

# ANALISIS SENTIMEN ISU ANCAMAN SIBER MENGGUNAKAN ALGORITMA MULTI-LAYER PERCEPTRON

Riezky Fauji Setiawan<sup>1</sup>, Chaerur Rozikin<sup>2</sup>, Arip Solehudin<sup>3</sup>

<sup>1,2,3</sup>Program Studi Teknik Informatika Fakultas Ilmu Komputer Universitas Singaperbangsa Karawang; Jl. HS.Ronggo Waluyo, Puseurjaya, Telukjambe Timur, Karawang, Jawa Barat 41361; (0267) 641177

## Keywords:

Analisis Sentimen, Ancaman Siber, Multi-Layer Perceptron, BERT, Pemrosesan Bahasa Alami (NLP)

## Correspondent Email:

riezkyfaujisetiawan@gmail.com



JITET is licensed under a Creative Commons Attribution-NonCommercial 4.0 International License

**Abstrak.** Ancaman siber di media sosial, khususnya di platform Twitter, semakin meningkat seiring dengan maraknya konten bernuansa ujaran kebencian, hinaan, dan serangan verbal yang dapat merugikan individu maupun kelompok. Mengingat tingginya volume data dan kecepatan pertumbuhan konten, deteksi manual menjadi tidak lagi efektif. Penelitian ini mengusulkan pendekatan berbasis pembelajaran mesin dengan fokus pada penggunaan metode Term Frequency-Inverse Document Frequency (TF-IDF) sebagai teknik utama representasi fitur teks. Data dikumpulkan melalui scraping menggunakan kata kunci terkait isu keamanan siber, kemudian dilakukan proses preprocessing untuk membersihkan dan menyiapkan data teks. Pelabelan sentimen dilakukan secara otomatis menggunakan model BERT Indonesia dari Hugging Face. Setelah data diolah, representasi TF-IDF diterapkan untuk mengubah teks menjadi fitur numerik yang kemudian digunakan sebagai input ke dalam model Multi-Layer Perceptron (MLP). Evaluasi performa dilakukan menggunakan metrik akurasi, presisi, recall, dan F1-score. Hasil evaluasi menunjukkan bahwa kombinasi TF-IDF dan MLP mampu mencapai akurasi klasifikasi sentimen sebesar 73.25%. Temuan ini mengindikasikan bahwa penggunaan TF-IDF sebagai teknik representasi fitur dalam analisis sentimen memiliki efektivitas yang tinggi, dan berpotensi dikembangkan lebih lanjut sebagai bagian dari sistem deteksi otomatis terhadap ancaman siber di media sosial.

**Abstract.** Cyber threats on social media, particularly on the Twitter platform, are increasingly prevalent due to the surge in content containing hate speech, insults, and verbal attacks that can potentially harm individuals and groups. Manual detection of such threats has become ineffective because of the high volume and rapid growth of content. This research proposes a machine learning-based approach focusing on the use of Term Frequency-Inverse Document Frequency (TF-IDF) for text feature representation in automated sentiment analysis of tweets related to cybersecurity issues. Data was collected through a scraping process using relevant keywords, followed by thorough preprocessing to clean the text and improve data quality. Sentiment labeling was performed automatically using a pre-trained Indonesian BERT model from Hugging Face. TF-IDF was utilized to convert the cleaned textual data into numerical features, which were then used as input to a Multi-Layer Perceptron (MLP) model. The performance of the model was evaluated using standard metrics such as accuracy, precision, recall, and F1-score. The results show that the TF-IDF-based approach combined with the MLP model achieved a sentiment classification accuracy of approximately 73.25%. These findings highlight the effectiveness of TF-IDF as a feature extraction

*technique and its potential to support the development of reliable and automated systems for detecting cybersecurity-related sentiment on social media platforms. identifying cyber threats on social media.*

---

## 1. PENDAHULUAN

Pesatnya pertumbuhan media sosial telah mendorong perubahan signifikan dalam cara masyarakat menyampaikan opini. Twitter, sebagai salah satu platform paling aktif, kerap digunakan untuk menyampaikan kritik, saran, maupun ujaran kebencian. Di tengah dinamika tersebut, maraknya ancaman siber yang muncul melalui konten berbahasa kasar menjadi perhatian serius karena dapat memicu konflik sosial hingga mengganggu stabilitas keamanan digital [1].

