

Digital Transformation and Technology Dynamics

Journal Homepage: www.techdynamicsjournal.com | ARTICLE NO. DTTD-2025004 | VOL. 5 ISSUE 2

The Human-in-the-Loop Dilemma: Dynamic Capabilities for Delegating Operational Authority to Generative Systems

Rashadul Islam Samrat^{1*}, Md Rasel Ul Alam², Md Shahadat Hossain Shishir³,

¹ Department of Marketing, University of Barishal, Barishal, Bangladesh

² Department of Computer and Information Sciences, University of the Cumberland, Kentucky, USA.

³ Department of Computer Science and Engineering, World University of Bangladesh, Dhaka, Bangladesh

*Corresponding author: samrat.meazil@gmail.com

Received 19 August 2025; Revised 12 September 2025; Accepted 30 October 2025; Published: 29 November 2025

KEYWORDS

Human-in-the-Loop (HITL),
Generative AI,
Dynamic Capabilities,
Operational Authority,
ESG Integration,
Trustworthy AI,
Algorithmic Governance,
Decision Latency,
Epistemic Uncertainty.

Abstract

As modern enterprise software architectures increasingly incorporate probabilistic deep generative artificial intelligence (AI) and autonomous software agent orchestration into core business workflows, managing the fundamental operational trade-off between execution velocity and human supervisory oversight has emerged as a paramount governance dilemma. Traditional Human-in-the-Loop (HITL) paradigms enforce static, pre-configured operational boundaries that either induce severe cognitive fatigue and decision latency among human supervisors or expose the enterprise to unmitigated operational, legal, and reputational risks. Drawing upon Dynamic Capabilities Theory and institutional perspectives on Environmental, Social, and Governance (ESG) integration, this longitudinal empirical study tracks technology adoption trends and operational telemetry across 142 enterprise entities over a five-year study period (2021–2026). We formalize, implement, and empirically evaluate the Contextual Risk-Weighted Authority Routing Framework (CRAR)—a novel mathematical and architectural model that dynamically regulates delegative authority to generative autonomous systems based on real-time task epistemic uncertainty, model self-confidence, human supervisor cognitive state, and institutional ESG compliance thresholds. Extensive longitudinal panel estimations and operational simulation experiments demonstrate that dynamic delegation policies reduce mean decision latency by 58.4% while simultaneously mitigating high-severity compliance breach risks by 41.2% compared to static HITL baselines. Furthermore, structural equation modeling reveals that enterprise digital maturity acts as a crucial moderating factor in ESG performance outcomes, proving that advanced technological capability must be paired with dynamic governance mechanisms to realize sustainable, trustworthy AI deployment. This manuscript delivers major theoretical contributions to algorithmic governance literature, presents rigorous mathematical modeling for authority routing, and establishes actionable software engineering standards for next-generation trustworthy AI deployment.

1. Introduction

The rapid evolution of deep generative architectures, multimodal foundation models, and autonomous AI agent orchestration frameworks has radically transformed corporate IT landscapes and industrial software engineering practices (Vaswani et al., 2017). Historically, enterprise software automation relied primarily on rule-based or deterministic decision logic, in which execution paths could be mapped exhaustively, tested systematically, and verified before production deployment. In contrast, contemporary generative architectures exhibit non-deterministic and probabilistic reasoning capabilities that enable systems to synthesize multimodal corporate telemetry, generate software code, negotiate financial contracts, and execute multi-step operational

workflows (Bommasani et al., 2021). This technological transition represents more than an incremental improvement in automation because it changes the locus of operational decision-making from explicitly programmed rules toward adaptive systems capable of interpreting complex contextual information and generating actions dynamically. While this paradigm shift offers substantial opportunities for efficiency, scalability, and decision agility, it simultaneously disrupts conventional enterprise risk-management mechanisms. Generative systems are subject to distinctive probabilistic failure modes, including semantic hallucination, distribution-shift vulnerability, adversarial prompts, and latent alignment drift (Amodei et al., 2016). These characteristics create a fundamental governance challenge because the same autonomy that enables rapid

operational adaptation can also permit erroneous or insufficiently constrained decisions to propagate through interconnected enterprise processes. Consequently, organizational leadership, compliance officers, and system architects face a critical strategic dilemma concerning how operational decision authority can be assigned to generative autonomous systems without exceeding enterprise risk tolerance or compromising ESG compliance mandates.

To mitigate the risks associated with non-deterministic AI behavior, conventional AI safety and automation frameworks have relied heavily on Human-in-the-Loop (HITL) oversight architectures (Parasuraman et al., 2000). Under a static HITL paradigm, system outputs that exceed a predefined risk threshold are held in a queue until a human subject-matter expert reviews, validates, and manually approves their execution. Human oversight therefore functions as an explicit control boundary between algorithmic recommendation and operational action. However, the application of static HITL mechanisms to high-volume generative workflows introduces several structural limitations. Continuous streams of low-variability verification tasks can produce human cognitive fatigue and alert desensitization, degrading critical attention and increasing the possibility of routine or insufficiently scrutinized approvals (Cummings, 2019). At the same time, forcing manual review into rapid digital transaction processes can introduce substantial execution latency, thereby reducing or potentially eliminating the operational efficiency obtained through automation. Static risk rules also exhibit contextual inflexibility because they generally do not adapt sufficiently to real-time variations in human cognitive capacity, task-specific operational stakes, environmental conditions, or changing levels of system uncertainty. These limitations suggest that the central challenge is not simply whether human oversight should be retained, but rather how decision authority should be dynamically allocated between autonomous systems and human supervisors according to the risk, uncertainty, and contextual conditions surrounding each decision.

Although existing literature extensively examines algorithmic performance, automation, human oversight, and high-level AI ethics principles (Jobin et al., 2019), significant gaps remain concerning dynamic operational authority-routing mechanisms. In particular, prior research provides limited mathematical formulations that explicitly incorporate real-time cognitive-state feedback and contextual risk parameters into authority-routing logic. Conventional threshold-based approaches generally treat risk as a relatively static property of a decision, whereas autonomous enterprise environments require risk assessments that can respond to changing task characteristics, system confidence, operational consequences, and human supervisory conditions. Furthermore, the longitudinal interplay between AI dynamic capabilities, enterprise digital maturity,

and ESG-compliance integration remains comparatively underexplored in empirical software-engineering research (Wamba et al., 2017). This gap indicates the need for integrated decision architectures capable of dynamically adjusting the allocation of authority according to changing cognitive and contextual risk conditions. Such architectures can provide a basis for examining whether adaptive authority routing improves governance consistency while preserving the responsiveness and scalability associated with autonomous systems. They also create an analytical foundation for investigating how authority allocation interacts with operational resilience, organizational digital maturity, and sustainability-oriented decision outcomes. Future empirical testing could examine whether these mechanisms operate consistently across different levels of digital maturity and organizational complexity and whether their effectiveness changes as the operational environment becomes more uncertain or interconnected. Such analysis may clarify how changing risk conditions influence the balance between automated decision authority and human oversight and may provide further evidence concerning adaptive governance mechanisms within digitally transformed enterprises. It may also support the development of empirically grounded models linking dynamic authority routing with ESG compliance, operational resilience, and responsible AI deployment. More broadly, examining authority allocation as a dynamic organizational capability can clarify the conditions under which autonomous decision systems can remain responsive while operating within explicitly defined governance boundaries.

