

Klasifikasi Serangan Siber menggunakan *Multilayer Perceptron*

Cyber Attack Classification using a Multilayer Perceptron

¹Inna T. Kalakmabin*, ²Irwan Sembiring

^{1,2}Program Studi Teknik Informatika, Fakultas Teknologi Informasi, Universitas Kristen Satya
Wacana, Salatiga, Indonesia

^{1,2}Jl. Diponegoro No.52-60, Salatiga, Jawa Tengah

*e-mail: 672021403@student.uksw.edu, irwan@uksw.edu

(received: 4 July 2026 2026, revised: 19 August 2026, accepted: 20 August 2026)

Abstrak

Serangan *Distributed Denial of Service* (DDoS) yang terus berkembang menjadi ancaman serius bagi keamanan jaringan. Banyak penelitian mengimplementasikan *machine learning* dan *deep learning* untuk deteksi DDoS, namun sebagian besar berfokus pada peningkatan akurasi keseluruhan tanpa mengatasi masalah *class imbalance* yang menjadi karakteristik utama dataset CICDDoS2019. Eksplorasi fungsi kerugian *Focal Loss* pada arsitektur *Multilayer Perceptron* (MLP) untuk klasifikasi multi-kelas juga masih terbatas. Penelitian ini bertujuan mengembangkan model klasifikasi serangan siber berbasis MLP dengan *Focal Loss* untuk mengatasi ketidakseimbangan kelas tersebut. Proses penelitian meliputi pra-proses data, normalisasi fitur, pembobotan kelas, serta pelatihan model dengan dua *hidden layer* menggunakan aktivasi ReLU dan *output Softmax*. Evaluasi dilakukan menggunakan metrik akurasi, presisi, recall, F1-score, dan analisis Confusion Matrix. Hasil eksperimen menunjukkan model MLP mencapai akurasi sebesar 98,80 % dengan *weighted F1-score* sebesar 0,99 dan *macro F1-score* sebesar 0,87. Meskipun *Random Forest* memperoleh akurasi numerik lebih tinggi (99,18 %), model MLP menunjukkan performa kompetitif dan stabil dalam menangani data tidak seimbang. Penelitian ini berkontribusi dalam menunjukkan indikasi awal potensi model MLP dengan *Focal Loss* untuk pengembangan Intrusion Detection System (IDS) berbasis deep learning meskipun memerlukan pengujian terisolasi lebih lanjut.

Kata kunci: CICDDoS2019, DDoS, deep learning, intrusion detection system, multilayer perceptron.

Abstract

The continued evolution of Distributed Denial-of-Service (DDoS) attacks poses a serious threat to network security. Numerous studies have employed machine learning and deep learning techniques for DDoS detection; however, most have focused on improving overall accuracy without adequately addressing the class imbalance characteristic of the CICDDoS2019 dataset. Moreover, the exploration of the Focal Loss function within a Multilayer Perceptron (MLP) architecture for multiclass classification remains limited. This study aims to develop an MLP-based cyberattack classification model using Focal Loss to address class imbalance. The research process includes data preprocessing, feature normalization, class weighting, and model training using two hidden layers with ReLU activation and a Softmax output layer. The model was evaluated using accuracy, precision, recall, F1-score, and confusion matrix analysis. The experimental results show that the MLP model achieved an accuracy of 98.80%, with a weighted F1-score of 0.99 and a macro F1-score of 0.87. Although Random Forest achieved a higher overall accuracy of 99.18%, the MLP model demonstrated competitive and stable performance in handling imbalanced data. This study provides preliminary evidence of the potential of an MLP model with Focal Loss for the development of a deep learning-based Intrusion Detection System (IDS), although further isolated testing is required.

Keywords: CICDDoS2019, DDoS, deep learning, intrusion detection system, multilayer perceptron

1 Pendahuluan

Perkembangan teknologi informasi dan komunikasi telah mendorong peningkatan penggunaan jaringan komputer dalam berbagai sektor, seperti pemerintahan, pendidikan, kesehatan, perbankan, dan industri. Ketergantungan yang semakin tinggi terhadap layanan digital menjadikan keamanan jaringan sebagai salah satu aspek yang sangat penting dalam menjaga keberlangsungan operasional suatu sistem. Di sisi lain, perkembangan teknologi juga diikuti oleh meningkatnya jumlah dan kompleksitas serangan siber yang menargetkan infrastruktur jaringan maupun layanan berbasis internet.

Salah satu bentuk serangan siber yang paling umum dan berdampak signifikan adalah Distributed Denial of Service (DDoS). Serangan DDoS dilakukan dengan membanjiri target menggunakan sejumlah besar trafik jaringan sehingga sumber daya sistem tidak mampu melayani permintaan pengguna yang sah. Serangan DDoS modern terus berkembang baik dari sisi volume, variasi metode serangan, pemanfaatan botnet seperti Mirai [1], pemanfaatan perangkat IoT [2], maupun kompleksitas lalu lintas data pada infrastruktur jaringan [3]. Akibatnya, layanan menjadi lambat, tidak stabil, bahkan tidak dapat diakses sama sekali. Peningkatan frekuensi serangan DDoS mendorong kebutuhan akan mekanisme deteksi yang mampu mengidentifikasi pola serangan secara cepat dan akurat. Salah satu pendekatan yang banyak digunakan adalah Intrusion Detection System (IDS) berbasis Machine Learning dan Deep Learning. Pendekatan ini memungkinkan sistem mempelajari karakteristik lalu lintas jaringan dan melakukan klasifikasi secara otomatis terhadap aktivitas normal maupun aktivitas yang mengindikasikan serangan [4].

