

A Balanced Enterprise Architecture Framework for AI Adoption

Bhagyeshkumar Chokhawala 

Comcast, Philadelphia, Pennsylvania, USA, bhagyeshkumar_chokhawala@comcast.com

Abstract: Artificial intelligence is increasingly central to enterprise transformation, yet many organizations still approach adoption through disconnected pilots, vendor-driven experimentation, or narrow technology programs. This paper presents a balanced enterprise architecture framework for AI adoption that integrates business, technology, developer experience, and user experience perspectives with explicit security and ethical AI considerations. Governance and operating model provide the control and coordination needed to scale AI safely, responsibly, and sustainably. The paper introduces a TOGAF-style capability map and a practical execution toolkit covering readiness assessment, use case prioritization, governance gates, reference implementation assets, stakeholder conflict resolution, phased roadmap delivery, threat modeling, data protection, ethical impact assessment, bias review, accountability, and incident response. The proposed framework helps enterprises translate AI ambition into structured, governable, secure, and value-oriented implementation. It argues that enterprise AI maturity should be measured not only by the number of models deployed, but also by the alignment among business capability, technical readiness, engineering enablement, security assurance, responsible governance, and trustworthy human outcomes.

Keywords: Enterprise architecture, artificial intelligence, AI adoption, governance, AI security, AI ethics, responsible AI, MLOps, developer experience, user trust, capability mapping

Introduction

Artificial intelligence has moved from isolated innovation programs into the center of enterprise strategy. Enterprises increasingly seek to use AI for knowledge access, decision support, workflow automation, software delivery, and customer engagement. However, AI adoption is rarely constrained by model availability alone. In practice, scale is limited by fragmented data, unclear accountability, weak controls, uneven engineering maturity, and low user trust. Enterprise architecture provides a suitable discipline for structuring these dependencies because it connects business intent, information flows, applications, technology platforms, and governance into a coherent operating model (Ross et al., 2006; Lankhorst, 2017; The Open Group, 2022; Tabassi et al., 2023).

The architectural challenge is therefore not merely how to integrate a model API, but how to institutionalize AI as an enterprise capability. A balanced strategy must align business value with technology readiness, developer enablement, trustworthy user outcomes, security assurance, and ethical accountability. When these dimensions are optimized independently, organizations often experience an innovation paradox: AI pilots appear promising at the edge of the enterprise while production adoption remains fragmented, costly, risky, and difficult to trust (ISO/IEC, 2023a, 2023b; Kreuzberger et al., 2023; Sculley et al., 2015; Tabassi et al., 2023). This paper addresses that problem by proposing a balanced enterprise architecture framework and a practical execution toolkit for roadmap design.

Security and ethics are design-time architecture constraints. Insecure prompt flows, unprotected knowledge stores, weak identity boundaries, biased data, opaque automated decisions, and insufficient human oversight can turn AI pilots into privacy, operational, and trust risks (Floridi et al., 2018; Jobin et al., 2019; OWASP Foundation, 2025; Tabassi et al., 2023).

Research and Practice Background

Enterprise architecture literature has long emphasized that effective business execution depends on aligning strategy, operating model, process standardization, integration, and technology platforms (Lankhorst, 2017; Ross et al., 2006). TOGAF reinforces this alignment by linking business, data, application, and technology architecture through an iterative Architecture Development Method and capability-based planning approach (The Open Group, 2022). These foundations remain relevant for AI adoption because AI amplifies, rather than replaces, the need for coherent architecture.

Recent AI governance and engineering literature adds four important considerations. First, AI introduces distinct risks related to safety, bias, security, explainability, and lifecycle uncertainty, which are reflected in the NIST AI Risk Management Framework and ISO standards for AI management and AI risk management (ISO/IEC, 2023a, 2023b; Tabassi et al., 2023). Second, AI systems create significant operational and maintenance burdens, including feedback loops, hidden dependencies, and technical debt (Sculley et al., 2015). Third, production AI requires disciplined lifecycle practices such as evaluation, observability, and deployment orchestration, which have been codified in the emerging MLOps literature (Kreuzberger et al., 2023). Fourth, AI security must address AI-specific attack surfaces such as prompt injection, insecure output handling, sensitive information disclosure, model supply-chain exposure, and misuse risks (OWASP Foundation, 2025).

