

Analisis Sentimen Terhadap Program Makan Bergizi Gratis Menggunakan Pendekatan Pseudo-Labeling dan Arsitektur Transformer

Wisnu Aji Sanjaya^{1*}, Yoka Romadani², Jeremy Marcello Waani³, Nisrina Hanifa Setiono⁴

¹Program Studi Sains Data
Telkom Univeristy Purwokerto

^{1*}wisnuajisanjaya@student.telkomuniversity.ac.id

²Program Studi Sains Data
Telkom University Purwokerto

²yokaromadani@student.telkomuniversity.ac.id

³Program Studi Sains Data
Telkom Univeristy Purwokerto

³jeremywaani@student.telkomuniversity.ac.id

⁴Program Studi Sains Data
Telkom Univeristy Purwokerto

⁴nisrinahanifas@telkomuniversity.ac.id

Abstrak— Analisis sentimen terhadap kebijakan publik sering kali terkendala oleh ketimpangan distribusi data (*class imbalance*) yang ekstrem, sehingga model klasifikasi cenderung mengalami bias dan mengabaikan opini minoritas. Penelitian ini bertujuan untuk membangun model klasifikasi sentimen yang mampu mengurangi kecenderungan bias kelas mayoritas pada studi kasus Program Makan Bergizi Gratis di media sosial Twitter. Metodologi yang diusulkan mengintegrasikan teknik *pseudo-labeling* menggunakan IndoRoBERTa (*Teacher Model*) untuk mengekstraksi dan menyaring 18.385 cuitan menjadi korpus *Silver Standard*. Untuk mengatasi ketimpangan kelas tanpa manipulasi data (*oversampling*), penelitian ini menerapkan *Cost-Sensitive Learning* melalui modifikasi bobot *Inverse Class Frequency* pada fungsi *Cross-Entropy Loss* saat melakukan *fine-tuning* pada *Student Model* (IndoBERT dan XLM-RoBERTa). Hasil eksplorasi data menunjukkan kelas Netral mendominasi (50,4%), sementara kelas Negatif menjadi minoritas ekstrem (21,4%) yang memuat kritik rasional terkait beban anggaran. Hasil komputasi membuktikan bahwa arsitektur XLM-RoBERTa menunjukkan kinerja paling optimal dengan Akurasi global 87% dan *Recall* makro 87%, lebih tangguh dibandingkan IndoBERT. Evaluasi *Confusion Matrix* menegaskan bahwa penerapan penalti bobot berhasil menekan rasio kegagalan prediksi (*False Negative*) pada kelas minoritas hingga menyentuh angka 13,47%. Kesimpulannya, kerangka kerja hibrida ini efektif mendisiplinkan jaringan algoritma untuk mengurangi kecenderungan bias kelas mayoritas, sehingga menghasilkan instrumen kecerdasan buatan yang lebih proporsional untuk mendukung evaluasi opini kebijakan publik.

Kata kunci— Analisis Sentimen, Cost-Sensitive Learning, Pseudo-Labeling, Transformer, Kebijakan Publik.

Abstract— Sentiment analysis of public policy is often hampered by extreme class imbalance in data distribution, meaning that classification models tend to be biased and overlook minority opinions. This study aims to develop a sentiment classification model capable of reducing the tendency for majority-class bias in a case study of the Free Nutritious Meals Programme on Twitter. The proposed methodology integrates a pseudo-labelling technique using IndoRoBERTa (Teacher Model) to extract and filter 18,385 tweets into a Silver Standard corpus. To address class imbalance without data manipulation (oversampling), this study applies Cost-Sensitive Learning by modifying the Inverse Class Frequency weights in the Cross-Entropy Loss function during fine-tuning of the Student Model (IndoBERT and XLM-RoBERTa). Data exploration results show that the Neutral class dominates (50.4%), whilst the Negative class is an extreme minority (21.4%), containing rational criticism regarding budgetary burdens. Computational results demonstrate that the XLM-RoBERTa architecture exhibits the most optimal performance with a global Accuracy of 87% and a macro Recall of 87%, outperforming IndoBERT. Confusion matrix evaluation confirms that the application of weight penalties successfully reduced the false negative rate for the minority class to 13.47%. In conclusion, this hybrid framework effectively disciplines the neural network to reduce the tendency towards majority-class bias, thereby producing a more balanced artificial intelligence tool to support the evaluation of public policy opinions.

