

Cultivating Collective Intelligence in Cyber Risk Governance: An Enterprise Architectural Approach

Sander Zeijlemaker, Francesca Venditti, Hannah Zabiepour, Kevin Zhang, Igor van Gemert, Angelica Marotta, Michael Siegel

Abstract

Cyber risk governance has become increasingly complex due to fragmented data, rapid technological change, and growing cognitive demands on decision-makers. Evidence shows that executives often make suboptimal cyber risk investment decisions, resulting in either excessive spending or insufficient protection. System dynamics (SD) provides a framework for analyzing such environments by modeling feedback structures, delays, and nonlinear interactions that shape cyber risk exposure and resilience. This study aims to strengthen cyber risk governance by augmenting human cognition through a collective intelligence architecture that integrates artificial intelligence with SD-based simulation. Using a design science approach, the architecture is validated through three proof-of-concept (PoC) implementations focusing on multi-agent decision-making, multi-agent data consolidation, and knowledge-based reasoning. Results indicate that AI-enhanced SD environments improve information integration, accelerate analysis, and support more consistent strategic decision-making while revealing new governance and validation challenges.

Keywords

Cybersecurity, artificial intelligence, enterprise architecture, governance, validation

Fragmented Data, Cognitive Overload, and Evolving Technologies Affecting Decision-Making

Cyber risk has become a central governance challenge. AI, cloud computing, Internet of Things, and automation increase organizational interconnectedness, expanding exposure to cyber threats and creating cascading financial, operational, and regulatory consequences (Bornet et al., 2021; Kotzias et al., 2023; Adekoya et al., 2025). Its governance is complicated by technical opacity, adaptive adversarial dynamics, and fragmented organizational data, which increase cognitive burdens and contribute to misaligned cyber investments (Armenia et al., 2018, 2021; Lowry et al., 2025; Lavoie & Liu, 2026; Zeijlemaker et al., 2022, 2025). Additionally, budgetary and expertise-related constraints further limit oversight by boards of directors, resulting in insufficient access to cybersecurity expertise (Coker, 2024; ISC2, 2025; World Economic Forum, 2024). To address these challenges, collective intelligence and human-machine collaboration offer opportunities to improve decision quality and speed (Malone, 2018; Malone & Bernstein, 2022).

Building on the demonstrated value of system dynamics (SD) for managerial foresight (Zeijlemaker & Siegel, 2023; Zeijlemaker et al., 2024), this study proposes integrating AI and agents with SD models through an enterprise architectural approach. Three proof-of-concept applications demonstrate how AI-enhanced modeling, multi-agent collaboration, and integrated data can strengthen strategic cyber risk governance.

Intelligence System Architecture Set-Up

Enterprise architecture provides a foundation for aligning strategy, organizational capabilities, and IT resources (The Open Group, 2019; Sherwood et al., 2009; Zachman, 2016). Building on this foundation, we propose an integrated intelligence architecture that combines system dynamics (SD), AI, data integration, knowledge-based reasoning, and human–computer interaction (HCI) to support cyber risk governance. Five interconnected components collectively operate as a virtual board advisor.

Data and Knowledge Foundation. The architecture integrates heterogeneous enterprise data through data-centric AI, including ingestion, normalization, semantic harmonization, and aggregation (Nittala, 2024; Akmal et al., 2025; Zha et al., 2025). Structured and unstructured information is mapped to unified schemas, ontologies, (Erraissi & Belangour, 2018) and SD model variables. A central knowledge base complements this layer by storing policies, regulations, ontologies, and historical cases. Hybrid retrieval and multi-hop reasoning provide traceable knowledge for decision-making (Batavia, 2025; Du et al., 2023).

Multi-Agent Orchestration and Reasoning. A distributed multi-agent system coordinates specialized agents for areas such as risk, finance, operations, and compliance (Anthropic, 2025; Jin et al., 2025; Trinkley et al., 2024). A central orchestrator decomposes tasks, coordinates agents, and aggregates their outputs. This structure enables parallel analysis while improving scalability, transparency, and robustness.

