

*This paper is shared under CC BY-ND license*

©2026 Bharath Yadla. This work is available under a Creative Commons Attribution-NoDerivatives 4.0 International License (CC BY-ND 4.0). You are free to share, copy, and redistribute the material in any medium or format, provided you give appropriate credit, do not distribute modified or derivative versions, and do not use it for commercial purposes without explicit permission. To view a copy of this license, visit <http://creativecommons.org/licenses/by-nd/4.0/>

## Becoming an Agentic Enterprise: A Practitioner Methodology & Frameworks for Human-AI Governance

Author & Presenter: Bharath Yadla

*Oxford Executive Diploma in AI for Business, Saïd Business School  
IE Business School | Forbes & Fast Company Contributor | Workato*

### Abstract

The transition from pilot projects to enterprise-scale autonomous AI systems represents one of the most significant governance challenges facing organizations today. While much attention has focused on the technical capabilities of large language models and agentic systems, comparatively little research has addressed the practical governance frameworks and methodology required to manage portfolios of agents at scale within an Enterprise. This creates a governance gap that keeps enterprises in long pilot cycles. For enterprises to confidently deploy, they need governing frameworks: Boards lack comprehensive frameworks to assess, classify, and oversee these increasingly sophisticated human-algorithm systems. Transformation leaders need a way to assess the human oversight required as agent complexity evolves.

This paper introduces the “**Becoming an Agentic Enterprise**” methodology developed by the author (s) as a comprehensive practitioner framework, grounded in direct engagement with enterprise transformation initiatives. The methodology comprises four interlocking components: **SCORE-AI** for board-level governance assessment, the **L1-L6 Agent Maturity Model** for capability classification, the **AURA** framework for human oversight calibration & the extent of integration required with existing systems and processes, and the **RRR** transformation sequence for phased implementation. We illustrate the application of this integrated methodology through a detailed case study of an enterprise software company that uses this framework, resulting in a blueprint for ~100 agents across eight business domains spanning the spectrum of agent maturity. The paper contributes to the human-AI interaction literature by demonstrating how layered governance structures can maintain meaningful human oversight while enabling operational autonomy, and by providing a replicable blueprint for enterprise-scale agentic transformation.

This paper addresses the HAI Workshop 2026 theme of 'Shaping the future of AI: innovation, ethics and impact' by:

- (1) Innovation: Proposing novel frameworks (SCORE-AI, L1-L6, AURA) for classifying and governing agentic AI systems across enterprises
- (2) Ethics: Embedding oversight through the SCORE-AI Ethics pillar and human oversight calibration at the agent level
- (3) AI Impact: Offering a practitioner-developed methodology for boards navigating AI governance.

### **Limitations and Future Research:**

This paper proposes a conceptual framework and illustrates its application in the author's current workplace. The frameworks emerge from practitioner observation, industry engagement, and synthesis of existing literature rather than systematic data collection. This paper presents the academic and industry communities with a framework for empirical studies and for further evolution as indicated in section 6.4 below.

**Keywords:** agentic AI, enterprise methodology, human-AI collaboration, agentic applications, organizational transformation, AI maturity models, Agentic Governance, algorithmic accountability, value-sensitive design, governance as strategy.

## **1. Introduction**

Let me start by acknowledging a fundamental tension in contemporary AI discourse. On one side, we have extraordinary technical capabilities emerging from frontier AI laboratories. On the other hand, we have enterprises struggling to move beyond isolated pilot projects toward meaningful operational deployment. In my experience working with Organizations across industries, this gap between capability and deployment represents the central challenge of our moment.

The challenge is not primarily technical. Most enterprises have access to the same foundation models, the same cloud infrastructure, the same integration platforms. What distinguishes Organizations that successfully deploy autonomous AI systems from those trapped in perpetual pilot programs is governance clarity. They have resolved the fundamental question that paralyzes others: How do we maintain meaningful human oversight over systems designed to operate autonomously?

### **The Three Agentic Paradoxes**

Despite massive investment and widespread AI adoption, enterprise deployments consistently reveal three structural paradoxes that this methodology is designed to address.

**The Board Paradox:** Boards hold fiduciary responsibility for AI-related outcomes but lack validated instruments to exercise that responsibility. Traditional metrics, EBITDA, and efficiency measures cannot account for the compounding, probabilistic nature of agentic intelligence. Boards approve AI investment without visibility into how agentic capability affects enterprise performance.

**The Deployment Paradox:** Pilots multiply; P&L impact stays minimal. 60% of AI projects never leave the pilot phase (ITPI, 2025); 95% of GenAI pilots deliver no measurable financial impact (MIT NANDA Study, 2025). The pressure to demonstrate AI usage has created a theatre of pilots that never reach operational scale.