Ethical AI literature also shows that transparency, fairness, accountability, privacy, non-maleficence, and human agency must be converted into operational controls (Floridi et al., 2018; Jobin et al., 2019; Shneiderman, 2020). For enterprise architecture, these requirements should be traceable to capabilities, data flows, decision rights, explanations, escalation paths, and review criteria.

User and developer concerns are equally consequential. Research on human-AI interaction emphasizes transparency, calibration of trust, override mechanisms, and context-sensitive guidance (Amershi et al., 2019; Shneiderman, 2020). Developer experience research shows that productivity, satisfaction, and sustainable performance are shaped not only by tools but also by clarity, workflow friction, and the quality of the working environment (Fagerholm & Munch, 2012; Greiler et al., 2023). These combined insights suggest that AI strategy should be framed as a socio-technical architecture problem rather than a model selection problem.

AI adoption also introduces persistent stakeholder tension because different actors optimize for different outcomes. Business sponsors emphasize time-to-value and measurable impact; architects emphasize interoperability and reuse; security and legal teams emphasize risk containment; product and delivery teams emphasize speed and autonomy; and end users emphasize convenience, transparency, and reliability. Stakeholder theory suggests that managerial attention is shaped by differences in power, legitimacy, and urgency (Freeman, 2010; Mitchell et al., 1997). In enterprise AI programs, these differences are not peripheral governance concerns; they are central design constraints that must be explicitly surfaced and negotiated.

Balanced Enterprise Architecture Framework for AI Adoption

The proposed framework organizes AI adoption into four primary perspectives supported by two cross-cutting enablers. The primary perspectives are business, technology, developer experience, and user experience. The cross-cutting enablers are governance, including responsible AI, privacy, security, compliance, ethics, and auditability, and operating model, including roles, funding, skills, change management, and measurement. Security and ethical AI operate as explicit design lenses within these enablers and across the four perspectives. This 4+2 formulation creates a balanced architecture model for enterprise AI.

