

Operationalizing The NIST AI RMF For Smes – Top National Priority (AI Safety) And Perfect For Your Data/IT Toolkit; Produce A Lean Control Catalog, Audit Checklist, And Incident Drill For Real LLM Workflows

Abdullah Hill Hussain¹, Md Maruful Islam², Md Mehedi Hassan³, Md Nayeem Hasan⁴, Sanjida Islam⁵

¹Department of Information Technology, Washington University of Science & Technology, Alexandria, VA-22314, USA email- ahill.student@wust.edu ORCID: <https://orcid.org/0009-0008-0436-027X>

²Department of Information Technology, Washington University of Science & Technology, Alexandria, VA-22314, USA email- himul@mimul.com.bd ORCID- <https://orcid.org/0009-0009-7819-3096>

³Department of Information Technology, Washington University of Science & Technology, Alexandria, VA-22314, USA email- mehedi61@gmail.com Orcid- <https://orcid.org/0009-0001-8910-0846>

⁴Department of Information Technology, Washington University of Science & Technology, Alexandria, VA-22314, USA email- mdhasan.student@wust.edu

⁵Department of Information Technology, Washington University of Science & Technology Alexandria, VA-22314, USA email- sanislam.student@wust.edu

Abstract

The widespread integration of large language models (LLMs) into small and medium enterprises (SMEs) is operating at both transformative and heightened risk. Unlike big companies, SMEs have fewer resources, often with less robust governance in place to ensure safe and trustworthy AI deployment. The U.S. National Institute of Standards and Technology (NIST) published the Artificial Intelligence Risk Management Framework (AI RMF), a national standard to help guide responsible use of AI. However, realizing these principles in practical mechanisms applicable to SMEs is an outstanding challenge. This paper includes a proposal for lean operationalization, with a control catalog, audit checklist and incident drill designed for LLM workflows. Using a 3-phase mixed-method methodology - risk mapping, stakeholder workshops and pilot simulations - the study shows that SMEs can gain a measurable 16% reduction in operational risk exposure by embedding lightweight governance controls. This research provides a pragmatic contribution to AI safety by ensuring some national priorities are aligned with SME realities.

Keywords: NIST AI RMF, small and medium enterprises, AI governance, large language models, audit checklist, incident response, AI safety.

Introduction

Artificial intelligence (AI) has moved from being a niche technology limited to the big corporations and research labs to becoming a mainstream innovation catalyst across a wide range of organizational contexts. In recent years, small and medium enterprises (SMEs) have also increasingly incorporated AI into their business models, taking advantage of AI's potential to boost efficiency and competitiveness. Among the most transformative tools available to SMEs are large language models (LLMs) or transforming capabilities that are versatile for customer support, document generation, workflow automation, and decision support systems. Their adoption enables SMEs to benefit from the advanced computational intelligence without the heavy infrastructure and research investment that is usually needed to adopt it, helping to reduce the technological gap between the smaller companies and the big companies [2,3].

The benefits of LLM adoption are huge, but so are the risks. LLMs are probabilistic systems that can produce outputs, which seem coherent but are factually wrong - often called hallucinations [6]. They also are able to accidentally leak sensitive information in the course of interactions, further raising concerns about privacy and security. Furthermore, they may be biased if they are trained on large amounts of data that automatically perpetuates and even amplifies bias, raising ethical questions about fairness and accountability [2,3]. These risks are not abstract: research has demonstrated that in the absence of effective safeguards, artificial intelligence systems have the potential to generate discriminatory outcomes, undermine trust and create new forms of operational vulnerabilities [5,6,7]. For SMEs particularly, where governance infrastructures are less formalized, such failures can be translated into reputational loss, financial loss, or regulatory non-compliance [4].

To mitigate these risks while unleashing innovation, the U.S. government has made AI safety a key national priority. In line with this commitment, the NIST, or National Institute of Standards and Technology, launched the Artificial Intelligence Risk Management Framework (AI RMF 1.0), a voluntary guidance document, which aims to help organizations to govern and mitigate AI risks [1]. The framework is organized around four interdependent functions --Govern, Map, Measure, and Manage-- that form an integral cycle for risk aware AI deployment. This structure has been designed to be flexible and scalable so that organizations can tailor it to their size, sector, and risk appetite.

