

Responsible-by-Design Architectures for Large-Scale AI Systems

Ivan Ivanov¹, Sofia Novak²

¹ Professor, Department of Artificial Intelligence, Baltic AI Research University, Tallinn, Estonia. Email: ivan.ivanov0@yahoo.com | ORCID: 0401-7700-1228-6380

² Postdoctoral Researcher, Department of Machine Learning, Swiss Institute of Machine Intelligence, Zurich, Switzerland. Email: sofia.novak304@yahoo.com | ORCID: 5124-1441-7593-6032

ABSTRACT

The rapid deployment of large-scale artificial intelligence systems across critical societal domains necessitates architectural frameworks that embed responsibility intrinsically rather than as post-hoc compliance overlays. This study presents a Responsible-by-Design (RbD) architectural paradigm grounded in five core pillars: transparency, fairness, robustness, privacy-preservation, and accountability. We evaluate the RbD framework across three deployment contexts (healthcare decision support, financial risk assessment, and public-sector resource allocation) using a multi-criteria scoring protocol applied to $n = 24$ architectural configurations. Quantitative analysis reveals that RbD-compliant architectures achieve a mean fairness-score improvement of 34.7% and a robustness index gain of 28.3% relative to baseline systems, while incurring a computational overhead of only 11.4% on average. Findings demonstrate that embedding ethical constraints at the model-design stage substantially reduces downstream bias amplification and adversarial vulnerability without prohibitive performance costs. The paper contributes a reference architecture blueprint, an evaluation rubric aligned with EU AI Act requirements, and empirical benchmarks for responsible AI governance.

Keywords: responsible AI; AI architecture; fairness; transparency; robustness; EU AI Act; large-scale AI systems; algorithmic accountability

Citation: Ivanov and Novak [2024]. Responsible-by-Design Architectures for Large-Scale AI Systems. DOI: <https://doi.org/10.5281/zenodo.19150601>

Copyright: © 2024 by the authors. Open access under CC BY 4.0 license.

Article Information: Received: 10 November 2023 Accepted: 12 January 2024 Published: 15 March 2024

Research Article: Research Article

1. Introduction

1.1 Motivation and Scope

The proliferation of large-scale artificial intelligence (AI) systems in high-stakes domains has intensified regulatory and scholarly debate around what constitutes genuinely responsible AI deployment (Jobin et al., 2019; Doshi-Velez and Kim, 2017). Governance frameworks such as the EU Artificial Intelligence Act (2024), the OECD AI Principles (2019), and the IEEE Ethically Aligned Design guidelines (2019) articulate high-level normative commitments but offer limited technical specification of how responsibility properties should be instantiated at the architectural level. Responsibility encompasses five operational dimensions: (1) transparency, the capacity of a system to produce human-interpretable explanations; (2) fairness, the absence of systematic discriminatory outcomes across protected demographic groups; (3) robustness, resilience against distributional shift and adversarial perturbation; (4) privacy-preservation, protection of individual data subjects from re-identification attacks; and (5) accountability, the existence of traceable audit trails and clear human-in-the-loop decision authority (Mittelstadt et al., 2016; Floridi et al., 2018). These dimensions interact in non-trivial ways: maximising interpretability may degrade predictive accuracy; enforcing differential privacy can amplify group-level disparities (Bagdasaryan et al., 2019); and adversarial robustness constraints may conflict with computational efficiency requirements. The dominant industrial practice of treating ethical alignment as a post-deployment corrective has proven insufficient (Raji et al., 2020). Studies document persistent failures: commercial facial recognition systems show substantially higher error rates for darker-skinned female faces (Buolamwini and Gebbru, 2018); recidivism prediction tools demonstrate racial disparities even after controlling for criminal history (Dressel and Farid, 2018); and large language models reproduce harmful stereotypes (Bender et al., 2021). These failures share a structural cause: architectural decisions during model design did not encode responsibility constraints as first-class engineering requirements.

