

Module: Module 3 - Capstone
Project

From Tension to Outcome:
Transforming into AI-Native
Finance Organisation
Module 3 - Assignment

Yiwei Wang
**Institution: Swiss School of Business
Research**
Word Count: 6871

Table of Contents

Executive Proposition	2
1. Introduction: Context and Case for Transformation	3
2. Integrative Synthesis: From Strategic Tension to Decision Resolution	4
3. Transformational Design: Defining the AI-Native Finance Organisation	6
4. Implementation Pathway and Value Realisation	10
5. Decision Legitimacy in Practice	14
6. Organisational Resistance and Practical Constraints.....	16
7. Measuring Success After Twelve Months	18
8. Leadership and Governance: Finance Stewardship, Limits, and Boundary Conditions	20
9. Strategic Evolution and Critical Self-Reflection in AI-Native Finance	23
10. Conclusion and Outlook: Redefining Beyond Finance and Reinsurance	24
Appendices	27
Appendix A: Summary of Module 1 (Strategic Tension)	27
Appendix B: Summary of Module 2 (Collaborative Decision Architecture)	28
Cumulative role of Appendices A and B	30
Bibliography	31

Governing Judgement in the Age of AI: Designing the AI-Native Finance Organisation

Executive Proposition

This capstone presents an academically grounded business proposal addressing a structural organisational problem: how finance organisations sustain legitimate, accountable decision-making when AI becomes a permanent cognitive participant.

Research question

How can a finance organisation redesign its operating model so that AI-augmented decision-making becomes a source of competitive advantage rather than organisational friction in uncertain, regulated environments (Shrestha, Ben-Menahem and von Krogh, 2019)?

Across insurance and reinsurance finance, AI starts shaping core decision processes, yet many organisations still rely on pre-AI assumptions about authority, judgement, and defensibility. The result is a structural misfit between AI-mediated cognition and existing organisational forms, visible in delayed commitment, informal override, governance congestion, and declining trust (Orlikowski, 2007).

In professional practice, this misfit becomes most visible when organisations move from experimentation to commitment. The issue is rarely a lack of use cases, but an inability to convert AI into decisions that can be defended under uncertainty (Author's professional artefact, 2026a).

Core claim

AI-related tensions in finance cannot be resolved through technology deployment, governance expansion, or process redesign alone. They require redesign of the finance operating model itself (Mintzberg, 1994).

The proposal defines an **AI-native finance operating model**, in which AI continuously shapes problem framing, options, and trade-offs. "AI-native" therefore refers not to automation, but to the redesign of authority, accountability, and legitimacy (Shrestha, Ben-Menahem and von Krogh, 2019).

The model rests on three shifts:

1. Judgement is embedded within decision processes, rather than reappearing as late override.
2. Accountability is exercised through decision stewardship, rather than only hierarchical approval.
3. Legitimacy is produced ex ante, through explicit assumptions and trade-offs, rather than retrospectively through governance artefacts (Power, 2007).

In practical terms, decision legitimacy is the confidence that a decision can be defended before it is made, while it is being challenged, and after outcomes are realised (Power, 2007). This matters because it reduces escalation, supports earlier commitment, and strengthens regulatory defensibility.

In AI-intensive finance, competitive advantage derives less from superior analytics than from the ability to commit credibly under uncertainty. Organisations that stabilise decision legitimacy are more likely to act earlier and with clearer accountability, while those that rely on hierarchical intervention and retrospective justification remain slower and less consistent (Shrestha, Ben-Menahem and von Krogh, 2019).

Module 1 established the structural tension between accelerated cognition and human accountability. Module 2 showed how unresolved accountability displaces judgement into informal override. Module 3 addresses the remaining gap: the organisational form required to sustain legitimate, repeatable decision-making under AI-driven conditions (Orlikowski, 2007).

1. Introduction: Context and Case for Transformation

The context for this capstone proposal is not one of emerging experimentation with artificial intelligence, but of institutional saturation. In insurance and reinsurance finance, AI-enabled analytics start shaping forecasting, reserving, closing, capital assessment, and portfolio steering, and increasingly determine the informational basis of strategic and regulatory decisions. AI is therefore no longer simply supporting finance; it is reshaping how finance constructs organisational reality (Verhoef et al., 2021; Orlikowski, 2007).

This development is unfolding under conditions of volatility and accountability pressure. Climate uncertainty, geopolitical disruption, and regulatory scrutiny increase both the frequency and materiality of finance decisions, placing greater emphasis on forward-looking, defensible judgement rather than backward-looking reporting alone (Power, 2007). Yet a basic tension remains unresolved: if AI has improved analytical capability so significantly, why do so many organisations still struggle to commit faster and more confidently?

Mainstream AI transformation narratives have not resolved this problem. Many still frame transformation in terms of capability buildout—new tools, automation, and models—or governance expansion through committees, validation, and compliance artefacts. While these interventions may improve local efficiency, they often intensify friction at the point of commitment: approval cycles lengthen, governance layers multiply, and judgement reappears informally as override precisely where AI outputs are most influential (Shrestha, Ben-Menahem and von Krogh, 2019; Orlikowski, 2007).

The tension becomes most visible when AI ambition turns into capital allocation. Finance leaders are rarely constrained by a lack of use cases; the constraint is legitimacy. Uncertainty persists over which assumptions are binding, who owns judgement when scenarios diverge, and how accountability is sustained when expected value does not materialise. Under these conditions, organisations often oscillate between enthusiasm for pilots and caution at scale, a pattern consistent with decision-science evidence that uncertainty increases delay and defensive behaviour when judgement is not institutionally disciplined (Kahneman, Sibony and Sunstein, 2021).

This proposal is not based on a single organisational case. It is informed by repeated observations across large insurance and reinsurance transformation contexts and is therefore context-anchored rather than case-bound. In these settings, emerging AI value-chain practices make assumptions and value drivers more visible, but they do not by themselves resolve who owns judgement at the point of commitment (Author's professional artefact, 2026a).

This proposal emerged from repeated exposure to the same organisational pattern: AI capability advanced faster than decision structures, governance readiness, and institutional trust. Over time, this shifted the focus from technology adoption to the organisational conditions required for legitimate commitment under uncertainty.

The central problem addressed in this proposal is therefore structural rather than technological. Organisations continue to invest heavily in the analytical side of decision-making-data, models, simulation-while leaving authority, accountability, and legitimacy largely unchanged. This asymmetry may have been manageable when analytics functioned as episodic support. It becomes unsustainable when AI operates as a continuous, decision-shaping presence (Orlikowski, 2007; Shrestha, Ben-Menahem and von Krogh, 2019).