Keywords— Sentiment Analysis, Cost-Sensitive Learning, Pseudo-Labeling, Transformer, Public Policy.

I. PENDAHULUAN

Media sosial seperti Twitter (kini X) merupakan salah satu sumber data yang paling banyak digunakan dalam penelitian untuk mengekstrak opini publik secara real-time melalui teknik

sentiment analysis dan opinion mining. [1], [2], [3]. Namun, teks yang berasal dari media sosial seperti Twitter memiliki karakteristik high-noise karena mengandung bahasa tidak baku, singkatan, slang, emoji, serta kesalahan ejaan, sehingga memerlukan

tahap preprocessing sebelum dilakukan analisis sentimen. [4], [5], [6]. Tantangan ini semakin rumit pada pemodelan Supervised Learning yang sangat bergantung pada ketersediaan data latih berlabel dalam jumlah besar. Proses anotasi data secara manual yang memakan waktu, membutuhkan biaya yang tinggi, serta rentan terhadap variasi dan subjektivitas antarpemilai (inter-annotator bias) [7], [8], [9].

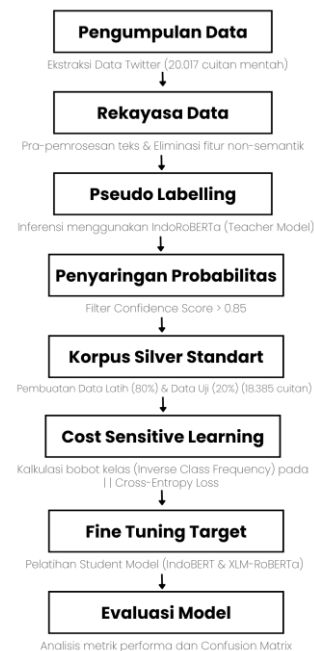
Selain keterbatasan jumlah label, dataset opini publik yang dikumpulkan secara organik dari media sosial umumnya memiliki distribusi kelas yang tidak seimbang (class imbalance), di mana satu atau beberapa kelas mendominasi jumlah sampel dibandingkan kelas lainnya [6], [10]. Dalam analisis sentimen terhadap respons masyarakat atas kebijakan publik, distribusi kelas yang tidak seimbang menyebabkan sentimen dominan lebih banyak muncul dibandingkan sentimen minoritas. Apabila model dilatih menggunakan data yang timpang, model cenderung mengalami *majority class bias*, yaitu kecenderungan memprediksi kelas mayoritas dengan akurasi tinggi, tetapi memiliki kinerja yang rendah dalam mengenali kelas minoritas [11], [12]. Penelitian terdahulu umumnya menyelesaikan ketimpangan ini dengan memanipulasi distribusi data, seperti menggunakan teknik *oversampling* atau *undersampling*. Pendekatan tersebut memiliki kelemahan fatal karena rentan memicu *overfitting*, terutama ketika diterapkan pada korpus data yang penuh anomali leksikal.

Penelitian ini bertujuan untuk membangun model analisis sentimen yang tangguh dan objektif terhadap Program Makan Bergizi Gratis dengan mengatasi hambatan anotasi dan ketimpangan kelas tersebut. Untuk mengatasi keterbatasan label, penelitian ini mengimplementasikan pendekatan pseudo-labeling menggunakan IndoRoBERTa. Model ini dipilih karena memanfaatkan tokenisasi Byte-Level BPE yang mampu merepresentasikan variasi karakter pada teks media sosial, sehingga lebih robust terhadap kata tidak baku, singkatan, dan kesalahan ejaan dibandingkan tokenisasi berbasis kata [13], [14]. Selanjutnya, penelitian ini menggunakan metode *Cost-Sensitive Learning* untuk menangani ketimpangan kelas tanpa merusak distribusi data asli. Pendekatan ini bekerja dengan memodifikasi fungsi objektif algoritma untuk memberikan penalti

matematis yang lebih berat terhadap kesalahan prediksi pada kelas minoritas. Secara spesifik, penelitian ini bermaksud mengevaluasi dan membandingkan performa dua arsitektur *Transformer* target, yakni IndoBERT dan XLM-RoBERTa, dalam mempertahankan akurasi klasifikasi pada kondisi data yang timpang.

