

ANALISIS SENTIMEN PENGGUNA MEDIA SOSIAL X TERHADAP KASUS PEJABAT KORUPSI: PERBANDINGAN METODE BI-LSTM, CNN, DAN BERT**Nadya Patricia Tanabora¹, Diah Aryani²**¹Program Studi Teknik Informatika, Universitas Esa Unggul, Jakarta, Indonesia²Program Studi Teknik Informatika, Universitas Esa Unggul, Jakarta, Indonesia**Abstract (English)**

Corruption involving public officials remains a serious problem in Indonesia and elicits diverse public reactions, particularly through social media, which is used as a means of openly expressing opinions. The large volume of data, along with the unstructured nature of text and the use of informal language, makes manual opinion analysis ineffective. This study aims to analyze public sentiment toward corruption cases on social media platform X and compare the performance of deep learning methods such as BiDirectional Long Short-Term Memory (Bi-LSTM), Convolutional Neural Network (CNN), and BiDirectional Encoder Representations from Transformers (BERT) in classifying sentiment. The research data consisted of Indonesian-language tweets collected through a crawling process using keywords related to cases of corrupt officials. The data then underwent preprocessing, sentiment labeling into three classes: positive, neutral, and negative, dataset division, and modeling and evaluation. The evaluation model used a confusion matrix and measurements of accuracy, precision, recall, and F1-score. The results of this study indicate that negative sentiment dominates public opinion regarding the official corruption case on social media platform X. Furthermore, a comparison of performance models indicates that BERT outperforms Bi-LSTM and CNN in classifying sentiment. The study concludes that transformer-based models are more effective for analyzing sentiment on complex and contextual social issues.

Article History

Submitted: 13 Juni 2026

Accepted: 16 Juli 2026

Published: 17 Juli 2026

Key Words

Sentiment Analysis, Social Media X, Official Corruption, Bi-LSTM, CNN, BERT

Abstrak (Indonesia)

Korupsi yang melibatkan pejabat publik masih menjadi permasalahan serius di Indonesia dan menimbulkan beragam reaksi dari masyarakat, khususnya melalui media sosial X yang digunakan sebagai sarana penyampaian opini secara terbuka. Besarnya volume data serta karakteristik teks yang tidak terstruktur dan menggunakan bahasa informal menyebabkan analisis opini secara manual menjadi tidak efektif. Penelitian ini bertujuan untuk menganalisis sentimen masyarakat terhadap kasus pejabat korupsi di media sosial X serta membandingkan kinerja metode deep learning Bidirectional Long Short-Term Memory (Bi-LSTM), Convolutional Neural Network (CNN), dan Bidirectional Encoder Representations from Transformers (BERT) dalam mengklasifikasikan sentimen. Data penelitian berupa tweet berbahasa Indonesia yang dikumpulkan melalui proses crawling menggunakan kata kunci terkait kasus pejabat korupsi. Data tersebut kemudian melalui tahapan prapemrosesan, pelabelan sentimen ke dalam tiga kelas yaitu positif, netral, dan negatif, pembagian dataset, serta pemodelan dan evaluasi. Evaluasi model dilakukan menggunakan confusion matrix dan metrik akurasi, presisi, recall, dan F1-score. Hasil penelitian menunjukkan bahwa sentimen negatif mendominasi opini masyarakat terhadap kasus pejabat korupsi di media sosial X. Meskipun BERT mencapai akurasi tertinggi (80,09%), model ini cenderung memprediksi seluruh data sebagai kelas mayoritas. Model CNN dan Bi-LSTM menunjukkan kemampuan klasifikasi multikelas yang lebih baik dengan akurasi 77,8% dan 74,5%. Penelitian ini menunjukkan pentingnya mempertimbangkan keseimbangan prediksi antar kelas dalam evaluasi model analisis sentimen, terutama pada dataset dengan ketidakseimbangan distribusi yang signifikan.

Sejarah Artikel

Submitted: 13 Juni 2026

Accepted: 16 Juli 2026

Published: 17 Juli 2026

Kata Kunci

Analisis Sentimen, Media Sosial X, Korupsi Pejabat, Bi-LSTM, CNN, BERT

1. PENDAHULUAN

Korupsi masih menjadi permasalahan serius dalam penyelenggaraan pemerintahan di Indonesia, khususnya yang melibatkan pejabat publik. Berdasarkan data Strategi Nasional Pencegahan Korupsi (2026), kasus korupsi yang melibatkan pejabat negara menunjukkan kecenderungan yang masih tinggi dan berdampak langsung terhadap stabilitas pemerintahan, kepercayaan publik, serta efektivitas pelayanan publik [1]. Kondisi ini menegaskan bahwa korupsi bukan hanya persoalan hukum, tetapi juga permasalahan sosial yang memengaruhi persepsi dan partisipasi masyarakat terhadap pemerintah.