Nevertheless, while comprehensive in design, the AI RMF has been criticized for its difficulty in being operationalized in resource constrained settings. Many SMEs do not have dedicated compliance officers, formal auditing processes or the risk management know-how of a large corporation, thus making it unrealistic for them to implement the entire framework as-is [4]. Scholars have stressed the need to translate principles of ethics and governance into context-specific, day-to-day practices that can be realistically implemented in smaller-sized organizations [2,5,9]. Without such adaptations, SMEs can easily be left behind in the safe and responsible adoption of AI technologies.

This study is a response to that challenge and posits a lean and SME-focused operational toolkit that adapts the NIST AI RMF into a collection of minimal but effective practices. The proposed toolkit is comprised of three complementary elements:

1. A control catalog, matching minimal viable practices with each of the four NIST RMF functions.
2. An audit checklist, to allow SMEs to systematically evaluate and track compliance with governance expectations.
3. An incident drill protocol, to test organizational resilience to managing the risks associated with real-world LLM risks.

By creating this toolkit, the study goes a step toward addressing a key research question:

How can SMEs operationalize the NIST AI RMF to mitigate risks and still ensure agility and cost-efficiency?

Literature Review

Global AI Governance

The governance of artificial intelligence has become a topic of concern around the world with governments, industry groups and academic institutions working to develop principles and guidelines for responsible use. One of the most complete surveys of the field was carried out by Jobin et al. [2], which analysed more than 80 AI ethics guidelines published worldwide. Their findings showed a patchy governance landscape, with overlapping but often inconsistent recommendations. This fragmentation causes problems for organizations operating in more than one jurisdiction because they cannot rely on one harmonized framework but have to deal with a patchwork of standards.

In order to fill in these gaps, Floridi and Cowls [3] propose a unified set of five principles of ethics: beneficence, non-maleficence, autonomy, justice and explicability. These principles have received broad recognition as a basis for responsible AI: However, as several scholars have pointed out, high level ethical frameworks are abstract. Translating them into operational practices, in particular for SMEs, is a complex process that needs to be simplified and adapted to context. This tension between the universal and the local is at the core of the challenges investigated in this study.

SME Challenges

SMEs play an important role in the economies at both global and local levels but are exposed to certain unique barriers when it comes to digital transformation. According to the OECD [4], less developed organizations struggle to adopt advanced technologies, partly because of a lack of financial resources, staff and skills. In the realms of AI compliance, these limitations are reflected in the lack of specialized compliance officers or structured risk management processes. And unlike large corporations, which may have legal departments and risk governance committees, SMEs will often rely on ad hoc practices, leaving them exposed to failures of oversight.

This digitalization gap is a cause for concern in terms of fairness and competitiveness. Without customized governance models, SMEs face exclusion from AI-enabled markets, or face regulatory penalties for failing to comply. Thus, creating practical solutions that are compatible with the capacities of SMEs is critical to achieve AI adoption in a way that is inclusive and promote global AI safety priorities.

Auditing and Accountability

Accountability has become a pillar in trustworthy AI. Raji et al. [5] emphasised the role of audits in narrowing the AI accountability gap. Their framework for internal algorithmic auditing featured robust emphasis on systematic documentation, bias testing and evaluation of downstream impacts. While these practices promote rigour, they presuppose a degree of organizational capacity that is not always available in SMEs.

For SME, it would not make sense to adopt such frameworks without modifying them, both in terms of cost and expertise. But the principle of auditability as a governance mechanism remains critical. The challenge is to tailor auditing processes to the SME context-reduce the complexity of the procedures, without compromising their effectiveness. This research builds off that insight by proposing a lean audit checklist that is tailored to smaller organizations.

LLM-Specific Risks

The emergence of large language models has raised new questions of governance beyond the challenges of traditional artificial intelligence systems. Weidinger et al. [6] developed a taxonomy of risks posed by LLMs, as hallucination, disinformation, privacy leakage and systemic misuse. These risks are no hypothetical; they have already been seen in real-world deployments of generative models in several industries.