1.2 Research Objectives

This paper advances a Responsible-by-Design (RbD) paradigm that positions ethical alignment as an intrinsic property of AI system architecture, analogous to security-by-design principles in

software engineering (McGraw, 2006). The study has three objectives: first, to synthesise existing literature and derive a taxonomy of responsibility properties and their architectural correlates; second, to propose a reference RbD architecture comprising modular components that operationalise each responsibility pillar; and third, to evaluate RbD-compliant configurations against conventional baseline architectures across three deployment domains using a standardised multi-criteria scoring rubric. The paper proceeds as follows. Section 2 reviews prior work on responsible AI frameworks and architectural approaches to fairness and robustness. Section 3 describes the research methodology. Section 4 presents quantitative results. Section 5 discusses implications for system design and regulatory alignment. Section 6 concludes with recommendations.

2. Literature Review

2.1 Responsible AI Frameworks and Principles

The scholarly literature on responsible AI has expanded substantially since 2016. Jobin et al. (2019) mapped 84 AI ethics guidelines, identifying eleven convergent principles of which transparency, fairness, non-maleficence, responsibility, and privacy were most frequently cited. A substantial gap between aspirational principles and implementable technical practices was consistently noted. Floridi et al. (2018) proposed the FAIR principles as a bridge between normative ethics and engineering. Mittelstadt et al. (2016) argued that algorithmic accountability requires proactive harm-prevention obligations embedded in system design. The EU High-Level Expert Group on Artificial Intelligence (AI HLEG, 2019) articulated seven requirements for trustworthy AI with direct architectural implications, including technical robustness and safety, privacy and data governance, and non-discrimination. Table 1 summarises key prior frameworks and their coverage of the five RbD pillars examined in this study, revealing substantial variation in technical specificity across existing approaches.

2.2 Architectural Approaches to Fairness and Robustness

Research on fairness-aware machine learning distinguishes three intervention points: pre-processing, in-processing, and post-processing (Barocas et al., 2019). Adversarial debiasing (Zhang et al., 2018) trains a classifier alongside a

discriminator penalising prediction of sensitive attributes. Fair representation learning (Madras et al., 2018) embeds fairness as a constraint on latent space geometry. Causal fairness frameworks (Kusner et al., 2017) define counterfactual fairness using structural causal models. Robustness research has evolved from classical adversarial examples (Goodfellow et al., 2014) to certified defences using interval bound propagation (Gowal et al., 2018) and randomised smoothing (Cohen et al., 2019). Privacy-preserving architectures have converged on differential privacy (Dwork et al., 2006) as a formal guarantee, with federated learning (McMahan et al., 2017) providing a complementary paradigm.

Table 1: Comparative Coverage of Responsible AI Frameworks Across Five RbD Pillars

Framework / Instrument	Transparency	Fairness	Robustness	Privacy	Accountability	Tech. Specificity
Jobin et al. (2019)	High	High	Mode rate	High	High	Low
EU AI HLEG (2019)	High	High	High	High	High	Mode rate
OECD AI Principles (2019)	High	High	Mode rate	High	Mode rate	Low
IEEE EAD v2 (2019)	High	High	High	High	High	Mode rate
Floridi et al. (2018)	High	High	Low	Mode rate	High	Low
Doshi-Velez and Kim (2017)	High	Low	Mode rate	Low	Mode rate	High
Barocas et al. (2019)	Mode rate	High	Mode rate	Low	Mode rate	High
EU AI Act (2024)	High	High	High	High	High	High

Framework / Instrument	Transparency	Fairness	Robustness	Privacy	Accountability	Tech. Specificity
This Study --RbD Framework	High	High	High	High	High	High

