

Zgłoszenie tematu pracy dyplomowej :: STUDIA II STOPNIA ::

na rok akademicki 2025/26

Promotor:	Dr hab. inż. Mateusz Muchacki, prof UKEN
Temat pracy inżynierskiej (j. polski oraz j. angielski):	Postrzeganie błędów i halucynacji generatywnej AI w opinii użytkowników <i>Users' Perception of Errors and Hallucinations in Generative AI</i>
Zakres pracy i oczekiwane rezultaty praktyczne:	Głównym założeniem pracy powinno być omówienie problematyki błędów i halucynacji w generatywnych modelach językowych zrealizowane na podstawie analizy aktualnej literatury z obszaru sztucznej inteligencji oraz interakcji człowiek-komputer. Część teoretyczna powinna zawierać analizę pojęcia zaufania technologicznego oraz innych czynników determinujących ocenę wiarygodności systemów AI przez użytkowników. Część badawcza będzie opierać się na opracowaniu i przeprowadzeniu badania ankietowego mającego na celu poznanie opinii użytkowników, które dotyczyć będzie występowania błędów i halucynacji w systemach generatywnej AI oraz ich wpływu na poziom zaufania i sposób korzystania z tych narzędzi. Analiza uwzględni różnice w percepcji w zależności od poziomu zaawansowania technicznego respondentów. Wyniki badania zostaną poddane analizie i interpretacji, a następnie zestawione z wnioskami z literatury, co pozwoli na sformułowanie praktycznych wniosków dotyczących świadomego korzystania z systemów AI.
Aspekt naukowy, problemowy, innowacyjny pracy:	Aspekt naukowy: analiza zjawiska błędów i halucynacji generatywnych modeli językowych oraz ich wpływu na zaufanie użytkowników do systemów AI. Zbadana zostanie percepcja użytkowników w kontekście wiarygodności i użyteczności generatywnej AI.
*Oprogramowanie, język programowania, środowisko systemowe:	Brak.
*Środowisko uruchomieniowe	Brak.
Dodatkowe wymagania i uwagi:	Brak uwag.
*Literatura:	A Survey on Hallucination in Large Language Models: Principles, Taxonomy, Challenges, and Open Questions, 2023 - Lei Huang, Weijiang Yu, Weitao Ma, Weihong Zhong. Siren's Song in the AI Ocean: A Survey on Hallucination in Large Language Models, 2023 - Yue Zhang, Yafu Li, Leyang Cui, Deng Cai. Formalizing Trust in Artificial Intelligence: Prerequisites, Causes and Goals of Human Trust in AI, 2021 - Alon Jacovi, Ana Marasović.

*pola opcjonalne