**The Builder Paradox:** 68% of agents are rebuilt within six months due to misaligned human-AI handoffs, absent audit trails, and no trust calibration (Gartner, 2025). Agents are built in isolation; governance is an afterthought. The root cause is structural: AI is stochastic, and business is deterministic. Agents decide at machine speed; business demands accountability, repeatability, and audit trails. These two worlds clash — and ungoverned AI compounds the collision with each new deployment.

*The problem isn't intelligence. It's governance. And the governance gap is widening with every pilot that fails to reach production.*

This paper presents a practitioner framework developed through direct engagement with enterprise transformation initiatives. We call this “Becoming An Agentic Enterprise” methodology, comprising four interlocking components designed to address governance across multiple organizational levels simultaneously. The methodology emerged not from theoretical abstraction but from the practical necessity of helping Organizations navigate the transition from experimental AI to operational AI.

The contribution of this work lies in its integrated approach. Existing frameworks tend to address individual dimensions of AI governance in isolation — technical capability assessment, ethical oversight, or implementation sequencing. The methodology presented here treats these as interdependent facets of a unified governance challenge, providing Organizations with a coherent approach to managing portfolios of autonomous agents at scale.

The remainder of this paper is organized as follows. Section 2 positions the methodology within existing governance literature. Section 3 presents the four component frameworks in detail. Section 4 illustrates application through a comprehensive case study of enterprise deployment. Section 5 examines a specific implementation of an autonomous agent to demonstrate governance principles in practice. Section 6 discusses implications and limitations.

## 2. Theoretical Positioning

The governance of autonomous AI systems has received substantial attention in recent literature, yet significant gaps remain between theoretical frameworks and practical implementation requirements. Hu's four-level AI transformation framework provides foundational concepts for understanding organizational AI maturity, emphasizing the progression from task-specific applications to enterprise-wide integration (Hu, 2024). This work builds upon and extends Hu's framework by providing operational governance mechanisms at each maturity level.

The challenge of human oversight in autonomous systems has been examined from multiple perspectives. Technical approaches focus on interpretability, constraint satisfaction, and fail-safe mechanisms. Organizational approaches emphasize governance structures, accountability frameworks, and the allocation of decision rights. The methodology presented here bridges these perspectives by treating human oversight as a design parameter that varies appropriately with agent capability and operational context.

A key insight from practitioner observation is that governance failures in AI deployments rarely stem from inadequate technical controls. More commonly, it results from misalignment between an agent's capability level and its oversight regime. Immature agents given excessive autonomy create operational risk. Mature agents subjected to excessive oversight create organizational friction and fail to deliver anticipated value. Generally, governance in enterprises is implemented as an afterthought, primarily through a GRC lens to address compliance checklists. Given the AI operational & decision risk, AI Governance, on the contrary, has to be the strategy, not a brake. The frameworks presented below are designed specifically to calibrate oversight to capability.

The transformation literature provides additional context for understanding enterprise AI adoption. The three-phase transformation sequence presented below draws on established change management principles while adapting them to the specific characteristics of agentic systems, where the deployment target is not merely new technology but new organizational capabilities that operate with varying degrees of autonomy.

## 2.1 Where This Work Sits in the Existing Landscape

The following table maps the most directly adjacent bodies of work to the BAE methodology, identifying the specific practitioner extension each enables. A central observation is that the academic literature has articulated the governance principles - traceability, value-sensitive design, oversight calibration to automation levels- that the BAE methodology operationalizes. The gap between articulation and implementation is precisely where the practitioner's contribution lies.

