

Random Forest and LLM Synergies Framework for Autonomous DDoS Mitigation

Romadhon Wiratama¹, Ananta Pirdhaus², Ellys Rahma Putri Bintoro^{*3}, Zamah Sari⁴, Syaifuddin⁵

^{1,2,3,4,5}Informatics, Universitas Muhammadiyah Malang, Indonesia

Email: ³ellysrahmaputri15@webmail.umm.ac.id

Received : Sep 19, 2025; Revised : Oct 2, 2025; Accepted : Oct 7, 2025; Published : Feb 15, 2026

Abstract

Modern Distributed Denial of Service (DDoS) attacks increasingly evade traditional defenses, and while Machine Learning (ML) has improved detection accuracy, a critical challenge remains in bridging detection with effective automated mitigation. This paper introduces a novel framework centered on a cognitive agent that synergistically combines high-speed ML detection with the advanced reasoning capabilities of a Large Language Model (LLM) for autonomous DDoS mitigation. The proposed architecture operates as a closed-loop security system. Following a data preprocessing phase that includes one-hot encoding and Standard Scaling (z-score normalization), a fine-tuned Random Forest model was identified as the optimal detector with 95.99% accuracy on the UNSW-NB15 dataset, which in turn triggers the LLM-based agent. This agent autonomously generates both human-readable incident explanations and machine-executable mitigation commands. Crucially, all generated commands undergo a syntax and logic validation step before execution to ensure operational integrity. Empirical results demonstrate the framework's efficacy, achieving a complete end-to-end detection-to-mitigation cycle in 24.20 seconds. This work validates that the unified approach presents a viable and transparent paradigm, contributing to the field of cybersecurity by enhancing automated mitigation and analytical processes through responsive and intelligent defense mechanisms.

Keywords : *Agentic AI Framework, Autonomous DDoS Mitigation, Closed-Loop Security, Cognitive Agent, Large Language Models, Machine Learning*

This work is an open access article and licensed under a Creative Commons Attribution-Non Commercial 4.0 International License



1. PENDAHULUAN

Serangan Distributed Denial of Service (DDoS) kini menjadi salah satu ancaman terbesar dalam keamanan jaringan. Baik di ranah publik maupun privat, pola serangan semakin kompleks seiring dengan meningkatnya jumlah perangkat Internet of Things (IoT) yang terhubung. Botnet berskala besar dengan mudah memanfaatkan perangkat IoT yang rentan untuk meluncurkan serangan. Contoh klasik adalah SYN Flood, yang mampu melumpuhkan server hanya dengan membanjirinya permintaan koneksi palsu [1]. Seiring waktu, serangan DDoS tidak lagi selalu tampil dengan volume besar. Beberapa varian, seperti Low Rate DDoS, dirancang agar beroperasi secara tersembunyi sehingga sulit dikenali oleh metode deteksi tradisional [2], [3]. Perkembangan ini menunjukkan bahwa model pertahanan lama berbasis tanda tangan tidak lagi memadai menghadapi ancaman modern.

Untuk menjawab tantangan tersebut, komunitas riset mulai memanfaatkan pendekatan berbasis Machine Learning (ML) dan deep learning (DL). Algoritma seperti Naïve Bayes dan Deep Neural Network (DNN) telah menunjukkan peningkatan akurasi deteksi secara signifikan dibandingkan metode tradisional [2], [4]. Hal ini membuka jalan baru dalam upaya melawan serangan yang semakin canggih. Model lain yang banyak digunakan dalam penelitian adalah Support Vector Machine (SVM), Decision Tree, dan Random Forest. Validasi menggunakan dataset CICDDoS2019 membuktikan bahwa model

ini efektif dalam mengklasifikasikan lalu lintas normal dan berbahaya [5]. Keunggulan Random Forest, misalnya, terletak pada kemampuannya mengurangi overfitting melalui pendekatan ensemble learning [6]. Selain itu, Convolutional Neural Network (CNN) yang biasanya dipakai untuk analisis citra juga berhasil diadaptasi dalam mendeteksi pola data lalu lintas jaringan. Dengan menangkap fitur spasial lokal pada paket data, CNN memberikan akurasi yang lebih baik dibanding sistem deteksi berbasis tanda tangan [7], [8]. Pendekatan ini semakin relevan ketika diiringi dengan rekayasa fitur, misalnya penggunaan Principal Component Analysis (PCA), untuk mengurangi dimensi data dan meningkatkan efisiensi komputasi [9], [10]. Meski demikian, pendekatan berbasis ML juga memiliki keterbatasan. Salah satunya adalah kebutuhan akan data pelatihan berkualitas tinggi, serta tingginya tingkat false positive yang dapat mengganggu layanan [11]. Sebaliknya, metode rule-based menawarkan transparansi lebih baik karena aturan yang jelas memudahkan proses investigasi. Namun, metode ini sulit beradaptasi dengan serangan baru yang tidak sesuai dengan aturan yang ada [11], [12], [13]. Banyak penelitian terkini menyarankan pendekatan hibrida yang mengkombinasikan kelebihan kedua metode tersebut agar sistem lebih adaptif dan tetap transparan [11].

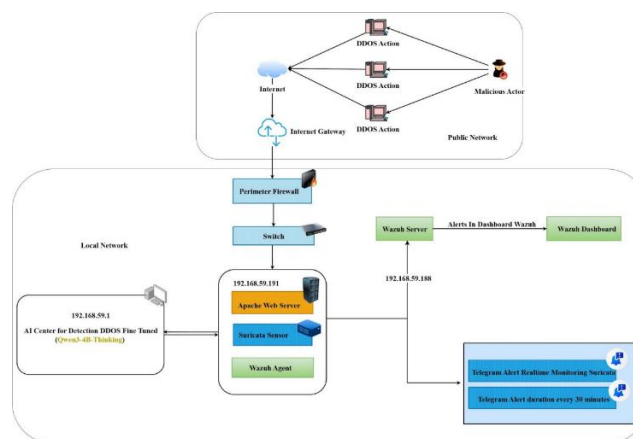
Selain deteksi, aspek mitigasi juga krusial. Berbagai teknologi telah diadopsi, mulai dari firewall berbasis Iptables hingga solusi perlindungan berbasis cloud seperti Cloudflare [14], [15]. Dampak serangan DDoS bukan hanya terganggunya layanan daring, tetapi juga kerugian finansial yang besar serta turunnya reputasi organisasi [16], [17], [18]. Teknologi Web Application Firewall (WAF) seperti milik Cloudflare terbukti mampu memblokir serangan secara proaktif [14], [18]. Inovasi lebih lanjut terlihat pada penggunaan Software-Defined Networking (SDN) sebagai fondasi mitigasi. Kerangka kerja seperti FMDADM dan CC-Guard memanfaatkan pemantauan lalu lintas secara real-time untuk mendeteksi dan mengalihkan arus data ketika pola serangan teridentifikasi [19], [20]. Penelitian lain menekankan pentingnya kemampuan adaptif dalam infrastruktur SDN agar sistem dapat menyesuaikan kebijakan lalu lintas secara dinamis [21]. Di ekosistem keamanan yang lebih luas, platform Security Information and Event Management (SIEM) juga memainkan peran penting. Wazuh dan Splunk, misalnya, mengumpulkan serta menganalisis data log untuk mendukung deteksi dini serangan DDoS [22], [23]. SIEM modern bahkan memungkinkan otomatisasi sebagian proses respons, sehingga mengurangi waktu reaksi terhadap insiden. Integrasi dengan Cyber Threat Intelligence (CTI) semakin memperkuat kemampuan deteksi secara real-time [22], [24].