Perkembangan teknologi informasi saat ini telah mengubah cara masyarakat mengakses dan menyebarkan informasi, terutama melalui media sosial. Media sosial tidak hanya digunakan sebagai sarana komunikasi, tetapi juga sebagai ruang publik digital tempat masyarakat dapat menyampaikan pendapat serta merespons berbagai isu secara cepat dan terbuka. Sosial media X telah berkembang menjadi platform diskusi digital yang aktif dan relevan dalam pembentukan sentimen masyarakat terhadap berbagai isu sosial, politik, maupun ekonomi [2]. Tingginya tingkat penggunaan platform ini di Indonesia menunjukkan bahwa media sosial X memiliki peran penting dalam membentuk opini publik [3].

Seiring meningkatnya perhatian masyarakat terhadap kasus pejabat korupsi, media sosial X menjadi ruang utama bagi masyarakat untuk mengekspresikan opini, kritik, dan penilaian terhadap perilaku pejabat pemerintah. Opini tersebut umumnya berbentuk data teks tidak terstruktur, berjumlah besar, serta menggunakan bahasa informal yang beragam. Penelitian sebelumnya menunjukkan bahwa opini masyarakat terhadap perilaku pejabat korupsi di media sosial didominasi oleh sentimen negatif dan memiliki karakteristik teks yang kompleks sehingga sulit dianalisis secara manual [4]. Oleh karena itu, diperlukan pendekatan analisis sentimen berbasis komputasi untuk mengolah data teks dari media sosial X secara sistematis dan objektif. Penelitian ini berfokus pada data teks berbahasa Indonesia yang bersifat informal dan berasal dari media sosial X, dengan tujuan mengklasifikasikan sentimen ke dalam tiga kelas, yaitu positif, netral, dan negatif [5].

Analisis sentimen merupakan proses untuk mengidentifikasi dan mengklasifikasikan opini yang tersaji dalam bentuk teks, dengan tujuan untuk menentukan apakah sentimen tersebut bersifat positif, negatif, atau netral [5]. Dalam beberapa tahun terakhir, metode deep learning seperti *Bidirectional Long Short-Term Memory* (Bi-LSTM), *Convolutional Neural Network* (CNN), dan *Bidirectional Encoder Representations from Transformers* (BERT) telah banyak diterapkan dalam analisis sentimen dan menunjukkan hasil yang menjanjikan [6][7][8].

Namun, sebagian penelitian sebelumnya masih berfokus pada satu metode analisis sentimen atau belum membandingkan beberapa arsitektur deep learning secara komprehensif pada konteks kasus pejabat korupsi di media sosial X. Pemilihan metode CNN, Bi-LSTM, dan BERT didasarkan pada perbedaan karakteristik arsitektur dalam memproses teks, sehingga memungkinkan dilakukan perbandingan kinerja metode secara objektif [4][9].

2. METODE PENELITIAN

Penelitian ini menggunakan pendekatan kuantitatif dengan metode analisis sentimen berbasis deep learning untuk mengklasifikasikan opini masyarakat terhadap kasus pejabat korupsi di media sosial X. Tahapan penelitian disusun secara kronologis, dimulai dari pengumpulan data, prapemrosesan teks, pemodelan, hingga evaluasi performa model, sehingga hasil penelitian dapat dianalisis secara sistematis dan objektif.

2.1 Pengumpulan Data

Data penelitian diperoleh melalui proses crawling menggunakan Google Colab dengan autentikasi Twitter Authentication Token. Proses crawling dilakukan dengan mengumpulkan tweet yang mengandung kata kunci "*pejabat korupsi*", "*korupsi pemerintah*", "*korupsi pejabat publik*", dan "*koruptor*" dalam rentang waktu 1 Januari 2022 hingga 21 Januari 2026 dari media sosial X. Hasil crawling disimpan dalam format CSV dengan atribut yang mencakup *conversation_id_str*, *created_at*, *favorite_count*, *full_text*, *id_str*, *image_url*, *in_reply_to_screen_name*, *lang*, *location*, *quote_count*, *reply_count*, *retweet_count*, *tweet_url*, *user_id_str*, dan *username*. Data yang terkumpul berjumlah 7.863 tweet. Setelah melalui proses filtering untuk menghilangkan duplikasi dan data tidak valid, diperoleh 6.904 tweet yang selanjutnya digunakan dalam penelitian.

2.2 Pemrosesan Data

Prapemrosesan data merupakan tahapan penting untuk meningkatkan kualitas data sebelum dilakukan pemodelan [10]. Tahapan prapemrosesan yang dilakukan dalam penelitian ini meliputi:

- **Cleaning:** Menghapus noise seperti URL, mention (@), hashtag (#), emoticon, angka, simbol, dan karakter khusus yang tidak relevan [10].
- **Case Folding:** Mengubah seluruh teks menjadi huruf kecil untuk menyeragamkan format [2].
- **Normalisasi:** Mengubah katakata tidak baku menjadi bentuk baku menggunakan kamus normalisasi bahasa Indonesia [11].
- **Tokenisasi:** Memecah teks menjadi unit-unit kata (token) yang lebih kecil [11].
- **Stopword Removal:** Menghapus kata-kata yang tidak memiliki makna signifikan seperti 'yang', 'dan', 'di', dll [11].
- **Stemming:** Mengubah kata berimbuhan menjadi kata dasar menggunakan library Sastrawi [11].