Kehadiran opini publik dalam bentuk kritik dan saran sejatinya dapat menjadi bahan evaluasi yang bermanfaat. Namun, keberadaan ujaran kebencian atau hujatan tanpa substansi sering kali mengaburkan pesan penting, sehingga menyulitkan pemangku kebijakan dalam mengambil keputusan yang berbasis data [2]. Di sisi lain, jumlah data yang terus meningkat di media sosial menjadikan deteksi manual terhadap ancaman semacam ini tidak lagi efektif. Oleh karena itu, diperlukan pendekatan otomatis berbasis pemrosesan bahasa alami (NLP) untuk melakukan analisis sentimen secara efisien dan akurat [3]. (NLP) atau pengolahan bahasa alami merupakan cabang dari kecerdasan buatan yang berhubungan dengan pemrosesan bahasa alami dan interpretasi komputer [12].

Penelitian ini mengembangkan sistem klasifikasi sentimen pada konten Twitter berbahasa Indonesia yang berkaitan dengan isu ancaman siber. Algoritma Multi-Layer Perceptron (MLP) dipilih karena kemampuannya dalam memodelkan hubungan non-linear antar fitur teks [4]. Untuk mendukung pelabelan data, digunakan pendekatan berbasis BERT dari model `ayameRushia/bert-base-indonesian-1.5G-sentiment-analysis-smsa`, yang mampu mengenali konteks semantik dalam bahasa Indonesia secara lebih akurat [5]. Selanjutnya, proses optimasi parameter dilakukan menggunakan Grid Search Cross-Validation

untuk meningkatkan performa model. Proses analisis mengikuti tahapan Knowledge Discovery in Databases (KDD) guna memastikan data yang digunakan relevan dan bersih, serta hasil klasifikasi dapat diandalkan [6].

Dengan pendekatan ini, diharapkan sistem yang dikembangkan mampu membedakan antara sentimen positif seperti kritik atau saran, dan sentimen negatif seperti ujaran kebencian, sehingga dapat mendukung upaya pencegahan penyebaran ancaman siber melalui media sosial.

## 2. TINJAUAN PUSTAKA

Penelitian mengenai analisis sentimen telah banyak dilakukan pada berbagai topik sosial dan isu kontemporer dengan pendekatan machine learning, khususnya menggunakan metode Multi-Layer Perceptron (MLP). Hendrawati dan Yanti [7] melakukan analisis sentimen terhadap cuitan pengguna Twitter mengenai pandemi Covid-19. Penelitian ini menggunakan metode MLP dengan algoritma backpropagation yang dioptimasi menggunakan Adam Optimizer. Model yang dibangun berhasil mencapai akurasi sebesar 70%, dengan nilai F1-score untuk kelas positif sebesar 0.77, negatif 0.75, dan netral 0.50. Hasil ini menunjukkan efektivitas metode MLP dalam mengklasifikasikan sentimen, meskipun masih terdapat tantangan dalam meningkatkan performa klasifikasi untuk kelas netral.

Behl et al. [8] mengusulkan sistem klasifikasi berbasis MLP untuk menganalisis sentimen tweet yang berkaitan dengan bencana alam dan pandemi Covid-19, guna mendukung upaya penanggulangan bencana. Model yang dikembangkan mencapai akurasi 83% dalam mengklasifikasikan tweet ke dalam tiga kategori, yaitu "resource needs", "resource availability", dan "others". Studi ini menekankan pentingnya penggunaan embedding berbasis domain dan fitur linguistik dalam meningkatkan akurasi model klasifikasi sentimen untuk aplikasi dunia nyata, khususnya dalam konteks manajemen bencana.

Dalam ranah sosial-politik, Sasmita et al. [9] menganalisis opini publik terhadap kontroversi putusan Mahkamah Konstitusi mengenai usia calon presiden dan wakil presiden. Penelitian ini menggunakan teknik SMOTE untuk mengatasi ketidakseimbangan data sentimen, serta metode MLP untuk klasifikasi. Hasil penelitian menunjukkan bahwa penggunaan TF-IDF dan SMOTE secara signifikan meningkatkan akurasi model hingga mencapai 99%, serta mengurangi masalah overfitting pada kelas mayoritas. Studi ini membuktikan efektivitas teknik balancing data dalam meningkatkan generalisasi model klasifikasi.

