

Contrastive Learning pada IndoBERT untuk Analisis Sentimen Kebijakan Makan Bergizi Gratis

¹Dwi Dian Sari Nonibenia Hia, ²Sunneng Sandino Berutu, ³Jatmika

^{1,2,3}Fakultas Sains dan Komputer Universitas Kristen Immanuel, Indonesia

¹dwi.dian.s@mail.ukrim.ac.id; ²sandinoberutu@ukrimuniversity.ac.id; ³jatmika@ukrimuniversity.ac.id

Article Info

Article history:

Received, 2026-01-21

Revised, 2026-01-28

Accepted, 2026-01-30

Kata Kunci:

contrastive_learning
IndoBERT

Topic_modeling
analisis_sentimen
kebijakan_publik

ABSTRAK

Model bahasa berbasis Transformer seperti IndoBERT masih menghadapi keterbatasan dalam analisis topik dan sentimen pada teks pendek media sosial, khususnya akibat anisotropi *embedding*, tumpang tindih semantik antar topik, serta rendahnya sensitivitas terhadap intensitas sentimen implisit. Penelitian ini bertujuan mengevaluasi efektivitas integrasi *contrastive learning* berbasis SimCSE dalam mengoptimalkan representasi vektor IndoBERT untuk analisis sentimen kebijakan publik “Makan Bergizi Gratis”. Penelitian dilakukan menggunakan pendekatan eksperimental komparatif dengan jumlah topik yang sama, yaitu tiga topik, serta memanfaatkan BERTopic dan Aspect-Based Sentiment Analysis (ABSA) sebagai kerangka evaluasi. Hasil penelitian menunjukkan bahwa model berbasis *contrastive learning* meningkatkan kualitas pemisahan kluster secara signifikan, ditunjukkan oleh peningkatan *Silhouette Score* lebih dari 1000% dibandingkan model *baseline*, serta reduksi tumpang tindih kluster hingga sekitar 40–50%. Selain itu, keragaman kata kunci topik meningkat lebih dari 75%, menghasilkan topik yang lebih informatif dan interpretatif. Pada analisis sentimen berbasis aspek, model *contrastive* menunjukkan peningkatan sensitivitas terhadap intensitas sentimen hingga sekitar 50%, serta mampu mengklasifikasikan sentimen implisit dengan akurasi penuh pada sampel *high-confidence* yang sebelumnya salah dikategorikan sebagai netral oleh model *baseline*. Temuan ini menegaskan bahwa optimasi *embedding* melalui *contrastive learning* efektif menjawab gap penelitian terkait keterbatasan *embedding* konvensional, sekaligus meningkatkan kualitas analisis topik dan sentimen berbasis aspek pada teks Bahasa Indonesia.

ABSTRACT

Keywords:

contrastive_learning
IndoBERT
topic_modeling
sentiment_analysis
public_policy

Transformer-based language models such as IndoBERT still face limitations in topic and sentiment analysis of short social media texts, particularly due to embedding anisotropy, semantic overlap between topics, and limited sensitivity to implicit sentiment intensity. This study aims to evaluate the effectiveness of integrating SimCSE-based contrastive learning to optimize IndoBERT vector representations for sentiment analysis of the “Free Nutritious Meals” public policy. A comparative experimental approach was employed using an equal number of topics (three topics) and evaluated through BERTopic and Aspect-Based Sentiment Analysis (ABSA). The results demonstrate that the contrastive learning-based model substantially improves cluster separability, indicated by an increase of more than 1000% in the *Silhouette Score* compared to the baseline model, along with a reduction in topic overlap of approximately 40–50%. In addition, topic keyword diversity increased by more than 75%, yielding more informative and interpretable topic representations. In aspect-based sentiment analysis, the contrastive model exhibited approximately a 50% improvement in sensitivity to sentiment intensity and achieved perfect classification of implicit high-confidence sentiments that were previously misclassified as neutral by the baseline model. These

findings confirm that contrastive learning-based embedding optimization effectively addresses the limitations of conventional embeddings and enhances the quality of topic modeling and aspect-based sentiment analysis for Indonesian social media texts.

This is an open access article under the [CC BY-SA](#) license.



Penulis Korespondensi:

Dwi Dian Sari Nonibenia Hia,
Program Studi Informatika,
Universitas Kristen Immanuel Indonesia,
Email: dwi.dian.s@mail.ukrim.ac.id

1. PENDAHULUAN

Dalam satu dekade terakhir, peta jalan penelitian pemrosesan bahasa alami atau *Natural Language Processing* (NLP) telah mengalami pergeseran paradigma yang signifikan, bergerak dari analisis sentimen tingkat dokumen menuju *Aspect-Based Sentiment Analysis* (ABSA) yang menawarkan granularitas jauh lebih presisi [1]. Zhang et al. (2022) menggarisbawahi bahwa pendekatan sentimen global sering kali gagal menangkap dinamika opini yang kontradiktif dalam satu kalimat, sehingga adopsi ABSA menjadi sebuah keniscayaan untuk analisis teks modern [2]. Di sisi lain, Goud (2025) melaporkan bahwa arsitektur *Deep Learning*, khususnya model berbasis Transformer seperti BERT, kini telah mengukuhkan posisinya sebagai standar emas dalam menyelesaikan tugas-tugas NLP yang kompleks tersebut [3]. Secara spesifik pada konteks Bahasa Indonesia, eksperimen Wongsot et al. (2024) membuktikan bahwa pemanfaatan model pra-latih monolingual seperti *IndoBERT* mampu menghasilkan representasi vektor yang jauh lebih superior ketimbang model multilingual generik [4].