2.3 Pelabelan Dan Pembagian Data

Setelah prapemrosesan, data tweet diberi label sentimen secara manual menjadi tiga kategori: positif (dukungan terhadap penegakan hukum/pemberantasan korupsi), negatif (kecaman/kritik terhadap korupsi dan pejabat korup), dan netral (informasi faktual tanpa opini yang jelas). Dataset yang telah diberi label kemudian dibagi menjadi dua bagian dengan proporsi 80% untuk data latih (5.523 tweet) dan 20% untuk data uji (1.381 tweet).

Tabel 2.1 Distribusi Label Sentimen pada Dataset

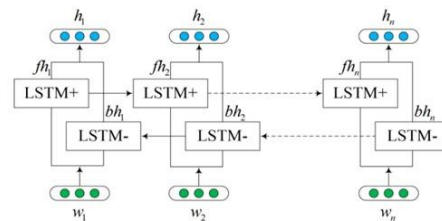
Label	Jumlah Data	Persentase
Negatif	5.527	80,05%
Netral	1.335	19,34%
Positif	42	0,61%
Total	6.904	100%

2.4 Model Klasifikasi

Penelitian ini mengimplementasikan tiga model deep learning untuk klasifikasi sentimen:

1. Bidirectional Long Short-Term Memory (Bi-LSTM)

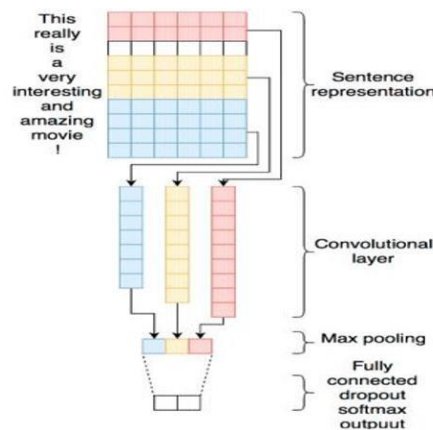
Model Bi-LSTM dikembangkan menggunakan arsitektur jaringan saraf yang mampu memproses teks dalam dua arah, yaitu forward dan backward, se-hingga model dapat memahami konteks kata berdasarkan kata sebelum dan sesudahnya. Proses pemodelan dimulai dari embedding layer yang mengubah kata menjadi vektor numerik, dilanjutkan dengan lapisan Bi-LSTM yang menangkap ketergantungan jangka panjang dalam teks. Dropout diterapkan sebagai regularisasi untuk mencegah overfitting. Lapisan output menggunakan softmax untuk menghasilkan prediksi kelas sentimen. Bi-LSTM dipilih karena memiliki performa yang baik dalam menganalisis teks pendek seperti tweet [11].



Gambar 2.1 Arsitektur Bi-LSTM

2. Convolutional Neural Network (CNN)

Model CNN digunakan untuk mengekstraksi fitur lokal berdasarkan pola n-gram dalam teks. Proses dimulai dengan embedding layer, kemudian teks di-proses melalui lapisan Conv1D yang mengekstraksi fitur berdasarkan filter konvolusi. Selanjutnya digunakan max pooling untuk menyaring fitur paling signifikan dari setiap filter. Fitur yang diperoleh kemudian diratakan (flatten) dan di-proses melalui dense layer untuk menghasilkan prediksi sentimen. CNN efektif untuk analisis sentimen karena mampu mengenali pola lokal dalam teks yang berulang [12].

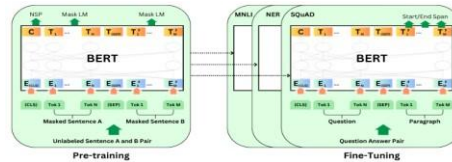


Gambar 2.2 Arsitektur CNN

3. Bidirectional Encoder Representations from

Transformers (BERT) Model BERT digunakan melalui proses fine-tuning terhadap model pra-latih Bahasa Indonesia seperti IndoBERT. Pada tahap awal, teks ditokenisasi menggunakan WordPiece agar sesuai dengan struktur input BERT. Selanjutnya token

diproses oleh transformer encoder yang menggunakan multi-head self-attention untuk memahami hubungan antar kata secara mendalam. Di bagian akhir ditambahkan classification layer untuk menghasilkan prediksi sentimen. BERT sering memberikan hasil terbaik dalam analisis sentimen karena kemampuannya memahami konteks dua arah secara komprehensif [12].



Gambar 2.3 Prosedur Pre-Training dan Fine-Tuning pada Arsitektur BERT

2.5 Evaluasi Model

Evaluasi dalam penelitian ini peneliti menggunakan Multiclass Confusion Matrix 3×3 yang membandingkan kelas sentimen aktual dengan hasil prediksi model. Karena sentimen diklasifikasikan ke dalam tiga kategori positif, netral, dan negatif

Metrix yang digunakan meliputi accuracy, precision, recall, dan F1-score. Rumus perhitungan metrik evaluasi adalah sebagai berikut:

- **Accuracy** Accuracy digunakan untuk menilai sejauh mana model mampu melakukan klasifikasi dengan benar [12].