Meskipun demikian, otomatisasi mitigasi masih menghadapi sejumlah tantangan. Serangan DDoS modern berkembang dengan cepat dan adaptif sehingga sistem mitigasi dituntut responsif sekaligus cerdas [11], [25]. Di sisi lain, integrasi dengan infrastruktur jaringan yang kompleks sering kali menimbulkan kendala teknis [15], [26], [27], sementara aspek regulasi dan perlindungan privasi juga menjadi isu penting dalam penerapan solusi otomatis [28]. Dalam konteks inilah, Large Language Models (LLM) mulai dilirik sebagai solusi baru karena kemampuannya memahami pola serangan melalui analisis log, menghasilkan rekomendasi mitigasi berbasis konteks, serta memperkuat intelijen ancaman [29], [30], [31]. Bahkan, penelitian terbaru menunjukkan potensi LLM dalam mitigasi ransomware dan pengelolaan keamanan berbasis Governance, Risk, and Compliance (GRC) [32], [33], [34], [35].

Secara umum, penelitian terdahulu umumnya masih bersifat konseptual dan belum mengarah pada implementasi secara *real-time*. Selain itu, aspek *latency* maupun penggunaan memori belum banyak dibahas secara eksplisit [36]. Oleh karena itu, dalam penelitian ini mengusulkan integrasi Machine Learning untuk proses deteksi DDoS dan LLM untuk proses mitigasi, sehingga membentuk sebuah Loop Otonom. Pendekatan ini juga memanfaatkan *machine learning* untuk meminimalkan potensi halusinasi yang dapat dihasilkan oleh LLM.

2. METODE

Metodologi penelitian ini disusun untuk mengembangkan dan menguji sebuah kerangka kerja agen kognitif yang menggabungkan deteksi serangan DDoS berbasis Machine Learning dengan mitigasi otomatis yang digerakkan oleh Large Language Models (LLM). Gagasan dasarnya adalah memadukan kecepatan klasifikasi algoritma dengan kemampuan penalaran LLM, sehingga tercipta sistem pertahanan siber yang lebih cerdas dan responsif. Kerangka kerja ini diharapkan mampu memangkas time-to-mitigate, yaitu jeda antara deteksi serangan dan penerapan respons pertahanan sebuah titik lemah umum pada arsitektur keamanan konvensional. Penelitian dilakukan secara terstruktur, mulai dari analisis dataset, pengembangan dan perbandingan model deteksi, hingga perancangan agen LLM dan evaluasi menyeluruh dalam lingkungan simulasi. Setiap tahap dirancang agar hasil penelitian valid secara ilmiah, dapat direplikasi, dan relevan untuk praktik nyata.



Gambar 1. Arsitektur jaringan untuk deteksi dan mitigasi serangan DDoS

Gambar 1. merupakan arsitektur jaringan untuk menguji kerangka kerja yang diusulkan, penelitian ini membangun sebuah lingkungan simulasi yang mereplikasi kondisi jaringan nyata. Topologi ini terdiri atas jaringan publik dan jaringan lokal, dengan server target (192.168.59.191) yang menjalankan layanan Apache Web Server serta dipantau oleh Suricata sebagai Intrusion Detection System (IDS) dan Wazuh Agent untuk keamanan. Komponen inti, yaitu AI Center for Detection DDoS (192.168.59.1), berfungsi sebagai pusat komando tempat agen kognitif berbasis fine-tuned LLM (Qwen3-4B-Thinking) beroperasi untuk menganalisis data sensor dan secara otomatis menghasilkan tindakan mitigasi. Seluruh peringatan dan insiden keamanan dikumpulkan oleh Wazuh Server (192.168.59.188) dan ditampilkan melalui dasbor serta dikirimkan sebagai notifikasi real-time via Telegram. Dengan arsitektur ini, serangan DDoS dari jaringan publik dapat disimulasikan secara realistis, sehingga efektivitas deteksi dan respons mitigasi kerangka kerja dapat diukur secara menyeluruh dari ujung ke ujung.

2.1. Dataset UNSW-NB15

Pemilihan dataset yang tepat menjadi fondasi penting dalam penelitian keamanan siber berbasis data. Pada studi ini digunakan UNSW-NB15, sebuah dataset yang dikembangkan oleh Australian Centre for Cyber Security menggunakan alat IXIA Perfect Storm. Dataset ini berisi kombinasi lalu lintas normal dan beragam jenis serangan kontemporer, termasuk DDoS, dengan 49 atribut yang kompleks dan merepresentasikan kondisi jaringan nyata. Karakteristik tersebut menjadikannya sangat relevan untuk mengembangkan serta menguji algoritma deteksi DDoS. Penelitian Al-Obaidi et al. bahkan menegaskan bahwa model yang divalidasi dengan dataset ini memiliki potensi generalisasi yang lebih baik di lingkungan produksi [37].

Sebelum digunakan untuk pelatihan model Machine Learning, dataset ini akan diproses melalui tahap pra-pemrosesan dan rekayasa fitur. Tahap ini mencakup pembersihan data dari nilai hilang dan anomali, konversi fitur kategorikal seperti proto, service, dan state ke bentuk numerik melalui one-hot encoding, serta normalisasi fitur numerik dengan metode Min-Max scaling agar skala data seragam. Selain itu, teknik rekayasa fitur berbasis spatio-temporal juga akan dieksplorasi untuk menangkap pola aliran data. Pendekatan ini terbukti efektif meningkatkan kinerja klasifikasi [38].

2.2. Pra-pemrosesan Data

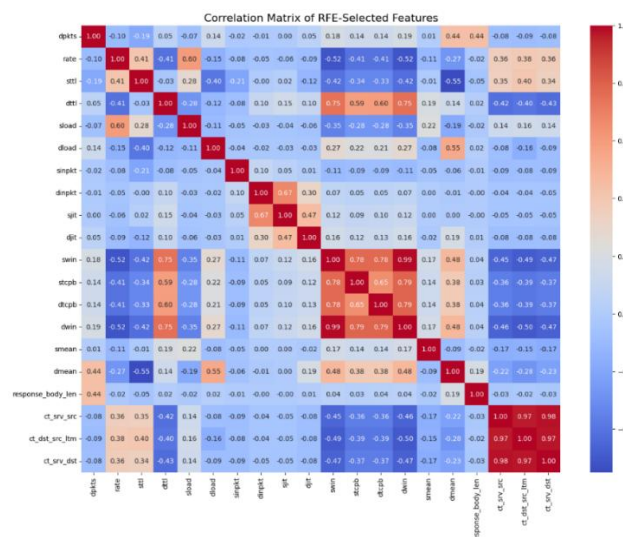
Langkah selanjutnya adalah mengolah data UNSW-NB15 melalui pra-pemrosesan. Tahapan ini mencakup pembersihan data, seleksi fitur, ekstraksi fitur dan standarisasi fitur agar dataset siap digunakan dalam pelatihan model Machine Learning.

a. Pembersihan Data

Tahap pembersihan data dilakukan untuk memastikan kualitas, kelengkapan, dan konsistensi data sebelum proses analisis. Proses ini meliputi identifikasi serta penanganan nilai yang hilang (missing values), duplikasi, dan data yang tidak valid. Nilai kosong (NaN) dihapus atau diganti dengan nilai representatif agar tidak menimbulkan bias selama pelatihan model. Selain itu, nilai placeholder seperti tanda “-” atau label “unknown” dikonversi ke format yang sesuai agar dapat dikenali oleh algoritma Machine Learning. Langkah ini bertujuan untuk menghasilkan dataset yang bersih dan siap digunakan dalam tahap pemodelan selanjutnya.

b. Seleksi Fitur

Tahap seleksi fitur bertujuan memilih atribut yang paling relevan terhadap klasifikasi. Penelitian ini menggunakan Recursive Feature Elimination (RFE), yaitu teknik yang secara bertahap menghapus fitur dengan kontribusi rendah hingga diperoleh subset yang optimal. Metode ini dipilih karena mampu meningkatkan akurasi model, mengurangi kompleksitas komputasi, dan meminimalkan risiko overfitting.