Dynamic Modeling and Strategic Foresight. SD provides the architecture with contextual and anticipatory reasoning by representing feedback, delays, nonlinear relationships, and adaptive behavior (Armenia et al., 2018, 2021; Jalali et al., 2019; Siegel & Zeijlemaker, 2023; Zeijlemaker et al., 2024, 2025). Agents support the construction and refinement of causal loop and stock-and-flow models (MetCalf, n.d.). Simulation, counterfactual analysis, and loop-dominance methods, including “Loops That Matter,” reveal future trajectories, systemic consequences, and potential leverage points (Schoenberg et al., 2020).

Human Interaction and Explainability. The HCI layer translates system intelligence into decision support through conversational interfaces, visualizations, and explainable AI (Reddy, 2024; Rong et al., 2024). Decision-makers can compare scenarios, examine uncertainty, modify assumptions, and trace recommendations to underlying data, knowledge, agent reasoning, and simulation mechanisms.

Architecture Validation. Together, these components create a modular architecture that connects enterprise data and knowledge with AI-enabled reasoning, dynamic modeling, and human decision-making. Three proof-of-concept applications focus on multi-agent coordination, data-centric integration, and knowledge-based inference—areas that remain comparatively underexplored in SD-based cyber risk governance. These applications provide an initial means of validating the proposed architecture.

Proof of Concept as Architectural Validation Mechanisms

In the development of complex socio-technical systems, early validation is essential for reducing uncertainty before large-scale implementation (Rodrigues Neto et al., 2018). A proof of concept (PoC) serves as an early validation mechanism for testing whether a proposed architectural idea or technological mechanism can function in principle (Kendig et al., 2016). Rather than producing a complete system, a PoC examines whether key architectural assumptions are technically viable under controlled conditions (Camburn et al., 2013).

The use of PoCs as feasibility demonstrations emerged in aerospace engineering in the 1960s (Jobin et al., 2020) and has since become common practice in domains such as biomedical research, public-sector innovation, and information technology. In these contexts, PoCs help identify opportunities and constraints early, enabling informed decisions before significant development resources are committed. As such, they represent an established approach to system design and development (Naumann & Jenkins, 1982; Janson & Smith, 1985).

In this research, PoCs serve as architectural validation mechanisms for the proposed collective intelligence architecture. They are used to assess the feasibility of core components, including multi-agent coordination, layered data-centric integration, and knowledge-based inference. Observing how these components perform and interact in practice allows the research to generate empirical insights into system behavior, identify unexpected challenges, and refine architectural assumptions. By combining feasibility testing with learning from real-world implementation, the PoC provides evidence that the proposed architectural approach is technically sound and practically applicable (Christie et al., 2012).

Lessons Learned from Three Proof-of-Concept Implementations

The three PoCs examine multi-agent coordination, data-centric integration, and knowledge-based inference. Each demonstrates a shift from traditional practices toward AI-enabled approaches, highlighting the architecture, real-world applications, and key implementation lessons.

Multi-Agent Architectures for Cyber Risk Governance: Results and Analytical Insights

Architectural Design. The PoC models executive roles as specialized AI agents with distinct KPIs, risk tolerances, and collaboration levels. A central chair agent coordinates deliberation and connects the agents to an established cyber risk simulation (Jalali et al., 2019; Kim et al., 2025; Zeijlemaker et al., 2025). A dashboard visualizes financial, risk, resilience, and agent-behavior outcomes, creating a dynamic environment for testing cyber investment strategies.

Experimental Results. Experiments varied collaboration and risk tolerance. High collaboration generated the strongest overall outcomes, reaching \$3.16 million in profit, 13% of systems at risk, and 89% availability under high risk tolerance. High collaboration with low risk tolerance reduced systems at risk to 5% and increased availability to 91%, but lowered profit to \$1.83 million. These results demonstrate clear trade-offs between profitability and resilience.

