



THE SYNERGY HALL

RESEARCH SERIES

AN INITIATIVE OF THE SHEYENE GERARDI FOUNDATION
STRATEGIC SECURITY · SCIENCE · SOCIETY

No. 3

August 2026

ISSN [pending]

The C2S Reliability Framework

Causal Artificial Intelligence and Institutional Decision Support in High-Consequence Domains

By Sheyene Gerardi

PUBLISHED BY THE SHEYENE GERARDI FOUNDATION

Research, Systems, and Institutional Studies

Washington, D.C.

ABOUT THE SYNERGY HALL RESEARCH SERIES

The Synergy Hall Research Series is the scholarly monograph series of the Sheyene Gerardi Foundation. It advances interdisciplinary research in institutional design, strategic security, artificial intelligence, systems science, governance, and collective human capability.

Each volume presents original analysis intended to connect scientific knowledge, technological development, and institutional design. The series examines how independent systems may be intentionally coordinated to produce outcomes greater than the sum of their individual contributions.

Each monograph is issued as a separately numbered volume within an ongoing research program intended to develop durable institutional models in service of human flourishing.

PUBLICATION DETAILS

Series title: Synergy Hall Research Series

Series type: Scholarly Research Monographs

Publisher: Sheyene Gerardi Foundation

Research Division: Research, Systems, and Institutional Studies

Place of publication: Washington, D.C., United States

Editor: Sheyene Gerardi

Frequency: Irregular

Language: English

ISSN: *[pending — U.S. ISSN Center, Library of Congress]*

Inaugural Volume: Monograph No. 3, August 2026

Contact: contact@sheyene.org · www.sheyene.org

COPYRIGHT

© 2026 Sheyene Gerardi Foundation. All rights reserved. The Synergy Hall Research Series is an independent scholarly publication of the Sheyene Gerardi Foundation. The views expressed in each monograph are solely those of the author and do not necessarily reflect the views of any government, international organization, academic institution, or private entity.

No part of this publication may be reproduced, stored in a retrieval system, or transmitted in any form or by any means without prior written permission from the publisher, except for brief quotations used for scholarly, educational, or review purposes with appropriate attribution.

SUGGESTED CITATION

Gerardi, S. (2026). *The C2S Reliability Framework: Causal Artificial Intelligence and Institutional Decision Support in High-Consequence Domains*. **Synergy Hall Research Series**, Monograph No. 3. Washington, D.C.: Sheyene Gerardi Foundation.

SYNERGY HALL RESEARCH SERIES

Intelligence acquires its highest value when it improves the institutions that serve humanity.



The C2S Reliability Framework

Causal Artificial Intelligence and Institutional Decision Support in High-Consequence Domains

By Sheyene Gerardi

Sheyene Gerardi Foundation · Washington, D.C. · August 2026

Artificial intelligence is becoming an essential component of institutional decision-making across security, healthcare, finance, and public administration. As these systems assume increasingly consequential advisory roles, the central challenge is no longer computational capability alone, but institutional trust.

Executive Summary

Artificial intelligence is becoming an essential component of institutional decision-making across security, healthcare, finance, and public administration. As these systems assume increasingly consequential advisory roles, the central challenge is no longer computational capability alone, but institutional trust. Decisions affecting human lives, national security, and strategic infrastructure require analytical systems that are not only accurate, but also transparent, explainable, and scientifically accountable.

This monograph introduces the **C2S Reliability Framework**, a conceptual model proposing that causal artificial intelligence can strengthen personnel reliability in high-consequence domains by supporting—rather than replacing—human institutional judgment. Building upon the principles established in *SynsID* and *The Survival Blueprint*, the framework argues that artificial intelligence should function as an evidence-integration system capable of organizing complex behavioral, developmental, physiological, and operational information into transparent decision-support models.

Particular attention is given to the emergence of foundation models capable of interpreting biological information, including the **Cell2Sentence (C2S-Scale)** family developed by Yale University, Google Research, and Google DeepMind. These technologies demonstrate the growing capacity of large language models to analyze complex biological systems. Their significance within this monograph, however, lies not in direct behavioral prediction but in illustrating how future multimodal AI systems may contribute to broader institutional evaluation frameworks.

The C2S Reliability Framework therefore rejects both biological determinism and fully autonomous decision-making. Instead, it proposes a human-centered model in which explainable

causal AI assists qualified professionals by integrating multiple independent sources of objective evidence while preserving accountability, transparency, ethical safeguards, and due process.