Table A: Adjacent Literature Landscape and BAE Practitioner Extensions

| Adjacent Literature                                                                                                                               | Core Insight                                                                                                                                                                  | Our Practitioner Extension                                                                                                                                                                                                                                                                                |
|---------------------------------------------------------------------------------------------------------------------------------------------------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| Cummings, M. (2014). Man versus Machine: The Future of Air Traffic Control. IEEE Intelligent Systems, 29(3).                                      | Human role should shift from passive oversight to active collaboration at high Levels of Automation (LOAs). Uncertainty demands human presence, not removal.                  | L1-L6 spectrum and oversight scale operationalises LOAs as an enterprise governance taxonomy, precisely indicating that higher agent maturity (L6) demands greater governance sophistication, the inverse of the naïve assumption that more autonomy requires less oversight.                             |
| Wieringa, M. (2020). What to account for when accounting for algorithms. ACM FAccT Conference Proceedings.                                        | Accountability requires traceable chains from decision to outcome. Algorithmic systems sever traditional accountability chains; new structures must reconstruct traceability. | SCORE-AI provides structured accountability traceability for board confidence. The Awareness dimension of AURA directly operationalises Wieringa's traceability requirement at the individual agent level.                                                                                                |
| Friedman, B., Kahn, P., & Borning, A. (2008). Value sensitive design and information systems. In The Handbook of Information and Computer Ethics. | Value must be embedded in design, not retrofitted post-deployment. Ethical consideration that arrives after architecture decisions is too late.                               | RRR institutionalizes this principle: governance architecture is a Relmagine-phase output, not a Relnforce-phase correction. Sequence matters as much as framework content; governance before capability is a design rule, not a preference.                                                              |
| IMDA (2026). Model AI Governance Framework for Agentic AI. Infocomm Media Development Authority, Singapore Government. May 2026.                  | L0-L5 autonomy levels defined by action-space and degree of human oversight. Makes humans meaningfully accountable; assesses and bounds risks upfront before deployment.      | BAE L1-L6 model was developed prior to IMDA 2026 and extends it: board-level SCORE-AI provides the organizational readiness context IMDA presupposes but does not supply. AURA operationalizes IMDA's principle of human accountability at the individual agent level with a scored, deployable protocol. |

Source: Author synthesis. Literature selected for direct adjacency to BAE governance challenges. Core insights are quoted from the respective works; practitioner extensions are the author's operationalizations.

## 2.2 The Governance Lens: From GRC to Strategic

Traditional governance frameworks were built for control. The dominant enterprise AI governance paradigm, derived from Governance, Risk and Compliance (GRC) frameworks, carries four characteristic limitations when applied to agentic systems: a checkbox mindset treating compliance as destination rather than floor; slow sequential approvals incompatible with agent deployment velocity; risk aversion over value creation, consistently converting governance into a deployment veto; and compliance as the end goal rather than a precondition for competitive value.

The BAE methodology introduces an additional lens: governance as a competitive advantage. Rather than replacing GRC governance, the Strategic Lens adds the dimension that GRC frameworks cannot provide. Where GRC asks 'are we compliant?', the Strategic Lens asks 'does our governance enable the autonomous operations our competitors cannot yet safely deploy?' This shift, from compliance as destination to governance as strategy, is the foundational reorientation that the four BAE frameworks operationalize.

### 3. The Becoming An Agentic Enterprise Methodology

The methodology comprises four interlocking frameworks, each addressing a distinct governance dimension while maintaining coherence with the others. We present these in sequence in this paper, though in practice they can operate concurrently and interdependently.

#### 3.1 SCORE-AI: Board-Level AI Governance

SCORE-AI provides a structured assessment framework for board-level evaluation of organizational AI readiness and governance maturity. The framework emerged from the observation that board discussions of AI initiatives frequently lack the structured vocabulary needed for effective oversight. Directors find themselves either deferring entirely to management or focusing on narrow technical details that obscure strategic implications.

SCORE-AI directly addresses the Board Paradox identified in Section 1: boards need an instrument panel for agentic AI decisions. It operationalizes Wieringa's (2020) accountability traceability requirement at the organizational level, and provides the board-level governance infrastructure that the IMDA Agentic AI Governance Framework (2026) presupposes but does not supply.

The framework operationalizes five governance dimensions, each weighted to reflect its relative importance in enterprise AI success:

**Strategy Alignment (30%):** This dimension assesses the degree to which AI initiatives align with the articulated business strategy. The weighting reflects practitioner observation that strategy-capability alignment is the strongest predictor of successful enterprise AI deployment. Organizations that deploy AI as a solution-seeking approach consistently struggle to demonstrate value, while those that begin with strategic problems and evaluate AI as a potential solution succeed.

Capability Readiness (25%): This dimension evaluates the organization's technical and operational readiness for AI deployment, including data infrastructure, integration architecture, and talent availability. The assessment explicitly avoids treating capability as a binary threshold and instead measures readiness across multiple sub-dimensions.

Operational Integration (20%): This dimension examines how deeply AI capabilities integrate into existing business processes, systems & data. Surface-level integration, where AI provides recommendations that humans must manually implement, delivers fundamentally different value than deep integration, where AI recommendations trigger automated downstream actions.

Risk Management (15%): This dimension assesses the maturity of AI-specific risk management practices, including model monitoring, bias detection, performance degradation identification, and incident response protocols. The weighting reflects the observation that while risk management is essential, an excessive focus on it often paralyzes deployment entirely.

Ethics and Compliance (10%): This dimension evaluates adherence to ethical guidelines and regulatory requirements. The relatively lower weighting reflects not diminished importance but rather the observation that ethics and compliance failures typically stem from inadequacies in the other dimensions rather than from ethics in isolation.

