

Ideological patterns in Large Language Model Behaviours

Tijl De Bie¹

University of Ghent, Belgium

`Tijl.DeBie@UGent.be`

Abstract. On 7th September 2026, Tijl De Bie delivered a keynote talk at the 14th Workshop on New Frontiers in Mining Complex Patterns in conjunction with the European Conference on Machine Learning and Principles and Practice of Knowledge Discovery in Databases (ECML-PKDD) in Naples, Italy. The following is the abstract of his talk and a short biography.

Keywords: Large Language Models · Ideological Bias

1 Abstract

As generative AI systems are increasingly used on social and traditional media, in political discourse, and in society at large, it becomes increasingly important to understand which ideological preferences these systems reflect. In this talk I will discuss some methodologies for achieving this understanding, highlight specific findings, and suggest ways to effectively mitigate the associated risks without causing collateral harm.

2 Short Biography

Tijl De Bie is a Professor of Artificial Intelligence at Ghent University. He previously studied and worked at KU Leuven, UC Berkeley, UC Davis, Southampton University and the University of Bristol. His past and current research has covered the foundations of machine learning, data science and pattern mining, as well as applications in biology, (social) media analysis, employment, and more. Most recently he has concentrated on AI safety and information integrity, funded by ERC Advanced Grant VIGILIA.