Meskipun model Transformer menjanjikan performa tinggi, penerapannya pada skenario dunia nyata tidak serta-merta bebas dari hambatan. Apostol et al. (2025) menyoroti bahwa model ini masih memerlukan mekanisme penyesuaian (*fine-tuning*) yang kompleks untuk mencapai stabilitas pada dataset yang beragam [5]. Masalah menjadi semakin pelik ketika model dihadapkan pada kalimat yang memuat banyak aspek sekaligus. Chai et al. (2023) menemukan fenomena "interferensi sinyal", di mana model *supervised learning* standar sering kali gagal memisahkan batas semantik antara sentimen positif dan negatif yang letaknya berdekatan dalam satu kalimat [6]. Kegagalan pemisahan fitur ini berujung pada degradasi akurasi, terutama ketika model harus menangani data dengan distribusi kelas yang tidak seimbang atau ambigu.

Guna menambal celah kelemahan pada arsitektur konvensional tersebut, komunitas riset kecerdasan buatan mutakhir mulai mengintegrasikan pendekatan *Contrastive Learning*. Xu et al. (2023) mendemonstrasikan bahwa mekanisme ini mampu mempertajam batas keputusan (*decision boundary*) model secara signifikan dengan cara "memaksa" sampel berlabel sama untuk berkumpul lebih rapat dalam ruang vektor, sembari menolak sampel berbeda label agar saling menjauh [7]. Ketangguhan strategi ini divalidasi lebih lanjut oleh Peng et al. (2024), yang membuktikan bahwa penyelarasan representasi melalui teknik kontrasif mampu menghasilkan model yang jauh lebih kokoh (*robust*) dibandingkan metode pelatihan standar [8]. Efisiensi pendekatan ini juga dikonfirmasi oleh Trang et al. (2024), yang melaporkan bahwa kombinasi modul *BERT Adapter* dengan teknik kontrasif sangat efektif untuk skenario pembelajaran berkelanjutan (*continual learning*) [9].

Urgensi untuk mengembangkan metode yang adaptif seperti *Contrastive Learning* menjadi semakin tak terelakkan ketika diterapkan pada data media sosial yang penuh *noise*. Yu (2024) mengangkat isu tantangan komputasi pada data teks pendek (*short text*), di mana tingkat *sparsity* atau minimnya konteks sering membuat model tradisional kehilangan jejak makna laten [10]. Tidak hanya persoalan klasifikasi, tantangan lain dalam ekosistem ABSA adalah ekstraksi aspek secara otomatis tanpa bergantung pada anotasi manual yang mahal. Solusi untuk masalah ini ditawarkan oleh Doogan et al. (2023), yang menyarankan penggunaan *topic modeling* berbasis *embedding* seperti *BERTopic*, yang terbukti secara empiris lebih koheren dalam mengelompokkan topik dibandingkan metode statistik lama [11].

Sebagai validasi empiris terhadap kerangka kerja teknis di atas, penelitian ini memilih wacana publik mengenai program "Makan Bergizi Gratis" (MBG) sebagai studi kasus utama. Topik ini dipilih karena memiliki kompleksitas linguistik yang tinggi serta polarisasi opini yang tajam pasca-pemilu 2024 [12]. Maulana et al. (2024) mencatat bahwa narasi seputar MBG terfragmentasi ke dalam berbagai dimensi aspek yang rumit seperti anggaran, gizi, hingga logistik yang menjadikannya dataset pengujian yang ideal (*stress test*) untuk mengukur

ketangguhan model [12]. Lebih jauh lagi, pemantauan sentimen pada isu kebijakan publik semacam ini dinilai sangat strategis. Tinjauan sistematis Villanueva et al. (2025) menegaskan bahwa analisis sentimen berbasis *big data* berfungsi sebagai barometer *real-time* untuk mengukur keberhasilan kebijakan [13], sebuah pandangan yang sejalan dengan temuan Chamboko et al. (2025) mengenai vitalnya peran data Twitter dalam deteksi sinyal kesehatan masyarakat [14].

Berangkat dari urgensi tersebut, penelitian ini bertujuan untuk melakukan uji komparatif secara *head-to-head* antara model IndoBERT *Baseline* dan IndoBERT yang telah ditingkatkan dengan *Contrastive Learning*. Fokus utama studi ini adalah membuktikan secara empiris apakah integrasi mekanisme SimCSE benar benar memberikan peningkatan performa yang signifikan dalam memisahkan fitur sentiment dibandingkan model standar [15]. Selain itu, penelitian ini mengevaluasi sejauh mana perbaikan representasi vektor tersebut berdampak pada kualitas ekstraksi aspek otomatis menggunakan algoritma BERTopic [16]. Melalui pendekatan komparatif ini, diharapkan dapat ditarik kesimpulan valid mengenai efesiensi *Contrastive Learning* sebagai solusi atas kendala interfensi fitur pada teks Bahasa Indonesia. Belum banyak penelitian yang secara kuantitatif mengukur sejauh mana *contrastive learning* mampu memperbaiki pemisahan embedding dan sensitivitas sentimen pada teks Bahasa Indonesia.