The SCORE-AI assessment produces a composite score calculated as:

$$\text{AI Readiness (Agentic Index)} = (S \times 0.30) + (C \times 0.25) + (O \times 0.20) + (R \times 0.15) + (E \times 0.10)$$

Each dimension is assessed on a 1–5 scale, where 1 represents nascent capability and 5 represents industry-leading maturity. The composite score thus ranges from 1.0 to 5.0, with scores below 2.5 indicating significant governance gaps requiring remediation before enterprise-scale deployment.

### 3.2 The L1-L6 Agent Maturity Model

While SCORE-AI assesses organizational readiness, the Agent Maturity Model classifies individual agents by capability and appropriate governance treatment. The six-level taxonomy provides a common vocabulary for discussing agent capabilities and their governance implications. It operationalizes Cummings' (2014) Levels of Automation analysis for the enterprise context: a critical, counterintuitive insight, consistent with Cummings, is that governance sophistication increases with agent maturity; higher autonomy demands more capable oversight, not less. This is how it differs from IMDA 2026, classification of the agents, where the autonomous nature of agent is argued for less human oversight.

Table 1: The L1-L6 Agent Maturity Model

| Level | Name | Capability Description | Governance Implication |
|-------|------|------------------------|------------------------|
|-------|------|------------------------|------------------------|

|    |                         |                                                                    |                                        |
|----|-------------------------|--------------------------------------------------------------------|----------------------------------------|
| L1 | Robotic                 | Rule-based execution of predefined tasks without adaptation        | Standard IT controls sufficient        |
| L2 | Task                    | Single-task automation with parameter-driven variation             | Periodic review of output quality      |
| L3 | Knowledge               | Context-aware assistance drawing on organizational knowledge bases | Human review of novel situations       |
| L4 | Process                 | End-to-end workflow automation with multi-step reasoning           | Exception-based human escalation       |
| L5 | Orchestrated            | Coordination of multiple agents toward shared objectives           | Strategic outcome monitoring           |
| L6 | Sentient/<br>Autonomous | Full autonomy with ethical reasoning and self-improvement          | Governance parity with human employees |

The maturity model serves three governance functions. First, it provides classification vocabulary that enables consistent discussion of agent capabilities across technical and non-technical stakeholders. Second, it establishes baseline governance expectations for each maturity level. Third, it creates pathways for agent advancement as capabilities improve and governance mechanisms mature.

A critical insight from practical application is that organizational distribution across maturity levels follows predictable patterns. Healthy agentic portfolios typically exhibit a bell-shaped distribution, with the majority of agents at L3 and L4, smaller populations at L2 and L5, and a limited presence at L1 and L6. Organizations with portfolios heavily skewed toward lower maturity levels are likely under-investing in agent capability. Those skewed toward higher maturity levels may be overextending governance capacity.

### 3.3 AURA: Human Oversight Calibration

The AURA framework addresses the central governance question: How much human oversight does a given agent require? Rather than treating oversight as a binary choice between full human control and full agent autonomy, AURA provides a multi-dimensional assessment that calibrates oversight to specific agent characteristics. It operationalizes the task-allocation principle articulated by Malone (2018): effective human-AI collaboration achieves an appropriate division of labor, neither maximizing human control nor agent autonomy in isolation.

This framework comprises four assessment dimensions:

Awareness (A): The degree to which an agent possesses visibility into business operations, including inputs processed, decisions made, and outputs generated. High awareness means that agent decision-making abilities have to be governed.

Understanding (U): The degree to which an agent needs to understand the knowledge, documents, and process knowledge to reason and act properly for the goal that the agent has.

This dimension distinguishes between agents that provide interpretable decision rationale and those that operate as black boxes. Interpretability is a governance requirement, not merely a technical desideratum: oversight that cannot be understood cannot be meaningfully exercised (Wieringa, 2020).

Reasoning Quality (R): Assessment of the agent's decision-making sophistication, including its ability to handle edge cases, recognize its own limitations, and escalate appropriately when facing situations beyond its competence. The higher an agent has reasoning ability or is expected to have higher reasoning abilities, the more governing it should be.

Actionability (Act): The degree to which agent outputs translate directly into organizational business actions - meaning, how many enterprise skills it possesses for execution, from pure recommendations requiring human interaction to fully autonomous implementation.

Each dimension is assessed on a 1–5 scale, producing a composite AURA score:

$$\text{AURA}\% = ((A + U + R + \text{Act}) / 20) \times 100$$

The AURA percentage directly informs governance treatment. Agents scoring below 50% require intensive human oversight with approval requirements for significant actions. Those scoring 50–75% operate under exception-based oversight. Those scoring above 75% operate under outcome-based oversight, in which humans monitor aggregate results rather than individual decisions, with kill switches used where needed.