II. METODOLOGI PENELITIAN

Penelitian ini menggunakan kerangka kerja terstruktur yang mengintegrasikan teknik pemrosesan bahasa alami (NLP) tingkat lanjut dengan modifikasi pada arsitektur *Transformer*. Secara umum, arsitektur *pipeline* penelitian divisualisasikan pada Gambar 1.



Gambar 1. Alur kerja integrasi pseudo-labeling dan cost-sensitive learning

Tahapan operasional dari kerangka kerja tersebut diuraikan secara rinci pada sub-bab berikut:

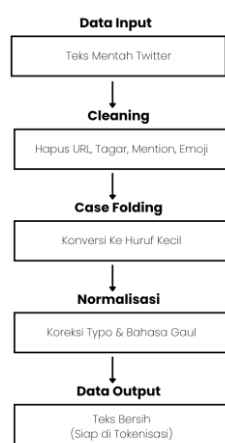
A. Pengumpulan dan Seleksi Fitur Dataset

Dataset yang dieksplorasi dalam penelitian ini bersumber dari cuitan (*tweet*) organik berbahasa Indonesia yang diekstraksi secara spesifik terkait kebijakan Program Makan Bergizi Gratis. Pengumpulan data dilakukan melalui metode penambangan data (*web scraping*) pada rentang periode Desember 2024 hingga Juni 2026 menggunakan parameter kata kunci utama "makan bergizi gratis". Proses ekstraksi menghasilkan 20.017 data mentah yang belum memiliki label

(*unlabeled data*). Data cuitan mentah dari Twitter memiliki 16 atribut bawaan (*metadata*). Untuk mengoptimalkan komputasi memori pada arsitektur *Transformer* dan menghindari masuknya elemen *noise* ke dalam pemodelan sentimen, dilakukan seleksi fitur.

Atribut yang mengukur metrik interaksi (seperti *retweetCount*, *replyCount*, *likeCount*, dan *viewCount*) dieliminasi karena tidak memiliki korelasi sintaksis dengan polaritas kalimat. Metadata identitas logistik pengguna (seperti *id*, *url*, *profilePicture*, dan *source*) juga direduksi. Atribut bahasa (*lang*) dipertahankan secara temporer untuk memfilter cuitan berbahasa Indonesia, menyisakan variabel *text* sebagai input utama atau fitur independen tunggal dalam pemrosesan teks ini.

B. Pra-pemrosesan Teks (Text Preprocessing)



Gambar 2. Tahapan pra-pemrosesan teks (Text Preprocessing)

Karakteristik teks pada media sosial berdimensi *high-noise*, penuh dengan format informal, dan tidak terstruktur. Oleh karena itu, variabel *text* mentah ditransformasikan menjadi teks normal melalui serangkaian teknik pra-pemrosesan. Proses ini mencakup pembersihan karakter non-alfanumerik ekstrem (seperti tautan, sebutan pengguna atau *mention*, dan tagar). Selanjutnya, dilakukan normalisasi bahasa gaul (*slang*) dan kesalahan ketik (*typo*) menggunakan dataset *Colloquial Indonesian Lexicon*. Untuk menjamin efisiensi alur kerja (*pipeline*) pra-pemrosesan data berskala besar, korpus kamus leksikon tersebut dipetakan ke dalam struktur memori berbasis *dictionary*. Pendekatan ini memungkinkan mesin melakukan pencarian dan substitusi kosa kata tidak baku ke dalam bahasa formal secara instan dengan kompleksitas waktu

komputasi konstan $O(1)$. Serangkaian transformasi tekstual ini sangat krusial untuk menekan angka kemunculan token di luar kosakata (*Out-Of-Vocabulary / OOV*) saat teks dipetakan ke dalam ruang vektor oleh algoritma *tokenizer*.

C. Ekstraksi Korpus Silver Standard via Pseudo-Labeling

Mengatasi hambatan ketiadaan data latih berlabel tanpa mengandalkan anotasi manual, penelitian ini mengimplementasikan pendekatan *pseudo-labeling* otomatis berbasis *Teacher Model*. Model pakar yang dipilih adalah IndoRoBERTa (*indonesian-roberta-base-sentiment-classifier*). Pemilihan model ini didasarkan pada ketangguhan arsitektur bawaan *Byte-Level Byte-Pair Encoding* (BPE) yang jauh lebih superior dibandingkan *WordPiece tokenizer* standar dalam menangani anomali dan *noise* karakter pada data Twitter.

