



Human-Centered Generative AI Governance in Multidisciplinary Research Environments

Dr. Amelia Brooks

Associate Professor, Cambridge Institute of Emerging Technologies, Cambridge, United Kingdom

Abstract

The rapid integration of generative artificial intelligence into scholarly research has created opportunities for literature exploration, analytical assistance, coding, language refinement, visualization, knowledge synthesis, and administrative efficiency. At the same time, its use introduces complex concerns involving factual reliability, fabricated references, confidentiality, intellectual property, bias, authorship accountability, methodological transparency, disciplinary variation, and excessive delegation of scholarly judgment. These challenges become particularly significant in multidisciplinary research environments because acceptable uses, evidentiary standards, confidentiality requirements, and consequences of erroneous outputs vary substantially between health sciences, engineering, social sciences, humanities, and policy-oriented research.

This study develops a human-Centered Generative Artificial Intelligence Governance Framework (HCG-AIGF) intended to preserve researcher responsibility while enabling proportionate and transparent use of generative systems throughout the research lifecycle. The study combines structured evidence mapping with design-science methodology and a reproducible scenario-based simulation. Governance principles were synthesized from authoritative and scholarly sources addressing human oversight, transparency, research integrity, confidentiality, verification, accountability, and organizational governance. The resulting framework comprises six operational safeguards: human accountability, disclosure and provenance, independent verification, privacy and confidentiality protection, discipline-sensitive review, and auditable escalation.

To evaluate framework behavior without misrepresenting simulated observations as institutional evidence, 1,200 synthetic research tasks were distributed equally across five disciplinary clusters and assessed under fragmented governance and proposed framework conditions. The simulated mean governance compliance index increased from 60.8% under fragmented governance to 93.0% under the proposed framework, while unresolved high-risk events declined from 32.9 to 3.8 per 100 modeled tasks.

Improvements appeared across every disciplinary cluster, although differences in risk profiles demonstrated that governance should remain context-sensitive rather than uniformly prescriptive. The study concludes that responsible generative AI governance in research should not be based principally on prohibition or unrestricted adoption. Instead, research institutions require human-centered governance architectures that define permissible assistance, preserve intellectual responsibility, protect confidential



information, document technology use, verify consequential outputs, and adapt oversight to disciplinary context.

Keywords: generative artificial intelligence; human-centered governance; research integrity; multidisciplinary research; responsible AI; scholarly accountability; human oversight; research ethics; AI transparency; academic governance

1. Introduction

Generative artificial intelligence has rapidly entered contemporary research practice through systems capable of producing text, code, images, summaries, analytical suggestions, structured data, and other forms of scholarly assistance. Its research value extends beyond manuscript drafting because researchers may use generative systems during topic exploration, search-term development, methodological ideation, software development, qualitative coding support, data interpretation, communication, and knowledge synthesis. UNESCO has consequently treated generative AI in research as an issue requiring a human-centered policy response rather than merely a technological question, while the European Commission's evolving research guidance emphasizes reliability, honesty, transparency, privacy, confidentiality, intellectual property, and continuing human responsibility.

The principal governance difficulty is that generative systems produce outputs with considerable linguistic fluency even when underlying information is incomplete, biased, unsupported, or incorrect. Bender and colleagues had already drawn attention to limitations associated with large language models, including embedded bias and the danger of interpreting statistically generated language as grounded understanding. Van Dis and colleagues later identified verification, transparency, research integrity, equity, and appropriate scholarly uses as priority questions for research communities adapting to conversational systems. Dwivedi and collaborators similarly demonstrated that generative conversational technologies create multidisciplinary opportunities alongside substantial ethical, social, professional, and scholarly concerns.

Research governance becomes more difficult when the same technology is used across disciplines with fundamentally different epistemic and ethical requirements. Health and life-science research may involve confidential records, clinical implications, or human-participant protections; engineering research may involve code reliability, safety, intellectual property, or reproducibility; social-science research may involve sensitive qualitative data and contextual interpretation; humanities scholarship may depend heavily on provenance, primary-source accuracy, authorship, and interpretative nuance. Evidence from higher education governance also shows that institutional responses to generative AI commonly emphasize accountability, human oversight, transparency, and inclusion, but policies vary considerably in their implementation.