Gambar 2. Matriks Korelasi Fitur yang Dipilih Menggunakan RFE

Gambar 2. menunjukkan matriks korelasi antarfitur yang dipilih menggunakan metode RFE. Visualisasi ini membantu mengidentifikasi tingkat hubungan antarfitur sehingga dapat diketahui potensi redundansi maupun keterkaitan atribut yang berpengaruh terhadap klasifikasi.

c. Ekstraksi Fitur

Tahap ekstraksi fitur dilakukan menggunakan metode *Principal Component Analysis* (PCA) untuk mereduksi dimensi data tanpa kehilangan informasi penting. PCA bekerja dengan mentransformasikan sekumpulan fitur asli menjadi himpunan fitur baru yang disebut *principal components*, yaitu kombinasi linier dari fitur awal yang paling banyak menjelaskan variasi data. Dalam penelitian ini, jumlah fitur utama dipangkas menjadi 15 komponen untuk menyederhanakan proses komputasi sekaligus mempertahankan sebagian besar informasi yang relevan. Penggunaan PCA dipilih karena mampu mengurangi kompleksitas data yang ditandai dengan tingginya jumlah fitur dan adanya korelasi antaratribut, sehingga mempercepat proses pelatihan model serta membantu mencegah *overfitting* dengan menghilangkan fitur yang *redundant* atau kurang informatif.

d. Standarisasi Fitur

Tahap standarisasi fitur dilakukan menggunakan metode *StandardScaler* untuk menyamakan skala antar variabel sebelum proses pelatihan model. *StandardScaler* bekerja dengan mengubah distribusi setiap fitur agar memiliki nilai rata-rata (*mean*) sebesar 0 dan standar deviasi sebesar 1. Langkah ini penting karena banyak algoritma *Machine Learning*, seperti SVM dan *Logistic Regression*, sensitif terhadap perbedaan skala antar fitur. Dengan standarisasi, setiap variabel memiliki kontribusi yang seimbang dalam proses pembelajaran, sehingga model dapat menghasilkan prediksi yang lebih akurat dan stabil.

Sebagai tahap akhir prapemrosesan, dataset yang telah dibersihkan, diseleksi, diekstraksi, dan distandarisasi akan dibagi menjadi 70% data pelatihan dan 30% data pengujian. Pembagian ini memastikan model belajar dari mayoritas data sekaligus dievaluasi secara objektif dengan data baru untuk mengukur kemampuan generalisasi dalam mendeteksi serangan DDoS.

2.3. Model Deteksi Machine Learning

Modul deteksi dalam kerangka kerja ini dibangun menggunakan model *Machine Learning* yang dioptimalkan untuk kecepatan dan akurasi. Penelitian ini menguji lima algoritma yang banyak digunakan dalam literatur, yaitu *Random Forest*, *Decision Tree*, *XGBoost*, *Deep Neural Network* (DNN), dan *Gated Recurrent Unit* (GRU).

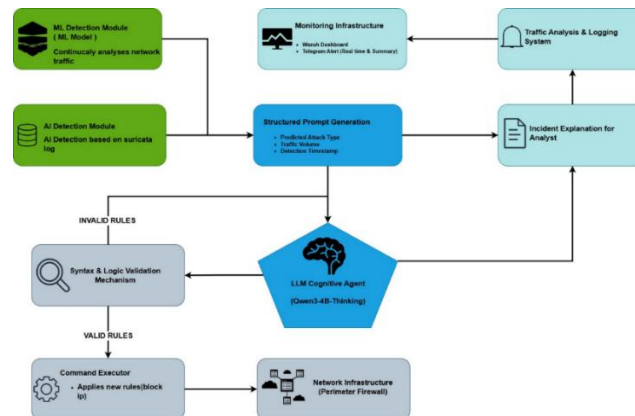
Random Forest dipilih karena kekuatan metode *ensemble*-nya dalam menangani data berdimensi tinggi sekaligus mengurangi risiko *overfitting*. *Decision Tree* digunakan sebagai *baseline* berkat interpretabilitasnya yang tinggi, sedangkan *XGBoost* dipertimbangkan karena efisiensinya yang terbukti unggul di banyak kompetisi ML. Di sisi pembelajaran mendalam, DNN dipilih untuk kemampuannya mempelajari representasi fitur non-linear yang kompleks. Sementara itu, GRU sebagai varian RNN diikutsertakan karena kemampuannya menangkap pola temporal pada data sekuensial, dan telah terbukti efektif dalam mendeteksi lalu lintas IoT [39].

Pemilihan kelima model ini didasarkan pada hipotesis bahwa tidak ada algoritma tunggal yang unggul di semua kondisi. Oleh karena itu, perbandingan empiris diperlukan untuk menentukan model paling sesuai bagi dataset UNSW-NB15.

2.4. Proses Mitigasi Berbasis Large Language Model

Modul mitigasi merupakan komponen inovatif dari kerangka kerja ini, yang memanfaatkan kemampuan penalaran dan generasi bahasa dari *Large Language Models* (LLM) untuk mengotomatiskan respons terhadap ancaman yang terdeteksi, seperti yang diilustrasikan pada Gambar 3. Proses integrasi LLM ke dalam pipeline keamanan akan mengikuti langkah-langkah teknis yang sistematis. Ketika Modul Deteksi ML (ML Model) dan Modul Deteksi AI (berdasarkan log Suricata) mengidentifikasi lalu lintas sebagai serangan dengan tingkat kepercayaan tinggi, informasi tersebut akan

dikirimkan ke modul Structured Prompt Generation. Modul ini bertanggung jawab untuk membuat prompt terstruktur yang berisi detail penting seperti jenis serangan yang diprediksi, volume lalu lintas, dan stempel waktu deteksi.



Gambar 3. Proses Mitigasi menggunakan LLM serta Proses Pelaporan menggunakan Telegram Bot dan Wazuh Dashboard

Gambar 3. merupakan proses mitigasi menggunakan LLM. Prompt tersebut dikirimkan ke LLM Cognitive Agent (Qwen3-4B-Thinking) yang menjalankan dua tugas utama. Pertama, menghasilkan penjelasan insiden dalam bahasa alami sehingga lebih mudah dipahami analis, sekaligus mengatasi masalah “kotak hitam” pada model ML. Penjelasan ini ditampilkan melalui Wazuh Dashboard, Telegram Alert, dan sistem analisis lalu lintas. Kedua, LLM merumuskan tindakan mitigasi konkret dalam format JSON, misalnya aturan firewall untuk memblokir alamat IP.

Untuk menjamin keamanan, setiap perintah yang dihasilkan akan melewati Syntax & Logic Validation sebelum dieksekusi oleh Command Executor. Hanya aturan yang valid yang akan diterapkan pada infrastruktur jaringan. Pendekatan ini sejalan dengan penelitian yang menunjukkan bahwa LLM dapat di-fine-tune atau dipandu dengan prompt agar menghasilkan output yang relevan di domain keamanan siber [40].