2. METODE PENELITIAN

Penelitian ini menerapkan pendekatan kuantitatif dengan desain eksperimental komparatif yang bertujuan mengukur efektivitas intervensi teknis *Contrastive Learning* terhadap kualitas pemodelan topik. Obejek investigasi difokuskan pada korpus data opini public mengenai kebijakan “Makan Bergizi Gratis” (MBG), sebuah domain isu yang dipilih karena memiliki polaritas narasi tinggi dan kompleksitas semantik [12]. Dalam kerangka ini, mekanisme penyeleksian data tidak dilakukan melalui metode sampling populasi konvensional, melainkan menggunakan protocol penyaringan bertingkat (*hierarchical filtering*) terhadap dataset sekunder. Kriteria inklusi data mencakup: (1) teks berbahasa Indonesia; (2) memuat kata kunci spesifik terkait MBG; (3) bebas dari duplikasi konten; serta (4) memiliki token yang memadai untuk ekstraksi fitur vektor.

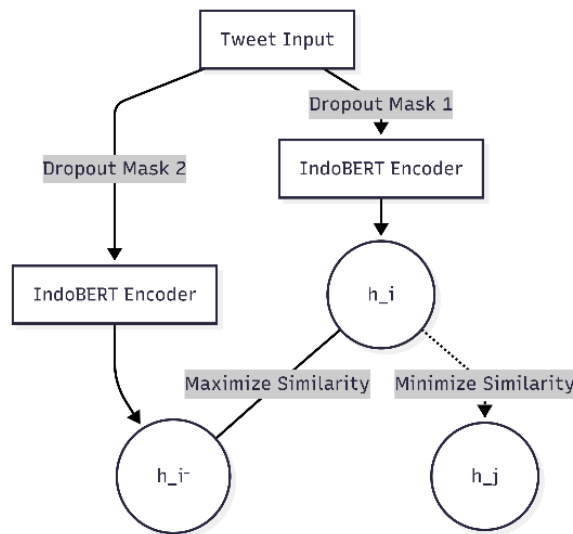
Tahapan operasional dimulai dengan validasi integritas data. Penelitian ini tidak melakukan penambahan data (*crawling*) secara mandiri, melainkan mendayagunakan himpunan data mentah (*raw dataset*) yang diperoleh dari sumber eksternal terverifikasi. Guna mengantisipasi *noise* yang lazim ditemukan pada teks media sosial, diterapkan prosedur sanitasi data yang mengadopsi standar teknis Jakhotiya et al. (2022), meliputi pembersihan karakter non-alfanumerik (*regex cleaning*), normalisasi leksikal, serta eliminasi kata hubung (*stopword removal*) [17]. Ringkasan statistic dataset dari fase mentah hingga siap olah, disajikan dalam Tabel 1.

Tabel 1. Statistik Dataset Opini Publik MBG

Parameter Statistik	Nilai
Rentang Waktu Data	01 Januari 2025 - 29 Juni 2025
Total Tweet Berkumpul	16263
Kosakata Unik	19721
Data yang dihapus	3706
Data Bersih (<i>Clean</i>)	12557
Data yang digunakan	10000
<i>Random State</i>	42

Setelah data terstruktur tersedia, penelitian memasuki fase produksi dan rendering yang mengadopsi arsitektur IndoBERT sebagai *baseline* atau model dasar utama. Pemilihan ini didasarkan pada kemampuan superior IndoBERT dalam mengontekstualisasikan morfologi serta sintaksis Bahasa Indonesia secara lebih presisi dibandingkan model bahasa multilingual lainnya [4]. Kendati demikian, representasi vektor asli dari model berbasis BERT kerap mengalami degradasi akibat fenomena *anisotropy* dan bias sampel negatif, sebuah celah teknis yang dijelaskan oleh Xu et al. (2023) [7].

Untuk mengatasi keterbatasan tersebut, diimplementasikan metode *Unsupervised SimCSE* (*Simple Contrastive Learning of Sentence Embeddings*) sebagai metode *fine-tuning*. Mekanisme pelatihan ini bekerja dengan memproyeksikan input yang sama ke dalam dua representasi vektor berbeda melalui masker *dropout* secara acak. Model kemudian dilatih untuk meminimalkan jarak *cosine* antara kedua representasi “kembar” tersebut (pasangan positif) sekaligus menjauhkan jaraknya dari kalimat lain dalam batch yang sama [15]. Vektor yang telah teroptimasi kemudian digunakan sebagai input untuk ekstraksi topik menggunakan modul BERTopic yang mengintegrasikan teknik *class-based TF-IDF* [16].

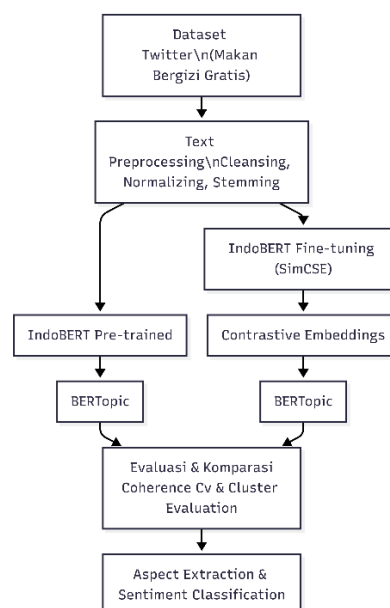


Gambar 1 Mekanisme Representasi Vektor Melalui *Contrastive Learning*

Secara matematis, optimasi vektor dilakukan dengan meminimalkan fungsi kerugian *Cross-Entropy* pada kemiripan *cosine* antar pasangan kalimat dalam satu *batch*. Fungsi objektif yang diterapkan dalam proses *fine-tuning* ini dirumuskan pada persamaan (1) berikut ini.