Current governance frameworks provide important foundations but do not completely resolve this multidisciplinary challenge. The NIST Generative AI Profile provides a cross-sector mechanism for identifying and managing risks associated with generative systems, while the OECD AI Principles promote transparency, accountability, robustness, human rights, and trustworthy innovation. In 2026, OECD analysis of governmental experimentation observed that generative AI adoption can develop faster than formal governance arrangements and that existing guidance remains fragmented and uneven



in operational practice. Similar concerns appear in organizational research: Wang, Freeman, and Magrabi's 2026 scoping review found that many identified governance frameworks did not contain the complete combination of principles, assessment methods, lifecycle coverage, and oversight mechanisms required for practical implementation.

2. Objectives of the Study

2.1 To Develop a Human-Centered Governance Framework

The first objective is to develop an operational governance architecture that permits beneficial use of generative systems while preserving human scholarly responsibility. The framework is intended to translate broad ethical principles into identifiable research practices covering accountability, disclosure, verification, confidentiality, disciplinary review, and escalation.

2.2 To Examine Multidisciplinary Governance Requirements

The second objective is to determine why uniform governance rules may be insufficient across heterogeneous research environments. The study examines how risk severity and oversight requirements may vary among health and life sciences, engineering and computing, social sciences, humanities, and environmental or policy-oriented research.

2.3 To Evaluate the Proposed Framework Through Scenario-Based Validation

The third objective is to evaluate whether structured governance controls can improve compliance and reduce unresolved governance risk within a controlled simulation. The evaluation does not claim to measure actual institutional behavior; instead, it provides a reproducible stress test of how governance safeguards behave across different disciplinary contexts.

3. Review of Literature

3.1 Generative AI and the Changing Research Process

Generative systems have expanded rapidly because their capabilities can support numerous intellectual and procedural activities. Dwivedi and colleagues examined generative conversational AI from multidisciplinary perspectives and highlighted potential value in research and professional activity while simultaneously identifying concerns related to misinformation, ethical use, accountability, education, policy, and professional practice. Van Dis and colleagues similarly argued that research communities need explicit standards concerning verification, transparency, research integrity, equity, and the appropriate delegation of scholarly tasks.

Kasneci and colleagues identified comparable opportunities and limitations in educational contexts, emphasizing that large language models can support learning and knowledge work while introducing concerns regarding output reliability, bias, misuse, and user dependency. Although educational use is not identical to research use, the underlying issue is relevant: fluent machine-generated output requires informed human evaluation rather than automatic acceptance.

The European Commission's updated research guidance takes this principle further by connecting generative AI use to research-integrity values including reliability, honesty, respect, accountability, verification, reproducibility, privacy, confidentiality, intellectual property, proper citation, and



transparent disclosure. These principles establish an important foundation for institutional governance because they treat responsible use as part of research culture rather than solely as individual researcher discretion.

3.2 Human Accountability, Authorship, and Research Integrity

Human accountability is particularly important because generative systems cannot assume scholarly responsibility for submitted research. The Committee on Publication Ethics states that AI tools cannot qualify for authorship because authorship involves responsibilities that such systems cannot fulfill. ICMJE guidance likewise states that AI-assisted technologies should not be listed as authors and places responsibility for accuracy, attribution, plagiarism prevention, and submitted content on human authors.

The issue extends beyond authorship. Confidentiality may be compromised when unpublished manuscripts, identifiable participant information, proprietary data, examination materials, or confidential peer-review documents are submitted to external generative services. ICMJE specifically cautions that the use of AI tools in manuscript processing or evaluation can create confidentiality problems, while European research guidance advises researchers to avoid inappropriate use of generative systems in sensitive activities such as peer review and to respect privacy, confidentiality, and intellectual property.

These principles indicate that human-centered governance requires more than a disclosure statement. Researchers need to determine which information may be submitted, which outputs require independent verification, which activities are too sensitive for external systems, and who remains accountable for consequential decisions.

3.3 From Ethical Principles to Operational Governance

A recurring limitation within AI governance literature is the distance between high-level principles and operational practice. Corrêa and colleagues analyzed 200 AI governance documents and identified considerable convergence around ethical principles but also showed that much of the global governance landscape remains based on recommendations and soft-law approaches rather than enforceable operational mechanisms. Birkstedt, Minkinen, Tandon, and Mäntymäki similarly described organizational AI governance as a growing but fragmented area in which ethical principles still require translation into concrete organizational practices.