Berbagai penelitian telah dilakukan untuk mengembangkan model klasifikasi serangan siber menggunakan algoritma machine learning maupun deep learning. Ahmed dkk. menunjukkan bahwa Multilayer Perceptron (MLP) mampu memberikan performa yang baik dalam mendeteksi serangan DDoS melalui kemampuannya mempelajari hubungan non-linear antar fitur jaringan [1]. Penelitian lain yang dilakukan oleh Mehmood dkk. juga mengombinasikan Convolutional Neural Network (CNN) dan Multilayer Perceptron (MLP) untuk meningkatkan kemampuan deteksi serangan DDoS pada lingkungan Software Defined Network (SDN) [3]. Selain itu, Kumar dkk. juga menunjukkan bahwa pendekatan deep learning memiliki tingkat akurasi yang tinggi dalam mendeteksi berbagai jenis serangan DDoS modern [4]. Meskipun demikian, salah satu tantangan utama dalam pengembangan model klasifikasi serangan siber adalah ketidakseimbangan distribusi data (*class imbalance*). Dataset keamanan siber umumnya memiliki jumlah data yang jauh lebih besar pada beberapa kelas dibandingkan kelas lainnya. Kondisi ini menyebabkan model cenderung mempelajari pola dari kelas mayoritas dan mengabaikan karakteristik kelas minoritas sehingga performa klasifikasi pada beberapa jenis serangan menjadi kurang optimal [5].

Permasalahan *class imbalance* juga ditemukan pada dataset CICDDoS2019 yang merupakan salah satu dataset benchmark yang banyak digunakan untuk penelitian deteksi DDoS. Dataset ini menyediakan berbagai jenis serangan DDoS modern dengan karakteristik lalu lintas jaringan kompleks, namun memiliki distribusi data yang tidak merata antar kelas. Ketidakseimbangan tersebut berpotensi menurunkan kemampuan model dalam mengenali jenis serangan yang jumlah datanya relatif sedikit [6].

Beberapa penelitian telah mencoba mengatasi permasalahan *class imbalance* melalui pendekatan oversampling, undersampling, maupun cost-sensitive learning. Penelitian yang dilakukan oleh Priatna, Purnomo, Iriani, Sembiring, dan Wellem menunjukkan bahwa optimasi Multilayer Perceptron menggunakan Cost-Sensitive Learning mampu meningkatkan performa klasifikasi pada data yang tidak seimbang dengan memberikan penalti lebih besar terhadap kesalahan klasifikasi pada kelas minoritas [7]. Hasil tersebut menunjukkan bahwa penanganan *class imbalance* merupakan faktor penting dalam meningkatkan performa model berbasis jaringan saraf tiruan. Meskipun demikian, pendekatan Cost-Sensitive Learning dan Focal Loss memiliki mekanisme penanganan ketidakseimbangan kelas yang berbeda sehingga efektivitas Focal Loss pada MLP untuk klasifikasi multi-kelas serangan DDoS masih perlu dievaluasi lebih lanjut.

Berdasarkan kajian penelitian terdahulu, sebagian besar penelitian masih berfokus pada peningkatan akurasi klasifikasi secara umum, sementara implementasi *Focal Loss* pada arsitektur Multilayer Perceptron untuk klasifikasi multi-kelas serangan siber menggunakan dataset CICDDoS2019 masih relatif terbatas. *Focal Loss* merupakan fungsi kerugian yang dirancang untuk

memberikan perhatian lebih besar terhadap sampel yang sulit diklasifikasikan sehingga berpotensi meningkatkan sensitivitas model terhadap kelas minoritas.

Oleh karena itu, penelitian ini mengusulkan penerapan Multilayer Perceptron (MLP) dengan *Focal Loss* untuk mengatasi permasalahan ketidakseimbangan kelas pada dataset CICDDoS2019. Penelitian ini bertujuan untuk mengembangkan model klasifikasi serangan siber berbasis Multilayer Perceptron dengan *Focal Loss* yang mampu mengidentifikasi berbagai jenis serangan DDoS pada kondisi distribusi data yang tidak seimbang.