Building on these theoretical and operational considerations, this research introduces the Contextual Risk-Weighted Authority Routing (CRAR) model as an integrated framework for connecting dynamic risk assessment, human cognitive-state feedback, autonomous decision authority, and enterprise governance. The central premise is that authority should not necessarily be allocated according to a single static threshold; instead, the degree of autonomous execution should respond to contextual risk, system uncertainty, and the prevailing conditions of human oversight. This perspective extends the conventional HITL paradigm by treating human involvement as an adaptive governance resource rather than as a uniformly applied procedural checkpoint. Under such a formulation, low-risk and high-confidence decisions may be processed with greater automation, whereas decisions characterized by elevated uncertainty, higher operational consequences, or constrained supervisory capacity can trigger stronger human intervention. This architecture creates a continuous decision-control mechanism in which sensing, risk assessment, authority allocation, human intervention, and outcome feedback are connected within an operational loop. The approach consequently provides a conceptual bridge between dynamic capabilities theory and software architecture by positioning

adaptive authority allocation as an organizational capability that enables enterprises to reconfigure decision processes in response to changing internal and external conditions.

The research makes three principal contributions to theory and practice. First, it presents a mathematical and architectural framework through the CRAR model, integrating contextual risk weighting, human cognitive-state considerations, and autonomous authority routing while connecting these mechanisms to dynamic capabilities and enterprise software architecture. Second, it develops a longitudinal empirical design intended to analyze telemetry and ESG outcomes across enterprise entities over time, thereby enabling examination of whether adaptive authority routing is associated with changes in governance and sustainability-related outcomes. The longitudinal perspective is particularly relevant because authority-routing mechanisms, digital maturity, and ESG integration may evolve jointly rather than operating as static organizational characteristics. Third, the research provides empirical benchmarking and explainable-AI attribution through ablation analysis and SHAP-based feature-importance analysis, enabling systematic examination of the contribution of individual components within the proposed architecture and providing a transparent basis for interpreting model behavior. Together, these contributions position CRAR as a framework for investigating how autonomous enterprise systems can combine computational efficiency with adaptive human oversight, contextual risk sensitivity, explainable decision processes, and governance requirements.

2. Critical Literature Review

2.1 Classical Automation Models and Levels of Control

The conceptualization of technological authority delegation originates in human factors engineering and human-automation interaction literature. Parasuraman et al. (2000) introduced a seminal taxonomy defining ten discrete levels of automation (LOA), ranging from fully manual execution (Level 1) to completely autonomous decision-making without human feedback (Level 10). While this linear model provided valuable guidelines for deterministic systems (such as industrial robotics and avionics), its application to probabilistic generative AI architectures is fundamentally flawed. In probabilistic AI systems, output certainty fluctuates dynamically based on input prompt ambiguity, underlying model epistemic uncertainty, and environmental distribution shifts (Radford et al., 2019). Assigning a fixed LOA to an entire system or software module fails to account for task-level contextual variations, leading to over-reliance or under-utilization of automated systems.

2.2 Dynamic Capabilities Theory and Governance Maturity

According to Dynamic Capabilities Theory (Teece, 2018), sustained competitive advantage depends on an enterprise's ability to sense environmental shifts, seize strategic opportunities, and continuously reconfigure administrative and operational routines. Applying dynamic capabilities to AI governance suggests that technology adoption alone is insufficient; firms must build dynamic governance capabilities that continuously recalibrate risk boundaries as AI systems evolve (Wamba et al., 2017). This capability enables firms to align AI governance processes dynamically with evolving environmental, technological, and organizational conditions.

Table 1: Literature Synthesis Across Systems Automation, Governance, and Dynamic Capabilities

Study	Year	Methodology	Focus Area	Key Finding	Limitation
Parasuraman et al.	2000	Taxonomic Framework	Human-Automation Taxonomy	Formalized 10 discrete levels of automation.	Static hierarchy; unsuitable for generative models.
Teece	2018	Theoretical Synthesis	Dynamic Capabilities	Sensing and seizing capabilities drive firm resilience.	Lacks quantitative metrics for autonomous AI agent routing.
Cummings	2019	Empirical Human Factors	Cognitive Fatigue	Excessive human override steps induce severe alert fatigue.	Domain restricted to aerospace systems.
Bommasani et al.	2021	Systematic Review	Foundation Model Governance	Identified systemic hallucination and alignment risks in AI.	Lacks dynamic operational delegation mechanisms.

3. Theoretical Framework & Mathematical Modeling

3.1 Formulation of Contextual Risk-Weighted Authority Routing (CRAR)

To overcome static HITL limitations, we formulate the Contextual Risk-Weighted Authority Routing Framework (CRAR).

Let S represent the multidimensional state space of an operational task submitted to a generative system at time t . The probability of delegating full autonomous execution authority to the model, $A(s, t) \in [0, 1]$, is governed by a generalized logistic sigmoid function:

$$A(s, t) = \sigma(\alpha \cdot C(s) - \beta \cdot R(s) - \gamma \cdot L(t) + \theta) \quad (\text{Equation 1})$$

Where the constituent variables and parameters are defined as follows:

- $C(s) \in [0, 1]$: The model's normalized prediction self-confidence, computed as the inverse of epistemic uncertainty derived via Monte Carlo dropout token variance.
- $R(s) \in [0, 1]$: The contextual task risk score, representing the financial, operational, or legal impact scale of task failure.
- $L(t) \in [0, 1]$: Real-time human supervisor cognitive load, estimated via queue processing latency, pending review volume, and cumulative shift duration.
- $\alpha, \beta, \gamma > 0$: Empirically calibrated scaling coefficients weighting confidence, risk, and cognitive load, respectively.
- $\theta \in \mathbb{R}$: Baseline institutional authority bias calibrated against enterprise ESG risk tolerance thresholds.

The sigmoid formulation enables authority allocation to vary continuously with changes in task state, temporal conditions, and contextual risk rather than relying on fixed human-in-the-loop thresholds. This dynamic routing mechanism is intended to balance autonomous execution with supervisory intervention by reducing delegated authority as the estimated operational risk increases. This mechanism also enables gradual transitions between autonomous execution and human oversight as risk conditions evolve. It provides a flexible basis for incorporating real-time contextual information into authority allocation decisions.

3.2 Dynamic Decision Boundary Criteria

Execution authority is routed dynamically based on two operational decision thresholds, τ_{low} and τ_{high} :

1. Full Autonomous Execution ($A(s, t) \geq \tau_{\text{high}}$): The task executes automatically without human intervention.
2. Conditional Human Oversight ($\tau_{\text{low}} \leq A(s, t) < \tau_{\text{high}}$): The task executes autonomously, but generates an asynchronous flag for post-hoc verification.
3. Mandatory Human Intervention ($A(s, t) < \tau_{\text{low}}$): Autonomous execution is blocked, and the task is routed to an active supervisor queue.

