

UNIVERZITET U SARAJEVU -
ELEKTROTEHNIČKI FAKULTET
S a r a j e v o
Zmaja od Bosne bb

Na osnovu čl. 5. i 6. Odluke Vijeća Univerziteta u Sarajevu - Elektrotehničkog fakulteta o definiranju procedure realizacije naučnoistraživačkih seminara na trećem ciklusu studija – doktorskom studiju (broj: 01-503/21 od 01.02.2021. godine) i Odluke Vijeća Univerziteta u Sarajevu - Elektrotehničkog fakulteta (broj: 01-1545/25 od 12.05.2025. godine), Univerzitet u Sarajevu - Elektrotehnički fakultet, daje

O B A V I J E S T
o odbrani seminara

Student trećeg ciklusa studija - dokorskog studija, Vahidin Hasić, magistar elektrotehnike - diplomirani inženjer elektrotehnike, branit će Naučnoistraživački seminar 1.1. pod naslovom "Towards Explainable Image Classification".

Seminar je izrađen u saradnji sa akademskom savjetnicom, dr. Senkom Krivić, docenticom Univerziteta u Sarajevu - Elektrotehničkog fakulteta.

Odbrana seminara održat će se 20. maja 2025. godine (utorak), s početkom u 13:00 sati, u prostorijama Univerziteta u Sarajevu - Elektrotehničkog fakulteta (Amfiteatar A3).

Odbrana seminara je javna.

Obavijest o odbrani i sažetak seminara, oglašavaju se na oglasnim pločama i internet stranici Univerziteta u Sarajevu - Elektrotehničkog fakulteta.

Oglašeno:
Sarajevo, 14.05.2025. godine



Akadska savjetnica: Doc. dr. Senka Krivić

Student: Vahidin Hasić, magistar elektrotehnike – diplomirani inženjer elektrotehnike

Naziv naučnoistraživačkog seminara:

Towards Explainable Image Classification

Sažetak:

Deep Neural Networks (DNNs) achieve remarkable performance across diverse applications, but their inherent opacity hinders trustworthiness. This lack of interpretability poses a significant challenge, particularly in critical domains like image classification. This paper surveys seminal and state-of-the-art XAI (Explainable AI) methods for DNNs, drawing from leading AI conferences and journals. Our analysis reveals that concept-based counterfactual and contrastive explanations have the most potential to achieve model explainability. Furthermore, we argue that incorporating sample based explanations is crucial for a holistic understanding of model behavior. This comprehensive approach to explainability can significantly improve the trustworthiness and reliability of DNN deployments.

