



Department of Digital Business

Journal of Artificial Intelligence and Digital Business (RIGGS)

Homepage: <https://journal.ilmudata.co.id/index.php/RIGGS>

Vol. 5 No. 2 (2026) pp: 12745-12756

P-ISSN: 2963-9298, e-ISSN: 2963-914X

Analisis Sentimen Opini Publik terhadap Inovasi Teknologi pada Administrasi Publik Indonesia Menggunakan Metode IndoBERT

Komang Tya Vinandita Kusuma Dewi¹, Putu Hendra Suputra², I Nyoman Tri Anindia Putra³

^{1,3}Sistem Informasi, Fakultas Teknik dan Kejuruan, Universitas Pendidikan Ganesha

²Ilmu Komputer, Fakultas Teknik dan Kejuruan, Universitas Pendidikan Ganesha

¹tya@student.undiksha.ac.id, ²hendra.suputra@undiksha.ac.id, ³tri.anindia@undiksha.ac.id

Abstrak

Transformasi digital dalam administrasi publik di Indonesia mendorong pemerintah untuk meningkatkan kualitas pelayanan melalui penerapan berbagai inovasi teknologi. Namun, implementasi layanan digital pemerintah masih menghasilkan persepsi yang beragam di masyarakat yang dapat dilihat melalui opini publik pada media sosial. Penelitian ini bertujuan untuk menganalisis sentimen opini publik terhadap inovasi teknologi pada administrasi publik Indonesia menggunakan metode IndoBERT, mengidentifikasi topik dominan yang dibahas masyarakat menggunakan Latent Dirichlet Allocation (LDA), serta menyajikan hasil analisis dalam bentuk dashboard interaktif. Penelitian ini menggunakan pendekatan kuantitatif dengan kerangka kerja Knowledge Discovery in Database (KDD) yang terdiri dari tahapan data selection, preprocessing, transformation, data mining, dan evaluation. Data diperoleh dari media sosial X dan TikTok sebanyak 11.787 data opini yang kemudian melalui proses seleksi dan preprocessing sehingga menghasilkan 5.541 data yang digunakan untuk analisis. Klasifikasi sentimen dilakukan menggunakan model IndoBERT melalui proses fine-tuning dengan beberapa kombinasi hyperparameter. Evaluasi model dilakukan menggunakan confusion matrix, akurasi, presisi, recall, F1-Score, dan 5-Fold Cross Validation. Hasil penelitian menunjukkan bahwa konfigurasi terbaik diperoleh pada epoch 4 dengan learning rate $2e-5$ yang menghasilkan akurasi sebesar 75% dan F1-Score sebesar 74%. Selain itu, hasil topic modeling menghasilkan lima topik utama pada masing-masing kelas sentimen yang menggambarkan isu dominan terkait penerapan inovasi teknologi administrasi publik. Seluruh hasil analisis divisualisasikan dalam dashboard berbasis website untuk membantu pengguna memahami pola persepsi masyarakat dan mendukung evaluasi layanan digital pemerintah.

Kata kunci: Analisis sentimen, Opini Publik, IndoBERT, Topic Modeling, K-fold, Inovasi Teknologi, Dashboard

1. Latar Belakang

Transformasi digital telah menjadi salah satu faktor utama dalam pengembangan administrasi publik untuk meningkatkan efisiensi, transparansi, dan kualitas pelayanan kepada masyarakat. Di Indonesia, upaya tersebut diwujudkan melalui implementasi Sistem Pemerintahan Berbasis Elektronik (SPBE) yang diatur dalam Peraturan Presiden Nomor 95 Tahun 2018. Berdasarkan laporan evaluasi SPBE tahun 2024, indeks SPBE nasional mencapai 3,12 dengan predikat baik, meningkat dibandingkan tahun 2023 yang memperoleh nilai 2,79. Peningkatan ini menunjukkan bahwa transformasi digital pemerintah di Indonesia terus mengalami perkembangan yang signifikan [1], [2]. Implementasi SPBE mendorong digitalisasi berbagai layanan publik, seperti pelaporan pajak melalui aplikasi M-Pajak, serta perizinan usaha melalui OSS. Selain itu, peningkatan skor *Online Service Index* (OSI) yang mencapai 0,8035 pada UN E-Gov Survey 2024 menunjukkan semakin luasnya pemanfaatan layanan digital pemerintah oleh masyarakat [3]. Namun, pengalaman pengguna terhadap layanan digital pemerintah tidak selalu sama. Berbagai kritik, keluhan dan aspirasi masyarakat banyak disampaikan melalui media sosial yang saat ini menjadi salah satu sarana utama dalam pembentukan opini publik [4], [5].

Opini yang disampaikan masyarakat melalui media sosial dapat dimanfaatkan sebagai sumber informasi untuk mengevaluasi kualitas layanan digital pemerintah. Analisis sentimen merupakan salah satu pendekatan yang banyak digunakan untuk mengidentifikasi dan mengklasifikasikan opini masyarakat ke dalam kategori sentimen tertentu berdasarkan data teks [6]. Dalam konteks administrasi publik, analisis sentimen dapat dimanfaatkan untuk mengevaluasi respons masyarakat terhadap berbagai inovasi teknologi yang diterapkan pemerintah. Selain mengidentifikasi kecenderungan sentimen positif dan negatif, penelitian ini juga menerapkan topic modeling untuk