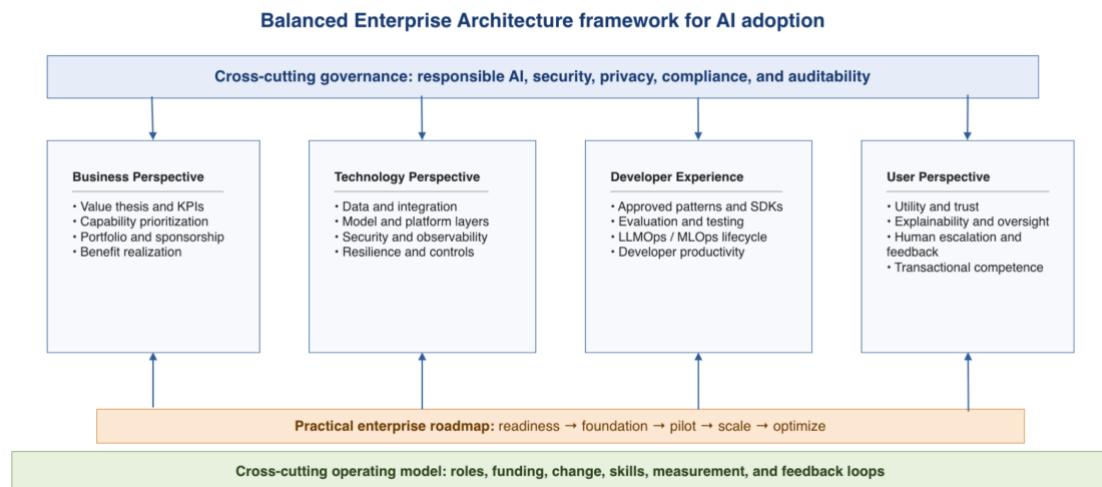


Figure 1. Balanced enterprise architecture framework for AI adoption

Business perspective

From a business perspective, AI adoption should be anchored in measurable value rather than technology enthusiasm. Suitable entry points include productivity improvement, cost reduction, cycle-time compression, service quality, risk reduction, and new digital offerings. Architecture teams should map candidate AI initiatives to business capabilities, value streams, and sponsoring owners, based on the premise that architecture exists to support business execution (Ross et al., 2006; The Open Group, 2022). In this view, the first design task is not model selection but capability prioritization.

Technology perspective

The technology perspective defines the structural conditions for safe scale. It includes data readiness, authoritative knowledge sources, integration patterns, platform services, observability, security controls, and resilience mechanisms. In many enterprises, retrieval-augmented generation and workflow orchestration are valuable patterns because they ground model outputs in enterprise knowledge and connect AI interactions to controlled business processes (Kreuzberger et al., 2023; Lewis et al., 2020). Architectural separation between systems of record, systems of engagement, and systems of intelligence is critical to avoid bypassing transactional integrity.

Developer experience perspective

The developer-experience perspective translates architectural intent into delivery capability. Teams need approved access patterns, reusable SDKs, templates, testing harnesses, versioned prompts or policies, and lifecycle automation. Without such enablement, AI delivery tends to fragment into isolated scripts and ad hoc integrations. Prior research on developer experience and delivery performance indicates that improvements in the development environment support better productivity, confidence, and consistency (Fagerholm & Munch, 2012; Greiler et al., 2023). For enterprise AI, this means the governed path must also be the easiest path.

User perspective

The user perspective focuses on utility, trust, and transactional completeness. Human-AI interaction research emphasizes that successful systems communicate capabilities and limitations, offer corrective pathways, and keep people appropriately in control (Amershi et al., 2019; Shneiderman, 2020). For enterprise settings, an AI-enabled interface should not be judged only by conversational fluency or reduced clicks. It must also preserve data quality, policy compliance,

and downstream operational executability. In other words, user experience and enterprise execution must remain coupled.

Security and ethical AI lens

Security and ethical AI should be treated as an embedded architecture lens across the four perspectives rather than an after-the-fact review. Security controls should include data classification, least-privilege access, secure model gateways, prompt and output protections, adversarial testing, monitoring, and incident response. Ethical controls should include fairness and bias assessment, explainability, human override, contestability, and clear accountability for decisions that affect users, employees, or customers (Floridi et al., 2018; Jobin et al., 2019; OWASP Foundation, 2025).

Cross-cutting enablers

Governance and operating model provide the connective tissue across all four perspectives. Governance establishes policies, security baselines, ethical AI requirements, risk tiers, impact assessments, review criteria, and monitoring expectations in alignment with AI risk and management standards (Floridi et al., 2018; ISO/IEC, 2023a, 2023b; Tabassi et al., 2023). The operating model defines decision rights, funding channels, product ownership, training, support, escalation paths, and value review cadences. Together, these enablers make the architecture actionable rather than aspirational and ensure that security and ethics are integrated into normal delivery decisions.

Competing priorities and stakeholder conflict resolution

Competing priorities are a normal condition of enterprise AI adoption. Different stakeholder groups often pursue partially incompatible objective functions: growth versus control, standardization versus local speed, automation versus accountability, and simplicity versus completeness. These tensions can intensify in AI programs because model behavior is probabilistic and because benefits and risks are distributed unevenly across stakeholders. For example, a business unit may favor rapid deployment of a copilot, while security seeks tighter data boundaries, architecture seeks reusable patterns, and operations seeks supportability. Without an explicit method for reconciling these priorities, disagreements tend to reappear as backlog churn, shadow implementations, architecture exceptions, or post-launch remediation (Boehm et al., 2001; Freeman, 2010; Mitchell et al., 1997).