$$accuracy = \frac{TP + TN}{TP + FN + FP + TN} \times 100\%$$

- **Precision** Precision digunakan untuk mengukur seberapa tepat model dalam memprediksi kelas positif, yaitu proporsi prediksi positif yang benar dari seluruh prediksi yang dikategorikan sebagai positif [12].

$$precision = \frac{TP}{TP + FP}$$

- **Recall**

Recall digunakan untuk mengetahui seberapa banyak data positif yang berhasil dikenali dengan benar oleh model dari seluruh data yang sebenarnya positif [12].

$$recall = \frac{TP}{TP + FN}$$

- **F1-score**

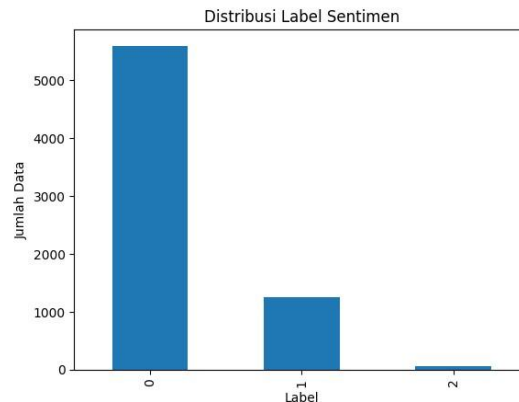
F1-Score digunakan untuk memberikan nilai keseimbangan antara precision dan recall, terutama ketika distribusi kelas tidak seimbang [12].

$$F1 - score = \frac{precision \times recall}{precision + recall}$$

3. HASIL DAN PEMBAHASAN

3.1 Analisis Distribusi Data

Dataset yang digunakan terdiri dari 6.904 tweet terkait kasus pejabat korupsi. Distribusi sentimen menunjukkan ketidakseimbangan kelas yang signifikan, dengan sentimen negatif mendominasi (80,05%), diikuti sentimen netral (19,34%), dan sentimen positif hanya 0,61%. Ketidakseimbangan ini mencerminkan realitas opini publik yang didominasi oleh kecaman dan kritik terhadap kasus korupsi pejabat.



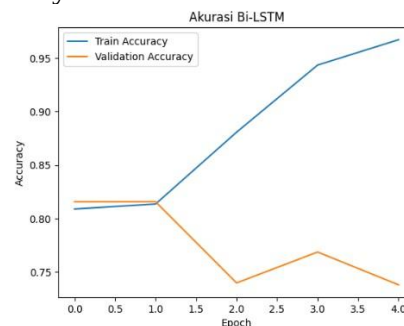
Gambar 3.1 Distribusi Label Sentimen

3.2 Hasil Pemodelan Evaluasi

Tiga model deep learning Bi-LSTM, CNN, dan BERT dilatih menggunakan data latih dan dievaluasi menggunakan data uji. Berikut adalah hasil evaluasi dari masing-masing model:

a. Model Bi-LSTM

Model Bi-LSTM dilatih dengan konfigurasi 128 unit LSTM dan dropout 0.5. Proses pelatihan dilakukan selama 10 epoch dengan batch size 32. Model Bi-LSTM terdiri dari lapisan embedding dengan dimensi vektor sebesar 128 dan panjang sekuens 50 kata, diikuti oleh satu lapisan Bidirectional LSTM dengan 64 unit, serta lapisan dense dengan fungsi aktivasi softmax untuk klasifikasi ke dalam tiga kelas sentimen. Total parameter yang dilatih pada model ini berjumlah 1.379.203 parameter, yang seluruhnya bersifat trainable.



Gambar 3.4 Grafik Akurasi Pelatihan Model Bi-LSTM

Tabel 3.1 Hasil Evaluasi Model Bi-LSTM

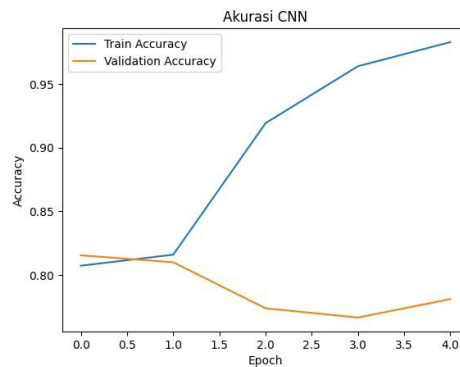
Evaluasi	Nilai
Accuracy	71,54%

Loss	1,01
------	------

Model Bi-LSTM menunjukkan performa yang cukup baik dengan akurasi 74,5%. Model ini memiliki performa terbaik pada kelas negatif dengan F1-score 0,85, yang sesuai dengan distribusi data yang didominasi oleh sentimen negatif. Namun, performa pada kelas positif masih rendah (F1-score 0.41) karena keterbatasan jumlah data pada kelas tersebut.

b. Model CNN

Model yang menggunakan filter konvolusi untuk mengekstraksi fitur lokal dari teks. Arsitektur CNN menggunakan embedding layer (dimensi 100), convolutional layer (128 filters, kernel size 5), max pooling, dropout (0.5), dan fully connected layers. Model dilatih selama 10 epoch dengan batch size 64.



