2Vec; IndoBERT

Abstract– The development of the Indonesian Capital City (IKN) in East Kalimantan is a national strategic policy that has generated diverse responses among the public. Social media (Twitter) is one of the main platforms for the public to express their views, both in support and criticism, regarding this policy. The focus of this research is to evaluate public perception of the establishment of the IKN by comparing the performance of the Word2Vec algorithm classified using a Support Vector Machine (SVM) and the IndoBERT Fine-Tuning model. The research data was obtained by crawling Indonesian-language tweets from 2022 to 2025, utilizing terms relevant to the development of the IKN. A total of 2,667 datasets that have undergone manual and automatic labeling stages and preprocessing, consisting of text cleaning, case folding, normalization, and tokenization, were used in the modeling process. The dataset was then divided into training data and test data with a ratio of 80:20. Model performance evaluation was carried out using a confusion matrix with accuracy, precision, recall, and F1-score metrics. The results showed that the IndoBERT Fine-Tuning model with manual labeling produced the best performance, with an accuracy value of 0.900 and an F1-score of 0.899, and outperformed the Word2Vec–SVM approach. This analysis confirmed that the Transformer-based model is more effective in understanding the Indonesian language context in social media data. The sentiment analysis results are then presented in an interactive dashboard to facilitate the interpretation of public opinion on the development of the Indonesian capital.

Keywords: Analysis sentiment; IKN; Deep Learning; Word2Vec; IndoBERT

1. PENDAHULUAN

Pembangunan nasional merupakan salah satu upaya yang dilakukan pemerintah untuk mewujudkan pemerataan kesejahteraan di seluruh wilayah Indonesia. Dalam konteks ini, pemindahan Ibu Kota Negara adalah bagian dari kebijakan besar yang bertujuan untuk menciptakan pusat pemerintahan yang lebih efisien, merata, dan berkelanjutan. Pemerintah Indonesia menetapkan Kalimantan Timur menjadi lokasi Ibu Kota Nusantara (IKN) mengingat wilayah ini memenuhi sejumlah kriteria strategis yang mendukung. Dari sisi infrastruktur, wilayah ini memiliki sarana transportasi utama berupa Pelabuhan dan bandara, lokasi yang dekat dengan dua kota besar, yaitu Balikpapan dan Samarinda, serta tingkat potensi konflik yang rendah [1]. Pertimbangan tersebut menjadikan Kalimantan Timur sebagai kawasan dengan aksesibilitas yang tinggi dan dinilai layak untuk dijadikan Ibu Kota Negara baru Indonesia.

Meski demikian, kebijakan IKN memunculkan beragam opini di masyarakat. Akibat dari fenomena ini, muncul berbagai sentimen di media sosial, termasuk X (Twitter). Platform ini termasuk salah satu media sosial yang populer dimanfaatkan publik untuk membagikan informasi terkini, menyuarakan pendapat, serta menanggapi isu-isu nasional yang sedang ramai diperbincangkan [2]. Melalui X, masyarakat dapat dengan bebas menyampaikan pandangan mereka, baik dalam bentuk pendapat positif maupun negatif terhadap kebijakan pemerintah [3]. Aktivitas tersebut menunjukkan bahwa media sosial kini berperan sebagai ruang interaktif bagi masyarakat untuk

berpartisipasi dalam pembahasan pembangunan nasional, khususnya dalam menanggapi kebijakan pemindahan dan pembangunan Ibu Kota Nusantara.

Beragam opini di X mencerminkan perbedaan cara pandang masyarakat terhadap pemindahan dan pembangunan IKN. Sebagian masyarakat menilai pembangunan besar-besaran ini menimbulkan potensi kerusakan ekosistem, ancaman terhadap keragaman hayati, sekaligus meningkatnya potensi bencana alam seperti banjir dan longsor [4]. Dalam perspektif lain, kebijakan ini dianggap sebagai strategi pemerataan pembangunan, pengembangan ekonomi lokal, serta pemanfaatan digitalisasi dan teknologi di wilayah baru sebagai upaya mendorong kemajuan bangsa.

Berbagai pandangan yang berkembang di X membuka peluang untuk mengidentifikasi dan mengelompokkan sentimen masyarakat terhadap kebijakan IKN. Analisis sentimen dilakukan dengan menerapkan teknik *Natural Language Processing (NLP)* dan *text mining* [5]. Untuk mengenali serta mengevaluasi keadaan emosional dalam teks, analisis sentimen dapat digunakan untuk memahami perilaku dan persepsi pengguna terhadap suatu topik, khususnya pada platform media sosial seperti Twitter. Melalui analisis sentimen, dapat diketahui sejauh mana opini masyarakat menunjukkan sentimen positif atau negatif terhadap kebijakan tersebut. Namun, data yang diperoleh dari Twitter bersifat tidak terstruktur, sehingga memerlukan tahapan *preprocessing* seperti pembersihan data, *embedding*, dan normalisasi data dalam analisis sentimen [6]. Menurut [7], *text data mining* meliputi tahapan *preprocessing*, representasi fitur, serta teknik klasifikasi dan *clustering* dalam menganalisis teks tidak terstruktur seperti media sosial.

Penerapan teknik *Natural Language Processing (NLP)* dan *text mining* dalam analisis sentimen telah banyak dilakukan pada penelitian sebelumnya dengan berbagai pendekatan algoritma klasifikasi. Penelitian terdahulu menganalisis ulasan produk pada *marketplace* menggunakan algoritma *Naive Bayes* dan *Support Vector Machine (SVM)*, di mana hasil penelitian menunjukkan bahwa *SVM* memperoleh akurasi sebesar 86,14%, lebih tinggi dibandingkan *Naive Bayes* yang mencapai 81,37%, sehingga membuktikan keunggulan *SVM* dalam klasifikasi sentimen teks [8]. Selanjutnya, penelitian oleh [9] mengombinasikan representasi fitur *Word2Vec* dengan model *Long Short-Term Memory (LSTM)* untuk menganalisis sentimen pada data Twitter dan menunjukkan bahwa penggunaan *Word2Vec* mampu meningkatkan kualitas representasi teks serta performa klasifikasi sentimen pada data media sosial.