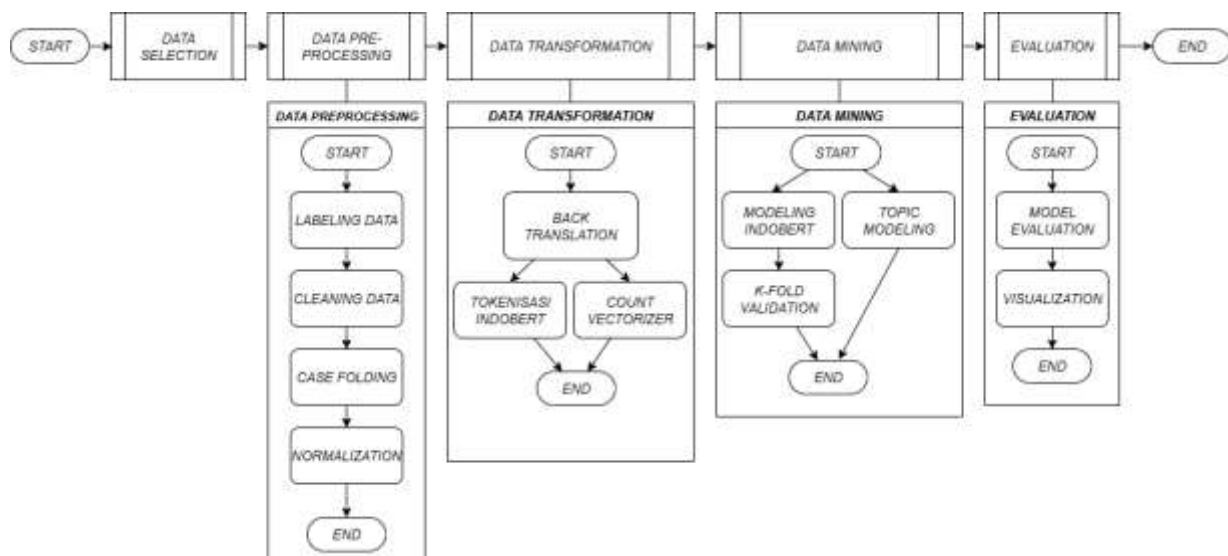
mengidentifikasi isu-isu utama yang menjadi perhatian masyarakat hingga menghasilkan pemahaman yang lebih komprehensif mengenai persepsi publik.

Berbagai penelitian sebelumnya yang telah menerapkan metode machine learning untuk analisis sentimen menunjukkan hasil yang cukup baik. Namun, metode tersebut masih memiliki keterbatasan dalam memahami konteks bahasa secara menyeluruh [7]. Seiring perkembangan teknologi pemrosesan bahasa alami, model berbasis transformer seperti Bidirectional Encoder Representations from Transformers (BERT) semakin banyak digunakan karena memiliki kemampuan memahami konteks kalimat yang lebih baik. Beberapa penelitian menunjukkan bahwa IndoBERT mampu menghasilkan performa yang lebih unggul dibandingkan dengan metode tradisional seperti Naive Bayes, SVM, dan Random Forest, maupun beberapa varian model BERT lainnya dalam tugas analisis sentimen berbahasa Indonesia [8], [9], [10].

Meskipun demikian, penelitian yang mengkaji persepsi publik terhadap inovasi teknologi pada administrasi publik Indonesia secara komprehensif masih relatif terbatas. Sebagian besar penelitian sebelumnya berfokus pada layanan atau aplikasi tertentu dan hanya melakukan klasifikasi sentimen tanpa mengidentifikasi isu dominan yang dibahas oleh masyarakat maupun menyajikan hasil analisis dalam bentuk visualisasi yang mudah dipahami. Oleh karena itu, penelitian ini bertujuan untuk menganalisis sentimen opini publik terhadap inovasi teknologi pada administrasi publik menggunakan metode IndoBERT, mengidentifikasi topik-topik dominan melalui topic modeling dan menyajikan hasil analisis dalam bentuk dashboard interaktif. Hasil penelitian diharapkan dapat memberikan gambaran yang lebih jelas mengenai persepsi masyarakat serta menjadi masukan bagi pemerintah dalam mengevaluasi dan meningkatkan kualitas layanan publik digital.

2. Metode Penelitian

Penelitian ini menggunakan pendekatan kuantitatif dengan kerangka kerja *Knowledge Discovery in Database* (KDD) yang merupakan rangkaian tahapan yang dimulai dari pemilihan data mentah, pembersihan, transformasi, penerapan algoritma, dan evaluasi hasil yang bertujuan untuk mendapatkan pengetahuan yang digunakan dalam pengambilan keputusan yang terdiri dari tahapan *data selection*, *data preprocessing*, *data transformation*, *data mining*, dan *evaluation* [11], [12]. Alur penelitian dapat dilihat pada Gambar 1.



Gambar 1. Alur Penelitian

2.1 Data Selection

Data penelitian diperoleh dari platform media sosial X dan TikTok yang dipilih karena banyak digunakan masyarakat untuk menyampaikan opini terkait layanan dan inovasi teknologi pemerintah. Untuk platform X, pengumpulan data dilakukan dengan teknik *crawling data*, sedangkan untuk platform TikTok pengumpulan data dilakukan dengan menggunakan teknik *scraping data* dengan masing-masing kata kunci yang berkaitan dengan inovasi teknologi administrasi publik yaitu SPBE, Inovasi teknologi pemerintah, Digitalisasi Administrasi, E-

Government, Transformasi Digital, dan Layanan Digital. Data yang diperoleh disimpan dalam format CSV yang memuat informasi berupa tanggal unggahan dan isi komentar.