Analytical Insights and Conclusion. Collaboration improved both financial and resilience outcomes, while risk tolerance shaped the balance between protection and profitability. Agent negotiation produced stronger system-level outcomes than isolated optimization, but simulations also revealed AI drift and short-term KPI optimization. The proof of concept demonstrates the potential of multi-agent governance while emphasizing the need for transparent metrics, system constraints, and sustained human oversight.

Multi-Agent Environment for Consolidation and Calibration

Architectural Design. This PoC replaces fragmented reporting with a three-layer multi-agent architecture. Domain-specific agents integrate heterogeneous data from finance, security, IT, IAM, HR, and incident response. A coordinating agent consolidates and dynamically calibrates these inputs against a cyber-risk simulation reference mode (Zeijlemaker & Siegel, 2023; Zeijlemaker et al., 2024). Calibration is optimized using Theil's U and Theil statistics (Sterman, 2000; Theil, 1966).

Experimental Results. Testing against 12 months of historical data showed that automated multi-agent calibration ($U = 0.406$) outperformed the semi-automated approach ($U = 0.578$) by approximately 30%, while reducing calibration time from weeks to minutes. The results also remained robust under noisy conditions. However, unconstrained agent processing caused excessive database growth, requiring architectural redesign.

Analytical Insights and Conclusion. The results demonstrate that multi-agent integration can improve the speed, accuracy, and consistency of cyber risk intelligence while reducing fragmented reporting. However, the database growth incident illustrates the risks of unconstrained agent optimization. Effective multi-agent architecture, therefore, requires explicit boundaries, controlled data processing, architectural safeguards, and human oversight to deliver reliable and scalable anticipatory cyber risk governance.

AI-Powered Advisor: From Static Knowledge Repositories to Reasoning Systems

Architectural Design. This PoC shifts cyber governance from static document interpretation toward an AI-enabled advisory system for navigating fragmented cybersecurity regulations (Marotta et al., 2023; Marotta & Madnick, 2023, 2025). The architecture evolved into a hybrid retrieval-augmented generation (RAG) system combining a curated knowledge base, local retrieval and reranking, and cloud-based LLM reasoning (Muniyandi, 2025; Polan, 2025). This design grounds recommendations in verified sources while enabling contextual reasoning and interactive analysis.

Experimental Results. Structured Q&A testing by internal and external cybersecurity experts drove three architectural revisions to improve accuracy and reduce hallucinations. The final system accelerated time-to-insight by synthesizing dispersed regulatory and technical information into structured guidance. It also produced more consistent recommendations across different frameworks. However, testing revealed rapid user trust in AI outputs, creating risks that non-experts may overlook inaccurate information.

Analytical Insights and Conclusion. The hybrid RAG approach can reduce cognitive burden and improve the speed and consistency of cyber governance decisions. However, increased reliance on AI may encourage cognitive offloading and create cognitive debt, weakening human expertise and critical evaluation. Effective deployment, therefore, requires transparent reasoning, expert validation, and governance safeguards to ensure that AI augments rather than replaces human judgment.

From Decision Support to Collective Intelligence in Cyber Risk Governance

This study examined how human-machine collaboration can strengthen cyber risk governance. The findings support a shift from regularly used (hierarchical) decision support toward integrated socio-technical intelligence that combines human judgment with multi-agent coordination, system dynamics, data integration, knowledge-based reasoning, and explainable HCI.

The three PoCs demonstrate complementary benefits: multi-agent simulations improve strategic trade-off analysis; multi-agent data integration strengthens forecasting and responsiveness; and AI-enabled advisory systems transform fragmented knowledge into consistent governance guidance. Together, these capabilities improve situational awareness, foresight, and decision quality.

However, agent drift, unconstrained optimization, and cognitive overreliance highlight the need for transparency, architectural safeguards, and sustained human oversight. Overall, enterprise-scale

collective intelligence can strengthen cyber risk governance by augmenting rather than replacing human strategic judgment.