This framework is presented as a conceptual contribution intended to stimulate interdisciplinary discussion regarding the future relationship between artificial intelligence, institutional governance, and human reliability in environments where individual decisions may carry irreversible consequences.

Part I. The Scientific Foundation

Artificial intelligence cannot improve institutional judgment unless the science upon which it relies is itself sound. Before AI can assist in evaluating human reliability, it must operate upon an accurate understanding of how human behavior emerges.

Contemporary behavioral genetics, neuroscience, and psychology demonstrate that character is neither random nor genetically predetermined. Instead, human behavior develops through continuous interaction among biology, development, environment, education, experience, and individual agency. Artificial intelligence must therefore integrate these dimensions rather than reduce them to isolated variables.

The following sections establish the scientific foundation upon which the C2S Reliability Framework is built.

1.1 Human Character as a Complex Adaptive System

Modern behavioral science increasingly understands personality as a complex adaptive system rather than the product of isolated biological factors. Twin studies consistently demonstrate that character is moderately heritable across populations, while molecular genetics reveals that this heritability emerges from highly distributed biological networks involving hundreds of interacting genetic variants rather than individual genes.¹

These discoveries resolve an important scientific misconception. Human behavior is not controlled by "genes for aggression," "genes for honesty," or any other isolated behavioral determinant. Instead, biological predispositions contribute probabilistically to the development of neural systems whose expression remains continuously influenced by education, culture, relationships, health, and lived experience.

Consequently, the objective of institutional evaluation cannot be the prediction of behavior from inherited biology. Rather, biological science contributes to understanding the origins of human behavior while observable performance remains the most reliable indicator of operational reliability.

This distinction provides the conceptual bridge between contemporary behavioral science and the institutional principles developed throughout the previous monographs.

1.2 From Biological Information to Institutional Knowledge

Scientific knowledge becomes institutionally valuable only when it improves decision-making without exceeding the limits of available evidence.

Recent advances in computational biology—including large-scale analyses of polygenic networks and foundation models capable of interpreting complex biological data—have substantially expanded humanity's ability to understand living systems. These technologies demonstrate extraordinary analytical capability while simultaneously illustrating an equally important principle: increasing computational power does not eliminate scientific uncertainty.

Within the C2S Reliability Framework, biological information therefore serves as one component of a broader evidence ecosystem. It contributes context rather than conclusions, informs judgment rather than replacing it, and strengthens institutional understanding without becoming an independent basis for certification decisions.

This principle preserves both scientific integrity and ethical legitimacy. Institutions benefit from advances in biological science while avoiding the errors of biological determinism, technological overconfidence, and reductionist approaches to human behavior.

1.3 Why Artificial Intelligence Requires Causality

The increasing adoption of artificial intelligence in institutional decision-making has exposed a fundamental limitation of conventional predictive models. Statistical prediction alone cannot adequately support decisions whose consequences involve national security, human rights, or irreversible public harm.

Most contemporary machine-learning systems identify patterns through correlation. Although often highly accurate, these models rarely explain why a prediction has been generated or which underlying mechanisms produced the observed result. In high-consequence domains, prediction without explanation is insufficient for institutional accountability.

The C2S Reliability Framework therefore distinguishes between predictive intelligence and causal intelligence. Predictive systems estimate the probability that an event may occur. Causal systems seek to explain the relationships that generate those outcomes, evaluate alternative scenarios, and support evidence-based intervention.

This distinction fundamentally changes the institutional role of artificial intelligence. Rather than functioning as an automated evaluator of human beings, AI becomes an analytical instrument that assists professionals in understanding complex interactions among biological, behavioral, developmental, environmental, and operational evidence.

The remainder of this monograph explores how causal AI may strengthen institutional judgment while preserving transparency, accountability, and meaningful human oversight.

Part II. Causal Artificial Intelligence

Artificial intelligence is often evaluated according to predictive accuracy. For institutional decision-making, however, accuracy alone is insufficient. Decisions affecting public safety, national security, or individual rights require systems capable of explaining their conclusions, identifying uncertainty, and supporting accountable human judgment.

Causal artificial intelligence represents an important evolution beyond conventional predictive analytics. Rather than identifying statistical associations alone, causal models seek to represent the mechanisms through which outcomes emerge, allowing institutions to evaluate evidence more transparently and to explore potential interventions before irreversible consequences occur.

The following sections examine how causal reasoning may contribute to institutional reliability while remaining subject to scientific, ethical, and legal constraints.