Perkembangan metode berbasis *deep learning* kemudian mendorong penggunaan *model transformer*, seperti *IndoBERT*, dalam analisis sentimen berbahasa Indonesia. Penelitian oleh [10] menggunakan model *IndoBERT fine-tuning* untuk menganalisis ulasan pengguna aplikasi layanan kesehatan di Google Play Store dan memperoleh nilai *accuracy* sebesar 96%, *precision* 95%, *recall* 96%, serta *F1-score* 95%, yang mengindikasikan efektivitas *IndoBERT* dalam klasifikasi sentimen teks berbahasa Indonesia. Selain itu, penelitian oleh [11] juga memanfaatkan *IndoBERT* untuk menganalisis sentimen publik terhadap program kebijakan pemerintah pada platform X (Twitter) dan menunjukkan bahwa *IndoBERT* memiliki kinerja yang lebih baik dibandingkan metode klasifikasi konvensional.

Meskipun analisis sentimen telah banyak dikembangkan dengan berbagai pendekatan metode, penelitian sebelumnya umumnya dilakukan pada ranah penelitian yang berbeda-beda, seperti ulasan produk *marketplace*, aplikasi digital, maupun isu sosial tertentu di media sosial. Namun, penelitian yang secara khusus mengkaji sentimen masyarakat terhadap kebijakan publik strategis, khususnya terkait pemindahan dan pembangunan Ibu Kota Nusantara (IKN), masih relatif terbatas. Selain itu, meskipun pendekatan *machine learning* dan *deep learning* telah banyak digunakan secara terpisah, belum ditemukan penelitian yang secara langsung membandingkan kinerja metode *Word2Vec-SVM* dengan *model transformer* berbasis *IndoBERT fine-tuning* dalam menganalisis sentimen masyarakat terhadap kebijakan pembangunan IKN pada platform X (Twitter).

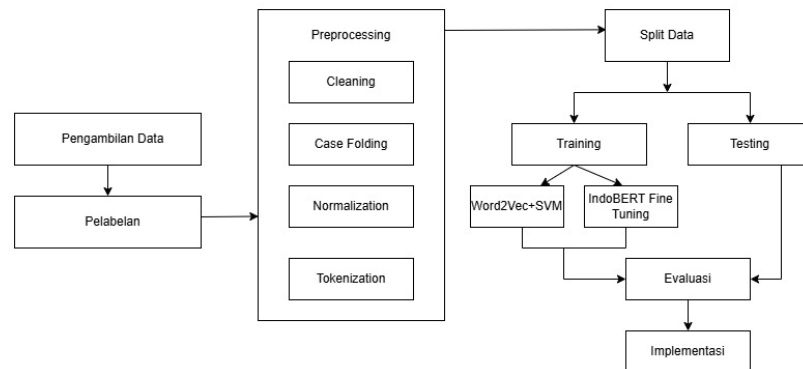
Berdasarkan penelitian terdahulu, algoritma *Support Vector Machine (SVM)* telah terbukti efektif dalam mengklasifikasikan sentimen teks yang direpresentasikan menggunakan *Word2Vec*, sedangkan *IndoBERT Fine-Tuning* menunjukkan kemampuan unggul dalam memahami konteks dan makna sentimen teks berbahasa Indonesia secara mendalam. Dengan demikian, penelitian ini dengan maksud untuk membandingkan kinerja metode *Word2Vec* yang diklasifikasikan menggunakan *SVM* dengan metode *IndoBERT Fine-Tuning* dalam menganalisis sentimen masyarakat terhadap kebijakan pembangunan Ibu Kota Nusantara (IKN) pada media sosial Twitter. Hasil analisis sentimen tersebut akan diimplementasikan dalam bentuk *dashboard* interaktif sederhana yang menampilkan distribusi sentimen positif dan negatif, sehingga mempermudah interpretasi hasil serta memberikan gambaran umum mengenai pandangan masyarakat untuk pembangunan IKN.

2. METODOLOGI PENELITIAN

2.1 Tahapan Penelitian

Penelitian ini meliputi atas beberapa tahapan. Proses yang pertama adalah pengambilan data yang dijalankan melalui proses *crawling* Twitter. Selanjutnya, data diberi label sesuai dengan kategori sentimen. Tahap *preprocessing* dilakukan guna memastikan data bersih dan layak dianalisis [12], dilakukan ekstraksi fitur untuk memperoleh representasi teks yang sesuai sebelum data dibagi menjadi data latih dan data uji. Pelatihan model dengan data latih dan pengujian dengan data uji dilakukan selanjutnya untuk mengukur kinerja model. Terakhir,

model dievaluasi berdasarkan nilai *accuracy*, *precision*, *recall*, dan *F1-score*, serta hasilnya divisualisasikan melalui *dashboard* sederhana [13]. Gambar berikut mempresentasikan tahapan dalam penelitian yang menjelaskan alur proses analisis sentimen masyarakat terhadap kebijakan pembangunan Ibu Kota Nusantara (IKN) pada media sosial Twitter, mulai dari pengambilan data hingga tahap implementasi. Penjelasan setiap tahapan adalah sebagai berikut:



Gambar 1. Tahapan Penelitian.

2.2 Pengambilan Data

Proses pengambilan data dilakukan dengan teknik *crawling* menggunakan Google Colab. Data yang diambil adalah komentar pengguna di Twitter menggunakan tools *Tweet Harvest* [14]. Pemrograman *Python* digunakan untuk mengambil tweet berbahasa Indonesia yang memuat kata kunci “Ibu Kota Nusantara, Pembangunan IKN, Dampak IKN, Anggaran IKN” pada rentang waktu 2022 hingga 2025.