Dalam penelitian ini, kualitas pseudo-label dijaga melalui penerapan confidence threshold sebesar 0,85 sehingga hanya prediksi dengan tingkat keyakinan tinggi yang digunakan sebagai korpus Silver Standard. Meskipun demikian, penelitian ini belum melakukan validasi manual terhadap sebagian sampel pseudo-label menggunakan anotator manusia. Oleh karena itu, kualitas label masih bergantung pada performa Teacher Model dan berpotensi mewarisi kesalahan prediksi apabila terjadi kesalahan sistematis pada proses pelabelan otomatis. Cuitan yang memiliki probabilitas keyakinan di bawah ambang batas tersebut dieliminasi, sehingga menghasilkan 18.385 cuitan bersih. Kumpulan data yang telah dianotasi oleh mesin ini kemudian ditetapkan sebagai korpus *Silver Standard* dan dibagi secara acak (*random splitting*) menjadi 80% data latih dan 20% data uji.

D. Kalibrasi Cost-Sensitive Learning

Eksplorasi terhadap korpus Silver Standard menunjukkan adanya ketimpangan distribusi kelas (*class imbalance*). Kelas sentimen Netral mendominasi distribusi sebesar 50,4% (9.262 cuitan), diikuti oleh kelas Positif sebesar 28,2% (5.188 cuitan), sedangkan kelas Negatif memiliki proporsi terendah sebesar 21,4% (3.935 cuitan). Meskipun distribusi tersebut tidak tergolong ekstrem, perbedaan proporsi antarkelas tetap berpotensi memengaruhi proses pembelajaran model dan menyebabkan kecenderungan terhadap kelas

mayoritas. Oleh karena itu, penelitian ini menerapkan pendekatan Cost-Sensitive Learning sebagai upaya untuk meningkatkan sensitivitas model terhadap kelas dengan jumlah sampel yang lebih sedikit tanpa melakukan manipulasi distribusi data melalui teknik oversampling.

Sebagai alternatif yang lebih robust dalam menangani ketimpangan distribusi kelas, penelitian ini mengimplementasikan pendekatan Cost-Sensitive Learning melalui modifikasi fungsi kerugian (*loss function*). Pendekatan ini memberikan bobot yang berbeda pada setiap kelas berdasarkan metode *Inverse Class Frequency*, sehingga kelas dengan jumlah sampel yang lebih sedikit memperoleh bobot yang lebih besar dibandingkan kelas mayoritas. Selanjutnya, bobot tersebut diintegrasikan secara langsung ke dalam fungsi *Cross-Entropy Loss* untuk meningkatkan sensitivitas model terhadap kelas minoritas selama proses pelatihan. Formulasi matematis dari fungsi *loss* yang telah dimodifikasi disajikan sebagai berikut.

$$\mathcal{L} = - \sum_{i=1}^c w_i y_i \log(\hat{y}_i) \quad (1)$$

Di mana $\mathcal{L}$ adalah nilai kerugian, C adalah jumlah kelas sentimen, y_i adalah label kelas aktual, $\hat{y}_i$ adalah probabilitas prediksi dari model, dan w_i mewakili bobot penalti untuk kelas ke- i . Melalui formulasi ini, kesalahan prediksi pada sentimen minoritas (Negatif) akan dikenakan penalti matematis yang jauh lebih besar, mendisiplinkan jaringan saraf (*neural network*) untuk lebih peka dalam mengekstraksi fitur semantik minoritas tanpa perlu mendistorsi jumlah data aslinya.

E. Skenario Fine-Tuning Arsitektur Target

Transfer pengetahuan dari korpus dilakukan melalui proses *fine-tuning* pada dua arsitektur *Student Model* target IndoBERT (merepresentasikan pemahaman monolingual) dan XLM-RoBERTa (merepresentasikan ruang representasi multilingual berskala besar). Untuk mencegah terjadinya degradasi gradien (*gradient vanishing*) dan fenomena kegagalan mengingat pola (*catastrophic forgetting*), hiperparameter komputasi dikonfigurasi secara identik untuk kedua model.