2.1 Beyond Prediction

Predictive artificial intelligence has transformed numerous fields by identifying patterns within large datasets. Financial fraud detection, medical imaging, language translation, and many other applications demonstrate the remarkable capabilities of modern machine learning. Yet prediction alone provides only part of the information required for responsible institutional decision-making.

High-consequence environments demand an additional level of understanding. Institutions must know not only that a particular pattern has been detected, but also how the conclusion was reached, which variables contributed most significantly, what uncertainties remain, and whether alternative interventions might change the projected outcome.

Causal artificial intelligence addresses these requirements by explicitly modeling relationships among variables rather than relying solely upon statistical association. Through causal graphs, structural causal models, and counterfactual reasoning, AI can assist decision-makers in distinguishing probable causes from coincidental correlations.

Within the C2S Reliability Framework, this capability transforms artificial intelligence from a predictive technology into an institutional decision-support system whose primary purpose is not autonomous judgment but improved human understanding.

2.2 Explainability as an Institutional Requirement

Institutional decisions must be understandable to the individuals affected by them as well as to those responsible for reviewing, auditing, and governing the decision-making process. Explainability is therefore not simply a desirable technical characteristic of artificial intelligence; it is a prerequisite for institutional legitimacy.

Black-box models may achieve impressive predictive performance while providing little insight into how their conclusions are generated. Such opacity may be acceptable in applications where incorrect predictions carry limited consequences. It becomes unacceptable, however, when AI contributes to decisions involving national security, personnel certification, or other responsibilities with potentially irreversible outcomes.

The C2S Reliability Framework therefore requires that AI-assisted evaluations remain transparent, traceable, and scientifically interpretable. Human reviewers must be capable of understanding the evidence integrated by the system, identifying the principal factors influencing each recommendation, and challenging conclusions whenever necessary.

Artificial intelligence should therefore function as an analytical partner rather than an unquestionable authority. Explainability transforms computational capability into institutional trust by ensuring that every recommendation remains open to scientific scrutiny, ethical evaluation, and human accountability.

2.3 Cell2Sentence as a Case Study

Recent advances in biological foundation models illustrate how rapidly artificial intelligence is expanding its capacity to interpret highly complex scientific information. Among the most significant examples is the **Cell2Sentence (C2S-Scale)** family of models developed through collaboration among Yale University, Google Research, and Google DeepMind.

These models transform single-cell gene expression data into representations that large language models can analyze, allowing the integration of complex biological information for tasks such as cell-type identification, perturbation prediction, and biomedical hypothesis generation. Their demonstrated success within cellular biology represents an important milestone in the application of foundation models to life sciences.

Within the C2S Reliability Framework, however, Cell2Sentence is not presented as a behavioral assessment system. Instead, it serves as an illustrative example of a broader technological transition: the emergence of artificial intelligence capable of integrating biological information at unprecedented scale.

This distinction is essential. The framework does not propose that existing biological language models can determine human reliability. Rather, it argues that advances represented by technologies such as Cell2Sentence provide insight into how future multimodal AI systems may contribute to institutional decision support when integrated responsibly with behavioral science, neuroscience, psychology, and objective longitudinal evidence.

Consequently, Cell2Sentence functions within this monograph as a scientific case study rather than as the foundation of the proposed certification framework itself.

2.4 The Limits of Artificial Intelligence

Artificial intelligence continues to advance at an extraordinary pace. Nevertheless, every AI system remains constrained by the quality of its underlying data, the assumptions embedded within its models, and the scientific validity of the evidence upon which its conclusions depend.

These limitations become especially important in the evaluation of human behavior. Behavioral reliability cannot be reduced to isolated biological measurements, historical records, or computational predictions. Human conduct remains dynamic, context-dependent, and

continuously influenced by new experiences, changing environments, education, health, and personal development.

For this reason, the C2S Reliability Framework explicitly rejects autonomous AI certification. Artificial intelligence may organize evidence, identify relationships, estimate uncertainty, and assist professional evaluation, but institutional responsibility must remain with accountable human decision-makers operating under transparent legal and ethical standards.

The objective is therefore not to replace human judgment with machines, but to improve the quality, consistency, and scientific foundation of institutional judgment through responsibly governed artificial intelligence.

Part III. The C2S Reliability Framework

Artificial intelligence reaches its greatest institutional value not when it replaces human judgment, but when it strengthens it.