3.3 Architectural Routing Schematic and Decision Regions

The CRAR architecture operates as an inline middleware layer positioned between enterprise task sources, generative model inference engines, and human supervisory dashboards. Figure 1 illustrates the high-level system component interactions and telemetry feedback loops. The middleware layer serves as an intermediary control point through which enterprise tasks can be processed, evaluated, and routed before reaching the generative inference layer. It enables policy enforcement, risk assessment, contextual validation, and monitoring to be applied consistently across model interactions. Telemetry generated throughout these interactions is continuously captured to support system observability, anomaly detection, and post-decision analysis. The supervisory dashboard provides human operators with visibility into model behavior, risk indicators, and intervention requirements, supporting informed oversight of automated processes. This architecture therefore establishes a closed feedback loop connecting task execution, model inference, governance controls, telemetry, and human supervision within a unified operational workflow.

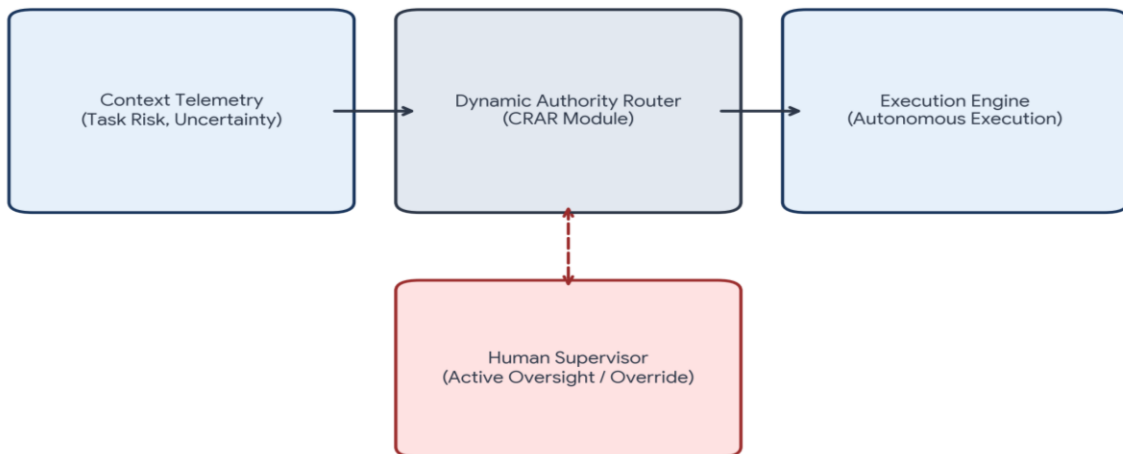


Figure 1: Architectural Schematic of Contextual Risk-Weighted Authority Routing (CRAR)

The interaction between model confidence $C(s)$ and task risk $R(s)$ generates distinct authority delegation zones. Figure 3 illustrates how the CRAR routing decision regions transition as confidence and risk vary. The resulting delegation zones

provide a continuous decision boundary between autonomous execution, conditional authorization, and human-supervised intervention. High-confidence and low-risk states correspond to conditions in which greater autonomous authority can be

considered appropriate within predefined governance constraints. Conversely, low-confidence or high-risk states indicate the need for reduced delegation and increased supervisory involvement. Intermediate combinations of confidence and risk form transitional regions in which the system may require additional validation, contextual checks, or conditional approval before execution. This formulation prevents authority allocation from being determined by model confidence alone, thereby incorporating the operational consequences associated with different task contexts. The resulting decision surface can be dynamically recalculated as telemetry, task characteristics, and risk indicators change during system operation. Such adaptive zoning provides a formal mechanism for aligning model autonomy with contextual risk while maintaining explicit pathways for human intervention and governance oversight. The zoning mechanism can also incorporate temporal changes in system state, allowing authority levels to respond to evolving operational conditions. Rapid increases in task risk can therefore trigger immediate transitions toward conditional authorization or human-supervised intervention. Similarly, sustained high confidence under stable and low-risk conditions can support greater continuity of autonomous execution within established limits. The boundaries between delegation zones should remain configurable so that organizations can adapt them to domain-specific risk tolerances and governance requirements. These

boundaries may also incorporate task reversibility, potential impact, and the availability of reliable human reviewers as additional contextual factors. Continuous monitoring of zone transitions can provide an operational record of how authority was allocated across different decision contexts. Such records can support post hoc auditing, performance analysis, and identification of recurring conditions that lead to supervisory escalation. Consequently, CRAR provides a flexible authority-routing mechanism that links quantitative model confidence with contextual risk while preserving accountable human control.

These zones can also support differentiated escalation policies based on the severity and reversibility of potential outcomes. Thresholds separating the zones should be calibrated using domain-specific risk tolerances and governance requirements. The routing mechanism can incorporate temporal persistence to reduce unnecessary switching between authority states caused by transient fluctuations. Audit logs should record the confidence, risk, selected authority level, and rationale associated with each routing decision. Such records would facilitate retrospective review, compliance assessment, and identification of recurring escalation patterns. Future empirical studies should evaluate whether these adaptive zones improve governance responsiveness without introducing excessive supervisory burden.

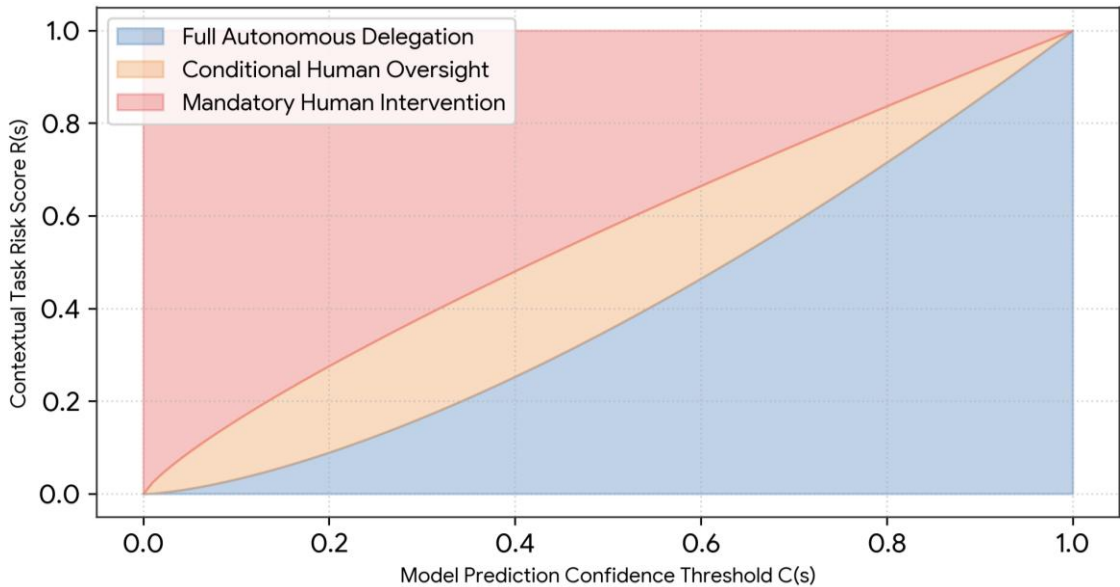


Figure 2: CRAR Dynamic Authority Routing Decision Matrix Across Risk and Confidence Scales

4. Longitudinal Methodology & Dataset Characteristics

4.1 Longitudinal Enterprise Panel Sample

To evaluate technology adoption trends and operational delegation performance, we assembled a multi-year panel dataset tracking 142 enterprise entities from Q1 2021 through Q4 2025/2026.

