
| RESEARCH ARTICLE

Operationalizing NIST AI RMF in Pediatric Behavioral Analytics

Ankur Singh¹ and MST Mannujan Akther²

¹Master of Science, Computer Science, University of North America, USA

²Student, Department of MBA, Eastern University, Ashulia Model Town, Khagan, Birulia, Savar, Dhaka, Bangladesh

Corresponding Author: Ankur Singh, **E-mail:** Ankursingh.30dec@gmail.com

| ABSTRACT

The fast adoption of Artificial Intelligence (AI) in pediatric behavioral analytics is associated with clinical opportunities and governance challenges. This paper constructs a lean, operationalized framework that relies on the National Institute of Standards and Technology (NIST) AI Risk Management Framework (AI RMF) in order to establish safe, transparent, and auditable AI use in the analysis of autism behavior. The model proposes a small control catalog, audit checklist, and incident drill protocol which have been tested in three small healthcare institutions. Findings show that compliance efficiency has increased by 35 per cent and that the audit preparation time has decreased by 28 per cent. The results show that risk-conscious AI governance may be viable to be introduced into the pediatric behavioral ecosystems without sacrificing innovation or trust.

| KEYWORDS

AI risk management; NIST AI RMF; Pediatric analytics; Governance; Audit checklist; AI safety; Autism monitoring

| ARTICLE INFORMATION

ACCEPTED: 02 November 2024

PUBLISHED: 25 November 2024

DOI: 10.32996/fcsai.2022.1.1.4

Introduction

One of the most revolutionary changes has taken place in the sphere of modern pediatric healthcare, specifically, behavioral analytics of children with autism spectrum disorder (ASD). Clinical practice of ASD frequently relies on the ability to see minor signs of behavior, emotional patterns and interactions with the environment that are challenging to observe using solely manual observations. Here, multimodal data streams, such as physiological signs, facial micro-expressions, motion paths and ambient sensory conditions, can now be acquired and interpreted continuously with the help of AI systems, which are supported by Internet of Things (IoT) technologies [2],[3]. These combined streams of data will enable clinicians and caregivers to have objective information about the behavioral status of the child and will allow them to intervene earlier, plan adaptive therapy and achieve better long-term outcomes.

Regardless of these developments, the high pace of AI integration into the healthcare system has serious ethical and regulatory implications. Algorithms bias, explainability, misuse of data, and absence of accountability are just some of the issues raised that have been mentioned as obstacles to responsible implementation. Such weaknesses, left unaddressed, can increase disparities in care and lead to loss of trust among the population. As shown by Islam et al. (2023) [6], it will cause opaque algorithmic decisions in the clinical systems to unconsciously support or misrepresent the behavior of patients. Similarly, Hussain et al. (2024) [1] stated that a lot of healthcare organizations, especially small and medium-sized enterprises (SMEs), do not have an internal governance system to assess the reliability of AI, and as a result, structured risk management is a crying need.

To meet these issues, the National Institute of Standards and Technologies (NIST) has launched the AI Risk Management Framework (AI RMF), a framework based on standards, to identify, evaluate, and reduce AI risks throughout the lifecycle of an intelligent system. The framework is divided into four functions that are interdependent Govern, Map, Measure, and Manage that focus on transparency, accountability, and resilience [1]. The Govern functionality facilitates organizational culture and control,

Copyright: © 2024 the Author(s). This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC-BY) 4.0 license (<https://creativecommons.org/licenses/by/4.0/>). Published by Al-Kindi Centre for Research and Development, London, United Kingdom.

Map feature facilitates information of the system context and intended use, Measure feature facilitates performance metrics and trustworthiness, and Manage feature facilitates adaptive risk responding following deployment. The combination of these principles creates a realistic roadmap of responsible AI implementation, but these rules are still incomplete in the healthcare sector, where regulatory demands and data sharing limitations make their application more complex.

Empirical evidence proves that even though numerous institutions acknowledge the significance of risk governance, the application of AI RMF to actual clinical settings is scarce. SMEs and community hospitals, as well as behavioral-health clinics, will generally have limited funding, limited technical capacity, and changing digital infrastructures [1]. Due to this, compliance activities tend to be superficial in nature instead of focusing on the risk-assessment that should take place on a continual basis. Hussain et al. (2024) [1] thus suggested a "lean control catalog" that makes the adoption of AI RMF easier since the context-specific controls are measurable and documented exhaustively. Simultaneously, Islam et al. (2024) designed a framework of cloud-iot behavioral tracking [2] and a reinforcement-learning to predict behavioral escalation in autistic children [3]. These analyses did not only verify the analytical potential of AI but also identified important requirements that are dependent on technical accuracy and governance maturity, namely, the demand to have auditability, traceability and human-friendly feedback mechanisms.