2 Tinjauan Literatur

Intrusion Detection System (IDS) merupakan sistem yang dirancang untuk mendeteksi aktivitas mencurigakan maupun serangan yang terjadi pada jaringan komputer. IDS berperan penting dalam menjaga keamanan jaringan dengan melakukan pemantauan terhadap lalu lintas jaringan dan mengidentifikasi pola aktivitas yang menyimpang dari kondisi normal. Seiring meningkatnya kompleksitas serangan siber, pendekatan berbasis machine learning dan deep learning semakin banyak digunakan karena mampu mempelajari pola serangan secara otomatis dari data historis serta memberikan kemampuan deteksi yang lebih adaptif dibandingkan pendekatan berbasis aturan statis [8].

Salah satu jenis serangan yang paling sering menjadi target penelitian adalah Distributed Denial of Service (DDoS). Serangan DDoS bertujuan mengganggu ketersediaan layanan dengan membanjiri target menggunakan sejumlah besar trafik berbahaya yang berasal dari berbagai sumber secara bersamaan. Serangan ini dapat menyebabkan penurunan performa layanan hingga penghentian layanan secara total. Kompleksitas pola serangan DDoS modern menyebabkan kebutuhan akan metode klasifikasi yang mampu mengidentifikasi berbagai jenis serangan secara cepat dan akurat [1].

Dalam beberapa tahun terakhir, algoritma deep learning menunjukkan performa yang menjanjikan dalam mendeteksi serangan siber. Salah satu arsitektur yang banyak digunakan adalah Multilayer Perceptron (MLP). MLP merupakan jaringan saraf tiruan yang terdiri dari lapisan input, satu atau lebih *hidden layer*, serta lapisan output yang saling terhubung melalui bobot (weight). Kemampuan MLP dalam memodelkan hubungan non-linear antar fitur menjadikan efektif untuk berbagai permasalahan klasifikasi, termasuk deteksi anomali dan serangan jaringan [9]. Penelitian yang dilakukan oleh Ahmed dkk. menunjukkan bahwa MLP mampu menghasilkan performa klasifikasi yang tinggi pada deteksi serangan DDoS dengan memanfaatkan karakteristik lalu lintas jaringan sebagai fitur masukan [1].

Dataset CICDDoS2019 merupakan salah satu dataset benchmark yang banyak digunakan dalam penelitian deteksi serangan DDoS. Dataset ini dikembangkan oleh Canadian Institute for Cybersecurity (CIC) dan menyediakan berbagai jenis serangan DDoS modern seperti DrDoS_DNS, DrDoS_NTP, DrDoS_SSDP, TFTP, UDP, UDPLag, Syn, Portmap, dan WebDDoS. Keberagaman jenis serangan tersebut memungkinkan peneliti mengevaluasi kemampuan model dalam melakukan klasifikasi multi-kelas pada lingkungan yang menyerupai kondisi jaringan nyata [10]. Namun demikian, distribusi data pada dataset CICDDoS2019 tidak merata sehingga menimbulkan permasalahan *class imbalance* yang dapat memengaruhi performa model klasifikasi.

Class imbalance terjadi ketika jumlah data pada suatu kelas jauh lebih besar dibandingkan kelas lainnya. Dalam kasus klasifikasi serangan siber, kondisi ini menyebabkan model lebih cenderung mempelajari pola dari kelas mayoritas dan mengabaikan karakteristik dari kelas minoritas. Akibatnya, model dapat menghasilkan nilai akurasi tinggi tetapi memiliki kemampuan yang rendah dalam mendeteksi jenis serangan tertentu yang jumlah datanya lebih sedikit [11]. Oleh karena itu, penanganan *class imbalance* menjadi salah satu aspek penting dalam pengembangan model deteksi serangan siber.

Berbagai pendekatan telah dikembangkan untuk mengatasi permasalahan *class imbalance*. Teknik oversampling dan undersampling merupakan metode yang umum digunakan untuk menyeimbangkan distribusi data sebelum proses pelatihan model. Selain itu, pendekatan cost-sensitive learning juga banyak diterapkan dengan memberikan penalti yang lebih besar terhadap kesalahan klasifikasi pada kelas minoritas. Penelitian yang dilakukan oleh Priatna, Purnomo, Iriani, Sembiring, dan Wellem [7] menunjukkan bahwa optimasi Multilayer Perceptron menggunakan Cost-Sensitive Learning mampu meningkatkan kemampuan model dalam menangani data yang tidak seimbang. Hasil penelitian tersebut menunjukkan peningkatan performa klasifikasi dibandingkan model MLP konvensional sehingga memperkuat pentingnya strategi penanganan *class imbalance* pada model deep learning.

Salah satu pendekatan lain yang berkembang untuk mengatasi *class imbalance* adalah *Focal Loss*. *Focal Loss* merupakan pengembangan dari *Cross Entropy Loss* yang memberikan bobot lebih besar terhadap sampel yang sulit diklasifikasikan dan mengurangi kontribusi sampel yang mudah diklasifikasikan selama proses pelatihan. Tantangan ketidakseimbangan kelas pada dataset keamanan siber telah memicu pengembangan fungsi kerugian yang lebih adaptif, dimana *Focal Loss* terbukti efektif memberikan bobot pembaruan lebih besar pada sampel yang sulit di klasifikasikan tanpa mengubah distribusi data asli [13],[14]. Dengan mekanisme tersebut, model menjadi lebih fokus dalam mempelajari kelas minoritas sehingga berpotensi meningkatkan kemampuan klasifikasi pada dataset yang memiliki distribusi data tidak seimbang.