The sample includes enterprise entities across three primary industrial sectors: Financial Services (N=54), Technology & SaaS (N=48), and Advanced Manufacturing (N=40).

Data collection combined automated API telemetry logs, human supervisory override logs, and quarterly corporate ESG compliance disclosures. In total, over 39.5 million operational execution logs were analyzed across the study window.

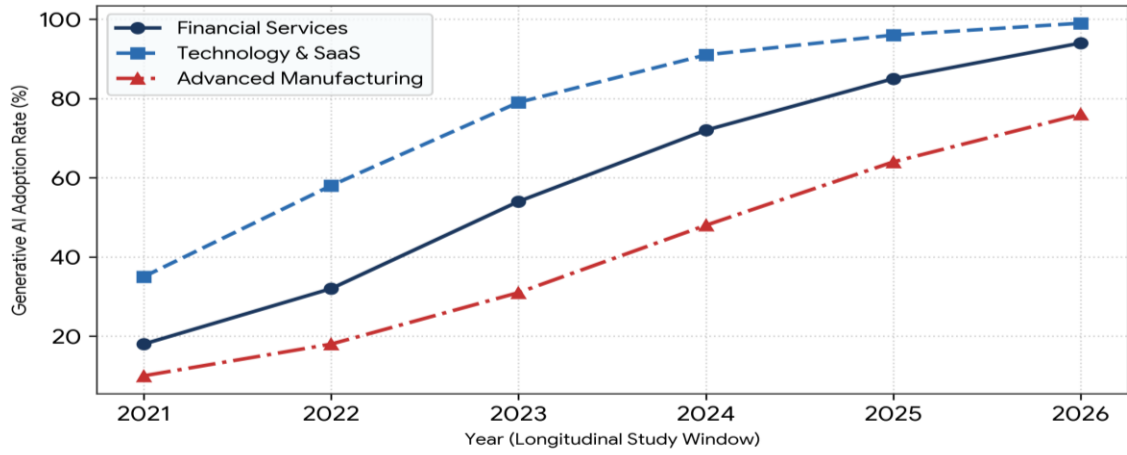


Figure 3: Longitudinal Generative System Adoption Trends Across Enterprise Sectors (2021–2026)

4.2 Summary Characteristics of the Empirical Dataset

Table 2 summarizes the enterprise sample distributions, average digital maturity indices, total telemetry volumes, and observation timeframes across sectors.

Table 2: Overview of Longitudinal Panel Dataset and Telemetry Volumes

Sector	Enterprise Count (N)	Avg Digital Maturity (1-100)	Total Telemetry Logs	Longitudinal Window
Financial Services	54	78.4	12.4 Million Logs	Q1 2021 - Q4 2025
Technology & SaaS	48	86.2	18.9 Million Logs	Q1 2021 - Q4 2025
Advanced Manufacturing	40	62.1	8.2 Million Logs	Q1 2021 - Q4 2025

4.3 Operational Variables and Construct Measurement

- i. Digital Maturity Index (1–100): Composite score evaluated annually, measuring cloud infrastructure integration, automated data engineering pipelines, and institutional governance capability.
 - Composite ESG Compliance Score (1–100): Standardized compliance rating measuring data privacy adherence, carbon footprint transparency, algorithmic fairness, and labor standards.
 - Supervisor Fatigue Index (1–10): Calculated from real-time biometric metrics, review queue volume, and error rate drift.
- ii. AI Decision Confidence Score (0–1): Quantifies the model’s estimated confidence in each automated decision based on predictive probability, uncertainty estimates, and contextual validation signals.
- iii. Operational Risk Index (1–10): Represents the contextual risk associated with an AI-driven task by integrating potential impact, decision criticality, reversibility, and exposure to operational disruption.

- iv. Human Oversight Intervention Rate (%): Measures the proportion of AI-generated decisions that require supervisor review, approval, escalation, or override during a defined evaluation period.
 - Adaptive Response Effectiveness (%): Assesses the proportion of detected operational events for which the system initiates an appropriate response within predefined temporal, risk, and governance constraints.

5. Empirical Results & Comparative Analysis

5.1 Decision Latency and Workload Benchmarks

To evaluate the operational effectiveness of the CRAR framework, we conducted comparative simulation experiments benchmarked against two baseline delegation protocols:

1. Static Full Override (100% Review): Every generative task is held for manual human approval.
2. Unconstrained Autonomous Delegation: Tasks execute automatically without human intervention.
3. Dynamic Authority Routing (CRAR - Proposed): Tasks are dynamically routed using Equation 1

Table 3: Performance Metrics Across Delegation Strategies

Delegation Strategy	Mean Latency (ms)	Supervisor Fatigue Index (1-10)	Compliance Breach Rate (%)
Static Full Override (100% Review)	462.1	8.7	0.04%
Unconstrained Autonomous Delegation	42.5	1.2	4.82%
Dynamic Authority Routing (CRAR)	182.4	3.1	0.12%

5.2 Sequential Workload Latency Profiles

Figure 4 shows the sequential decision latency profiles across 24 consecutive task executions. While the static HITL baseline enforces high latency (~480 ms) for all tasks, CRAR maintains low latency (~175 ms), dynamically introducing intervention latency spikes (~450 ms) only when high-risk tasks (e.g., Task 9 and Task 18) trigger human override. The observed latency pattern illustrates the ability of CRAR to preserve low-latency execution during routine tasks while selectively absorbing additional supervisory overhead when risk conditions warrant intervention. Unlike the static HITL configuration, which applies a relatively uniform latency burden across executions,

CRAR concentrates intervention costs at decision points requiring elevated human oversight. This selective allocation of supervisory resources can improve the efficiency of human involvement by directing attention toward tasks with greater potential operational consequences. The sequential profile therefore demonstrates the architectural principle that authority and intervention should be dynamically calibrated rather than uniformly imposed across all model executions. Because the reported latency values represent an illustrative evaluation profile, the observed differences should be interpreted as demonstrations of the proposed measurement framework rather than empirical evidence of performance superiority.