Implications for SD Model Validation and Governance of Intelligent Systems

The three PoCs reveal a broader governance challenge that extends beyond cybersecurity decision support. As organizations introduce intelligent systems capable of autonomous analysis, reasoning, and coordination, traditional governance frameworks and SD validation frameworks may become increasingly insufficient.

They demonstrate that intelligent systems introduce new forms of systemic risk, including behavioral drift, optimization beyond intended constraints, and the gradual erosion of human expertise through cognitive offloading. Consequently, the governance and validation of SD models must evolve to account for systems that learn, adapt, and act with a degree of autonomy. This shift requires rethinking both the principles of oversight and the mechanisms used to ensure reliability, accountability, and alignment with organizational objectives.

Governing AI: Lessons Learned from the Proofs of Concept

The three PoCs demonstrate that (Agentic) AI must be governed as a dynamic socio-technical system, addressing autonomy, behavioral change, human–AI interaction, and the preservation of human strategic judgment. Three governance shifts emerge.

Governance of Autonomous Decision-Making. Agentic AI redistributes decision authority by enabling autonomous analysis, negotiation, and recommendations. The experiments revealed risks of agents optimizing beyond intended organizational objectives. Governance should, therefore, define clear decision rights, autonomy boundaries, approval thresholds, and human accountability.

Continuous Monitoring and Escalation. Agent drift, uncontrolled optimization, and cognitive overreliance demonstrate that periodic oversight is insufficient (Parasuraman & Riley, 1997; Langosco et al., 2022). Three proposed indicators are: (i) the Human Strategic Autonomy Index, (ii) the AI Intent Drift Score, and (iii) the Adversarial Resilience Rating, which support continuous monitoring. Graduated escalation protocols should align human intervention with the impact and reversibility of AI decisions.

From Compliance to Resilience. Effective AI governance requires moving beyond static compliance toward continuous organizational resilience. Agentic systems can improve foresight, integration, and decision speed, but safeguards must prevent behavioral drift and the erosion of human expertise. Ultimately, AI should augment collective intelligence while humans retain strategic direction, accountability, and final authority.

Validation of System Dynamics Models in the Age of Artificial Intelligence

Traditional SD validation assesses structural integrity, behavioral consistency, parameter validity, extreme conditions, and sensitivity to ensure credible decision support (Barlas, 1996; Forrester & Senge, 1980; Sterman, 2000). However, AI changes the validation context by supporting automated causal discovery, model generation, data integration, calibration, and agent-based interaction.

These capabilities introduce new risks, including AI-generated structural bias, semantic data–model misalignment, adaptive calibration drift, agent-induced distortions, and human overreliance. Validation must, therefore, extend beyond the SD model to the broader socio-technical intelligence system.

We propose six complementary validation dimensions: (i) structural validation, (ii) data–model alignment, (iii) behavioral validation, (iv) adaptive calibration validation, (v) agent–model interaction validation, and (vi) human–AI interpretability validation. Together, these extend established SD practices to ensure that AI-enhanced models remain transparent, reliable, and suitable for strategic decision-making.

Acknowledgements

This research is funded by Cygent.