2.2 Data Preprocessing

Tahapan preprocessing bertujuan untuk menghasilkan data yang siap untuk digunakan dalam proses analisis. Tahapan ini meliputi :

- a. Pelabelan data
Pelabelan data dilakukan secara manual yang dibantu oleh ahli bahasa Indonesia, kedalam dua kelas sentimen yaitu positif dan negatif. Data dikategorikan sebagai sentimen negatif apabila mengandung kritikan, kontra, dan permasalahan yang dialami oleh pengguna. Sedangkan data dikategorikan sebagai sentimen positif apabila mengandung dukungan dan apresiasi terhadap layanan yang diberikan.
- b. Cleaning data
Tahapan cleaning bertujuan untuk membersihkan data dari unsur yang tidak diperlukan seperti tanda baca, emoji, simbol, angka, spasi yang berlebih, URL, serta variable yang tidak diperlukan pada data.
- c. Case Folding
Case folding merupakan tahapan mengubah semua huruf dalam teks menjadi lowercase atau huruf kecil agar kata-kata yang memiliki bentuk huruf berbeda namun memiliki makna yang sama dapat dianggap sebagai satu entitas.
- d. Normalisasi
Tahapan ini bertujuan untuk menyamakan penulisan kata yang memiliki arti yang sama tetapi bentuknya berbeda seperti singkatan atau kata baku. Tahapan ini membantu model memahami setiap kata dalam kalimat opini.

2.3 Data Transformation

Data transformation bertujuan untuk mengubah data bersih menjadi bentuk yang dapat diperoleh oleh model. Data yang telah melalui proses preprocessing dibagi menjadi data latih dan data uji dengan rasio 80:20. Setelah dibagi data akan digunakan pada beberapa tahapan untuk siap melalui proses pemodelan, diantaranya yaitu :

- a. *Back Translation*
Back Translation merupakan teknik augmentasi data yang bekerja dengan menerjemahkan data teks yang telah diterjemahkan ke dalam bahasa sasaran kemudian diterjemahkan kembali ke bahasa aslinya. Proses ini digunakan untuk mengurangi ketidakseimbangan data antar kelas sentimen dengan memperbanyak varian data latih tanpa mengubah makna utama teks. Bahasa sasaran yang digunakan adalah Bahasa Inggris dan Bahasa Jerman. Data yang berbahasa Indonesia diterjemahkan ke Bahasa Inggris lalu ke Bahasa Jerman dan kembali ke Bahasa Indonesia dengan menggunakan model *Neural Machine Translation* yaitu *MarianMT*.
- b. Tokenisasi
Tokenisasi merupakan proses memecah teks menjadi unit-unit yang lebih kecil dengan memanfaatkan tokenizer dari model *IndoBERT* menggunakan *WordPiece Tokenizer*.
- c. *Count Vectorizer*
Count vectorizer merupakan tahapan dari proses *topic modeling*. Tahapan ini tidak menggunakan data latih, namun seluruh data untuk direpresentasikan kedalam bentuk vektor numerik berdasarkan frekuensi kemunculan kata dalam dokumen [13].

2.4 Data Mining

Data mining dilakukan untuk memperoleh dua jenis pengetahuan dari data yaitu analisis sentimen dan *topic modeling*. Analisis sentimen dilakukan dengan menggunakan model *IndoBERT* dengan model pre-trained *IndoBenchmark/IndoBERT-base-p1*. *IndoBERT* merupakan metode *deep learning* yang menggunakan teknologi model BERT dan dilatih menggunakan data dari berbagai sumber untuk memastikan kemampuannya dalam menangkap makna dan konteks teks yang lebih mendalam [10]. dan dilakukan fine tuning dengan menggunakan dataset penelitian untuk mengklasifikasikan sentimen kedalam kelas positif dan negatif. Sedangkan untuk *topic modeling* menerapkan metode *Latent Dirichlet Allocation* (LDA) untuk menemukan topik-topik utama

berdasarkan distribusi probabilitas kata dalam dokumen dengan menggunakan representasi teks dari *count vectorizer*.

2.5 Evaluation

Evaluation merupakan tahapan mengevaluasi performa model dengan menggunakan confusion matrix yang merupakan metode evaluasi untuk mengetahui performa model dengan membandingkan label aktual dan label hasil prediksi [14]. Confusion matrix terdiri dari empat komponen yaitu *True Positif* (TP), *True Negatif* (TN), *False Positive* (FP), dan *False Negatif* (FN) seperti ditunjukkan pada Tabel 1.

Tabel 1. Contoh Confusion Matrix

	Predicted Positive	Predicted Negatif
Actual Positif	TP	FN
Actual Negatif	FP	TN

Berdasarkan Tabel 1, TP merupakan jumlah data dengan kelas positif yang berhasil diprediksi dengan benar, sedangkan TN merupakan jumlah data dengan kelas negatif yang berhasil diprediksi dengan benar. Sebaliknya, FP merupakan data negatif yang salah diprediksi sebagai positif dan FN merupakan data positif yang salah diprediksi sebagai negatif. Nilai-nilai yang diperoleh dari matrix tersebut kemudian digunakan menghitung performa model seperti berikut.

a. Akurasi

Merupakan proporsi prediksi yang benar terhadap total data yang diprediksi [14]. Akurasi dihitung dengan rumus 1.

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

b. Presisi

Merupakan nilai ketepatan sistem mengenai informasi data positif dan negatif [15], untuk menghitung presisi dapat digunakan rumus 2.

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

c. Recall

Merupakan proporsi data positif yang berhasil diprediksi dengan benar oleh model [14]. *Recall* dihasilkan dari jumlah TP dibandingkan dengan nilai aktual positif dengan rumus 3.

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

d. F1-Score

Merupakan perbandingan rata-rata dari presisi dan *recall* [14]. Nilai ini memberikan gambaran performa model secara keseluruhan dengan mempertimbangkan *trade-off* antara presisi dan *recall* dengan rumus 4.

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

Selain itu, penelitian ini menerapkan 5-Fold Cross Validation pada data latih untuk mengukur konsistensi performa model. *5-fold Validation* merupakan metode evaluasi model yang umum digunakan untuk mengevaluasi kinerja suatu model dengan membagi dataset menjadi k subset (fold). Satu subset digunakan sebagai data uji dan k-1 subset lainnya sebagai data latih. Proses ini diulang sebanyak 5 kali sehingga setiap subset digunakan sebagai data uji sekali. Setiap data mendapat kesempatan untuk berperan sebagai data uji yang meningkatkan akurasi dan generalisasi model [16]. Hasil analisis sentimen dan *topic modeling* kemudian divisualisasikan dalam bentuk dashboard berbasis website yang menampilkan distribusi sentimen, distribusi topik, performa model, serta wordcloud untuk mendukung interpretasi hasil penelitian.

