

Studi Komparatif Teknik Explainable Artificial Intelligence dalam Outlier Detection = Comparative Study of Explainable Artificial Intelligence Techniques in Outlier Detection

Ahmad Haulian Yoga Pratama, author

Deskripsi Lengkap: <https://lib.ui.ac.id/detail?id=9999920567052&lokasi=lokal>

Abstrak

Penerapan teknik Explainable AI (XAI) telah menjadi fokus utama penelitian dalam upaya untuk meningkatkan interpretabilitas dan kepercayaan dalam model AI, khususnya pada bidang outlier detection. Penelitian ini bertujuan untuk mengungkapkan proses pengambilan keputusan yang kompleks di balik proses outlier detection, serta untuk memberikan pemahaman yang lebih dalam tentang faktor-faktor yang mempengaruhi keputusan tersebut. Dalam penelitian ini, diselidiki berbagai teknik XAI yang dapat digunakan dalam konteks outlier detection. Penelitian ini memberikan evaluasi komprehensif tentang aplikasi XAI dalam outlier detection, dengan mengevaluasi kelebihan dan kelemahan dari setiap teknik yang digunakan. Hasil eksperimen menunjukkan bahwa penerapan XAI dalam outlier detection dapat memberikan wawasan yang berharga tentang faktor-faktor yang mempengaruhi keputusan model, dan dapat meningkatkan interpretabilitas dan kepercayaan dalam model outlier detection.

.....The application of Explainable AI (XAI) techniques has been the main focus of research to improve interpretability and trust in AI models, particularly in the field of outlier detection. This study aims to uncover the complex decision-making process behind outlier detection and provide a deeper understanding of the factors influencing these decisions. Various XAI techniques that can be used in outlier detection are investigated in this research. This study provides a comprehensive evaluation of XAI applications in outlier detection by assessing the strengths and weaknesses of each technique used. The experimental results indicate that the implementation of XAI in outlier detection can provide valuable insights into the factors influencing model decisions and can enhance the interpretability and trustworthiness of outlier detection models.