Gambar 3.8 Grafik Akurasi Pelatihan Model CNN

Evaluasi	Nilai
Accuracy	76,03%
Loss	0,94

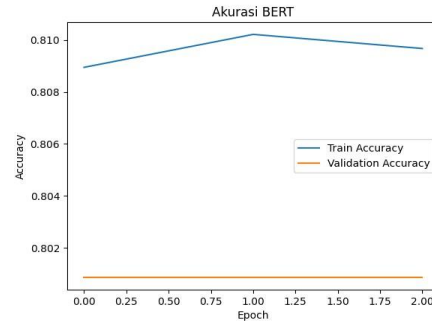
Tabel 3.2 Hasil Evaluasi Model CNN

Model CNN menunjukkan peningkatan performa dibandingkan Bi-LSTM dengan akurasi 77,8%. Model ini juga memiliki performa yang sangat baik pada kelas negatif dengan F1-score 0,88 dan recall 0,91. Performa pada kelas netral juga mengalami peningkatan dengan F1score 0,59. Namun, performa pada kelas positif masih perlu ditingkatkan.

c. Model BERT

Model BERT yang digunakan dalam penelitian ini adalah IndoBERT-base-p1 yang telah dilatih sebelumnya menggunakan korpus bahasa Indonesia. Proses fine-tuning dilakukan dengan menambahkan lapisan klasifikasi pada bagian akhir model untuk menyesuaikan dengan tugas klasifikasi sentimen ke dalam tiga kelas, yaitu (negatif, netral, dan positif). Input teks terlebih dahulu ditokenisasi menggunakan tokenizer BERT dengan skema WordPiece serta dilengkapi dengan attention mask agar sesuai dengan format masukan model.

Model dilatih selama 3 epoch dengan batch size sebesar 8 dan learning rate sebesar $2e-5$. Berdasarkan hasil pelatihan, diperoleh nilai akurasi data latih sebesar 80,89% dan akurasi validasi sebesar 80,09%, dengan nilai loss validasi sebesar 0,54. Hasil ini menunjukkan bahwa model BERT mampu mempelajari representasi kontekstual teks dengan baik dan menghasilkan performa klasifikasi yang stabil.



Gambar 3.12 Grafik Akurasi Model BERT

Evaluasi	Nilai
Accuracy	80,09%
Loss	0,54

Tabel 3.2 Hasil Evaluasi Model BERT

Berdasarkan hasil pengujian menggunakan multiclass confusion matrix, diperoleh bahwa model BERT cenderung mengklasifikasikan seluruh data uji ke dalam kelas negatif, sehingga tidak menghasilkan prediksi pada kelas netral dan positif. Hal ini mengindikasikan adanya bias model terhadap kelas mayoritas, yang disebabkan oleh ketidakseimbangan distribusi data pada dataset, di mana kelas negatif mendominasi sebesar 80,05%, sedangkan kelas netral dan positif hanya sebesar 19,34% dan 0,61%.

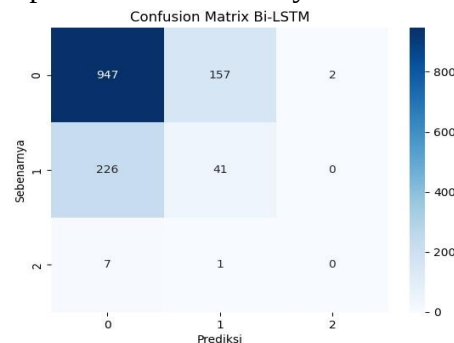
Meskipun nilai akurasi yang diperoleh tergolong tinggi, hasil ini belum sepenuhnya mencerminkan kemampuan model dalam melakukan klasifikasi multikelas secara seimbang. Oleh karena itu, diperlukan penerapan strategi penanganan ketidakseimbangan kelas, seperti oversampling, undersampling, atau pemberian class weight, guna meningkatkan kemampuan model dalam mengenali kelas minoritas, khususnya sentimen positif.

3.3 Analisis Confusion Matrix

Confusion matrix digunakan untuk memberikan gambaran yang lebih rinci mengenai kemampuan model dalam mengklasifikasikan setiap kelas sentimen, yaitu negatif (0), netral (1), dan positif (2). Matriks ini menunjukkan jumlah data yang diprediksi dengan benar maupun salah oleh model pada masing-masing kelas.

a. Model Bi-LSTM

Berdasarkan confusion matrix Bi-LSTM, jumlah data uji yang digunakan adalah 1.381 data, dengan jumlah prediksi benar sebanyak 988 data.



Gambar 3.13 Confusion Matrix Bi-LSTM

Prediksi/Prediksi	0	1	2
Negatif (0)	947	157	2
Netral (1)	226	41	0
Positif (2)	7	1	0

Tabel 3.4 Confusion Matrix Bi-LSTM

Tabel ini menunjukkan perbandingan antara kelas aktual dan kelas prediksi pada klasifikasi sentimen menggunakan model BiLSTM. Baris merepresentasikan kelas aktual, sedangkan kolom menunjukkan hasil prediksi model untuk tiga kelas sentimen, yaitu (negatif, netral, dan positif).