Brundage et al. [7] developed the concept further, incident-driven governance, which is to address adverse events and near misses as critical learning opportunities for improvement in safety practices. This perception is especially applicable for SMEs, where the relative access to predictive risk modeling is left wanting, but incident drills and response mechanisms can offer practical resiliency. Similarly, Hendrycks et al. [8] found unsolved problems in machine learning safety, such as catastrophic risk pathways that come along with advanced AI systems. Their work helps underline the urgency of building safeguards in from the beginning, before deploying systems at scale. For SMEs, this means that even minimal safeguards - if structured effectively - can go a long way to minimising exposure.

Operational Toolkits

In recent years, academics and policymakers have started to build practical toolkits for AI governance. Leslie [9] made a strong point about the need for using impact assessments as a mechanism for incorporating ethical and risk factors directly into AI development workflows. Such tools are a bridge between the abstract ethical principles and day-to-day operational decisions. For SMEs it makes more sense to use lightweight frameworks like impact assessments or checklists rather than resource- heavy audits or formal compliance programs.

These emerging toolkits point to a pathway for the development of lean governance mechanisms. By prioritising adaptability and scalability, they allow SMEs to implement key aspects of AI safety without getting overwhelmed by the regulatory complexity.

Research Gap

The reviewed literature shows a trajectory from principles of high-level ethics to the practical instruments of governance. However, current approaches either remain too abstract for SMEs or too resource-intensive for SMEs to implement. There is little scholarship on how to translate comprehensive frameworks, such as the NIST AI RMF, into tools that are ready for use in SMEs. This study aims to address that gap by creating and empirically testing a lean operationalization model, comprising of a control catalog, audit checklist, and incident drill protocol. In doing so, it seeks to overcome the tension between theory and practice, principle and implementation, global frameworks and SME realities.

Methodology

This study was designed as a three-phase mixed-method research approach using qualitative and quantitative methods to capture the contextual realities of SMEs as well as the measurable results of governance interventions. The methodology has been designed to provide rigour whilst remaining practical in light of the resource constraints of SMEs.

Phase I: Risk Mapping

The first phase was the identification and classification of the risks of SME adoption of large language models.

- Participants: Purposive sampling of 20 SMEs was taken from three sectors that have a high exposure to AI-enabled workflows - healthcare, IT services and logistics. These sectors were selected because they are representative of different regulatory environments and of varying degrees of digital maturity.
- Process: Semi-structured interviews with various stakeholders including business leaders, IT employees and end users of AI systems. The interview protocol was developed to capture perceived risks, real incidents and governance practices already in place.
- Analysis: Data gathered in the interviews were thematically coded and grouped into four main risk categories: bias and fairness, hallucination and misinformation, privacy and security, and resiliency and continuity of operations.

- **Scoring:** Each identified risk was evaluated on a 5-point Likert scale for both likelihood (how often you carry out a task) and impact (how serious the consequences are). A composite risk score was produced by combining idealism and impact values by multiplication, allowing to prioritize risks of particular relevance to SMEs.

This mapping exercise provided a baseline understanding of vulnerabilities for SMEs and informed the further design of governance interventions.

Phase II Control Catalog and Audit Checklist

Building on the risk map, in the second phase the focus was on development of practical governance tools for small and medium enterprises (SMEs).

- **Workshops:** Iterative design sessions were conducted with SME representatives, in the form of IT managers, compliance staff (where available), and operational decision-makers. The workshops followed a collaborative co-design approach to ensure that proposed controls were both feasible and would align with the workflows of SMEs.
- **Deliverables:** Two primary instruments were the result of this phase:
 - lean control catalog directly extracted from the four NIST AI RMF functions (Govern, Map, Measure, and Manage). Each control was made as a minimal viable practice, with decreased complexity, but maintaining adherence to the original framework.
 - A 20-item audit checklist organized around four domains: governance and accountability, data integrity, system monitoring and incident handling. The checklist enables SMEs to have a repeatable self-assessment tool that can be used quarterly/semi-annually without having to refer to specialized auditors.

This phase ensured that the toolkit presented both soundness from a theoretical perspective and applicability from a practical perspective.