Fadillah et al. [10] menerapkan metode MLP dalam menganalisis sentimen terhadap ulasan restoran menggunakan dataset dari Kaggle. Penelitian ini menggabungkan teknik Word2Vec sebagai representasi vektor kata untuk meningkatkan pemahaman semantik teks. Evaluasi dilakukan terhadap berbagai konfigurasi arsitektur MLP, termasuk fungsi aktivasi dan learning rate. Model terbaik yang dihasilkan memiliki akurasi 97.39%, membuktikan kemampuan MLP dalam memproses data ulasan pelanggan untuk klasifikasi sentimen secara efektif.

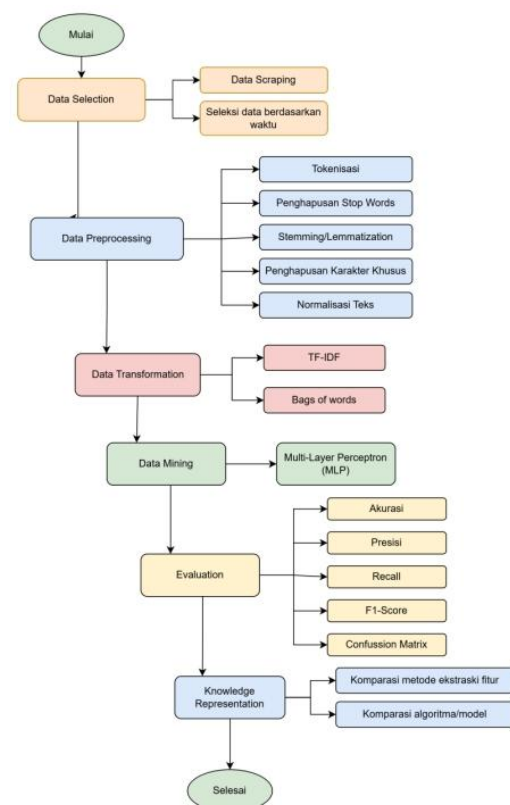
Sementara itu, Makmur et al. [11] mengkaji sentimen masyarakat terhadap isu penghapusan skripsi sebagai tugas akhir mahasiswa dengan mengumpulkan data dari media sosial TikTok dan YouTube. Data dianalisis menggunakan TF-IDF dan diklasifikasikan menggunakan MLP. Hasilnya, kombinasi preprocessing berupa case folding dan stemming, serta proporsi data latih 90% menghasilkan akurasi 94%, precision 96%, recall 93%, dan F1-score 94%. Penelitian ini menunjukkan potensi MLP dalam menangani opini publik pada isu-isu pendidikan tinggi.

Secara umum, berbagai studi di atas memperlihatkan bahwa metode MLP merupakan pendekatan yang fleksibel dan efektif dalam klasifikasi sentimen pada berbagai konteks. Peningkatan performa dapat dicapai melalui pemilihan fitur yang tepat, teknik balancing data, serta pemrosesan teks yang baik seperti penggunaan TF-IDF.

### 3. METODE PENELITIAN

Dalam penelitian ini menggunakan metodologi penelitian Knowledge Discovery in Database (KDD) karena penelitian ini hanya

berfokus pada modelling dan data mining. KDD adalah suatu proses sistematis untuk mengekstraksi pengetahuan yang bermakna dari kumpulan data dalam jumlah besar. Proses ini tidak hanya mencakup analisis statistik atau pelatihan model, tetapi juga mencakup serangkaian tahapan mulai dari pemilihan data, pembersihan, transformasi, eksplorasi pola melalui metode data mining, hingga penyajian pengetahuan (knowledge representation) secara informatif dan dapat dimanfaatkan untuk pengambilan keputusan. Dengan demikian, KDD merupakan pendekatan end-to-end dalam pengolahan data yang memungkinkan transformasi dari data mentah menjadi pengetahuan yang dapat diinterpretasikan. Pada gambar 1 merupakan gambaran terkait metodologi penelitian KDD.



**Gambar 1.** Metodologi penelitian KDD pada artikel ini

Untuk memahami lebih lanjut tahapan KDD ini, berikut merupakan poin penjelasan tahapan per tahapan dalam metodologi ini:

#### 3.1. Data Selection

Data dikumpulkan dari platform media sosial Twitter melalui teknik web scraping, dengan rentang waktu pengambilan data dari Agustus 2024 hingga Januari 2025. Hanya tweet berbahasa Indonesia yang memuat kata kunci terkait kebocoran data dan ancaman siber yang disertakan dalam dataset. Atribut yang dikumpulkan meliputi isi tweet, tanggal dan waktu publikasi, jumlah retweet dan likes, serta metadata akun (tanpa data pribadi untuk menjaga privasi).