- **Perhitungan Accuracy**

$$Accuracy = \frac{988}{1381} \times 100\% = 71,54\%$$

- **Perhitungan Precision**

$$Precision\ Negatif = \frac{947}{947 + 226 + 7} = \frac{947}{1180} = 0,80$$

$$Precision\ Netral = \frac{41}{157 + 41 + 1} = \frac{41}{199} = 0,21$$

$$Precision\ positif = \frac{0}{2 + 0 + 0} = 0$$

- **Perhitungan Recall**

$$Recall\ Negatif = \frac{947}{947 + 157 + 2} = \frac{947}{1106} = 0,86$$

$$Recall\ Netral = \frac{41}{226 + 41 + 0} = \frac{41}{267} = 0,15$$

$$Recall\ Positif = \frac{0}{7 + 1 + 0} = 0$$

- **Perhitungan F1-Score**

$$F1 - Score\ Negatif = \frac{2(0,90 \times 0,86)}{0,90 + 0,86} = 0,83$$

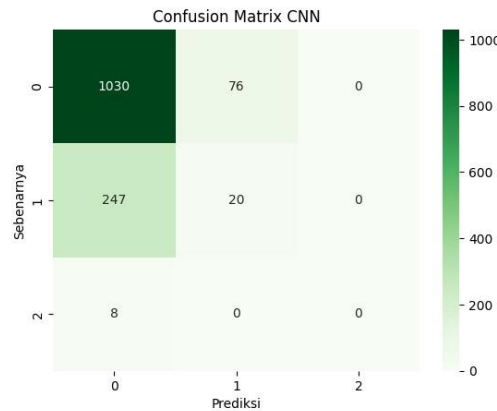
$$F1 - Score\ Netral = \frac{2(0,21 \times 0,15)}{0,21 + 0,15} = 0,17$$

$$F1 - Score\ Positif = 0$$

Berdasarkan hasil evaluasi, model BiLSTM menunjukkan performa yang baik dalam mengklasifikasikan sentimen negatif, namun belum optimal dalam mengenali sentimen netral dan positif. Hal ini disebabkan oleh ketidakseimbangan jumlah data antar kelas, di mana kelas negatif mendominasi data uji.

b. Model CNN

Confusion matrix digunakan untuk memberikan gambaran rinci mengenai kemampuan model CNN dalam mengklasifikasikan sen-timen negatif (0), netral (1), dan positif (2).

**Gambar 3.14** *Confusion Matrix CNN*

Aktual/Prediksi	0	1	2
Negatif (0)	1030	76	0
Netral (1)	247	20	0
Positif (2)	8	0	0

Tabel 3.5 *Confusion Matrix CNN*

Berdasarkan tabel tersebut, terlihat bahwa model CNN mampu mengklasifikasikan sebagian besar data sentimen negatif dengan benar. Namun, model masih mengalami kesalahan yang cukup besar pada kelas netral dan sama sekali tidak mampu mengenali kelas positif.

- Perhitungan Accuracy**

$$Accuracy = \frac{1030 + 20 + 0}{1381} \times 100\% = 76,03\%$$

- Perhitungan Precision**

$$Precision \text{ Negatif} = \frac{1030}{1030 + 247 + 8} = \frac{1030}{1285} = 0,80$$

$$Precision \text{ Netral} = \frac{20}{76 + 20 + 0} = \frac{20}{96} = 0,21$$

$$Precision \text{ positif} = \frac{0}{0 + 0 + 0} = 0,00 = 0,00$$

- Perhitungan Recall**

$$Recall \text{ Negatif} = \frac{1030}{1030 + 76} = \frac{1030}{1106} = 0,93$$

$$Recall \text{ Netral} = \frac{20}{247 + 20} = \frac{20}{267} = 0,07$$

$$Recall\ Positif = \frac{0}{8 + 0} = 0,00$$

- **Perhitungan F1-Score**

$$F1 - Score\ Negatif = 2 \times \frac{(0,80 \times 0,93)}{0,80 + 0,93} = \frac{0,744}{1,73} 0,86$$

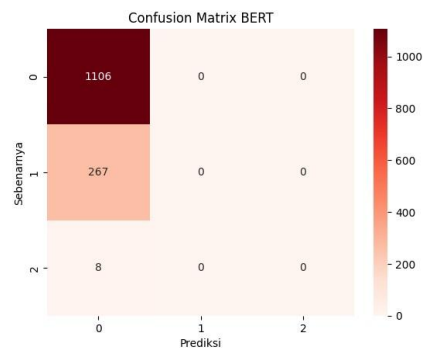
$$F1 - Score\ Netral = 2 \times \frac{(0,21 \times 0,07)}{0,21 + 0,07} = \frac{0,0147}{0,28} 0,11$$

$$F1 - Score\ Positif = 0,00$$

c. Model BERT

Model BERT yang digunakan dalam penelitian ini adalah IndoBERT-base-p1 yang telah dilatih sebelumnya menggunakan korpus bahasa Indonesia. Proses fine-tuning dilakukan dengan menambahkan lapisan klasifikasi pada bagian akhir model untuk menyesuaikan dengan tugas klasifikasi sentimen ke dalam tiga kelas, yaitu negatif, netral, dan positif. Input teks terlebih dahulu ditokenisasi menggunakan tokenizer BERT dengan skema WordPiece dan dilengkapi dengan attention mask.

