

BAB I

Pendahuluan

1.1 Latar Belakang

Ancaman keamanan siber pada era digital saat ini terus berkembang, seiring dengan meningkatnya jumlah perangkat yang terhubung ke internet. Salah satu bagian yang sering terdampak oleh serangan siber ini adalah lapisan jaringan. Kerentanan ini dapat dipicu oleh berbagai faktor, seperti lemahnya kebijakan keamanan, kurangnya perangkat pengaman, serta kurangnya kesadaran akan tanda-tanda serangan. Sebagai upaya menghadapi ancaman tersebut, *Intrusion Detection System* (IDS) hadir sebagai salah satu mekanisme pertahanan yang berfungsi untuk mendeteksi serangan (Ananda Febian dkk., 2024).

Secara umum IDS diklasifikasikan menjadi dua kategori utama berdasarkan cara deteksinya, yaitu berbasis *signature* (tanda tangan) dan berbasis *anomaly* (anomali). Pendekatan berbasis *signature* bekerja dengan cara membandingkan pola pada trafik jaringan dengan pola serangan yang telah diketahui, sedangkan pendekatan berbasis *anomaly* bekerja dengan cara membuat profil trafik normal, lalu mengidentifikasi penyimpangan berdasarkan profil tersebut (Chiba dkk., 2019; Sontan & Samuel, 2024). Pendekatan berbasis *signature* memiliki kelebihan dalam mendeteksi serangan yang polanya sudah diketahui dengan nilai *false positive rate* yang rendah, sedangkan pendekatan berbasis *anomaly* memiliki kelebihan dalam mendeteksi serangan yang polanya tidak diketahui (Ananda Febian dkk., 2024; Ouiazzane dkk., 2022).

Penelitian Xu melakukan perbandingan antara model *supervised* dan *unsupervised* dalam mendeteksi serangan pada skenario pengujian *known attack* dan *unknown attack*. Data yang digunakan pada penelitian tersebut adalah CIC-IDS2017 yang terdiri atas 2.830.743 baris data. Berdasarkan pengujian yang dilakukan, model *supervised* lebih unggul daripada model *unsupervised* pada skenario pengujian *known attack*. Sebaliknya, pada skenario pengujian *unknown attack*, model *supervised* tidak dapat mendeteksi serangan dengan baik. Hal ini ditunjukkan dengan nilai *recall* tertinggi model *supervised* yang hanya sebesar 19,54%. Sebaliknya, model *unsupervised* seperti *Local Outlier Factor* (LOF) dan *One-Class Support Vector Machine* (OCSVM) memiliki performa yang baik pada skenario pengujian *unknown attack*. Nilai *recall* tertinggi model *unsupervised* pada skenario pengujian *unknown attack* mencapai 83,70%. Pencapaian nilai *recall* ini datang dengan bayaran *false positive rate* yang cukup tinggi (Xu & Liu, 2025).

Penelitian Verdejo menggunakan tiga model IDS berbasis *signature* untuk mendeteksi serangan. Penelitian dilakukan terhadap tujuh *dataset* publik dengan tiga model IDS, yaitu Snort, ModSecurity, dan Nemesida. Berdasarkan percobaan yang dilakukan, penggunaan keputusan kooperatif (*union*) berhasil meningkatkan *precision* ataupun *detection rate*. Nilai rata-rata *detection rate* tertinggi yang diperoleh melalui metode *union* mencapai 96,7%. Meskipun *union* meningkatkan *detection rate* terhadap serangan, karakter IDS berbasis *signature* yang sangat bergantung pada *database signature* membuat IDS tersebut tidak efektif untuk mendeteksi serangan yang tidak diketahui (Díaz-Verdejo dkk., 2022; Nawaal dkk., 2023).

Berdasarkan beberapa penelitian terdahulu, IDS dengan pendekatan *signature* mengalami kesulitan dalam mendeteksi serangan yang tidak diketahui. Sementara itu, IDS dengan pendekatan *anomaly* mampu mengidentifikasi serangan yang tidak diketahui, tetapi cenderung menghasilkan tingkat *false positive* yang tinggi (Díaz-Verdejo dkk., 2022; Nawaal dkk., 2023; Xu & Liu, 2025). Untuk mengatasi kekurangan metode *signature* dalam mendeteksi serangan yang tidak diketahui, penelitian ini mengusulkan model hibrida dengan menambahkan komponen deteksi *anomaly* berbasis *machine learning*. Pada model hibrida, komponen *signature* nantinya akan berperan sebagai filter awal untuk mengidentifikasi serangan yang polanya sudah diketahui. Dengan demikian, hanya trafik yang tidak sesuai dengan *signature* yang akan diproses lebih lanjut oleh komponen deteksi *anomaly* berbasis *machine learning*.

