

1. Pendahuluan

1.1 Latar Belakang

Meningkatnya integrasi perangkat IoT meningkatkan kenyamanan dan menghadirkan tantangan keamanan siber, terutama yang berkaitan dengan serangan botnet. Serangan ini mengeksploitasi perangkat IoT yang disusupi untuk aktivitas jahat seperti Serangan *Denial of Service* (DDoS) dan pencurian data. Botnet terkenal, seperti MIRAI dan BASHLITE, telah menyebabkan gangguan yang signifikan. *Machine Learning* (ML) telah meningkatkan deteksi botnet IoT dengan mengidentifikasi pola dalam kumpulan data besar yang sering tidak terdeteksi oleh manusia [1] [2].

Explainable Artificial Intelligence (XAI) telah memfasilitasi deteksi serangan siber yang lebih efektif di IoT. XAI memberikan wawasan yang komprehensif tentang alasan di balik keputusan yang dibuat oleh model *Artificial Intelligence* (AI), meningkatkan keandalan dan transparansi kerangka kerja deteksi [3]. Algoritma *ML Ensemble Trees* (ET) mengintegrasikan beberapa pohon keputusan untuk menghasilkan model yang lebih akurat dan andal.

XAI adalah teknik yang memungkinkan penjelasan sistem deteksi intrusi berbasis ML. Hal ini memungkinkan model untuk memberikan alasan mendasar di balik prediksinya, dengan membuka "*Black Box*" model AI dan menawarkan pemahaman yang jelas tentang cara kerja model [4]. Dua pendekatan utama untuk XAI adalah transparansi desain dan penjelasan *post-hoc*. Transparansi desain menjelaskan bagaimana model dirancang. Ini melibatkan analisis kontribusi setiap pohon keputusan dalam ansambel terhadap keputusan akhir, struktur keseluruhan model, dan pohon keputusan yang membentuk model. XAI menawarkan visualisasi fitur yang memiliki pengaruh signifikan, seperti *Shapley Additive Explains* (SHAP) dan *Local Interpretable Model-agnostic Explains* (LIME). Visualisasi ini membantu mengidentifikasi faktor-faktor utama yang paling memengaruhi hasil prediksi model.

Penelitian ini mengevaluasi kedua metode tersebut secara komprehensif dalam konteks deteksi botnet IoT. Penelitian ini dilakukan dengan menggunakan beberapa model machine learning, antara lain *Gradient Boosting with Early Stopping*, *Catboost* dan *Histogram Gradient Boosting*. Berbeda dengan penelitian sebelumnya yang sering menggunakan berbagai metrik kuantitatif, penelitian ini berfokus pada penerapan metode explainability untuk memahami keputusan model. Hal ini bertujuan untuk mengisi kesenjangan dalam penelitian tentang efektivitas metode *explainability* lokal dalam model machine learning untuk aplikasi keamanan IoT. Dengan menggunakan SHAP dan LIME sebagai metode utama, penelitian ini memberikan wawasan awal mengenai peran kemampuan menjelaskan dalam meningkatkan kepercayaan pada model, tanpa bergantung pada metrik kuantitatif.

Struktur penelitian ini adalah sebagai berikut. Bagian II menyajikan tinjauan pustaka tentang topik yang dibahas. Bagian III menjelaskan metode yang diusulkan dalam penelitian. Bagian IV menyajikan hasil penelitian dan pembahasan yang dihasilkan, dan Bagian V menyajikan simpulan penelitian ini.

1.2 Rumusan Masalah

Berdasarkan latar belakang yang telah dijabarkan sebelumnya, maka rumusan masalah yang akan dibahas adalah sebagai berikut:

1. Bagaimana performansi *Ensemble Tree-Based* (*CatBoost*, *HGBD*, *Gradient Boosting with Early Stopping*) dalam deteksi serangan botnet pada perangkat IoT yang heterogen?
2. Bagaimana XAI dapat mentransparansikan kinerja model *Ensemble Tree-Based* (*CatBoost*, *HGBD*, *Gradient Boosting with Early Stopping*) dalam deteksi serangan botnet pada perangkat IoT yang heterogen?

1.3 Tujuan

Adapun tujuan dilakukannya penelitian ini, antara lain:

1. Menerapkan metode *Ensemble Tree-Based* (*CatBoost*, *HGBD*, *Gradient Boosting with Early Stopping*) dalam deteksi serangan botnet.
2. Menganalisis hasil interpretasi yang diberikan oleh LIME dan SHAP untuk memahami faktor-faktor utama yang mempengaruhi serangan deteksi botnet.