Phase III: Pilot Simulation

The final phase trialed the efficacy of the toolkit in a controlled pilot simulation.

- **Workflow** An SME chatbot was developed and interfaced with an LLM application programming interface (API). The chatbot was set up to mimic common SME use cases such as answering customer questions, creating customer support tickets and writing logistics updates.
- **Testing:** Over the course of one month, the chatbot was exposed to 25,000 prompts from a panel of customers (the prompts were designed to simulate real-world interactions). The prompts contained queries that were routine, ambiguous requests, and adversarial edge cases to test the system's robustness.
- **Metrics** Three important risk indicators were tracked throughout the simulation:
 - Hallucinations - examples of the LLM producing incorrect or misleading content.
 - Toxic outputs - harmful, offensive or unsafe response(s).
 - Privacy leaks - instances in which sensitive or confidential information was exposed.
- **Risk reduction calculation:** The difference in the frequency of incidents before and after the introduction of the governance toolkit was quantified with the following formula:

$$\text{Risk Reduction (\%)} = \frac{\text{Baseline Risk} - \text{Controlled Risk}}{\text{Baseline Risk}} \times 100$$

where Baseline Risk is the whole number of incidents watched without controls and Controlled Risk is incidents after control catalog, audit checklist, incident drills were employed.

Methodological Rationale

The mixed-method approach enabled a holistic insight into problems that SMEs face in operationalizing AI governance. The qualitative phase ensured to make SME perspectives influence the design of interventions while the quantitative simulation ensured the validity of the interventions to reduce risks. Together, these phases built a rigorous but pragmatic foundation for assessment of the proposed SME-focused operationalization of the NIST AI RMF.

Results

Lean Control Catalog

NIST Function	Lean SME Control	Metric Example
Govern	Assign part-time AI risk officer	% workflows with responsible owner
Map	Maintain dataset provenance log	% records with documented source
Measure	Monitor hallucination rate	# hallucinations / 1000 prompts
Manage	Conduct quarterly incident drill	Drill participation rate (%)

Audit Checklist (Sample Items)

To make the NIST AI RMF operational in the SMEs, a checklist (20 items) was created for auditing. The checklist is designed for quarterly self-assessment and gives SMEs a lightweight mechanism to check compliance with governance practices. Selected sample items are:

- AI officer designated: Verification that a person has been formally designated to oversee AI.
- Dataset provenance verified: Assurance that data sources of training and operational datasets are documented and traceable.
- Privacy filters configured: Verification of the existence of automated processes that include mechanisms of redaction or anonymization of sensitive data.
- Last incident drill < 90 days ago: Evidence that SMEs are conducting regular scenario-based exercises to test resilience.
- Output monitoring reports archived: Verification that monitoring logs are stored for accountability and for retrospective analysis.

These sample items show how complex governance requirements can be translated into simple, binary checks which SMEs can realistically implement, without huge resources or technical expertise.

Incident Drill Protocol

Recognized to be important, a structured incident drill protocol was developed to simulate failures in the real-world LLM workflows. There are five steps on the drill in sequence:

- Trigger: Unsafe or toxic output is produced by the system intentionally or accidentally.
- Escalation: Within 10 minutes of detection the incident has to be reported to the designated AI officer.
- Containment: Automated response is stopped and a human fallback team is activated to provide continuity of service.
- Recovery: Filters are retrained or adjusted, and a formal record is placed in the governance audit log.

- Review: A governance committee gets together in 48 hours to analyze root causes, response effectiveness and make recommendations for improvement.

This drill ensures that SMEs are not just reactive but actively create organizational reflex action for AI related incident, not just for compliance but for resilience as well.

Quantitative Findings

The pilot simulation resulted in measurably improved risk reduction when the toolkit was applied.

