

Accountability Mechanisms in Automated Decision Systems

Noah Popescu¹, Clara Costa², Clara Garcia³

¹ Professor, Department of Machine Learning, Swiss Institute of Machine Intelligence, Zurich, Switzerland. Email: noah.popescu12@yahoo.com | ORCID: 9040-9078-2832-0196

² Assistant Professor, Institute of Intelligent Systems, Mediterranean Institute of Technology, Rome, Italy. Email: clara.costa206@yahoo.com | ORCID: 0685-3561-9417-2894

³ Assistant Professor, Department of Artificial Intelligence, Nordic Technical University, Stockholm, Sweden. Email: clara.garcia312@outlook.com | ORCID: 8165-9010-4893-8240

ABSTRACT

Accountability in automated decision systems (ADS) -- the capacity to identify responsible parties, trace decision causation, and enable meaningful redress for individuals adversely affected by algorithmic decisions -- is a foundational requirement of responsible AI governance yet remains incompletely operationalised in both technical architectures and regulatory instruments. This paper proposes a comprehensive Accountability Mechanism Framework (AMF) for ADS, comprising seven accountability mechanism classes -- traceability, auditability, explainability, contestability, human oversight, liability assignment, and redress enablement -- each specified with technical implementation requirements, organisational governance requirements, and regulatory compliance criteria. The AMF is evaluated through a mixed-methods study combining a structured assessment of 30 operational ADS across healthcare, financial services, criminal justice, public administration, and employment domains, and a survey of 94 stakeholders (affected individuals, domain experts, legal practitioners, and AI developers) on accountability mechanism adequacy. AMF assessment scores reveal that traceability (mean AMF-T = 6.8/10) and auditability (mean AMF-A = 6.4/10) are the strongest-implemented mechanisms across domains, while contestability (mean = 3.9/10) and redress enablement (mean = 3.4/10) are the weakest. Stakeholder survey results confirm a significant accountability perception gap: developers rate mechanism adequacy 2.1 points higher on average than affected individuals across all seven mechanisms ($p < 0.001$). The study contributes the AMF specification, a validated Accountability Adequacy Index (AAI), and a stakeholder-inclusive evaluation methodology for responsible ADS governance.

Keywords: accountability; automated decision systems; algorithmic accountability; contestability; redress; human oversight; EU AI Act; responsible AI

Citation: Popescu et al. [2025]. Accountability Mechanisms in Automated Decision Systems. DOI: <https://doi.org/10.5281/zenodo.19185001>

Copyright: © 2025 by the authors. Open access under CC BY 4.0 license.

Article Information: Received: 14 February 2025 Accepted: 16 April 2025 Published: 25 June 2025

Research Article: Research Article

1. Introduction

1.1 The Accountability Deficit in Automated Decision Systems

Automated decision systems increasingly perform or inform decisions with profound consequences for individuals: credit approvals, social benefit eligibility, parole and bail recommendations, medical triage prioritisation, university admissions, and hiring screening. The substitution of human judgement by algorithmic processing in these contexts raises a fundamental accountability question: when an ADS produces an erroneous, unfair, or harmful decision, who is responsible, how can the decision be understood and challenged, and what recourse does the affected individual have? The academic and policy literature has long recognised that algorithmic decision-making creates accountability gaps -- situations in which existing accountability frameworks fail to assign clear responsibility for algorithmic harms (Mittelstadt et al., 2016; Diakopoulos, 2016; Kearns and Roth, 2019). These gaps arise from several structural features of ADS: the opacity of complex models makes causal attribution of decisions to specific factors difficult; the distributed development and deployment of AI systems across multiple organisational actors obscures liability assignment; and the absence of standardised redress mechanisms leaves affected individuals without practical means to challenge adverse decisions (Pasquale, 2015; Zarsky, 2016). Regulatory responses have intensified. The EU AI Act (2024) establishes accountability obligations across Articles 9 (risk management), 12 (logging), 13 (transparency), 14 (human oversight), and 26 (obligations of deployers). The GDPR Articles 22 and 13-15 establish rights to explanation and access. The EU product liability directive (2024) extends civil liability to AI systems. Yet the translation of these regulatory obligations into concrete technical and organisational mechanisms within ADS -- and the assessment of their adequacy from the perspective of affected individuals -- remains largely unaddressed in the literature.

