

BAB I PENDAHULUAN

1.1 Latar Belakang Masalah

Serangan *Zero-Day* kini menjadi ancaman paling kritis dalam keamanan siber, sebagaimana dibuktikan laporan *Mandiant* dan *Google Threat Analysis Group* (TAG) tahun 2023. Tercatat 97 eksploitasi *zero-day*, melonjak 50% dibandingkan tahun sebelumnya, yang mayoritas menargetkan platform konsumen dan teknologi perusahaan [1]. Tren peningkatan drastis ini menuntut adanya mekanisme pertahanan yang lebih cerdas daripada sekadar mengandalkan pembaruan perangkat lunak yang seringkali terlambat.

Sifat serangan *zero-day* yang memanfaatkan celah keamanan tak dikenal membuat vendor atau pengembang tidak memiliki *patch* solusi saat serangan pertama terjadi [1]. Hal ini membuat sistem rentan menjadi sasaran, dapat menyebabkan kerugian atau gangguan operasional [2]. Celah waktu antara eksploitasi awal dan perbaikan sistem inilah yang menjadi titik lemah utama yang dieksploitasi oleh peretas.

Metode deteksi tradisional seringkali gagal mengenali ancaman ini karena hanya efektif terhadap pola serangan yang sudah tersimpan dalam basis data [3]. Ketidakmampuan mengenali pola baru (*unseen patterns*) memaksa pergeseran menuju pendekatan yang lebih dinamis. Oleh karena itu, sistem deteksi intrusi modern beralih pada pendekatan berbasis *Machine Learning* (ML) yang secara teoritis mampu mempelajari anomali perilaku jaringan secara mandiri.

Penerapan ML pada data jaringan modern sering terkendala dimensi tinggi yang menurunkan akurasi [3], [4]. Untuk mengatasinya, penelitian ini menggunakan *Autoencoder* (AE) sebagai pengekstraksi fitur dan membandingkan efektivitas arsitektur hibrida AE dengan *Random Forest* (RF) dan *Multi-Layer Perceptron* (MLP). Penelitian ini diharapkan dapat menentukan model terbaik dan membuktikan signifikansi fitur *Reconstruction Error* (RE) dalam mendeteksi serangan *zero-day* secara akurat dan efisien.

1.2 Rumusan Masalah

1. Bagaimana perbandingan tingkat deteksi serangan *Zero-Day* antara model yang menggunakan fitur asli (*raw features*) dengan model yang menggunakan fitur hasil ekstraksi *Autoencoder*?
2. Manakah yang menghasilkan performa lebih baik (akurasi dan *throughput*) antara arsitektur *Autoencoder - Random Forest* (AE-RF) dan *Autoencoder - Multi Layer Perceptron* (AE-MLP) dalam mendeteksi serangan *Zero-Day*?
3. Bagaimana pengaruh penambahan fitur *Reconstruction Error* (RE) terhadap tingkat deteksi serangan *Zero-Day* pada model *Autoencoder - Random Forest* dibandingkan dengan model *Autoencoder-Multi Layer Perceptron*?

1.3 Batasan Masalah

1. Penelitian ini berfokus pada perbandingan performa antara dua arsitektur hibrida: *Autoencoder - Random Forest* dan *Autoencoder - MLP*. Penelitian ini tidak akan membandingkan dengan *classifier* lain (seperti SVM, K-NN) atau arsitektur *deep learning* lain (seperti CNN atau LSTM).
2. *Dataset* yang digunakan terbatas pada CIC-IDS-2018 dalam format *Parquet*. Simulasi serangan *zero-day* dilakukan dengan menahan lima kelas serangan (*FTP-BruteForce*, *SSH-Bruteforce*, *Brute Force-Web*, *Brute Force-XSS*, dan *SQL Injection*) dari data latih sehingga kelas-kelas tersebut tidak pernah dilihat model saat pelatihan.
3. Pengukuran kecepatan prediksi (*throughput*) dilakukan secara *offline* pada file statis dengan ukuran *chunk* 25 000 baris. Penelitian ini tidak mencakup pengujian *real-time* pada jaringan langsung maupun implementasi pada perangkat *embedded/edge device*.
4. Fokus utama adalah pada pengaruh penambahan fitur *Reconstruction Error* (RE) sebagai fitur ke-17 dan perbandingan *classifier* RF vs MLP. Optimasi *hyperparameter* yang ekstensif (*grid search*, *random search*, *Bayesian*

optimization) serta eksperimen dengan dimensi laten *Autoencoder* yang berbeda-beda berada di luar cakupan penelitian ini.

5. Evaluasi performa difokuskan pada metrik efektivitas deteksi (*Recall* dan *False Positive Rate*) berdasarkan *confusion matrix* serta efisiensi waktu (*throughput*). Metrik efisiensi perangkat keras seperti konsumsi memori, konsumsi daya, atau *latency per flow* tidak diukur secara spesifik.