2.3 Pelabelan

Pada tahap pelabelan, setiap komentar dari Twitter diklasifikasikan ke dalam dua kelas, yaitu negatif (0) dan positif (1). Penentuan label ini menjadi dasar untuk mengidentifikasi kecenderungan opini pada setiap komentar. Proses pelabelan data dilakukan secara manual dan otomatis.

2.4 Preprocessing

Setelah data diklasifikasikan berdasarkan label sentimen, tahap selanjutnya adalah *preprocessing* untuk membersihkan elemen-elemen yang tidak relevan [15]. Tahap ini bertujuan agar data menjadi lebih terstruktur, sehingga hasil analisis sentimen dapat diperoleh dengan lebih akurat. *Preprocessing* adalah langkah awal yang penting dalam menyiapkan data karena berperan dalam mengurangi *noise* serta menjaga konsistensi dan kualitas data [16]. Berikut tahapan dalam melakukan *preprocessing* data:

2.4.1 Cleaning

Merupakan tahapan untuk menghapus komponen yang tidak berguna, seperti kolom yang tidak diperlukan, tautan, *mention*, *hashtag*, tanda baca, *user ID*, serta data duplikat [17].

2.4.2 Case Folding

Merupakan tahapan untuk mengubah seluruh huruf kapital menjadi huruf kecil (*lowercase*) [18].

2.4.3 Normalization

Merupakan tahapan untuk mengubah bentuk kata dengan mengganti kata tidak baku atau salah ejaan menjadi bentuk yang sesuai kaidah bahasa agar mudah dipahami [19].

2.4.4 Tokenization

Merupakan tahapan untuk memecah teks dibagi ke dalam unit-unit kecil, yang dikenal sebagai token (kata, frasa, atau simbol), sehingga data menjadi lebih terstruktur dan mudah diolah. [20].

2.5 Split Data

Proses *split data* dilakukan dengan membagi dataset menjadi data latih (*training data*) dan data uji (*testing data*) dengan tujuan memisahkan proses pelatihan dan pengujian model agar terhindar dari *overfitting* [21]. Pembagian dilakukan dengan perbandingan 80% data latih dan 20% data uji, di mana data latih digunakan untuk membangun model, sedangkan data uji digunakan untuk mengukur kemampuan generalisasi model terhadap data

baru [22]. Tahap ini dirancang untuk menjamin bahwa model yang dibangun tidak hanya mampu mengenali pola dari data pelatihan, melainkan juga mampu menghasilkan prediksi yang akurat pada data baru yang sebelumnya tidak diperkenalkan selama pelatihan, sehingga hasil klasifikasi menjadi lebih valid.

2.6 Pemodelan Word2Vec-SVM

Pada pemodelan pertama, digunakan metode *Word2Vec* untuk mengubah teks menjadi representasi vektor numerik, yang merepresentasikan hubungan semantik antara kata pada teks [23]. Hasil representasi *Word2Vec* kemudian digunakan sebagai input bagi algoritma *Support Vector Machine (SVM)* dalam melakukan klasifikasi sentimen.

2.7 Pemodelan IndoBERT Fine-Tuning

Pemodelan kedua menggunakan *IndoBERT*, model *pre-trained language* model berbasis BERT yang dikembangkan khusus untuk Bahasa Indonesia dan dilatih menggunakan korpus Indo4B [24]. Model ini kemudian dilakukan *fine-tuning* menggunakan dataset tweet yang telah diberi label, sehingga model dapat menyesuaikan parameter internalnya agar mampu memprediksi sentimen positif dan negatif secara lebih akurat.

2.8 Evaluasi.

Setelah proses klasifikasi selesai, performa model dievaluasi dengan menggunakan *confusion matrix*. Matriks ini membantu mengidentifikasi empat jenis keluaran prediksi, yaitu *True Positive (TP)*, *True Negative (TN)*, *False Positive (FP)*, serta *False Negative (FN)*. Dari keempat nilai tersebut kemudian dihitung metrik *accuracy*, *precision*, *recall*, dan *F1-Score* untuk memperoleh gambaran menyeluruh mengenai kemampuan model [25]. Bentuk *confusion matrix* dapat dilihat pada Tabel 1.

Tabel 1. *Confusion Matrix*

	Actual Positive	Actual Negative
Predicted Positive	True Positive (TP)	False Positive (FP)
Predicted Negative	False Negative (FN)	True Negative (TN)

Berdasarkan tabel diatas dapat dituliskan rumus sebagai berikut :

- a) *Accuracy*, menunjukkan seberapa besar bagian prediksi model yang tepat, baik pada kelas positif maupun negatif, dibandingkan dengan seluruh prediksi yang dihasilkan.

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

- b) *Precision*, mengukur jumlah sampel yang berhasil dideteksi dengan benar, baik positif maupun negatif.

$$Precision = \frac{TP}{TP + FP} \quad (2)$$

- c) *Recall*, menilai ketepatan prediksi yang berhasil diidentifikasi dengan benar, baik untuk kelas positif maupun negatif.

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

- d) *F1-Score*, menilai kinerja keseluruhan dari model klasifikasi.

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

2.9 Implementasi.

Hasil akhir dipaparkan dalam bentuk aplikasi web dikembangkan menggunakan *Flask* untuk menjalankan model klasifikasi sentimen secara langsung. *Flask* digunakan sebagai kerangka kerja ringan yang memfasilitasi pengelolaan permintaan dan tampilan antarmuka, memungkinkan pengguna memperoleh hasil klasifikasi dengan cepat dan interaktif.

3. HASIL DAN PEMBAHASAN

3.1 Pengambilan Data

Data dikumpulkan melalui teknik *crawling* pada platform X menggunakan *Python* yang dijalankan di Google Colab, dengan tool *TweetHarvest*. Proses pengambilan data dilakukan menggunakan kata kunci terkait Ibu Kota

Nusantara, mencakup periode 1 Januari 2022 hingga 30 September 2025. Data mentah yang diperoleh kemudian melalui tahap *filtering*, yang meliputi pembersihan duplikasi data, penghapusan tweet yang tidak relevan dengan topik, serta penyaringan terhadap tweet yang tidak memuat informasi teks. Setelah seluruh proses penyaringan tersebut dilakukan, diperoleh 2.667 data akhir yang kemudian digunakan pada tahap pengolahan dan analisis.

