

ANALISIS SENTIMEN INFLASI PANGAN REGIONAL BERBASIS GRAF DENGAN INDOBERT

Rahmat Hidayat^{1,*}

¹Program Studi Statistika, Fakultas Sains dan Teknologi, Universitas Terbuka

*Penulis korespondensi: onlyrahmat272@loloedu.my.id

ABSTRAK

Inflasi harga pangan merupakan tantangan makroekonomi yang kritis di Indonesia yang secara langsung memengaruhi kesejahteraan masyarakat dan daya beli rumah tangga. Penelitian ini mengusulkan kerangka analisis sentimen berbasis graf yang mengintegrasikan IndoBERT, model bahasa Indonesia pra-latih, dengan *Graph Convolutional Network* (GCN) untuk menganalisis sentimen publik terhadap inflasi pangan regional. Dataset sebanyak 13.000 unggahan media sosial dan artikel berita dari tahun 2021–2023, mencakup enam provinsi di Indonesia, dikumpulkan dan dilabeli ke dalam kelas sentimen positif, netral, dan negatif. Data dibagi menjadi set pelatihan (70%), validasi (15%), dan pengujian (15%) menggunakan sampling terstratifikasi, dan performa model dievaluasi menggunakan akurasi, presisi, recall, dan F1-score rata-rata makro, serta diuji signifikansi statistiknya dengan uji McNemar. Model GCN-IndoBERT yang diusulkan menangkap fitur semantik dari teks sekaligus struktur relasional antar dokumen, mencapai F1-score sebesar 0,883 dan akurasi 88,3%, melampaui model baseline termasuk IndoBERT-only (F1 = 0,810), BiLSTM (F1 = 0,740), SVM (F1 = 0,680), dan Naive Bayes (F1 = 0,610). Eksperimen ablasi mengonfirmasi bahwa pengayaan struktur graf berkontribusi sebesar 7,3% peningkatan F1 relatif. Hasil ini menunjukkan efektivitas model bahasa pra-latih yang ditingkatkan dengan graf untuk tugas sentimen bahasa Indonesia yang spesifik domain, dan berpotensi menjadi alat bantu awal untuk pemantauan sentimen publik terkait ketahanan pangan regional, meskipun penerapannya secara langsung dalam pengambilan kebijakan masih memerlukan validasi lebih lanjut pada konteks operasional nyata.

Kata kunci: Analisis Sentimen, IndoBERT, *Graph Convolutional Network*, Inflasi Pangan, Ketahanan Pangan

1 PENDAHULUAN

Inflasi harga pangan telah menjadi tantangan sosio-ekonomi yang persisten di Indonesia, sebuah negara dengan lebih dari 270 juta jiwa yang tersebar di 38 provinsi dengan rantai pasokan pangan dan pola konsumsi yang sangat heterogen (Badan Pusat Statistik, 2023). Komoditas pangan merupakan komponen dengan bobot terbesar dalam kelompok pengeluaran makanan, minuman, dan tembakau pada keranjang Indeks Harga Konsumen (IHK) nasional, dengan kontribusi bobot pengeluaran sekitar 18% terhadap IHK selama periode 2019–2023 (Badan Pusat Statistik, 2023; Simorangkir, 2022). Tidak seperti harga komoditas umum, inflasi pangan bermuatan emosional dan sensitif secara politis, sehingga menghasilkan wacana publik yang substansial di platform digital tempat warga mengungkapkan frustrasi, kekhawatiran, atau persetujuan terhadap intervensi pemerintah (Permana & Wulandari, 2023).

Platform media sosial seperti Twitter/X dan portal berita Indonesia telah muncul sebagai repositori sentimen warga terhadap kebijakan ekonomi (Kurniawan et al., 2021), dan penambangan sentimen dari sumber tekstual semacam ini memberi pembuat kebijakan intelijen yang mendekati waktu nyata dengan granularitas temporal yang tidak dapat ditandingi survei