Batas panjang tensor input (*max sequence length*) dikalibrasi pada 128 token, yang merupakan ukuran

optimal untuk mengakomodasi batasan karakter Twitter tanpa membuang sumber daya komputasi pada *padding*. Kapasitas sampel per iterasi propagasi balik (*batch size*) diatur sebesar 16, dengan laju pembelajaran moderat (*learning rate*) 2×10^{-5} . Pelatihan dibatasi hingga 3 siklus (*epoch*) untuk mencegah *overfitting*, dan dioptimasi menggunakan algoritma *AdamW* yang memiliki kapabilitas regularisasi parameter (penurunan bobot atau *weight decay*) secara dinamis.

F. Metrik Evaluasi Kinerja

Luaran prediksi model target dievaluasi menggunakan instrumen *Confusion Matrix* untuk memetakan keberhasilan True Positive, True Negative, False Positive, dan False Negative secara granular. Instrumen ini krusial untuk memvalidasi seberapa jauh angka kebocoran prediksi pada kelas minoritas setelah diinjeksi *Cost-Sensitive Learning*. Selanjutnya, komputasi metrik evaluasi diformulasikan ke dalam empat parameter standar pengukuran Akurasi, Presisi, Recall, dan F1-Score.

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN}$$

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

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

$$F1-Score = 2 \times \frac{Precision \times Recall}{Precision + Recall}$$

III. HASIL DAN PEMBAHASAN

A. Analisis Eksplorasi Data (EDA)

Tahap pertama dalam penelitian ini adalah menganalisis karakteristik data pada korpus Silver Standard hasil Pseudo Labelling. Dari total 18.385 cuitan yang lolos seleksi probabilitas, distribusi kelas sentimennya disajikan pada Tabel 1.

Tabel 1. Distribusi Kelas Sentimen pada Korpus Silver Standard

Kelas Sentimen	Jumlah Cuitan	Persentase
----------------	---------------	------------

Netral	9.262	50,4%
Positif	5.188	28,2%
Negatif	3.935	21,4%

Tabel 1 menunjukkan adanya ketimpangan distribusi kelas (*class imbalance*). Kelas sentimen Netral memiliki proporsi terbesar sebesar 50,4%, sedangkan kelas sentimen Negatif memiliki proporsi paling rendah sebesar 21,4%. Meskipun distribusi tersebut tidak tergolong ekstrem, perbedaan proporsi antarkelas tetap berpotensi memengaruhi proses pembelajaran model dan menyebabkan kecenderungan terhadap kelas mayoritas. Oleh karena itu, penelitian ini menerapkan *Cost-Sensitive Learning* sebagai salah satu pendekatan untuk meningkatkan sensitivitas model terhadap kelas dengan jumlah sampel yang lebih sedikit tanpa memanipulasi distribusi data asli.

Selain menganalisis distribusi jumlah, penelitian ini juga menganalisis isi teks (leksikon) pada kelas minoritas (Negatif) menggunakan visualisasi *Word Cloud*.



Gambar 3. Word Cloud representasi leksikal pada kelas sentimen Negatif.

Berdasarkan Gambar 3, opini negatif masyarakat sama sekali tidak didominasi oleh makian kasar atau ujaran kebencian. Kata yang paling menonjol di luar entitas nama program itu sendiri adalah "anggaran", "pemerintah", "rakyat", dan "negara".

Selain itu, tingginya frekuensi kata negasi dan pertentangan seperti "tapi", "tidak", dan "bukan" menunjukkan penolakan bersyarat atau keraguan publik. Munculnya kata spesifik seperti "korupsi", "proyek", dan "pendidikan" di dalam kluster ini mengonfirmasi bahwa sentimen negatif berakar pada skeptisisme struktural. Publik mengkhawatirkan transparansi alokasi "anggaran", potensi "korupsi", serta urgensi program ini dibandingkan kebutuhan dasar "pendidikan".

Temuan tersebut menunjukkan bahwa keberadaan sentimen minoritas memiliki nilai strategis dalam

proses evaluasi kebijakan publik karena dapat merepresentasikan peringatan dini terhadap potensi permasalahan atau penolakan masyarakat yang belum tercermin pada sentimen dominan. Oleh karena itu, model analisis sentimen perlu mampu mengidentifikasi opini minoritas secara akurat dan tidak terpengaruh oleh ketimpangan distribusi data. Kondisi ini semakin menegaskan pentingnya penerapan arsitektur pendeteksi yang robust terhadap permasalahan *class imbalance* sehingga mampu memberikan representasi opini publik yang lebih komprehensif.