References

- Adekoya OA, Atlam HF, Lallie HS. 2025. Quantifying the multidimensional impact of cyber attacks in digital financial services: A systematic literature review. *Sensors* **25**(14): 4345. <https://doi.org/10.3390/s25144345>
- Akmal MU, Asif S, Koval L, Grossmann D. 2025. Layered data-centric AI to streamline data quality practices for enhanced automation. In *Artificial Intelligence and Data Science*. Springer, xx–xx. https://doi.org/10.1007/978-3-031-81542-3_11
- Anthropic. 2025. How we built our multi-agent research system. <https://www.anthropic.com/engineering/multi-agent-research-system>
- Armenia S, Angelini M, Nonino F, Palombi G, Schlitzer MF. 2021. A dynamic simulation approach to support the evaluation of cyber risks and security investments in SMEs. *Decision Support Systems* **147**: 113580. <https://doi.org/10.1016/j.dss.2021.113580>
- Armenia S, Ferreira Franco E, Nonino F, Spagnoli E, Medaglia CM. 2018. Towards the definition of a dynamic and systemic assessment for security risks. *System Research and Behavioral Science*. <https://doi.org/10.1002/sres.2556>
- Barlas Y. 1996. Formal aspects of model validity and validation in system dynamics. *System Dynamics Review* **12**(3): 183–210.
- Batavia A. 2025. Understanding knowledge-based agents in AI. Available from Nurix AI: <https://www.nurix.ai/blogs/understanding-knowledge-based-agents-ai>
- Bornet P, Barkin I, Wirtz J. 2021. *Intelligent Automation: Welcome to the World of Hyperautomation. Learn How to Harness Artificial Intelligence to Boost Business & Make Our World More Human*.
- Camburn, B., Dunlap, B., Kuhr, R., Viswanathan, V., Linsey, J., Jensen, D., Crawford, R., Otto, K., & Wood, K. (2013). Methods for prototyping strategies in conceptual phases of design: Framework and experimental assessment. In **Proceedings of the ASME International Design Engineering Technical Conferences (IDETC)**. <https://doi.org/10.1115/DETC2013-12065>
- Christie, E., Jensen, D., Buckley, R., & Ziegler, K. (2012). Prototyping strategies: Literature review and identification of critical variables. In **Proceedings of the ASEE Annual Conference & Exposition**. <https://peer.asee.org/21076>
- Coker J. 2024. US state CISOs warn of insufficient cybersecurity funding. Available from Infosecurity Magazine: <https://www.infosecurity-magazine.com/news/us-state-cisos-insufficient-funding/>
- Du Z, Zhou C, Yao J, Tu T, Cheng L, Yang H, Zhou J, Tang J. 2023. CogKR: Cognitive graph for multi-hop knowledge reasoning. *IEEE Transactions on Knowledge and Data Engineering* **35**(2): 1283–1295. <https://doi.org/10.1109/TKDE.2021.3104310>