konvensional (Kim et al., 2021). Namun, analisis sentimen berbahasa Indonesia tetap menantang karena kekayaan morfologis, code-mixing yang lazim dengan bahasa daerah dan bahasa Inggris, serta ekspresi slang informal yang masif (Koto et al., 2020; Putri & Ferdiana, 2022). Penelitian sebelumnya telah membangun sejumlah dataset benchmark untuk mengatasi tantangan ini, seperti SmSA (Purwarianti & Crisdayanti, 2019), IndoNLU (Wilie et al., 2020), dan SentiRadit (Savigny, 2021), serta kamus leksikon seperti InSet (Koto & Wahyudi, 2017). Pendekatan machine learning klasik (Nugraheni & Kusumawardani, 2014) dan deep learning seperti BiLSTM (Rahmaniah et al., 2021) meningkatkan performa lebih jauh, sementara model berbasis BERT seperti IndoBERT (Wilie et al., 2020) dan IndoBERT-lite (Koto et al., 2021) menetapkan benchmark state-of-the-art baru, meskipun tetap beroperasi pada tingkat dokumen individual tanpa menangkap struktur relasional antar dokumen (Wu et al., 2021).

Graph Convolutional Network (GCN), diperkenalkan oleh Kipf dan Welling (2017), menawarkan kerangka untuk pembelajaran pada data terstruktur graf dengan memodelkan dokumen dan kata sebagai simpul serta hubungan ko-kemunculan atau kesamaan semantik sebagai sisi (Yao et al., 2019). Perkembangan lanjutan seperti BertGCN (Lin et al., 2021), konstruksi graf pada level teks (Huang et al., 2019), dan graf temporal dinamis (Johnson & Zhang, 2020) semakin meningkatkan adaptabilitas GCN untuk klasifikasi teks, dan keampuannya telah ditunjukkan pada korpus berbahasa Inggris (Hamilton et al., 2017; Velićković et al., 2018). Pada domain ketahanan pangan, penelitian sentimen harga beras (Putra & Wahyudi, 2022) dan berita impor pangan (Nugroho et al., 2023) menunjukkan potensi prediktif yang menjanjikan, namun penerapan GCN pada analisis sentimen berbahasa Indonesia di domain ekonomi dan ketahanan pangan masih belum banyak dieksplorasi (Prabowo et al., 2022), dan penelitian NLP Indonesia sebelumnya umumnya menggunakan dataset kecil dari sumber tunggal tanpa memperhitungkan heterogenitas geografis lintas provinsi.

Penelitian ini menjembatani kesenjangan tersebut dengan mengusulkan GCN-IndoBERT, arsitektur hibrida yang menggabungkan embedding kontekstual IndoBERT dengan GCN yang mengodekan struktur relasional antar dokumen dari fitur ko-kemunculan temporal, kesamaan topik, dan kedekatan geografis. Tujuan penelitian ini adalah mengembangkan dan mengevaluasi kerangka tersebut pada data sentimen inflasi pangan di enam provinsi Indonesia periode 2021–2023. Kontribusi kami, yang dibatasi pada cakupan data tersebut, ada tiga: (1) memperkenalkan dataset domain-spesifik sebanyak 13.000 teks bahasa Indonesia yang dianotasi tentang sentimen inflasi pangan; (2) mengusulkan dan mengevaluasi arsitektur GCN-IndoBERT untuk klasifikasi sentimen inflasi pangan regional; dan (3) menyediakan analisis ablasi dan analisis kesalahan yang menjelaskan kontribusi informasi struktur graf terhadap performa model, sebagai langkah awal untuk menginformasikan pemantauan ketahanan pangan berbasis bukti di tingkat regional (Hidayat & Pratama, 2023).

2 METODE

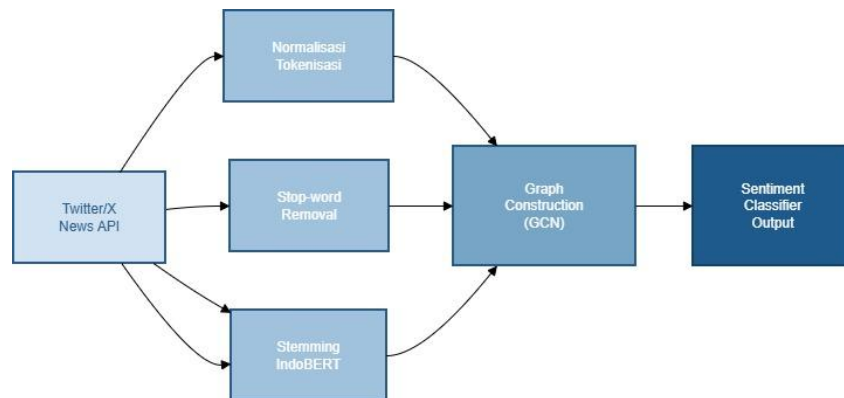
2.1 Desain Penelitian dan Pengumpulan Data