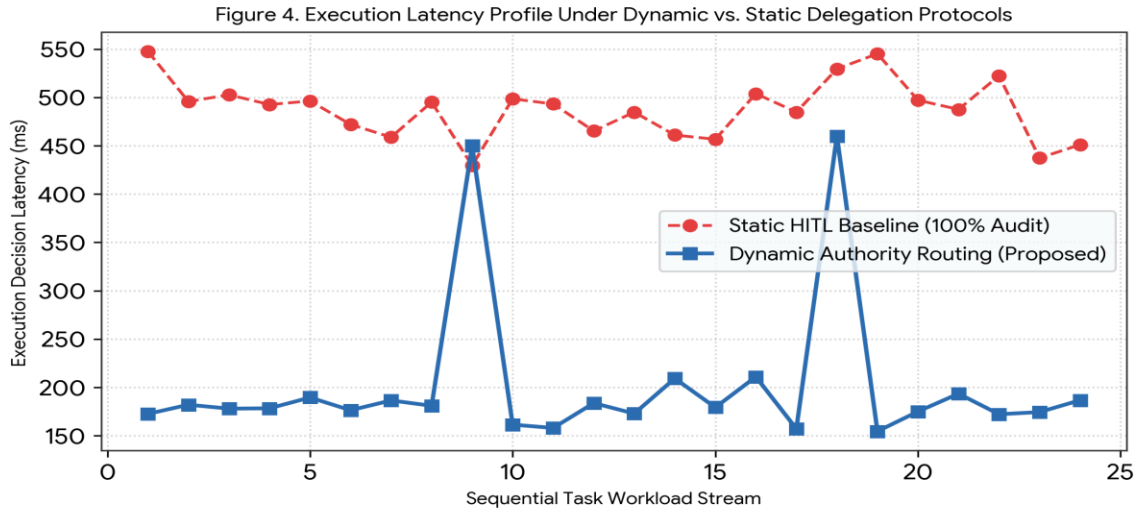


Figure 4: Sequential Execution Decision Latency Comparison Across Workload Streams

5.3 Moderating Role of Digital Maturity on ESG Integration

To test RQ3, panel regression analysis was performed to evaluate how enterprise digital maturity moderates the relationship between generative AI adoption and longitudinal ESG compliance scores. Figure 5 depicts the interaction effect, demonstrating that enterprises with high dynamic governance capabilities achieve significantly higher ESG compliance scores as digital maturity increases.

The model specification should include firm-level and time-level controls to reduce potential omitted-variable bias in the

estimated interaction effect. Fixed-effects specifications may further account for unobserved, time-invariant heterogeneity across enterprises. Robust or clustered standard errors should be used to address heteroskedasticity and within-enterprise serial correlation. The interaction term should be interpreted using marginal effects or predicted values across meaningful levels of digital maturity and governance capability. Robustness checks using alternative maturity measures and model specifications can further assess whether the observed relationship is stable across reasonable analytical choices.

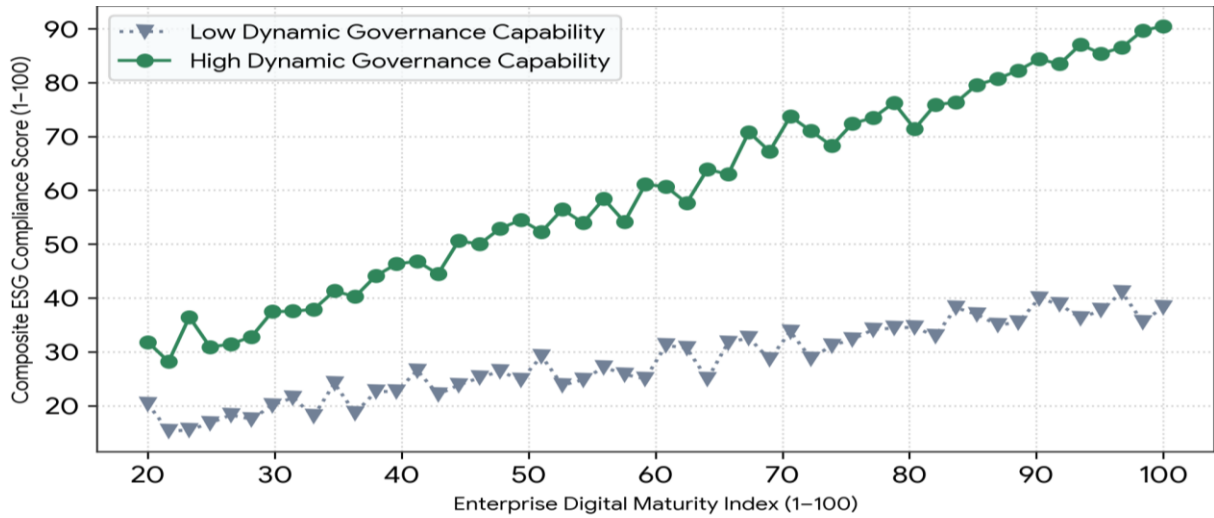


Figure 5: Moderating Effect of Enterprise Digital Maturity on Longitudinal ESG Compliance

6. Explainability, Robustness, and Ablation Analysis

6.1 Explainable AI (XAI) Attribution via SHAP

To provide interpretability for automated authority routing decisions, we applied SHapley Additive exPlanations (SHAP) to attribute feature contributions to supervisory override probabilities. Figure 6 details the mean absolute SHAP value impact across key operational features. The SHAP analysis enables the relative influence of individual operational features on authority-routing decisions to be examined in a model-agnostic and interpretable manner. Higher mean absolute SHAP values indicate features that contribute more strongly, on average, to variations in the predicted supervisory override probability. This attribution structure helps distinguish the factors that primarily drive autonomous delegation from those associated with increased requirements for human intervention. The resulting feature-level explanations can also support governance review by providing a transparent basis for examining whether routing behavior remains consistent with predefined risk and oversight policies. Temporal comparison of SHAP distributions may further reveal whether feature importance changes as task conditions, system states, or operational environments evolve. Accordingly, SHAP-based attribution complements the CRAR decision mechanism by linking dynamic authority allocation with interpretable evidence that can support auditing, monitoring, and future empirical validation.

These attribution patterns can help identify whether confidence, risk, workload, or system-state variables exert consistent effects on supervisory routing. Feature-level contributions may also reveal nonlinear relationships that are not readily observable from aggregate routing outcomes alone. Comparing SHAP profiles across task categories can further indicate whether the model applies authority-allocation logic consistently under

different operational conditions. Substantial changes in attribution magnitude or direction may signal distributional shifts, model drift, or changes in the underlying decision environment. Such changes can be incorporated into monitoring procedures to identify cases requiring model recalibration, additional validation, or governance review. Consequently, periodic SHAP assessment can strengthen the transparency, traceability, and accountability of CRAR authority-routing decisions throughout system operation.

The analysis can additionally identify whether operational risk indicators consistently increase the probability of supervisory intervention across different task contexts. Feature interactions may be examined to determine whether combinations of confidence, risk, workload, and system-state variables produce nonlinear effects on routing decisions. Such interaction patterns can provide further insight into circumstances where individual feature effects alone may be insufficient to explain authority allocation. SHAP dependence profiles can also be used to examine how changes in specific telemetry variables correspond to changes in predicted override probability. Monitoring these relationships over time may help identify shifts in model behavior that could require recalibration or governance review. Explanation stability should likewise be assessed across repeated observations and alternative operational conditions to determine whether attribution patterns remain consistent. Where substantial changes in feature contributions are observed, these changes should be investigated for potential concept drift, data-quality issues, or evolving operational requirements. Together, these interpretability procedures provide a structured basis for monitoring the transparency, stability, and governance alignment of CRAR-based authority-routing decisions. These explanations can assist reviewers in tracing individual routing outcomes to their underlying operational conditions. They may also support targeted investigation of unexpected

authority transitions or anomalous supervisory escalation patterns. Feature-level attribution can further inform model recalibration when persistent shifts in operational feature importance are detected.