Untuk memperjelas *research gap* serta posisi kebaruan (novelty) dari penelitian ini dibandingkan dengan penelitian terdahulu, berikut pada Tabel 1 perbandingan singkat dengan beberapa penelitian sebelumnya.

Tabel 1 Perbandingan

Penelitian	Dataset	Model/Metode	Teknik Penanganan <i>Class imbalance</i>	Hasil Utama (Accuracy-F1)	Keterbatasan (Research Gap)
Ahmed dkk.(2023) [1]	CICDDoS2019	Multilayer Perceptron (MLP)	Tidak difokuskan secara eksplisit	Akurasi >98%	Mengabaikan evaluasi mendalam pada kelas minoritas berukuran ekstrim.
Mehmood dkk. (2025) [3]	SDN Dataset	CNN-MLP Hybrid	Optimization-based featured selection	Akurasi 98,2%	Kompleksitas komputasi tinggi dan belum menguji efektivitas <i>Focal Loss</i> .
Priatna dkk. (2024) [7]	Credit Card Fraud	MLP	Cost-Sensitive Learning	Weighted F1-Score tinggi	Fokus pada klasifikasi biner dan domain finansial.
Penelitian ini	CICDDoS2019	MLP (2 Hidden Layer)	<i>Focal Loss</i> ($\gamma = 2.0, \alpha = 0.25$)	98,80%, Macro F1 0.87, Weighted F1 0.99	Mengkaji/menerapkan kombinasi <i>Focal Loss</i> dan class weight khusus pada klasifikasi multi kelas (18 kelas) serangan DDoS <i>Imbalanced</i> .

Berbeda dengan penelitian-penelitian terdahulu yang umumnya berfokus pada peningkatan akurasi keseluruhan atau pengolahan data biner, kebaruan penelitian ini terletak pada eksplorasi integrasi fungsi kerugian *Focal Loss* pada arsitektur Multilayer Perceptron (MLP) untuk mengatasi ketidakseimbangan kelas ekstrem pada dataset multi-kelas CICDDoS2019 tanpa modifikasi distribusi asli data melalui teknik *resampling*. Pendekatan ini diharapkan mampu meningkatkan sensitivitas model terhadap kelas minoritas sekaligus mempertahankan performa klasifikasi secara keseluruhan sehingga dapat mendukung pengembangan Intrusion Detection System yang lebih efektif.

3 Metode Penelitian

Penelitian ini menggunakan pendekatan kuantitatif dengan tahapan pra-proses data, perancangan model, pelatihan, dan evaluasi. Tahap pra-proses data dilakukan untuk mempersiapkan dataset agar dapat digunakan dalam proses klasifikasi. Selanjutnya, model dirancang dengan menggunakan algoritma Multilayer Perceptron (MLP) untuk melakukan klasifikasi terhadap jenis serangan siber. Tahap pelatihan dilakukan dengan menggunakan data latih yang telah melalui pra-proses dan pembagian data. Setelah pelatihan selesai, model dievaluasi untuk mengetahui kemampuan model dan mengklasifikasikan data serangan siber.

Dataset CICDDoS2019, yang berisi berbagai jenis serangan Distributed Denial of Service (DDoS) serta trafik normal. Dataset tersebut digunakan sebagai data utama dalam proses pengembangan dan pengujian model klasifikasi serangan siber. Setelah proses pembersihan dan seleksi data, diperoleh sebanyak 345.096 data digunakan sebagai data latih dan 86.275 data digunakan sebagai data uji. Pembagian data dilakukan menggunakan teknik *stratified split* untuk menjaga distribusi setiap kelas tetap proporsional pada data latih dan data uji.

Rincian distribusi jumlah sampel untuk masing-masing kelas pada dataset CICDDoS2019 sebelum dan sesudah pembagian data disajikan pada Tabel 2. Tabel tersebut menunjukkan jumlah data yang digunakan sebagai data latih dan data uji pada setiap kelas. Pembagian data dilakukan dengan mempertahankan proporsi masing-masing kelas agar data latih dan uji tetap merepresentasikan distribusi kelas pada dataset. Informasi jumlah sampel pada setiap kelas juga digunakan untuk melihat perbedaan distribusi data sebelum dan sesudah proses pembagian. Dengan demikian, Tabel 2 memberikan gambaran mengenai distribusi kelas dataset yang digunakan dalam proses pelatihan dan pengujian model.