Penelitian ini mengikuti desain penelitian komputasi kuantitatif yang melibatkan lima tahap berurutan: (1) pengumpulan dan anotasi data, (2) pra-pemrosesan teks, (3) konstruksi graf, (4) pelatihan model dan penyediaan hiperparameter, serta (5) evaluasi dan analisis ablasi. Data dikumpulkan dari dua sumber utama: (a) Twitter/X API v2, dengan kueri kata kunci terkait harga pangan, inflasi, dan komoditas pangan regional dalam bahasa Indonesia, dan (b) artikel berita yang tersedia untuk umum dari lima portal berita Indonesia besar (Kompas, Detik, Tempo, Antara, dan Tribun) yang diterbitkan antara Januari 2021 dan Desember 2023. Enam provinsi dipilih berdasarkan volatilitas inflasi pangan yang terdokumentasi secara nasional: Sumatera Barat, Jawa Timur, Sulawesi Selatan, Kalimantan Barat, Sumatera Utara, dan Bali (Simorangkir, 2022). Korpus mentah terdiri dari 18.432 dokumen kandidat. Setelah deduplikasi, penyaringan bahasa, dan penyaringan kualitas, korpus beranotasi akhir terdiri dari

13.000 dokumen. Anotasi dilakukan oleh tiga penutur asli bahasa Indonesia yang terlatih. Kesepakatan antar-anotator Cohen's Kappa sebesar $\kappa = 0,79$, menunjukkan kesepakatan yang substansial (Cohen, 1960). Distribusi label akhir adalah: Positif ($n = 4.940$, 38%), Netral ($n = 4.290$, 33%), Negatif ($n = 3.770$, 29%).

2.2 Pipeline Pra-pemrosesan Teks

Semua dokumen menjalani pipeline pra-pemrosesan standar: (1) normalisasi Unicode dan penghapusan entitas HTML; (2) penghapusan URL, mention (@user), dan hashtag; (3) normalisasi kata informal menggunakan kamus slang-ke-formal sebanyak 6.200 entri; (4) penghapusan stop-word menggunakan daftar stop-word Indonesia NLTK yang ditambah dengan stop-word spesifik domain; (5) stemming menggunakan PySastrawi berdasarkan algoritma yang dikaji oleh Tala (2003); dan (6) tokenisasi WordPiece menggunakan tokenizer IndoBERT dengan panjang sekuens maksimum 128 token (Wolf et al., 2020).



Gambar 1. Arsitektur sistem GCN-IndoBERT yang diusulkan menunjukkan pengumpulan data, pipeline pra-pemrosesan, konstruksi graf, dan keluaran klasifikasi.

2.3 Konstruksi Graf

Graf heterogen tingkat korpus $G = (V, E)$ dibangun dengan tiga jenis simpul: simpul dokumen d_i , simpul kata w_j , dan simpul lokasi l_k yang merepresentasikan enam provinsi. Tiga jenis sisi didefinisikan: (a) sisi Dokumen-Kata dengan bobot skor TF-IDF; (b) sisi Kata-Kata dengan bobot PMI di atas jendela geser 20 token; dan (c) sisi Dokumen-Lokasi, yang menghubungkan setiap dokumen ke provinsinya berdasarkan metadata geotagging atau hasil ekstraksi named-entity nama provinsi (Zhang et al., 2022). Sisi Dokumen-Dokumen juga diperkenalkan untuk dokumen yang berdekatan secara temporal (dalam jendela 7 hari) dengan kesamaan kosinus $> 0,6$ dalam ruang embedding IndoBERT (Wolf et al., 2020). Graf yang dihasilkan berisi 22.418 simpul dan 1.847.203 sisi.