3.2 Pelabelan

Tahap selanjutnya setelah pengumpulan data adalah pelabelan, yang dilakukan pada seluruh 2.667 dataset. Proses labeling dilakukan dengan dua metode, yaitu manual dan otomatis, di mana setiap komentar diklasifikasikan menjadi dua kategori, yaitu positif (1) atau negatif (0). Pelabelan manual dilakukan dengan membaca setiap komentar untuk menentukan arah opini berdasarkan konteks kalimat. Sementara itu, pelabelan otomatis dilakukan menggunakan model *IndoBERT* melalui *HuggingFace Transformers*. Kedua metode pelabelan, manual dan otomatis, diterapkan pada dua model (*SVM-Word2Vec*) dan *IndoBERT Fine-Tuning*), sehingga menghasilkan empat kombinasi pemodelan yang dibandingkan untuk menentukan performa terbaik. Hasil dari proses pelabelan data keseluruhan ditampilkan pada Tabel 2.

Tabel 2. Hasil pelabelan data

No.	Keterangan	Jumlah Data	Total Data
1.	Positif (Manual)	1.494	2.667
2.	Negatif (Manual)	1.173	
3.	Positif (Otomatis)	796	2.667
4.	Negatif (Otomatis)	1.871	

3.3 Preprocessing

Pada tahap ini bertujuan agar data lebih terstruktur dan bersih sebelum melakukan tahap proses selanjutnya. Tahap ini dilakukan mulai dari *cleaning*, *case folding*, *normalization*, dan *tokenization*.

a) Cleaning

Pada tahap ini proses pembersihan data yang bertujuan menghapus elemen-elemen tidak sesuai, termasuk *tautan*, *mention*, *hashtag*, tanda baca, serta informasi lain yang tidak memengaruhi hasil analisis. Contoh hasil *cleaning* ditampilkan pada Tabel 3.

Tabel 3. Hasil *Cleaning*

Sebelum	Sesudah
@nalar_logis @UmarSyadatHsb__ Yakan bener klo anggaran sebesar itu dibuat memajukan kota2 yg ada di Indonesia bukan cuma di Kalimantan doang alasan pembangunan IKN kan katanya utk pemerataan klo bangun kota cuma di Kalimantan dimana pemerataan nya.	Yakan bener klo anggaran sebesar itu dibuat memajukan kota2 yg ada di Indonesia bukan cuma di Kalimantan doang alasan pembangunan IKN kan katanya utk pemerataan klo bangun kota cuma di Kalimantan dimana pemerataan nya.

b) Case Folding

Tahap ini dilakukan untuk mengubah seluruh kalimat dalam teks menjadi huruf kecil (*lowercase*) untuk menyamakan bentuk penulisan. Contoh hasil *case folding* ditampilkan pada Tabel 4.

Tabel 4. Hasil *Case Folding*

Sebelum	Sesudah
Yakan bener klo anggaran sebesar itu dibuat memajukan kota2 yg ada di Indonesia bukan cuma di Kalimantan doang alasan pembangunan IKN kan katanya utk pemerataan klo bangun kota cuma di Kalimantan dimana pemerataan nya.	yakan bener klo anggaran sebesar itu dibuat memajukan kota2 yg ada di indonesia bukan cuma di kalimantan doang alasan pembangunan iKN kan katanya utk pemerataan klo bangun kota cuma di kalimantan dimana pemerataan nya.

c) Normalization

Pada tahap ini bertujuan untuk proses penyelarasan kata dengan mengubah bentuk tidak baku menjadi baku yang sesuai kaidah bahasa, seperti contoh “utk” menjadi “untuk”, “bener” menjadi “benar”, dan “yg” menjadi “yang”. Contoh hasil *normalization* dilihat pada Tabel 5.

Tabel 5. Hasil *Normalization*

Sebelum	Sesudah
yakan bener klo anggaran sebesar itu dibuat memajukan kota2 yg ada di indonesia bukan cuma di kalimantan doang alasan pembangunan ikn kan katanya utk pemerataan klo bangun kota cuma di kalimantan dimana pemerataan nya.	ya kan benar kalau anggaran sebesar itu dibuat memajukan kota-kota yang ada di indonesia, bukan hanya di kalimantan saja. alasan pembangunan ikn kan katanya untuk pemerataan. kalau membangun kota hanya di kalimantan, di mana pemerataannya.

d) Tokenization

Tahap ini digunakan untuk memecah kalimat menjadi unit-unit kecil berupa kata, frasa, atau token agar lebih mudah dianalisis pada pengolahan data. Contoh hasil ditampilkan pada Tabel 6.

Tabel 6. Hasil *Tokenization*

Sebelum	Sesudah
ya kan benar kalau anggaran sebesar itu dibuat memajukan kota-kota yang ada di indonesia, bukan hanya di kalimantan saja. alasan pembangunan ikn kan katanya untuk pemerataan. kalau membangun kota hanya di kalimantan, di mana pemerataannya.	['ya', 'kan', 'benar', 'kalau', 'anggaran', 'sebesar', 'itu', 'dibuat', 'memajukan', 'kota', 'kota', 'yang', 'ada', 'di', 'indonesia', 'bukan', 'hanya', 'di', 'kalimantan', 'saja', 'alasan', 'pembangunan', 'ikn', 'kan', 'katanya', 'untuk', 'pemerataan', 'kalau', 'membangun', 'kota', 'hanya', 'di', 'kalimantan', 'di', 'mana', 'pemerataannya']

3.4 Split data Train-Test

Sebelum model dilatih, dataset terlebih dahulu dipisahkan menjadi dua data menggunakan metode *train_test_split*, yaitu 80% untuk pelatihan dan 20% untuk pengujian. Bagian pelatihan digunakan untuk membangun model, sedangkan bagian pengujian dimanfaatkan untuk mengevaluasi kemampuan model saat menghadapi data yang belum pernah dilihat sebelumnya. Proses pemisahan dilakukan secara *stratified* agar proporsi kelas positif dan negatif tetap seimbang pada kedua data tersebut. Detail hasil pemisahan data ditampilkan pada tabel 7.