The critical governance insight embedded in AURA is that appropriate oversight depends not on agent capability alone but on the interaction between capability and organizational trust. An L4 agent with perfect technical capability but poor interpretability may require more intensive oversight than an L3 agent with transparent reasoning. The framework makes this calibration explicit and manageable.

### 3.4 RRR: The Transformation Sequence

The final component addresses the sequencing of the transformation journey. The RRR framework, **ReImagine**, **ReWire**, **ReInforce** - provides a phased approach to enterprise AI transformation that maintains governance integrity throughout the transition. The sequence institutionalizes Friedman et al.'s (2008) value-sensitive design principle at the organizational level: governance architecture is designed before agents are deployed, not retrofitted after deployment reveals its absence.

Phase 1: ReImagine (2–3 months). This initial phase focuses on strategic clarity before technical deployment. Organizations assess their current state using SCORE-AI, define their target agentic portfolio using the L1-L6 taxonomy, and establish governance baselines using AURA. The critical output is not a technology roadmap but a governance design specifying oversight requirements for each agent category. Based on it, the existing process/workflow for a specific case within an Organization has to be completely reimagined to ensure Agentic flow/orchestration fits the goal/outcome. If skipped or compressed, this phase consistently encounters governance failures in subsequent phases.

Phase 2: ReWire (6–12 months). This phase involves actual deployment of agentic capabilities, beginning with lower-maturity agents in lower-risk domains and progressively expanding both capability and scope. Each agent deployment follows a standard sequence: deployment with intensive oversight, validation of performance and governance effectiveness, gradual reduction of oversight as trust accumulates, and eventual transition to steady-state governance appropriate to agent maturity.

Phase 3: ReInforce (ongoing). This continuous phase focuses on governance refinement, agent capability advancement, and organizational learning. Unlike traditional technology deployments with defined endpoints, agentic transformation requires ongoing governance attention. Agents advance through maturity levels, requiring adjustments to governance. New use cases emerge, requiring extension of governance frameworks. Organizational trust evolves, enabling recalibration of oversight.

The RRR sequence embeds a fundamental principle: governance before capability. Organizations are encouraged to resist the temptation to deploy advanced agents before governance frameworks can accommodate them. The short-term efficiency gains from premature deployment are consistently outweighed by the organizational disruption when governance failures occur.

## 4. Case Illustration: Enterprise Blueprint Development

To illustrate the integrated application of the methodology, the author (s) present a blueprint for an enterprise software company deploying autonomous agents across its operations. This blueprint was developed using the framework, codified as a ~8000-line Claude Skill that produces the consumable blueprint. This organization provides integration and automation technology to other enterprises, making its own agentic transformation both an operational initiative and a demonstration of capability to customers.

### 4.1 Organizational Context and SCORE-AI Assessment

The organization operates across eight functional domains: Product Engineering, Customer Success, Sales and Revenue Operations, Marketing and Growth, People Operations, Finance and Operations, IT and Security, and Partner Ecosystem Management. Prior to the transformation initiative, AI deployment was fragmented, with individual teams experimenting with various tools without centralized governance.

Initial SCORE-AI assessment revealed the following profile:

Table 2: Initial SCORE-AI Assessment

| Dimension      | Score | Key Findings                                               |
|----------------|-------|------------------------------------------------------------|
| Strategy (S)   | 4     | Clear executive mandate with board-level sponsorship       |
| Capability (C) | 4     | Strong technical infrastructure with integration expertise |
| Operations (O) | 3     | Fragmented deployment requiring consolidation              |
| Risk (R)       | 3     | Standard IT controls but limited AI-specific monitoring    |
| Ethics (E)     | 4     | Established ethics guidelines needing operationalisation   |

The composite SCORE-AI score of 3.6 indicated organizational readiness for enterprise-scale deployment with identified areas for improvement. The gap between Strategy and Capability scores (both strong) versus Operational Integration (moderate) confirmed the diagnosis of fragmented deployment requiring consolidation. The Risk Management dimension was flagged for enhancement as part of the transformation initiative.

4.2 Blueprint Architecture: ~100 Agents Across Eight Domains

The ReImagine phase produced a comprehensive blueprint specifying 100 autonomous agents distributed across the eight functional domains. Distribution was not arbitrary but reflected strategic priorities, operational complexity, and governance capacity within each domain.

