

Zgłoszenie tematu pracy magisterskiej :: **STUDIA II STOPNIA** ::

na rok akademicki 2025/26

Promotor:	dr hab. inż. Mateusz Muchacki, prof. UKEN
Temat pracy magisterskiej (j. polski oraz j. angielski):	Analiza kryteriów oceny autentyczności treści w Internecie: badanie zdolności użytkowników do rozróżniania opinii produktowych tworzonych przez ludzi i modele LLM. <i>Analysis of authenticity assessment criteria for online content: A study on users' ability to distinguish between human-written and LLM-generated product reviews.</i>
Zakres pracy i oczekiwane rezultaty praktyczne:	Celem pracy jest zidentyfikowanie cech (heurystyk), którymi kierują się użytkownicy przy ocenie wiarygodności tekstów w e-commerce. Część praktyczna obejmuje przygotowanie zbioru danych (datasetu) składającego się przykładowo z trzech grup opinii dla wybranych produktów: 1. Opinie autentyczne (Human): Pobranie rzeczywistych opinii z platform e-commerce. 2. Opinie AI Basic: Wygenerowane przez model LLM w trybie domyślnym (Zero-shot). 3. Opinie AI Advanced: Wygenerowane z użyciem inżynierii promptów (nadanie osoby, celowe błędy, stylizacja). Głównym aspektem badawczym jest przeprowadzenie eksperymentu z udziałem użytkowników, w którym badani będą nie tylko klasyfikować opinie, ale również wskazywać przyczyny swojej decyzji
Aspekt naukowy, problemowy, innowacyjny pracy:	Problem badawczy dotyczy percepcji "człowieczeństwa" w tekście cyfrowym. Wyniki badań powinny wskazywać na te atrybuty tekstu (stylistyczne, gramatyczne, semantyczne), które są dla użytkowników kluczowymi wyznacznikami autentyczności. Badanie pomoże sprawdzić, czy nowoczesne modele LLM potrafią skutecznie symulować te atrybuty (mimikrę) oraz na ile są w stanie zakłócić ocenę intuicyjną człowieka.
*Oprogramowanie, język programowania, środowisko systemowe:	Generowanie i przetwarzanie danych: Python i LLM APIs. Formularz ankietowy (np. Google Forms). Analiza statystyczna wyników: Python (Matplotlib, Seaborn, Scikit-learn)/Excel.
*Środowisko uruchomieniowe	-
Dodatkowe wymagania i uwagi:	Brak uwag.
*Literatura:	1. Jakesch, M., Hancock, J. T., & Naaman, M. (2023). Human Heuristics for AI-Generated Language Are Flawed. Proceedings of the National Academy of Sciences (PNAS) 2. Clark, E., et al. (2021). All That's 'Human' Is Not Gold: Evaluating

Zgłoszenie tematu pracy magisterskiej :: **STUDIA II STOPNIA** ::

na rok akademicki 2025/26

	Human Evaluation of Generated Text. Proceedings of the 59th Annual Meeting of the Association for Computational Linguistics. 3. Adelani, D. I., et al. (2020). Generating Sentiment-Preserving Fake Online Reviews Using Neural Language Models. In Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing (EMNLP). 4. Automated Journalism: The Effects of AI Authorship and Evaluative Information on the Perception of a Science Journalism Article Angelica Lermann Henestrosa, Hannah Greving, Joachim Kimmerle 2022
--	--

*pola opcjonalne