The C2S Reliability Framework is proposed as a conceptual model through which causal artificial intelligence assists institutional decision-making by integrating diverse sources of scientifically relevant evidence into transparent, explainable, and accountable evaluations. Its purpose is not to automate certification, but to improve the quality and consistency of human institutional judgment.

The framework builds upon two principles established in the previous monographs. First, reliable certification must evaluate the expressed individual rather than inherited biological predisposition. Second, institutional reliability increases as independent sources of evidence become intentionally integrated rather than evaluated in isolation.

The following sections describe how these principles may be operationalized through a human-centered model of AI-assisted institutional decision support.

3.1 Principles of AI-Assisted Institutional Judgment

The C2S Reliability Framework is organized around a simple principle: artificial intelligence should assist institutions in integrating evidence, not in replacing institutional responsibility.

Every recommendation generated by an AI system should remain understandable, independently reviewable, and supported by objective evidence. Human decision-makers retain full responsibility for evaluating recommendations, considering contextual information unavailable to computational models, and ensuring that every decision remains ethically and legally defensible.

This principle distinguishes institutional decision support from automated decision-making. Artificial intelligence contributes computational capacity, consistency, and analytical depth, while human institutions provide judgment, accountability, proportionality, and moral responsibility.

Institutional trust therefore emerges not from artificial intelligence alone, but from the deliberate integration of computational capability and accountable human governance.

3.2 The Evidence Integration Model

Rather than evaluating isolated indicators, the C2S Reliability Framework organizes multiple independent sources of information into a unified evidence architecture.

These sources may include verified developmental history, behavioral assessment, professional performance, psychological evaluation, physiological monitoring, operational records, and—where scientifically appropriate—validated biological information. No single category functions as an independent determinant of institutional decisions.

Artificial intelligence contributes by identifying relationships among these heterogeneous forms of evidence, highlighting inconsistencies, estimating uncertainty, and presenting transparent analytical summaries for human review.

The objective is not prediction for its own sake. The objective is to improve institutional understanding by ensuring that complex evidence is evaluated consistently, comprehensively, and proportionally to the responsibilities entrusted to each individual.

3.3 Human Oversight and Institutional Responsibility

The integration of artificial intelligence into institutional decision-making does not diminish human responsibility; it increases it. As computational systems become more capable of organizing complex information, institutions assume a greater obligation to ensure that every recommendation remains scientifically justified, ethically defensible, and legally accountable.

Within the C2S Reliability Framework, artificial intelligence never functions as the final decision-maker. Human professionals remain responsible for interpreting evidence, considering contextual factors beyond computational analysis, and exercising judgment consistent with institutional values and legal standards.

This principle preserves a fundamental distinction between intelligence and authority. Artificial intelligence may contribute analytical capability, but institutional authority must remain inseparable from human accountability. Decisions affecting careers, public safety, or national security ultimately require moral responsibility that cannot be delegated to computational systems.

Accordingly, the framework envisions AI as an enduring partner in institutional reasoning rather than as an autonomous institutional actor.

Part IV. Ethical and Institutional Governance

Artificial intelligence cannot strengthen public trust unless the institutions governing its use are themselves worthy of trust. Scientific capability alone is insufficient. Legitimate institutional decision-making requires transparency, accountability, proportionality, due process, and respect for human dignity.

The C2S Reliability Framework therefore places governance at the center of AI-assisted decision support. Ethical safeguards are not secondary constraints imposed upon technological systems; they constitute part of the framework itself.

The following principles establish the institutional conditions necessary for the responsible integration of artificial intelligence into high-consequence decision-making.

4.1 From Deterrence to Certification

Institutional decisions must remain understandable to those responsible for making them, reviewing them, and living with their consequences. Artificial intelligence therefore has an obligation to produce explanations rather than conclusions alone.

Explainability serves several institutional purposes simultaneously. It enables scientific validation, supports legal review, facilitates independent auditing, strengthens public confidence, and allows individuals affected by institutional decisions to understand the evidence supporting those decisions.

Transparency also protects institutions themselves. Systems whose reasoning cannot be examined cannot be effectively corrected when errors occur. Explainable AI therefore improves not only fairness but institutional resilience by allowing continuous refinement as scientific knowledge evolves.

Within the C2S Reliability Framework, transparency is not treated as a desirable technical feature. It is a governing institutional principle.

4.2 Ethics Beyond Compliance

Ethical governance requires more than compliance with existing regulations. Laws establish minimum standards of acceptable conduct, whereas institutional ethics seeks to ensure that scientific capability remains aligned with human dignity and the broader public interest.