Table 3: Agent Distribution by Domain

| Domain              | Agent Count | Representative Examples                                                                      |
|---------------------|-------------|----------------------------------------------------------------------------------------------|
| Product Engineering | 13          | Platform Intelligence Engine, Recipe Quality Orchestrator, Connector Development Accelerator |
| Customer Success    | 13          | Customer Health Orchestrator, Onboarding Orchestrator, Support Intelligence Agent            |
| Sales & Revenue     | 13          | Revenue Intelligence Engine, Competitive Intelligence Agent, Deal Coaching Assistant         |

|                      |    |                                                                                           |
|----------------------|----|-------------------------------------------------------------------------------------------|
| IT & Security        | 13 | SOC Intelligence Orchestrator, Identity Governance Engine, Threat Detection Agent         |
| Marketing & Growth   | 12 | Marketing Intelligence Engine, Content Personalisation Agent, Demand Generation Optimiser |
| People Operations    | 12 | Talent Intelligence Engine, Recruiting Orchestrator, Employee Experience Agent            |
| Finance & Operations | 12 | Financial Intelligence Engine, Revenue Recognition Agent, Billing Optimisation Agent      |
| Partner Ecosystem    | 12 | Partner Intelligence Engine, Co-Sell Matching Agent, Partner Onboarding Orchestrator      |

### 4.3 Maturity Level Distribution and Governance Implications

Application of the L1-L6 taxonomy to the 100-agent portfolio revealed the following distribution:

Table 4: Portfolio Maturity Distribution

| Level | Class        | Count / % | Governance Treatment                                |
|-------|--------------|-----------|-----------------------------------------------------|
| L6    | Sentient     | 9 (9%)    | Full autonomy with ethical boundary monitoring      |
| L5    | Orchestrated | 17 (17%)  | Multi-agent coordination with outcome tracking      |
| L4    | Process      | 26 (26%)  | End-to-end automation with exception escalation     |
| L3    | Knowledge    | 28 (28%)  | Context-aware assistance with human verification    |
| L2    | Task         | 17 (17%)  | Rule-based automation with periodic review          |
| L1    | Robotic      | 3 (3%)    | Basic scripted operations with standard IT controls |

The distribution approximates the healthy bell curve, with 54% of agents at L3-L4, representing the operational core of the portfolio. The relatively high proportion at L5 and L6 (26% combined) reflects the organization's technical sophistication and established governance capacity. The small L1 population indicates that basic robotic process automation was largely superseded by more capable agents.

This distribution informed governance resource allocation. The 26 L4 agents required the most significant governance investment, as exception-based oversight demands sophisticated escalation protocols and trained human reviewers. The 9 L6 agents, while requiring less frequent intervention, demanded the highest-skill governance personnel capable of evaluating ethical reasoning and strategic alignment.

## 4.4 AURA Assessment and Oversight Calibration

Each of the 100 agents received an individual AURA assessment, with scores informing specific oversight configurations. To illustrate the assessment methodology, consider the Customer Health Orchestrator, an L6 agent responsible for autonomous customer health monitoring with predictive churn prevention.

Table 5: AURA Assessment — Customer Health Orchestrator (L6)

| Dimension     | Score | Assessment Rationale                                                                      |
|---------------|-------|-------------------------------------------------------------------------------------------|
| Awareness     | 5     | Complete visibility into all inputs, decisions, and outputs through comprehensive logging |
| Understanding | 5     | All predictions include interpretable explanations with contributing factor weights       |
| Reasoning     | 5     | Demonstrated ability to recognise novel situations and escalate appropriately             |
| Actionability | 5     | Direct integration with customer success workflows for automated intervention             |

The perfect AURA score (100%) enabled outcome-based governance for this agent, with human review focused on aggregate retention metrics rather than individual predictions. Contrast this with the Build Pipeline Monitor, an L1 agent scoring 35% on AURA due to limited interpretability and narrow actionability. This agent operates under continuous human monitoring despite its technical simplicity, because its outputs lack the transparency that would enable lighter oversight.

## 5. Implementation Deep Dive: The Competitive (Market) Intelligence Agent

To demonstrate governance principles in specific implementation, the author (s) in this paper examine one agent from the portfolio in detail: the Competitive Intelligence Agent (CIA) deployed within the Sales and Revenue domain. This agent exemplifies the governance challenges and solutions characteristic of L4-class autonomous systems.

### 5.1 Agent Specification and Governance Design

The Competitive Intelligence Agent operates as a pre-meeting research assistant for sales personnel, autonomously gathering and synthesizing intelligence about prospective customers and competitive threats. The agent integrates multiple data sources including the organisation's CRM system, internal communication channels, historical conversation records, and external market intelligence.