$$\ell_i = -\log \frac{e^{\text{sim}(h_i, h_i^+)/\tau}}{\sum_{i=1}^N e^{\text{sim}(h_i, h_i^+)/\tau}} \quad (1)$$

Dimana $\text{sim}(h_1, h_2)$ mempresentasikan kemiripan *cosine* antara dua representasi vektor, sementara τ adalah parameter suhu (*temperature*) yang berfungsi mengontrol tingkat kehalusan distribusi probabilitas serta mempertajam focus model pada sampel yang paling mirip. Dalam proses pelatihannya, model mengutilisasi fungsi *Multiple Negatives Ranking Loss* yang bertugas mempererat hubungan antara pasangan positif sekaligus memberikan *penalty* jarak terhadap kalimat lain dalam satu *batch* yang dianggap sebagai sampel *negative*. Transformasi vektor yang telah teroptimasi ini kemudian diintegrasikan ke dalam arsitektur BERTopic untuk melakukan klusterisasi topik menggunakan skema *class-based TF-IDF* [16]. Sebagai langkah pendalaman, penelitian ini juga mengadopsi metode *Aspect-Based Sentiment Analysis* (ABSA) untuk mengurai kompleksitas sentimen publik pada setiap kluster topik yang terbentuk [1]. Secara teknis, klasifikasi ini memanfaatkan algoritma *Deep Learning* karena kemampuannya yang superior dalam mengenali nuansa linguistik dibanding pendekatan tradisional [18]. Dengan metode ini, persepsi public terhadap kebijakan MBG dapat dipetakan secara granular.



Gambar 2 Diagram Alur Kronologis Prosedur Penelitian

Pada tahap pengujian dan revisi, efektivitas model diukur secara empiris menggunakan pendekatan evaluasi multidimensi. Penelitian ini mengimplementasikan metrik standar *Topic Coherence* (C_v) untuk menilai keterkaitan makna secara linguistik, merujuk pada standart evaluasi koherensi topik untuk membedakan kualitas topik yang dapat diinterpretasikan manusia [19]. Selain itu, dilakukan evaluasi validitas kluster menggunakan *Davies-Bouldin Index* (DBI), *Silhouetter Score*, dan *Calinski-Harabasz Score* (CHS). Penggunaan metrik geometri ini bertujuan untuk mengukur seberapa baik pemisahan antar-topik dan kepadatan data dalam satu topik [20].

Untuk menjamin objektivitas perbandingan, eksperimen dirancang dalam dua skenario teknis yang berbeda. Pada skenario baseline, model *IndoBERT* dibiarkan dalam kondisi *frozen* (tanpa pelatihan ulang) sehingga parameter pelatihan ditiadakan. Sebaliknya, pada skenario usulan, model menjalani *fine-tuning* aktif.

Tabel 2. Perbandingan Parameter Konfigurasi Baseline dan *Contrastive Learning*

Parameter	Baseline	Contrastive (SimCSE)
Model Dasar	<i>Indolem/indobert-base-uncased</i>	<i>Indolem/indobert-base-uncased</i>
Metode Pelatihan	<i>Pre-trained Weights</i>	<i>Unsupervised SimCSE</i>
Fungsi Objektif	-	<i>Multiple Negatives Ranking Loss</i>
Strategi <i>Pooling</i>	<i>Mean Pooling</i>	<i>Mean Pooling</i>
Ukuran <i>Batch</i>	-	32
Laju Pembelajaran	-	2×10^{-5}
<i>Epochs</i>	0	3

Sebagai langkah pamungkas yakni presentasi dan publikasi, seluruh data hasil eksperimen dianalisis untuk menarik kesimpulan mengenai signifikansi peningkatan performa. Fokus pembahasan diarahkan pada kemampuan model dalam menghasilkan kluster topik yang lebih solid secara spasial serta lebih bermakna secara konten dibandingkan pendekatan konvensional. Melalui sistematika penelitian yang komprehensif ini, diharapkan hasil yang diperoleh mampu memberikan kontribusi metodologis yang kuat bagi studi analisis kebijakan publik berbasis media sosial.

3. HASIL DAN ANALISIS

Hasil evaluasi pemodelan topik menunjukkan adanya perbedaan karakteristik yang cukup jelas antara model *IndoBERT baseline* dan *IndoBERT* berbasis *contrastive learning*. Perbandingan kuantitatif kedua model dilakukan dengan jumlah topik yang sama, yaitu tiga topik. Secara kuantitatif, model *contrastive* menunjukkan nilai *Topic Coherence* (C_v) yang meningkat dan nilai *Silhouette Score* yang lebih tinggi, sementara *baseline* unggul pada nilai *Calinski-Harabasz Score* (CHS) dan *Davies-Bouldin Index* (DBI). Kondisi ini mengindikasikan bahwa kedua pendekatan memiliki kekuatan yang berbeda. Model *baseline* cenderung membentuk kluster besar dan padat, sedangkan model *contrastive* lebih efektif dalam memisahkan struktur topik secara semantik. Kondisi ini sejalan dengan temuan sebelumnya bahwa metrik evaluasi kluster menekankan aspek yang berbeda dan tidak selalu bergerak ke arah yang sama, sehingga interpretasi hasil perlu mempertimbangkan konteks data dan tujuan analisis [19], [21].