2.5. Evaluasi Kinerja Kerangka Kerja

Evaluasi efektivitas kerangka kerja dilakukan dengan menilai baik komponen individu maupun sistem secara keseluruhan. Pertama, lima model deteksi Machine Learning diuji secara kuantitatif menggunakan matrix precision, recall, dan F1-score. Precision penting untuk menekan alarm palsu yang dapat mengganggu layanan [41], sementara recall memastikan serangan yang sebenarnya tidak terlewat [42]. F1-score dipakai sebagai ukuran seimbang, terutama karena dataset keamanan siber umumnya tidak seimbang [43], [44]. Untuk mengurangi risiko overfitting, digunakan Stratified K-Fold Cross-Validation, yang menjaga distribusi kelas pada setiap fold agar tetap konsisten [45].

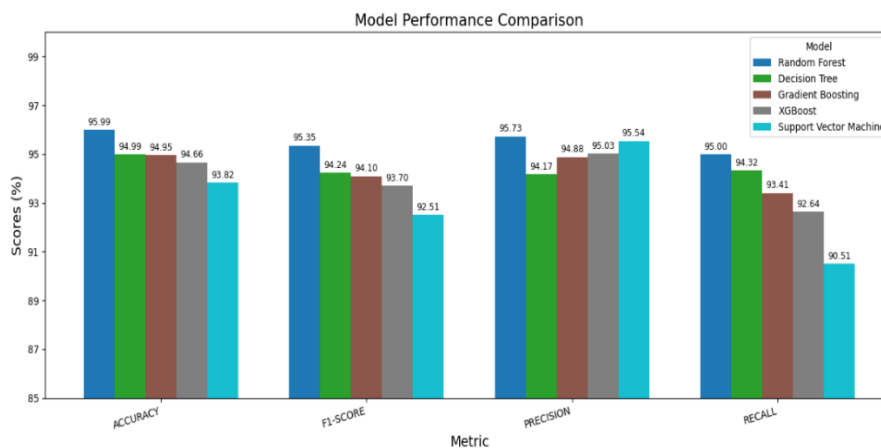
Selanjutnya, kinerja agen mitigasi berbasis LLM dinilai dari dua aspek, yaitu akurasi mitigasi dan kualitas penjelasan. Akurasi dilihat dari ketepatan perintah yang dihasilkan dalam menanggulangi serangan, sedangkan kualitas penjelasan dievaluasi secara kualitatif oleh pakar untuk memastikan kejelasan dan kegunaannya bagi analis keamanan.

Terakhir, performa sistem secara end-to-end diukur melalui metrik Time-to-Mitigate, yaitu waktu sejak paket berbahaya pertama terdeteksi hingga LLM menghasilkan perintah mitigasi yang valid. Pengujian dilakukan dalam simulasi dengan memutar ulang serangan dari dataset UNSW-NB15, sehingga latensi respons dapat diukur secara akurat. Hasil dari evaluasi ini akan menjadi bukti empiris mengenai kemampuan kerangka kerja dalam mengotomatiskan siklus pertahanan, mulai dari deteksi hingga mitigasi.

3. HASIL

Bab ini menyajikan hasil empiris dari evaluasi yang dilakukan untuk menilai efektivitas dan kinerja kerangka kerja deteksi mitigasi DDoS otonom. Hasil penelitian dibagi ke dalam tiga bagian utama. Pertama, evaluasi kuantitatif terhadap kinerja model Machine Learning pada tugas deteksi anomali. Kedua, analisis holistik sistem dengan membandingkan metrik operasional antara detektor mandiri (standalone) dan kerangka terintegrasi ML+LLM. Ketiga, analisis kualitatif mengenai kemampuan agen kognitif berbasis LLM dalam menghasilkan tindakan mitigasi yang tepat serta penjelasan insiden yang mudah dipahami.

3.1. Perbandingan Model Machine Learning



Gambar 4. Diagram Perbandingan Model Machine Learning

Gambar 4. menunjukkan diagram perbandingan model Machine Learning. Hasil validasi empiris pada dataset UNSW-NB15 menunjukkan bahwa Random Forest adalah algoritma dengan performa terbaik di antara lima model yang diuji, dengan akurasi 95,99%, F1-score 95,35%, presisi 95,73%, dan recall 95,00%. Hasil ini menegaskan kemampuannya membedakan lalu lintas normal dan anomali dengan tingkat kesalahan yang sangat rendah. Sebagai perbandingan, Decision Tree mencatat akurasi 94,99% dan F1-score 94,24%, namun rentan terhadap overfitting pada data kompleks. Model Gradient Boosting dan XGBoost juga cukup kompetitif, dengan akurasi masing-masing 94,95% dan 94,66%, meski kinerja XGBoost sedikit lebih rendah dibanding metode ensemble lainnya. Sementara itu, Support Vector Machine (SVM) menunjukkan presisi tertinggi (95,54%), tetapi recall terendah (90,51%), yang membuatnya berisiko melewati sebagian serangan.

Berdasarkan perbandingan ini, Random Forest dipilih sebagai model deteksi paling optimal untuk diintegrasikan ke dalam kerangka kerja agen kognitif, terutama karena F1-score-nya yang seimbang antara presisi dan recall dalam dataset yang tidak seimbang.

3.2. Performa Deteksi DDoS menggunakan Machine Learning dan Framework Machine Learning dengan AI

Tabel 1. Perbandingan Standart Detection dan Agentic Mitigation

Metode	Throughput Packet	Memory Usage	Latency
Deteksi DDoS Machine Learning	1309	981MB	13.8 detik
Integrasi Machine Learning dengan AI untuk Mitigasi DDoS	1379	1278MB	15.2 detik

Tabel 1. menunjukkan analisis kinerja sistem secara holistik, yang membandingkan arsitektur deteksi ML mandiri dan kerangka kerja terintegrasi ML dengan LLM, mengungkapkan adanya trade-off yang terukur antara kecerdasan otonom dan overhead sumber daya. Sebagaimana disajikan dalam Tabel Perbandingan Performa, integrasi agen kognitif LLM secara inheren memperkenalkan peningkatan pada penggunaan sumber daya komputasi. Penggunaan memori tercatat meningkat dari 981 MB pada sistem deteksi ML mandiri menjadi 1278 MB pada kerangka kerja terintegrasi. Peningkatan sebesar 297 MB ini merupakan konsekuensi yang dapat diantisipasi dari pemuatan model bahasa skala besar (LLM) beserta dependensinya, yang diperlukan untuk menjalankan fungsi penalaran dan generasi bahasa. Serupa dengan itu, latensi sistem yang diukur sebagai time-to-mitigate mengalami peningkatan dari 13.8 detik menjadi 15.2 detik. Tambahan latensi sebesar 1.4 detik ini merepresentasikan waktu yang dibutuhkan oleh agen LLM untuk memproses input dari modul deteksi, menganalisis konteks ancaman, menghasilkan perintah mitigasi dalam format JSON, dan menyusun penjelasan insiden dalam bahasa alami. Meskipun terjadi peningkatan overhead, kerangka kerja terintegrasi juga menunjukkan sedikit peningkatan pada throughput paket dari 1309 menjadi 1379 paket per unit waktu. Hal ini mengindikasikan bahwa penambahan komponen LLM tidak menciptakan bottleneck yang signifikan pada pemrosesan lalu lintas jaringan, dan sistem mampu mempertahankan kinerja pemrosesan paket yang tinggi sambil menyediakan kapabilitas respons otonom. Peningkatan latensi yang relatif kecil ini dapat dianggap sebagai pertukaran yang sangat dapat diterima untuk memperoleh kemampuan mitigasi otomatis yang cerdas, yang secara drastis mengurangi jeda waktu dibandingkan intervensi manual oleh analis keamanan.