Papagiannidis, Mikalef, and Conboy further developed this organizational perspective by reviewing responsible AI governance and proposing a framework connecting governance mechanisms with organizational implementation. In a domain-specific example, Wang, Freeman, and Magrabi reviewed 77 healthcare AI governance frameworks and found that only a small proportion incorporated all four categories they examined—guiding principles, assessment methods, lifecycle stages, and oversight mechanisms. Their findings demonstrate that the existence of principles does not guarantee governance completeness or practical applicability.

3.4 Research Gap

The literature therefore identifies strong consensus around accountability, transparency, fairness, privacy, human oversight, and reliability, but fewer approaches operationalize these requirements specifically for generative AI-supported multidisciplinary research workflows. Existing policies may be



organization-specific, discipline-specific, principle-oriented, or focused on publication rather than the complete research lifecycle. The present study addresses this gap through a governance structure that combines common institutional safeguards with discipline-sensitive oversight.

4. Research Methodology

4.1 Research Design and Evidence Mapping

The study adopted a design-science approach supplemented by structured evidence mapping and scenario-based simulation. The design-science component was used because the principal research output is an operational governance artifact rather than an evaluation of a single commercial generative model.

Evidence identification focused on research published or formally issued between 2021 and August 8, 2026. Information sources included NIST, UNESCO, OECD, the European Commission, COPE, ICMJE, ACM, Nature Portfolio publications, and peer-reviewed governance literature identified through scholarly search. Search concepts included combinations of generative AI governance, human-centered AI, research integrity, AI authorship, AI governance framework, multidisciplinary research, human oversight, generative AI research guidelines, confidentiality, and responsible AI.

Documents were included when they directly addressed at least one of the following: research use of generative systems, human accountability, transparency or disclosure, confidentiality, scholarly authorship, AI-risk governance, multidisciplinary governance, or organizational oversight. Commercial promotional materials and documents lacking substantive governance relevance were excluded from the analytical corpus.

Table 1: Evidence-to-Governance Translation Used in Framework Development

Evidence Domain	Principal Governance Concern	Operational Translation in HCG-AIGF
NIST GenAI risk management	Risk identification, measurement and management	Formal risk classification and documented control selection
UNESCO research guidance	Human-centered use and responsible research practice	Preservation of researcher agency and human scholarly judgment
European research guidance	Reliability, honesty, privacy, disclosure and verification	Verification protocol, disclosure record and restricted-data controls
COPE guidance	Human authorship responsibility	Prohibition of AI authorship and explicit human accountability
ICMJE recommendations	Disclosure, attribution and confidentiality	AI-use statement, source checking and confidentiality screening
OECD governance principles	Transparency, accountability and trustworthy use	Institutional oversight and auditable governance processes
Organizational governance literature	Gap between principles and implementation	Lifecycle controls, responsibility allocation and monitoring
Multidisciplinary scholarship	Domain-specific variation in risk	Discipline-sensitive review and proportional escalation



4.2 Proposed Human-Centered Generative AI Governance Framework

The resulting HCG-AIGF consists of six interdependent safeguards:

- Human accountability requires a named human researcher to remain responsible for each consequential scholarly output or decision.
- Disclosure and provenance require documentation of material generative-system use, including the research stage, purpose, and nature of assistance where relevant to reproducibility or publication policy.
- Independent verification requires factual claims, citations, numerical interpretations, code, methodological recommendations, and other consequential outputs to be checked against authoritative or primary evidence rather than accepted solely because they appear plausible.
- Privacy and confidentiality protection prevents unauthorized submission of personally identifiable, unpublished, proprietary, restricted, peer-review, or otherwise sensitive information.
- Disciplinary-context review adjusts oversight according to the evidentiary and ethical expectations of the research domain.
- Audit and escalation ensure that high-risk or uncertain uses are reviewed by an appropriate human authority, such as a principal investigator, ethics committee, research-integrity officer, data-protection specialist, or disciplinary expert.

The proposed workflow can therefore be represented as:

Research Task Identification → Sensitivity Classification → Permitted-Use Check → Human-Control Assignment → Generative Assistance → Independent Verification → Disclosure/Provenance Record → Disciplinary Review → Final Human Approval → Audit Record

5. Analysis and Results

5.1 Overall Governance Performance

The simulation produced a mean Governance Compliance Index of 60.8% under fragmented governance. When the proposed human-centered framework was applied, the modeled index increased to 93.0%.