### **3.2. Data Preprocessing**

Tahap ini bertujuan untuk membersihkan dan menyiapkan data sebelum dianalisis. Proses yang dilakukan meliputi tokenisasi, penghapusan stop words, stemming, penghapusan karakter khusus dan simbol non-teks, serta normalisasi kata tidak baku. Hasil dari preprocessing adalah korpus teks bersih yang siap diproses lebih lanjut.

### **3.3. Data Transformation**

Tweet yang telah dibersihkan kemudian diberi label sentimen menggunakan model pra-latih BERT berbahasa Indonesia, yaitu *ayameRushia/bert-base-indonesian-1.5G-sentiment-analysis-smsa*. Sentimen diklasifikasikan ke dalam dua kelas: positif dan negatif. Setelah pelabelan, teks dikonversi ke dalam representasi numerik menggunakan dua teknik: TF-IDF dan word embedding (misalnya Word2Vec atau FastText), guna mempersiapkan input untuk proses klasifikasi.

### **3.4. Data Mining**

Proses klasifikasi dilakukan menggunakan algoritma Multi-Layer Perceptron (MLP). Arsitektur MLP terdiri atas tiga komponen utama: input layer, satu atau lebih hidden layer, dan output layer. Fungsi aktivasi ReLU digunakan pada hidden layer, sedangkan fungsi softmax digunakan pada output layer untuk klasifikasi biner. Model dilatih menggunakan Adam Optimizer dan dropout regularization untuk mencegah overfitting. Selain itu, hyperparameter tuning dilakukan menggunakan grid search dengan cross-validation untuk menentukan konfigurasi terbaik dari jumlah neuron, lapisan tersembunyi, dan learning rate.

### **3.5. Evaluation**

Evaluasi model dilakukan menggunakan metrik akurasi, presisi, recall, dan F1-score. Selain itu, digunakan confusion matrix untuk mengidentifikasi distribusi kesalahan prediksi antar kelas. Validasi silang juga diterapkan untuk memastikan generalisasi model terhadap data baru. Evaluasi menyeluruh ini bertujuan untuk menilai efektivitas model dalam mengklasifikasikan sentimen publik terhadap isu kebocoran data.

### **3.6. Knowledge Representation**

Tahap Knowledge Representation dalam kerangka KDD bertujuan untuk menyajikan hasil temuan dari proses klasifikasi sentimen ke dalam bentuk yang mudah dipahami dan diinterpretasikan oleh pihak terkait, seperti pembuat kebijakan, analis keamanan siber, maupun masyarakat umum. Pada tahap ini, hasil klasifikasi ditampilkan dalam bentuk visualisasi data seperti bar chart, pie chart, dan confusion matrix yang menggambarkan distribusi sentimen positif dan negatif terhadap topik kebocoran data.

Visualisasi ini dilengkapi dengan informasi temporal (misalnya tren sentimen bulanan) untuk menunjukkan dinamika opini publik dari waktu ke waktu. Selain itu, word cloud juga digunakan untuk merepresentasikan kata-kata yang paling sering muncul pada masing-masing kategori sentimen, sehingga memudahkan dalam mengidentifikasi isu yang paling banyak diperbincangkan oleh masyarakat.

Hasil representasi ini tidak hanya menjadi bukti validasi terhadap efektivitas model klasifikasi yang dikembangkan, tetapi juga sebagai bentuk transformasi data ke dalam wawasan (insight) yang actionable. Dengan demikian, knowledge representation menjadi bagian krusial dalam menghubungkan hasil teknis dari proses data mining dengan nilai praktis dan strategis dalam konteks pengambilan keputusan terhadap isu kebocoran data dan ancaman siber.

## **4. HASIL DAN PEMBAHASAN**

Proses pemilihan data pada penelitian ini dilakukan secara strategis untuk memastikan relevansi dan kelengkapan informasi terkait isu ancaman keamanan siber di Indonesia. Data dikumpulkan menggunakan alat bantu tweet-harvest yang dikembangkan oleh Helmi Satria,

dengan kata kunci seperti “Ancaman Keamanan Siber”, “Kebocoran Data”, dan “Serangan Cyber”. Pemilihan kata kunci ini dirancang agar mencakup berbagai bentuk ekspresi publik terkait isu keamanan digital. Rentang waktu pengambilan data dari September 2024 hingga Februari 2025 memungkinkan peneliti menangkap dinamika sentimen dalam periode yang cukup panjang. Hasil scraping menghasilkan total 1.437 tweet dalam format mentah, dengan atribut `full_text` dipilih sebagai komponen utama dalam proses analisis sentimen karena mengandung konten teks yang relevan dan lengkap.