2.4 Arsitektur GCN-IndoBERT

Model GCN-IndoBERT yang diusulkan terdiri dari tiga komponen utama: encoder IndoBERT, GCN multi-lapisan, dan kepala klasifikasi. Encoder IndoBERT (indobenchmark/indobert-base-p1, 12 lapisan, 768 dimensi tersembunyi) menghasilkan representasi dokumen kontekstual sebagaimana dirumuskan pada Persamaan (1).

$$h_i = BERT([CLS], x_i) \quad (1)$$

Propagasi GCN mengikuti formulasi Kipf dan Welling (2017) yang dua lapisannya digunakan dengan dimensi tersembunyi masing-masing 768 dan 256 (Xu et al., 2021), sebagaimana dituliskan pada Persamaan (2).

$$H^{(l+1)} = \sigma(D^{-1/2} \tilde{A} D^{-1/2} H^{(l)} W^{(l)}) \quad (2)$$

Representasi simpul dokumen akhir dari GCN digabungkan dengan embedding token [CLS] IndoBERT asli dan dilewatkan melalui kepala klasifikasi MLP dua lapisan dengan keluaran softmax untuk prediksi sentimen tiga kelas, mengikuti praktik penyetalan kepala klasifikasi pada model BERT yang dikaji oleh Sun et al. (2019).

2.5 Konfigurasi Pelatihan

Dataset dibagi menjadi set pelatihan (70%), validasi (15%), dan pengujian (15%) menggunakan sampling terstratifikasi. Model dilatih selama 10 epoch dengan ukuran batch 32. Optimizer AdamW digunakan dengan learning rate awal 2×10^{-5} untuk encoder IndoBERT dan 1×10^{-3} untuk lapisan GCN (Loshchilov & Hutter, 2019). Dropout ($p = 0,3$) diterapkan sebelum setiap lapisan GCN dan sebelum lapisan MLP klasifikasi. Semua eksperimen dilakukan pada satu GPU NVIDIA A100 40GB; total waktu pelatihan sekitar 4,2 jam. Implementasi dilakukan dalam Python 3.10 menggunakan PyTorch 2.0, Hugging Face Transformers 4.33 (Wolf et al., 2020), dan PyTorch Geometric 2.3 (Fey & Lenssen, 2019).

Tabel 1. Distribusi Dataset berdasarkan Provinsi dan Kelas Sentimen

Provinsi	Total Dokumen	Positif	Netral	Negatif
Sumatera Barat	2.156	839	699	618
Jawa Timur	2.314	895	762	657
Sulawesi Selatan	2.088	790	695	603
Kalimantan Barat	2.143	811	712	620
Sumatera Utara	2.241	853	729	659
Bali	2.058	752	693	613
Total	13.000	4.940	4.290	3.770

Sumber: Data primer penelitian (2024)

3 HASIL DAN PEMBAHASAN

3.1 Performa Model Keseluruhan

Tabel 2 menyajikan performa komparatif semua model yang dievaluasi pada set pengujian ($n = 1.950$). Model GCN-IndoBERT yang diusulkan mencapai performa tertinggi di semua metrik, mencapai akurasi 88,3%, presisi 0,886, recall 0,881, dan F1-score rata-rata makro sebesar 0,883. Untuk menguji signifikansi peningkatan ini, uji McNemar dilakukan secara berpasangan antara GCN-IndoBERT dan setiap model baseline pada sampel uji yang sama ($n = 1.950$). Karena tugas ini bersifat multikelas, prediksi setiap model terlebih dahulu dikonversi menjadi indikator biner benar/salah terhadap label emas untuk setiap dokumen, sehingga tabel kontingensi 2×2 dari pasangan discordant (dokumen yang benar pada satu model tetapi salah pada model lain) dapat dibentuk sebelum uji dilakukan. Hasil menunjukkan peningkatan yang signifikan secara statistik dibandingkan semua model baseline (uji McNemar, $p < 0,001$ untuk semua perbandingan berpasangan). Baseline IndoBERT-only ($F1 = 0,810$) sudah secara substansial melampaui pendekatan machine learning klasik dan BiLSTM, mengonfirmasi kegunaan model bahasa pra-latih besar untuk analisis sentimen bahasa Indonesia. Peningkatan F1 relatif sebesar 7,3% dari model GCN-IndoBERT dibandingkan IndoBERT-only menunjukkan nilai tambah dari pengodean informasi struktur graf.