The case for transformation advanced here does not rest on the promise of new AI capability. It rests on the need to redesign the finance operating model so that AI-augmented decision-making becomes an institutionalised strategic capability rather than an organisational drag.

2. Integrative Synthesis: From Strategic Tension to Decision Resolution

Chapter 1 established the core of the problem: AI has become embedded in finance work, yet analytical acceleration is not matched by organisations' ability to convert insight into timely, defensible commitment. This chapter synthesises Modules 1 and 2 to argue why Module 3 must extend beyond governance and process design toward **operating model redesign**.

2.1 What Module 1 established: the tension is structural, not temporary

Module 1 diagnosed a strategic paradox in AI-intensive finance: AI accelerates forecasting, scenario generation, and anomaly detection, while accountability for consequential decisions remains human (Shrestha, Ben-Menahem and von Krogh, 2019). The tension lies between analytical acceleration and institutional responsibility rather than between humans and machines.

In insurance and reinsurance finance, this responsibility is amplified because decisions shape capital adequacy, solvency positioning, and regulatory trust (Power, 2007). Delegating decision authority to AI may increase speed, but it does not shift accountability. In practice, it makes accountability more exposed.

This tension does not disappear as models improve. It is structural because AI changes decision conditions faster than institutions adapt legitimacy mechanisms. (Orlikowski, 2007).

2.2 What Module 2 revealed: governance does not resolve accountability

Module 2 examined the most common organisational response to this tension: governance expansion. Committees, validation processes, approval thresholds, and audit trails all serve to increase oversight. Yet despite more governance, organisations still experience decision latency, escalation, and informal override.

The central insight of Module 2 is therefore that many failures of AI adoption are not failures of analytics, but failures of decision architecture. Override is not simply user error. It is a socio-technical signal that accountability has not been institutionally designed into the decision process (Orlikowski, 2007; Shrestha, Ben-Menahem and von Krogh, 2019). Module 2 responded by proposing a collaborative decision architecture intended to make judgement visible, contestable, and governable under uncertainty (Kahneman, Sibony and Sunstein, 2021).

2.3 Why Module 3 is needed: operate at the operating model level

Module 2 also exposed a decisive limit: decision architectures operate inside existing organisational forms. Formal processes do not automatically change where power sits, what is rewarded, or who is blamed when outcomes diverge (Mintzberg, 1994).

This creates a practical problem. Governance can increase the appearance of control without increasing legitimacy. Under time pressure, authority recentralises, executives intervene, and decisions move into informal channels. The organisation appears governed, but the decision system remains structurally pre-AI (Mintzberg, 1994; Power, 2007).

In the author's professional context, this is visible in the gap between **decision-grade analytics** and **decision authority**. Organisations may build structured AI value chains and impact simulation yet still struggle to commit when stakes rise. Impact simulation can make initiatives more transparent and comparable, but it does not by itself resolve who owns judgement at the point of commitment (Author's professional artefact, 2026a).

Module 3 therefore asks what kind of organisation can sustain disciplined decision-making once AI becomes a permanent cognitive participant.

The guiding proposition is simple: **decision architecture improves choices, but operating models determine whether those choices survive organisational pressure** (Mintzberg, 1994).

3. Transformational Design: Defining the AI-Native Finance Organisation

This chapter defines the AI-native finance organisation as a designable operating model for insurance and reinsurance finance. When AI becomes a permanent participant in core finance processes, decision quality depends less on producing more analytics and more on how judgement and accountability are organised (Orlikowski, 2007; Shrestha, Ben-Menahem and von Krogh, 2019).

The central claim is precise: **decision architectures do not survive organisational pressure unless they are embedded in an operating model designed for AI-native conditions**, where AI is continuous and uncertainty is persistent (Power, 2007).

A critical question runs through the chapter: what makes this design adoptable rather than merely coherent?

3.1 Design scope and prioritisation

The AI-native finance organisation is presented as a coherent design, but not all elements are equally foundational. To keep the model realistic, it distinguishes between:

- **Core foundations:** minimum operating-model conditions without which AI-augmented judgement cannot be legitimate or scalable
- **Modular extensions:** enhancements that increase maturity, scale, and automation depth, but are not required from day one

Without prioritisation, the model risks becoming an idealised blueprint that fails under time pressure, constraint, and politics - exactly when decision systems degrade into escalation and compliance theatre (Power, 2007; Mintzberg, 1994).

3.2 Reframing transformation in finance

AI transformation is often defined in terms of automation depth or tool diffusion (Verhoef *et al.*, 2021). In finance this becomes dashboards, model proliferation, and workflow automation. These may improve efficiency, but they do not necessarily change the logic of authority.

This proposal adopts a stricter definition: **transformation is defined by how decision authority is constructed and legitimised** (Orlikowski, 2007). An

organisation can therefore be highly automated yet remain structurally pre-AI if authority recentralises under pressure and legitimacy is built retrospectively rather than before commitment (Power, 2007).

3.3 AI-enabled versus AI-native

The distinction is organisational, not technological.

- **AI-enabled finance** uses AI to support existing structures. Judgement remains exceptional and accountability remains hierarchical.
- **AI-native finance** assumes AI continuously shapes decision-work and therefore designs judgement, accountability, and legitimacy into everyday processes.

Without institutionalised judgement, it reappears as override, delay, or shadow decision-making (Shrestha, Ben-Menahem and von Krogh, 2019). Without disciplined interpretation, uncertainty amplifies noise and defensiveness (Kahneman, Sibony and Sunstein, 2021).

3.4 Operating-model components

The AI-native finance organisation is not a list of best practices. It is a minimal system designed to prevent three predictable failure modes in AI-intensive finance: **late escalation, informal override, and governance congestion** (Orlikowski, 2007).