Data mentah hasil scraping memerlukan pembersihan sebelum dapat digunakan dalam proses analisis yang valid. Tahapan preprocessing mencakup penghapusan komponen yang tidak relevan seperti URL, mention, dan hashtag yang tidak memberikan kontribusi terhadap makna sentimen. Selanjutnya, dilakukan konversi emoji ke dalam bentuk teks guna mempertahankan ekspresi emosi dalam format yang dapat dibaca oleh sistem. Kata-kata umum atau stopwords dalam bahasa Inggris dan Indonesia juga dihapus untuk mengurangi noise. Tahap akhir preprocessing mencakup proses stemming dan lemmatization untuk mengembalikan kata ke bentuk dasarnya, sehingga model dapat mengenali variasi kata sebagai satu entitas semantik. Proses preprocessing ini menghasilkan representasi teks yang lebih bersih, ringkas, dan siap untuk diekstraksi fiturnya. Contoh hasil preprocessing adalah sebagai berikut

Teks asli : *“belakangan ini jd sadar banyaaak skills yg belakangan ini dicari banget & gaada salahnya buat belajar ke sana. salah satunya cybersecurity soalnya sekarang semua orang organisasi perusahaan butuh keamanan digital. <https://t.co/giyUFfYxI3>”*

Teks hasil preprocessing : *“jd sadar banyaaak skill yg dicari banget & gaada salahnya belajar . salah satunya cybersecur orang organisasi perusahaan butuh keamanan digit .”*

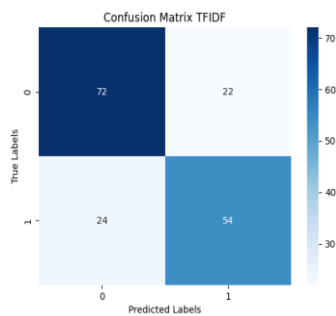
Setelah data preprocessing, dilakukan labelling menggunakan model BERT dari `ayameRushia/bert-base-indonesian-1.5G-sentiment-analysis-smsa`, yang mampu mengenali konteks semantik dalam bahasa Indonesia secara lebih akurat [5] dan akan

menghasilkan label “positif”, “negatif” dan “netral”. Namun, dalam penelitian ini hanya akan menggunakan label “positif” dan “negatif” sehingga label “netral” bisa diseleksi dan dihapus dari dataset.

Transformasi data dilakukan untuk mengubah teks yang telah dibersihkan menjadi format numerik yang dapat diproses oleh algoritma pembelajaran mesin. pendekatan representasi fitur yang digunakan adalah Term Frequency-Inverse Document Frequency (TF-IDF). metode ini mengkonversi teks menjadi vektor dengan struktur yang mencerminkan keberadaan atau pentingnya kata dalam dokumen. Untuk menyederhanakan proses klasifikasi, hanya dua label digunakan: positif (1) dan negatif (0), sementara label netral disingkirkan dari dataset. Transformasi ini memungkinkan teks untuk direpresentasikan dalam bentuk angka dan klasifikasi menjadi biner, yang lebih mudah untuk dianalisis oleh model klasifikasi.

Pada tahap ini, dilakukan proses pelatihan model menggunakan algoritma Multi-Layer Perceptron (MLP), yang dikenal memiliki kemampuan dalam menangani data non-linear. Parameter yang digunakan pada MLP meliputi fungsi aktivasi tanh, struktur hidden layer berukuran (100, 50), serta metode optimasi sgd (stochastic gradient descent). Dataset dibagi menjadi data latih dan data uji untuk menghindari overfitting dan memastikan kemampuan generalisasi model terhadap data baru. Fungsi aktivasi tanh dipilih karena memiliki output zero-centered yang mempermudah pembelajaran, sementara penggunaan sgd memungkinkan pembaruan bobot yang efisien dan cepat meskipun memerlukan perhatian lebih pada pengaturan learning rate.

Evaluasi performa model dilakukan dengan menggunakan metrik confusion matrix, akurasi, precision, recall, dan f1-score. Pada gambar 2, merupakan confusion matrix dari hasil evaluasi yang dilakukan.