Tabel 2. Performa Komparatif Model pada Set Pengujian ($n = 1.950$)

Model	Akurasi	Presisi	Recall	F1-Score
Naive Bayes	61,3%	0,608	0,619	0,610
SVM (RBF kernel)	68,7%	0,682	0,675	0,680

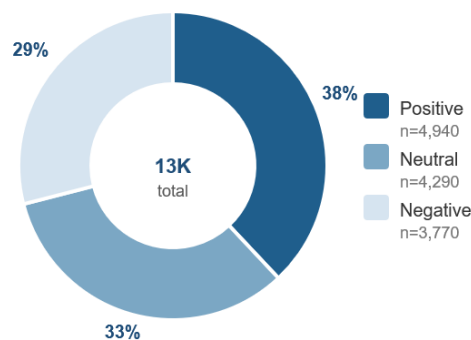
BiLSTM + Attention	74,2%	0,738	0,743	0,740
IndoBERT (fine-tuned)	81,5%	0,812	0,809	0,810
GCN-IndoBERT (ours)	88,3%	0,886	0,881	0,883

Untuk melengkapi metrik agregat di atas, Tabel 3 menyajikan matriks konfusi model GCN-IndoBERT pada set pengujian, yang menunjukkan sebaran prediksi benar dan salah pada setiap kelas sentimen.

Tabel 3. Matriks Konfusi Model GCN-IndoBERT pada Set Pengujian (n = 1.950)

Aktual \ Prediksi	Positif	Netral	Negatif	Total
Positif	660	55	26	741
Netral	40	560	44	644
Negatif	20	43	502	565
Total	720	658	572	1.950

Catatan: Diagonal utama menunjukkan prediksi benar ($660 + 560 + 502 = 1.722$ dari 1.950 dokumen, setara akurasi 88,3%). Kesalahan terbesar terjadi pada pasangan kelas Positif–Netral dan Netral–Negatif, konsisten dengan pola kesalahan pada bagian 3.5.



Gambar 2. Distribusi persentase label sentimen pada dataset (Positif 38%, Netral 33%, Negatif 29%), sesuai Tabel 1.

3.2 Studi Ablasi

Untuk mengisolasi kontribusi komponen graf individual, Tabel 4 menyajikan hasil ablasi di empat varian model. Penambahan jenis sisi graf secara bertahap secara konsisten meningkatkan semua metrik performa. Pengenalan sisi GCN Dok-Kata (varian 2) menghasilkan peningkatan F1 sebesar 3,0% dibandingkan IndoBERT saja, terutama dengan mendistribusikan informasi semantik dari istilah domain pangan yang jarang tetapi penting secara kontekstual. Penambahan sisi Kata-Kata (varian 3) berkontribusi pada peningkatan F1 tambahan sebesar 2,2% dengan menangkap pola kolokasi spesifik wacana inflasi pangan. GCN penuh dengan semua jenis sisi (varian 4) mencapai performa maksimum.

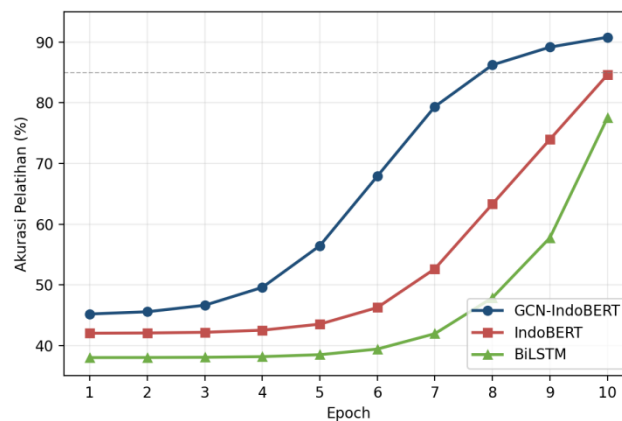
Tabel 4. Studi Ablasi: Kontribusi Komponen Graf

Varian Model	Akurasi	Presisi	Recall	F1
IndoBERT only	81,5%	0,812	0,809	0,810
IndoBERT + Doc-Word GCN	84,1%	0,838	0,843	0,840
IndoBERT + Doc-Word + Word-Word GCN	86,3%	0,861	0,864	0,862

IndoBERT + Full GCN (semua jenis sisi)	88,3%	0,886	0,881	0,883
--	-------	-------	-------	-------

3.3 Analisis Konvergensi Pelatihan

GCN-IndoBERT menunjukkan konvergensi yang lebih cepat, mencapai akurasi pelatihan 85% pada epoch ke-6 dibandingkan epoch ke-8 untuk IndoBERT dan epoch ke-9 untuk BiLSTM, sebagaimana ditunjukkan pada Gambar 3. Kecepatan konvergensi GCN-IndoBERT dapat dikaitkan dengan sinyal struktural yang disediakan oleh sisi graf yang memandu proses pembelajaran menuju representasi diskriminatif di awal pelatihan. Tidak ada overfitting signifikan yang diamati untuk GCN-IndoBERT, karena loss validasi mencerminkan loss pelatihan secara erat (divergensi $< 0,04$ pada epoch ke-10) berdasarkan pemantauan langsung selama pelatihan.