- Erraissi A, Belangour A. 2018. Data sources and ingestion big data layers: Meta-modeling of key concepts and features. *International Journal of Engineering & Technology* 7(4): 3607–3612. <https://doi.org/10.14419/ijet.v7i4.19653>
- Forrester J, Senge P. 1980. Tests for building confidence in system dynamics models. *Studies in the Management Sciences*: 209–228.
- ISC2. 2025. 2025 ISC2 cybersecurity workforce study. <https://www.isc2.org/Insights/2025/12/2025-ISC2-Cybersecurity-Workforce-Study>
- Jalali MS, Siegel M, Madnick S. 2019. Decision-making and biases in cybersecurity capability development: Evidence from a simulation game experiment. *The Journal of Strategic Information Systems* 28(1): 66–82. <https://doi.org/10.1016/j.jsis.2018.09.003>
- Janson, M. A., & Smith, C. P. (1985). *Prototyping for systems development: A critical appraisal*. *MIS Quarterly*, 9(4), 305–316. <https://doi.org/10.2307/248951>
- Jobin, C., Masson, D., & Weil-Dubuc, P. (2020). *What does the proof-of-concept (POC) really prove? A historical perspective and a cross-domain analytical study*. [https://www.semanticscholar.org/paper/What-does-the-proof-of-concept-\(POC\)-really-prove-A-Jobin-Masson/1f3fd2b21edd6e84289809e893f80bd77fe20bb0](https://www.semanticscholar.org/paper/What-does-the-proof-of-concept-(POC)-really-prove-A-Jobin-Masson/1f3fd2b21edd6e84289809e893f80bd77fe20bb0)
- Jin W, Du H, Zhao B, Tian X, Shi B, Yang G. 2025. A comprehensive survey on multi-agent cooperative decision-making: Scenarios, approaches, challenges, and perspectives. *arXiv*. <https://arxiv.org/abs/2503.13415>
- Kendig, C. E. (2016). *What is proof of concept research and how does it generate epistemic and ethical categories for future scientific practice?* *Science and Engineering Ethics*, 22, 735–753. <https://doi.org/10.1007/s11948-015-9654-0>
- Kim G, Zeijlemaker S, Proudfoot J, Pal R, Siegel M. 2025. Balancing risk and reward in cybersecurity investment decisions. In *AMCIS 2025 Proceedings*. Association for Information Systems. https://aisel.aisnet.org/amcis2025/sig_sec/sig_sec/4
- Kotzias K, Bukhsh FA, Arachchige JJ, Daneva M, Abhishta A. 2023. Industry 4.0 and healthcare: Context, applications, benefits, and challenges. *IET Software* 17(3): 195–248.
- Langosco LD, Koch J, Sharkey LD, Pfau J, Krueger D. 2022. Goal misgeneralization in deep reinforcement learning. In *Proceedings of the 39th International Conference on Machine Learning* (Vol. 162, pp. 12004–12019). PMLR. <https://proceedings.mlr.press/v162/langosco22a.html>
- Lavoie JT, Liu X. 2026. Advancing risk management with cybersecurity intelligence: A design science approach. *Journal of Systems and Information Technology*. <https://doi.org/10.1108/JSIT-05-2025-0223>
- Lowry M, Vance A, Vance MD. 2025. Inexpert supervision: Field evidence on boards' oversight of cybersecurity. *Management Science*. <https://doi.org/10.1287/mnsc.2023.04147>
- Malone TW. 2018. *Superminds: The Surprising Power of People and Computers Thinking Together*. Little Brown Spark.
- Malone TW, Bernstein MS. (eds.). 2022. *Handbook of Collective Intelligence*. MIT Press.
- Marotta A, Madnick S. 2025. Analyzing and categorizing emerging cybersecurity regulations. *ACM Computing Surveys* 58(2): Article 51. <https://doi.org/10.1145/3757318>
- Marotta A, Madnick BF, Madnick S. 2023. Navigating the wave of cybersecurity regulations: A systematic analysis of emerging regulatory developments. In *Proceedings of the Americas Conference on Information Systems*. AMCIS 2023. https://aisel.aisnet.org/amcis2023/sig_sec/sig_sec/12