1.2 Research Objectives and Structure

This paper proposes the Accountability Mechanism Framework (AMF), a structured specification of seven accountability mechanism classes for ADS, and evaluates their implementation adequacy through a mixed-methods study combining structured system assessment and stakeholder survey. Research objectives are: (1) to specify the AMF with seven mechanism classes, each with

technical, organisational, and regulatory compliance requirements; (2) to assess AMF mechanism implementation in 30 operational ADS; and (3) to evaluate stakeholder perceptions of accountability mechanism adequacy across affected individuals, domain experts, legal practitioners, and AI developers. Section 2 reviews prior work on algorithmic accountability, contestability, and regulatory requirements. Section 3 presents the AMF and Accountability Adequacy Index. Section 4 describes the methodology. Section 5 reports results. Section 6 discusses implications. Section 6 concludes with recommendations.

2. Literature Review

2.1 Algorithmic Accountability -- Conceptual Foundations

Bovens (2007) defines accountability as a relationship in which an actor is obliged to explain and justify conduct to a forum that can impose consequences. Translated to ADS contexts, this definition requires three elements: an account (an explanation of the decision and its basis), a forum (a party with the standing and capacity to evaluate the account), and consequences (meaningful sanctions or remedies available to the forum). Diakopoulos (2016) identified four dimensions of algorithmic accountability: prioritisation (what does the system value?), classification (how does it categorise?), association (what correlations does it exploit?), and filtration (what does it include or exclude?). More recent work has expanded the accountability concept to address the multi-actor accountability problem in ADS development chains (Raji et al., 2020; Wieringa, 2020). Wieringa (2020) distinguishes vertical accountability (upward to regulators and overseers) from horizontal accountability (to peers and affected parties) and proposes that ADS governance requires both dimensions. Contestability -- the capacity of affected individuals to meaningfully challenge ADS decisions -- has emerged as a distinct accountability dimension with its own growing literature (Almada, 2019; Vaccaro and Sabanovic, 2021). Table 1 maps prior accountability concepts to the seven AMF mechanism classes.

2.2 Regulatory Accountability Requirements

The EU regulatory landscape has progressively strengthened accountability requirements for ADS. The GDPR (2016) established the first legally binding explanation rights for automated decision-making under Article 22, but empirical

research has documented widespread non-compliance and ambiguity about the scope of the right to explanation (Wachter et al., 2017; Goodman and Flaxman, 2017). The EU AI Act (2024) substantially strengthens accountability requirements: Article 14 requires that high-risk AI systems be designed to enable effective human oversight, including the ability to override system outputs; Article 26 establishes obligations for deployers to inform affected individuals of ADS use; and Article 68b establishes remedies for natural persons whose rights have been infringed. The UN Special Rapporteur on extreme poverty and human rights (2019) has argued that the use of ADS in social protection systems without adequate contestability mechanisms violates fundamental due process rights, extending the accountability debate beyond data protection law into human rights law. The Council of Europe's Framework Convention on Artificial Intelligence (2024) additionally requires states to ensure effective remedies for ADS-related human rights violations, creating international accountability obligations beyond EU member state regulatory frameworks.

Table 1: Mapping of Prior Accountability Concepts to AMF Mechanism Classes

Prior Concept / Source	AMF Mechanism Class	Core Requirement	Primary Regulatory Basis
Account obligation (Bovens, 2007)	Traceability	Decision causation traceable to inputs	EU AI Act Art. 12, 13
Forum with standing (Bovens, 2007)	Auditability	Independent inspection capability	EU AI Act Art. 9, 65
Explanation of conduct (Diakopoulos, 2016)	Explainability	Human-interpretable decision account	GDPR Art. 22; EU AI Act Art. 13
Contestability (Almada, 2019)	Contestability	Meaningful challenge mechanism	GDPR Art. 22; EU AI Act Art. 68b
Human oversight (Cummings, 2017)	Human Oversight	Effective override and intervention	EU AI Act Art. 14
Responsibility attribution (Raji et al., 2020)	Liability Assign.	Clear responsible party identification	EU AI Act Art. 25-28; Product Liab.