The C2S Reliability Framework therefore rejects approaches based upon surveillance, biological determinism, opaque algorithmic authority, or irreversible classification. Instead, it emphasizes proportional evaluation, procedural fairness, continuous review, and meaningful opportunities for human oversight and appeal.

Ethics is consequently understood as an enabling component of institutional reliability rather than as an external limitation upon technological innovation. Responsible governance strengthens public confidence while preserving the legitimacy required for long-term institutional effectiveness.

4.3 International Governance and Shared Standards

Many of the institutions responsible for managing high-consequence risks operate across national boundaries. Nuclear security, civil aviation, biosecurity, artificial intelligence, cybersecurity, and critical infrastructure increasingly depend upon cooperation among governments, scientific institutions, and international organizations.

The C2S Reliability Framework is therefore presented as a conceptual foundation for international dialogue rather than as a proposal for centralized global authority. Individual nations retain full responsibility for implementing policies consistent with their constitutional traditions, legal systems, and institutional priorities.

International cooperation nevertheless offers substantial advantages. Shared scientific standards, interoperable evaluation methodologies, independent auditing, collaborative research, and transparent governance mechanisms can strengthen institutional reliability while preserving national sovereignty.

The objective is not institutional uniformity, but institutional compatibility. Reliable cooperation emerges when independent institutions operate according to shared principles of scientific integrity, transparency, accountability, and respect for human dignity.

Conclusion

Artificial intelligence represents one of the most significant institutional technologies of the twenty-first century. Its greatest contribution, however, will not arise from replacing human judgment, but from improving the quality of institutional decision-making through the responsible integration of scientific knowledge.

The C2S Reliability Framework proposes that causal artificial intelligence should function as an institutional decision-support system rather than as an autonomous decision-maker. By integrating behavioral science, neuroscience, psychology, operational evidence, and—where scientifically appropriate—validated biological information into transparent analytical models, AI may strengthen the consistency, accountability, and scientific foundation of high-consequence personnel evaluation.

At the same time, this monograph emphasizes that technological capability must remain subordinate to institutional responsibility. Explainability, human oversight, due process, proportionality, and ethical governance are not optional safeguards added after technological development; they constitute the conditions under which artificial intelligence becomes worthy of institutional trust.

The framework presented here is therefore offered not as an operational program but as a conceptual contribution to the evolving relationship between artificial intelligence and institutional governance. Its broader objective is to encourage interdisciplinary collaboration among scientists, engineers, policymakers, legal scholars, ethicists, and public institutions seeking to develop responsible models of AI-assisted decision-making.

Ultimately, the future of artificial intelligence will be determined less by the sophistication of its algorithms than by the wisdom of the institutions that choose how those algorithms are used.

Disclaimer

This monograph presents a conceptual framework intended to stimulate interdisciplinary discussion regarding the responsible integration of artificial intelligence into institutional decision-making in high-consequence domains. It does not constitute legal, governmental, regulatory, or operational policy.

References to third-party technologies—including the Cell2Sentence (C2S-Scale) family of models developed by Yale University, Google Research, and Google DeepMind—are limited to their documented scientific capabilities within cellular biology. Any discussion of their potential future relevance to personnel reliability or institutional decision support represents the author's conceptual analysis and should not be interpreted as an endorsement by, affiliation with, or representation of those organizations.

The C2S Reliability Framework explicitly rejects biological determinism and fully autonomous AI decision-making. Throughout this monograph, artificial intelligence is presented exclusively as a decision-support technology intended to strengthen accountable human institutional judgment under appropriate scientific, ethical, legal, and procedural safeguards.

The views expressed are solely those of the author and are intended to encourage evidence-based dialogue concerning artificial intelligence, institutional governance, and human reliability.