Tabel 7. Pembagian Data *Train-Test*

Dataset	Train (80%)	Test (20%)
Manual	2.133	534
Otomatis	2.133	534

3.5 Pemodelan Word2Vec-SVM

3.5.1 Word2Vec-SVM dengan manual labeling.

Pada model pertama, data hasil pelabelan manual diekstraksi menggunakan *Word2Vec* yang mengubah teks menjadi bentuk vektor numerik. *Word2Vec* bekerja berdasarkan kedekatan antar-kata dalam kalimat sehingga menghasilkan representasi makna kata dalam bentuk vektor. Setelah embedding diperoleh, klasifikasi dilakukan menggunakan algoritma *Support Vector Machine (SVM)*, yang dipilih karena kemampuannya dalam mengelompokkan data menjadi dua kategori secara akurat. Karena jumlah data pada tiap label tidak seimbang, diterapkan teknik *SMOTE* agar proporsi data menjadi lebih seimbang sebelum proses pelatihan dilakukan. Hasil pembagian data setelah penerapan *SMOTE* pada dataset dengan label manual ditampilkan pada Tabel 8.

Tabel 8. Hasil data setelah dilakukan *SMOTE*.

Label	Sebelum <i>SMOTE</i>	Sesudah <i>SMOTE</i>
1	1.195	1.195
0	938	1.195

3.5.2 *Word2Vec-SVM* dengan *automatic* labeling.

Pada model kedua, pendekatan pemodelannya tetap sama, namun dataset yang digunakan merupakan dataset dengan label otomatis. Proses ekstraksi fitur dilakukan menggunakan *Word2Vec*, kemudian hasilnya diklasifikasikan dengan *SVM* serta diseimbangkan menggunakan *SMOTE*. Secara umum, performa model ini lebih rendah dibandingkan model pertama. Penurunan tersebut dipengaruhi oleh kualitas label yang digunakan, karena *Word2Vec* cukup sensitif terhadap ketepatan data latih. Dengan demikian, hasil pada model ini lebih dipengaruhi oleh kualitas data dibandingkan oleh proses pemodelannya. Hasil pembagian data setelah penerapan *SMOTE* pada dataset dengan label otomatis ditampilkan pada Tabel 9.

Tabel 9. Hasil data setelah dilakukan *SMOTE*.

Label	Sebelum <i>SMOTE</i>	Sesudah <i>SMOTE</i>
0	1.496	1.496
1	637	1.496

3.6 Pemodelan *IndoBERT Fine-Tuning*

3.6.1 *IndoBERT Fine-Tuning* (Manual).

Pada Model ketiga menggunakan pendekatan *fine-tuning* pada *IndoBERT* dengan dataset berlabel manual. Berbeda dari model sebelumnya, *IndoBERT* tidak memerlukan proses ekstraksi fitur terpisah karena representasi teks dipelajari langsung oleh arsitektur *Transformer*. Dataset yang digunakan memiliki ketidakseimbangan kelas, sehingga dilakukan penyeimbangan menggunakan *random oversampling*. Teknik ini dipilih karena *IndoBERT* memiliki representasi konteks yang kompleks sehingga pembuatan sampel sintetis berbasis interpolasi, seperti pada *SMOTE*, berpotensi mengubah makna teks. Dengan *oversampling*, karakteristik asli teks tetap terjaga namun distribusi kelas menjadi lebih seimbang. Pembagian data setelah *oversampling* dapat dilihat pada Tabel 10.

Tabel 10. Hasil data setelah diterapkan *Oversampling*.

Label	Sebelum <i>Oversampling</i>	Sesudah <i>Oversampling</i>
1	1.195	1.195
0	938	1.195

3.6.2 *IndoBERT Fine-Tuning* (Otomatis).

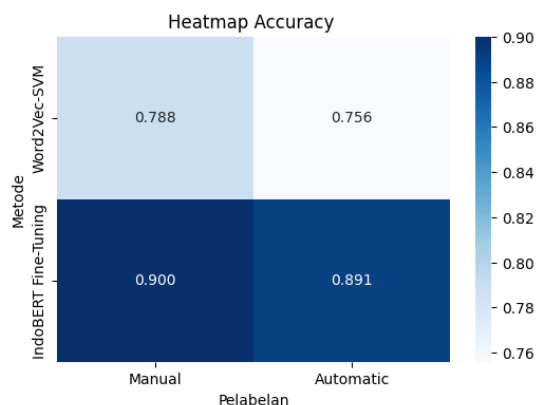
Pada Model ketiga menggunakan pendekatan *fine-tuning* pada *IndoBERT* dengan dataset berlabel manual. Berbeda dari model sebelumnya, *IndoBERT* tidak memerlukan proses ekstraksi fitur terpisah karena representasi teks dipelajari langsung oleh arsitektur *Transformer*. Dataset yang digunakan memiliki ketidakseimbangan kelas, sehingga dilakukan penyeimbangan menggunakan *random oversampling*. Teknik ini dipilih karena *IndoBERT* memiliki representasi konteks yang kompleks sehingga pembuatan sampel sintetis berbasis interpolasi, seperti pada *SMOTE*, berpotensi mengubah makna teks. Dengan *oversampling*, karakteristik asli teks tetap terjaga namun distribusi kelas menjadi lebih seimbang. Pembagian data setelah *oversampling* dapat dilihat pada Tabel 11.

Tabel 11. Hasil data setelah dilakukan *Oversampling*

Label	Sebelum <i>Oversampling</i>	Sesudah <i>Oversampling</i>
0	1.496	1.496
1	637	1.496

3.7 Evaluasi

Kinerja model kemudian diukur menggunakan empat indikator utama *accuracy*, *precision*, *recall*, dan *F1-score* yang berfungsi untuk melihat seberapa baik model dapat mengenali serta membedakan kedua kategori sentimen secara tepat dan konsisten pada kedua kelas sentimen.