3.3. Hasil Deteksi DDoS Secara Langsung Menggunakan AI dan ML



Gambar 5. Deteksi DDoS Secara Real Time

Gambar 5. menunjukkan hasil keluaran dari sistem deteksi intrusi Suricata yang secara real-time mengidentifikasi adanya aktivitas jaringan yang berpotensi berbahaya. Pada kasus ini, sistem mendeteksi adanya indikasi serangan bertipe SYN Flood, yaitu salah satu bentuk serangan Distributed Denial of Service (DDoS) yang berupaya membanjiri server dengan permintaan koneksi palsu (TCP SYN requests) dalam jumlah besar. Tujuan utama dari serangan ini adalah untuk menguras sumber daya jaringan sehingga layanan menjadi tidak responsif atau berhenti beroperasi.

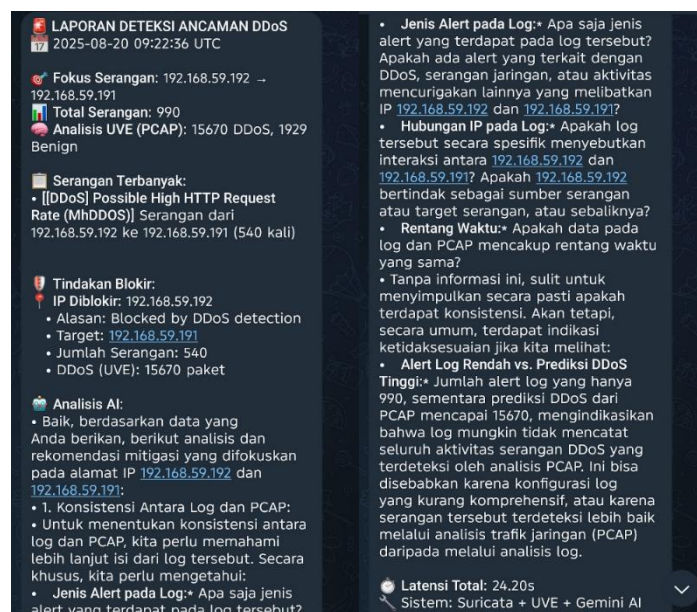
Notifikasi yang dihasilkan oleh Suricata berisi sejumlah informasi penting yang direpresentasikan secara terstruktur, meliputi:

- Jenis Peringatan (Signature) : Label “DDoS Possible SYN Flood Detected” menunjukkan bahwa Suricata mengenali pola lalu lintas jaringan yang menyerupai karakteristik serangan SYN Flood. Deteksi ini dilakukan melalui analisis pola paket TCP yang dikirim secara masif tanpa menyelesaikan proses handshake secara normal.
- Waktu Kejadian (Timestamp) : Waktu 2025-06-26T16:20:43.400053+0000 menunjukkan momen terjadinya deteksi dalam format UTC. Informasi waktu ini penting untuk keperluan korelasi insiden, pelacakan sumber serangan, serta rekonstruksi timeline aktivitas jaringan.

- c. Tingkat Keparahan (Severity) : Nilai 3 menunjukkan bahwa peringatan berada pada tingkat keparahan sedang (medium severity).
- d. Alamat Sumber dan Tujuan (Source dan Destination IP) : Alamat sumber 192.168.59.192:38888 diidentifikasi sebagai entitas yang menginisiasi lalu lintas mencurigakan, sedangkan tujuan 192.168.59.191:80 merupakan target serangan pada port 80, yang umumnya digunakan untuk layanan web (HTTP).
- e. Protokol Komunikasi (Protocol) : Jenis protokol yang digunakan adalah TCP, yang merupakan lapisan transport utama dalam komunikasi jaringan.

Secara keseluruhan, Gambar 5. merepresentasikan tahap awal dalam siklus respons otonom yang diusulkan dalam penelitian ini. Notifikasi mentah dari Suricata berfungsi sebagai pemicu awal bagi agen kognitif berbasis LLM untuk melakukan analisis lanjutan. Agen tersebut kemudian mengonversi data log yang masih mentah menjadi wawasan yang dapat ditindaklanjuti, seperti mengidentifikasi tingkat risiko, menentukan pola serangan, serta menghasilkan rekomendasi mitigasi secara otomatis. Dengan demikian, Gambar 5. memperlihatkan hubungan sinergis antara sistem deteksi tradisional (Suricata) sebagai penghasil data utama dan data *feedback* dari Random Forest yang kemudian digunakan oleh LLM sebagai data *inference*.

3.4. Hasil Deteksi DDoS Secara Langsung Menggunakan AI dan ML



Gambar 6. Hasil Analisis Integrasi AI, Saran dan Tindakan Mitigasi IP Address yang Terdeteksi DDoS

Gambar 6. menampilkan hasil akhir dari proses kognitif yang dilakukan oleh LLM, yang secara efektif menjembatani kesenjangan antara deteksi dan mitigasi. Agen LLM tidak hanya mengkonfirmasi serangan dengan menganalisis data PCAP mengidentifikasi 15.670 paket DDoS dibandingkan 1.929 paket benign tetapi juga secara otonom melakukan sintesis informasi dari berbagai sumber (log dan analisis PCAP). Kapabilitas penalaran agen LLM terwujud dalam kemampuannya untuk tidak hanya mengeksekusi tindakan mitigasi seperti pemblokiran alamat IP sumber serangan 192.168.59.192 dengan alasan Blocked by DDoS detection tetapi juga menyajikan analisis mendalam dalam format bahasa alami. Analisis ini mencakup evaluasi konsistensi antara data log dan data PCAP, bahkan menyoroti potensi diskrepansi, seperti jumlah peringatan log yang lebih rendah dibandingkan prediksi serangan dari PCAP, dan memberikan hipotesis logis atas fenomena tersebut. Kemampuan untuk memberikan penjelasan kontekstual ini merupakan nilai tambah yang signifikan, mengubah sistem dari sekadar

mekanisme pertahanan "kotak hitam" menjadi mitra analitik yang transparan bagi analis keamanan. Latensi total untuk satu siklus penuh dari deteksi hingga mitigasi dan pelaporan, seperti yang tercatat pada contoh insiden ini, adalah 24.20 detik, memberikan bukti empiris yang kuat mengenai efektivitas kerangka kerja *closed-loop* dalam mengotomatiskan siklus pertahanan siber secara end-to-end.

4. DISKUSI

Analisis empiris atas kerangka kerja yang diusulkan memberikan wawasan pada dua level, yaitu kinerja model deteksi dan efektivitas sistem *closed-loop* secara keseluruhan. Uji coba pada dataset UNSW-NB15 menempatkan Random Forest (RF) sebagai model deteksi terbaik dengan akurasi 95,99% dan F1-score 95,35%. Keunggulan RF berasal dari sifat ensemble-nya yang mampu mengurangi varians dan meningkatkan generalisasi pada data kompleks. Decision Tree memberikan hasil kompetitif, tetapi lebih rentan terhadap overfitting, sedangkan SVM mencatat presisi tertinggi (95,54%) namun recall terendah (90,51%), sehingga berisiko melewatkan serangan [5], [46], [47].