Works Cited

1. Zwi I, Arnedo J, Del-Val C, et al. Uncovering the complex genetics of human character. *Molecular Psychiatry*. 2020;25:2295–2312. Accessed November 2, 2025. <https://www.nature.com/articles/s41380-018-0263-6>
2. Zwi I, Arnedo J, Del-Val C, et al. Uncovering the complex genetics of human temperament. *Molecular Psychiatry*. 2020;25:2275–2294. Accessed November 2, 2025. <https://www.nature.com/articles/s41380-018-0264-5>
3. Caspi A, McClay J, Moffitt TE, et al. Role of genotype in the cycle of violence in maltreated children. *Science*. 2002;297:851–854. Accessed November 2, 2025. <https://www.science.org/doi/10.1126/science.1072290>
4. Byrd AL, Manuck SB. MAOA, childhood maltreatment, and antisocial behavior: meta-analysis of a gene-environment interaction. *Biological Psychiatry*. 2014;75(1):9–17. Accessed November 2, 2025. https://moffittcaspi.trinity.duke.edu/sites/moffittcaspi.trinity.duke.edu/files/file-attachments/MAOA_2013_ByrdManuck.pdf
5. Explaining heritable variance in human character (methodological commentary on the PGMRA model). *bioRxiv preprint*. 2018. Accessed November 2, 2025. <https://www.biorxiv.org/content/10.1101/446518.full.pdf>
6. Gene Expression Networks Regulated by Human Personality. PMC, National Institutes of Health. Accessed November 2, 2025. <https://pmc.ncbi.nlm.nih.gov/articles/PMC11408262/>
7. Cell2Sentence-Scale Gemma 2 27B (C2S-Scale): model card. Yale University, Google Research, and Google DeepMind. Hugging Face. Accessed November 2, 2025. <https://huggingface.co/vandijlab/C2S-Scale-Gemma-2-27B>
8. Cell2Sentence: Teaching Large Language Models the Language of Biology. van Dijk Lab, Yale University. GitHub. Accessed November 2, 2025. <https://github.com/vandijlab/cell2sentence>

9. Google's Gemma AI Model Helps Discover a New Potential Cancer Therapy Pathway (Cell2Sentence-Scale 27B). Google. 2025. Accessed June 4, 2026. <https://blog.google/innovation-and-ai/products/google-gemma-ai-cancer-therapy-discovery/>
 10. Bridging Biology and AI: Yale and Google's Collaborative Breakthrough in Single-Cell RNA Analysis. Yale School of Medicine. 2025. Accessed June 4, 2026. <https://medicine.yale.edu/news-article/bridging-biology-and-ai-yale-and-googles-collaborative-breakthrough-in-single-cell-analysis/>
 11. The Explainability Challenge of Generative AI and LLMs. OCEG. Accessed November 2, 2025. <https://www.oceg.org/the-explainability-challenge-of-generative-ai-and-llms/>
 12. LLMs for Explainable AI: A Comprehensive Survey. arXiv:2504.00125. Accessed November 2, 2025. <https://arxiv.org/html/2504.00125v1>
 13. Harnessing AI in Multi-Modal Omics Data Integration: Paving the Path for the Next Frontier in Precision Medicine. PMC, National Institutes of Health. Accessed November 2, 2025. <https://pmc.ncbi.nlm.nih.gov/articles/PMC11972123/>
 14. A Zero-Shot LLM-Based Framework for Descriptive Gene-Gene Interaction Network Generation. OpenReview. Accessed November 2, 2025. <https://openreview.net/pdf?id=TKOKINCZNE>
 15. Personality Traits in Large Language Models. arXiv:2307.00184. Accessed November 2, 2025. <https://arxiv.org/html/2307.00184v4>
 16. Risk Profiling and Modulation for LLMs. arXiv:2509.23058. Accessed November 2, 2025. <https://arxiv.org/abs/2509.23058>
 17. Challenges and Opportunities for Developing More Generalizable Polygenic Risk Scores. PMC, National Institutes of Health. Accessed June 4, 2026. <https://pmc.ncbi.nlm.nih.gov/articles/PMC9828290/>
 18. Researchers Optimize Genetic Tests for Diverse Populations to Tackle Health Disparities. National Human Genome Research Institute. 2024. Accessed June 4, 2026. <https://www.genome.gov/news/news-release/researchers-optimize-genetic-tests-for-diverse-populations-to-tackle-health-disparities>
 19. DoD Instruction 5210.42: DoD Nuclear Weapons Personnel Reliability Assurance. U.S. Department of Defense. Accessed June 4, 2026. <https://www.esd.whs.mil/Portals/54/Documents/DD/issuances/dodi/521042p.pdf>
 20. Nuclear Security Series. International Atomic Energy Agency. Accessed June 4, 2026. <https://www.iaea.org/resources/nuclear-security-series>
 21. Application of Machine Learning in Predicting Aggressive Behaviors from Hospitalized Patients with Schizophrenia. PMC, National Institutes of Health. Accessed June 4, 2026. <https://pmc.ncbi.nlm.nih.gov/articles/PMC10067917/>
- 5, 2026. <https://arxiv.org/abs/2601.22401>