B. Komparasi Kinerja Model

Tahap kedua mengevaluasi hasil pelatihan pada dua algoritma target, yaitu IndoBERT dan XLM-RoBERTa. Pengujian ini tidak sekadar mengejar akurasi, tetapi membuktikan model mana yang lebih kebal terhadap bias setelah diberikan penalti bobot kelas (*Cost-Sensitive Learning*). Hasil evaluasi kedua model disajikan pada Tabel 2.

Tabel 2. Komparasi Kinerja IndoBERT dan XLM-RoBERTa

Model	Metrik Evaluasi			
	Akurasi	Presisi	Recall	F1-Score
IndoBERT	87%	86%	86%	86%
XLM-RoBERTa	87%	86%	87%	86%

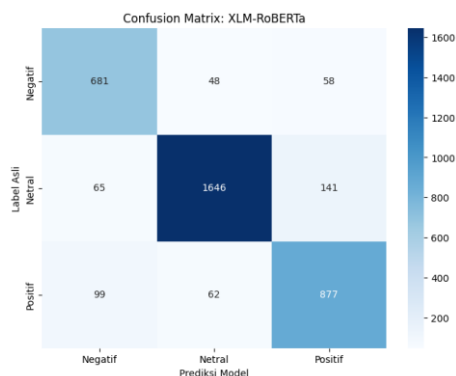
Berdasarkan Tabel 2, kedua model mampu mencapai Akurasi global (87%) dan F1-Score makro (86%) yang identik. Tingginya akurasi ini membuktikan bahwa arsitektur *Transformer* dapat belajar dengan optimal tanpa perlu memanipulasi jumlah data asli (seperti *oversampling*).

Namun, indikator utama keberhasilan pada data yang timpang bukanlah akurasi, melainkan metrik *Recall*. Pada indikator ini, XLM-RoBERTa terbukti lebih tangguh dengan skor *Recall* makro sebesar 87%, unggul satu persen atas IndoBERT (86%). Keunggulan ini menegaskan bahwa arsitektur representasi multilingual (XLM-RoBERTa) memiliki kepekaan yang lebih tajam dalam mengekstraksi dan mengenali pola bahasa pada kelas minoritas (Negatif) dibandingkan model monolingual biasa.

C. Analisis Confusion Matrix

Tahap terakhir adalah melacak letak kesalahan prediksi model XLM-RoBERTa menggunakan matriks kebingungan (*Confusion Matrix*). Evaluasi

ini membuktikan secara langsung apakah algoritma murni mengenali pola bahasa, atau sekadar menebak kelas mayoritas secara membabi buta demi mengejar skor akurasi global.



Gambar 4. Confusion Matrix pada evaluasi model XLM-RoBERTa.

Hasil pada Gambar 4 membuktikan bahwa metode *Cost-Sensitive Learning* bekerja dengan sangat efektif dalam mencegah model mengabaikan kelas minoritas. Secara spesifik, rasio kegagalan algoritma dalam mendeteksi sentimen Negatif (*False Negative Rate*) berhasil ditekan pada angka 13,47%. Perhitungan ini diperoleh berdasarkan total data aktual kelas Negatif yang berjumlah 787 data. Dari populasi tersebut, model gagal memprediksi 106 data secara tepat (*False Negative*), di mana 48 data meleset diprediksi sebagai Netral dan 58 data meleset diprediksi sebagai Positif. Nilai persentase 13,47% tersebut merupakan hasil pembagian antara total kesalahan prediksi (106) terhadap total data aktual kelas Negatif (787).

D. Diskusi dan Implikasi Temuan

Eksperimen ini memberikan pembuktian yang solid, baik dari sisi ketangguhan arsitektur komputasi maupun urgensinya terhadap evaluasi kebijakan publik. Secara teknis, hasil pengujian menegaskan bahwa masalah ketimpangan data ekstrem tidak harus diselesaikan dengan memanipulasi distribusi dataset asli. Implementasi penalti matematis melalui *Cost-Sensitive Learning* terbukti jauh lebih aman dari risiko *overfitting* dibandingkan teknik *oversampling*, sekaligus sukses memperbaiki model untuk tetap mengenali kelas minoritas secara presisi.