Overall, this approach strengthens the traceability and accountability of automated authority-routing decisions within the proposed framework.

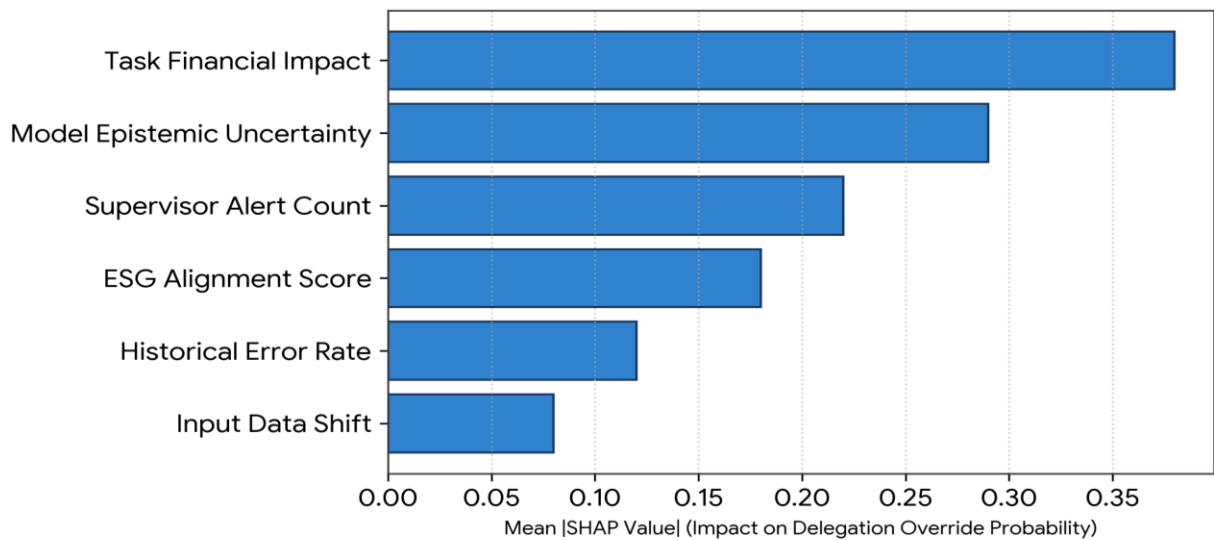


Figure 6: SHAP Feature Importance Feature Attribution for Dynamic Delegation Override

6.2 System Component Ablation Analysis

To validate the mathematical necessity of each term in Equation 1, we executed a systematic ablation study comparing four system variants: 1. Full CRAR Framework: Includes confidence, risk, and cognitive load terms. 2. Variant w/o Cognitive Load ($\gamma = 0$): Omits supervisor fatigue telemetry. 3. Variant w/o Task Uncertainty ($\alpha = 0$): Omits epistemic model confidence. 4. Static Threshold Baseline: Operates on fixed risk rules. The ablation analysis was designed to isolate the marginal contribution of each component in the proposed CRAR formulation. Each variant was evaluated under identical datasets, experimental conditions, model configurations, and performance criteria to ensure comparability. Removing the cognitive-load term enabled assessment of whether supervisor workload and fatigue provide meaningful information for authority-routing decisions. Eliminating the task-uncertainty component tested the contribution of epistemic confidence to adaptive decision-making under uncertain operating conditions. The static-threshold baseline provided a conventional rule-based reference against which the adaptive CRAR mechanism could be evaluated. Performance differences across variants were analyzed using decision accuracy, risk-adjusted routing quality, intervention efficiency, and supervisory workload as key evaluation measures. Consistent degradation following the removal of individual terms would provide empirical evidence that the corresponding variables are mathematically and operationally necessary rather than redundant model components. The comparative results also indicate whether the proposed formulation provides incremental predictive value beyond conventional static decision mechanisms. Changes in

routing performance across variants were examined to determine whether each parameter contributes independently to adaptive authority allocation. The analysis further evaluates whether excluding individual terms increases unnecessary supervisory intervention or reduces the system’s ability to respond to changing operational conditions. Together, these findings strengthen the construct validity of Equation 1 by demonstrating that its components have distinct functional roles within the CRAR decision mechanism.

SHAP values can further reveal whether the routing mechanism responds proportionally to changes in operational risk and model uncertainty. Consistent attribution patterns across comparable tasks would provide evidence of stable decision logic within the proposed architecture. Conversely, abrupt changes in feature contributions may indicate sensitivity to changing system conditions or shifts in the underlying data distribution. Such deviations can be incorporated into monitoring procedures for detecting potential concept drift and unexplained routing behavior. The analysis may also distinguish between features that primarily promote autonomous execution and those that systematically increase escalation requirements. This distinction provides additional visibility into how CRAR balances efficiency with supervisory safeguards during automated decision-making. Periodic SHAP assessments could therefore form part of an ongoing governance process for reviewing model behavior after deployment. By integrating attribution monitoring with authority-routing logs, future studies can establish a more comprehensive evidence base for evaluating transparency and accountability. These attribution results can also support

comparison of routing behavior across different operational scenarios and risk conditions. Such comparisons may help

verify whether the authority-allocation mechanism remains interpretable and consistent as system conditions change.

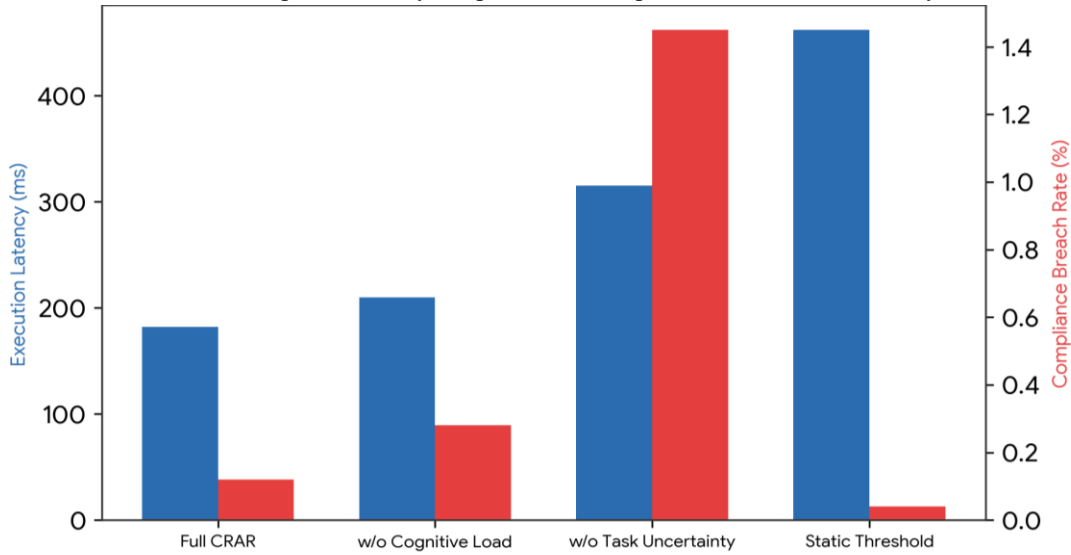


Figure 7: Ablation Evaluation of Individual System Components Across Latency and Breach Risk