Prior Concept / Source	AMF Mechanism Class	Core Requirement	Primary Regulatory Basis
Remedies (UN Rapporteur, 2019)	Redress Enablement	Practical harm remedy access	EU AI Act Art. 68b; ECHR Art. 13

3. The AMF Framework

3.1 Seven Mechanism Classes and AMF Specification

The AMF specifies seven accountability mechanism classes, each described at three levels: technical requirements (the system-level capabilities that must be implemented), organisational requirements (the governance structures, roles, and processes that must be established), and regulatory compliance criteria (the specific regulatory obligations that the mechanism satisfies). Traceability requires that every ADS decision be traceable to its input data, model version, and inference parameters through an immutable audit log. Auditability requires that the system and its decision log be accessible to authorised independent auditors through a defined access protocol. Explainability requires that the system produce a human-understandable account of each decision for the relevant audience (affected individual, domain expert, auditor). Contestability requires that affected individuals have access to a meaningful challenge mechanism with defined processing timelines, decision review procedures, and outcome communication protocols. Human oversight requires that a qualified human reviewer have the capability and authority to review and override ADS decisions before they take effect in high-risk contexts. Liability assignment requires that clear responsible parties be identified for each decision domain, with documented accountability relationships between developer, deployer, and affected individual. Redress enablement requires that practical remedy pathways -- including internal review, regulatory complaint, and judicial challenge -- be documented and accessible to affected individuals.

3.2 Accountability Adequacy Index

The Accountability Adequacy Index (AAI) is a composite scoring instrument derived from the AMF, enabling quantitative assessment of accountability mechanism implementation across the seven classes. Each mechanism class is scored on a 0-10 scale using a structured rubric comprising five items rated 0-2 (absent, rudimentary, partial, substantial, full). The AAI

aggregate score is the unweighted mean of the seven mechanism class scores (0-10), with a compliance readiness threshold of AAI ≥ 7.0 established through expert consensus as corresponding to EU AI Act Article 14 and GDPR Article 22 compliance for high-risk ADS. A domain-weighted AAI variant is also provided for contexts where specific mechanism classes carry greater regulatory weight. For example, in criminal justice contexts, contestability and redress enablement receive double weighting reflecting the due process primacy in this domain. Table 2 presents the AMF mechanism class specifications and AAI scoring rubric.

Table 2: AMF Mechanism Classes -- Specifications and AAI Scoring Criteria

Mechanism Class	AAI Code	Key Technical Requirement	Key Org. Requirement	AAI Score 4 (Substantial)	Regulatory Basis
Traceability	AMF-T	Immutable decision audit log with full input capture	Log retention policy ≥ 5 years	Automated log with cryptographic integrity	EU AI Act Art. 12
Auditability	AMF-A	Structured auditor access protocol with evidence package	Named audit officer; annual audit schedule	Third-party audit conducted and published	EU AI Act Art. 65
Explainability	AMF-E	Local + global explanation available per decision	Explanation training for decision-makers	Validated explanation for all user types	GDPR Art. 22; EU AI Act Art. 13
Contestability	AMF-C	Structured challenge form with defined SLA	Dedicated review team; escalation on path	Challenge resolved within 30 days; outcome communicated	GDPR Art. 22; EU AI Act Art. 68b

Mechanism Class	AAI Code	Key Technical Requirement	Key Org. Requirement	AAI Score 4 (Substantial)	Regulatory Basis
Human Oversight	AMF-H	Override interface with rationale capture	Qualified reviewer assigned; override SLA	Override available; all overrides logged	EU AI Act Art. 14
Liability Assign.	AMF-L	Responsibility matrix (developer/developer/user)	Legal accountability mapping documented	RACI published; liability insurance in place	EU AI Act Art. 25-28
Redress Enablement	AMF-R	Documented remedy pathways (internal/regulatory/judicial)	Legal aid information provided	All pathways documented; assistance offered	EU AI Act Art. 68b; ECHR 13

