

From Opportunity to Compliance: the AEQUITAS framework for Fair AI

Roberta Calegari¹

¹University of Bologna

Abstract

AI offers unprecedented opportunities—but also carries the risk of discrimination and unfair treatment embedded in data, design, or decision-making processes. The AEQUITAS framework addresses this challenge through a holistic approach that embeds a controlled experimentation environment to identify, measure, and mitigate bias before deployment. By integrating legal, ethical, and technical perspectives and translating them into concrete, actionable methods, AEQUITAS enables organisations to detect discriminatory patterns, apply fairness mitigation strategies, and validate compliance with the EU AI Act, ensuring that AI systems are developed and deployed in a fair, transparent, and accountable manner.

Biography

Roberta Calegari is a researcher and professor at the Department of Computer Science and at the Alma Mater Research Institute for Human-Centered Artificial Intelligence at the University of Bologna. Her research field is related to trustworthy and explainable systems, distributed intelligent systems, software engineering, multi-paradigm languages and AI law. She is the coordinator of the project Horizon Europe 2020 (G.A. 101070363) about Assessment and engineering of equitable, unbiased, impartial and trustworthy AI systems. The project aims to provide an experimentation playground to assess and repair bias in AI. She has been part of the EU Horizon 2020 Project “PrePAI” (G.A. 101083674) working on the definition of requirements and mechanisms that ensure all resources published on the future AIonDemand platform can be labelled as trustworthy and in compliance with the future AI regulatory framework. Her research interests lie within the broad area of knowledge representation and reasoning in AI for trustworthy and explainable AI and in particular focus on symbolic AI including computational logic, logic programming, argumentation, logic-based multi-agent systems, non-monotonic/defeasible reasoning. She is Member of the Editorial Board of ACM Computing Surveys for the area of Artificial Intelligence. She is the author of more than 90 papers in peer-reviewed international conferences and journals. She is leading many European, Italian and regional projects and she is responsible for collaborations with industries.

Declaration on Generative AI

The author has not employed any Generative AI tools.

Identity-Aware AI workshop at 28th European Conference on Artificial Intelligence, October 25, 2025, Bologna, Italy

 <https://apice.unibo.it/xwiki/bin/view/RobertaCalegari/> (R. Calegari)

 0000-0003-3794-2942 (R. Calegari)



© 2025 Copyright for this paper by its authors. Use permitted under Creative Commons License Attribution 4.0 International (CC BY 4.0).