

BAB V PENUTUP

5.1 Kesimpulan

Berdasarkan hasil eksperimen dan analisis yang telah dilakukan dalam membandingkan arsitektur *Autoencoder-Random Forest* (AE-RF) dan *Autoencoder-Multi Layer Perceptron* (AE-MLP) untuk deteksi serangan *zero-day*, dapat ditarik kesimpulan sebagai berikut:

- a. Implementasi ekstraksi fitur menggunakan *Autoencoder* terbukti secara signifikan lebih efektif dibandingkan penggunaan fitur asli (*raw features*) dalam mendeteksi serangan *Zero-Day*. Pada model *Baseline* yang menggunakan 78 fitur asli, sistem mengalami kegagalan total dengan tingkat deteksi sebesar 0%, karena model mengalami *overfitting* pada dimensi tinggi dan gagal mengenali anomali baru. Sebaliknya, penggunaan 16 fitur laten (*bottleneck*) yang dihasilkan *Autoencoder* berhasil mereduksi *noise* dan menonjolkan karakteristik esensial data, sehingga memungkinkan model klasifikasi mengenali pola serangan yang belum pernah dipelajari sebelumnya.
- b. Dalam perbandingan antara dua arsitektur *hybrid*, model berbasis *Autoencoder-Multi Layer Perceptron* (AE-MLP) menunjukkan performa yang jauh lebih unggul dibandingkan *Autoencoder-Random Forest* (AE-RF), baik dari sisi akurasi maupun efisiensi komputasi. Model *Autoencoder - Multi Layer Perceptron* yang dioptimasi dengan akselerasi GPU mampu mencapai tingkat deteksi hingga 99,84% dengan *throughput* tertinggi sebesar 178.917 baris per detik. Sementara itu, model *Autoencoder - Random Forest* mampu menyamai performa model terbaik pada deteksi serangan protokol terbuka seperti FTP (100%), namun gagal mempertahankan konsistensi tersebut pada serangan terenkripsi seperti SSH karena keterbatasan algoritma berbasis pohon dalam memisahkan anomali halus di ruang laten, dan memiliki kecepatan pemrosesan yang lebih rendah, yaitu 78.488 baris per detik karena keterbatasan eksekusi pada CPU.

- c. Penambahan fitur *Reconstruction Error* (RE) memberikan dampak yang sangat kontras antara kedua arsitektur. Pada model *Autoencoder - Random Forest*, integrasi RE terbukti tidak efektif karena tingkat deteksi stagnan di angka 51,22% dan justru menurunkan kecepatan *throughput*. Sebaliknya, pada model *Autoencoder - Multi Layer Perceptron*, fitur RE menjadi faktor determinan yang meningkatkan performa secara drastis dari 44,19% menjadi 99,84%. Hal ini membuktikan bahwa arsitektur jaringan saraf tiruan (MLP) mampu mengeksplorasi informasi selisih rekonstruksi sebagai sinyal anomali serangan Zero-Day jauh lebih baik dibandingkan algoritma berbasis pohon (*tree-based*).

5.2 Saran

Berdasarkan temuan dan keterbatasan dalam penelitian ini, beberapa saran yang dapat diajukan untuk pengembangan penelitian selanjutnya adalah:

- a. Selama proses pengerjaan, kendala utama adalah terbatasnya kapasitas memori (RAM) dan seringnya terjadi reset *runtime* karena volume data yang sangat besar. Peneliti selanjutnya disarankan menggunakan perangkat dengan spesifikasi RAM yang lebih tinggi atau layanan *cloud computing* berbayar agar proses pengolahan data ribuan baris dapat berjalan lebih stabil tanpa kendala teknis.
- b. Penelitian ini berfokus pada serangan *Brute Force* (FTP dan SSH). Disarankan bagi peneliti berikutnya untuk menguji model yang sama pada jenis serangan lain yang tersedia di *dataset* CIC-IDS-2018, seperti serangan *DDoS* atau *Infiltration*, untuk melihat apakah fitur *Reconstruction Error* memberikan efektivitas yang sama kuatnya pada jenis serangan tersebut.
- c. Meskipun model saat ini sudah sangat cepat, penelitian selanjutnya dapat mencoba mengurangi jumlah fitur asli (dari 78 fitur) menjadi hanya fitur-fitur yang paling berpengaruh saja sebelum dimasukkan ke *Autoencoder*. Hal ini bertujuan untuk melihat apakah sistem bisa bekerja lebih ringan lagi tanpa mengurangi tingkat deteksi 99%.

- d. Mengingat tingginya sensitivitas model *Autoencoder - Multi Layer Perceptron* dengan RE yang memicu banyak alarm palsu, penelitian masa depan sebaiknya tidak hanya mengandalkan MLP untuk mempelajari bobot *Reconstruction Error* (RE) secara implisit. Disarankan untuk menerapkan mekanisme *Adaptive Thresholding* dinamis pada nilai RE. Dengan menentukan ambang batas RE yang bergerak menyesuaikan statistik lalu lintas jaringan *real-time*, diharapkan sistem dapat menyeimbangkan antara deteksi serangan *zero-day* (*Recall*) dan kenyamanan operasional (FPR rendah).