1.4 Tujuan Penelitian

Tujuan yang akan dicapai dalam penelitian ini

1. Menganalisis urgensi penggunaan *Autoencoder* dengan membandingkan performa model berbasis fitur hasil ekstraksi AE melawan model *baseline* (fitur asli).
2. Membandingkan performa *classifier* RF dan MLP dalam arsitektur hibrida berdasarkan metrik tingkat deteksi, False Positive Rate (FPR), dan kecepatan *throughput*.
3. Menganalisis dampak penambahan fitur *Reconstruction Error* (RE) terhadap performa deteksi serangan Zero-Day pada model *Autoencoder - Random Forest* dibandingkan dengan model *Autoencoder - Multi Layer Perceptron*.

1.5 Manfaat Penelitian

1. Manfaat Teoritis:
 - a. Menyumbangkan bukti empiris mengenai perbandingan performa arsitektur hibrida *Autoencoder-Random Forest* dan *Autoencoder-Multi Layer Perceptron* pada tugas deteksi serangan *zero-day*, sehingga dapat menjadi referensi bagi penelitian lanjutan di bidang sistem deteksi intrusi berbasis pembelajaran mendalam.
 - b. Memberikan wawasan ilmiah tentang kontribusi fitur *Reconstruction Error* (RE) sebagai *discriminator* tambahan dalam meningkatkan kemampuan deteksi ancaman yang belum pernah dilihat sebelumnya

(*zero-day attacks*).

- c. Menghasilkan analisis mendalam mengenai efektivitas strategi penanganan ketidakseimbangan kelas (*class imbalance*) dan ekstraksi fitur tak-terawasi pada *dataset* lalu lintas jaringan berdimensi tinggi.

2. Manfaat Praktis:

- a. Memberikan rekomendasi berbasis data kepada praktisi keamanan siber, pengelola SOC (*Security Operation Center*), dan pengembang NIDS dalam memilih arsitektur yang paling sesuai antara *Autoencoder - Random Forest* (lebih stabil pada CPU) atau *Autoencoder - Multi Layer Perceptron* dengan *Reconstruction Error* (deteksi *zero-day* tertinggi dengan akselerasi GPU).
- b. Menyediakan nilai *benchmark throughput* (baris/detik) yang realistis pada *dataset* CIC-IDS-2018, sehingga dapat digunakan sebagai acuan dalam perencanaan kapasitas sistem deteksi intrusi di lingkungan produksi.
- c. Menghasilkan metodologi teruji yang dapat diadopsi oleh peneliti atau institusi lain yang menggunakan *dataset* publik CIC-IDS-2018 atau *dataset* serupa.

1.6 Sistematika Penulisan

BAB I PENDAHULUAN

Bab ini memuat latar belakang masalah terkait ancaman serangan *zero-day* yang terus meningkat, rumusan masalah, batasan masalah, tujuan penelitian, manfaat penelitian baik secara teoritis maupun praktis, serta sistematika penulisan skripsi.

BAB II TINJAUAN PUSTAKA

Bab ini menguraikan landasan teori dan studi literatur yang meliputi konsep dasar keamanan jaringan, sistem deteksi intrusi berbasis jaringan (NIDS), serangan *zero-day*, *dataset* CIC-IDS-2018, teknik ekstraksi fitur dengan *Autoencoder*, klasik, prinsip kerja *Random Forest* dan *Multi Layer Perceptron*, serta tinjauan terhadap

penelitian terdahulu yang relevan pada periode 2020–2025.

BAB III METODE PENELITIAN

Bab ini menjelaskan secara rinci objek penelitian (*dataset* CIC-IDS-2018), alur penelitian secara keseluruhan, spesifikasi perangkat keras dan perangkat lunak yang digunakan, tahapan *preprocessing* dan konversi ke format *Parquet*, strategi simulasi *zero-day attack*, desain dan arsitektur semua model eksperimen (*Autoencoder - Random Forest*, *Autoencoder - Multi Layer Perceptron*, *ablation study*), serta prosedur pelatihan, evaluasi, dan pengukuran performa.

BAB IV HASIL DAN PEMBAHASAN

Bab ini menyajikan hasil eksperimen secara lengkap, meliputi performa masing-masing model terhadap serangan *known* dan *zero-day*, analisis *confusion matrix*, visualisasi t-SNE dan distribusi *Reconstruction Error*, perbandingan *throughput*, pembahasan *ablation study*, serta interpretasi temuan utama dibandingkan dengan *state-of-the-art* terkini.

BAB V PENUTUP

Bab ini merangkum kesimpulan yang menjawab seluruh rumusan masalah serta memberikan saran pengembangan lebih lanjut bagi penelitian berikutnya maupun implementasi praktis di industri keamanan siber.