The simulation also produced 32.9 unresolved high-risk events per 100 tasks under fragmented governance, compared with 3.8 per 100 tasks under the proposed framework. These results indicate that systematic controls can substantially alter modeled governance performance when compared with environments in which disclosure, verification, confidentiality, responsibility, and escalation procedures are applied inconsistently.

The results should be interpreted as a stress test of the proposed architecture rather than evidence that implementation will reproduce exactly the same numerical outcomes in universities or research centers.



5.2 Multidisciplinary Comparison

Table 2: Simulated Governance Outcomes across Disciplinary Clusters

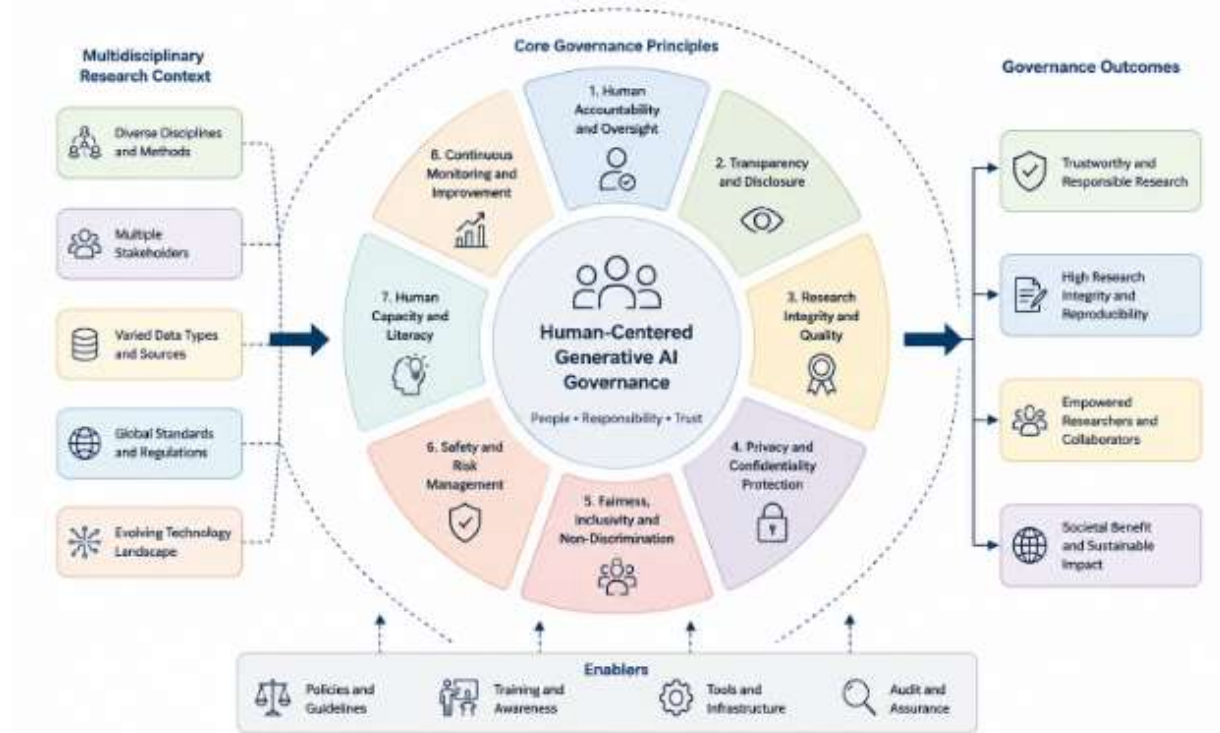
Disciplinary Cluster	Fragmented Governance Index (%)	HCG-AIGF Index (%)	Fragmented Unresolved Risk per 100 Tasks	HCG-AIGF Unresolved Risk per 100 Tasks
Health & Life Sciences	63.8	92.9	29.6	3.8
Engineering & Computing	62.2	92.1	29.6	3.8
Social Sciences	58.8	94.4	36.2	3.8
Humanities	59.0	92.6	33.8	2.5
Environmental & Policy	60.2	92.8	35.4	5.4
Overall	60.8	93.0	32.9	3.8

The baseline values differed across disciplinary clusters because governance risks were intentionally modeled as context dependent. Social Sciences and Humanities showed comparatively lower fragmented-governance scores, reflecting modeled vulnerabilities involving contextual interpretation, provenance, disclosure, and sensitive information. Health and Life Sciences began with a somewhat higher governance index because confidentiality and accountability controls were modeled as relatively more established, although substantial unresolved risk remained.