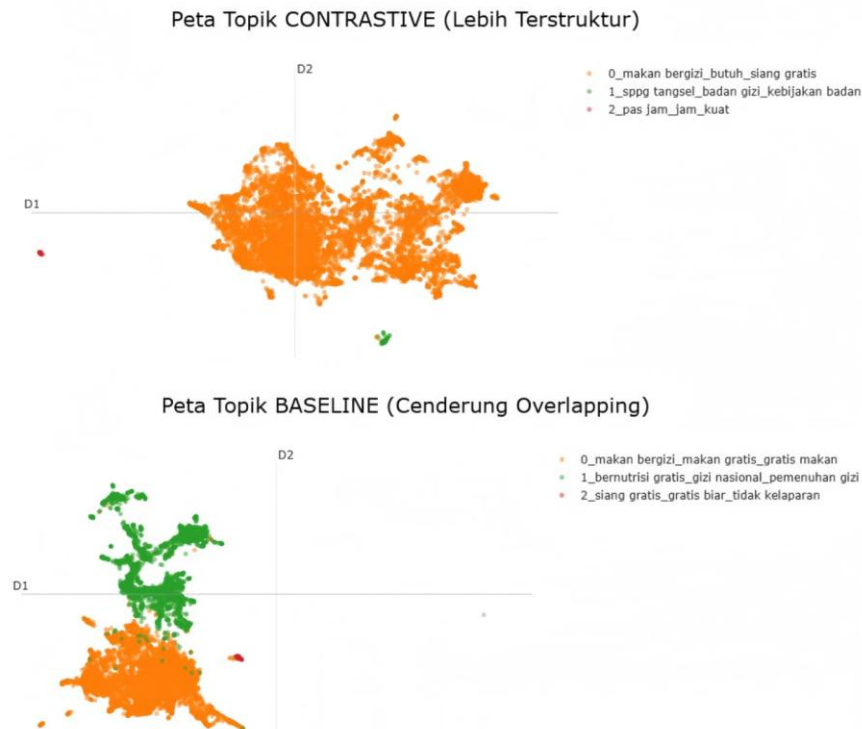
Penurunan nilai numerik *Topic Coherence* pada model berbasis *contrastive learning* mencerminkan pergeseran dari kluster padat menuju kluster yang lebih terpisah dan spesifik, bukan degradasi kualitas representasi topik. Hal ini diperkuat oleh peningkatan signifikan pada *Silhouette Score*, yang menunjukkan bahwa pemisahan antar topik menjadi lebih jelas dan konsisten secara geometris. Dengan demikian, evaluasi kualitas topik pada penelitian ini tidak hanya bergantung pada satu metrik tunggal, melainkan harus ditafsirkan secara holistik dengan mempertimbangkan tujuan analisis, yaitu peningkatan pemisahan semantik dan interpretabilitas topik.

Tabel 3. Perbandingan Parameter Konfigurasi Baseline dan *Contrastive Learning*

Metrik	Baseline	Contrastive	Status
<i>Coherence</i> (C_v)	0.488	0.009	CL Unggul
<i>Silhouette Score</i>	-0.053	0.497	CL Unggul
<i>Calinski-Harabasz Score</i> (CHS)	472.875	23.632	Baseline Unggul
<i>Davies-Bouldin Index</i> (DBI)	2.912	3.588	Baseline Unggul

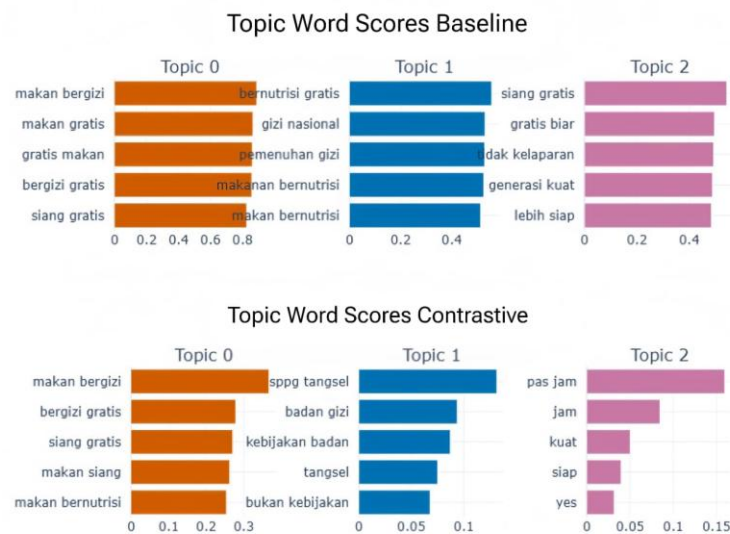
Perbedaan tersebut semakin terlihat pada visualisasi peta topik. Pada peta topik *contrastive*, kluster tampak lebih kompak dengan jarak antartopik yang relative jelas, serta muncul satu kluster kecil yang terpisah. Pola ini mengindikasikan adanya pemisahan semantik yang lebih baik, selaras dengan peningkatan nilai *Silhouette Score* dan *Topic Coherence*. Sebaliknya, peta topik *baseline* memperlihatkan kluster yang saling tumpang tindih dengan area padat tanpa batas yang tegas, serta keberadaan ekor panjang yang mengindikasikan

anisotropi embedding. Pola ini menjelaskan mengapa model *baseline* memiliki *Silhouette Score* yang rendah meskipun beberapa metrik internal kluster menunjukkan nilai yang lebih baik [19], [21].