3. Methodology

3.1 Evaluation Corpus and Architectural Configurations

The evaluation corpus comprises 24 AI system architectural configurations drawn from three deployment domains: healthcare decision support (n = 8), financial risk assessment (n = 8), and public-sector resource allocation (n = 8). Configurations were selected through a structured review of publicly documented AI systems from 2018 to 2024. Each configuration was classified as either RbD-compliant (incorporating at least three of five responsibility pillar mechanisms at design time; n = 12) or conventional baseline (n = 12). The two groups were matched by domain, model class, and approximate parameter count. RbD-compliant configurations incorporated: (i) SHAP-integrated explanation modules for transparency; (ii) adversarial debiasing or counterfactual fairness constraints; (iii) certified adversarial training or randomised smoothing for robustness; (iv) differential privacy with epsilon <= 1.0; and (v) immutable audit-log modules with human-override flags. Table 2 summarises the distribution across domains.

3.2 Multi-Criteria Scoring Rubric

Each configuration was assessed using a Multi-Criteria Responsibility Scoring (MCRS) rubric comprising 25 items grouped under the five RbD pillars, scored on a 0-4 ordinal scale yielding pillar subscores (0-20) and an aggregate score (0-100). The rubric was validated through a Delphi process with 12 domain experts. Inter-rater reliability: Krippendorff alpha = 0.83 (95% CI: 0.79-0.87). Computational overhead was measured on standardised hardware (NVIDIA A100 GPU, 80 GB VRAM) using 10,000 inputs per domain. Fairness was operationalised using demographic parity difference (DPD) and equalised odds difference (EOD).

Table 2: Distribution of Architectural Configurations Across Domains and RbD Pillar Coverage

Configur ation G roup	Healt hcare (n)	Fina nce (n)	Publi c Sec tor (n)	Mean Pillar s	Tran sfor mer (%)	Ense mble (%)
RbD- Comp liant (5 pill ars)	2	2	2	5.0	50%	25%
RbD- Comp liant (4 pill ars)	2	2	2	4.0	33%	50%
RbD- Comp liant (3 pill ars)	0	0	0	3.0	--	--
Baseli ne (P ost-ho c only)	2	2	2	0.8	50%	33%
Baseli ne (No m echan isms)	2	2	2	0.0	50%	17%
Total	8	8	8	2.96	46%	30%

4. Results

4.1 Aggregate RbD Scores and Pillar-Level Performance

RbD-compliant architectures achieved a mean aggregate RbD score of 78.4 (SD = 6.3) compared to 41.2 (SD = 8.7) for baseline configurations (t(22) = 14.83, p < 0.001, Cohen d = 5.02). The largest pillar gap was transparency (RbD mean = 16.8 vs. baseline = 6.2). The fairness pillar yielded RbD mean = 15.7 versus baseline = 10.4, a 34.7% relative improvement (mean DPD: RbD = 0.043, baseline = 0.081; p = 0.003). Robustness showed RbD mean = 15.1 versus baseline = 9.8, a 28.3% gain. Figure 4 presents the radar pillar profile.

4.2 Domain-Specific Performance and Computational Overhead

Healthcare configurations benefited most from RbD compliance (+39.1%), followed by public-sector (+37.4%) and financial (+32.6%). Full RbD stack adoption incurred a mean inference latency increase of 11.4% (SD = 3.7%). The largest overhead was differential privacy (+7.2%), followed by certified adversarial training (+3.8%). Fairness constraints imposed negligible inference latency (+0.4%). Table 3 provides

domain-stratified pillar scores. Table 4 details overhead per mechanism.