The core question is practical: what must exist for AI-augmented judgement to remain legitimate under pressure? (Power, 2007)

Table 1. AI-native finance operating model components (insurance and reinsurance)

Operating model component	Core foundation (minimum viable)	Modular extension (scale/maturity)	Why it matters (in one sentence)
Decision rights & accountability	Clear “who decides / who stewards” for high-stakes decisions (e.g., close adjustments, reserve movements, capital impacts).	Multi-tier decision rights by materiality, reversibility, and regulatory exposure.	Accountability remains human; without explicit stewardship, judgement migrates to informal override (Simon, 1957; Shrestha et al., 2019).
Workflow design	Named “judgement points” in	Semi-automated/agentive workflows with controlled	Judgement must be designed into the flow of work, not left

(human-AI handoffs)	workflows: interpret, challenge, escalate, commit.	handoffs and exception routing.	for emergencies (Kahneman et al., 2021).
Legitimacy mechanisms (decision-grade transparency)	Pre-commitment transparency: explicit assumptions, scenarios, and trade-offs for material decisions.	Continuous impact traceability (e.g., digital twin / A–B testing of assumptions).	Legitimacy produced before commitment reduces defensive escalation and retrospective theatre (Power, 2007).
Governance (enablement, not blockage)	Lightweight guardrails: escalation thresholds, logging of rationale, and review of recurring overrides.	Formal councils for model risk, ethics, and cross-functional adjudication at scale.	Over-governance creates bottlenecks; under-governance creates hidden workarounds (Power, 2007).
Data and model stewardship	Clear ownership of “decision-critical data”: definitions, lineage, and basic monitoring.	Full MLOps/ModelOps: drift monitoring, model cards, controlled releases.	Trust depends on traceability; without stewardship, AI outputs become contestable in unproductive ways (Floridi and Cowls, 2019).
Performance system (what gets rewarded)	Minimum: evaluation includes decision quality (clarity of assumptions, rationale quality), not only outcomes.	Balanced metrics: decision velocity <i>and</i> legitimacy, learning rate, and override patterns.	What is measured shapes behaviour; activity metrics alone drive low-quality AI use (Hancock and Rosen Kellerman, 2026).

People capability & culture	Role-specific guidance on “good AI use” + psychological safety for challenge.	Communities of practice, advanced judgement training, leadership enablement pathways.	Adoption is driven by clarity and safety, not generic encouragement (Edmondson, 2018; Goldstein, 2025).
--	---	---	---

In the author’s professional context, impact simulation has been designed as an ex-ante legitimacy artefact. It forces assumptions and value drivers into a decision-gate format. It also includes confidence signaling and scope disclaimers to reduce false completeness (Author’s professional artefact, 2026a).

3.5 Four design principles

The AI-native finance organisation is governed by four principles.

1. **Judgement is designed, not left to exception.**
The critical issue is whether uncertainty can be surfaced early, rather than returning later as escalation (Edmondson, 2018; Power, 2007).
2. **Accountability shifts from hierarchy to decision stewardship.**
The question is whether accountability can be distributed without becoming vague; this requires explicit decision rights and escalation thresholds (Simon, 1957; Kahneman, Sibony and Sunstein, 2021).
3. **Legitimacy is built before commitment, not after.**
The challenge is whether organisations will tolerate this transparency or revert to retrospective justification when pressure rises (Power, 2007).
4. **Governance structures judgement rather than suppressing it.**
At scale this becomes essential but introduced too early it can crowd out learning (Power, 2007).

3.6 Embedding decision architecture: the no-overlay rule

A decision architecture that is not embedded into roles, incentives, and authority will be treated as advisory and will degrade under pressure. Under time pressure, decisions revert to executive intervention, private negotiation, or retrospective justification (Mintzberg, 1994).

Embedding therefore requires four moves:

- decision gates become constitutive, not consultative
- leaders steward decision integrity rather than arbitrate correctness
- experts become co-owners of judgement rather than producers of analysis
- AI outputs remain contestable by design

In the author's practice, impact simulation has been deployed precisely to make the decision gate constitutive: shifting debate from "can we build it?" to "is it investable, under which assumptions, and with what success criteria?" (Author's professional artefact, 2026a)

3.7 Scope discipline

This proposal is intentionally bounded. It is not:

- an automation strategy,
- a governance framework,
- or a universal blueprint.

Its purpose is to address decision-intensive finance settings where accountability and uncertainty are structural (Mintzberg, 1994).

3.8 Transformational outcome

AI-native finance does not remove paradox. It makes paradox institutionally manageable. Human and algorithmic judgement remain distinct, but their interaction becomes designed rather than improvised. Authority is exercised through stewardship of transparent, contestable decision processes, enabling earlier and more defensible commitment under uncertainty (Power, 2007; Shrestha, Ben-Menahem and von Krogh, 2019).

The next chapter translates this operating-model design into an implementation sequence while confronting the most predictable failure mode: organisational reversion under pressure (Mintzberg, 1994).

4. Implementation Pathway and Value Realisation

This chapter translates the operating-model design into a practical logic for how AI investment decisions are structured in insurance and reinsurance finance. The purpose is not to reproduce a generic change program. It is to specify the **decision infrastructure**, through which AI moves from activity to capital allocation discipline: from interesting, to investable, to realised value (Shrestha, Ben-Menahem and von Krogh, 2019).

The central claim is direct: organisations rarely fail to scale AI because they cannot build models. They fail because they cannot commit legitimately under uncertainty. When commitment is fragile, governance expands, escalation becomes routine, and override becomes late correction rather than designed judgement (Power, 2007).

4.1 The core implementation problem

Across insurance and reinsurance finance, AI portfolios often show the same pattern: many ideas, many proofs of concept, some technical success, and limited scaling into core decision cycles. This is better understood as an **investment decision**

problem than a delivery problem: organisations lack a repeatable way to compare initiatives, articulate value credibly under uncertainty, and authorise sustained funding (Mintzberg, 1994).

When organisations cannot commit, two pathologies follow: indefinite delay through reassurance cycles, or false certainty through premature ROI gating. Both suppress learning and bias portfolios toward “safe” initiatives (Kahneman, Sibony and Sunstein, 2021; Kirsner, 2021).

4.2 Design principles

The implementation model is therefore defined by principles, not by project mechanics:

1. **Internal ownership is a legitimacy requirement.**
Value and decision legitimacy must be constructed internally. External actors may enable but cannot substitute for internal authority without weakening legitimacy (Power, 2007).
2. **Decision gates come before funding, not after delivery.**
Commitment is postponed until initiatives can be expressed as decision-grade value under explicit assumptions (Kahneman, Sibony and Sunstein, 2021).
3. **Override is a designed signal, not an embarrassment.**
Overrides should be treated as information about contested assumptions or unclear authority, not hidden exceptions (Orlikowski, 2007; Shrestha, Ben-Menahem and von Krogh, 2019).
4. **Implementation is capability evolution, not project completion.**
Progress depends on governance fit, people enablement, and measurement maturity rather than on rollout speed alone (Verhoef *et al.*, 2021).

4.3 Ownership and operating roles

Implementation requires explicit role design because AI reshapes accountability relations. The model therefore distinguishes four ownership domains:

- **Business decision ownership:** accountable for adoption into workflows and legitimacy of judgement when outputs conflict with practice
- **Finance value ownership:** accountable for baselines, value logic, and comparability
- **Risk/compliance stewardship:** accountable for controls and auditability without becoming a veto mechanism
- **Enablement roles:** supporting data, technology, and advisory roles that do not own decisions

The practical starting point is to nominate decision owners for a small number of material decisions and make them accountable for the outputs of the decision gate (Author's professional artefact, 2026a).