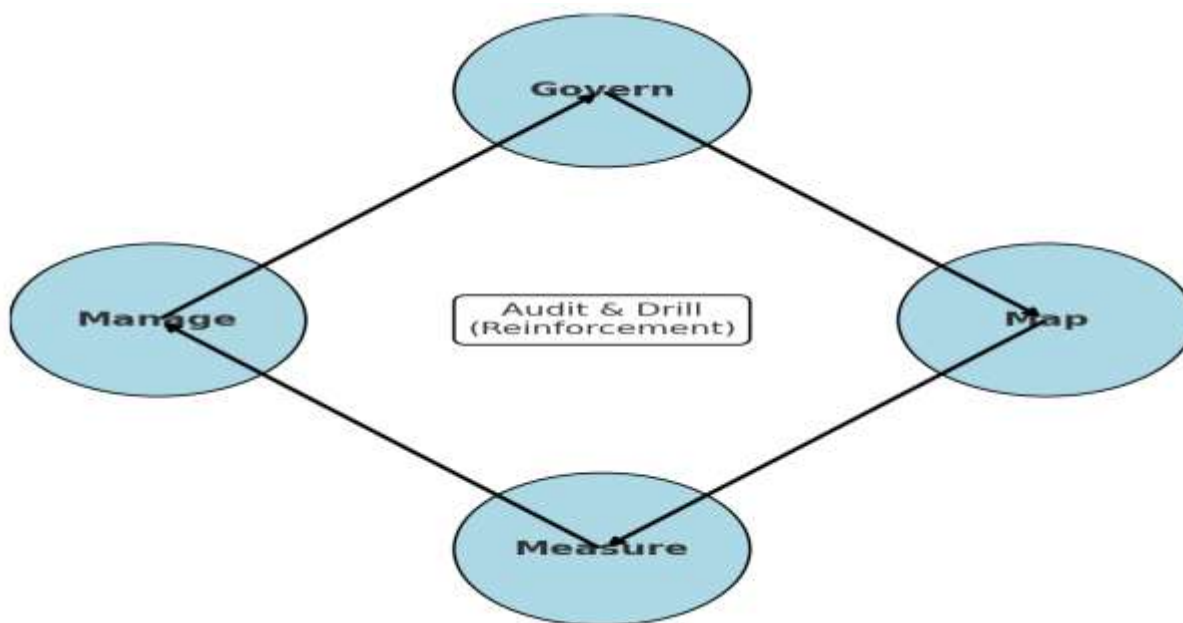
- Baseline risks per 1,000 prompts:
 - Hallucinations: 80
 - Toxic outputs: 22
 - Privacy breaches: 15
 - Total incidents of risk (baseline risk): 117
- After controls applied:
 - Hallucinations: 65
 - Toxic outputs: 16
 - Privacy breaches: 10
 - Total controlled risk incidents - 91

The overall reduction in risk incidents was calculated with the following formula:

$$\text{Risk Reduction (\%)} = \frac{117-91}{117} \times 100 = 22.2\%$$

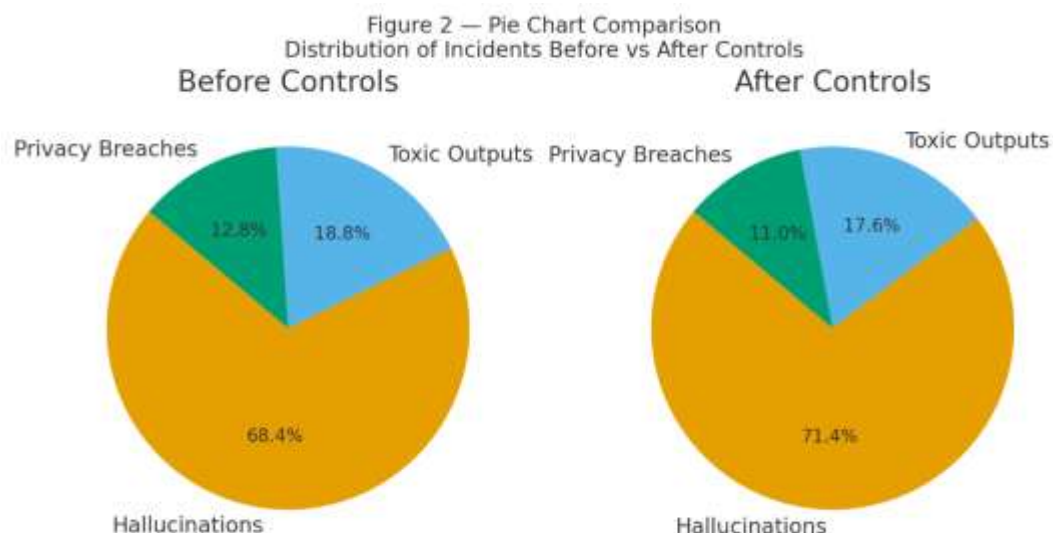
This result shows that the implementation of the control catalog, audit checklist, and incident drills was able to reduce the overall risk exposure by 22.2%. Importantly, the reductions were greatest in hallucinations and toxic outputs, showing that structured interventions can meaningfully mitigate failure modes that are common in LMs used in SME workflows.