**Gambar 2.** Hasil confusion matrix dari metode transformasi data TF-IDF

Dari hasil confusion matrix diatas, kita dapat menghitung beberapa metrics seperti akurasi, presisi, recall, dan F1-Score. Karena kedua nya memiliki hasil confusion matrix yang sama, maka hasil perhitungannya pun sama. Untuk hasil perhitungannya dapat dilihat pada tabel 1.

**Tabel 1.** Hasil perhitungan metric dari MLP dengan metode TFIDF

| Metric    | 0                  | 1    | accuracy | macro avg | weighted avg |
|-----------|--------------------|------|----------|-----------|--------------|
| precision | 0.75               | 0.71 | 0.73     | 0.73      | 0.73         |
| recall    | 0.77               | 0.69 |          | 0.73      | 0.73         |
| f1-score  | 0.76               | 0.70 |          | 0.73      | 0.73         |
| support   | 94                 | 78   | 172      | 172       | 172          |
| Accuracy  | 0.7325581395348837 |      |          |           |              |

untuk algoritma MLP dengan ekstraksi fitur TF-IDF dan split training 80:20, kita dapat menganalisis performa model secara rinci. Confusion matrix menunjukkan bahwa model berhasil mengidentifikasi 72 sampel sebagai True Negatives (kelas 0 yang diprediksi benar sebagai kelas 0) dan 54 sampel sebagai True Positives (kelas 1 yang diprediksi benar sebagai kelas 1). Namun, terdapat 22 sampel yang merupakan False Positives (kelas 0 diprediksi salah sebagai kelas 1) dan 24 sampel yang merupakan False Negatives (kelas 1 diprediksi salah sebagai kelas 0). Total sampel yang dievaluasi adalah 172, yang terbagi menjadi 94 sampel aktual dari kelas 0 dan 78 sampel aktual dari kelas 1.

merinci metrik evaluasi yang lebih dalam. Nilai precision untuk kelas 0 adalah 0.75, menunjukkan bahwa dari semua sampel yang diprediksi sebagai kelas 0, 75% di antaranya memang benar kelas 0. Untuk kelas 1, precision adalah 0.71, yang berarti 71% dari semua sampel yang diprediksi sebagai kelas 1 memang benar kelas 1. Macro average precision dan weighted average precision yang keduanya 0.73 mengindikasikan konsistensi precision di seluruh kelas, baik dengan atau tanpa mempertimbangkan ketidakseimbangan kelas.

Pada metrik recall, model menunjukkan performa yang sedikit berbeda antara kelas. Untuk kelas 0, recall mencapai 0.77, yang berarti 77% dari total sampel aktual kelas 0 berhasil diidentifikasi dengan benar oleh model. Namun, untuk kelas 1, recall sedikit lebih rendah yaitu 0.69, menandakan bahwa hanya 69% dari total sampel aktual kelas 1 yang berhasil dideteksi. Perbedaan ini tercermin dalam macro average recall sebesar 0.73 dan weighted average recall sebesar 0.73, yang menunjukkan bahwa model secara keseluruhan cukup baik dalam menangkap instansi positif, meskipun ada ruang untuk peningkatan dalam mengidentifikasi kelas 1.

Terakhir, metrik f1-score dan accuracy memberikan gambaran komprehensif tentang keseimbangan antara precision dan recall, serta performa keseluruhan model. F1-score untuk kelas 0 adalah 0.76 dan untuk kelas 1 adalah 0.70, menunjukkan keseimbangan yang lebih baik untuk kelas 0. Macro average f1-score dan weighted average f1-score masing-masing adalah 0.73. Akurasi model, yang dihitung sebagai  $(54+72)/172$ , adalah sekitar 0.7325. Nilai ini juga secara eksplisit tertera di tabel, menegaskan bahwa model secara keseluruhan mampu melakukan prediksi yang benar pada sekitar 73.25% dari total sampel. Dengan demikian, meskipun ada sedikit disparitas dalam kemampuan recall antar kelas, algoritma MLP dengan ekstraksi fitur TF-IDF menunjukkan performa yang memadai secara keseluruhan..