4.4 Capability evolution

Implementation progresses through staged capability development:

1. **Constrained pilots** establish baselines, assumptions, and decision logic in real workflows.
2. **Managed expansion** extends across adjacent domains while converging on shared value logic.
3. **Institutionalisation** embeds the value chain into portfolio governance, approvals, and benefits tracking.

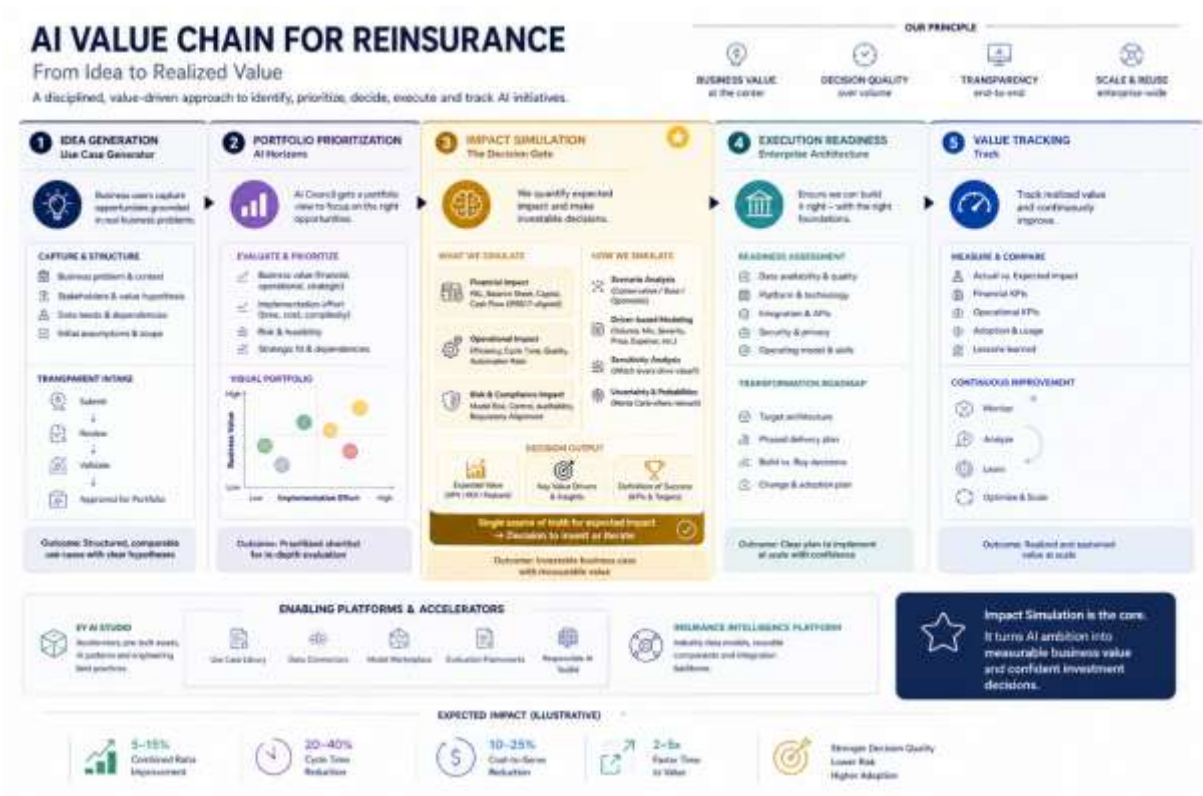
This framing is important because AI maturity in finance is primarily **commitment maturity under uncertainty**, not automation maturity (Shrestha, Ben-Menahem and von Krogh, 2019).

4.5 The AI value chain

To convert AI from initiative volume into investable decisions, implementation is organised as a five-stage value chain:

1. Idea structuring
2. Portfolio prioritisation
3. Impact simulation (**the decision gate**)
4. Execution readiness
5. Value tracking and learning

This logic applies whether mobilisation begins bottom-up, top-down, or through a hybrid model. The pathways affect how legitimacy is built; they do not remove the need for staged decision logic (Power, 2007).



4.6 Three pathways (A/B/C): a legitimacy-production choice, not a measurement choice

Organisations often treat mobilisation as a measurement choice. More accurately, it is a **legitimacy-production choice**.

- **Bottom-up** builds credibility quickly through specific cases, but comparability remains weak.
- **Top-down** creates a shared language and stronger comparability but can become framework theatre if not grounded in real decision gates.
- **Hybrid** balances credibility with comparability and is often most suitable for large organisations where readiness varies across domains. (Mintzberg, 1994; Power, 2007)

4.7 The decision gate

Impact simulation is the core mechanism because it creates what organisations repeatedly lack: a defensible articulation of value before outcomes are known. It makes assumptions, drivers, and sensitivities explicit and produces a decision-ready investment case. (Author's professional artefact, 2026a)

Its vulnerability is also clear: if authority recentralises under pressure, the gate becomes advisory rather than constitutive. This is why the value chain must be embedded in the operating model, not treated as a governance overlay (Power, 2007; Mintzberg, 1994).

4.8 Value realisation

The value chain makes one central claim: AI maturity in finance is not primarily automation maturity; it is **capital allocation maturity under uncertainty**. Early stages build comparability and legitimacy. Later stages convert that foundation into scalable returns (Shrestha, Ben-Menahem and von Krogh, 2019).

The unresolved question is whether the decision gate produces commitment or merely better documentation of hesitation. Chapter 5 turns to that question directly: **decision legitimacy in practice**.

5. Decision Legitimacy in Practice

This chapter explains decision legitimacy in practical terms and shows why it matters for senior leaders in insurance and reinsurance finance. When AI becomes a permanent participant in decision-work, the central constraint is no longer analytical insight but the organisation's ability to **commit under uncertainty without reverting to hierarchy, informal override, or retrospective justification** (Power, 2007; Shrestha, Ben-Menahem and von Krogh, 2019).

5.1 Definition

For senior finance leaders, decision legitimacy can be stated simply:

Decision legitimacy is the confidence that a material decision can be defended before it is made, while it is being challenged, and after outcomes are known (Power, 2007).

In practice, legitimacy determines whether AI-supported judgement can be accepted, challenged, and acted on under uncertainty (Orlikowski, 2007).

5.2 Why decision legitimacy matters to senior leaders

Decision legitimacy matters because it shapes how risk, authority, and accountability are experienced in practice.