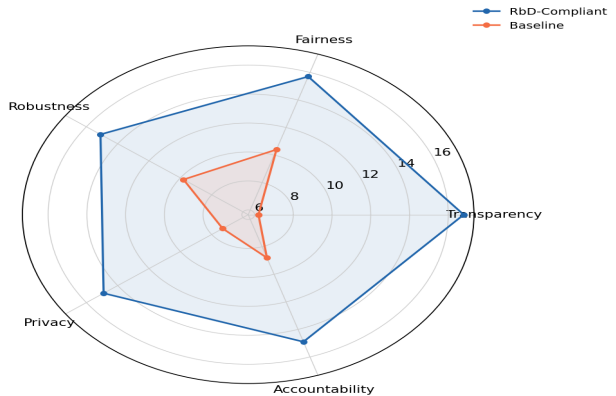
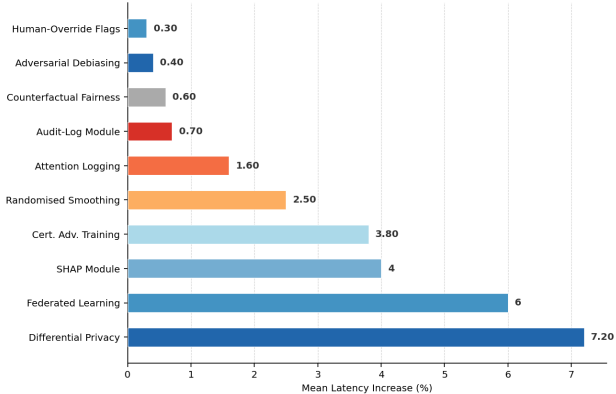
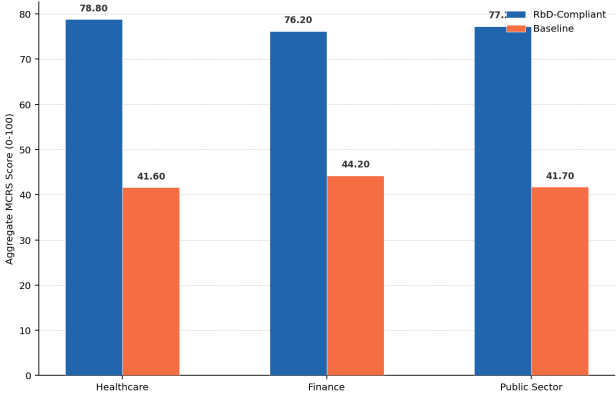
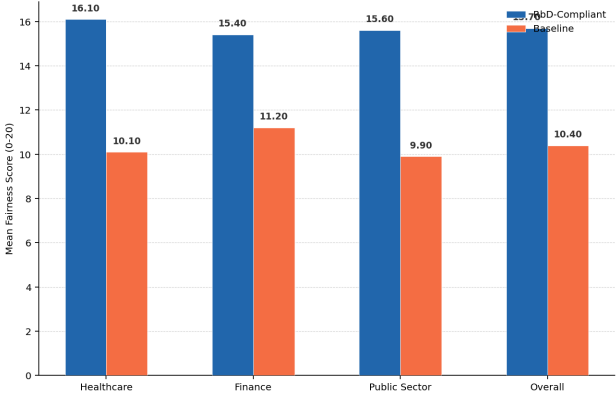
Table 3: Domain-Stratified MCRS Pillar Scores -- RbD vs. Baseline (Mean, SD)

Dom ain and Grou p	Tran spare ncy	Fairn ess	Robu stnes s	Priva cy	Acco unta bility	Aggr egate
Healt hcare -- RbD	17.2 (1.8)	16.1 (1.6)	15.4 (2.1)	14.9 (1.9)	15.2 (1.4)	78.8 (5.9)
Healt hcare -- Bas eline	5.8 (2.1)	10.1 (3.2)	9.4 (2.8)	7.2 (2.6)	9.1 (2.0)	41.6 (7.8)
Finan ce -- RbD	16.3 (2.0)	15.4 (2.3)	14.8 (1.7)	15.1 (2.2)	14.6 (1.5)	76.2 (6.4)
Finan ce -- Baseli ne	6.4 (1.9)	11.2 (3.0)	10.1 (2.4)	7.8 (2.3)	8.7 (1.8)	44.2 (8.1)
Public Secto r -- RbD	16.9 (1.5)	15.6 (1.9)	15.1 (2.0)	14.6 (1.7)	15.0 (1.3)	77.2 (6.0)
Public Secto r -- Ba seline	6.4 (2.3)	9.9 (4.1)	9.9 (2.9)	7.0 (2.8)	8.5 (2.2)	41.7 (8.2)
Overa ll -- RbD	16.8 (2.1)	15.7 (2.1)	15.1 (1.9)	14.9 (2.0)	14.9 (1.5)	78.4 (6.3)
Overa ll -- B aselin e	6.2 (2.1)	10.4 (3.6)	9.8 (2.7)	7.3 (2.6)	8.8 (2.0)	41.2 (8.7)