Komponen deteksi *anomaly* berbasis *machine learning* pada model hibrida memiliki kekurangan berupa tingginya tingkat *false positive*. Untuk membatasi dampak kekurangan tersebut, penelitian ini mengusulkan penerapan konsep *gating* (kebijakan pada keputusan) pada level *decision* dengan memanfaatkan mekanisme *watchlist* (kumpulan sumber data yang memiliki risiko tinggi). Tanpa *gating* yang tepat, kombinasi antarmodel berisiko meningkatkan *false positive* dan menurunkan kinerja model (Elmubarak dkk., 2019; Ogwara dkk., 2025; Wolsing dkk., 2024). Selain itu, untuk meningkatkan kinerja model hibrida, penelitian ini juga melakukan optimasi terhadap *hyperparameter* dari komponen deteksi *anomaly* berbasis *machine learning*. Berdasarkan percobaan beberapa studi lain, optimasi terhadap *hyperparameter* telah terbukti dapat meningkatkan kinerja beberapa

model *machine learning*. Peningkatan kinerja tersebut berdampak pada beberapa nilai, seperti akurasi, presisi, *recall*, dan *F1-score* (Alinda Rahmi & Defit, 2024; Nurcahyo & Sasongko, 2023).

Pada penelitian ini, evaluasi model akan dilakukan melalui dua skenario pengujian, yaitu *known attack* dan *unknown attack*. Skenario pengujian *known attack* didefinisikan untuk serangan yang memiliki *signature* aktif, sedangkan skenario pengujian *unknown attack* didefinisikan untuk serangan yang tidak memiliki *signature* aktif. Melalui pemisahan ini, *trade-off* kinerja model hibrida pada kondisi ketika *signature* tersedia dan tidak tersedia dapat lebih terukur.

1.2 Rumusan Masalah

Penjelasan pada bagian latar belakang menunjukkan bahwa pendekatan berbasis *signature* sangat baik dalam mendeteksi serangan yang polanya sudah diketahui, namun memiliki keterbatasan dalam mengenali serangan baru. Sementara itu, pendekatan berbasis *anomaly* mampu mengenali serangan baru, tetapi menghasilkan banyak *false positive*. Oleh karena itu, diperlukan suatu model yang dapat memanfaatkan keunggulan sekaligus mengendalikan kelemahan dari kedua pendekatan tersebut. Berdasarkan hal tersebut, rumusan masalah dalam penelitian ini adalah sebagai berikut.

1. Bagaimana efektivitas model hibrida yang menggabungkan pendekatan berbasis *signature* dan *anomaly* dalam mendeteksi serangan?
2. Bagaimana model hibrida membatasi dampak *false positive* yang dihasilkan oleh komponen deteksi anomaly berbasis machine learning pada *output* sistem?

1.3 Tujuan Penelitian

Tujuan dari penelitian ini adalah sebagai berikut.

1. Menganalisis efektivitas model hibrida yang menggabungkan pendekatan berbasis *signature* dan *anomaly* dalam mendeteksi serangan.
2. Mengevaluasi kinerja model hibrida dalam membatasi dampak *false positive* yang dihasilkan oleh komponen deteksi *anomaly* berbasis *machine learning* pada *output* sistem.

1.4 Manfaat Penelitian

Manfaat penelitian ini adalah sebagai berikut.

1. Memberikan kontribusi dalam pengembangan IDS dengan pendekatan hibrida yang mampu bekerja dengan baik dalam mendeteksi serangan yang diketahui maupun yang tidak diketahui.
2. Mengukur kemampuan masing-masing model dalam mendeteksi serangan, yaitu *signature*, *anomaly*, hibrida, dan hibrida dengan *watchlist*.
3. Mengukur *trade-off* kinerja model hibrida pada skenario pengujian *known attack* dan *unknown attack*.

1.5 Batasan Masalah

Batasan-batasan masalah dalam penelitian ini adalah sebagai berikut.

1. Data yang digunakan adalah CIC-IDS2017, yang mencakup beragam jenis serangan jaringan.
2. Simulasi tidak dilakukan secara *real-time*, melainkan dilakukan secara *offline*.

3. Performa model dievaluasi melalui akurasi, presisi, *recall*, *false positive rate* (FPR), dan *F1-score*.
4. Pengujian dilakukan di lingkungan simulasi berbasis perangkat keras umum (CPU dan GPU standar) dan lingkungan virtual.
5. Deteksi serangan difokuskan pada pola serangan berbasis jaringan dan tidak mencakup analisis serangan berbasis file, seperti virus atau ransomware.
6. Simulasi *unknown attack* dilakukan dengan cara mematikan *signature* terkait pada IDS berbasis *signature* dan mengisolasi data terkait dari data yang digunakan untuk melatih model.
7. Penelitian menggunakan Suricata sebagai IDS berbasis *signature* dan *Isolation Forest* sebagai metode deteksi *anomaly* berbasis *machine learning*.