Tabel 2 Distribusi kelas dataset CICDDoS2019 sebelum dan sesudah pembagian data (80:20)

No.	Nama Kelas /Label	Jumlah Data Latih (80%)	Jumlah Data Uji (20%)	Total Sampel	Persentase (%)
1.	Benign	78.265	19.566	97.831	22,68%
2.	DrDoS_DNS	2.934	734	3.668	0,85%
3.	DrDoS_LDAP	1.152	288	1.440	0,33%
4.	DrDoS_MSSQL	4.968	1.242	6.210	1,44%
5.	DrDoS_NTP	97.094	24.274	121.368	28,14%
6.	DrDoS_NetBIOS	480	120	600	0,14%
7.	DrDoS_SNMP	2.172	543	2.715	0,63%
8.	DrDoS_UDP	8.336	2.084	10.420	2,42%
9.	LDAP	1.524	381	1.905	0,44%
10.	MSSQL	6.820	1.705	8.525	1,98%
11.	NetBIOS	516	129	645	0,15%
12.	Portmap	548	137	685	0,16%
13.	Syn	39.498	9.875	49.373	11,45%
14.	TFTP	79.133	19.784	98.917	22,93%
15.	UDP	14.472	3.618	18.090	4,19%

16.	UDP-Lag	7.096	1.774	8.870	2,06%
17.	UDPLag	44	11	55	0,01%
18.	WebDDoS	40	10	50	0,01%
Total	345.096	86.275	431.371	100,00%	

Tahapan pra-proses data meliputi pembersihan nilai kosong (*missing value*), pengkodean label menggunakan label *encoding*, normalisasi fitur numerik, serta pembobotan kelas (*class weighting*) bersama penggunaan Focal Loss sebagai fungsi kerugian.

Secara matematis, *Focal Loss (FL)* dirancang untuk mengatasi *class imbalance* dengan menekan kerugian (*loss*) dari sampel yang mudah diklasifikasikan (*well-classified samples*) dan memfokuskan pembaruan bobot pada sampel yang sulit. Persamaan *Focal Loss* didefinisikan sebagaimana persamaan (1) berikut:

$$FL(p_i) = -\alpha(1-p_i)^\gamma \log(p_i) \quad (1)$$

Pada persamaan diatas, $FL(p_i)$ menyatakan nilai Focal Loss untuk suatu sampel. Sedangkan p_i menyatakan probabilitas prediksi model terhadap kelas target yang benar. Parameter α merupakan faktor pembobotan yang digunakan dalam fungsi Focal Loss, sedangkan γ merupakan focusing parameter yang mengatur tingkat pengurangan kontribusi sampel yang mudah diklasifikasikan. Nilai γ yang lebih besar menyebabkan fungsi loss memberikan fokus lebih besar terhadap sampel yang memiliki probabilitas prediksi rendah terhadap kelas target. Dalam penelitian ini, digunakan $\alpha = 0,25$ dan $\gamma = 2,0$ sesuai dengan parameter yang ditetapkan pada implementasi Focal Loss.

Pada proses pelatihan model, Focal Loss diterapkan pada data multikelas dengan menggunakan representasi one-hot. Perhitungan Focal Loss seluruh kelas dilakukan dengan persamaan (2) berikut:

$$Loss_Focal = 1/N \sum_{i=1}^N \sum_{c=1}^C [-\alpha y_{i,c} (1-\hat{y}_{i,c})^\gamma \log(\hat{y}_{i,c})] \quad (2)$$

Pada persamaan diatas, **Loss_Focal** menyatakan nilai rata-rata Focal Loss, N menyatakan jumlah sampel yang dihitung dalam satu batch, dan i merupakan indeks sampel. Simbol c menyatakan jumlah kelas yang diklasifikasikan oleh model, yaitu 18 kelas. Dalam penelitian ini terdapat 18 kelas yang diklasifikasikan oleh model. Variabel $y_{i,c}$ menyatakan nilai label sebenarnya untuk sampel ke- i pada kelas ke- c dalam bentuk one-hot encoding, sedangkan $\hat{y}_{i,c}$ menyatakan probabilitas prediksi model sampel ke- i pada kelas ke- c . Parameter α merupakan faktor pembobotan pada Focal Loss dengan nilai 0,25, sedangkan γ merupakan focusing parameter dengan nilai 2,0. Simbol **log** menyatakan fungsi logaritma natural.

Model utama yang digunakan adalah Multilayer Perceptron (MLP). Arsitektur model terdiri atas dua *hidden layer*, yaitu *hidden layer* pertama dengan 128 neuron dan fungsi aktivasi ReLU, serta *hidden layer* kedua dengan 64 neuron dan fungsi aktivasi ReLU. Untuk meningkatkan stabilitas proses pelatihan digunakan *Batch Normalization*, sedangkan *Dropout* sebesar 0,3 diterapkan untuk mengurangi risiko *overfitting*. Pada lapisan keluaran digunakan fungsi aktivasi Softmax untuk menghasilkan probabilitas klasifikasi multi-kelas.

Pemilihan dua *hidden layer* bertujuan agar model mampu menangkap pola non-linear yang kompleks tanpa meningkatkan kompleksitas model secara berlebihan. Jumlah neuron 128 dan 64 dipilih untuk memberikan keseimbangan antara kapasitas representasi dan efisiensi komputasi. Sementara itu, penerapan *Dropout* sebesar 0,3 bertujuan untuk meningkatkan kemampuan generalisasi model terhadap data yang belum pernah dilihat sebelumnya.