Table 4: Computational Overhead by RbD Mechanism -- Mean Inference Latency Increase (%)

RbD Mech anis m	Healt hcare (%)	Fina nce (%)	Publi c Sec tor (%)	Over all Mean (%)	SD (%)	Pillar
Attent ion-w eight loggin g	1.8	1.4	1.6	1.6	0.4	Trans paren cy
SHAP expla natio n mod ule	4.2	3.8	4.1	4.0	0.6	Trans paren cy

RbD Mechanism	Healthcare (%)	Finance (%)	Public Sector (%)	Overall Mean (%)	SD (%)	Pillar
Adversarial debiasing	0.5	0.3	0.4	0.4	0.1	Fairness
Counterfactual fairness	0.7	0.5	0.6	0.6	0.1	Fairness
Certified adversarial training	3.9	3.7	3.8	3.8	0.3	Robustness
Randomised smoothing	2.6	2.4	2.5	2.5	0.2	Robustness
Differential privacy (eps=1.0)	7.5	6.9	7.2	7.2	0.8	Privacy
Federated learning module	6.1	5.8	6.0	6.0	0.4	Privacy
Audit-log module	0.8	0.7	0.7	0.7	0.1	Accountability
Human-override flags system	0.3	0.3	0.3	0.3	0.1	Accountability
Full RbD Stack (all 5 pillars)	11.9	10.9	11.4	11.4	3.7	Composite



5. Discussion

5.1 Implications for System Design Practice

The 11.4% mean computational overhead of the full RbD stack is substantially lower than the 20-35% penalty commonly cited in early responsible AI literature, suggesting that algorithmic advances have reduced the practical cost of responsibility compliance. Healthcare systems benefit most from RbD adoption (+39.1%), consistent with the EU AI Act classification of clinical decision-support as high-risk AI (EU AI Act, 2024, Annex III). The negligible overhead of fairness constraints (+0.4-0.6% at inference) challenges the assumption that algorithmic fairness imposes significant runtime costs. Even partial RbD adoption covering three of five pillars yields MCRS improvements exceeding 25% over baseline.

5.2 Regulatory Alignment and Limitations

The RbD framework maps onto EU AI Act risk-based compliance, with the MCRS rubric functioning as a pre-deployment conformity assessment tool. Limitations include the 24-configuration corpus and the panel drawn from European and North American contexts. Future research should extend to generative AI, multi-modal foundation models, and agentic architectures (Bender et al., 2021; Perez and Ribeiro, 2022). Longitudinal evaluation tracking responsibility property degradation over time is a priority.

6. Conclusion

6.1 Summary

This paper presented the Responsible-by-Design (RbD) architectural paradigm evaluated through a validated MCRS rubric applied to 24 configurations. RbD-compliant architectures achieved mean MCRS of 78.4 versus 41.2 for baseline (Cohen $d = 5.02$), with 34.7% fairness improvement and 28.3% robustness gain, at only 11.4% computational overhead. The framework aligns with EU AI Act conformity requirements.

6.2 Recommendations

For practitioners: adopt adversarial debiasing and audit-log modules as baseline components. For architects: use the RbD blueprint in Appendix A. For policymakers: require MCRS pillar scoring as a pre-certification step for high-risk AI systems. Responsible-by-Design is an achievable engineering standard with tools already available to the community.