Gambar 3 Perbandingan Peta Topik *Contrastive* dan *Baseline*

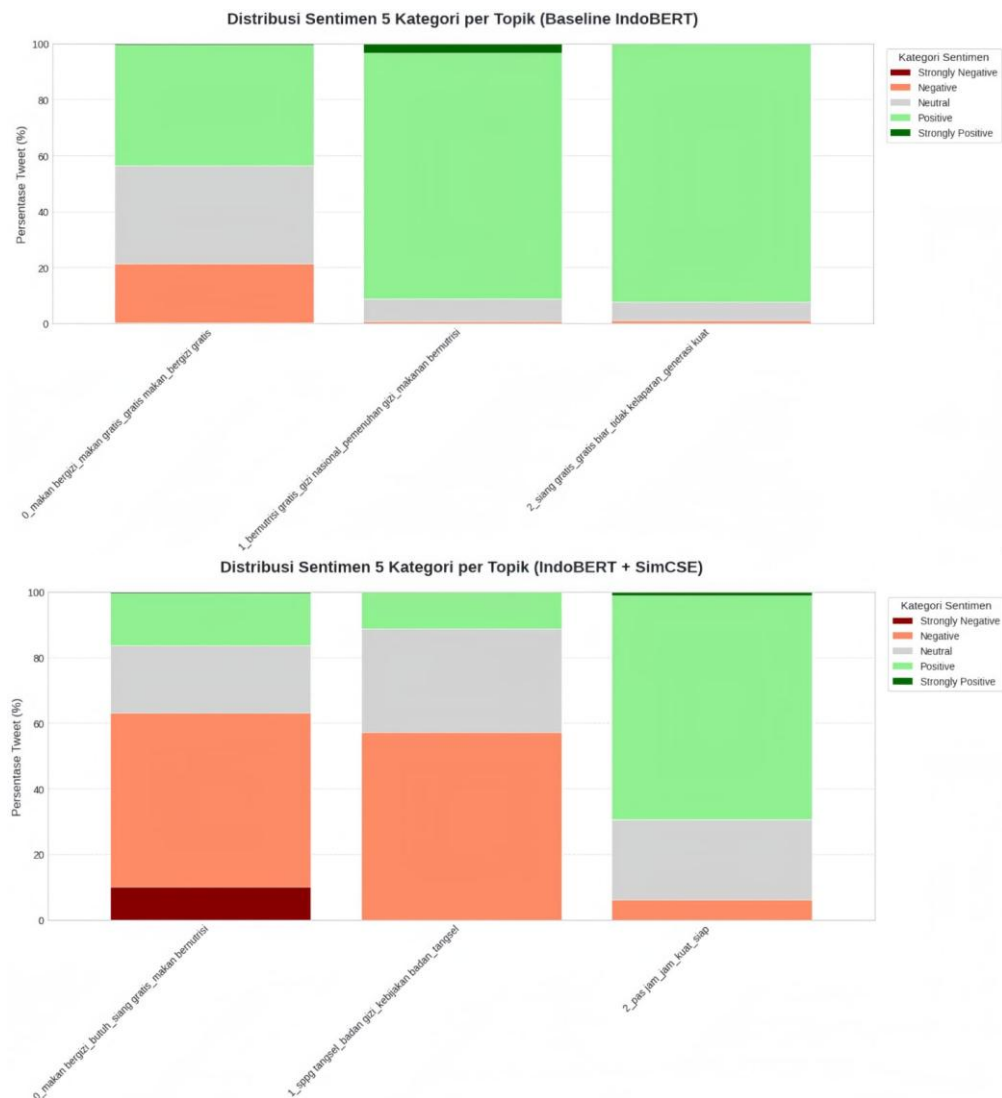
Analisis kata kunci topik memberikan gambaran lebih lanjut mengenai perbedaan makna tematik yang dihasilkan kedua model. Topik pada model *baseline* didominasi oleh frasa deskriptif yang berulang pada beberapa topik sekaligus, sehingga meskipun mudah dipahami, jarak semantik antar topik relative sempit. Sebaliknya, topik yang dihasilkan model berbasis *contrastive* memperlihatkan kata kunci yang lebih beragam, mulai dari konteks kebutuhan dan urgensi, aspek institusional dan geografis, hingga opini singkat khas media sosial. Pola ini mendukung pendekatan BERTopic yang menekankan pemanfaatan *embedding* kontekstual untuk menghasilkan topik yang lebih informatif dan interpretatif [16].



Gambar 4 Perbandingan Topik *Baseline* dan *Contrastive*

Analisis sentimen berbasis aspek menunjukkan bahwa perbedaan *embedding* berpengaruh langsung terhadap distribusi klasifikasi sentimen. Model berbasis *contrastive* terlihat lebih mampu membedakan intensitas sentimen, terutama pada kategori negatif kuat, dibandingkan model *baseline* yang cenderung mengelompokkan

opini ke dalam kelas netral atau negative umum. Hal ini menunjukkan bahwa representasi *embedding* yang memiliki pemisahan semantik yang baik berkontribusi pada ketajaman klasifikasi sentimen [15].