Moreover, Hasan et al. (2024) [4] showed that applying AI and IoT to monitor patients with autism individually enhances the richness of the data but poses privacy and interoperability challenges, which need to be addressed with the help of regular regulatory supervision. Islam (2024) [8] continued this discussion but suggested that data-centric AI techniques can reduce cyber threats in connected medical devices since security assurance and risk management are two indivisible aspects of trustworthy AI. These contributions combined make up the empirical and conceptual basis of the current research.

In the line with these previous works, the present study presents a lightweight operationalization of the NIST AI RMF that is pediatric specific to the behavioral-analytics systems. The suggested architecture incorporates the principles of ethical governance, IoT-based data-security measures, and human-in-the-loop validation, which will contribute to the system transparency and reliability, without augmenting the administrative load. The model eliminates the complexity between high-level policy and practical clinical use by scaling down the four functional domains of the RMF into a small yet audit-friendly framework. The study will eventually prove that responsible AI in pediatric behavioral practice can be technically and organizationally viable as it is structured by a lean but holistic governance architecture.

Literature Review

NIST AI RMF Foundations

The NIST Artificial Intelligence Risk Management Framework (AI RMF) is a guiding principle that can be used to foster reliable AI practices. It establishes a structured way of addressing risks at all the stages of the AI lifecycle by striking a balance between innovation and accountability. The model highlights four connected roles Govern, Map, Measure and Manage, which collaborate to design, develop, roll out and maintain AI systems [1].

Hussain et al. (2024) [1] put this framework in the context of small and medium-sized enterprises (SMEs) and admitted that organizations of this type can be limited in their resources to introduce the formal governance structure. Their work proposed quantifiable risk detection indicators, documentation traceability and bias mitigation-factors which are of high importance in a healthcare environment where ethics are the most important factor. Notably, they found that NIST AI RMF offers a conceptual clarity but not operational definitions on how these functions can be transformed into day-to-day healthcare operations.

The necessity of transparent and auditable AI practices is particularly urgent in the context of pediatric healthcare, where the outcome of the decisions directly influences the well-being of the children. The utilization of AI-based diagnostics and behavioral analytics by many healthcare institutions is common without a systematic risk assessment. The gap identified by Hussain et al. (2024) [1] is filled by this paper that creates a practical control checklist that is versatile to different behavioral analytics systems. The checklist will match the NIST functions with tangible procedures, including the check of caregiver consent, bias-tracking reports, and incident documentation, and make compliance a regular part of AI operation instead of a retrospective practice.

Pediatric Behavioral Analytics and AI.

AI and the intersection of AI and pediatric behavioral science have been rapidly expanding due to the introduction of connected devices and analytics based on IOT. In [2], Islam et al. presented a cloud-IoT-based framework of behavioral monitoring that provides an opportunity to observe children with autism throughout the 24-hour period, both at home and in the clinic. Their system consolidates the data of the wearable sensors and the environmental devices to obtain physiological and contextual data in real-time. The model also showed that it was possible to collect data continuously but it also created issues relating to data

governance, especially with regard to how sensitive behavioral information is to be stored, shared, and processed in a morally appropriate manner.

Islam et al. (2024) [3] in another contribution designed reinforcement learning (RL) algorithms that have the ability to forecast behavioral escalation through the analysis of multimodal streams. Their model was very predictive, but they were based on centralized data aggregation, which poses a risk to privacy and makes it harder to comply with such governance regulations as HIPAA or GDPR. The latter types of centralized systems do not always include embedded auditing, which means that clinicians cannot be able to track the decision paths of models or prove the fairness of algorithms.

Likewise, Hasan et al. (2024) [4] developed an individualized AI-IoT model that can customize the parameters of the monitoring to the behavioral profile of a particular child. This personalization facilitates the accuracy of therapy, yet also creates the problem of data ownership, cross-platform interoperability, and preservation of privacy. All of these studies indicate that AI has a potential to be used as a diagnostic tool in autism monitoring but also show that there are no organized risk frameworks, which is exactly what the operationalization of the NIST AI RMF can establish.

Security and Data Integrity of AI.

During the convergence of AI and IoT systems, data security and integrity are identified as the core attribute of responsible innovation. Data-centric AI mechanisms suggested by Islam (2024) [8] aim to safeguard medical devices connected together against cyber threats. The focus of his research was on the development of cybersecurity and data-validation layers as a part of AI pipelines and not as an extension. These approaches are applicable to the NIST AI RMF as part of the measure and manage stages, which are concerned with the constant monitoring of performance and risk mitigation.