Model dilatih selama 30 *epoch* dengan *batch size* 256 menggunakan optimizer Adam. Optimizer Adam dipilih karena kemampuannya dalam melakukan adaptasi *learning rate* secara otomatis sehingga dapat mempercepat konvergensi selama pelatihan model.

Sebagai pembandingan, digunakan dua model *baseline*, yaitu Logistic Regression yang merepresentasikan model klasifikasi linear dan Random Forest yang merepresentasikan model *ensemble* berbasis pohon keputusan. Perbandingan dilakukan untuk mengevaluasi efektivitas model MLP yang diusulkan dalam melakukan klasifikasi serangan DDoS pada dataset CICDDoS2019.

Evaluasi model dilakukan menggunakan metrik Accuracy, Precision, Recall, dan F1-Score. Selain itu, digunakan *Confusion Matrix* untuk menganalisis distribusi kesalahan prediksi pada setiap kelas sehingga dapat diketahui kemampuan model dalam menangani kelas mayoritas maupun kelas minoritas.

Seluruh proses implementasi dilakukan menggunakan bahasa pemrograman Python dengan dukungan framework TensorFlow/Keras dan Scikit-learn.

4 Hasil dan Pembahasan

4.1 Evaluasi Model Multilayer Perceptron

Hasil evaluasi performa model MLP dengan Focal Loss secara menyeluruh disajikan pada Tabel 3.

Tabel 3 Hasil evaluasi performa model MLP dengan Focal Loss

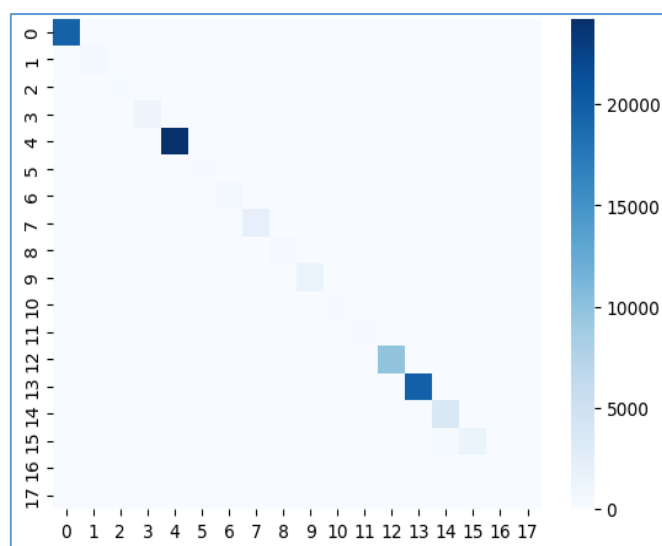
Metrik Evaluasi	Macro Average	Weighted Average
Accuracy	98.80%	98.80%
Precision	0.89	0.99
Recall	0.86	0.99
F1-Score	0.87	0.99

Hasil evaluasi secara rinci menunjukkan bahwa model mencapai performa yang sangat tinggi pada kelas mayoritas seperti Benign, DrDDoS_NTP, dan TFTP dengan nilai precision, Recall, dan F1-Score mencapai 1.00. Namun, pada kelas minoritas berukuran sedang seperti UDP-lag, diperoleh nilai Precision 1.00, Recall 0.75, dan F1-Score 0.85. Untuk kelas minoritas ekstrem seperti WebDDoS (10 sampel uji), model memperoleh Precision 0.80, Recall 0.40, dan F1-Score 0.53. Sementara itu, pada kelas UDPLag (11 sampel uji), model belum mampu memprediksi sampel dengan benar sehingga menghasilkan F1-Score 0.00.

Perbedaan performa pada kelas minoritas inilah yang menyebabkan nilai Macro Precision (0.89) dan Macro Recall (0.86) lebih rendah dibandingkan Weighted Average (0.99).

4.2 Analisis Confusion Matrix

Untuk memahami perilaku model secara lebih dalam, dilakukan analisis menggunakan *confusion matrix* pada data uji. Hasil visualisasi matriks disajikan pada Gambar 1.



Gambar 1 Confusion matrix model multilayer perceptron

<http://sistemasi.ftik.unisi.ac.id>

Berdasarkan visualisasi *confusion matrix*, model menunjukkan tingkat prediksi yang sangat tinggi pada kelas mayoritas seperti DrDDoS, TFTP, dan UDP yang ditandai dengan dominasi nilai diagonal utama. Namun demikian, pada kelas minoritas seperti UDPLag dan WebDDoS masih ditemukan sejumlah kesalahan klasifikasi, terutama ketika pola trafiknya memiliki kemiripan dengan kelas serangan lain berbasis UDP atau protokol serupa.