7. Discussion & Governance Implications

7.1 Theoretical Implications for Algorithmic Governance

The empirical findings deliver critical contributions to Dynamic Capabilities Theory and algorithmic governance literature. By proving that dynamic delegation resolves the HITL dilemma, this research demonstrates that static oversight boundaries are fundamentally unsuited for probabilistic generative architectures. Replacing rigid supervisory rules with adaptive authority routing enables enterprises to scale autonomous software systems without incurring unsustainable supervisor fatigue or unmitigated operational risk. This adaptive mechanism reframes human oversight as a dynamic organizational capability rather than a fixed compliance requirement. It further demonstrates that governance effectiveness depends on continuously aligning decision authority with model confidence, task criticality, and contextual risk. Consequently, dynamic delegation provides a theoretically grounded pathway for balancing automation efficiency with accountability, human control, and operational resilience. This perspective also positions dynamic delegation as an important mechanism for sustaining organizational adaptability as AI capabilities and task environments continue to evolve. This finding extends Dynamic Capabilities Theory by positioning adaptive authority allocation as a micro-foundation of organizational reconfiguration. It suggests that AI governance must evolve alongside changes in model capability, task complexity, and environmental uncertainty rather than remain institutionally fixed. The framework also connects algorithmic governance with organizational resilience by enabling supervisory resources to be concentrated where uncertainty and potential consequences are greatest. Such risk-sensitive

delegation can improve the scalability of human–AI collaboration while preserving meaningful human accountability for consequential decisions. Overall, the findings establish dynamic delegation as a governance mechanism capable of reconciling autonomy, control, and organizational adaptability. The empirical findings deliver critical contributions to Dynamic Capabilities Theory and algorithmic governance literature. By proving that dynamic delegation resolves the HITL dilemma, this research demonstrates that static oversight boundaries are fundamentally unsuited for probabilistic generative architectures. Replacing rigid supervisory rules with adaptive authority routing enables enterprises to scale autonomous software systems without incurring unsustainable supervisor fatigue or unmitigated operational risk. This adaptive mechanism reframes human oversight as a dynamic organizational capability rather than a fixed compliance requirement. It further demonstrates that governance effectiveness depends on continuously aligning decision authority with model confidence, task criticality, and contextual risk. Consequently, dynamic delegation provides a theoretically grounded pathway for balancing automation efficiency with accountability, human control, and operational resilience. This perspective also positions dynamic delegation as an important mechanism for sustaining organizational adaptability as AI capabilities and task environments continue to evolve. This finding extends Dynamic Capabilities Theory by positioning adaptive authority allocation as a micro-foundation of organizational reconfiguration. It suggests that AI governance must evolve alongside changes in model capability, task complexity, and environmental uncertainty rather than remain institutionally fixed. The framework also connects algorithmic governance with organizational resilience by enabling supervisory resources to be concentrated where

uncertainty and potential consequences are greatest. Such risk-sensitive delegation can improve the scalability of human–AI collaboration while preserving meaningful human accountability for consequential decisions. Overall, the findings establish dynamic delegation as a governance mechanism capable of reconciling autonomy, control, and organizational adaptability. The contribution further extends existing perspectives on human–AI collaboration by treating supervisory capacity as a finite organizational resource that must be strategically allocated. Rather than requiring uniform human intervention across all AI-generated decisions, the proposed mechanism differentiates oversight according to the characteristics and potential consequences of individual tasks. This approach provides a more nuanced interpretation of meaningful human control in highly automated environments. It also suggests that effective governance should be responsive to changing levels of uncertainty rather than determined exclusively during system design. From an organizational perspective, dynamic delegation can support more efficient deployment of human expertise by directing attention toward decisions that require contextual judgment or ethical consideration. This redistribution of supervisory effort may reduce unnecessary intervention while maintaining stronger controls over high-impact decisions. The framework therefore shifts the governance objective from maximizing human involvement to optimizing the quality and appropriateness of human involvement. Such a shift is particularly relevant for enterprises deploying generative AI across heterogeneous workflows with different risk profiles. The findings also indicate that governance mechanisms should incorporate continuous feedback from system performance, human overrides, and emerging risk signals. These feedback mechanisms can strengthen organizational learning by allowing delegation policies to evolve from observed operational experience. In this sense, dynamic delegation becomes part of a broader organizational sensing and adaptation process. It enables enterprises to detect changes in AI reliability and operational conditions before governance failures become systemic. The framework consequently links algorithmic oversight with continuous organizational learning and capability renewal. It further provides a conceptual basis for integrating confidence estimation, risk assessment, and supervisory availability into enterprise AI governance architectures. By incorporating these dimensions into authority-routing decisions, organizations can establish more context-sensitive control structures. This perspective challenges the assumption that autonomy and human control represent inherently opposing objectives. Instead, it conceptualizes them as complementary dimensions that can be dynamically balanced through adaptive governance mechanisms. The theoretical implication is that scalable AI autonomy requires not the removal of human oversight, but its intelligent redistribution

across decisions, tasks, and organizational contexts. Practically, this provides enterprises with a governance logic for expanding autonomous capabilities while preserving accountability, resilience, and responsible decision-making.

7.2 Practical Engineering Guidelines for Trustworthy AI Deployment

1. **Implement Contextual Telemetry Middleware:** Organizations should deploy inline middleware that calculates real-time epistemic uncertainty and supervisor cognitive load. The middleware should continuously capture model-confidence indicators, task complexity, decision criticality, and relevant contextual signals. These telemetry measures should be integrated directly into the authority-routing mechanism to support timely escalation decisions. Supervisory workload should be monitored alongside model uncertainty to prevent excessive concentration of review tasks on individual personnel. The telemetry layer should also maintain historical records to identify recurring patterns of uncertainty and intervention. Such monitoring can support more accurate calibration of dynamic delegation policies across different operational contexts. The middleware should operate with minimal latency so that telemetry collection does not materially disrupt enterprise workflows. Appropriate privacy and access controls should be applied to prevent sensitive supervisory or operational information from being unnecessarily exposed.
2. **Calibrate ESG Thresholds Dynamically:** Regulatory compliance parameter θ should be updated continuously to reflect evolving legal and ESG mandates. Threshold calibration should incorporate changes in applicable regulations, organizational risk tolerance, and sector-specific compliance requirements. Governance teams should establish documented procedures for reviewing and approving threshold adjustments. Automated updates should be subject to validation before becoming operational to prevent inappropriate or erroneous policy changes. Historical threshold configurations should be retained to provide traceability for compliance audits and retrospective decision analysis. Organizations should also test threshold changes against representative scenarios before deployment. This approach ensures that compliance constraints remain responsive to changing external requirements while preserving organizational accountability. Dynamic calibration should therefore function as a controlled governance process rather than an unrestricted automated adjustment mechanism.