Proses pelatihan dilakukan selama 3 epoch dengan batch size 8 dan learning rate $2e-5$. Berdasarkan hasil pelatihan, diperoleh nilai akurasi data latih sebesar 80,89% dan akurasi validasi sebesar 80,09%, dengan nilai loss validasi sebesar 0,54. Hasil ini menunjukkan bahwa model mampu mempelajari representasi kontekstual teks dengan cukup baik.



Gambar 3.15 *Confusion Matrix BERT*

Berdasarkan hasil pengujian, Evaluasi lebih lanjut dilakukan menggunakan multiclass confusion matrix 3×3 untuk menganalisis performa klasifikasi pada masing-masing kelas sentimen. Confusion matrix bertujuan untuk melihat distribusi prediksi model terhadap kelas aktual, sehingga dapat diketahui sejauh mana model mampu mengenali setiap kategori sentimen.

Tabel 3.6 *Confusion Matrix BERT*

Aktual/Prediksi	0	1	2
Negatif (0)	1106	0	0
Netral (1)	267	0	0

Positif (2)	8	0	0
-------------	---	---	---

Berdasarkan confusion matrix tersebut, terlihat bahwa seluruh data uji diprediksi ke dalam kelas negatif, sedangkan tidak terdapat prediksi pada kelas netral dan positif. Hal ini menunjukkan bahwa model BERT mengalami bias terhadap kelas mayoritas, yaitu sentimen negatif, yang mendominasi distribusi dataset.

Kondisi ini disebabkan oleh ketidakseimbangan distribusi data, di mana kelas negatif mencapai 80,05%, sedangkan kelas netral hanya 19,34%, dan kelas positif hanya 0,61%. Ketidakseimbangan ini menyebabkan model cenderung mempelajari pola dominan pada kelas negatif, sehingga mengabaikan karakteristik kelas minoritas.

- **Perhitungan Accuracy**

$$Accuracy = \frac{1106}{1381} \times 100\% = 80,07\%$$

- **Perhitungan Precision**

$$Precision \text{ Negatif} = \frac{1106}{1106 + 267 + 8} = \frac{1106}{1381} = 0,80$$

$$Precision \text{ Netral} = \frac{0}{0} = \frac{0}{0} = 0,00$$

$$Precision \text{ positif} = \frac{0}{0 + 0 + 0} = 0,00$$

- **Perhitungan Recall**

$$Recall \text{ Negatif} = \frac{1106}{1106 + 0} = 1,00$$

$$Recall \text{ Netral} = \frac{0}{267 + 0} = 0,00$$

$$Recall \text{ Positif} = \frac{0}{8 + 0} = 0,00$$

- **Perhitungan F1-Score**

$$F1 - Score \text{ Negatif} = 2 \times \frac{(0,80 \times 1,00)}{0,8 + 1,00} = 0,89$$

$$F1 - Score \text{ Netral} = 0,00$$

$$F1 - Score \text{ Positif} = 0,00$$

Hasil evaluasi menunjukkan bahwa model BERT memiliki performa yang sangat baik dalam mengenali sentimen negatif, ditunjukkan oleh nilai recall sebesar 1,00 dan F1-score sebesar 0,89. Namun, model ini tidak mampu mengklasifikasikan sentimen netral dan positif, sehingga seluruh metrik evaluasi pada kedua kelas tersebut bernilai nol.

3.4 Pembahasan

Tabel 4.3 Perbandingan Hasil Evaluasi Model

Model	Accuracy (%)	Loss
Bi-LSTM	71,54	1,01
CNN	76,03	0,94

BERT	80,09	0,54
------	-------	------

Berdasarkan hasil pengujian, model BERT memperoleh nilai akurasi tertinggi sebesar 80,09%. Namun, hasil confusion matrix menunjukkan bahwa model BERT cenderung mengklasifikasikan seluruh data uji ke dalam satu kelas, yaitu negatif. Hal ini mengindikasikan bahwa nilai akurasi yang tinggi belum sepenuhnya mencerminkan kemampuan klasifikasi multikelas secara optimal.

4. KESIMPULAN DAN SARAN

4.1 Kesimpulan

Berdasarkan hasil penelitian mengenai analisis sentimen pengguna media sosial X terhadap kasus pejabat korupsi menggunakan metode Bi-LSTM, CNN, dan BERT, dapat disimpulkan bahwa mayoritas opini masyarakat didominasi oleh sentimen negatif, yang mencerminkan tingginya tingkat kekecewaan dan kritik publik terhadap praktik korupsi yang melibatkan pejabat negara.

Hasil evaluasi menunjukkan bahwa model BERT menghasilkan nilai akurasi tertinggi sebesar 80,09% dengan nilai loss terendah dibandingkan model BiLSTM dan CNN. Namun, berdasarkan analisis confusion matrix, model BERT cenderung mengklasifikasikan seluruh data uji ke dalam kelas negatif akibat ketidakseimbangan distribusi data, sehingga performa klasifikasi multikelas belum optimal.