- Marotta A, Madnick S. 2023. Regulating cyber incidents: A review of recent reporting requirements. In *Proceedings of the 20th International Conference on Security and Cryptography (SECRYPT)* (pp. 410–416). SciTePress. <https://doi.org/10.5220/0012086000003555>
- Metcalf S. n.d. BEAMS: Behavioral, economic, and agent-based modeling & simulation. Available from University at Buffalo: <https://www.buffalo.edu/ai-data-science/research/beams.html>
- Muniyandi V. 2025. RAG architecture design patterns balancing retrieval depth and generative coherence. *International Journal of Computer Applications* **187**(12): 34–38. <https://doi.org/10.5120/ijca2025925142>
- Naumann, J. D., & Jenkins, A. M. (1982). *Prototyping: The new paradigm for systems development*. *MIS Quarterly*, **6**(3), 29–44. <https://doi.org/10.2307/248989>
- Nittala EP. 2024. AI-powered multimodal data integration in ERP systems for holistic enterprise analytics. *International Journal of Artificial Intelligence, Data Science, and Machine Learning* **5**(2): 107–115. <https://doi.org/10.63282/3050-9262.IJAIDSML-V5I2P112>
- The Open Group. 2019. TOGAF® version 9.1 (Open Group Standard G116). <https://pubs.opengroup.org/architecture/togaf91-doc/arch/>
- Parasuraman R, Riley V. 1997. Humans and automation: Use, misuse, disuse, abuse. *Human Factors* **39**(2): 230–253. <https://doi.org/10.1518/001872097778543886>
- Polan Ş. 2025. Retrieval-Augmented Generation: Architecture, Techniques, and Evaluations. *Journal of Modern Technology and Engineering* **10**(1): 42–56. <https://doi.org/10.62476/jmte.10142>
- Reddy STA. 2024. Human–computer interaction techniques for explainable artificial intelligence systems. *Research & Review: Machine Learning and Cloud Computing* **3**(1): 1–7. <https://doi.org/10.46610/RTAIA.2024.v03i01.001>
- Rodrigues Neto, A. J., Borges, M. M., & Roque, L. (2018). *A preliminary study of proof of concept practices and their connection with information systems and information science*. In **Proceedings of the Sixth International Conference on Technological Ecosystems for Enhancing Multiculturality (TEEM 2018)** (pp. 270–275). ACM. <https://doi.org/10.1145/3284179.3284226>
- Rong Y, Wang W, Xu Z, Tan Q, Liu Y. 2024. Towards human-centered explainable AI: A survey of user studies for model explanations. *IEEE Transactions on Pattern Analysis and Machine Intelligence* **46**(4): 2104–2122. <https://doi.org/10.1109/TPAMI.2023.3339266>
- Schoenberg W, Davidsen P, Eberlein R. (2020). Understanding model behavior using loops that matter (Version 2). *arXiv*. <https://doi.org/10.48550/arXiv.1908.11434>
- Sherwood J, Clark A, Lynas D. 2009. *Enterprise Security Architecture (SABSA White Paper)*. SABSA Limited.
- Sterman J. 2000. *Business Dynamics: System Thinking and Modeling for a Complex World*. Irwin McGraw-Hill.
- Theil H. 1966. *Applied Economic Forecasting*. Rand McNally.
- Trinkley KE, Blalock SJ, Bender BG, et al. 2024. Leveraging artificial intelligence to advance implementation science. <https://pmc.ncbi.nlm.nih.gov/articles/PMC10880216/>
- World Economic Forum. 2024. The cybersecurity industry faces a talent shortage. Here's how to close the gap. <https://www.weforum.org/stories/2024/04/cybersecurity-industry-talent-shortage-new-report/>
- Zachman JA. 2016. The framework for enterprise architecture: Background, description, and utility. <https://www.zachman.com/resources/ea-articles-reference/327-the-framework-for-enterprise-architecture-background-description-and-utility-by-john-a-zachman>
- Zha D, Cheng L, Yin H, Peng Y, Lu J. 2023. Data-centric artificial intelligence: A survey. *ACM Computing Surveys* **56**(7): Article 147.

- Zeijlemaker S, Rouwette EAJA, Cunico G, Armenia S, von Kutzschenbach M. 2022. Decision-makers' understanding of cyber-security's systemic and dynamic complexity: Insights from a board game for bank managers. *Systems* **10**(2):49. <https://doi.org/10.3390/systems10020049>
- Zeijlemaker S, Pal R, Proudfoot J, Siegel M. 2025. Advancing cyber risk by reducing strategic control gaps. In *AMCIS 2025 Proceedings*. Association for Information Systems. https://aisel.aisnet.org/amcis2025/sig_sec/sig_sec/28
- Zeijlemaker S, Pal R, Siegel M. 2024. Strengthening managerial foresight to defeat cyber threats. In *Proceedings of the 30th Americas Conference on Information Systems (AMCIS 2024)*. Association for Information Systems. <https://aisel.aisnet.org/amcis2024/security/security/18>
- Zeijlemaker S, Siegel M. 2023. Capturing the dynamic nature of cyber risk: Evidence from an explorative case study. In *Proceedings of the 56th Hawaii International Conference on System Sciences (HICSS-56)*. Association for Information Systems. <https://aisel.aisnet.org/hicss-56/os/cybersecurity/5>