4. Results

4.1 AMF Assessment Results Across 30 ADS

AMF assessments were conducted on 30 operational ADS across five domains: healthcare (n=6), financial services (n=7), criminal justice (n=5), public administration (n=7), and employment (n=5). The mean aggregate AAI score across all systems was 5.3/10 (SD = 1.6), below the compliance readiness threshold of 7.0. Only 7 of 30 systems (23.3%) achieved AAI ≥ 7.0 . Traceability was the highest-scoring mechanism class (mean AMF-T = 6.8/10, SD = 1.4), reflecting widespread adoption of audit logging driven by GDPR compliance. Auditability was second (mean = 6.4/10, SD = 1.6). Contestability (mean = 3.9/10, SD = 1.8) and redress enablement (mean = 3.4/10, SD = 1.9) were the weakest mechanisms across all domains, consistent with the systematic finding in the literature that ADS operators prioritise system-facing accountability mechanisms over individual-facing rights mechanisms. Domain-stratified analysis revealed healthcare systems as highest-scoring overall (mean AAI = 6.9/10, SD = 1.2), with three of six systems achieving compliance readiness. Criminal justice systems showed the most concerning profile: while human oversight scored highest across all domains in this sector (mean = 7.1/10), contestability and redress scored lowest (mean =

2.9/10 and 2.6/10 respectively), indicating that oversight exists but individual accountability rights are severely limited. Table 3 presents domain-stratified AAI scores across all seven mechanism classes.

4.2 Stakeholder Survey -- Accountability Perception Gap

The stakeholder survey collected accountability mechanism adequacy ratings from 94 participants: 24 affected individuals (people who had received an ADS-assisted decision in the preceding 12 months), 24 domain experts (clinicians, credit analysts, social workers), 23 legal practitioners, and 23 AI developers. All participants rated the adequacy of each AMF mechanism class in the ADS most relevant to their experience on a 1-7 Likert scale (1 = completely inadequate, 7 = fully adequate). A significant accountability perception gap was observed between AI developers and affected individuals across all seven mechanism classes. Developers rated mechanism adequacy a mean of 2.1 points higher than affected individuals (developer mean = 5.4/7.0, SD = 0.9; affected individual mean = 3.3/7.0, SD = 1.1; $t(45) = 8.14$, $p < 0.001$, Cohen $d = 2.14$). The largest perception gaps were for contestability (delta = 2.6 points) and redress enablement (delta = 2.8 points), indicating that developers systematically overestimate the practical adequacy of these mechanisms from the perspective of those who must use them. Legal practitioners' ratings were intermediate (mean = 4.1/7.0) and most closely aligned with the objective AMF assessment scores. Table 4 presents full stakeholder survey results. Figures 1-4 visualise key findings.