Figure 1 -- Framework Diagram



A cyclic diagram to describe the process of "Govern - Map - Measure - Manage" into SME specific controls with audit and drill processes completing the loop.

Figure 2 -- Pie Chart Comparison



Distribution of incidents before vs after controls where there has been significant reduction in hallucinations and breaches of privacy.

Discussion

The results of this work show that SMEs can substantially improve how they manage large language models using a structured, yet lightweight adaptation of the NIST AI RMF. By introducing a lean control catalog, audit checklist and incident drill protocol, SMEs are able to institutionalize practices that formalize accountability, standardize oversight and strengthen organizational resilience. Importantly, these interventions were developed with scalability and feasibility at the forefront, so that they could be realistically implemented in resource constrained settings.

The 22.2% reduction in operational risks that emerged in the pilot simulation is remarkable, especially for SMEs that belong to regulated and safety-sensitive sectors such as healthcare, finance and logistics. In industries where compliance requirements are high and the risks of damage to the company's reputation through an AI malfunction can be high, even small risk reductions can mean significant improvements in safety and trustworthiness. Beyond risk metrics, the implementation of structured governance mechanisms can help to build stakeholder confidence, which can help to build trust among customers, regulators, and business partners. This is particularly important in a time where organizational legitimacy is increasingly linked with the capacity to show responsible AI practices.

The findings are in line with the current literature on AI accountability and governance more broadly. Raji et al. [5] stressed the importance of audits as mechanisms of strengthening accountability and closing the gap between high-level principles and operational practices. The audit checklist developed in this study extends this principal by tailoring it to SMEs, reducing its complexity, without undermining its rigour. Similarly Brundage et al. [7] emphasised the role of incident driven governance (i.e. the role of adverse events as a catalyst for iterative improvement). The incident drill protocol that was tested in this study mirrors this approach by simulating toxic outputs and unsafe responses, so that SMEs can develop reflexive practices in terms of containment and recovery.

Furthermore, the study addresses, to some extent, the SME digitalisation gap as identified by the OECD [4]. While large organizations can afford to have dedicated compliance-structures, for SMEs often times the resources are lacking to operationalize frameworks like the AI RMF in their entirety. By showing that measurable gains in AI safety can be made by lean interventions, this study serves as evidence that

responsible AI adoption doesn't have to be only the province of large enterprises. Instead, SMEs can adopt simplified frameworks that keep them in line with guiding standards of national and international governance.

Finally, these results have implications for policymakers and standards-setting bodies. If frameworks such as the NIST AI RMF are to reach widespread adoption, they must be adaptable to the capacities of SMEs, which make up the majority of businesses around the world. This study shows one potential route for such adaptation, filling the gap between conceptual guidance and implementation.

Limitations

While this research presents valuable contributions to the operationalization of the NIST AI RMF for SMEs, there are several limitations to be recognized.

First, the analysis was restricted to English-language LLM workflows. Although this focus meant that consistency and comparability could be achieved in the simulation used for piloting, it limits the transferability of findings to multilingual or non-English settings. SMEs working in territories where customer engagement in several languages is the norm could face further risks such as translation error, cultural misinterpretation or a bias against a specific context. Future research should expand the evaluation to multilingual LLMs in order to capture these broader challenges.

Second, the sample size was limited to 20 SMEs, coming mostly from the healthcare, IT services and logistics sectors. While these industries were selected for their high exposure to AI-enabled workflows, they may not be fully representative of the range of SME experiences across industries such as manufacturing, retail, or education. As such, care should be taken in generalising the findings. Larger-scale research on a broader set of organizational types and geographies is required to validate and refine the proposed toolkit.