3. Hasil dan Diskusi

3.1. Data Selection

Tahapan awal penelitian dilakukan dengan mengumpulkan data opini publik melalui platform media sosial yaitu X dan Tiktok terkait inovasi teknologi pada administrasi publik di Indonesia. Proses *crawling* dilakukan dengan

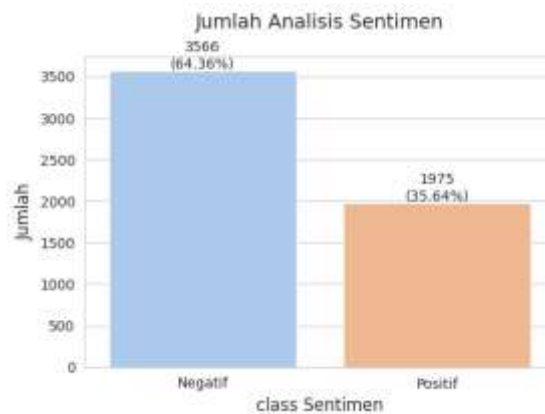
DOI: <https://doi.org/10.31004/riggs.v5i2.10149>

Lisensi: Creative Commons Attribution 4.0 International (CC BY 4.0)

menggunakan tools *tweet-harvest* dan *scraping data* menggunakan layanan *Apify Client* yang menghasilkan sebanyak 11.787 data opini. Data tersebut kemudian diseleksi berdasarkan relevansi topik penelitian sehingga diperoleh 5.541 data yang digunakan untuk tahap analisis selanjutnya.

3.2. Data preprocessing

Tahapan preprocessing dilakukan untuk meningkatkan kualitas data sebelum digunakan dalam proses pemodelan. Tahap pelabelan dilakukan secara manual yang dibantu oleh tiga orang ahli bahasa Indonesia yaitu Guru Bahasa Indonesia. Dari hasil pelabelan seluruh data didapatkan sebanyak 3.566 data bersentimen negatif dan 1.975 data bersentimen positif yang dapat dilihat pada Gambar 2.



Gambar 2. Hasil Distribusi Sentimen

Data yang telah dilabeling, kemudian melalui tahapan cleaning untuk memastikan data benar benar bersih dari tanda baca, emoji, symbol, serta variable yang tidak diperlukan. Setelah itu data akan melalui proses case folding untuk menyeragamkan bentuk huruf pada kalimat. Hasil dari case folding dapat dilihat pada Tabel 2.

Tabel 2. Hasil Case Folding

Data	Hasil Case Folding
Btw aku baca buku ini di aplikasi iPusnas ya Buat kalian yang pengen baca novel terkenal tapi belum ada budget buat beli bukunya bisa pinjem di sini aja Gratis dan legal Tapi ini aplikasi pemerintah ya jadi jangan kaget kalau banyak kurangnya	btw aku baca buku ini di aplikasi ipusnas ya buat kalian yang pengen baca novel terkenal tapi belum ada budget buat beli bukunya bisa pinjem di sini aja gratis dan legal tapi ini aplikasi pemerintah ya jadi jangan kaget kalau banyak kurangnya

Selanjutnya, data akan dinormalisasi untuk mengubah kata yang tidak sesuai dengan KBBI seperti singkatan, slang, dan kata tidak baku. Proses ini memanfaatkan *resource* dari GitHub riochr17 pada tautan <https://github.com/riochr17/Analisis-Sentimen-ID> seperti yang ditunjukkan pada Tabel 4.

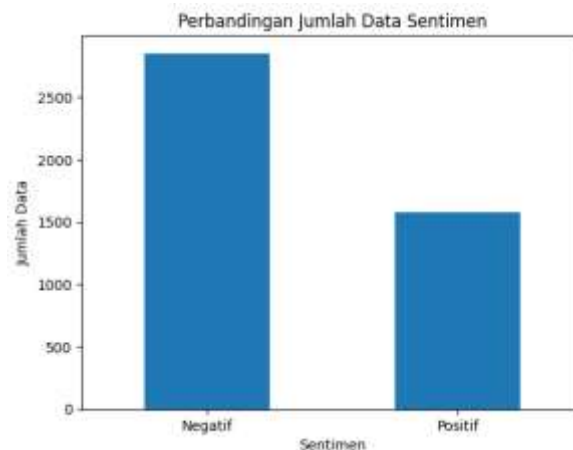
Tabel 3. Hasil Normalisasi

Data	Hasil Normalisasi
btw aku baca buku ini di aplikasi ipusnas ya buat kalian yang pengen baca novel terkenal tapi belum ada budget buat beli bukunya bisa pinjem di sini aja gratis dan legal tapi ini aplikasi pemerintah ya jadi jangan kaget kalau banyak kurangnya	ngomong ngomong aku baca buku ini di aplikasi ipusnas iya buat kalian yang ingin baca novel terkenal tapi belum ada budget buat beli bukunya bisa pinjem di sini saja gratis dan legal tapi ini aplikasi pemerintah iya jadi jangan kaget kalau banyak kurangnya

Hasil dari tahapan preprocessing menghasilkan dataset yang lebih terstruktur dengan menghilangkan berbagai unsur yang tidak diperlukan sehingga proses klasifikasi dapat dilakukan secara optimal.