A balanced framework, therefore, needs explicit mechanisms for conflict resolution. Competing priorities should be surfaced at the level of goals, constraints, and decision rights before teams harden solution choices, because negotiation is more effective when it focuses on interests and assumptions rather than on already-defended implementation preferences (Boehm et al., 2001). In practice, this means using stakeholder maps, objective matrices, documented decision rights, architecture review forums with both business and control stakeholders, and escalation rules that distinguish non-negotiable controls from locally optimizable design choices. The objective is not perfect consensus; it is bounded alignment around shared outcomes, acceptable risk, and documented trade-offs.

TOGAF-Style Capability Map

To translate the balanced framework into enterprise planning language, Table 1 presents a TOGAF-style capability map. The capability map is designed to bridge strategic ambition and implementation execution by organizing AI adoption across architecture, engineering, governance, security, ethics, and operations domains (Lankhorst, 2017; The Open Group, 2022).

Table 1. TOGAF-style capability map for enterprise AI adoption

Domain	Capability	Illustrative scope
Strategy and Governance	AI vision and strategy	Principles, value thesis, investment logic, target outcomes
Strategy and Governance	Portfolio management	Use case intake, prioritization, benefits tracking, roadmap decisions
Strategy and Governance	AI governance	Policies, responsible AI controls, approval workflows, auditability
Strategy and Governance	Risk and compliance	Privacy, legal review, security posture, regulatory alignment
Business Architecture	Capability alignment	Business capability mapping, value streams, operating model fit
Business Architecture	Process transformation	Workflow redesign, augmentation targets, exception handling
Business Architecture	Decision intelligence	Decision support, recommendations, predictive insights
Data Architecture	Data readiness	Quality, lineage, access, stewardship, reference data
Data Architecture	Knowledge management	Enterprise content, retrieval, grounding, provenance
Application Architecture	AI-enabled applications	Assistants, copilots, summarization, recommendation services
Application Architecture	Workflow orchestration	Tool routing, policy checks, human-in-the-loop pathways
Technology Architecture	AI platform services	Model gateways, vector stores, orchestration runtimes
Technology Architecture	Security and operations	IAM, encryption, logging, reliability, incident response
Delivery and DevEx	Engineering enablement	SDKs, templates, reference architectures, standards
Delivery and DevEx	Evaluation and LLMOps	Prompt testing, regression, deployment lifecycle, rollback
Operations and Service	Service management and value realization	Runbooks, support, KPI review, optimization

Toolkit for Execution

The framework becomes operational through a compact toolkit that can be embedded into architecture governance, portfolio review, and product delivery. The toolkit below is intentionally lightweight: it is designed to support decision-making without adding excessive procedural friction.

Readiness assessment

A practical readiness assessment should score at least eight dimensions: business alignment, data readiness, platform readiness, governance maturity, security readiness, ethical impact management, developer enablement, and user or operational readiness. A three-level maturity scale, such as emerging, defined, and scaled, is usually sufficient for early planning. The assessment should identify gaps in identity controls, data protection, monitoring, human oversight, bias review, and incident response (ISO/IEC, 2023a, 2023b; Kreuzberger et al., 2023; OWASP Foundation, 2025; Sculley et al., 2015; Tabassi et al., 2023).

Use case prioritization matrix

A simple weighted score can be used to separate quick wins, strategic investments, and defer-or-redesign candidates. The most effective portfolios usually balance one or two highly visible productivity wins with one or two structurally important platform or process use cases.