## 5. KESIMPULAN

- Untuk mengklasifikasikan sentimen ancaman siber pada tweet, digunakan pendekatan pembelajaran mesin dengan algoritma Multi-Layer Perceptron

- (MLP). Data tweet diproses menggunakan teknik ekstraksi fitur seperti TF-IDF sebelum dimasukkan ke dalam model MLP.
- b. Algoritma MLP menunjukkan performa yang menjanjikan dalam menganalisis sentimen pada dataset tweet. Terbukti dari akurasi yang relatif tinggi, yaitu sekitar 0.7325 dengan ekstraksi fitur TF-IDF. Akurasi ini menunjukkan bahwa MLP memiliki kemampuan yang baik dalam mengenali pola sentimen terhadap ancaman siber dalam teks berbahasa alami, serta konsisten dalam mengolah kedua bentuk representasi fitur tersebut.
  - c. Berdasarkan penelitian ini, menunjukkan bahwa mayoritas tweet yang mengandung kata-kata terkait ancaman siber memiliki sentimen negatif, yang berarti publik cenderung menyuarakan kekhawatiran, ketakutan, atau pandangan buruk terhadap isu-isu terkait keamanan siber di Twitter.

#### UCAPAN TERIMA KASIH

Penulis mengucapkan terima kasih kepada pihak-pihak terkait yang telah memberi dukungan terhadap penelitian ini.

#### DAFTAR PUSTAKA

- [1] "Social Media – Hate Speech – Hate Crime", *Connections: The Quarterly Journal*, Volume 20, Issue 2, pp. 57-73, 2021.
- [2] "The Use of Twitter in Social Media Marketing: Evidence from Hotels in Asia", *Asia Pacific Management and Business Application*, Vol. 11, No. 2, August 2022.
- [3] "Hate Speech Detection in Indonesian Twitter using Contextual Neural Language Models", *Jurnal Ilmu Komputer dan Sistem Informasi*, Universitas Gadjah Mada, 2023.
- [4] "Machine Learning Model for Language Classification: Bag-of-words and Multilayer Perceptron", *Journal of Information Technology and Engineering*, OJS UMA, 2023.
- [5] ayameRushia, "bert-base-indonesian-1.5G-sentiment-analysis-smsa", Hugging Face, 2023. Model fine-tuned dari caha/bert-base-indonesian-1.5G pada dataset indonlu.
- [6] U. Fayyad, G. Piatetsky-Shapiro, and P. Smyth, "Knowledge Discovery and Data Mining: Towards a Unifying Framework", AAAI, 1996.
- [7] T. Hendrawati and C. P. Yanti, "Analysis of Twitter users sentiment against the Covid-19 outbreak using the backpropagation method with Adam optimization," *Journal of Electrical, Electronics and Informatics*, vol. 5, no. 1, pp. 1–4, 2021.
- [8] S. Behl, A. Rao, S. Aggarwal, S. Chadha, and H. S. Pannu, "Twitter for disaster relief through sentiment analysis for COVID-19 and natural hazard crises," *International Journal of Disaster Risk Reduction*, vol. 55, p. 102101, 2021. doi: 10.1016/j.ijdr.2021.102101
- [9] S. Sasmita, R. N. J. S. Intam, D. F. Surianto, and M. F. B. Fajar, "Analisis Sentimen Terhadap Kontroversi Putusan MK Mengenai Usia Capres-Cawapres Menggunakan Multi-Layer Perceptron Dengan Teknik SMOTE," *Faktor Exacta*, vol. 17, no. 2, pp. 188, 2024. doi: 10.30998/faktorexacta.v17i2.22442
- [10] S. Fadillah, W. Witanti, and E. Ramadhan, "Analisis Sentimen Terhadap Ulasan Restoran Menggunakan Metode Multi Layer Perceptron," *Seminar Nasional Corisindo*, 2024.
- [11] H. Makmur, W. Wulandari, D. F. Surianto, and M. F. B. Fajar, "Analisis Sentimen Penghapusan Skripsi sebagai Tugas Akhir Mahasiswa Menggunakan Metode Multi-Layer Perceptron," *Komputika: Jurnal Sistem Komputer*, vol. 13, no. 2, pp. 213–221, 2024. doi: 10.34010/komputika.v13i2.12402
- [12] H. Rahmawati and A. Sudrajat, "Implementasi chatbot pada penerimaan mahasiswa baru di Politeknik TEDC Bandung menggunakan natural language processing," *JITET (Jurnal Informatika dan Teknik Elektro Terapan)*, vol. 13, no. 1, pp. 1–13, Jan. 2025, doi: 10.23960/jitet.v13i1.5456.