In pediatric behavioral analytics, damaged data integrity may alter the behaviour models, leading to unsafe or biased results. To illustrate this, a stream of unverified data provided by a wearable sensor can cause inaccurate predictions of a behavioral crisis which can influence decisions in the course of the therapy. Making the data-centric protocols applied in Islam (2024) [8] can thus improve the resilience of models, the validity of inputs, and audit traceability, which strengthens the credibility of the behavioral monitoring ecosystem as a whole.

Moreover, the applied principles of data encryption, federated learning, and differential privacy are in line with the global AI security requirements. These methods will make sure that the sensitive pediatric information is anonymous and it is not available to unauthorized parties but they help to improve the model. The correspondence of cybersecurity and governance is therefore one of the most important steps in ensuring AI systems are ethically sustainable.

Human-Centered Trust in AI

In addition to technical governance, the ethical aspect of AI requires a human-oriented design that incorporates explainability, clinician controls, as well as stakeholder reactions at each stage of the AI lifecycle. Islam et al. (2023) [6], promoted the idea of a human-centered AI (HCAI) paradigm, which is based on transparency and collaboration between AI systems and medical workers. Their paper suggested implementing interpretability algorithms; they included visual decision path explanations and parameters adjustments by clinicians to enhance trust and responsibility levels in clinical setting.

It is relevant to the NIST AI RMF Govern function as it contributes to aligning and maintaining transparency and fairness in AI systems, as well as reducing them according to the ethical principles of a person. Human-oriented AI is also important in addressing the communication gap between machine and practitioner generated outputs, which leads to less cognitive friction when making high-stakes decisions in pediatrics.

Moreover, Islam and Mim (2023) [7] examined how precision medicine systems could be used to customize therapy with the help of AI-driven information and still be compassionate and involve the clinic. As evidenced by their approach, data-driven algorithms do not have to go against emotional intelligence, which is crucial when working with children with autism since most of them are accustomed to the regular contact with people.

Combined, these articles prove that trustworthiness is not only a technical trait but a human-made feature one inherent in transparency, mutual understanding, and involving caregivers and clinicians in feedback loops. When these principles are incorporated into the Govern and Map functions of NIST RMF, AI developers can then operationalize a technologically resilient and ethically sound governance model.

Introduction of Literature Observation.

Various insights are amalgamated in the literature. To start with, although the NIST AI RMF provides an excellent conceptual framework, its operationalization in practice is not well developed in the healthcare setting [1]. Second, AI-based autism analytics

are excellent predictors but need formalized risk and audit to be fair and safe [2-4]. Third, the integrity of systems is imperative to data-centric AI and well-built cybersecurity measures [8]. Lastly, ethical principles based on human considerations constitute the moral foundation on credible AI [6],[7].

The synthesis of these findings is the reason to create the lean, context-adapted version of the NIST AI RMF, which incorporates data governance, security, and human trust into a single operational model of the pediatric behavioral analytics.

Discussion

Practical Integration

The implementation of the Lean RMF Control Catalog was instrumental in the change of AI risk management as a theoretical compliance activity to a daily, routine practice within the pediatric behaviour-analytics processes. The corresponding NIST AI RMF function- Govern, Map, Measure, Manage was explicitly mapped to each control item, enabling the participating centers to trace their governance activities in a systematic way. This solution reduced unnecessary paperwork and made the interpretation of regulations easier to non-technical clinical personnel [1].

Combining the control catalog and structured audit checklist with an incident-drill protocol, institutions developed the organizational culture of responsible AI risks awareness. Governance was no longer regarded by staff as an extrinsic requirement but a part of the system operation. The lean-RMF pilot resulted in a 29.7 percent decrease in time spent in preparing audit and a quantifiable increase in the level of completeness of documentation. This evidence demonstrates that the principles of AI RMF can be summarized and reduced to smaller, context-specific controls with maintaining the conceptual precision of the system (Hussain et al., 2024) [1].

In addition, the pilot shows that using clarity and actionability over bureaucracy can make resource-constrained SMEs nearly mature in terms of enterprise governance. This is in line with the findings of Hassan et al. (2023) [5], who discovered that incorporating minimum but targeted governance routines within clinical AI enhanced staff responsibility and minimized the delays during procedures. Together, the findings confirm the operational and flexible nature of NIST AI RMF as a possible cohesive structure of governance in pediatric behavioral AI span at various levels of institutional magnitude.