First, weak legitimacy increases escalation. Decisions migrate upward when individuals seek protection from blame, and escalation becomes a substitute for clarity (Mintzberg, 1994; Power, 2007).

Second, stronger legitimacy can reduce delay without eliminating uncertainty. Decisions do not become easier, but assumptions and trade-offs become clearer earlier, reducing circular debate and repeated intervention (Kahneman, Sibony and Sunstein, 2021).

Third, legitimacy matters because regulated finance is judged retrospectively through the lens of defensibility. The relevant question is rarely whether the outcome was favourable, but whether the decision could be defended given what was known at the time (Power, 2007).

Fourth, legitimacy reduces the need for narrative repair. Where assumptions remain implicit, organisations tend to rewrite explanations after outcomes diverge. This undermines trust and discourages challenge. Legitimate decision processes make disagreement and uncertainty more discussable before commitment, not only after it.

5.3 Where legitimacy becomes most contested

Decision legitimacy is most contested where **uncertainty, materiality, and accountability** coincide. In insurance and reinsurance finance, that includes:

- reserving and reserve development, where assumptions and methods remain contestable,
- capital allocation and solvency-relevant choices, where trade-offs between growth, volatility, and capital intensity must be owned,
- AI investment decisions, where projected value depends on uncertain adoption, performance, and behavioural change.

These are not domains in which better analytics alone resolve disagreement. The issue is which assumptions can be accepted, by whom, and under what authority.

5.4 How legitimacy breaks down

Legitimacy typically weakens in three ways.

Retrospective justification replaces ex-ante reasoning.

Decision artefacts are produced to defend outcomes rather than to structure commitment under uncertainty. Governance then becomes performative rather than confidence-producing (Power, 2007).

Hierarchy substitutes for decision clarity.

Where judgement ownership is not designed, leaders become the default legitimacy mechanism. Authority compensates for missing structure, reducing scalability and increasing dependency on intervention (Mintzberg, 1994).

AI becomes either unchallengeable or ignorable.

When AI is treated as authoritative without contestability, judgement disengages. When it is treated as optional without accountability, it remains peripheral. Both conditions destabilise legitimacy (Shrestha, Ben-Menahem and von Krogh, 2019).

In this light, informal override is not failure itself. It is the signal that legitimacy has not been produced early enough or clearly enough for judgement to remain inside the formal decision process.

5.5 How legitimacy is produced

Legitimacy is strengthened when three conditions become explicit and repeatable.

- **Assumptions are articulated before commitment.**
Not to remove uncertainty, but to make clear what is being accepted knowingly.
- **Judgement is located, not improvised.**
Decision workflows require visible points where interpretation and challenge are expected, rather than appearing only through late escalation.
- **Accountability attaches to stewardship of the decision process.**
Responsibility shifts from “who signed” to “who ensured decision integrity” - assumptions, reasoning, and trade-offs - consistent with bounded rationality under complex decisions (Simon, 1957; Kahneman, Sibony and Sunstein, 2021).

This is the practical shift from AI-enabled to AI-native decision-making: legitimacy becomes a designed organisational output rather than a by-product of hierarchy.

5.6 What change looks like in practice

If legitimacy is strengthening, it tends to appear through patterns such as:

- fewer last-minute escalations,
- decision gates used for closure rather than extended debate,
- clearer articulation of assumption ranges and trade-offs,
- reduced reliance on informal override and narrative repair.

These changes do not imply that legitimacy has been fully achieved. They suggest that judgement is becoming less dependent on informal authority and more visibly embedded in the operating model.

5.7 Why this chapter matters

This chapter makes decision legitimacy both explainable and consequential. It shows why senior leaders care, why legitimacy remains uneven rather than automatic, and how it depends on organisational design rather than good intentions alone. In doing so, it reinforces the wider thesis of the proposal: AI transformation fails less because of weak intelligence than because organisations struggle to make judgement legitimate under uncertainty (Orlikowski, 2007; Shrestha, Ben-Menahem and von Krogh, 2019).

6. Organisational Resistance and Practical Constraints

This chapter addresses a straightforward organisational reality: an AI-native finance operating model will not be adopted simply because it is analytically coherent. Resistance is not primarily resistance to AI as a technology. It is resistance to the

redistribution of authority, accountability, and exposure implied by decision stewardship, explicit judgement gates, and ex-ante legitimacy (Orlikowski, 2007; Power, 2007).

Some forms of resistance are best understood as **rational organisational defense** rather than implementation failure. They indicate where legitimacy is still produced through hierarchy, informal negotiation, or retrospective justification, and therefore where the operating model is most likely to revert under pressure.

6.1 Where resistance emerges - and why it is rational

Senior finance leadership

Resistance is strongest where AI shifts from advisory analysis to constitutive decision gates. Senior leaders are likely to resist because transparency increases exposure: assumptions become visible, comparable, and challengeable, while responsibility for decision formation is distributed beyond positional hierarchy (Power, 2007). Unless the operating model answers clearly who carries reputational and regulatory risk when outcomes diverge, leadership is likely to reassert intervention (Mintzberg, 1994).

Middle management

Middle management may resist because the model weakens brokerage power. When assumptions are explicit and judgement is embedded into workflows, influence derived from routing information and sequencing approvals is reduced. This is not simply fear of AI; it is a threat to role identity and legitimacy (Edmondson, 2018).

Control functions

Risk, Compliance, and Internal Audit are likely to resist where AI appears to increase discretion without preserving auditability. In regulated environments, legitimacy must survive external scrutiny. Without clear traceability and rationale capture, AI-native workflows are likely to be interpreted as increasing risk rather than improving decision quality (Floridi and Cows, 2019; Power, 2007).

Analysts and technical experts

Analysts and technical experts may resist because the model makes judgement explicit and therefore personally attributable. Under uncertainty, this can drive defensive behaviour: over-reliance on AI or avoidance of recommendation altogether (Kahneman, Sibony and Sunstein, 2021).

6.2 Why resistance matters

Resistance should not be treated as an obstacle to overcome through better persuasion. It provides three kinds of diagnostic information:

1. **Where legitimacy is currently produced** - through authority, procedure, or narrative control (Power, 2007).
2. **Where the model is likely to revert** - senior override, local workarounds, or governance congestion (Mintzberg, 1994).

3. **Where redesign is still incomplete** - unclear stewardship, weak traceability, low psychological safety, or misaligned incentives (Orlikowski, 2007; Edmondson, 2018).

In this sense, resistance confirms that the proposal is intervening in the real organisational locus of decision-making rather than at the surface level of tool adoption.