Sementara itu, model CNN dan BiLSTM menunjukkan kemampuan yang lebih stabil dalam mengenali lebih dari satu kelas sentimen, meskipun memiliki nilai akurasi yang lebih rendah. Dengan demikian, penelitian ini membuktikan bahwa pemilihan model analisis sentimen tidak hanya bergantung pada nilai akurasi, tetapi juga harus mempertimbangkan keseimbangan prediksi antar kelas, terutama pada dataset dengan distribusi label yang tidak seimbang.

4.2 Saran

Berdasarkan hasil penelitian yang telah dilakukan, terdapat beberapa saran yang dapat dijadikan acuan untuk pengembangan penelitian selanjutnya, yaitu:

1. Penelitian selanjutnya disarankan untuk menerapkan teknik penanganan data tidak seimbang, seperti oversampling, undersampling, atau pemberian class weight, agar model dapat mengenali kelas minoritas dengan lebih baik, khususnya sentimen positif.
2. Perlu dilakukan eksperimen dengan variasi hyperparameter dan arsitektur model, terutama pada model BERT, guna meningkatkan performa klasifikasi multikelas.
3. Penelitian lanjutan dapat menambahkan metode ensemble atau hybrid model untuk memperoleh hasil klasifikasi yang lebih optimal dan stabil.
4. Disarankan untuk memperluas cakupan data dengan menambahkan periode waktu yang lebih panjang atau memperluas kata kunci pencarian agar representasi opini masyarakat menjadi lebih komprehensif.

DAFTAR PUSTAKA

- [1] Strategi Nasional Pencegahan Korupsi, "PENCEGAHAN KORUPSI," Jakarta, 2026.
- [2] E. Setyaningtyas and K. Nugroho, "Analisis Sentimen Media Sosial Pada Pengguna Twitter Terhadap Pemilu 2024 Menggunakan Metode LSTM," *J. Ris. Sist. Inf. Dan Tek. Inform.*, vol. 9, no. 2, pp. 673–683, 2024, [Online]. Available: <https://tunasbangsa.ac.id/ejurnal/index.php/jurasik>

- [3] T. Hamim Aftoni, W. H. Sugiharto, and M. I. Ghazali, “Analisis Sentimen Terhadap Tagar #Kaburajadulu Di Media Sosial X Menggunakan Metode Naïve Bayes,” *J. Komput. dan Teknol.*, vol. 4, no. 2, pp. 154–167, 2025, doi: 10.64626/jukomtek.v4i2.443.
- [4] A. Syakir and F. N. Hasan, “Analisis Sentimen Masyarakat Terhadap Perilaku Korupsi Pejabat Pemerintah Berdasarkan Tweet Menggunakan Naive Bayes Classifier,” vol. 7, pp. 1796–1805, 2023, doi: 10.30865/mib.v7i4.6648.
- [5] B. Pang and L. Lee, “Opinion mining and sentiment analysis,” vol. 2, no. 1, 2008.
- [6] T. Pemilu, M. Khadapi, and V. M. Pakpahan, “Analisis Sentimen Berbasis Jaringan LSTM dan BERT terhadap Diskusi,” vol. 6, pp. 130–137, 2024.
- [7] A. He and M. Abisado, “Text Sentiment Analysis of Douban Film Short Comments Based on BERT-CNN- BiLSTM-Att Model,” *IEEE Access*, vol. PP, p. 1, 2024, doi: 10.1109/ACCESS.2024.3381515.
- [8] L. Khan, A. Amjad, K. M. Afaq, and H. T. Chang, “Deep Sentiment Analysis Using CNN-LSTM Architecture of English and Roman Urdu Text Shared in Social Media,” *Appl. Sci.*, vol. 12, no. 5, 2022, doi: 10.3390/app12052694.
- [9] Y. Sibaroni and S. S. Prasetyowati, “Jurnal Resti,” *Resti*, vol. 1, no. 1, pp. 19–25, 2018.
- [10] A. Elhan, M. Kusuma, D. Hardhienata, Y. Herdiyeni, S. H. Wijaya, and J. Adisantoso, “Analisis Sentimen Pengguna Twitter terhadap Vaksinasi COVID-19 di Indonesia menggunakan Algoritme Random Forest dan BERT Sentiment Analysis of Twitter Users on COVID-19 Vaccines in Indonesia using Random Forest and BERT Algorithms,” vol. 9.
- [11] A. N. Hidayati and Y. Yamasari, “Analisis Sentimen Masyarakat Terhadap Fenomena Sandwich Generation Pada Aplikasi X Menggunakan Metode Naïve Bayes dan Teknik Pelabelan Lexicon Inset,” vol. 07, pp. 85–94, 2025.
- [12] M. R. F. Kamarula and N. Rochmawati, “Perbandingan CNN dan Bi-LSTM pada Analisis Sentimen dan Emosi Masyarakat Indonesia Di Media Sosial Twitter Selama Pandemi Covid-19 yang Menggunakan Metode Word2vec,” *J. Informatics Comput. Sci.*, vol. 04, pp. 219–228, 2022, doi: 10.26740/jinacs.v4n02.p219-228.