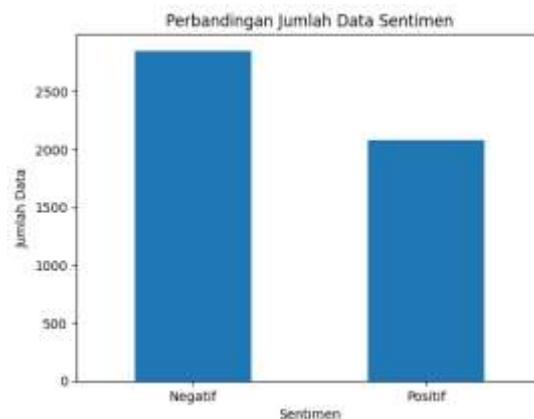
3.3 Data Transformation

Pada tahapan ini dataset dibagi menjadi data latih dan data uji dengan rasio 80:20. Data latih digunakan dalam berbagai tahapan seperti pada *back translation* (BT) dan tokenisasi. Pada data latih terdapat perbandingan jumlah data positif dan negatif yang sangat tinggi dapat dilihat pada Gambar 3.



Gambar 3. Perbandingan Jumlah Sentimen Sebelum Proses BT

Karena perbandingan ini, data latih melalui proses BT dengan menggunakan 500 data positif untuk diterjemahkan. Proses penerjemahan dilakukan melalui 4 kali penerjemahan mulai dari Bahasa Indonesia – Bahasa Inggris, Inggris – Jerman, Jerman – Inggris, dan Inggris – Indonesia dengan menggunakan model MarianaMT. Pemilihan bahasa ini dilakukan untuk menghindari terjadinya kesamaan struktur teks opini sehingga proses BT dapat menghasilkan opini yang lebih beragam namun masih dalam makna yang sama. Hasil proses ini dapat dilihat pada Gambar 4.



Gambar 4. Perbandingan Jumlah Sentimen Setelah BT

Setelah proses BT, perbandingan antara data positif dan negatif menjadi lebih kecil seperti yang ditunjukkan pada Gambar 4. Jumlah data positif meningkat yang sebelumnya berjumlah 1.580 data menjadi 2080 data. Proses tidak dilakukan dengan menggunakan 100% data positif untuk tetap mempertahankan mayoritas dari data asli. Setelah melalui proses BT, kemudian data latih ditokenisasi untuk membantu model dalam memahami konteks bahasa Indonesia dengan lebih baik. Data dalam bentuk kata diubah ke dalam bentuk numerik atau token token subkata dengan menggunakan tokenizer bawaan model *indoBERT* yang hasilnya akan digunakan sebagai input dalam pemodelan. Sementara itu, pada proses topic modeling data yang digunakan berasal dari keseluruhan data opini

yang telah melalui proses preprocessing. Data ini ditransformasikan menggunakan metode Count Vectorizer yang menghasilkan representasi numerik berdasarkan frekuensi kemunculan kata. Tahapan ini diawali dengan proses stopword removal untuk menghapus kata umum yang tidak diperlukan. Hasil dari tahapan ini dapat dilihat pada Tabel 4.

Tabel 4. Hasil Count Vectorizer

Kata	Frekuensi
Aplikasi	1685
Login	529
Susah	347

Hasil representasi tersebut kemudian digunakan sebagai input pada metode *Latent Dirichlet Allocation* (LDA) untuk menemukan topik topik utama yang terdapat dalam opini masyarakat.

3.4 Data Mining

Tahapan data mining dilakukan melalui dua proses yaitu klasifikasi sentiment dengan menggunakan IndoBERT dan Topic Modeling dengan menggunakan LDA. Pengujian model dilakukan dengan menggunakan kombinasi hyperparameter yang terdapat pada Tabel 5.

Tabel 5. Kombinasi Hyperparameter

Percobaan	Epoch	Learning rate
I	3	1e-5
II	3	2e-5
III	4	1e-5
IV	4	2e-5

Setiap konfigurasi dilatih menggunakan data latih dan menghasilkan model yang kemudian digunakan untuk melakukan prediksi terhadap data uji. Model yang digunakan adalah model yang memberikan performa terbaik dari keempat percobaan. Tahapan data mining juga dilakukan untuk menemukan topik utama menggunakan LDA dari masing-masing kelas sentimen. Hasil pemodelan menghasilkan lima topik dominan yang menggambarkan isu utama yang dibahas masyarakat untuk setiap kelas sentimen.

3.5 Evaluation

Tahapan evaluasi dilakukan untuk mengetahui kemampuan model IndoBERT dalam melakukan klasifikasi sentimen. Evaluasi dilakukan dengan menggunakan metrik akurasi, presisi, recall, dan F1-Score terhadap data uji. Adapun hasil evaluasi untuk setiap percobaan berdasarkan weighted avg yang dapat dilihat pada Tabel 6.

Tabel 6. Hasil Tiap Percobaan

Percobaan	Epoch	Learning rate	Akurasi	Presisi	Recall	F1-Score
I	3	1e-5	74%	74%	74%	73%
II	3	2e-5	72%	73%	72%	72%
III	4	1e-5	73%	73%	73%	73%
IV	4	2e-5	75%	75%	75%	74%