4.3 Perbandingan dengan Model Baseline

Akurasi performa model MLP utama kemudian dibandingkan dengan dua model dasar klasik lainnya yang diringkas pada Tabel 4. Perbandingan ini dilakukan untuk mengukur efektivitas arsitektur *deep learning* MLP dibandingkan dengan algoritma *machine learning* linier dan berbasis *ensemble*. Berdasarkan hasil pengujian, algoritma *Logistic Regression* memperoleh nilai akurasi terendah yaitu sebesar 96,91% karena keterbatasannya dalam memetakan batas keputusan nonlinear pada fitur lalu lintas jaringan yang kompleks. Di sisi lain, *Random Forest* mencatatkan akurasi numerik tertinggi yaitu sebesar 99,18% berkat kemampuannya menangani struktur data berdimensi tinggi melalui pembentukan banyak pohon keputusan. Meskipun demikian, model MLP yang diusulkan berhasil menunjukkan performa yang sangat kompetitif dengan akurasi 98,80%. Keunggulan MLP dengan *Focal Loss* terletak pada kemampuannya untuk tetap stabil dan konsisten dalam mempelajari karakter fitur abstrak tanpa mengalami ketergantungan berlebih (*overfitting*) pada kelas mayoritas.

Tabel 4 Pembandingan performa model klasifikasi

Model	Nilai
Logistic Regression	96,91%
Random Forest	99,18%
Multilayer Perceptron	98,80%

4.4 Pembahasan

Nilai akurasi sebesar 98,80% menunjukkan bahwa model MLP mampu mengklasifikasikan sebagian besar sampel dengan sangat baik. Namun, perbedaan antara *macro F1-Score* (0,87) dan *weighted F1-Score* (0,99) menunjukkan adanya pengaruh distribusi kelas yang tidak seimbang. *Macro f1-Score* menghitung rata-rata performa setiap kelas tanpa mempertimbangkan jumlah sampel, sehingga nilai yang lebih rendah mengindikasikan bahwa performa kelas minoritas ekstrem belum optimal. Sebaliknya, *weighted F1-Score* yang tinggi menunjukkan bahwa model sangat unggul dalam mengenali kelas mayoritas.

Pada pengujian model pembandingan, Random Forest menghasilkan akurasi sebesar 99,18%, sedikit lebih tinggi secara numerik dibandingkan MLP (98,80%). Keunggulan Random Forest ini disebabkan oleh kemampuannya memisahkan ruang fitur tabular berdimensi tinggi secara non-linear melalui gabungan pohon keputusan (*ensemble decision trees*) tanpa bergantung pada pembobotan fungsi kerugian. Walaupun Random Forest unggul tipis secara numerik, model MLP tetap menunjukkan performa yang sangat kompetitif dan memiliki fleksibilitas tinggi sebagai arsitektur *deep learning* yang dapat dikembangkan lebih lanjut untuk deteksi berbasis aliran trafik (*streaming*) secara real-time. Di sisi lain, Logistic Regression menghasilkan akurasi terendah (96,91%), yang mengindikasikan keterbatasan model linear dalam memodelkan pola serangan DDoS yang kompleks.

4.5 Analisis Perilaku Model

Hasil pengujian menunjukkan bahwa kombinasi *Focal Loss* dan pembobotan kelas berpotensi membantu menjaga stabilitas performa pada kelas mayoritas dan menengah. Namun demikian, pernyataan bahwa *Focal Loss* terbukti efektif secara mandiri masih memerlukan kehati-hatian, sebab penelitian ini belum menyertakan eksperimen baseline terisolasi (seperti membandingkan MLP + Cross Entropy standar secara langsung dengan MLP + *Focal Loss* tanpa pembobotan kelas). Oleh karena itu, diperlukannya eksperimen pembandingan lebih lanjut untuk mengisolasi dampak tunggal dari masing-masing metode penanganan ketidakseimbangan kelas tersebut.

Penelitian ini dilaporkan berdasarkan satu kali pengujian pemisahan data (*single-test split* 80:20). Penggunaan *single split* ini diakui menjadi salah satu keterbatasan dalam mengevaluasi tingkat stabilitas statistik model. Untuk meningkatkan validitas hasil di masa mendatang, disarankan untuk melakukan pengujian berulang (*multi-run*) menggunakan variasi *random seed* berbeda atau skema *K-Fold Cross-Validation* guna melaporkan nilai rata-rata (*mean*) serta standar deviasi (*standard deviation*) dari metrik akurasi dan Macro F1-Score).