Table 2. Illustrative use case prioritization matrix

Criterion	Question	Suggested weighting
Business value	Will the use case improve revenue, cost, quality, speed, or risk?	High
Strategic fit	Does it support priority capabilities or transformation themes?	High
Data readiness	Are trusted data and knowledge sources available?	High
Technical feasibility	Can the solution be integrated and operated with acceptable complexity?	Medium
Governance risk	What level of privacy, safety, or regulatory risk is introduced?	High
Adoption readiness	Will users and operators trust and use the solution?	Medium
Security exposure	Could the use case expose sensitive data, privileged actions, external attack paths, or unsafe model behavior?	High
Ethical impact	Could outputs or decisions affect fairness, autonomy, transparency, accessibility, or user well-being?	High

Governance gates

Governance gates should be tied to risk rather than to organizational hierarchy. A workable sequence includes a strategic fit review, a readiness review, a risk and control review, a security and ethical impact review, a production readiness review, and a post-launch value review. This aligns with the lifecycle orientation of NIST AI RMF, ISO/IEC 42001, and ISO/IEC 23894, while remaining usable in an enterprise delivery context (ISO/IEC, 2023a, 2023b; Tabassi et al., 2023).

Security and ethics controls

The execution toolkit should include a control checklist before a pilot proceeds to production. Security controls should cover identity and access management, data classification, encryption, model gateway policies, retrieval source control, prompt-injection testing, output handling, logging, and incident response. Ethical controls should cover affected stakeholders, fairness and bias risks, transparency notices, human oversight, appeal or correction paths, accessibility, and evidence that the system remains aligned with its stated purpose (ISO/IEC, 2023a, 2023b; OWASP Foundation, 2025; Tabassi et al., 2023).

Reference implementation assets

Architects and platform teams should provide a starter pack for product teams: model access libraries, retrieval patterns, approved integration methods, prompt and policy templates, security guardrails, ethical impact checklists, evaluation checklists, observability standards, runbooks, and sample human-in-the-loop flows. These assets convert policy into reusable engineering capability and reduce the probability of uncontrolled shadow implementations (Greiler et al., 2023; Kreuzberger et al., 2023; Sculley et al., 2015).

Stakeholder alignment and conflict resolution toolkit

The execution toolkit should include a small set of conflict-resolution artifacts. These include a stakeholder salience map, an objective-trade-off matrix, a decision-rights or RACI view, a conflict

log with escalation thresholds, and a decision register that records rationale, owner, date, and review trigger. Used together, these artifacts make competing priorities visible early, reduce rediscovery of prior decisions, and help teams separate fixed controls from negotiable design freedoms. They are especially useful when AI initiatives span product, data, legal, security, architecture, and operations groups, each with different success measures (Boehm et al., 2001; Freeman, 2010; Mitchell et al., 1997).

Practical Enterprise Roadmap

A practical roadmap should move from readiness to scale through a staged sequence that treats foundational controls, security assurance, ethical impact management, and delivery enablement as prerequisites for expansion. It should also include explicit stakeholder-alignment checkpoints because many AI delays stem not from model performance but from unresolved differences in objectives, risk tolerance, ownership, accountability, and operational expectations among participating groups. Roadmap phases should therefore validate both technical readiness and the quality of cross-stakeholder agreements needed for execution.

Table 3. Practical enterprise roadmap for balanced AI adoption.

Phase	Primary goal	Key activities	Core outputs
1. Strategy and readiness	Define vision and assess maturity	Executive alignment, current-state assessment, first-wave use cases	Principles, maturity scorecard, prioritized backlog
2. Foundation and guardrails	Establish reusable controls and services	Model access, data patterns, logging standards, approval workflows	Reference architecture, starter toolkit, governance controls
3. Pilot and validate	Demonstrate value with controlled scope	Build 2-5 pilots, evaluate outcomes, capture lessons	Pilot reports, pattern refinements, scale decisions
4. Scale and industrialize	Expand reuse across teams and domains	Shared platform onboarding, support model, training, KPI dashboards	Enterprise rollout plan, reusable services, runbooks
5. Optimize and transform	Embed AI into normal planning and operations	Portfolio tuning, cost-performance optimization, continuous governance	Improved portfolio, updated maturity targets, next-wave roadmap

A useful architectural principle is that each phase should have explicit exit criteria. For example, pilot completion should not authorize broad scale-out unless governance, observability, user trust, and operational support are demonstrably adequate. This avoids the common error of scaling promising prototypes before the enterprise is prepared to operate them.