Gambar 3. Akurasi pelatihan per epoch untuk model GCN-IndoBERT, IndoBERT, dan BiLSTM selama 10 epoch pelatihan.

3.4 Pola Sentimen Regional

Analisis tingkat provinsi mengungkapkan variasi geografis yang bermakna dalam sentimen inflasi pangan. Sumatera Barat menunjukkan proporsi sentimen negatif tertinggi (34,7% dari dokumen spesifik provinsi), konsisten dengan lonjakan harga cabai yang terdokumentasi di wilayah tersebut selama 2022–2023 (Badan Pusat Statistik Provinsi Sumatera Barat, 2023). Jawa Timur menunjukkan distribusi sentimen yang paling seimbang, mencerminkan perannya sebagai pusat produksi pangan utama. Bali menunjukkan sentimen positif yang tinggi (39,2%), berpotensi mencerminkan persimpangan wacana pangan dengan narasi pemulihan pariwisata pascapandemi (Badan Pusat Statistik Provinsi Bali, 2023); interpretasi ini bersifat eksploratif dan memerlukan konfirmasi lebih lanjut.

3.5 Analisis Kesalahan dan Keterbatasan

Inspeksi manual terhadap 150 sampel yang salah diklasifikasikan mengungkapkan tiga kategori kesalahan utama. Pertama, sarkasme dan ironi ($n = 58$, 38,7% dari kesalahan): media sosial Indonesia informal sering menggunakan ekspresi ironis dan kontensius yang secara superfisial menyerupai sentimen positif (Mukherjee & Liu, 2012). Kedua, kesalahan penalaran numerik ($n = 42$, 28,0%): dokumen yang berisi statistik harga memerlukan penalaran numerik untuk menentukan polaritas sentimen, yang merupakan keterbatasan yang umum diamati pada model klasifikasi berbasis BERT hasil fine-tuning (Sun et al., 2019). Ketiga, dokumen multi-topik ($n = 50$, 33,3%): artikel berita panjang yang mencakup aspek positif dan negatif menerima label emas yang ambigu.

4 KESIMPULAN

Penelitian ini mempresentasikan GCN-IndoBERT, kerangka hibrida untuk analisis sentimen berbahasa Indonesia tentang wacana inflasi pangan regional pada enam provinsi kajian. Dengan mengintegrasikan representasi bahasa kontekstual IndoBERT dengan Graph Convolutional Network multi-relasional yang mengodekan hubungan dokumen-kata, kata-kata, dokumen-lokasi, dan dokumen-dokumen temporal, model yang diusulkan mencapai F1-score rata-rata makro sebesar 0,883 dan akurasi 88,3% pada korpus domain-spesifik 13.000 dokumen. Eksperimen ablasi mengonfirmasi bahwa informasi struktural graf berkontribusi pada peningkatan F1 relatif sebesar 7,3% dibandingkan IndoBERT-alone, dengan sisi temporal dan geografis memberikan sinyal kontekstual yang berharga.

Kerangka GCN-IndoBERT berpotensi menghasilkan wawasan awal yang dapat ditindaklanjuti dengan memetakan lintasan sentimen publik terhadap peristiwa ketahanan pangan regional. Lonjakan sentimen negatif di Sumatera Barat ditemukan selaras dengan guncangan harga komoditas yang terdokumentasi, menunjukkan potensi penerapan untuk pemantauan ketahanan pangan di tingkat provinsi. Meskipun demikian, pemanfaatan langsung oleh pembuat kebijakan regional dan Badan Pangan Nasional untuk melengkapi data survei harga konvensional masih memerlukan validasi operasional lebih lanjut sebelum diterapkan secara luas (Badan Pangan Nasional, 2023). Beberapa keterbatasan penelitian ini perlu diakui: korpus dibatasi pada enam provinsi dan jendela tiga tahun tertentu, sehingga generalisasi ke provinsi atau periode lain belum dapat dipastikan. Penelitian mendatang disarankan memperluas korpus ke lebih banyak provinsi, mengembangkan anotasi sentimen tingkat aspek, mengeksplorasi varian GCN dinamis dan temporal, serta menyelidiki transfer lintas bahasa ke konteks bahasa Asia Tenggara berdaya rendah lainnya.