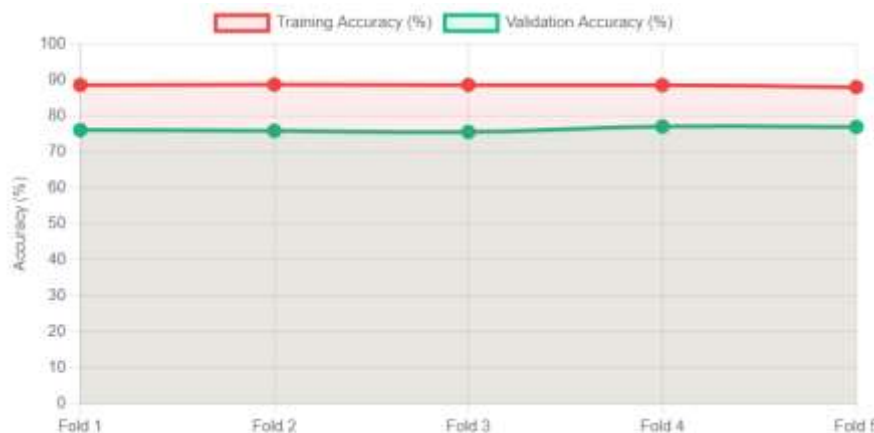
Berdasarkan dari evaluasi pada Tabel 6, konfigurasi terbaik diperoleh pada percobaan ke IV dengan penggunaan epoch 4 dan learning rate 2e-5. Model tersebut menghasilkan nilai akurasi sebesar 75% dengan nilai F1-Score 74%. Hasil ini menunjukkan bahwa model IndoBERT mampu melakukan klasifikasi sentiment dengan performa yang cukup baik pada data opini publik terkait inovasi teknologi administrasi publik. Setelah diperoleh konfigurasi model terbaik, evaluasi tambahan dilakukan untuk mengetahui kestabilan performa model dan kemampuan generalisasi terhadap data yang berbeda dengan menggunakan *5-Fold Cross Validation*, dimana dataset pelatihan dibagi menjadi lima bagian besar. Proses pelatihan dan validasi dilakukan sebanyak lima kali dengan satu bagian digunakan sebagai data validasi dan bagian lainnya sebagai data pelatihan. Pengujian 5-Fold juga dilakukan dengan

beberapa varian jumlah epoch yaitu 4, 3 dan 2 untuk mengetahui pengaruh jumlah epoch terhadap stabilitas model. Hasil pengujian 5-Fold dapat dilihat pada Tabel 7.

Tabel 7. Hasil 5-Fold Cross Validation

Epoch	Fold	Train Accuracy	Valid Accuracy	Precision	Recall	F1-Score
4	1	0.973	0.768	0.767	0.768	0.767
	2	0.971	0.752	0.751	0.752	0.750
	3	0.974	0.776	0.776	0.776	0.772
	4	0.970	0.761	0.762	0.761	0.758
	5	0.971	0.746	0.744	0.746	0.745
3	1	0.941	0.780	0.779	0.780	0.779
	2	0.940	0.759	0.758	0.759	0.758
	3	0.936	0.767	0.766	0.767	0.764
	4	0.937	0.767	0.770	0.767	0.763
	5	0.936	0.758	0.757	0.758	0.756
2	1	0.886	0.761	0.760	0.761	0.760
	2	0.887	0.759	0.758	0.759	0.758
	3	0.886	0.756	0.756	0.756	0.749
	4	0.886	0.771	0.773	0.771	0.767
	5	0.881	0.770	0.769	0.770	0.768

Berdasarkan hasil pengujian pada tiap epoch, epoch 2 menunjukkan performa yang lebih stabil dibandingkan epoch 3 dan 4. Hal ini ditunjukkan dari gap antara *train accuracy* dan *validation accuracy* yang lebih kecil yaitu 12,2% dibandingkan dengan epoch lainnya. Meskipun epoch 4 menghasilkan performa klasifikasi tertinggi, hasil 5-Fold menunjukkan bahwa epoch 2 memiliki kestabilan yang lebih baik berdasarkan gapnya. Perbandingan gap *train accuracy* dan *validation accuracy* dapat dilihat pada Gambar 5.



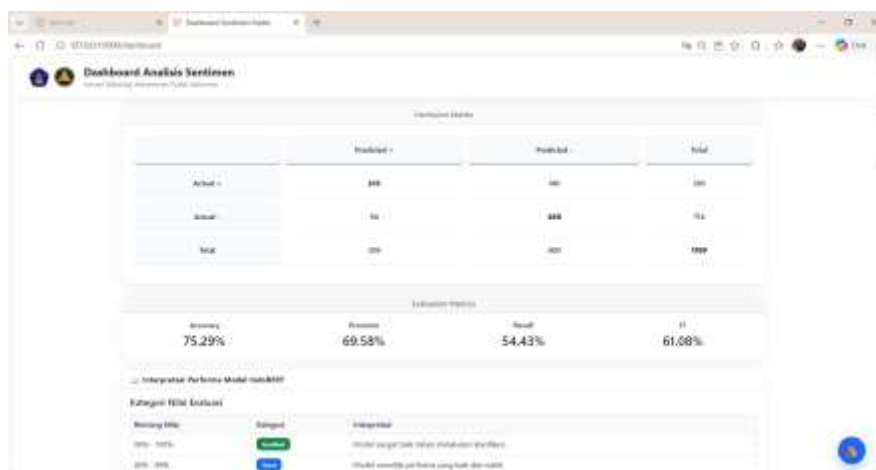
Gambar 5. Perbandingan *Train Accuracy* dan *Validation Accuracy*

Tahapan akhir dari penelitian adalah memvisualisasikan hasil sentiment dan topic modeling dalam bentuk dashboard berbasis website. Visualisasi ini bertujuan untuk menyajikan informasi hasil pengolahan data secara lebih informatif sehingga pola persepsi masyarakat terhadap inovasi teknologi dapat dipahami dengan lebih mudah. Dashboard dikembangkan dengan menampilkan beberapa informasi utama seperti distribusi sentimen masyarakat, hasil evaluasi model, performa K-fold serta topik dominan yang diperoleh dari proses topic modeling. Tampilan dashboard dapat dilihat pada Gambar 6.



Gambar 6. Tampilan Dashboard Analisis Sentimen

Dashboard menampilkan jumlah opini berdasarkan kategori sentimen positif dan negatif. Informasi ini memberikan gambaran mengenai kecenderungan persepsi masyarakat terhadap penerapan inovasi teknologi administrasi public di Indonesia. Selain itu visualisasi distribusi sentiment membantu pengguna dalam melihat perbandingan opini masyarakat secara lebih cepat. Selanjutnya, dashboard juga menyajikan hasil evaluasi model IndoBERT berupa confusion matrix dan nilai matrix evaluasi seperti pada Gambar 7.