Gambar 2. Heatmap *accuracy* kedua model.

Berdasarkan *heatmap accuracy*, *IndoBERT Fine-Tuning* menjadi model dengan performa terbaik. Pada data berlabel manual, model ini mencapai *accuracy* sebesar 0.900, sedangkan pada data berlabel otomatis memperoleh *accuracy* 0.891. Nilai tersebut lebih tinggi dibandingkan *Word2Vec-SVM*, di mana pelabelan manual hanya menghasilkan *accuracy* 0.788 dan pelabelan otomatis sebesar 0.756. Hasil ini menunjukkan bahwa *IndoBERT* lebih efektif dalam memanfaatkan konteks bahasa Indonesia, sementara *Word2Vec-SVM* cenderung mengalami penurunan performa ketika ketidakonsistenan pelabelan.



Gambar 3. Bar chart Perbandingan *Precision*, *Recall*, *F1-Score*

Selanjutnya, berdasarkan *bar chart precision*, *recall*, dan *F1-score*, *IndoBERT Fine-Tuning* dengan label manual kembali menunjukkan performa paling optimal dengan nilai *precision* 0.898544, *recall* 0.901352, dan *F1-score* 0.899685. Kedekatan nilai *precision* dan *recall* menunjukkan bahwa model mampu mengklasifikasikan kedua kelas sentimen secara seimbang. Pada *IndoBERT Fine-Tuning* dengan label otomatis, nilai *precision*, *recall*, dan *F1-score* masing-masing sebesar 0.879259, 0.855648, dan 0.866115 yang menunjukkan adanya penurunan performa akibat penggunaan label otomatis.

Sementara itu, model *Word2Vec-SVM* menunjukkan performa yang lebih rendah. *Word2Vec-SVM* dengan label manual menghasilkan *precision* 0.800385 *recall* 0.800107 dan *F1-score* 0.788389 sedangkan pada label otomatis nilainya menurun menjadi *precision* 0.709573, *recall* 0.712553 dan *F1-score* 0.711005. Penurunan ini menegaskan bahwa *Word2Vec-SVM* lebih sensitif terhadap ketidakonsistenan pelabelan dibandingkan *IndoBERT*.

Secara keseluruhan, hasil evaluasi menunjukkan bahwa *IndoBERT Fine-Tuning* dengan label manual merupakan model terbaik dalam klasifikasi sentimen masyarakat terhadap pembangunan Ibu Kota Nusantara, baik dari sisi *accuracy* maupun keseimbangan nilai *precision*, *recall*, dan *F1-score*.

3.8 Perbandingan hasil dengan penelitian terdahulu

Hasil penelitian ini menunjukkan bahwa *IndoBERT Fine-Tuning* memberikan performa paling stabil untuk analisis sentimen opini publik terkait IKN. Penelitian ini tidak sepenuhnya sejalan dengan beberapa penelitian terdahulu yang justru melaporkan bahwa model klasik seperti *SVM* dapat memberikan performa lebih tinggi pada dataset tertentu.

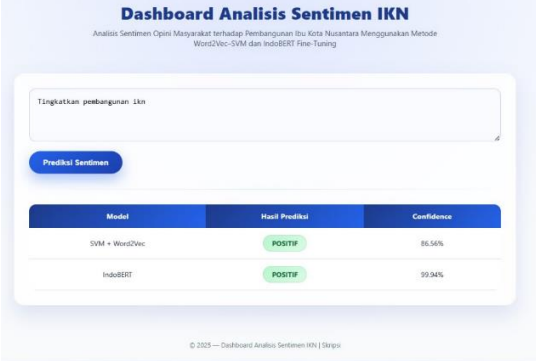
Penelitian oleh [26] menunjukkan bahwa *Word2Vec*–*SVM* mencapai akurasi sangat tinggi pada ulasan aplikasi investasi yang bahasanya lebih terstruktur. Selain itu, penelitian oleh [27] menunjukkan bahwa *SVM* unggul pada teks media sosial dengan distribusi kelas yang relatif seimbang. Selanjutnya penelitian komparatif oleh [28] menunjukkan bahwa *BERT-Base* memberikan performa yang lebih baik daripada *SVM*, pada teks yang sangat informal dan penuh variasi bahasa.

Penelitian ini sejalan dengan karakteristik teks IKN, yang kompleks, penuh ironi, dan bervariasi secara linguistik, sehingga representasi statis *Word2Vec* dan model *SVM* kurang optimal. Sebaliknya, kemampuan *IndoBERT* menangkap konteks semantik lintas-kalimat membuatnya lebih efektif dalam menganalisis opini publik yang dinamis dan tidak terstruktur.

Oleh karena itu, meskipun penelitian terdahulu menunjukkan keunggulan *SVM* pada konteks tertentu, hasil penelitian ini menegaskan bahwa transformer seperti *IndoBERT* lebih tepat digunakan untuk menganalisis opini publik pada isu sosial-politik yang kompleks, termasuk pembangunan IKN.

3.9 Implementasi

Pada penelitian ini dikembangkan aplikasi web untuk menilai persepsi publik mengenai pembangunan Ibu Kota Nusantara (IKN). Pengguna dapat memasukkan teks melalui kolom yang telah disediakan pada antarmuka aplikasi, kemudian sistem memproses masukan tersebut menggunakan dua model yang telah dilatih, yaitu *Word2Vec*–*SVM* dan *IndoBERT Fine-Tuning*. Setelah tombol “Prediksi Sentimen” ditekan, hasil dari masing-masing model disajikan dalam tabel yang berisi nama model, kategori sentimen, dan nilai *confidence*. Dari Tampilan aplikasi menunjukkan bahwa kedua model dapat menghasilkan prediksi secara bersamaan sehingga memudahkan proses evaluasi dan perbandingan performa model.