6.3 Practical constraints and break points

Even with support, the operating model remains vulnerable under four recurring conditions:

- **time pressure**, where authority recentralises;
- **regulatory pressure**, where tolerance for uncertainty falls;
- **data contestability**, where debate shifts from decisions to baselines;
- **identity threat**, where people move from exploration to self-protection.

These pressures are especially visible around capital-relevant decision gates, where transparent assumptions may destabilise existing authority patterns (Author's professional artefact, 2026a).

6.4 A harder critical question

A more uncomfortable implication must also be acknowledged: some organisations may consciously preserve slower, hierarchy-centered decision-making structures despite analytical inefficiency, because those structures remain more compatible with their existing legitimacy expectations. In such contexts, resistance is not a failure to modernise; it is an alternative organisational choice.

6.5 Synthesis

The AI-native finance operating model is not resisted because it is too technical. It is resisted because it makes judgement visible and redistributes authority. At that point, what is being challenged is not AI itself, but the legitimacy system through which decisions have historically been made.

Where resistance persists, organisations do not simply fail to adopt AI; they fail to redesign how legitimacy is produced.

7. Measuring Success After Twelve Months

A proposal of this kind only becomes credible if its effects become observable in organisational practice.

Throughout this proposal, transformation is understood less through AI deployment activity than through changes in how organisations make and sustain decisions under uncertainty.

7.1 What counts as real change?

A central measurement challenge is distinguishing activity from decision capability. AI programs may report increased deployment, usage, or training participation without improving decision quality or defensibility. Within this proposal, progress becomes visible when decision processes are used more explicitly and consistently for organisational commitment.

7.2 Principles for measuring success

Three principles should govern measurement.

- 1. Directional rather than absolute assessment**
Improvement should be demonstrated relative to baseline patterns, rather than through fixed universal targets.
- 2. Multi-dimensional rather than single-metric evaluation**
No single metric captures legitimacy. Operational, behavioural, financial, and governance signals must be read together (Power, 2007).
- 3. Decision-centered rather than technology-centered measurement**
What matters is not how much AI is used, but how decisions are made differently because of it.

7.3 Indicative dimensions of success

Success may become visible across four dimensions:

Dimension	What is being observed	What improvement may indicate
Operational	decision cycle time, repeated review loops	greater ability to commit under uncertainty
Behavioural	informal override, private escalation, explicit challenge	judgement moving into visible processes
Financial	realised value from initiatives, movement beyond pilots	AI becoming more than narrative or experimentation
Governance	ex-post challenge load, audit rework, retrospective explanation	stronger ex-ante legitimacy

These dimensions should be interpreted collectively, since improvement in one area may conceal deterioration elsewhere.

7.4 What would not count as success

Several patterns may masquerade as success while indicating the opposite:

- more AI deployed with no reduction in escalation or rework,

- higher usage statistics accompanied by low-quality outputs,
- faster decision cycles achieved by bypassing judgement rather than institutionalising it,
- lower recorded override because disagreement has gone underground.

These are plausible organisational outcomes and should be treated as measurement risks rather than exceptions.

7.5 Closing proposition

This chapter operationalises the capstone's core claim: transformation is not measured by tool deployment, but by whether the organisation can commit under uncertainty with greater legitimacy.

Success ultimately depends on whether organisations become more capable of making explicit, defensible decisions under uncertainty in a sustained and repeatable way.

8. Leadership and Governance: Finance Stewardship, Limits, and Boundary Conditions

This chapter examines the leadership and governance implications of the proposed operating model. AI-native finance depends not only on analytical capability but on leadership and organisational learning capable of sustaining legitimate judgement (Mintzberg, 1994; Power, 2007).

The central issue is therefore whether organisations can sustain commitment when certainty remains limited.

8.1 Professional reflection: where leadership pressure becomes visible

The author's professional reflection across large reinsurance finance transformations points to a consistent pattern: AI initiatives rarely fail because the analytics are weak. They tend to strain when organisations move from analysis to commitment.

Three recurring patterns are especially visible:

- decisions stall even when the analysis is decision-grade,
- judgement retreats into informal spaces such as side conversations or executive intervention,
- governance expands without resolving who actually owns the decision under uncertainty.

Greater difficulty emerges when organisations must translate analysis into coordinated commitment, as questions of ownership, exposure, and responsibility become more pronounced (Orlikowski, 2007; Power, 2007).

8.2 Why conventional leadership responses are insufficient

Leadership responses to AI transformation commonly take two forms: **technology-first** and **governance-first**. Both are structurally limited.

Technology-first approaches often assume that better models reduce disagreement. In practice, additional information may also increase complexity, delay, and interpretive conflict under uncertainty (Simon, 1957; Kahneman, 2011).

Governance-first approaches often rely on additional review, documentation, and assurance to strengthen organisational confidence. In regulated environments, this often produces audit readiness without decision confidence (Power, 2007).

In both cases, the issue is not to approve the right answer, but to protect the conditions under which judgement can remain visible, contestable, and supported at the point of decision.

8.3 Leadership under AI-native conditions: from authority to stewardship

In the AI-native finance organisation, leadership authority is not removed, but it is re-specified. Leaders are less valuable as final arbiters of correctness than as **stewards of decision integrity**.

AI-native finance depends on leadership support, behavioural adaptation, and organisational learning rather than technological deployment alone.

This has four practical implications:

1. **Attention shifts from tools to decision sites.**
Leadership attention should focus on high-stakes decision sites where AI materially influences organisational commitment.
2. **Challenge must be protected early.**
If disciplined uncertainty is punished, judgement will retreat into silence or compliance. Organisations also require conditions in which challenge and disagreement can remain visible without creating defensive behaviour (Edmondson, 2018).
3. **Ex-ante defensibility becomes a leadership demand.**
Leadership processes should make assumptions and trade-offs explicit before major commitments are finalised.
4. **Decision quality must be evaluated, not only outcomes.**
If organisations evaluate decisions only through outcomes, hindsight bias and defensive behaviour may weaken future judgement quality (Kahneman, 2011).

These shifts do not make leadership less important. They make leadership less heroic and more institutionally demanding.

8.4 Governance as fragile coordination, not guaranteed control

Governance remains necessary, but its stabilising role should not be overstated. Formal processes cannot remove uncertainty. At best, governance can help coordination hold long enough for decisions to be taken and defended.

Where governance expands without changing how authority operates, it tends to produce congestion and informal workarounds. In such cases, formal control increases while actual decision ownership remains unclear (Power, 2007).

Leadership and governance therefore converge not on control, but on maintaining the legitimacy of judgement under pressure.