Discussion

The proposed framework does not argue for heavy centralization. Instead, it argues for architectural coherence under conditions of plural stakeholder objectives. Product teams should retain flexibility within bounded standards, while platform, security, ethics, and governance teams provide shared services, risk controls, and decision support. This balance is especially important in AI because model behavior is probabilistic, data quality is uneven, user trust can be lost quickly after visible failures, security vulnerabilities can create enterprise-wide exposure, and stakeholder conflict can otherwise turn ordinary design trade-offs into organizational deadlock (Amershi et al., 2019; Boehm et al., 2001; Freeman, 2010; ISO/IEC, 2023a, 2023b; Mitchell et al., 1997; OWASP Foundation, 2025; Shneiderman, 2020; Tabassi et al., 2023).

The framework also highlights that AI adoption is not purely a technology modernization issue. Business sponsorship without technical readiness leads to stalled pilots. Platform investment without use-case discipline leads to underutilized infrastructure. User-centric design without transactional rigor leads to downstream operational defects. DevEx neglect leads to fragmented implementation. Balanced adoption requires all four primary perspectives to move together, even if not at an identical speed.

For enterprise architects, the practical value of this model is traceability. It connects strategic intent to capability gaps, governance expectations, security controls, ethical requirements, delivery assets, and measurable business outcomes. That traceability improves portfolio decisions, architecture reviews, and executive communication because AI is framed as an enterprise capability with explicit dependencies rather than as a generic innovation theme.

Adding security and ethics also changes how AI success is measured. A mature AI solution should demonstrate secure access, protected data flows, audit trails, explainable outputs, fair treatment of affected users, human accountability, and a maintained path for correction or rollback.

Conclusions

AI adoption should be treated as an enterprise architecture problem as much as a model or tooling problem. A balanced strategy integrates business value, technology readiness, developer enablement, user trust, security assurance, and ethical accountability, while governance and operating model provide the discipline needed to scale safely. The framework, capability map, toolkit, and roadmap proposed here offer a practical approach to translating AI ambition into secure, ethical, and governable execution.

In this sense, enterprise AI maturity is not defined solely by how many models are deployed, but by how effectively an organization aligns its architecture, risk controls, engineering systems, human outcomes, security posture, and ethical commitments to sustain value creation. Enterprises that adopt this balanced approach are more likely to move beyond experimentation toward durable and trusted AI capability.