Model	Hasil Prediksi	Confidence
SVM + Word2Vec	POSITIF	85.56%
IndoBERT	POSITIF	99.94%

Gambar 4. Tampilan antarmuka aplikasi web flask.

4. KESIMPULAN

Berdasarkan hasil penelitian dan analisis yang telah dilaksanakan, dapat disimpulkan bahwa pengidentifikasian sentimen masyarakat terkait kebijakan pembangunan Ibu Kota Nusantara (IKN) di media sosial X dapat dilakukan secara efektif dengan memanfaatkan metode *Word2Vec*–*SVM* serta model *IndoBERT fine-tuning*. Dari 2.667 data tweet yang telah melalui tahapan *preprocessing*, pemodelan, dan evaluasi, model *IndoBERT Fine-Tuning* dengan label manual menunjukkan kinerja terbaik dengan nilai *accuracy* dan *F1-score* tertinggi dibandingkan model lainnya. Hal ini membuktikan bahwa model berbasis *Transformer* lebih mampu memahami konteks semantik teks berbahasa Indonesia secara mendalam dibandingkan metode berbasis *Word2Vec* yang diklasifikasikan menggunakan *SVM*. Faktor kualitas dalam proses labeling data turut menentukan kemampuan model melakukan klasifikasi, dan metode manual memberikan performa yang lebih presisi serta lebih seimbang. Berdasarkan kesimpulan tersebut, penelitian selanjutnya dianjurkan untuk meningkatkan jumlah dan variasi data dengan memasukkan lebih banyak platform media sosial agar hasil analisis sentimen menjadi lebih representatif. Selain itu, pengembangan metode pelabelan yang lebih sistematis serta penerapan klasifikasi multikelas dapat dilakukan untuk memperoleh gambaran opini publik yang lebih komprehensif. Dari sisi implementasi, pengembangan dashboard analisis sentimen dengan fitur visualisasi yang lebih interaktif dan analisis temporal juga disarankan agar hasil penelitian dapat dimanfaatkan secara optimal sebagai dasar pertimbangan dalam pembuatan keputusan kebijakan pembangunan Ibu Kota Nusantara. Meskipun penelitian ini menghasilkan kinerja model yang baik, terdapat beberapa keterbatasan yang perlu diperhatikan. Data yang digunakan dalam penelitian ini hanya bersumber dari media sosial X dengan jumlah 2.667 tweet, sehingga hasil analisis sentimen belum sepenuhnya merepresentasikan keseluruhan opini masyarakat terhadap kebijakan pembangunan Ibu Kota Nusantara. Selain itu, proses pelabelan manual memiliki potensi subjektivitas peneliti dalam menentukan kategori sentimen. Penelitian ini juga masih menggunakan pendekatan klasifikasi sentimen

biner, sehingga belum mampu menangkap variasi sentimen yang lebih kompleks. Rentang waktu pengambilan data yang terbatas turut menjadi kendala dalam menganalisis dinamika perubahan sentimen masyarakat secara jangka panjang.

