

Abstrak

Deteksi anomali jaringan sangat penting dalam bidang keamanan siber. Fenomena ini dapat dikaitkan dengan meningkatnya kompleksitas lalu lintas jaringan dan kecanggihan ancaman siber kontemporer. Namun, opasitas inheren dalam model pembelajaran mendalam, seperti jaringan saraf tiruan (JST), menimbulkan tantangan signifikan dalam memahami dan mempercayakan keputusan sistem. Studi ini mengatasi masalah ini dengan mengevaluasi dan membandingkan interpretabilitas visualisasi SHAP (SHapley Additive exPlanations) global dan lokal yang diterapkan pada sistem deteksi anomali jaringan berbasis JST (jaringan saraf tiruan). Dalam studi ini, kami menggunakan dataset NF-UNSW-NB15 v2 untuk melakukan serangkaian eksperimen yang memanfaatkan lima visualisasi SHAP: bar, beeswarm, waterfall, cohort attack, dan decision plot. Evaluasi dilakukan dengan menggunakan berbagai metrik kuantitatif dan kualitatif, yang mencakup waktu untuk mendapatkan wawasan, kompleksitas rendering, entropi, kepadatan tepi, dan ukuran berkas. Temuan menunjukkan bahwa plot batang mencapai waktu pemahaman tercepat (0,0999 detik lokal, 0,2012 detik global), rendering paling sederhana (17 detik lokal, 31 detik global), dan entropi paling substansial (0,7646 detik lokal, 0,8353 detik global). Akibatnya, plot batang paling efektif dalam interpretabilitas global dan lokal. Sebaliknya, plot kohort mencatat ukuran berkas dan kepadatan tepi tertinggi, sehingga menyoroti kompleksitasnya. Temuan ini membangun kerangka kerja untuk pemilihan metode visualisasi dalam aplikasi keamanan siber waktu nyata (real-time) dan berdimensi tinggi.