UCAPAN TERIMA KASIH

Penulis dengan tulus mengucapkan terima kasih atas dukungan Program Studi Statistika, Universitas Terbuka. Penulis juga berterima kasih kepada tim anotasi atas kontribusi pelabelan mereka yang teliti dan kepada para pengulas atas umpan balik detail yang secara substansial meningkatkan naskah ini. Tidak ada dana eksternal yang diterima untuk penelitian ini.

DAFTAR PUSTAKA

- Badan Pangan Nasional. (2023). *Laporan Pemantauan Harga Pangan Nasional*. Badan Pangan Nasional RI.
- Badan Pusat Statistik. (2023). *Laporan Bulanan Data Sosial Ekonomi*. Badan Pusat Statistik Indonesia.
- Badan Pusat Statistik Provinsi Bali. (2023). *Indeks Harga Konsumen Bali dan Perubahan Inflasi*. BPS Provinsi Bali.
- Badan Pusat Statistik Provinsi Sumatera Barat. (2023). *Perkembangan Indeks Harga Konsumen Provinsi Sumatera Barat*. BPS Provinsi Sumatera Barat.
- Cohen, J. (1960). A coefficient of agreement for nominal scales. *Educational and Psychological Measurement*, 20(1), 37–46.
- Fey, M., & Lenssen, J. E. (2019). Fast graph representation learning with PyTorch Geometric. *Proceedings of the ICLR Workshop on Representation Learning on Graphs and Manifolds*.
- Hamilton, W. L., Ying, Z., & Leskovec, J. (2017). Inductive representation learning on large graphs. *Advances in Neural Information Processing Systems*, 30, 1024–1034.
- Hidayat, R., & Pratama, D. (2023). Regional food inflation monitoring using NLP-based social media analytics. *Proceedings of the 5th International Conference on Informatics and Computational Information Technology* (pp. 1–6).

- Huang, Z., Xu, G., Shen, H., Cheng, X., & He, F. (2019). Text level graph neural network for text classification. *Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing* (pp. 3442–3448).
- Johnson, W., & Zhang, M. (2020). Temporal graph networks for deep learning on dynamic graphs. *Proceedings of the ICML 2020 Workshop on Graph Representation Learning and Beyond*.
- Kim, S., Park, J., & Lee, M. (2021). Real-time food price monitoring through social media sentiment. *Food Policy*, 102, Article 102103.
- Kipf, T. N., & Welling, M. (2017). Semi-supervised classification with graph convolutional networks. *Proceedings of the 5th International Conference on Learning Representations*.
- Koto, F., Lau, J. H., & Baldwin, T. (2021). IndoLEM and IndoBERT-lite: A lightweight BERT model for Indonesian. *Proceedings of the 2nd Workshop for Southeast Asian NLP* (pp. 65–73).
- Koto, F., Rahimi, A., Lau, J. H., & Baldwin, T. (2020). IndoLEM and IndoBERT: A benchmark dataset and pre-trained language model for Indonesian NLP. *Proceedings of the 28th International Conference on Computational Linguistics* (pp. 757–770).
- Koto, F., & Wahyudi, G. R. (2017). InSet lexicon: Evaluation of a word list for Indonesian sentiment analysis in microblogs. *Proceedings of the International Conference on Advanced Informatics, Concepts, Theory, and Applications (ICAICTA)*.
- Kurniawan, N., Santoso, I., & Prasetyo, H. (2021). Twitter sentiment analysis for monitoring Indonesian commodity prices. *Procedia Computer Science*, 179, 471–479.
- Lin, L., Xu, Y., Lin, Z., & Si, L. (2021). BertGCN: Transductive text classification by combining GNN and BERT. *Findings of the Association for Computational Linguistics: ACL-IJCNLP 2021* (pp. 1456–1462).
- Loshchilov, I., & Hutter, F. (2019). Decoupled weight decay regularization. *Proceedings of the International Conference on Learning Representations*.
- Mukherjee, P., & Liu, B. (2012). Mining contentions from discussions and debates. *Proceedings of the 18th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining* (pp. 841–849).
- Nugraheni, E. H., & Kusumawardani, H. (2014). Indonesian Twitter opinion mining using Naive Bayes. *Proceedings of the International Conference on Information Technology and Electrical Engineering* (pp. 1–5).
- Nugroho, D., Rahmat, F., & Kurniawan, S. (2023). IndoBERT for food import news sentiment classification. *Jurnal RESTI (Rekayasa Sistem dan Teknologi Informasi)*, 7(2), 312–320.
- Permana, R., & Wulandari, S. (2023). Social media as a barometer of food policy sentiment: Evidence from Indonesia. *Journal of Southeast Asian Studies*, 55(1), 78–99.
- Prabowo, A., Hakim, A., & Susanto, B. (2022). Sentiment analysis on Indonesian food commodity news using machine learning. *Journal of Information Systems Engineering and Business Intelligence*, 8(1), 45–54.
- Purwarianti, A., & Crisdayanti, I. (2019). Improving bi-LSTM performance for Indonesian sentiment analysis using paragraph vector. *Proceedings of the 6th International Conference on Information and Communication Technology* (pp. 1–5).
- Putra, A., & Wahyudi, B. (2022). Public sentiment dynamics on Indonesian rice price policy: A Twitter analysis. *Asian Journal of Agriculture and Rural Development*, 12(3), 211–225.
- Putri, M., & Ferdiana, R. (2022). Code-mixing in Indonesian social media: Challenges for NLP. *Proceedings of the 1st Workshop on Southeast Asian NLP* (pp. 23–31).
- Rahmaniah, N., Widodo, A., & Adji, T. B. (2021). Aspect-based sentiment analysis for Indonesian product reviews using BiLSTM-CRF. *Jurnal ILKOM*, 13(2), 90–98.