Dari sisi operasional, integrasi agen LLM menghasilkan trade-off yang terukur. Konsumsi memori meningkat sekitar 297 MB dan latensi time-to-mitigate bertambah 1,4 detik. Namun, peningkatan ini tidak menimbulkan bottleneck, throughput paket justru naik dari 1309 menjadi 1379 paket per unit waktu. Studi kasus juga menunjukkan siklus end-to-end dari deteksi hingga mitigasi selesai dalam 24,20 detik, membuktikan kelayakan sistem ini dibandingkan respons manual [29]. Kontribusi utama kerangka kerja ini adalah menjembatani deteksi berbasis data (ML/IDS) dengan mitigasi melalui agen kognitif berbasis prompt terstruktur. Tidak seperti literatur SIEM yang lebih menekankan korelasi log [23] atau penelitian LLM yang fokus pada analisis strategis [30], [31], sistem ini mampu melangkah lebih jauh dari “peringatan” menuju “tindakan tervalidasi”. Agen LLM terbukti dapat menggabungkan log Suricata dengan analisis PCAP untuk mengonfirmasi pola serangan SYN Flood serta memberikan penjelasan kontekstual yang menjadikannya co-pilot analitik, bukan sekadar generator perintah.

Dibandingkan penelitian sebelumnya, pendekatan ini lebih maju karena mengintegrasikan deteksi dan mitigasi dalam satu mekanisme *closed-loop*. Sebelumnya, model seperti SVM, Decision Tree, atau Random Forest hanya digunakan untuk klasifikasi serangan [5], sementara mitigasi masih dipisahkan. Sistem ini menutup celah tersebut dengan menghasilkan serta memvalidasi tindakan mitigasi [29], [48], [49].

Maka dari itu, kami mengusulkan pendekatan integrasi antara Machine Learning dan LLM dalam satu framework, dengan LLM berperan sebagai executor dan Random Forest sebagai komponen guidance yang memberikan arahan dalam proses pengambilan keputusan. Pendekatan ini secara efektif menonjolkan keunggulan Random Forest dalam tahap deteksi serta peran LLM sebagai agen mitigasi kontekstual dalam sistem *closed-loop*, yang menjembatani kesenjangan antara deteksi berbasis ML dan tindakan otomatis.

Hasil pengujian menunjukkan adanya peningkatan kecil pada latensi (sekitar 1,4 detik), namun peningkatan tersebut sepadan dengan kemampuan mitigasi otonom yang diperoleh. Jika dibandingkan dengan penelitian ML/IDS sebelumnya maupun berbagai kerangka kerja LLM di bidang keamanan siber seperti [36], [50], [51], [52] pendekatan ini menunjukkan keunggulan operasional melalui integrasi validasi sintaks dan logika sebelum eksekusi tindakan. Meskipun demikian, penelitian lanjutan tetap diperlukan untuk mengoptimalkan trade-off antara latensi dan kinerja, misalnya melalui strategi kompresi model dengan metode kuantisasi, *edge-level deployment*, serta *prompt caching* agar sistem tetap responsif tanpa mengorbankan akurasi mitigasi.

Dari segi keandalan, sistem dilengkapi dengan validasi sintaks dan logika sebelum aturan dijalankan, menjaga integritas dan keamanan. Mekanisme ini dapat diperkuat dengan allowlist/denylist, pembatasan akses, dan jejak audit melalui Wazuh [23]. Namun, keterbatasan tetap ada, seperti

keterwakilan dataset UNSW-NB15 yang tidak langsung merepresentasikan trafik spesifik (misalnya IoT) serta kebutuhan metrik tambahan (ROC-AUC, PR-AUC, MCC) untuk memperkaya evaluasi. Untuk pengembangan lebih lanjut, beberapa aspek perlu diperkuat, termasuk pengendalian risiko over-automation, perlindungan terhadap prompt injection, serta mekanisme adaptasi *concept drift*. Validasi pada trafik nyata juga penting untuk memastikan konsistensi *time-to-mitigate*.

5. KESIMPULAN

Studi ini sukses mendemonstrasikan integrasi antara kecepatan deteksi Machine Learning (ML) dan kecakapan pengambilan keputusan Large Language Model (LLM) untuk mitigasi serangan secara otonom. Dengan model Random Forest yang mencapai akurasi 95,99%, sistem closed-loop ini mampu menuntaskan siklus respons dalam 24.20 detik, meskipun terdapat trade-off latensi sebesar 1.4 detik akibat peran LLM. Temuan ini menegaskan bahwa sinergi kedua teknologi tersebut efektif menjembatani kesenjangan antara deteksi dan tindakan otomatis, sekaligus meningkatkan transparansi serta situational awareness [29]. Penelitian ini membuka peluang eksplorasi lebih lanjut, terutama dalam upaya meminimalkan trade-off latensi. Fokus pada metode seperti kuantisasi model LLM dan prompt caching dapat menjadi langkah berikutnya yang menjanjikan, diikuti dengan validasi sistem menggunakan trafik nyata untuk memperkuat skalabilitas dan keandalannya.

REFERENSI

- [1] S. M. Syifa Munawarah, Kurniabudi, and Eko Arip Winanto, "Deteksi Serangan DDoS SYN Flood Pada Jaringan Internet of Things (IoT) Menggunakan Metode Deep Neural Network (DNN)," *J. Inform. Dan Rekayasa Komputer JAKAKOM*, vol. 4, no. 1, pp. 982–991, Apr. 2024, doi: 10.33998/jakakom.2024.4.1.1710.
- [2] D. Firdaus, F. Fahira, and R. Rianti, "DETEKSI ANOMALI DAN SERANGAN LOW RATE DDOS DALAM LALU LINTAS JARINGAN MENGGUNAKAN NAIVE BAYES," *Naratif J. Nas. Ris. Apl. Dan Tek. Inform.*, vol. 5, no. 2, pp. 140–148, Dec. 2023, doi: 10.53580/naratif.v5i2.208.
- [3] M. Ilman Aqilaa, D. Firdaus, and N. Naofal, "Identifikasi Serangan Lowrate Distributed Denial Of Services Dalam Jaringan Dengan Menggunakan Algoritma Adaboost," *Simpatik J. Sist. Inf. Dan Inform.*, vol. 3, no. 1, pp. 34–41, June 2023, doi: 10.31294/simpatik.v3i1.1829.
- [4] S. Joses, S. Quinevera, R. Mardianto, D. Yulvida, and A. M. Shiddiqi, "Pendekatan Metode Ensemble Learning untuk Deteksi Serangan DDoS menggunakan Soft Voting Classifier," *J. Edukasi Dan Penelit. Inform. JEPIN*, vol. 10, no. 1, p. 79, Apr. 2024, doi: 10.26418/jp.v10i1.73241.
- [5] Y. Xie, "Machine learning-based DDoS detection for IoT networks," *Appl. Comput. Eng.*, vol. 29, no. 1, pp. 99–107, Dec. 2023, doi: 10.54254/2755-2721/29/20230972.
- [6] A. A. Alahmadi *et al.*, "DDoS Attack Detection in IoT-Based Networks Using Machine Learning Models: A Survey and Research Directions," *Electronics*, vol. 12, no. 14, p. 3103, July 2023, doi: 10.3390/electronics12143103.
- [7] M. A. Owaid and A. S. Hammoodi, "Evaluating Machine Learning and Deep Learning Models for Enhanced DDoS Attack Detection," *Math. Model. Eng. Probl.*, vol. 11, no. 2, pp. 493–499, Feb. 2024, doi: 10.18280/mmep.110221.
- [8] M. S. Raza, M. N. A. Sheikh, I.-S. Hwang, and M. S. Ab-Rahman, "Feature-Selection-Based DDoS Attack Detection Using AI Algorithms," *Telecom*, vol. 5, no. 2, pp. 333–346, Apr. 2024, doi: 10.3390/telecom5020017.
- [9] B. Goparaju and Dr. B. S. Rao, "A DDoS Attack Detection using PCA Dimensionality Reduction and Support Vector Machine," *Int. J. Commun. Netw. Inf. Secur. IJCNIS*, vol. 14, no. 1s, pp. 01–08, Jan. 2023, doi: 10.17762/ijcnis.v14i1s.5586.
- [10] Zerir Hasan Sahosh, Azraf Faheem, Marzana Bintay Tuba, Md. Istiaq Ahmed, and Syed Anika Tasnim, "A Comparative Review on DDoS Attack Detection Using Machine Learning Techniques," *Malays. J. Sci. Adv. Technol.*, pp. 75–83, Mar. 2024, doi: 10.56532/mjsat.v4i2.208.