The agent produces a structured intelligence briefing comprising six components — the ABCDE framework for competitive intelligence delivery:

- Attack Line of Thought: Strategic analysis of account context, deal history, and recommended engagement vectors derived from competitive positioning.
- Battle Cards: Detailed competitive comparison matrices for each identified competitor, including differentiation points and suggested positioning language.
- Competitive Insights and Alerts: Threat assessment with severity classification (Critical, High, Medium, Low) and strategic white space identification.
- Differentiation Framework: Mapping of organizational differentiators to specific account context, with industry-specific positioning guidance.
- Executive Digest: Condensed strategic summary with recommended actions, risk register, and suggested meeting agenda structure.
- References: Complete citation of all sources utilized, enabling human verification of intelligence claims.

## 5.2 Governance Mechanisms in Practice

The CIA agent operates at L4 maturity, indicating end-to-end process automation with exception-based human escalation. Its AURA profile reveals the specific governance configuration:

Awareness (Score: 4): All data sources accessed and intelligence synthesized are logged with timestamps, enabling full auditability of the agent's research process. The agent cannot access data sources without generating audit trails.

Understanding (Score: 4): Each intelligence claim in the briefing links to its source data, enabling human reviewers to trace conclusions to underlying evidence. The agent provides confidence scores for predictive assessments.

Reasoning (Score: 4): The agent applies explicit threat classification criteria, making its severity assessments interpretable. When evidence is ambiguous, the agent flags uncertainty rather than generating false precision.

Actionability (Score: 4): Briefings are delivered directly to sales personnel without intermediate human review, though the presence of the References section enables verification when stakes are high.

The composite AURA score of 80% places this agent in the outcome-based governance category. Human oversight focuses on aggregate metrics, whether sales personnel report that briefings are accurate and useful, rather than reviewing individual briefings. The agent's governance design includes automatic escalation triggers for specific conditions: sensitive internal intelligence is flagged prominently; Critical-severity competitive threats are routed to sales leadership; and intelligence gaps are explicitly acknowledged rather than filled with speculative content.

### 5.3 Governance Lessons from Implementation

First, appropriate governance enables rather than constrains agent value. The audit logging, source citation, and escalation protocols do not reduce the agent's usefulness. They increase organizational trust, enabling the agent to operate with greater autonomy than would otherwise be acceptable. Governance investment yields autonomy returns.

Second, governance design must anticipate failure modes. The sensitive information flagging mechanism emerged from recognition that the agent would inevitably encounter information that should not be casually disclosed. Rather than preventing access to confidential information, which would cripple the agent's intelligence-gathering capability, the governance design accommodates sensitive information with appropriate handling protocols.

Third, interpretability enables accountability without approval bottlenecks. Because each intelligence claim traces to source data, human reviewers can efficiently verify concerning outputs without reviewing routine outputs.

## 6. Discussion and Implications

### 6.1 Governance as the Strategy, not a restraint

Contrary to prevailing assumptions that governance constrains AI deployment, the framework demonstrates how thoughtful governance enables deployment that would otherwise be unacceptable. Organizations that invest in governance capacity can deploy more capable agents, grant them greater autonomy, and extract more value than Organizations that view governance as overhead to be minimized. The AURA framework makes this relationship explicit by showing how governance investment translates into reduced oversight.

### 6.2 Calibration Over Standardization

A central insight from practical application is that one-size-fits-all governance approaches fail. An L1 agent and an L6 agent require fundamentally different oversight regimes. The maturity model and AURA framework provide mechanisms for appropriate calibration, ensuring that governance resources are allocated where they deliver value rather than distributed uniformly.

### 6.3 Transformation as Ongoing Process

The RRR framework's emphasis on continuous ReInforce activity reflects recognition that agentic transformation is not a project with a defined endpoint but an ongoing organizational capability. Agents advance through maturity levels, requiring adjustments to governance. Organizational trust evolves, enabling recalibration of oversight. New use cases emerge, requiring an extension of the framework. Organizations that treat agentic transformation as a project consistently struggle when the project officially ends, but governance demands continue.

### 6.4 Limitations, Open Questions, and Future Directions

The methodology presented here emerged from practitioner observation rather than controlled empirical study. While the framework has demonstrated practical utility across multiple enterprise implementations, formal validation would strengthen its claims of general applicability. Four specific open questions emerge from implementation experience that practice alone cannot address, offered here as a research invitation to the academic community:

#1 AURA Reliability: Can the four dimensions be scored with acceptable inter-rater reliability across assessors from different organizational backgrounds? The current framework assigns equal weights across dimensions; are these empirically defensible?

#2 L1-L6 Validity: Does the six-level maturity spectrum hold across industry sectors? The current taxonomy was developed and validated primarily in enterprise software and automation contexts.

#3 Board Behavior Change: Does SCORE-AI change board decision-making in practice, or do boards reach similar governance conclusions independently? How can we empirically distinguish SCORE-AI's causal contribution?