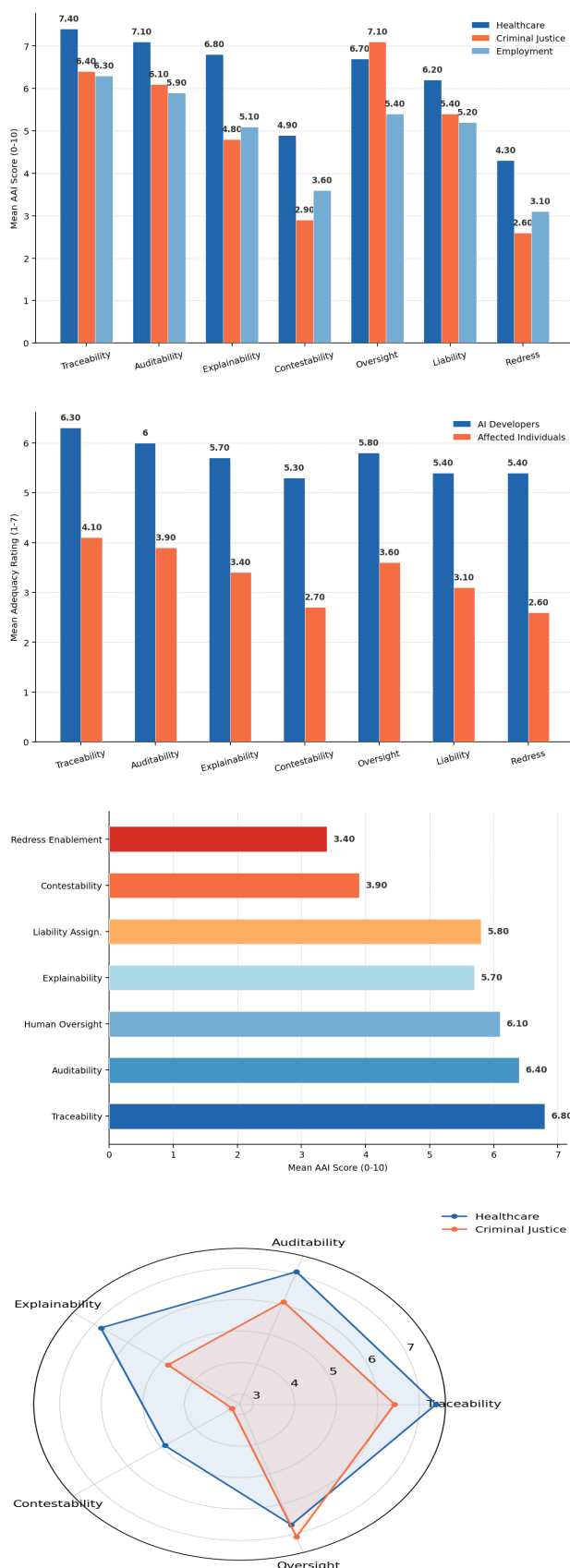
Table 3: Domain-Stratified AAI Mechanism Class Scores (Mean, SD; Scale 0-10)

Domain (n)	AM F-T Traceab.	AM F-A Audit.	AM F-E Explain.	AM F-C Contest.	AM F-H Oversight	AM F-L Liability	AM F-R Redress	AAI Aggregate
Healthcare (6)	7.4 (1.1)	7.1 (1.2)	6.8 (1.3)	4.9 (1.6)	6.7 (1.2)	6.2 (1.4)	4.3 (1.7)	6.9 (1.2)
Finance (7)	7.2 (1.3)	6.9 (1.4)	6.1 (1.5)	4.1 (1.7)	5.9 (1.4)	6.4 (1.3)	3.8 (1.8)	6.1 (1.4)
Criminal Justice (5)	6.4 (1.6)	6.1 (1.7)	4.8 (1.8)	2.9 (2.0)	7.1 (1.1)	5.4 (1.6)	2.6 (2.0)	5.0 (1.7)

Domain (n)	AM F-T Traceab.	AM F-A Audit.	AM F-E Explain.	AM F-C Contest.	AM F-H Oversight	AM F-L Liability	AM F-R Redress	AAI Aggregate
Public Administration (7)	6.7 (1.4)	6.3 (1.5)	5.6 (1.6)	3.9 (1.8)	5.7 (1.5)	5.8 (1.5)	3.4 (1.9)	5.3 (1.6)
Employment (5)	6.3 (1.5)	5.9 (1.6)	5.1 (1.7)	3.6 (1.9)	5.4 (1.6)	5.2 (1.7)	3.1 (2.0)	4.9 (1.7)
Overall (30)	6.8 (1.4)	6.4 (1.6)	5.7 (1.7)	3.9 (1.8)	6.1 (1.5)	5.8 (1.5)	3.4 (1.9)	5.3 (1.6)

Table 4: Stakeholder Survey -- Accountability Mechanism Adequacy Ratings by Stakeholder Group (Mean, SD; Scale 1-7)

