

## BAB 5

### KESIMPULAN DAN SARAN

#### 5.1 Kesimpulan

Penelitian ini berhasil mengembangkan sistem deteksi ancaman internal berbasis GNN dengan mengimplementasikan GraphSAGE untuk klasifikasi perilaku pengguna, serta GraphSVX sebagai metode interpretasi model. Pendekatan ini dirancang untuk mengidentifikasi serta menjelaskan pola perilaku mencurigakan pada pengguna dalam dataset CERT Insider Threat versi r6.2, dengan representasi graf heterogen yang menggabungkan entitas user, pc, dan url.

Model yang dibangun menggunakan metode GraphSAGE mampu mengklasifikasikan pengguna sebagai normal atau insider dengan akurasi rata-rata sebesar 99,2% dan F1-score sebesar 96%, bahkan mencapai skor sempurna dalam beberapa skenario pengujian. Hasil ini menunjukkan bahwa pendekatan berbasis graf memiliki potensi besar dalam memetakan hubungan kompleks antar entitas perilaku pengguna.

Kemudian, metode interpretasi GraphSVX berhasil digunakan untuk menjelaskan hasil klasifikasi model dengan mengidentifikasi fitur yang paling berkontribusi terhadap prediksi. Fitur-fitur dominan yang ditemukan mencakup login di luar jam kerja, akses file dalam jumlah besar, dan kunjungan ke situs cloud dan pencarian kerja, yang merupakan indikator kuat terhadap aktivitas berisiko tinggi.

Dari hasil analisis risiko, metode GraphSVX juga mampu mengidentifikasi pengguna-pengguna dengan skor probabilitas ancaman tertinggi, sehingga dapat dimanfaatkan untuk peringatan dini (*early warning*) dan audit internal secara selektif. Dengan demikian, metode yang dikembangkan memiliki potensi tinggi untuk diadopsi sebagai bagian dari mekanisme pengambilan keputusan dalam pengelolaan keamanan siber berbasis profil perilaku pengguna.

Namun, perlu dicatat bahwa dataset yang digunakan bersifat simulatif. Oleh karena itu, performa dan efektivitas model dalam lingkungan dunia nyata dapat

dipengaruhi oleh faktor-faktor lain seperti dinamika perilaku pengguna, kualitas data log, serta adanya atribut-atribut non-teknis yang tidak tersedia dalam dataset ini.

## 5.2 Saran

Berdasarkan proses dan hasil yang telah dicapai, berikut adalah saran untuk pengembangan penelitian lebih lanjut:

1. Perluasan cakupan data: Penelitian ini hanya menggunakan dua skenario (2 dan 5) dari dataset CERT Insider Threat r6.2. Untuk memperoleh generalisasi yang lebih baik, disarankan untuk mengikutsertakan seluruh skenario atau menggunakan dataset dari lingkungan organisasi nyata apabila tersedia.
2. Perbandingan dengan model GNN lainnya: GraphSAGE merupakan salah satu pendekatan GNN yang populer. Penelitian selanjutnya dapat membandingkan performanya dengan model lain seperti GAT, HAN, atau HeteroGraph Transformer untuk menentukan pendekatan paling tepat dalam konteks deteksi ancaman internal.
3. Integrasi validasi silang (*cross-validation*): Karena pelatihan model saat ini hanya dilakukan satu kali per konfigurasi, validasi silang sangat disarankan untuk meningkatkan keandalan hasil dan mengurangi risiko overfitting.
4. Uji coba pada sistem nyata: Mengingat penelitian ini bersifat eksperimental, langkah selanjutnya adalah mengimplementasikan model pada sistem monitoring organisasi nyata dengan mempertimbangkan aspek privasi, latensi sistem, dan integrasi secara *real-time*.