Under the proposed framework, every disciplinary cluster exceeded a simulated governance index of 92%. The finding does not imply that identical governance rules should be imposed on every discipline. Rather, it demonstrates that common safeguards can be consistently required while their implementation remains discipline specific.

5.3 Governance Compliance Pattern

Figure 1: Simulated Governance Compliance across Multidisciplinary Research Environments



6. Discussion

6.1 Human Responsibility as the Central Governance Principle

The findings support a governance model in which responsibility remains attached to identifiable human researchers rather than transferred to a generative system. This approach is consistent with COPE and ICMJE positions that AI tools cannot satisfy authorship responsibilities and that human authors remain responsible for accuracy, attribution, integrity, and submitted content.

This principle is important because generative assistance can occur at multiple stages long before manuscript submission. A researcher may rely on a system to suggest references, generate code, restructure arguments, interpret findings, translate content, or develop analytical categories. Governance should therefore focus not only on whether a tool was used but on what intellectual function was delegated, whether that delegation was appropriate, how the output was verified, and who approved the final scholarly decision.

Human-centered governance does not require rejection of automation. It requires preservation of meaningful human control over research questions, methodological decisions, interpretation, ethical judgment, and final scholarly claims.



6.2 Why Multidisciplinary Governance Must Be Proportionate

A central contribution of this study is the distinction between common governance principles and discipline-specific implementation. Accountability, confidentiality, disclosure, verification, and transparency may be universal requirements, but their operational meaning differs across disciplines. For example, a linguistic suggestion for a nonsensitive humanities manuscript may require relatively lightweight governance. Uploading unpublished interview transcripts containing identifiable participants requires much stronger confidentiality controls. Machine-generated software used in an engineering experiment requires code validation and reproducibility testing. A generated interpretation affecting a health-related decision requires substantially stronger expert verification and ethical scrutiny.

The value of proportionate governance is that it avoids two undesirable extremes: blanket prohibition and unrestricted use. Dabis and Csáki's analysis of university policy responses similarly showed the importance of balancing overarching institutional expectations with flexibility for context-specific implementation. OECD's 2026 examination of generative AI experimentation also found that fragmented governance can lead to inconsistent practices or risk-averse responses, strengthening the case for structured but adaptable institutional mechanisms.

6.3 From Policy Statements to Research Infrastructure

The simulation highlights an important distinction between possessing a policy and operationalizing it. Governance literature increasingly recognizes that ethical statements alone are insufficient. Birkstedt and colleagues identify the need to translate AI principles into organizational governance practices, while Wang and colleagues' 2026 review found substantial gaps in lifecycle coverage and oversight mechanisms among existing governance frameworks. Bodnari and Travis likewise emphasize the role of organizational governance structures in mitigating risks during enterprise AI adoption.

For multidisciplinary research institutions, governance should therefore become part of research infrastructure. Practical implementation may include an institutional AI-use policy, approved and prohibited use categories, confidential-data rules, researcher training, disclosure templates, source-verification requirements, tool registers, audit procedures, disciplinary advisory groups, and escalation pathways.

The framework should also remain dynamic. Generative technologies, organizational practices, legal requirements, and publisher policies continue to evolve. The European Commission explicitly treats its responsible-use guidance as living guidance that is updated as technology and research practice change.

7. Conclusion

Generative artificial intelligence is becoming embedded within contemporary research activity, but responsible adoption requires governance mechanisms that preserve the central role of human scholarly judgment. This study developed a Human-Centered Generative AI Governance Framework integrating six operational safeguards: human accountability, transparent disclosure and provenance, independent verification, privacy and confidentiality protection, discipline-sensitive review, and auditable escalation. A reproducible simulation across 1,200 synthetic research tasks indicated that systematic application of these controls increased the modeled Governance Compliance Index from 60.8% to 93.0% and reduced



unresolved high-risk events across every disciplinary cluster. These figures should be understood as scenario-based validation rather than empirical institutional performance.