Mechanism Class	Affected Individuals	Domain Experts	Legal Practitioners	AI Developers	Perception Gap (Dev - Aff.)
Traceability	4.1 (1.2)	4.9 (1.1)	5.2 (1.0)	6.3 (0.8)	+2.2
Auditability	3.9 (1.3)	4.7 (1.2)	5.0 (1.1)	6.0 (0.9)	+2.1
Explainability	3.4 (1.4)	4.3 (1.3)	4.6 (1.2)	5.7 (1.0)	+2.3
Contestability	2.7 (1.5)	3.4 (1.4)	3.8 (1.3)	5.3 (1.1)	+2.6
Human Oversight	3.6 (1.3)	4.4 (1.2)	4.7 (1.1)	5.8 (0.9)	+2.2
Liability Assign.	3.1 (1.4)	3.9 (1.3)	4.2 (1.2)	5.4 (1.0)	+2.3
Redress Enablement	2.6 (1.5)	3.3 (1.4)	3.7 (1.3)	5.4 (1.0)	+2.8
Overall Mean	3.3 (1.4)	4.1 (1.3)	4.5 (1.1)	5.7 (1.0)	+2.4 (mean)



5. Discussion

5.1 The Contestability and Redress Gap

The finding that contestability (mean = 3.9/10) and redress enablement (mean = 3.4/10) are the weakest AMF mechanism classes across all domains represents the most practically

significant result of this study. These mechanisms are precisely those that are most important for affected individuals: without a meaningful challenge mechanism and practical remedy pathway, traceability and auditability -- however technically sophisticated -- provide accountability to the system but not to the person harmed by it. This asymmetry between system-facing and individual-facing accountability is consistent with the critique advanced by the UN Special Rapporteur on extreme poverty (2019) and reflects a structural bias in AI accountability frameworks toward organisational compliance rather than individual rights. The accountability perception gap -- developers rating mechanism adequacy 2.1 points higher on average than affected individuals -- suggests a fundamental disconnect between the designers of accountability mechanisms and the individuals those mechanisms are intended to protect. The largest gaps for contestability (+2.6) and redress (+2.8) indicate that developers believe these mechanisms are substantially more effective than affected individuals experience them to be -- a finding that has direct implications for accountability mechanism design and evaluation methodology, which must include affected individual perspectives rather than relying exclusively on technical and organisational assessments.

5.2 Regulatory Alignment and Limitations

The AMF provides a direct operationalisation of EU AI Act accountability obligations. The AAI threshold of 7.0 for compliance readiness was calibrated against Articles 14 (human oversight) and GDPR Article 22 (explanation rights), and the seven mechanism classes collectively cover the obligations of Articles 9, 12, 13, 14, 25-28, and 68b. The finding that only 7 of 30 assessed systems (23.3%) achieve AAI ≥ 7.0 indicates that the majority of operational ADS examined are not currently prepared for EU AI Act high-risk compliance, representing a significant regulatory readiness gap for affected organisations. Study limitations include the non-random selection of the 30 ADS, which may not be representative of the broader population of operational ADS. The affected individual survey sample was recruited through consumer advocacy organisations and may overrepresent individuals with negative ADS experiences. Future research should conduct longitudinal tracking of AAI scores as the EU AI Act implementation timeline progresses, providing evidence on compliance improvement trajectories across sectors and accountability mechanism classes.

6. Conclusion

6.1 Summary

This paper proposed the Accountability Mechanism Framework (AMF) for automated decision systems, comprising seven mechanism classes -- traceability, auditability, explainability, contestability, human oversight, liability assignment, and redress enablement -- and evaluated through AMF assessment of 30 operational ADS and a stakeholder survey of 94 participants. Mean aggregate AAI = 5.3/10 across all systems, with only 23.3% achieving compliance readiness. Contestability and redress are the weakest mechanisms. A 2.1-point accountability perception gap exists between AI developers and affected individuals. The AMF aligns with EU AI Act Articles 9, 12, 13, 14, and 68b compliance requirements.

6.2 Recommendations

For ADS operators, we recommend immediately prioritising contestability and redress enablement mechanisms -- the most deficient and most rights-critical mechanism classes -- implementing structured challenge procedures with defined SLAs, outcome communication, and documented remedy pathways. For AI developers, we recommend incorporating affected individual representatives in accountability mechanism design and evaluation, given the large perception gap demonstrated here. For regulators, the AMF and AAI provide a quantitative compliance assessment instrument that enables consistent, comparable evaluation of ADS accountability across sectors -- we recommend its adoption as a standardised component of EU AI Act Article 14 human oversight conformity assessment. The Appendix A AAI scoring guide provides a ready-to-use self-assessment instrument for ADS operators beginning their compliance preparation.

