

ChatGPT 4 and personality prediction

Zahra Eslami, Marcus Cheetham, Seyed Abolfazl Valizadeh*

First author: Eslami - Department of Biomedical engineering, Shahid Beheshti University, Tehran, I.R. Iran

Second author: Marcus Cheetham – Department of Internal Medicine, University Hospital Zurich, Zurich, Switzerland

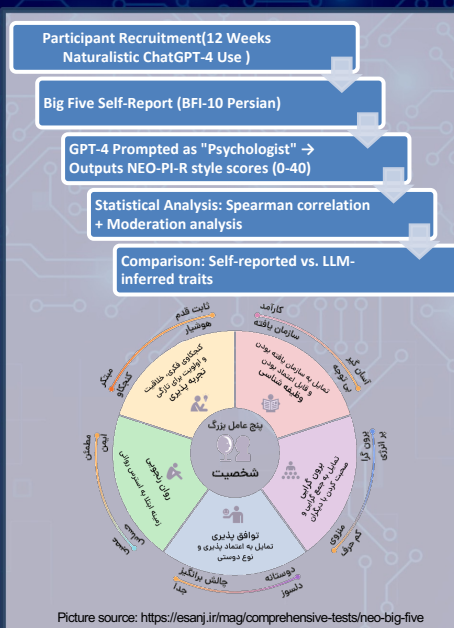
*Corresponding author: Seyed Abolfazl Valizadeh – Department of Biomedical engineering, Shahid Beheshti University, Tehran, I.R. Iran

Emails zahraeslami2@gmail.com, Marcus.Cheetham@usz.ch, s_valizadeh@sbu.ac.ir

Abstract

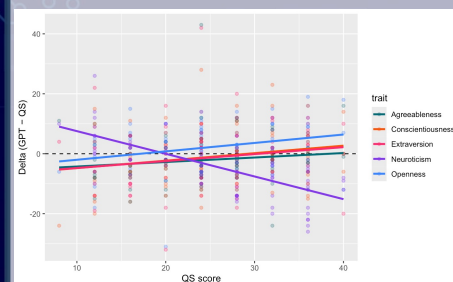
Large Language Models (LLMs) such as ChatGPT-4 have been successfully applied to infer the Big Five personality traits (openness, conscientiousness, extraversion, agreeableness, and neuroticism) from general text. However, their use for personality inference from LLM user interaction text (e.g., queries and retrieved information) remains underexplored. We hypothesize that conscientiousness acts as a moderator, increasing the predictive accuracy of LLM-inferred traits due to the linguistic alignment between characteristic conscientious text patterns and LLM training data. Participants (N=36) completed a validated self-report of the Big Five and used ChatGPT-4 for over 12 weeks. ChatGPT-4 was used to generate personality inferences from their text inputs. Correlational and moderation analyses revealed that the correlation between self-reported and LLM-inferred traits was significantly stronger for individuals with higher self-reported conscientiousness. This effect was most notable for the trait of conscientiousness itself, with modest, consistent moderation effects observed across the other Big Five traits. Predictive accuracy remained moderate overall, aligning with existing computational personality research. These findings suggest that the linguistic consistency of conscientious individuals facilitates more accurate LLM-based personality inference, emphasizing the need for improved model calibration and input assessment.

Research method



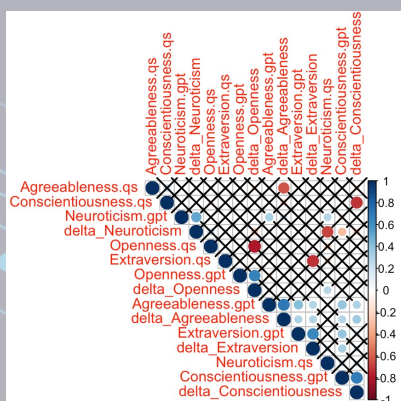
conclusion

ChatGPT-4 infers personality from naturalistic text but shows strong central tendency bias, overestimating low traits and underestimating high ones. Inferred traits lack discriminant validity, with excessive cross-dimension correlations. Despite moderate neuroticism agreement, LLM-based personality estimates are not psychometrically equivalent to validated questionnaires and require rigorous calibration before applied use.



The purpose of the research

This study examined whether ChatGPT-4 accurately infers Big Five personality traits from naturalistic, longitudinal user text inputs and tested the hypothesis that conscientiousness moderates predictive accuracy, with higher conscientiousness enhancing alignment between self-reported and LLM-inferred traits.



Practical or simulation results

Spearman correlations between self-reported and GPT-inferred traits were significant only for Neuroticism ($p=0.401$, $p=0.015$). Regression revealed strong central tendency bias across all traits (slopes: -0.664 to -0.967 ; $R^2=0.673$ for Openness, $p=8.96 \times 10^{-10}$). GPT-inferred Extraversion, Agreeableness, and Conscientiousness showed excessive inter-correlations (p up to 0.654).

Trait	Spearman p (QS vs. GPT)	P-Value
Neuroticism	0.401	0.015
Openness	0.200	0.241
Extraversion	0.033	0.849
Agreeableness	0.168	0.326
Conscientiousness	0.137	0.425

References

Foundational evidence that language use reflects personality traits

Pennebaker, J. W., & King, L. A. (1999). Linguistic styles: language use as an individual difference. *Journal of Personality and Social Psychology*, 77(6), 1296–1312.

ChatGPT-4's emergent capability to estimate Big Five traits from text

Piastra, M., & Catellani, P. (2025). On the emergent capabilities of ChatGPT 4 to estimate personality traits. *Frontiers in Artificial Intelligence*, 8, 1484260.

GPT-4 infers personality from free-form user interactions with accuracy varying by conversational context

Peters, H., Cerf, M., & Matz, S. C. (2024). Large language models can infer personality from free-form user interactions. *arXiv preprint arXiv:2405.13052*.

Limitation: Small sample size, cultural specificity, and privacy constraints limit generalizability and preclude fine-grained linguistic analysis.