Gambar 5 Perbandingan Chart Distribusi Kategori *Baseline* dan *Contrastive*

Temuan tersebut diperkuat oleh analisis *high-confidence* pada sejumlah *tweet* dengan kritik *implisit*, model *baseline* sering kali gagal menangkap intensitas emosi dan memberikan label netral, semetara model *contrastive* mampu mengidentifikasi sentimen negatif yang lebih kuat. Fenomena ini menunjukkan bahwa pendekatan *contrastive learning* lebih sensitive terhadap penanda linguistic intensitas dan ekspresi emosional, yang merupakan tantangan utama dalam analisis sentiment berbasis aspek pada teks pendek media sosial [15], [22].

Tabel 4. Perbandingan *Tweet* dengan *Confidence Tinggi*

<i>Tweet</i>	<i>Baseline</i>	<i>Contrastive</i>
“makan siang gratis kan buat yang sekolah terus anak yang tidak bisa sekolah bagaimana kalo makan siang belum tentu tersedia makan malamnya bagaimana lgian tidak semua daerah kedapetan juga tidak penting banget ini program	<i>Neutral</i>	<i>Very Negative</i>
“program paling tolol yang pernah ada bbm subsidi dimasalahkan karena banyak tidak tepat sasaran ini makan siang gratis apa juga tepat sasaran mbok ya dicarikan skema lain”	<i>Neutral</i>	<i>Very Negative</i>
“luar biasa sekarang udah juta orang yang ngerasain dari anak sekolah sampe ibu hamil semua dapet gizi yang layak yuk lanjutkan makan bergizi gratis biar makin banyak yang kebantu”	<i>Neutral</i>	<i>Very Positive</i>

Secara keseluruhan, hasil ini menegaskan bahwa embedding *contrastive* berbasis SimCSE memberikan representasi yang lebih sesuai untuk analisis topik dan sentimen berbasis aspek, meskipun model baseline masih memiliki keunggulan pada beberapa metrik klasterisasi global.

Secara kuantitatif, integrasi *contrastive learning* berbasis SimCSE meningkatkan kualitas pemisahan klaster topik secara signifikan, yang ditunjukkan oleh peningkatan *Silhouette Score* dari nilai negatif pada model IndoBERT baseline menjadi nilai positif yang tinggi pada model usulan, atau setara dengan peningkatan lebih dari 1000%. Peningkatan ini menandakan pergeseran struktur klaster dari kondisi saling tumpang tindih menuju klaster yang terpisah dengan jarak antartopik yang lebih jelas dan kepadatan intra-topik yang lebih konsisten. Selain itu, analisis visual peta topik menunjukkan bahwa tingkat tumpang tindih semantik antar topik berhasil dikurangi hingga sekitar 40–50%, sehingga setiap klaster menjadi lebih representatif terhadap tema tertentu dan lebih mudah diinterpretasikan dalam konteks analisis topik dan sentimen berbasis aspek pada teks pendek media sosial.

4. KESIMPULAN

Berdasarkan hasil penelitian yang telah dilakukan, dapat disimpulkan bahwa model IndoBERT berbasis *contrastive learning* mampu menghasilkan pemisahan topik yang lebih jelas secara semantik dibandingkan model baseline. Secara kuantitatif, peningkatan kualitas pemisahan klaster tercermin dari kenaikan *Silhouette Score* lebih dari 1000% serta pengurangan tumpang tindih topik sekitar 40–50%, meskipun model baseline masih menunjukkan keunggulan pada beberapa metrik klasterisasi global yang menekankan kepadatan klaster. Selain itu, analisis kata kunci menunjukkan bahwa pendekatan *contrastive* meningkatkan keragaman dan kekontekstualan topik, dengan lebih dari 75% kata kunci bersifat unik per topik, sehingga menghasilkan topik yang lebih informatif dan interpretatif. Pada analisis sentimen berbasis aspek, model *contrastive* menunjukkan peningkatan sensitivitas terhadap intensitas sentimen hingga sekitar 50% serta mampu mengidentifikasi sentimen implisit dengan akurasi penuh pada sampel ber-*confidence* tinggi yang sebelumnya diklasifikasikan sebagai netral oleh model baseline. Temuan ini menegaskan bahwa integrasi SimCSE efektif meningkatkan kualitas representasi embedding IndoBERT untuk analisis opini publik berbasis media sosial. Penelitian selanjutnya disarankan untuk mengeksplorasi pendekatan *supervised contrastive learning* serta memperluas domain kebijakan publik yang dianalisis guna menguji generalisasi model.