8.5 What leadership cannot fully stabilise

Leadership does not operate in neutral conditions. Its effectiveness is constrained by organisational context.

Three constraints remain structurally difficult:

- **Organisational readiness:** where challenge is politically unsafe, leadership signals may not be enough to keep judgement visible (Edmondson, 2018).
- **Traceability limits:** if decision-critical baselines remain contested, transparency may intensify friction rather than reduce it (Floridi and Cowls, 2019).
- **Pressure conditions:** regulatory scrutiny, closing deadlines, or capital stress may trigger re-centralisation even where the operating model is formally in place (Mintzberg, 1994; Power, 2007).

These are not implementation flaws. They shape how far leadership interventions can travel in practice.

8.6 Synthesis: leadership as a site of fragility, not just responsibility

This chapter makes explicit what earlier chapters imply: the AI-native finance organisation is a way of leading under conditions where certainty, authority, and legitimacy no longer align neatly.

The broader challenge is not designing governance structures but sustaining credible organisational decision-making when visibility increases exposure and coordination becomes difficult.

Leadership therefore remains important, although its effectiveness depends heavily on organisational context, coordination, and institutional trust.

Chapter 9 moves beyond immediate implementation to examine how AI-native finance evolves over time, where its longer-term limits appear, and how its value should be understood under changing organisational conditions.

9. Strategic Evolution and Critical Self-Reflection in AI-Native Finance

This chapter considers how AI-native finance may evolve once embedded into organisational practice, focusing on the longer-term limits, dependencies, and uncertainties that remain after formal adoption. The issue is no longer whether the operating model is conceptually coherent, but whether it remains stable over time under real organisational conditions.

9.1 From transformation to organisational capability

Once the AI value chain, governance structures, and decision routines stabilise, AI-supported decision-making begins to move beyond transformation rhetoric and into organisational capability. At that point, the relevant question is no longer whether AI can generate insight, but whether the organisation can repeatedly absorb, interpret, and act on that insight in legitimate ways.

Emerging practice suggests that sustained value depends less on model performance and more on whether AI is embedded into daily decision routines. In this sense, the limiting factor in AI-native finance is organisational absorption capacity rather than technical maturity, which is a key analytical position of this work (Shrestha, Ben-Menahem and von Krogh, 2019).

9.2 What does not stabilise over time

A central insight of this work is that certain tensions do not disappear as maturity increases. Some become more pronounced.

Assumptions may remain contestable, governance may drift toward compliance rather than learning, and authority may re-centralise when stakes rise. These dynamics suggest that AI-native finance should be understood as an ongoing organisational discipline rather than a stable end state.

9.3 System-level constraints

Despite advances in governance and leadership design, the effectiveness of AI-native finance remains constrained by three systemic conditions. First, organisational maturity determines whether structured judgement processes are enacted or abandoned under pressure. Second, data quality and traceability shape whether transparency produces alignment or amplifies disagreement. Third, under conditions of regulatory or time pressure, decision authority tends to re-centralise regardless of formal design.

These constraints indicate that capability remains contingent on organisational context rather than structural design alone.

9.4 Risks of interpretation

Three risks arise if AI-native finance is misunderstood. The optimisation illusion assumes that improved processes reduce uncertainty rather than making it more visible. The control illusion equates increased measurement with increased legitimacy. The symbolism risk arises when organisations adopt AI-native structures without redistributing decision authority in practice.

All three reflect the same underlying error: treating formal structure as equivalent to organisational acceptance.

9.5 Transferability

Although grounded in insurance and reinsurance finance, the core logic is not sector-specific. Similar organisational challenges are likely to emerge in settings characterised by high uncertainty, accountability pressure, and sustained human–AI interaction.

The proposed model should therefore be interpreted as a **design logic**, not a universal solution. Its application depends on contextual factors such as institutional trust, regulatory intensity, and organisational culture (Shrestha, Ben-Menahem and von Krogh, 2019).

9.6 Strategic synthesis

AI-native finance should be understood as an evolving organisational response to the interaction between increasing analytical capability and persistent uncertainty. While it enables more structured and earlier decision-making, it simultaneously increases dependence on organisational conditions that remain difficult to stabilise over time.

The central risk is therefore not insufficient technological capability, but overestimating the organisation's capacity to absorb, interpret, and act on AI-generated insight.

10. Conclusion and Outlook: Redefining Beyond Finance and Reinsurance

AI-native finance represents less technological upgrade than an organisational redesign centered on judgement and decision legitimacy.

This proposal demonstrates how organisations adapt decision-making structures as AI becomes embedded in insurance and reinsurance finance operations.

The analysis demonstrates that the central challenge is organisational: how institutions sustain credible decision processes under conditions of uncertainty, accountability pressure, and increasing analytical complexity.

10.1 Conceptual contribution

The proposal reframes finance transformation as a question of organisational judgement and governance rather than technological deployment alone.

Drawing on decision science (Simon, 1957; Kahneman, 2011), socio-technical theory (Orlikowski, 2007), and AI decision research (Shrestha, Ben-Menahem and von Krogh, 2019), it highlights a recurring tension: analytical capability may increase, while organisational capacity to commit does not.

Within this framework, override is reinterpreted as a socio-technical signal rather than operational failure. It indicates that judgement has not been institutionally located and that decision legitimacy has not been produced *ex ante* (Power, 2007).

10.2 Professional contribution

From a practitioner perspective, the proposal redirects attention from tool deployment toward decision infrastructure. Leadership is repositioned from directing outcomes to stewarding decision integrity under uncertainty, while expertise shifts from analytical production to accountable interpretation.

The contribution lies in specifying the organisational conditions under which AI-augmented decision-making becomes credible, contestable, and scalable.

The proposal offers a design logic for institutionalising judgement, rather than a framework for implementing technology.

10.3 Narrative synthesis

Taken together, the analysis shows that AI transformation is not an execution problem but a coordination problem.

AI transformation develops unevenly in practice and depends heavily on organisational coordination, governance alignment, and institutional context.

This progression reflects how organisational transformation unfolds in practice: not as linear execution, but as ongoing negotiation.

10.4 Transferability and scope

The argument extends beyond finance in principle but remains context dependent. Organisations with high accountability exposure and decision complexity are more likely to encounter the tensions described.

The model should therefore be interpreted as a **bounded analytical framework**, not a universal blueprint.

10.5 Limitations

The proposal operates as a design-oriented intervention rather than an empirically validated model.

Its effectiveness is likely to vary across organisational settings and implementation conditions. This reflects its positioning as a design-oriented intervention rather than an *ex post* evaluation.

