



ID de Contribution: 13

Type: **Contribution orale**

Apport des motifs gradual dans pour l'IA explicable

jeudi 26 février 2026 14:00 (20 minutes)

Les modèles ensemblistes et les réseaux de neurones profonds démontrent de très bons résultats dans les tâches de classification. Cependant, leur nature « boîte noire » empêche leur déploiement généralisé dans des domaines critiques comme la santé. L'IA explicable vise à rendre ces modèles plus compréhensibles. Dans la littérature, les résultats de classification sont expliqués principalement par attribution de l'importance des caractéristiques ou par génération de contrefactuelles basées sur des voisinages générés aléatoirement, par algorithmes génétiques ou à partir des connaissances d'experts. Dans cet article, nous montrons comment les motifs graduels permettent de générer des voisinages plausibles sans expertise, produisant des explications mieux adaptées aux instances individuelles. Nous étendons notre méthode avec une analyse complète des résultats expérimentaux comparant notre approche à LIME, LORE et SHAP.

Auteur: KAMGA NGUIFO, Durande Berluskoni (Université clermont Auvregne)

Classification de Session: Foundations and Sustainability in Machine Learning