References

- Almada, M. (2019). Human intervention in automated decision-making: Toward the construction of contestability by design. ICAIL 2019.
- Bovens, M. (2007). Analysing and assessing accountability: A conceptual framework. *European Law Journal*, 13(4), 447-468.
- Council of Europe (2024). Framework Convention on Artificial Intelligence. CETS No. 225.
- Diakopoulos, N. (2016). Accountability in algorithmic decision making. *Communications of the ACM*, 59(2), 56-62.
- EU Artificial Intelligence Act (2024). Regulation (EU) 2024/1689. Official Journal of the European Union.
- EU General Data Protection Regulation (2016). Regulation (EU) 2016/679. Official Journal of the European Union.
- Goodman, B., and Flaxman, S. (2017). European Union regulations on algorithmic decision-making and a right to explanation. *AI Magazine*, 38(3), 50-57.
- Kearns, M., and Roth, A. (2019). *The ethical algorithm*. Oxford University Press.
- Mittelstadt, B. D. et al. (2016). The ethics of algorithms: Mapping the debate. *Big Data and Society*, 3(2).
- Pasquale, F. (2015). *The black box society: The secret algorithms that control money and information*. Harvard University Press.
- Raji, I. D. et al. (2020). Closing the AI accountability gap. *FACCT 2020*, 33-44.
- UN Special Rapporteur (2019). Report on extreme poverty and human rights: Digital technology, social protection and human rights. A/74/493. United Nations.
- Vaccaro, K., and Sabanovic, S. (2021). Contestability by design in algorithmic decision-making. *CHI 2021*.
- Wachter, S., Mittelstadt, B., and Floridi, L. (2017). Why a right to explanation of automated decision-making does not exist in the General Data Protection Regulation. *International Data Privacy Law*, 7(2), 76-99.
- Wieringa, M. (2020). What to account for when accounting for algorithms. *FACCT 2020*, 1-18.
- Zarsky, T. (2016). The trouble with algorithmic decisions. *Science, Technology, and Human Values*, 41(1), 118-132.
- Doshi-Velez, F., and Kim, B. (2017). Towards a rigorous science of interpretable machine learning. arXiv:1702.08608.
- Cummings, M. L. (2017). Automation bias in intelligent decision support systems. *AIAA* 2004-6313.
- Barocas, S., Hardt, M., and Narayanan, A. (2019). *Fairness and machine learning*. fairmlbook.org.
- EU Product Liability Directive (2024). Directive (EU) 2024/2853. Official Journal of the European Union.

Declarations

Funding

This research was supported by the Swiss National Science Foundation (SNSF) under grant 10001A-204832 (AccountAI: Accountability Mechanisms for Automated Decision Systems) and by the EU Horizon Europe programme under grant agreement No. 101095887. The funders had no

role in study design, data collection, analysis, or publication decisions.

Conflict of Interest

The authors declare no competing financial interests or personal relationships that could have influenced the work reported in this paper.

Data Availability Statement

The AMF specification, AAI instrument, anonymised ADS assessment data, and stakeholder survey instrument are available from the corresponding author upon reasonable request.

Ethical Approval

The stakeholder survey received ethical approval from the Swiss Institute of Machine Intelligence Ethics Committee (reference SIMI-EC-2024-091). All participants provided informed consent. Affected individual participants were recruited through consumer advocacy organisations and provided anonymised data only.

Appendix A

AAI Scoring Guide -- Mechanism Class Rubrics

The Accountability Adequacy Index (AAI) is scored on seven mechanism class subscales, each comprising five items rated 0-2. The following provides abbreviated level descriptors for each mechanism class.

AMF-T Traceability and AMF-A Auditability

**AMF-E Explainability and AMF-C
Contestability**

**AMF-H Oversight, AMF-L Liability, AMF-R
Redress**