References

- Amershi, S., Weld, D., Vorvoreanu, M., Fourney, A., Nushi, B., Collisson, P., Suh, J., Iqbal, S., Bennett, P. N., Inkpen, K., Teevan, J., Kikin-Gil, R., & Horvitz, E. (2019). Guidelines for human-AI interaction. In *Proceedings of the 2019 CHI Conference on Human Factors in Computing Systems* (Paper no. 3, pp. 1-13). Association for Computing Machinery. <https://doi.org/10.1145/3290605.3300233>
- Boehm, B., Gruenbacher, P., & Briggs, R. O. (2001). Developing groupware for requirements negotiation: Lessons learned. *IEEE Software*, 18(3), 46-55. <https://doi.org/10.1109/52.922725>
- Brynjolfsson, E., Li, D., & Raymond, L. R. (2025). Generative AI at work. *Quarterly Journal of Economics*, 140(2), 889-942. <https://doi.org/10.1093/qje/qjae044>
- Fagerholm, F., & Munch, J. (2012). Developer experience: Concept and definition. In *2012 International Conference on Software and System Process* (pp. 73-77). IEEE. <https://doi.org/10.1109/ICSSP.2012.6225984>
- Floridi, L., Cowls, J., Beltrametti, M., Chatila, R., Chazerand, P., Dignum, V., Luetge, C., Madelin, R., Pagallo, U., Rossi, F., Schafer, B., Valcke, P., & Vayena, E. (2018). AI4People - An ethical framework for a good AI society: Opportunities, risks, principles, and recommendations. *Minds and Machines*, 28, 689-707. <https://doi.org/10.1007/s11023-018-9482-5>
- Freeman, R. E. (2010). *Strategic management: A stakeholder approach*. Cambridge University Press. <https://doi.org/10.1017/CBO9781139192675>
- Greiler, M., Storey, M.-A., & Noda, A. (2023). An actionable framework for understanding and improving developer experience. *IEEE Transactions on Software Engineering*, 49(4), 1888-1907. <https://doi.org/10.1109/TSE.2022.3175660>
- ISO/IEC. (2023a). *ISO/IEC 23894:2023 Information technology - Artificial intelligence - Guidance on risk management*. International Organization for Standardization.
- ISO/IEC. (2023b). *ISO/IEC 42001:2023 Information technology - Artificial intelligence - Management system*. International Organization for Standardization.

- Jobin, A., Ienca, M., & Vayena, E. (2019). The global landscape of AI ethics guidelines. *Nature Machine Intelligence*, 1(9), 389-399. <https://doi.org/10.1038/s42256-019-0088-2>
- Kreuzberger, D., Kuhl, N., & Hirschl, S. (2023). Machine learning operations (MLOps): Overview, definition, and architecture. *IEEE Access*, 11, 31866-31879. <https://doi.org/10.1109/ACCESS.2023.3262138>
- Lankhorst, M. (2017). *Enterprise architecture at work: Modeling, communication and analysis* (4th ed.). Springer. <https://doi.org/10.1007/978-3-662-53933-0>
- Lewis, P., Perez, E., Piktus, A., Petroni, F., Karpukhin, V., Goyal, N., Kuttler, H., Lewis, M., Yih, W.-T., Rocktaschel, T., Riedel, S., & Kiela, D. (2020). Retrieval-augmented generation for knowledge-intensive NLP tasks. In *Advances in Neural Information Processing Systems* 33, 9459-9474.
- Mitchell, R. K., Agle, B. R., & Wood, D. J. (1997). Toward a theory of stakeholder identification and salience: Defining the principle of who and what really counts. *Academy of Management Review*, 22(4), 853-886. <https://doi.org/10.5465/AMR.1997.9711022105>
- OWASP Foundation. (2025). *OWASP Top 10 for large language model applications 2025*. <https://owasp.org/www-project-top-10-for-large-language-model-applications/>
- Ross, J. W., Weill, P., & Robertson, D. C. (2006). *Enterprise architecture as strategy: Creating a foundation for business execution*. Harvard Business Press.
- Sculley, D., Holt, G., Golovin, D., Davydov, E., Phillips, T., Ebner, D., Chaudhary, V., Young, M., Crespo, J.-F., & Dennison, D. (2015). Hidden technical debt in machine learning systems. In *Advances in Neural Information Processing Systems* 28, 2503-2511.
- Shneiderman, B. (2020). Human-centered artificial intelligence: Reliable, safe and trustworthy. *International Journal of Human-Computer Interaction*, 36(6), 495-504. <https://doi.org/10.1080/10447318.2020.1741118>
- Tabassi, E., Burns, K. J., Hadjimichael, M., Molina-Markham, A. D., & Sexton, J. T. (2023). *Artificial intelligence risk management framework* (AI RMF 1.0) (NIST AI 100-1). National Institute of Standards and Technology. <https://doi.org/10.6028/NIST.AI.100-1>
- The Open Group. (2022). *TOGAF Standard* (10th ed.). The Open Group.