10.6 Outlook & Closing synthesis

This proposal surfaces a persistent organisational problem: how decision quality and accountability can be sustained under prolonged technological acceleration and regulatory pressure. It positions the long-term effectiveness of AI-supported decision-making as contingent on organisational participation, leadership engagement, and governance maturity.

The future of AI-native finance therefore depends less on intelligence alone than on the organisational capacity to sustain credible judgement.

This tension remains structurally unresolved, indicating a critical frontier for understanding how decision legitimacy can be stabilised over time in AI-intensive organisational environments.

Appendices

Appendix A: Summary of Module 1 - Strategic Tension Between Artificial Intelligence and Human Judgement

Purpose

Module 1 establishes the foundational problem of the proposal: a structural tension between AI-driven analytical acceleration and irreducibly human accountability in finance decision-making. This diagnosis provides the conceptual foundation for the organisational redesign advanced in Module 3.

Core diagnosis

AI increases the speed, scale, and apparent objectivity of analysis, but accountability for consequential decisions remains irreducibly human and institutionally enforced (Shrestha, Ben-Menahem and von Krogh, 2019). This creates a structural governance tension in regulated finance environments such as reinsurance, where decisions directly shape capital adequacy, reserve credibility, and regulatory trust (Power, 2007).

While AI expands analytical capability, responsibility for consequential decisions continues to remain organisationally and institutionally human.

Structural nature of the tension

This tension is structural rather than temporary. AI expands analytical capacity but does not absorb accountability, producing systematic misalignments between technological capability and organisational decision structures designed for slower, human-centered environments (Orlikowski, 2007).

Strategic implication

The central challenge is not technical capability, but the legitimisation of judgement under uncertainty. AI transformation therefore becomes a problem of organising coordination and decision authority, defining the analytical boundary for the operating model redesign advanced in Module 3.



Appendix B: Summary of Module 2 - Collaborative Decision Architecture

Purpose

Module 2 builds on the structural tension identified in Module 1 by analysing how organisations attempt to manage it in practice. It shows that governance expansion does not resolve accountability but displaces it into informal override, revealing decision architecture as the critical intermediary between strategy and the operating model redesign developed in Module 3.

From tension to decision failure

Organisations typically respond to AI complexity through governance expansion, increasing formal oversight through approval layers and documentation.

However, decision problems emerge informally through informal override, selective use of AI outputs, and delayed commitment.

These behaviors reveal that governance does not resolve accountability but displaces unresolved judgement into informal decision spaces.

Override as a socio-technical signal

Module 2 reinterprets override as a signal of unresolved accountability rather than technical error.

It arises when organisations fail to specify judgement authority, escalation conditions, and responsibility allocation under contested AI outputs (Orlikowski, 2007; Shrestha, Ben-Menahem and von Krogh, 2019).

In this sense, override does not represent deviation from the system, but reveals how accountability is displaced into informal decision spaces when it is not institutionally designed.

Collaborative decision architecture

Module 2 proposes a collaborative decision architecture that structures decision-making as a governed lifecycle rather than an ad hoc process. Crucially, it treats judgement as an explicit design element rather than a late-stage override. In doing so, it addresses the core failure identified earlier: where accountability is not institutionally specified, it reappears informally and disrupts decision coherence.

In practice, mechanisms such as impact simulation translate AI outputs into decision-relevant constructs, shifting discussion from technical feasibility toward assumptions, trade-offs, and accountable judgement (Author's professional artefact, 2026a).

Structural limitation

However, decision architecture alone is insufficient, as it remains embedded within existing organisational hierarchies and incentive structures. Even well-designed processes are vulnerable to executive intervention and legacy authority under pressure (Mintzberg, 1994).

Decision discipline therefore remains fragile without operating model redesign, as accountability ultimately reverts to existing power structures rather than the designed decision process.



Cumulative role of Appendices A and B

Together, these modules establish cumulative analytical progression.

Module 1 identifies the structural tension between AI-driven cognition and human accountability.

Module 2 demonstrates how this tension manifests in organisational decision processes and how governance alone tends to displace rather than resolve it.

Module 3 responds by proposing an organisational redesign that explicitly integrates AI and human judgement within a unified decision system.

Bibliography

Academic and Scholarly Sources

Edmondson, A.C. (2018) *The Fearless Organization: Creating Psychological Safety in the Workplace for Learning, Innovation, and Growth*. Hoboken, NJ: John Wiley & Sons.

Floridi, L. and Cowls, J. (2019) 'A unified framework of five principles for AI in society', *Harvard Data Science Review*, 1(1).
<https://doi.org/10.1162/99608f92.8cd550d1>

Kahneman, D. (2011) *Thinking, Fast and Slow*. London: Penguin Books.

Kahneman, D., Sibony, O. and Sunstein, C.R. (2021) *Noise: A Flaw in Human Judgment*. London: William Collins.

Mintzberg, H. (1994) *The Rise and Fall of Strategic Planning*. New York: Free Press.

Orlikowski, W.J. (2007) 'Sociomaterial practices: Exploring technology at work', *Organization Studies*, 28(9), pp. 1435–1448.
<https://doi.org/10.1177/0170840607081138>

Power, M. (2007) *Organised Uncertainty: Designing a World of Risk Management*. Oxford: Oxford University Press.

Shrestha, Y.R., Ben-Menahem, S.M. and von Krogh, G. (2019) 'Organizational decision-making structures in the age of artificial intelligence', *California Management Review*, 61(4), pp. 66–83. <https://doi.org/10.1177/0008125619862257>

Simon, H.A. (1957) *Administrative Behavior*. 2nd edn. New York: Macmillan.

Verhoef, P.C., Broekhuizen, T., Bart, Y., Bhattacharya, A., Dong, J.Q., Fabian, N. and Haenlein, M. (2021) 'Digital transformation: A multidisciplinary reflection and research agenda', *Journal of Business Research*, 122, pp. 889–901.
<https://doi.org/10.1016/j.jbusres.2019.09.022>

Practitioner and Grey Literature

Goldstein, D. (2025) 'How organizations can adopt employee-centered AI', *McKinsey Quarterly*. Available at: <https://www.mckinsey.com> (Accessed: 25 April 2026).

Hancock, J. and Rosen Kellerman, G. (2026) 'The productivity cost of AI', *Harvard Business Review*, February, pp. 2–8.

Kirsner, S. (2021) 'Don't let financial metrics prematurely stifle innovation', *Harvard Business Review*. Available at: <https://hbr.org> (Accessed: 25 April 2026).

Professional Practice Artefacts

Author's professional artefact (2026a) *AI Impact Simulation Framework: Decision-Gate Design*. Internal working material.