Oleh karena itu, integrasi *pseudo-labeling* dan *Cost-Sensitive Learning* pada model XLM-

RoBERTa tidak bisa dipandang sekadar sebagai pencapaian metrik komputasi. Kerangka kerja hibrida ini membuktikan dirinya sebagai instrumen analitik yang objektif, kebal terhadap bias data, dan sangat reliabel untuk mengaudit respons publik terhadap program pemerintah secara proporsional. Meskipun demikian, pendekatan *pseudo-labeling* memiliki keterbatasan karena seluruh label pada korpus Silver Standard dihasilkan secara otomatis oleh Teacher Model. Kondisi ini berpotensi menimbulkan error propagation atau confirmation bias apabila Teacher Model menghasilkan kesalahan prediksi secara konsisten. Oleh sebab itu, penelitian selanjutnya disarankan melakukan validasi manual terhadap sebagian sampel data atau melibatkan beberapa anotator manusia untuk mengevaluasi kualitas *pseudo-label* yang dihasilkan.

IV. KESIMPULAN

Penelitian ini berhasil membuktikan bahwa ketimpangan distribusi kelas pada analisis opini publik dapat ditangani melalui penerapan pendekatan *pseudo-labeling* dan *Cost-Sensitive Learning* tanpa perlu melakukan manipulasi terhadap distribusi dataset asli. Integrasi metode *pseudo-labeling* dan *Cost-Sensitive Learning* terbukti efektif dalam mendisiplinkan jaringan algoritma. Berdasarkan hasil evaluasi komparatif, arsitektur XLM-RoBERTa menunjukkan keunggulan representasi multilingual yang lebih presisi dibandingkan IndoBERT. Model tersebut mampu mencatatkan skor *Recall* makro sebesar 87% dan sukses menekan rasio kegagalan prediksi pada kelas minoritas (*False Negative*) hingga menyentuh angka 13,47%. Pencapaian teknis ini menunjukkan bahwa pendekatan yang diusulkan mampu mengurangi kecenderungan bias terhadap kelas mayoritas sekaligus meningkatkan kemampuan model dalam mengenali kelas dengan jumlah sampel yang lebih sedikit.

Dalam penerapannya, pendekatan algoritma tersebut memiliki peran penting dalam mendukung evaluasi kebijakan publik yang lebih akurat dan berimbang. Kemampuan model dalam mendeteksi kelas sentimen minoritas secara presisi memastikan bahwa opini kritis rasional dari masyarakat seperti skeptisisme terhadap alokasi anggaran dan potensi

utang negara pada Program Makan Bergizi Gratis tidak terabaikan oleh sistem komputasi. Kerangka kerja analitik yang dihasilkan melalui penelitian ini menawarkan instrumen kecerdasan buatan bagi pemerintah untuk melakukan audit opini publik secara proporsional, transparan, dan tidak terdistorsi oleh dominasi suara mayoritas.

Penelitian ini masih memiliki keterbatasan karena korpus Silver Standard sepenuhnya diperoleh melalui proses pseudo-labeling tanpa validasi manual terhadap sebagian sampel data. Oleh karena itu, penelitian selanjutnya perlu mengombinasikan pseudo-labeling dengan anotasi manusia untuk mengevaluasi kualitas label dan meminimalkan potensi confirmation bias.

Sebagai langkah pengembangan di masa mendatang, arsitektur yang telah divalidasi ini dapat diekspansi untuk mengakomodasi analisis yang lebih kompleks. Penelitian selanjutnya direkomendasikan untuk menerapkan kerangka kerja ini pada pemrosesan aliran data secara langsung (*real-time streaming*) guna memantau pergeseran opini publik dari waktu ke waktu. Selain itu, pengembangan menuju Klasifikasi Sentimen Berbasis Aspek (*Aspect-Based Sentiment Analysis*) sangat disarankan agar sistem tidak hanya mampu mendeteksi polaritas emosi, tetapi juga memetakan secara spesifik aspek kebijakan mana yang paling banyak disorot atau dikritik oleh masyarakat.

UCAPAN TERIMA KASIH

Penulis mengucapkan terima kasih kepada Program Studi Sains Data, Telkom University Purwokerto, atas dukungan dan fasilitas yang telah diberikan selama proses penyusunan penelitian ini. Penulis juga menyampaikan apresiasi kepada dosen pembimbing serta seluruh pihak yang telah memberikan arahan, masukan, dan dukungan dalam penyelesaian penelitian berjudul “Analisis Sentimen Terhadap Program Makan Bergizi Gratis Menggunakan Pendekatan Pseudo-Labeling dan Arsitektur Transformer”. Semoga hasil penelitian ini dapat memberikan kontribusi secara praktis dan akademis dalam pengembangan sistem intelijen kebijakan publik serta pemantauan opini masyarakat berbasis kecerdasan buatan.