REFERENCES

- [1] M. K. Saraswati and E. A. W. Adi, "Pemindahan Ibu Kota Negara Ke Provinsi Kalimantan Timur Berdasarkan Analisis SWOT," *JISIP (Jurnal Ilmu Sos. dan Pendidikan)*, vol. 6, no. 2, pp. 4042–4052, 2022, doi: 10.58258/jisip.v6i2.3086.
- [2] A. N. Sativa, A. Rizky, Imelda Putri, J. A. Putri, Akhmad Irsyad, and Islamiyah, "Analisis Sentimen Twitter Ibu Kota Negara Nusantara Menggunakan Algoritma Naive Bayes, Logistic Regression dan K-Nearest Neighbors," *Adopsi Teknol. dan Sist. Inf.*, vol. 3, no. 2, pp. 40–46, 2024, doi: 10.30872/atasi.v3i2.1371.
- [3] M. David Angelo, R. Widyatna Harwenda, I. Budi, A. Budi Santoso, and P. Kresna Putra, "Sentiment Analysis and Topic Modeling of Public Opinion on Indonesia New Capital City Development Policies," *J. Eduvest*, vol. 5, no. 5, p. 2025, 2025, [Online]. Available: <http://eduvest.greenvest.co.id>
- [4] A. Yusuf, A. Rizani, R. Fitri, K. Nursyaiful Priyo Pamungkas, and W. Arifha Saputra, "Sentimen Positif Atau Negatif: Perspektif Masyarakat Terhadap Pemindahan Ibu Kota Negara Positive or Negative Sentiment: Public Perspectives on the Relocation of the National Capital," *J. Masy. Indones.*, vol. 50, no. 2, pp. 277–300, 2024, doi: 10.55981/jmi.2024.8842.
- [5] W. Widayat, "Analisis Sentimen Movie Review menggunakan Word2Vec dan metode LSTM Deep Learning," *J. MEDIA Inform. BUDIDARMA*, vol. 5, p. 1018, 2021, doi: 10.30865/mib.v5i3.3111.
- [6] L. Chaudhary, N. Girdhar, D. Sharma, J. Andreu-Perez, A. Doucet, and M. Renz, "A Review of Deep Learning Models for Twitter Sentiment Analysis: Challenges and Opportunities," *IEEE Trans. Comput. Soc. Syst.*, vol. 11, no. 3, pp. 3550–3579, 2024, doi: 10.1109/TCSS.2023.3322002.
- [7] C. Zong, R. Xia, and J. Zhang, *Text data mining*, vol. 711. Singapore: Springer, 2021.
- [8] N. Z. B. Jannah and K. Kusnawi, "Comparison of Naïve Bayes and SVM in Sentiment Analysis of Product Reviews on Marketplaces," *Sinkron*, vol. 8, no. 2, pp. 727–733, 2024, doi: 10.33395/sinkron.v8i2.13559.
- [9] S. Khoirunnisa and E. B. Setiawan, "Sentiment Analysis on Social Media Using Long Short- Term Memory and Word2Vec Feature Expansion Methods with Adam Optimization," vol. 11, no. 1, 2025.
- [10] H. Imaduddin, F. Y. A'la, and Y. S. Nugroho, "Sentiment Analysis in Indonesian Healthcare Applications using IndoBERT Approach," *Int. J. Adv. Comput. Sci. Appl.*, vol. 14, no. 8, pp. 113–117, 2023, doi: 10.14569/IJACSA.2023.0140813.
- [11] A. Z. M. B. Hasmawati, "Sentiment Analysis of Indonesia 's Free Nutritious Meal Program on Platform X (Formerly Twitter) Using IndoBERT," vol. 9, no. 3, pp. 884–893, 2025, doi: 10.30829/zero.v9i3.27629.
- [12] A. Hermawan and I. Jowensen, "Implementasi Text-Mining untuk Analisis Sentimen pada Twitter dengan Algoritma Support Vector Machine," vol. 12, no. 1, pp. 129–137, 2023.
- [13] D. A. Pramudita, H. Imaduddin, S. N. Afiana Azizah, and I. Muslihah, "Sentiment Analysis Of Reviews On The Chatgpt Application Using Long Shortterm Memory Method," *J. Locus Penelit. dan Pengabd.*, vol. 4, no. 9, pp. 8814–8820, 2025, doi: 10.58344/locus.v4i9.4812.
- [14] M. Fahreza, A. Luthfiarta, M. Indrawan, and A. Nugraha, "Analisis sentimen: pengaruh jam kerja terhadap kesehatan mental generasi z," *J. Appl. Comput. Sci. Technol.*, vol. 5, no. 1, pp. 16–25, 2024.
- [15] A. A. Nh, E. W. Pamungkas, and S. Akhtar, "Artificial Intelligence Systems and Its Applications (AISA)," vol. 1, no. 1, pp. 1–30, 2025.
- [16] A. S. Safitri, I. Wijayanto, and S. Hadiyoso, "Improving Classification Accuracy With Preprocessing Techniques For Sentiment Analysis," in *2024 International Conference on Data Science and Its Applications (ICoDSA)*, 2024, pp. 487–490. doi: 10.1109/ICoDSA62899.2024.10651657.
- [17] H. T. Duong and T. A. Nguyen-Thi, "A review: preprocessing techniques and data augmentation for sentiment analysis," *Comput. Soc. Networks*, vol. 8, no. 1, pp. 1–16, 2021, doi: 10.1186/s40649-020-00080-x.
- [18] Tiara Danirmala and Y. S. Nugroho, "Analisis Sentimen Terhadap Topik Kenaikan Harga Bahan Bakar Minyak (BBM) pada Media Sosial Twitter," *Indones. J. Comput. Sci.*, vol. 12, no. 3, pp. 1258–1268, 2023, doi: 10.33022/ijcs.v12i3.3199.
- [19] M. I. Abidin and E. W. Pamungkas, "Analisis Sentimen Terhadap Timnas Indonesia Di Piala Asia 2023 Dengan Model Transformer Berbahasa Indonesia," *Rabit J. Teknol. dan Sist. Inf. Univrab*, vol. 10, no. 2, pp. 482–496, 2025, doi: 10.36341/rabit.v10i2.6142.
- [20] Muhammad Davit Hilal Fahri and D. Gunawan, "Analisis Sentimen Pengguna X terhadap Perempuan di Lingkungan Kerja Menggunakan Algoritma Machine Learning," *J. Technol. Informatics*, vol. 7, no. 2, pp. 134–146, 2025, doi: 10.37802/joti.v7i2.1087.
- [21] I. A. Fahrezi, Rudiman, and N. A. Verdikha, "Analisis Sentimen Twitter Atas Isu Hak Angket Menggunakan Pembobotan TF-IDF dan Algoritma SVM," vol. 3, pp. 179–192, 2024.
- [22] F. Rafiandi Andhika, W. Witanti, and P. N. Sabrina, "Analisis Sentimen Menggunakan Metode IndoBERT pada Ulasan Aplikasi Zoom Menggunakan Fitur Ekstrasi GloVe," vol. 9, p. 2025, 2025, doi: 10.47002/metik.v9i2.1098.
- [23] Verra Budhi Lestari, Ema Utami, and Hanafi, "Combining Bi-LSTM And Word2vec Embedding For Sentiment Analysis Models Of Application User Reviews," *Indones. J. Comput. Sci.*, vol. 13, no. 1, pp. 312–326, 2024, doi: 10.33022/ijcs.v13i1.3647.
- [24] P. Sayarizki and H. Nurrahmi, "Implementation of IndoBERT for Sentiment Analysis of Indonesian Presidential Candidates," *J. Comput.*, vol. 9, no. 2, pp. 61–72, 2024, doi: 10.34818/indoic.2024.9.2.934.
- [25] J. O. Leandro and M. I. Fianty, "Evaluation of Sentiment Analysis Methods for Social Media Applications: A



- Comparison of Support Vector Machines and Naïve Bayes,” *Int. J. Informatics Vis.*, vol. 9, no. 2, pp. 796–807, 2025, doi: 10.62527/joiv.9.2.2905.
- [26] F. Rifaldy, Y. Sibaroni, and S. S. Prasetyowati, “Effectiveness of word2vec and tf-idf in sentiment classification on online investment platforms using support vector machine 1.,” vol. 10, no. 2, pp. 863–874, 2025.
- [27] J. Saputra, L. Maryani, D. Wulandari, W. Eka, P. T. Informatika, and P. T. Komputer, “Analisis Performa Naive Bayes dan SVM terhadap Sentimen Teks Media Sosial dengan Word2Vec dan SMOTE,” vol. 10, no. 1, pp. 143–155, 2025.
- [28] R. L. Mulianingrum and E. Y. Hidayat, “Comparative Performance of SVM and BERT-Base Using Hybrid Preprocessing for Fast Fashion Sentiment Analysis,” vol. 9, no. 6, pp. 3464–3478, 2025.