-
- [11] A. Purnomo, A. Kurniasih, A. Nuarminah, and S. Hartati, "Peran Artificial Intelligence dalam Deteksi Dini Ancaman Keamanan Jaringan," *J. Minfo Polgan*, vol. 13, no. 2, pp. 2044–2048, Dec. 2024, doi: 10.33395/jmp.v13i2.14356.
- [12] K. Handayani and E. Erni, "PENERAPAN LIGHT GRADIENT BOOSTING DALAM PREDIKSI RASIO KLIK TAYANG," *JATI J. Mhs. Tek. Inform.*, vol. 7, no. 1, pp. 13–18, Jan. 2023, doi: 10.36040/jati.v7i1.6010.
- [13] B. Pernama and H. D. Purnomo, "Analisis Risiko Pinjaman dengan Metode Support Vector Machine, Artificial Neural Network dan Naïve Bayes," *J. JTIK J. Teknol. Inf. Dan Komun.*, vol. 7, no. 1, pp. 92–99, Jan. 2023, doi: 10.35870/jtik.v7i1.693.
- [14] N. Mamuriyah, S. E. Prasetyo, and A. O. Sijabat, "Rancangan Sistem Keamanan Jaringan dari serangan DDoS Menggunakan Metode Pengujian Penetrasi," *J. Teknol. Dan Sist. Inf. Bisnis*, vol. 6, no. 1, pp. 162–167, Jan. 2024, doi: 10.47233/jteksis.v6i1.1124.
- [15] I. Rahmadaniar, D. A. A. Tondang, B. S. Fernando, and A. Setiawan, "Implementasi Firewall Menggunakan Iptables untuk Melindungi Server dari Serangan DDoS," *J. Internet Softw. Eng.*, vol. 1, no. 3, p. 10, June 2024, doi: 10.47134/pjise.v1i3.2564.
- [16] L. Bagdadi and B. Messabih, "Distributed denial of service attacks classification system using features selection and ensemble techniques," *Indones. J. Electr. Eng. Comput. Sci.*, vol. 34, no. 3, p. 1868, June 2024, doi: 10.11591/ijeecs.v34.i3.pp1868-1878.
- [17] E. DeniZ and S. Serttas, "Deep learning-based distributed denial of service detection system in the cloud network," *J. Sci. Rep.-A*, no. 055, pp. 16–33, Dec. 2023, doi: 10.59313/jsr-a.1333839.
- [18] USA and A. Aluwala, "Mitigating DDoS Attacks via AI Detection and SDN Response," *J. Artif. Intell. Cloud Comput.*, vol. 2, no. 4, pp. 1–5, Dec. 2023, doi: 10.47363/JAICC/2023(2)E151.
- [19] W. I. Khedr, A. E. Gouda, and E. R. Mohamed, "FMDADM: A Multi-Layer DDoS Attack Detection and Mitigation Framework Using Machine Learning for Stateful SDN-Based IoT Networks," *IEEE Access*, vol. 11, pp. 28934–28954, 2023, doi: 10.1109/ACCESS.2023.3260256.
- [20] J. Wang, L. Wang, and R. Wang, "A Method of DDoS Attack Detection and Mitigation for the Comprehensive Coordinated Protection of SDN Controllers," *Entropy*, vol. 25, no. 8, p. 1210, Aug. 2023, doi: 10.3390/e25081210.
- [21] Jain(Deemed-to-be University) and S. Pandey, "Surveying Emerging Trends in DDoS Defense," *INTERANTIONAL J. Sci. Res. Eng. Manag.*, vol. 08, no. 05, pp. 1–5, May 2024, doi: 10.55041/IJSREM34483.
- [22] M. Nas, F. Ulfiah, and U. Putri, "Analisis Sistem Security Information and Event Management (SIEM) Aplikasi Wazuh pada Dinas Komunikasi Informatika Statistik dan Persandian Sulawesi Selatan," *J. Teknol. Elekterika*, vol. 20, no. 2, p. 92, Nov. 2023, doi: 10.31963/elekterika.v20i2.4536.
- [23] W. P. Putra, R. Burjulius, M. A. Al Hilmi, and A. Sumarudin, "Implementasi Sistem Manajemen Log untuk Penanggulangan Serangan Server dengan SIEM," *IKRA-ITH Inform. J. Komput. Dan Inform.*, vol. 8, no. 3, pp. 23–30, Oct. 2024, doi: 10.37817/ikraith-informatika.v8i3.4359.
- [24] M. R. T. Hidayat, N. Widiyasono, and R. Gunawan, "OPTIMASI DETEKSI MALWARE PADA SIEM WAZUH MELALUI INTEGRASI CYBER THREAT INTELLIGENCE DENGAN MISP DAN DFIR-IRIS," *J. Inform. Dan Tek. Elektro Terap.*, vol. 13, no. 1, Jan. 2025, doi: 10.23960/jitet.v13i1.5686.
- [25] P. Purwanti, "Visualisasi Data Cyber Security Attack Dengan Fitur Prediksi Serangan Dan Mitigasi Risiko: Perspektif Generative Gemini AI," *J. Minfo Polgan*, vol. 13, no. 2, pp. 2340–2350, Jan. 2025, doi: 10.33395/jmp.v13i2.14453.
- [26] M. Q. Syahputra, D. R. Akbi, and D. Risqiwati, "Deteksi Dan Mitigasi Serangan DDoS Pada Software Defined Network Menggunakan Algoritma Decision Tree," *J. Repos.*, vol. 2, no. 11, p. 1491, Dec. 2020, doi: 10.22219/repositor.v2i11.795.
- [27] A. Tanton, M. T. A. Zaen, and Y. Yuliadi, "PENERAPAN VLAN DALAM MITIGASI SERANGAN DDOS PADA OLT HSGQ DAN ROUTER MIKROTIK," *J. Inform. Teknol. Dan Sains Jinteks*, vol. 6, no. 2, pp. 298–305, June 2024, doi: 10.51401/jinteks.v6i2.4137.
- [28] D. Aryani and E. D. Absharina, "Mitigasi Risiko Cybercrime Terhadap Keamanan Sistem Komputasi Awan pada Perusahaan," *J. Cakrawala Akad.*, vol. 1, no. 4, pp. 1365–1373, Dec. 2024, doi: 10.70182/JCA.v1i4.27.
-

