

BAB V PENUTUP

5.1 Kesimpulan

Penelitian ini telah berhasil menganalisis dan mengevaluasi tiga metode XAI, yaitu SHAP, LIME, dan PDP, dalam menjelaskan prediksi model deteksi *phishing* berbasis *machine learning*. Berdasarkan hasil evaluasi, diketahui bahwa metode LIME memiliki nilai *fidelity* tertinggi, yaitu sebesar 1.0, diikuti oleh SHAP dengan nilai 0.989, dan PDP sebesar 0.501. Hal ini menunjukkan bahwa LIME dan SHAP memberikan penjelasan yang lebih sesuai dengan *output* model dibandingkan PDP. Dari sisi efisiensi waktu, metode SHAP menunjukkan waktu komputasi paling tinggi, yaitu sekitar 150 detik, sedangkan LIME dan PDP jauh lebih efisien dengan waktu komputasi di bawah 1 detik. Secara khusus, LIME dalam skripsi ini mencatat waktu komputasi sebesar 0,608 detik, yang lebih cepat dibandingkan dengan nilai waktu komputasi pada penelitian yang dilakukan oleh Shafin untuk model KNN, yaitu sebesar 0,622 detik. [14]

Berdasarkan segi stabilitas, PDP menunjukkan performa paling tinggi, namun konsistensinya dalam menjelaskan hasil prediksi relatif rendah dibandingkan SHAP dan LIME. SHAP terbukti memiliki nilai *fragility* terendah, menunjukkan ketahanannya terhadap gangguan kecil pada input data. Selain itu, visualisasi yang ditawarkan oleh SHAP dan LIME tergolong sangat baik untuk memahami pengaruh fitur terhadap keputusan model, sedangkan PDP unggul dalam menyampaikan hubungan global antar fitur. Dengan demikian, penelitian ini menyimpulkan bahwa tidak ada satu metode XAI yang unggul dalam semua aspek. SHAP sangat direkomendasikan untuk aplikasi yang menuntut akurasi dan ketelitian dalam interpretasi, meskipun membutuhkan sumber daya komputasi yang lebih tinggi. LIME lebih cocok digunakan dalam sistem interaktif yang membutuhkan interpretasi cepat, sedangkan PDP lebih relevan untuk analisis global terhadap perilaku model.

5.2 Saran

Untuk penelitian selanjutnya disarankan untuk memperluas cakupan dengan menyertakan metode XAI lainnya seperti *Anchorx*, *Integrated Gradients*, dan *DeepLIFT* guna memperoleh perspektif interpretabilitas yang lebih luas. Penelitian juga dapat diperkaya dengan pendekatan evaluasi berbasis pengguna (*user-centered evaluation*), sehingga kualitas penjelasan dapat dinilai tidak hanya secara kuantitatif, tetapi juga dari aspek pemahaman dan kepercayaan pengguna terhadap model.