#4 Accountability at L6: At L6, accountability formally stays with humans, but the causal chain from human decision to agent outcome is long. When does this chain length make meaningful human accountability untenable?

Additionally, the current framework addresses governance of agentic systems operating within defined organizational boundaries. As agents increasingly interact across organizational boundaries, governance frameworks will need to be extended to address cross-organizational coordination and accountability.

## 7. Conclusion

The transition to enterprise-scale autonomous AI systems represents a governance challenge at least as significant as the technical challenges of building capable agents. This paper has presented the Becoming An Agentic Enterprise methodology, an integrated framework comprising SCORE-AI for organizational readiness assessment, the L1-L6 model for agent capability classification, AURA for human oversight calibration, and RRR for transformation sequencing, with Human-AI governance.

The case illustration demonstrates that enterprise deployment of 100 autonomous agents across eight functional domains is achievable with an appropriate governance architecture. The Competitive Intelligence Agent implementation shows how governance mechanisms can be embedded in specific agent designs, enabling autonomy while maintaining accountability.

The central message of this work is that governance enables rather than constrains agentic transformation. Organizations that invest in governance capacity will deploy more capable agents, grant them greater autonomy, and extract more value than Organizations that treat

governance as overhead. As autonomous AI systems become central to enterprise operations, the frameworks presented here offer one path toward managing this transformation responsibly.

## References

- Cummings, M. (2014). Man versus machine: The future of air traffic control. *IEEE Intelligent Systems*, 29(3), 2–5.
- Friedman, B., Kahn, P. H., & Borning, A. (2008). Value-sensitive design and information systems. In K. E. Himma & H. T. Tavani (Eds.), *The Handbook of Information and Computer Ethics* (pp. 69–101). Wiley.
- Gartner. (2025). Human-AI handoff realignment study. Gartner Research.
- Hu, K. (2024). Four-level AI transformation framework. *WAIC UP Magazine*.
- IMDA. (2026). Model AI governance framework for agentic AI. Infocomm Media Development Authority, Singapore Government. May 2026.
- ITPI. (2025). Governance gap report: Enterprise AI pilot attrition. IT Policy Institute.
- Malone, T. W. (2018). *Superminds: The surprising power of people and computers thinking together*. Little, Brown and Company.
- MIT NANDA Study. (2025). GenAI pilot P&L impact analysis. MIT Sloan Center for Information Systems Research.
- Raisch, S., & Krakowski, S. (2021). Artificial intelligence and management: The automation-augmentation paradox. *Academy of Management Review*, 46(1), 192–210.
- Shrestha, Y. R., Ben-Menahem, S. M., & von Krogh, G. (2019). Organizational decision-making structures in the age of artificial intelligence. *California Management Review*, 61(4), 66–83.
- Von Krogh, G. (2018). Artificial intelligence in Organizations: New opportunities for phenomenon-based theorizing. *Academy of Management Discoveries*, 4(4), 404–409.
- Wieringa, M. (2020). What to account for when accounting for algorithms: A systematic literature review. In *Proceedings of the 2020 ACM Conference on Fairness, Accountability, and Transparency (FAccT '20)*, 1–18.
- Yeomans, M., Shah, A., Mullainathan, S., & Kleinberg, J. (2019). Making sense of recommendations. *Journal of Behavioral Decision Making*, 32(4), 403–414.

## Author Bio

Bharath Yadla is a transformation strategist and the creator of the “**Becoming an Agentic Enterprise**” methodology and creator of SCORE-AI (board-level AI governance), the Agent Maturity Model (L1-L6), AURA (agent capability assessment), and the RRR Framework (transformation pathway). He has been a 3-time entrepreneur, with 28 years of experience in GTM Strategy, Product Strategy & Corporate Strategy, growing business extensively in two distinctive growth zones - Zero To One & 10x. He has been with Workato for the past 7.5 years, in various roles on the executive team, growing the company to a \$5.7B valuation. Currently, he works as a Strategic Advisor & thought partner for the CEO as a part of the Office of the CEO at Workato. He is also associated with Oxford University, by pursuing an Executive AI Business Program at Saïd Business School, University of Oxford. His thought leadership spans 15+ articles & white papers, including Forbes & FastCo.

*This paper is shared under CC BY-ND license*

©2026 Bharath Yadla. This work is available under a Creative Commons Attribution-NoDerivatives 4.0 International License (CC BY-ND 4.0). You are free to share, copy, and redistribute the material in any medium or format, provided you give appropriate credit, do not distribute modified or derivative versions, and do not use it for commercial purposes without explicit permission. To view a copy of this license, visit <http://creativecommons.org/licenses/by-nd/4.0/>