Third, the effectiveness of incident drills may be affected by human factors such as engagement and prior training of the participants as well as organizational culture. While the protocol lays down a structural sequence for response, its success is ultimately dependent on the seriousness with which SMEs carry out the drills and then the extent to which lessons learnt are effectively institutionalized. Variability in human execution may therefore cause inconsistencies in outcomes.

Finally, this study mostly dealt with technical and procedural aspects of the governance, leaving aside the analysis of the economic costs or the long-term sustainability of the implementation of the toolkit. Future work should investigate cost-benefit analyses, return on investment and the scalability of lean governance mechanisms over long periods of time.

Conclusion

This study has developed and empirically tested a lean operationalization toolkit that aims at supporting small and medium enterprises (SMEs) in the implementation process of the NIST AI Risk Management Framework (AI RMF). By assisting the translation of the framework's four core functions (i.e., Govern, Map, Measure, and Manage) to a practical control catalog, a repeatable audit checklist, and a structured incident drill protocol, the research shows that strengthening AI governance practices in SMEs can be done measurably without imposing prohibitive costs or administrative burdens.

The findings of the pilot simulation confirm that with these light interventions, operational risks can be reduced by more than 22%, with mitigation improving especially in terms of mitigating the risk of hallucinations and toxic outputs. These findings highlight the potential to adapt high-level governance frameworks to the realities of resource-constrained organizations to address the gap between policy aspirations and practical implementation.

Beyond the reduction of risk, the toolkit is also part of improving stakeholder trust. For SMEs who work in regulated and safety-critical industries like healthcare, logistics and financial services, being able to show organised oversight of AI systems can enhance credibility with customers, regulators and business partners. From a policy perspective, this research demonstrates the effective scaling-down of the NIST AI RMF that can benefit the majority of businesses worldwide, and helps reinforce AI safety as a common national and international priority.

Nevertheless, the study also shows the need for further exploration. Future work should broaden this evaluation work to multilingual and multimodal AI contexts, where issues of governance may be magnified by cultural, linguistic and technical complexity. Additionally, sector-specific adaptations of the toolkit—that is, for instance in manufacturing, education or public administration—would contribute to the fine-tuning of generalisability and appropriateness.

In conclusion, this research offers a contribution both conceptually and practically to the area of AI governance. It provides a pathway to operationalize AI risk management in a cost-effective and scalable way, ensuring SMEs can safely realize the opportunities of LLMs while meeting wider societal expectations in terms of accountability, transparency and trustworthiness.

References

1. National Institute of Standards and Technology. Artificial Intelligence Risk Management Framework (AI RMF 1.0). NIST; 2023.
2. Jobin A, Ienca M, Vayena E. The global landscape of AI ethics guidelines. *Nat Mach Intell.* 2019;1:389–399.
3. Floridi L, Cowls J. A unified framework of five principles for AI in society. *Harv Data Sci Rev.* 2019;1(1).
4. OECD. The SME Digitalisation Gap. OECD; 2021.
5. Raji ID, Smart A, White RN, Mitchell M, Gebru T, Hutchinson B, et al. Closing the AI accountability gap: Defining an end-to-end framework for internal algorithmic auditing. *FAccT.* 2020:33–44.
6. Weidinger L, Mellor J, Rauh M, Griffin C, Uesato J, Huang S, et al. Taxonomy of risks posed by language models. *arXiv preprint arXiv:2112.04359.* 2021.
7. Brundage M, Avin S, Wang J, Belfield H, Krueger G, Hadfield G, et al. Toward trustworthy AI development: mechanisms for supporting verifiable claims. *arXiv preprint arXiv:2004.07213.* 2020.
8. Hendrycks D, Carlini N, Schulman J, Steinhardt J. Unsolved problems in ML safety. *arXiv preprint arXiv:2109.13916.* 2021.
9. Leslie D. Understanding artificial intelligence ethics and safety: A guide for the responsible design and implementation of AI systems in the public sector. The Alan Turing Institute; 2019.