REFERENSI

- [1] X. Wang, “Topic Modeling in Finance : A Review of Methods , Applications , and Challenges,” *J. Risk Financ. Manag.*, vol. 19, no. 6, pp. 1–43, 2026, doi: 10.3390/jrfm19060399.
- [2] C. Shen, “Social media marketing and digital influence for visitor flow management in sustainable heritage tourism,” *Sci. Rep.*, vol. 15, pp. 1–14, 2025, doi: 10.1038/s41598-025-28555-9.
- [3] E. Fundisi and Q. P. Dlamini, “Ten Years on : A Revisit on the # FeesMustFall Movement Discourse on Twitter / X,” *Journal. Media*, vol. 7, no. 2, pp. 1–12, 2026, doi: 10.3390/journalmedia7020098.
- [4] M. A. Palomino, “Evaluating the Effectiveness of Text Pre-Processing in Sentiment Analysis,” *Appl. Sci.*, vol. 12, no. 17, p. 8765, 2022, doi: 10.3390/app12178765.
- [5] M. Ahmad, M. Waqas, A. Hamza, S. Usman, I. Batyrshin, and G. Sidorov, “UA-HSD-2025 : Multi-Lingual Hate Speech Detection from Tweets Using Pre-Trained Transformers,” *Computers*, vol. 14, pp. 1–24, 2025, doi: 10.3390/computers14060239.
- [6] A. Albladi, M. Islam, and C. Seals, “Sentiment Analysis of Twitter Data Using NLP Models : A Comprehensive Review,” *IEEE Access*, vol. 13, no. January, pp. 30444–30468, 2025, doi: 10.1109/ACCESS.2025.3541494.
- [7] T. Murugan *et al.*, “Research Advances in Maize Crop Disease Detection Using Machine Learning and Deep Learning Approaches,” *Computers*, vol. 15, no. 2, pp. 1–53, 2026, doi: 10.3390/computers15020099.
- [8] H. Abdul, A. Ariffin, and J. Singh, “Artificial Intelligence and Machine Learning solutions-based oral cancer screening and detection : A scoping review,” *Artif. Intell. Med.*, vol. 180, no. April 2025, p. 103480, 2026, doi: 10.1016/j.artmed.2026.103480.
- [9] I. Spiero, A. M. Leeuwenberg, J. Kamperman, L. Hooft, K. G. M. Moons, and J. A. A. Damen, “Using large language models for automated assessment of reporting quality and completeness of prediction model studies,” *J. Clin. Epidemiol.*, vol. 196, p. 112320, 2026, doi: 10.1016/j.jclinepi.2026.112320.
- [10] H. Iwida, “Integrating Large Language Models and Graph Neural Networks for enhanced Arabic sentiment analysis,” *Discov Comput.*, vol. 29, no. 376, 2026, doi: 10.1007/s10791-026-10211-z.
- [11] N. Al-milli, “Emotion detection and sentiment analysis of textual data,” *Egypt. Informatics J.*, vol. 34, no. April, p. 100976, 2026, doi: 10.1016/j.eij.2026.100976.
- [12] H. Jakha, S. Tbaikhi, S. El Houssaini, M. El Houssaini, and S. Ajjaj, “A Hybrid SBERT – WGAN Framework with Ensemble Learning for Sentiment Analysis in Imbalanced Datasets,” *Appl. Syst. Innov.*, vol. 9, no. 5, pp. 1–22, 2026, doi: 10.3390/asi9050103.
- [13] S. Liu, W. Kong, Z. Liu, and J. Sun, “Dual-view cross attention enhanced semi-supervised learning method for discourse cognitive engagement classification in online course discussions,” *Expert Syst. Appl.*, vol. 278, no. March, p. 127339, 2025, doi: 10.1016/j.eswa.2025.127339.
- [14] X. Jing, Z. Wu, and D. Mu, “Enhancing text classification with neural label embedding and weakly-supervised learning,” *Expert Syst. Appl.*, vol. 298, no. PB, p. 129569, 2026, doi: 10.1016/j.eswa.2025.129569.