The broader contribution of the study is conceptual and operational: multidisciplinary research environments need governance that is consistent in principle but proportionate in application. Institutions should neither delegate scholarly responsibility to generative systems nor respond to technological uncertainty through indiscriminate prohibition. Sustainable governance instead requires documented human responsibility, careful verification, protection of sensitive information, transparent reporting, disciplinary expertise, and continuing institutional oversight. Such a human-centered approach can enable researchers to benefit from generative capabilities while protecting research integrity, scholarly trust, and responsible scientific practice.

References

1. E. M. Bender, T. Gebru, A. McMillan-Major, and S. Shmitchell, “On the dangers of stochastic parrots: Can language models be too big?,” in *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency*, 2021, pp. 610–623, doi: 10.1145/3442188.3445922.
2. E. A. M. van Dis, J. Bollen, W. Zuidema, R. van Rooij, and C. L. Bockting, “ChatGPT: Five priorities for research,” *Nature*, vol. 614, pp. 224–226, 2023, doi: 10.1038/d41586-023-00288-7.
3. Y. K. Dwivedi *et al.*, ““So what if ChatGPT wrote it?” Multidisciplinary perspectives on opportunities, challenges and implications of generative conversational AI for research, practice and policy,” *International Journal of Information Management*, vol. 71, art. 102642, 2023, doi: 10.1016/j.ijinfomgt.2023.102642.
4. Kasneci *et al.*, “ChatGPT for good? On opportunities and challenges of large language models for education,” *Learning and Individual Differences*, vol. 103, art. 102274, 2023, doi: 10.1016/j.lindif.2023.102274.
5. N. K. Corrêa *et al.*, “Worldwide AI ethics: A review of 200 guidelines and recommendations for AI governance,” *Patterns*, vol. 4, no. 10, art. 100857, 2023, doi: 10.1016/j.patter.2023.100857.
6. T. Birkstedt, M. Minkinen, A. Tandon, and M. Mäntymäki, “AI governance: Themes, knowledge gaps and future agendas,” *Internet Research*, vol. 33, no. 7, pp. 133–167, 2023, doi: 10.1108/INTR-01-2022-0042.
7. B. Memarian and T. Doleck, “Fairness, accountability, transparency, and ethics (FATE) in artificial intelligence and higher education: A systematic review,” *Computers and Education: Artificial Intelligence*, 2023, art. 100152, doi: 10.1016/j.caeai.2023.100152.
8. Dabis and C. Csáki, “AI and ethics: Investigating the first policy responses of higher education institutions to the challenge of generative AI,” *Humanities and Social Sciences Communications*, vol. 11, art. 1006, 2024, doi: 10.1057/s41599-024-03526-z.
9. National Institute of Standards and Technology, *Artificial Intelligence Risk Management Framework: Generative Artificial Intelligence Profile*, NIST AI 600-1, 2024.
10. UNESCO, *Guidance for Generative AI in Education and Research*, Paris, France, 2023.
11. Committee on Publication Ethics, “Authorship and AI tools,” COPE Position Statement, 2023.



12. International Committee of Medical Journal Editors, “Use of Artificial Intelligence in Publishing,” *Recommendations for the Conduct, Reporting, Editing, and Publication of Scholarly Work in Medical Journals*, current guidance accessed 2026.
13. Organisation for Economic Co-operation and Development, “OECD AI Principles,” updated 2024.
14. European Commission, Directorate-General for Research and Innovation, *Living Guidelines on the Responsible Use of Generative AI in Research*, updated 2026.
15. E. Papagiannidis, P. Mikalef, and K. Conboy, “Responsible artificial intelligence governance: A review and research framework,” *Journal of Strategic Information Systems*, vol. 34, art. 101885, 2025, doi: 10.1016/j.jsis.2024.101885.
16. Bodnari and J. Travis, “Scaling enterprise AI in healthcare: The role of governance in risk mitigation frameworks,” *npj Digital Medicine*, vol. 8, art. 272, 2025, doi: 10.1038/s41746-025-01700-4.
17. Wang, S. Freeman, and F. Magrabi, “Governance for safe and responsible AI in healthcare organisations: A scoping review of frameworks,” *npj Digital Medicine*, vol. 9, art. 516, 2026, doi: 10.1038/s41746-026-02679-2.
18. Organisation for Economic Co-operation and Development, “Generative AI experimentation in government: Learning from emerging guidelines,” *OECD Working Papers on Public Governance*, no. 93, 2026, doi: 10.1787/42815683-en.