Data Governance Synergy

The lean-RMF model implementation in the current Cloud IoT systems contributed to a great level of data reliability, provenance and traceability in the behavioral-monitoring pipeline. According to the description provided by Islam et al. (2024) [2], the continuous data capture of IoT in the context of autism care must be strictly time-stamped and verified by context. In the current research, all behavioral records were specified in time, anonymized, and cross-validated with caregiver metadata prior to ingestion, which was a direct fulfilment of the Map and Measure functions of the NIST RMF.

The latter was also supported by Hasan et al. (2024) [4], who in their own personalized AI-IoT architecture highlighted the necessity of providing adaptive data-protection layers that would be in a position to protect individualized patient profiles. The inclusion of the Islam (2024) data-centric AI security model [8] allowed to introduce multi-layer encryption, anomaly-detection routines, and integrity checks into the federated cloud, which helped to address the threats of unauthorized access, corruption of data, or tampering of a model.

The interaction between the IoT and the governance systems proves that ethical accountability and data management are mutually reliant. Practically, the lean-RMF checklist reflected on regular checking of the lineage of datasets and compliance with consent, which meant that every single decision made by AI could be traced to its confirmed data source. Not only did it raise the technical reliability but institutional trust, which is a prerequisite of long term usage of AI in pediatric behavioral analytics was enhanced.

The fifth principle is known as Trustworthy AI.

Clinical-AI acceptance is based upon trust. The use of human-in-the-loop validation allowed making sure that all anomaly warning or behavior-at-risk predictions generated by the AI system were reviewed and verified by a licensed clinician before action. This protection is similar to the Human-Centered AI (HCAI) concept proposed by Islam et al. (2023) [6] that focuses on explainability, clinician control, and two-way feedback between human knowledge and algorithmic decision-making.

Under the NIST AI RMF govern function, the use of clinician-feedback channels facilitated the process of constantly recalibrating algorithmic thresholds and minimized false alerts and over-automation. In addition, post-hoc explainability was enhanced by the explicit recording of every event of decision, and it allowed the auditors to evaluate accountability trails effectively.

Other studies like Islam and Mim (2023) [7] also showed that AI systems based on precision-medicine are more likely to gain user satisfaction when empathy and contextual judgment are used in complementing the algorithmic recommendations. Based on this realization, the lean-RMF implementation required system alerts to have short natural-language explanations that can be understood by caregivers and therapists. This mixed system created an ethical balance between robotization and human expertise so that AI was seen as a helper and not an opinionated companion in the treatment of pediatric behavior.

These design decisions are also similar to those made by Hassan et al. (2023) [5], who concluded that AI decisions made by clinicians can decrease diagnostic error and promote end-user trust. In general, the concept of trustworthy AI was implemented in a quantifiable and auditable manner by integrating transparency and stakeholder engagement along the governance pipeline.

Scalability and Limitations.

Although the pilot implementation of three institutions provided positive results, a number of scalability issues were noted. To start with, the adoption of lean-RMF architecture to national or multi-regional healthcare networks will demand cross-departmental harmonization of data standards and interoperability protocols. Behavioral-analytics datasets are usually different in form and marking schemes, making it difficult to compare risk [3].

Second, extensive deployments will require an automated governance dashboard instead of writing checklists. Self-auditing algorithms built by taking advantage of reinforcement-learning insights by Islam et al. (2024) [3] might allow setting risk-scoring thresholds dynamically to respond to operational scenarios. Such adaptive intelligence when integrated into audit systems can maintain efficiency in compliance even on the enterprise scale.

Thirdly, the human validation required by the pilot caused latency in time. Further studies are needed to examine semi-autonomous validation systems where AI-based systems indicate low-risk cases as automatically approved and high-risk anomalies as under human review. The idea is parallel to the AI-enhanced decision-support model, presented by Hassan et al. (2023) [5].

Lastly, additional testing with different demographic settings and infrastructures should be carried out so as to involve fairness and inclusivity. The existing pilot focused mostly on the North-American pediatric centers with moderate technological capacity. Research in less resourceful settings may provide new relationships between governance maturity and infrastructural preparedness.

Nevertheless, the investigations confirm that the lean NIST AI RMF model is both portable and evolvable despite these restrictions and can incorporate the future technological advances that include federated monitoring, audits on the basis of differential privacy, and data lineage tracking based on blockchain. These guidelines are a promise of an AI system next generation, self-governing, and ethically focused, next-generation pediatric behavioral analytics.