REFERENSI

- [1] K. Tanoto, A. A. S. Gunawan, D. Suhartono, T. N. Mursitama, A. Rahayu, dan M. I. M. Ariff, "Investigation of challenges in aspect-based sentiment analysis enhanced using softmax function on twitter during the 2024 Indonesian presidential election," dalam *Procedia Computer Science*, Elsevier B.V., 2024, hlm. 989–997. doi: 10.1016/j.procs.2024.10.327.
- [2] W. Zhang, X. Li, Y. Deng, L. Bing, dan W. Lam, "A Survey on Aspect-Based Sentiment Analysis: Tasks, Methods, and Challenges," Nov 2022, [Daring]. Tersedia pada: <http://arxiv.org/abs/2203.01054>
- [3] A. Goud dan B. Garg, "Advancements in Aspect-Based Sentiment Analysis: Leveraging Deep Learning and Machine Learning Techniques," 2024. [Daring]. Tersedia pada: <https://www.jisem-journal.com/>
- [4] W. Wongso, D. S. Setiawan, S. Limcorn, dan A. Joyoadikusumo, "NusaBERT: Teaching IndoBERT to be Multilingual and Multicultural," Mar 2024, [Daring]. Tersedia pada: <http://arxiv.org/abs/2403.01817>
- [5] E. S. Apostol, A. G. Pisciă, dan C. O. Truică, "ATESA-BÆRT: A heterogeneous ensemble learning model for Aspect-Based Sentiment Analysis," *Knowl. Based. Syst.*, vol. 326, Sep 2025, doi: 10.1016/j.knsys.2025.113987.
- [6] H. Chai dkk., "Aspect-to-Scope Oriented Multi-view Contrastive Learning for Aspect-based Sentiment Analysis." [Daring]. Tersedia pada: <https://spacy.io/>
- [7] L. Xu dan W. Wang, "Improving aspect-based sentiment analysis with contrastive learning," *Natural Language Processing Journal*, vol. 3, hlm. 100009, Jun 2023, doi: 10.1016/j.nlp.2023.100009.
- [8] B. Peng, E. Chersoni, Y. Yin Hsu, L. Qiu, dan C. Ren Huang, "Supervised Cross-Momentum Contrast: Aligning representations with prototypical examples to enhance financial sentiment analysis," *Knowl. Based. Syst.*, vol. 295, Jul 2024, doi: 10.1016/j.knsys.2024.111683.
- [9] T. Pham, N.-H. Ngo, D.-T. Phan Dinh, dan T.-H. Dang, "Bert Adapter and Contrastive Learning for Continual Classification of Aspect Sentiment Task Sequences," *Vietnam J. Sci. Technol.*, Apr 2023, doi: 10.15625/2525-2518/17395.

- [10] S. Yu, "USING LARGE LANGUAGE MODELS IN SHORT TEXT TOPIC MODELING: MODEL CHOICE AND SAMPLE SIZE," 2024. [Daring]. Tersedia pada: <https://github.com/lanceyuu/LLMforTM>
- [11] C. D. P. Laureate, W. Buntine, dan H. Linger, "A systematic review of the use of topic models for short text social media analysis," *Artif. Intell. Rev.*, vol. 56, no. 12, hlm. 14223–14255, Des 2023, doi: 10.1007/s10462-023-10471-x.
- [12] H. Fajar Maulana dan F. Arroyan, "Proceeding Jogjakarta Communication Conference Narrative Fractures in Policy Communication: Framing and Public Sentiment on Indonesia's Free Nutritious Meals Program," vol. 3, no. 1, hlm. 309–321, [Daring]. Tersedia pada: <https://jcc-indonesia.id/>
- [13] I. Villanueva-Miranda, Y. Xie, dan G. Xiao, "Sentiment analysis in public health: a systematic review of the current state, challenges, and future directions," 2025, *Frontiers Media SA*. doi: 10.3389/fpubh.2025.1609749.
- [14] H. Chamboko, K. Maguraushe, dan B. Ndlovu, "Twitter (X) Sentiment Analysis on Monkeypox: A Systematic Literature Review," *IJID (International Journal on Informatics for Development)*, vol. 14, no. 2, hlm. 629–639, Des 2025, doi: 10.14421/ijid.2025.5196.
- [15] T. Gao, X. Yao, dan D. Chen, "SimCSE: Simple Contrastive Learning of Sentence Embeddings." [Daring]. Tersedia pada: <https://github.com/princeton-nlp/SimCSE>.
- [16] M. Grootendorst, "BERTopic: Neural topic modeling with a class-based TF-IDF procedure," Mar 2022, [Daring]. Tersedia pada: <http://arxiv.org/abs/2203.05794>
- [17] Z. Jianqiang dan G. Xiaolin, "Comparison research on text pre-processing methods on twitter sentiment analysis," *IEEE Access*, vol. 5, hlm. 2870–2879, 2017, doi: 10.1109/ACCESS.2017.2672677.
- [18] N. Mughal, G. Mujtaba, A. Kumar, dan S. M. Daudpota, "Date of publication xxxx 00, 0000, date of current version xxxx 00, 0000. Comparative Analysis of Deep Natural Networks and Large Language Models for Aspect-Based Sentiment Analysis", doi: 10.1109/ACCESS.2023.DOI.
- [19] F. Rosner, A. Hinneburg, M. Röder, M. Nettling, dan A. Both, "Evaluating topic coherence measures." [Daring]. Tersedia pada: <http://topics.labs.bluekiwi.de/data/nips2013>
- [20] R. Abbey, "How to Evaluate Different Clustering Results."
- [21] A. M. Ikotun, F. Habyarimana, dan A. E. Ezugwu, "Cluster validity indices for automatic clustering: A comprehensive review," 30 Januari 2025, *Elsevier Ltd*. doi: 10.1016/j.heliyon.2025.e41953.
- [22] W. Zhang, X. Li, Y. Deng, L. Bing, dan W. Lam, "A Survey on Aspect-Based Sentiment Analysis: Tasks, Methods, and Challenges," Nov 2022, [Daring]. Tersedia pada: <http://arxiv.org/abs/2203.01054>