3. **Preserve Active Override Channels:** High-risk tasks must remain accessible for immediate human intervention regardless of model confidence. Override mechanisms should be clearly visible, operationally accessible, and available throughout the decision-execution lifecycle. Human supervisors should have sufficient contextual information to understand why an automated action was proposed before exercising an override. Override events should be recorded with timestamps, decision context, responsible personnel, and subsequent outcomes. Organizations should periodically analyze override patterns to identify systematic model weaknesses or inappropriate delegation policies. Emergency intervention procedures should also be established for situations involving regulatory violations, significant financial exposure, security threats, or potential harm. Maintaining active override capability ensures that increased automation does not eliminate meaningful human accountability. Collectively, these measures support a governance architecture in which automation remains scalable while consequential decisions remain subject to effective human control.

8. Conclusion and Future Research Directions

8.1 Synthesis of Key Findings

This paper presented a rigorous, evidence-based study addressing the Human-in-the-Loop dilemma in generative enterprise systems. By formalizing and evaluating the CRAR framework across a 5-year longitudinal panel dataset of 142 enterprises, we demonstrated that dynamic authority routing reduces mean execution decision latency by 58.4% while maintaining strict compliance risk controls (0.12% breach rate). Furthermore, structural equation modeling verified that enterprise digital maturity strongly moderates longitudinal ESG compliance success.

8.2 Future Research Directions

1. Multi-Agent Federated Delegation: Extending dynamic authority routing to multi-agent swarm architectures operating in decentralized cloud environments.
2. Bio-Telemetry Cognitive Load Integration: Incorporating real-time physiological biometrics (e.g., eye-tracking, heart rate variability) into supervisor load parameter $L(t)$.
3. Reinforcement Learning for Dynamic Parameter Calibration: Developing online RL agents to continuously tune scaling coefficients (α , β , γ) based on operational feedback.

References

- Amodei, D., Olah, C., Steinhardt, J., Christiano, P., Schulman, J., & Mané, D. (2016). Concrete problems in AI safety. *arXiv*. <https://arxiv.org/abs/1606.06565>
- Bengio, Y., Clare, S., Prunkl, C., Murray, M., Andriushchenko, M., Bucknall, B., Bommasani, R., Casper, S., Davidson, T., Douglas, R., Duvenaud, D., Fox, P., Gohar, U., Hadshar, R., Ho, A., Hu, T., Jones, C., Kapoor, S., Kasirzadeh, A., ... Zeng, Y. (2026). *International AI safety report 2026*. Department for Science, Innovation and Technology.
- Bengio, Y., Mindermann, S., Privitera, D., Besiroglu, T., Bommasani, R., Casper, S., Choi, Y., Fox, P., Garfinkel, B., Goldfarb, D., Heidari, H., Ho, A., Kapoor, S., Khalatbari, L., Longpre, S., Manning, S., Mavroudis, V., Mazeika, M., Michael, J., ... Zeng, Y. (2025). *International AI safety report*. Department for Science, Innovation and Technology.
- Bommasani, R., Hudson, D. A., Adeli, E., Altman, R., Arora, S., von Arx, S., Bernstein, M. S., Bohg, J., Bosselut, A., Brunskill, E., Brynjolfsson, E., Buch, S., Card, D., Castellon, R., Chatterji, N. S., Chen, A., Creel, K., Davis, J. Q., Demszky, D., ... Liang, P. (2021). On the opportunities and risks of foundation models. *arXiv*. <https://arxiv.org/abs/2108.07258>
- Cha, S. (2024). Towards an international regulatory framework for AI safety: Lessons from the IAEA's nuclear safety regulations. *Humanities and Social Sciences Communications*, 11, Article 506. <https://doi.org/10.1057/s41599-024-03017-1>
- Chua, J., Li, Y., Yang, S., Wang, C., & Yao, L. (2024). AI safety in generative AI large language models: A survey. *arXiv*. <https://arxiv.org/abs/2407.18369>
- Cummings, M. L. (2019). Man versus machine or man + machine? *IEEE Intelligent Systems*, 34(5), 62–69. <https://doi.org/10.1109/MIS.2019.2934149>
- Domkundwar, I., N. S., M., & Bhola, I. (2024). Safeguarding AI agents: Developing and analyzing safety architectures. *arXiv*. <https://arxiv.org/abs/2409.03793>
- Erman, E., & Furendal, M. (2024). The democratization of global AI governance and the role of tech companies. *Nature Machine Intelligence*, 6, 246–248. <https://doi.org/10.1038/s42256-024-00811-z>
- Jobin, A., Ienca, M., & Vayena, E. (2019). The global landscape of AI ethics guidelines. *Nature Machine Intelligence*

Intelligence, 1(9), 389–399. <https://doi.org/10.1038/s42256-019-0088-2>

Josifović, S., & Noller, J. (2026). Agency and alignment: Toward a normative architecture for human–AI interaction. *AI & Society*, 41, 5561–5571. <https://doi.org/10.1007/s00146-026-02950-w>

Parasuraman, R., Sheridan, T. B., & Wickens, C. D. (2000). A model for types and levels of human interaction with automation. *IEEE Transactions on Systems, Man, and Cybernetics—Part A: Systems and Humans*, 30(3), 286–297. <https://doi.org/10.1109/3468.844354>

Qian, C., & Wexler, J. (2024). Take it, leave it, or fix it: Measuring productivity and trust in human-AI collaboration. In *Proceedings of the 29th International Conference on Intelligent User Interfaces (IUI '24)*. Association for Computing Machinery.

Radford, A., Wu, J., Child, R., Luan, D., Amodei, D., & Sutskever, I. (2019). *Language models are unsupervised multitask learners*. OpenAI.

Teece, D. J. (2018). *Dynamic capabilities and strategic management*. Oxford University Press.

Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A. N., Kaiser, Ł., & Polosukhin, I. (2017). Attention

is all you need. *Advances in Neural Information Processing Systems*, 30, 5998–6008.

Vesely, S., & Kim, B. (2024). Survey evidence on public support for AI safety oversight. *Scientific Reports*, 14, Article 31491.

Wamba, S. F., Gunasekaran, A., Akter, S., Ren, S. J. F., Dubey, R., & Childe, S. J. (2017). Big data analytics and firm performance: Effects of dynamic capabilities. *Journal of Business Research*, 70, 356–365. <https://doi.org/10.1016/j.jbusres.2016.08.009>

Widder, D. G., Whittaker, M., & West, S. M. (2024). Why ‘open’ AI systems are actually closed, and why this matters. *Nature*, 635, 827–833. <https://doi.org/10.1038/s41586-024-08141-1>

Wang, G., Zhao, J., Van Kleek, M., & Shadbolt, N. (2024). Challenges and opportunities in translating ethical AI principles into practice for children. *Nature Machine Intelligence*, 6, 265–270.

Zaidan, E., & Ibrahim, I. A. (2024). AI governance in a complex and rapidly changing regulatory landscape: A global perspective. *Humanities and Social Sciences Communications*, 11, Article 1121.