-
- [29] V. Gustina Dm and A. Ananda, "Kecerdasan Buatan untuk Security Orchestration, Automation and Response: Tinjauan Cakupan," *J. Komput. Terap.*, vol. 10, no. 1, pp. 36–47, June 2024, doi: 10.35143/jkt.v10i1.6247.
- [30] A. Januantoro and S. Supangat, "DETEKSI SERANGAN JARINGAN KOMPUTER BERBASIS SNORT DENGAN INTEGRASI NOTIFIKASI REAL-TIME MELALUI TELEGRAM," *J. Mnemon.*, vol. 8, no. 1, pp. 100–105, Mar. 2025, doi: 10.36040/mnemonic.v8i1.12175.
- [31] T. Syaiful Huda and S. Subektiningsih, "Analisis Keamanan Jaringan Komputer Menggunakan Metode IDS dan IPS dengan Notifikasi Telegram: Computer Network Security Analysis Using IDS and IPS Methods with Telegram Notifications," *Indones. J. Comput. Sci.*, vol. 13, no. 1, Jan. 2024, doi: 10.33022/ijcs.v13i1.3505.
- [32] H. Alturkistani and S. Chuprat, "Artificial Intelligence and Large Language Models in Advancing Cyber Threat Intelligence: A Systematic Literature Review," Nov. 27, 2024, *In Review*. doi: 10.21203/rs.3.rs-5423193/v1.
- [33] Z. Li, X. Wang, and Q. Zhang, "Evaluating the Quality of Large Language Model-Generated Cybersecurity Advice in GRC Settings," June 21, 2024, *In Review*. doi: 10.21203/rs.3.rs-4608321/v1.
- [34] S. M. Taghavi and F. Feyzi, "Using Large Language Models to Better Detect and Handle Software Vulnerabilities and Cyber Security Threats," May 21, 2024, *In Review*. doi: 10.21203/rs.3.rs-4387414/v1.
- [35] F. Wang, "Using Large Language Models to Mitigate Ransomware Threats," Nov. 08, 2023, *Open Science Framework*. doi: 10.31219/osf.io/mzsnh.
- [36] T. Wang, X. Xie, L. Zhang, C. Wang, L. Zhang, and Y. Cui, "ShieldGPT: An LLM-based Framework for DDoS Mitigation," in *Proceedings of the 8th Asia-Pacific Workshop on Networking*, Sydney Australia: ACM, Aug. 2024, pp. 108–114. doi: 10.1145/3663408.3663424.
- [37] I. K. Prasetya, M. Ahsan, M. Mashuri, and M. H. Lee, "Multivariate Robust MRCD Based Hotelling's T2 Control Chart with Bootstrap Control Limit for Intrusion Detection," Dec. 27, 2023, *Computer Science and Mathematics*. doi: 10.20944/preprints202312.2062.v1.
- [38] X. Tong, C. Zhang, J. Wang, Z. Zhao, and Z. Liu, "Dark-Forest: Analysis on the Behavior of Dark Web Traffic via DeepForest and PSO Algorithm," *Comput. Model. Eng. Sci.*, vol. 135, no. 1, pp. 561–581, 2023, doi: 10.32604/cmescs.2022.022495.
- [39] Y. S. Kuruba Manjunath, "Network Traffic Classification for Internet of Things Based on Deep Learning Models". Toronto Metropolitan University, 29-Aug-2023, doi: 10.32920/24050748.v1.
- [40] C.-N. Hang, P.-D. Yu, R. Morabito, and C.-W. Tan, "Large Language Models Meet Next-Generation Networking Technologies: A Review," *Future Internet*, vol. 16, no. 10, p. 365, Oct. 2024, doi: 10.3390/fi16100365.
- [41] Y. Alotaibi and M. Ilyas, "Ensemble-Learning Framework for Intrusion Detection to Enhance Internet of Things' Devices Security," *Sensors*, vol. 23, no. 12, p. 5568, June 2023, doi: 10.3390/s23125568.
- [42] M. Hebbaka Shivanajappa, R. Maidanahalli Seetharamaiah, B. Viswaraju Sai, A. Jakkannahally Siddegowda, and V. Kuppanna Rajuk, "An Intrusion Detection System against RPL-based Routing Attacks for IoT Networks," *Indones. J. Electr. Eng. Comput. Sci.*, vol. 34, no. 2, p. 1324, May 2024, doi: 10.11591/ijeecs.v34.i2.pp1324-1335.
- [43] T. T. Atmojo and Y. N. Kunang, "Machine Learning-Based E-Archive for Archives Management of South Sumatra Province," *J. Inf. Syst. Inform.*, vol. 5, no. 4, pp. 1491–1507, Dec. 2023, doi: 10.51519/journalisi.v5i4.566.
- [44] V. Z. Mohale and I. C. Obagbuwa, "A systematic review on the integration of explainable artificial intelligence in intrusion detection systems to enhancing transparency and interpretability in cybersecurity," *Front. Artif. Intell.*, vol. 8, p. 1526221, Jan. 2025, doi: 10.3389/frai.2025.1526221.
- [45] C. Chen *et al.*, "Application of GA-WELM Model Based on Stratified Cross-Validation in Intrusion Detection," *Symmetry*, vol. 15, no. 9, p. 1719, Sept. 2023, doi: 10.3390/sym15091719.
- [46] J. Shaikh, Y. A. Butt, and H. F. Naqvi, "Effective Intrusion Detection System Using Deep Learning for DDoS Attacks," *Asian Bull. Big Data Manag.*, vol. 4, no. 1, Mar. 2024, doi: 10.62019/abbdm.v4i1.113.
-

-
- [47] C.-S. Shieh, F.-A. Ho, M.-F. Horng, T.-T. Nguyen, and P. Chakrabarti, "Open-Set Recognition in Unknown DDoS Attacks Detection With Reciprocal Points Learning," *IEEE Access*, vol. 12, pp. 56461–56476, 2024, doi: 10.1109/ACCESS.2024.3388149.
- [48] U. D. Maiwada, K. U. Danyaro, A. Sarlan, M. S. Liew, A. Taiwo, and U. I. Audi, "Energy efficiency in 5G systems: A systematic literature review," *Int. J. Knowl.-Based Intell. Eng. Syst.*, vol. 28, no. 1, pp. 93–132, Mar. 2024, doi: 10.3233/KES-230061.
- [49] H. M. S. Saleeh, H. Marouane, and A. Fakhfakh, "A Novel Deep Learning Approach for Detecting Types of Attacks in the NSL-KDD Dataset," *Babylon. J. Netw.*, vol. 2024, pp. 171–181, Sept. 2024, doi: 10.58496/BJN/2024/017.
- [50] J. F. Loevenich, E. Adler, T. Hürten, F. Spelter, D. Roncevic, and R. R. F. Lopes, "Automating Cyber Threat Intelligence and Attack Chain Generation using Cyber Security Knowledge Graphs and Large Language Models," in *2025 International Conference on Military Communication and Information Systems (ICMCIS)*, Oerias, Portugal: IEEE, May 2025, pp. 1–10. doi: 10.1109/ICMCIS64378.2025.11047951.
- [51] K. N. Meda *et al.*, "Integrating Prompt Structures Using LLM Embeddings for Cybersecurity Threats," in *Proceedings of the 2025 ACM Southeast Conference*, Southeast Missouri State University Cape Girardeau MO USA: ACM, Apr. 2025, pp. 180–187. doi: 10.1145/3696673.3723069.
- [52] M. T. Sandaruwan, J. Wijayanayake, and J. Senanayake, "Integrating Large Language Models for Automated Vulnerability Scanning and Reporting in Network Hosts," in *2025 International Research Conference on Smart Computing and Systems Engineering (SCSE)*, Colombo, Sri Lanka: IEEE, Apr. 2025, pp. 1–7. doi: 10.1109/SCSE65633.2025.11031059.