Gambar 7. Tampilan Performa Model

Gambar 7 merupakan tampilan dashboard yang menampilkan performa dari model indoBERT dalam melakukan klasifikasi sentimen pada data. Terdapat *confusion matrix* dari model dan *evaluation matrix* sehingga pengguna dapat memahami tingkat performa model indoBERT dalam menganalisis sentimen opini publik. Untuk mempermudah pengguna dalam mengetahui apakah model bekerja dengan baik atau tidak berdasarkan dari performa yang ditampilkan, terdapat indikator informasi dibagian bawah evaluation matrix sehingga pengguna mengetahui akurasi dari model bisa dianggap baik atau tidak. Selain evaluasi model, hasil topic modeling juga divisualisasikan untuk melihat topik-topik utama yang muncul dalam opini masyarakat seperti pada Gambar 8 dan Gambar 9.

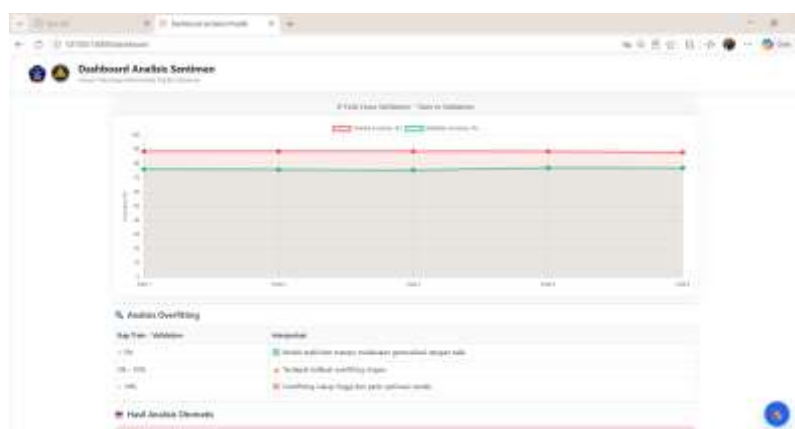


Gambar 8. Visualisasi Topic Modeling Sentimen Positif



Gambar 9. Visualisasi Topic Modeling Sentimen Negatif

Hasil visualisasi menunjukkan beberapa isu dominan yang sering dibahas masyarakat. Pada topic modeling menghasilkan 5 topik utama yang sering dibahas untuk setiap sentimen yang terdiri dari 10 kata yang paling sering muncul pada setiap dokumen. Informasi ini dapat digunakan sebagai bahan evaluasi untuk mengetahui aspek layanan digital yang masih memerlukan perbaikan. Komponen lainnya yang ditampilkan dashboard adalah visualisasi hasil K-Fold dalam bentuk diagram grafik seperti pada Gambar 10.



Gambar 10. Visualiasi K-Fold

Gambar 10 merupakan tampilan hasil *K-Fold Cross Validation*. Dashboard menunjukkan hasil akurasi dari setiap fold dan rata rata keseluruhannya. Selain itu, terdapat juga diagram perbandingan antara train akurasi dan valid

akurasi. Untuk mempermudah pengguna dalam memahami perbandingan dan arti dari gap yang terdapat pada diagram, pada dashboard ditambahkan informasi terkait rentang gap yang dianggap aman dan *overfitting*.

4. Kesimpulan

Berdasarkan hasil penelitian yang telah dilakukan mengenai analisis sentimen opini publik terhadap inovasi teknologi pada administrasi publik Indonesia menggunakan metode IndoBERT, dapat disimpulkan bahwa penerapan teknologi pemrosesan bahasa alami berbasis transformer mampu digunakan untuk mengidentifikasi persepsi masyarakat berdasarkan opini yang disampaikan melalui media sosial. Penelitian ini berhasil mengumpulkan sebanyak 11.787 data opini dari platform X dan Tiktok yang berkaitan dengan inovasi teknologi administrasi publik. Setelah melalui tahapan seleksi data dan preprocessing data yang meliputi pelabelan, cleaning, case folding, dan normalisasi diperoleh sebanyak 5.541 data yang digunakan untuk proses analisis selanjutnya. Pada tahap klasifikasi sentimen model IndoBERT dilakukan fine tuning dengan menggunakan beberapa kombinasi hyperparameter untuk memperoleh konfigurasi model terbaik. Berdasarkan hasil pengujian, konfigurasi dengan epoch 4 dan learning rate $2e-5$ memberikan performa terbaik dengan nilai akurasi sebesar 75%, persisi 75% recall 75% dan F1-score 74%. Hasil tersebut menunjukkan bahwa model IndoBERT mampu memahami pola bahasa pada opini masyarakat dan melakukan klasifikasi sentimen ke dalam kategori positif dan negatif dengan performa yang cukup baik. Selain itu evaluasi pengujian tambahan menggunakan 5-Fold dilakukan untuk mengetahui konsistensi dan kemampuan generalisasi model. Hasil pengujian menunjukkan bahwa model memiliki rata rata akurasi sebesar 76%. Berdasarkan analisis terhadap variasi epoch, epoch 2 menghasilkan performa yang lebih stabil karena memiliki gap train dan validation accuracy yang lebih kecil dibandingkan dengan epoch lainnya sehingga dapat mengurangi resiko terjadinya *overfitting* pada model. Selain klasifikasi sentimen, penelitian ini juga menerapkan topic modeling menggunakan metode *Latent Dirichlet Allocation* atau LDA untuk mengetahui topik topik dominan yang muncul pada masing masing kelas sentimen yang menggambarkan berbagai aspek yang sering dibahas oleh masyarakat baik berupa apresiasi maupun keluhan terhadap penerapan inovasi teknologi administrasi publik. Hasil ini menunjukkan bahwa analisis sentimen tidak hanya dapat memberikan informasi mengenai kecenderungan opini masyarakat tetapi juga dapat membantu memahami isu utama yang menjadi perhatian publik. Seluruh hasil analisis kemudian divisualisasikan dalam bentuk dashboard berbasis website yang menampilkan distribusi sentimen, hasil evaluasi model, confusion matrix, visualisasi topic modeling, serta performa 5-Fold. Dashboard tersebut dapat membantu pengguna dalam memahami hasil analisis secara lebih mudah dan informatif. Dengan adanya penelitian ini, diharapkan pendekatan analisis sentimen berbasis IndoBERT dan topic modeling dapat menjadi salah satu metode yang digunakan dalam mengevaluasi persepsi masyarakat terhadap inovasi teknologi pada administrasi publik serta memberikan gambaran mengenai aspek layanan digital yang perlu dipertahankan maupun ditingkatkan.