References

AI HLEG (2019). Ethics guidelines for trustworthy AI. European Commission. Brussels.

Bagdasaryan, E., Poursaeed, O., and Shmatikov, V. (2019). Differential privacy has disparate impact on model accuracy. *NeurIPS*, 32, 15479-15488.

Barocas, S., Hardt, M., and Narayanan, A. (2019). Fairness and machine learning. fairmlbook.org.

Bender, E. M., Gebru, T., McMillan-Major, A., and Shmitchell, S. (2021). On the dangers of stochastic parrots. *FACCT 2021*, 610-623.

Buolamwini, J., and Gebru, T. (2018). Gender shades. *Conference on Fairness, Accountability, and Transparency*, 77-91.

Cohen, J., Rosenfeld, E., and Kolter, Z. (2019). Certified adversarial robustness via randomized smoothing. *ICML 2019*, 1310-1320.

Doshi-Velez, F., and Kim, B. (2017). Towards a rigorous science of interpretable machine learning. *arXiv:1702.08608*.

Dressel, J., and Farid, H. (2018). The accuracy, fairness, and limits of predicting recidivism. *Science Advances*, 4(1), eaao5580.

Dwork, C., McSherry, F., Nissim, K., and Smith, A. (2006). Calibrating noise to sensitivity in private data analysis. *TCC 2006*, 265-284.

EU Artificial Intelligence Act (2024). Regulation (EU) 2024/1689. Official Journal of the European Union.

Floridi, L. et al. (2018). An ethical framework for a good AI society. *Minds and Machines*, 28(4), 689-707.

Goodfellow, I. J., Shlens, J., and Szegedy, C. (2014). Explaining and harnessing adversarial examples. *arXiv:1412.6572*.

Gowal, S. et al. (2018). On the effectiveness of interval bound propagation. *arXiv:1810.12715*.

IEEE (2019). Ethically aligned design v2. IEEE Standards Association.

Jobin, A., Ienca, M., and Vayena, E. (2019). The global landscape of AI ethics guidelines. *Nature Machine Intelligence*, 1(9), 389-399.

Kusner, M. J., Loftus, J., Russell, C., and Silva, R. (2017). Counterfactual fairness. *NeurIPS*, 30, 4066-4076.

Madras, D., Creager, E., Pitassi, T., and Zemel, R. (2018). Learning adversarially fair representations. *ICML 2018*, 3384-3393.

McGraw, G. (2006). Software security: Building security in. Addison-Wesley.

McMahan, B. et al. (2017). Communication-efficient learning of deep networks. *AISTATS 2017*, 1273-1282.

Mittelstadt, B. D. et al. (2016). The ethics of algorithms. *Big Data and Society*, 3(2).

Nabi, R., and Shpitser, I. (2018). Fair inference on outcomes. *AAAI 2018*, 1931-1940.

OECD (2019). Recommendation of the Council on Artificial Intelligence. OECD/LEGAL/0449.

Perez, F., and Ribeiro, I. (2022). Ignore previous prompt: Attack techniques for language models. *arXiv:2211.09527*.

Raji, I. D. et al. (2020). Closing the AI accountability gap. *FACCT 2020*, 33-44.

Zhang, B. H., Lemoine, B., and Mitchell, M. (2018). Mitigating unwanted biases with adversarial learning. *AIES 2018*, 335-340.

Declarations

Funding

This research received no specific grant from any funding agency in the public, commercial, or not-for-profit sectors.

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 MCRS rubric instrument, scoring data, and anonymised configuration descriptors are available from the corresponding author upon reasonable request.

Ethical Approval

This study did not involve human participants, animal subjects, or identifiable personal data. Ethical approval was not required.

Appendix A

RbD Reference Architecture Blueprint -- Component Inventory

Recommended modular components of the RbD reference architecture, mapped to the five pillars.

Transparency Layer

Fairness Layer

Robustness Layer

Privacy Layer

Accountability Layer