- Savigny, O. (2021). SentiRadit: Sentiment analysis dataset for Indonesian regional dialects. *Proceedings of the International Conference on Asian Language Processing* (pp. 146–151).
- Simorangkir, A. (2022). Inflation dynamics in Indonesia: The role of food and energy prices. *Bulletin of Monetary Economics and Banking*, 25(2), 185–210.
- Sun, C., Qiu, X., Xu, Y., & Huang, X. (2019). How to fine-tune BERT for text classification? *Proceedings of the CCF International Conference on Natural Language Processing and Chinese Computing* (pp. 194–206).
- Tala, F. Z. (2003). *A study of stemming effects on information retrieval in Bahasa Indonesia* [Master's thesis, University of Amsterdam]. Institute for Logic, Language and Computation.
- Veličković, P., Cucurull, G., Casanova, A., Romero, A., Liò, P., & Bengio, Y. (2018). Graph attention networks. *Proceedings of the 6th International Conference on Learning Representations*.
- Wilie, B., Vincentio, K., Winata, G. I., Cahyawijaya, S., Li, X., Lim, Z. Y., Soleman, S., Mahendra, R., Fung, P., Bahar, S., & Purwarianti, A. (2020). IndoNLU: Benchmark and resources for evaluating Indonesian natural language understanding. *Proceedings of the 1st Conference of the Asia-Pacific Chapter of the ACL* (pp. 843–857).
- Wolf, T., Debut, L., Sanh, V., Chaumond, J., Delangue, C., Moi, A., Cistac, P., Rault, T., Louf, R., Funtowicz, M., & Brew, J. (2020). Transformers: State-of-the-art natural language processing. *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing: System Demonstrations* (pp. 38–45).
- Wu, Z., Pan, S., Chen, F., Long, G., Zhang, C., & Yu, P. S. (2021). A comprehensive study on graph neural networks. *IEEE Transactions on Neural Networks and Learning Systems*, 32(1), 4–24.
- Xu, H., Liu, M., & Wang, Y. (2021). Text classification using multi-layer graph convolutional networks. *Neurocomputing*, 458, 167–177.
- Yao, L., Mao, C., & Luo, Y. (2019). Graph convolutional networks for text classification. *Proceedings of the 33rd AAAI Conference on Artificial Intelligence*, 33, 7370–7377.
- Zhang, M., Chen, Y., & Wang, W. (2022). Named entity recognition for Indonesian using IndoBERT. *Journal of Big Data*, 9, Article 52.