Referensi

- [1] Pemerintah Pusat, "Peraturan Presiden Nomor 95 Tahun 2018 tentang Sistem Pemerintahan Berbasis Elektronik," *Menteri Huk. Dan Hak Asasi Mns. Republik Indones.*, p. 110, 2018.
- [2] K. RI, "Laporan Pelaksanaan Evaluasi SPBE Tahun 2024," 2024.
- [3] D. of E. and S. Affairs, *E-Government Survey 2024*. 2024. [Online]. Available: <https://www.un-ilibrary.org/content/books/9789211067286>
- [4] J. Juleha et al., "Peran Media Sosial Dalam Dinamika Opini Publik dan Partisipasi Politik Era Digital," *Concept J. Soc. Humanit. Educ.*, vol. 3, no. 1, pp. 38–45, 2024, [Online]. Available: <https://doi.org/10.55606/concept.v3i1.951>
- [5] F. S. Pratiwi, "Peran Komunikasi Digital dalam Pembentukan Opini Publik : Studi Kasus Media Sosial," pp. 293–315, 2024.
- [6] M. R. Fahlevvi, "Analisis Sentimen Terhadap Ulasan Aplikasi Pejabat Pengelola Informasi Dan Dokumentasi Kementerian Dalam Negeri Republik Indonesia Di Google Playstore Menggunakan Metode Support Vector Machine," *J. Teknol. dan Komun. Pemerintah.*, vol. 4, no. 1, pp. 1–13, 2022, doi: 10.33701/jtkp.v4i1.2701.
- [7] A. Farhan, A. Y. Rahman, T. Informatika, U. W. Malang, and G. P. Store, "ANALISIS SENTIMEN ULASAN APLIKASI IDENTITAS KEPENDUDUKAN DIGITAL DI GOOGLE PLAY STORE DENGAN BERT," *JATI (Jurnal Mhs. Tek. Inform.*, vol. 9, no. 3, pp. 3776–3783, 2025.
- [8] A. Farhan, A. Y. Rahman, T. Informatika, U. W. Malang, and G. P. Store, "DIGITAL DI GOOGLE PLAY STORE DENGAN BERT," vol. 9, no. 3, pp. 3776–3783, 2025.
- [9] Taufiq Dwi Purnomo and Joko Sutopo, "Comparison of Pre-Trained Bert-Based Transformer Models for Regional Language Text Sentiment Analysis in Indonesia," *Int. J. Sci. Technol.*, vol. 3, no. 3, pp. 11–21, 2024, doi: 10.56127/ijst.v3i3.1739.
- [10] C. J. L. Tobing, IGN Lanang Wijayakusuma, and Luh Putu Ida Harini, "Perbandingan Kinerja IndoBERT dan MBERT Untuk Deteksi Berita Hoaks Politik dalam Bahasa Indonesia," *JST (Jurnal Sains dan Teknol.*, vol. 14, no. 1, pp. 114–123, 2025, doi: 10.23887/jstundiksha.v14i1.92126.
- [11] A. F. Panjalu, S. Alam, and M. I. Sulisty, "Moba Game Review Sentiment Analysis Using Support Vector Machine Algorithm," *JIKO (Jurnal Inform. dan Komputer)*, vol. 6, no. 2, pp. 131–137, 2023, doi: 10.33387/jiko.v6i2.6388.
- [12] T. Safitri, Y. Umaidah, and I. Maulana, "Analisis Sentimen Pengguna Twitter Terhadap Grup Musik BTS Menggunakan Algoritma Support Vector Machine," *J. Appl. Informatics Comput.*, vol. 7, no. 1, pp. 28–35, 2023, doi: 10.30871/jaic.v7i1.5039.
- [13] A. Averina, H. Hadi, and J. Siswanto, "Analisis Sentimen Multi-Kelas Untuk Film Berbasis Teks Ulasan Menggunakan Model

- Regresi Logistik,” *Teknika*, vol. 11, no. 2, pp. 123–128, 2022, doi: 10.34148/teknika.v11i2.461.
- [14] N. M. K. Sedana, I. N. S. Wijaya, and I. K. R. Artana, “ANALISIS SENTIMEN BERBAHASA INGGRIS DENGAN METODE LSTM STUDI KASUS BERITA ONLINE PARIWISATA BALI,” *J. Teknol. Inf. dan Ilmu Komput.*, vol. 11, no. 6, 2024, doi: 10.25126/jtiik.2024118792.
- [15] N. A. Salsabila, “ANALISIS SENTIMEN PADA MEDIA SOSIAL TWITTER TERHADAP TOKOH GUS DUR MENGGUNAKAN METODE NAÏVE BAYES DAN SUPPORT VECTOR MACHINE (SVM),” *Braz Dent J.*, vol. 33, no. 1, pp. 1–12, 2022.
- [16] R. Maheri, F. N. Salisah, and F. Muttakin, “Analisis sentimen ulasan aplikasi m-paspor menggunakan,” vol. 10, no. 1, pp. 448–458, 2025.