5 Kesimpulan

Penelitian ini berhasil mengimplementasikan model *Multilayer Perceptron* (MLP) dengan *Focal Loss* dan *class weighting* untuk klasifikasi multi-kelas serangan siber pada dataset CICDDoS2019 yang mengalami *class imbalance*, dengan pencapaian akurasi sebesar 98,80%, *weighted F1-Score* sebesar 0,99, dan *macro F1-Score* 0,87. Hasil pengujian menunjukkan bahwa kombinasi *Focal Loss* dan *class weighting* berpotensi membantu mengurangi dominasi kelas mayoritas pada beberapa kelas terdistribusi sedang. Namun, demikian, ketidakmampuan model mengenali kelas minoritas ekstrem (seperti UDPLag) mengindikasikan bahwa penggunaan *Focal Loss* belum cukup kuat jika diimplementasikan tanpa kombinasi teknik *resampling*. Oleh karena itu, diperlukan penelitian lanjutan dengan eksperimen pembandingan terisolasi untuk mengukur kontribusi pasti dari *Focal Loss* secara mandiri. Secara teoritis, penelitian ini memberikan kontribusi bagi ilmu pengetahuan berupa pemahaman empiris mengenai batas kemampuan *Focal Loss* pada arsitektur jaringan saraf tiruan sederhana dalam menangani ketidakseimbangan kelas ekstrem. Secara praktis di dunia nyata, hasil penelitian ini memberikan dampak positif bagi pengembang sistem keamanan jaringan (*Intrusion Detection System*) dalam merancang mekanisme deteksi dini serangan DDoS yang lebih adaptif, efisien, dan stabil. Untuk meningkatkan validitas dan stabilitas statistik hasil di masa mendatang, disarankan pula untuk menerapkan pengujian berulang (*multi-run*) menggunakan variasi *random seed* berbeda atau skema *K-Fold Cross-Validation*.

Referensi

- [1] A. Affinio et al., "The Evolution of Mirai Botnet Scans Over a Six-Year Period," *Journal of Network and Computer Applications*, Vol. 225, p. 103756, 2024.
- [2] S. Ahmed, Z. A. Khan, S. M. Mohsin, S. Latif, S. Aslam, H. Mujlid, and Z. Najam, "Effective and Efficient DDoS Attack Detection using Deep Learning Algorithm, Multi-Layer Perceptron," *Future Internet*, Vol. 15, No. 2, p. 76, 2023.
- [3] D. Akgun, S. Hizal, and U. Cavusoglu, "A New DDoS Attacks Intrusion Detection Model based on Deep Learning for Cybersecurity," *Computers & Security*, Vol. 118, p. 102746, 2022.
- [4] E. O. Nasution and A. Basuki, "Implementasi Algoritme C5.0 untuk Klasifikasi Serangan DDoS," *Jurnal Pengembangan Teknologi Informasi dan Ilmu Komputer*, Vol. 5, No. 1, pp. 389–395, 2021.
- [5] S. Mehmood et al., "Distributed Denial of Service (DDoS) Attack Detection in SDN using Optimizer-Equipped CNN-MLP," *PLOS ONE*, Vol. 20, No. 1, p. e0312425, 2025.
- [6] D. Kumar, R. K. Pateriya, R. K. Gupta, V. Dehalwar, and A. Sharma, "DDoS Detection using Deep Learning," *Procedia Computer Science*, Vol. 218, pp. 2420–2429, 2023.
- [7] W. Chimphlee and S. Chimphlee, "Network Intrusion Detector using Multilayer Perceptron (MLP) Approach," *Turkish Journal of Computer and Mathematics Education*, Vol. 13, No. 03, pp. 488–499, 2022.
- [8] Z. S. Arrak and R. J. S. Al-Janabi, "Detecting DDoS Attacks using Machine Learning: Survey," *Journal of Al-Qadisiyah for Computer Science and Mathematics*, Vol. 16, No. 2, pp. 118–134, 2024.
- [9] W. Priatna, H. D. Purnomo, A. Iriani, I. Sembiring, and T. Wellem, "Optimizing Multilayer Perceptron with Cost-Sensitive Learning for Addressing Class Imbalance in Credit Card Fraud Detection," *Jurnal RESTI (Rekayasa Sistem dan Teknologi Informasi)*, Vol. 8, No. 4, pp. 563–570, 2024.

- [10] I. Maulana and A. Alamsyah, "Optimalisasi Deteksi Serangan DDoS menggunakan Algoritma *Random Forest*, *SVM*, *KNN* dan *MLP* pada Jaringan Komputer," *Indonesian Journal of Mathematics and Natural Sciences*, Vol. 46, No. 2, 2023.
- [11] P. Chang, "Multi-Layer Perceptron Neural Network for Improving Detection Performance of Malicious Phishing URLs without Affecting Other Attack Types Classification," arXiv preprint arXiv:2203.00774, 2022.
- [12] K.Saluja et al., "Exploring Robust DDoS Detection: A Machine Learning Analysis with the *CICDDoS2019 Dataset*," in *Proc. IEEE INDISCON*, 2024, pp. 1–6.
- [13] C. Singh et al., "A Comprehensive Survey on DDoS Attacks Detection & Mitigation," *Computer Networks*, Vol. 242, p. 110137, 2024.
- [14] F. A. Setiawan, A. S. Ahmad, and R. A. Asmara, "Handling Class Imbalance in Network Intrusion Detection Systems Using Focal Loss," *Journal of Computer Science and Information Technology*, Vol. 10, No. 1, pp. 45–54, 2023.
- [15] M. A. Setitra, M. Fan, B. L. Y. Agbley, and Z. E. A. Bensalem, "Optimized MLP-CNN Model to Enhance Detecting DDoS Attacks in SDN Environment," *Network*, Vol. 3, No. 4, pp. 538–562, 2023.