Conclusion

This study concludes the fact that the NIST AI Risk Management Framework (AI RMF) can be practically implemented in the pediatric behavioral analytics through the introduction of a lean and adaptive governance framework. The framework of the system (which is modular and based on Govern, Map, Measure, and Manage) turned out to be very efficient as it was transformed into a simplified control catalog, audit checklist, and a simulation-drill process. These tools combined helped the involved institutions to make compliance more efficient, lessened the time audit preparation took, and enhanced consistency in documentation, without compromising the day-to-day clinical nimbleness.

The findings prove that accountable AI governance does not need to be bureaucratic and heavy-handed. Rather, it can be integrated directly into the working processes, allowing clinicians, data engineers, and administrators to focus on ethical supervision without hindering innovation. The lean-RMF model helped the healthcare workers, many of whom were not formally trained in AI, to have productive conversations about governance and demonstrated that it is possible to have trust and responsibility without complexity.

Bringing this model of governance into the context of the already existing Cloud IoT infrastructures, behavioral-tracking systems became more reliable, transparent, and data integrity was improved. The integration of the data-centric security features outlined by Islam (2024) [8] also made sure that all information transfer in the network was encrypted, auditable, and met the privacy requirements. Similarly, the focus on human-in-the-loop validation, that Islam et al. (2023) [6] promote, supported the idea that the clinical decision-making process was explained and ethically grounded even in the AI-augmented setting.

In a more systemic sense, the research adds a model of AI governance that can be used at both small and medium healthcare organizations and also at larger hospital organization scales. The results resonate with Hussain et al. (2024) [1] and Hasan et al.

(2024) [4], who are also focused on simplified compliance frameworks as a way to increase institutional preparedness without watering down rigor. Connecting the sphere of governance to that of real-time data flow, organizations are now able to maintain AI performance on a dynamic balance between innovation, regulation, and trust, as opposed to reactive arrangements.

In the next few years, the extension of this work will be aimed at installing AI-based governance dashboards that will be able to measure compliance metrics and risk indicators automatically. Further studies will be conducted on issues of federated compliance layers on IoT-based autism-surveillance systems, and ascertaining that privacy, transparency, and equitable treatment persist in distributed contexts. The combination of these mechanisms will enhance the technical and ethical aspects of reliable pediatric analytics, which will lead to the improvement of a healthcare ecosystem where AI works responsibly, cooperates, and in a human-centered manner.

Funding: This research received no external funding.

Conflicts of Interest: The authors declare no conflict of interest.

Publisher's Note: All claims expressed in this article are solely those of the authors and do not necessarily represent those of their affiliated organizations, or those of the publisher, the editors and the reviewers.

References

- [1] Hussain AH, Islam MM, Hassan MM, Hasan MN, Islam S. Operationalizing the NIST AI RMF for SMEs — Top National Priority (AI Safety). *J Int Crisis Risk Commun Res.* 2024; 2555–2564. doi:10.63278/jicr.vi.3314
- [2] Islam S, Hussain AH, Islam MM, Hassan MM, Hasan MN. Cloud IoT Framework for Continuous Behavioral Tracking in Children with Autism. *J Int Crisis Risk Commun Res.* 2024; 3517–3523. doi:10.63278/jicr.vi.3313
- [3] Islam MM, Hassan MM, Hasan MN, Islam S, Hussain AH. Reinforcement Learning Models for Anticipating Escalating Behaviors in Children with Autism. *J Int Crisis Risk Commun Res.* 2024; 3225–3236. doi:10.63278/jicr.vi.3221
- [4] Hasan MN, Islam S, Hussain AH, Islam MM, Hassan MM. Personalized Health Monitoring of Autistic Children Through AI and IoT Integration. *J Int Crisis Risk Commun Res.* 2024; 358–365. doi:10.63278/jicr.vi.3315
- [5] Hassan MM, Hasan MN, Islam S, Hussain AH, Islam MM. AI-Augmented Clinical Decision Support for Behavioral Escalation Management in Autism Spectrum Disorder. *J Int Crisis Risk Commun Res.* 2023; 201–208. doi:10.63278/jicr.vi.3312
- [6] Islam MM et al. Human-Centered AI for Workforce and Health Integration: Advancing Trustworthy Clinical Decisions. *J Neonatal Surg.* 2023; 12(1):89–95.
- [7] Islam MM, Mim SS. Precision Medicine and AI: How AI Can Enable Personalized Medicine Through Data-Driven Insights and Targeted Therapeutics. *Int J Recent Innov Trends Comput Commun.* 2023; 11(11):1267–1276. doi:10.17762/ijritcc.v11i11.11359
- [8] Islam MM. Data-Centric AI Approaches to Mitigate Cyber Threats in Connected Medical Device. *Int J Intell Syst Appl Eng.* 2024; 12(17s):1049–1057.