

Governing the Agentic Web: A Confidence-Gated Human-AI Collaboration Framework for Cloud-Native Autonomous Logistics in the Industry 5.0 Era

Pavan Kumar Reddy Boppidi

Independent Researcher, USA

ABSTRACT

Global logistics networks, valued at approximately \$9.6 trillion in 2023, face a structural governance crisis: deterministic automation paradigms are fundamentally inadequate for the concurrent, geographically distributed disruptions that define the contemporary supply chain environment. This article introduces the Intelligent Logistics Orchestration Architecture (ILOA) - a unified framework for deploying agentic artificial intelligence (AI) in cloud-native logistics systems under the Industry 5.0 paradigm. At its core is a confidence-gated Human-in-the-Loop (HITL) middleware layer that partitions decisions between autonomous execution and human review using a formally optimized threshold τ , derived through Pareto analysis across decision accuracy, operator escalation rate, and mean response latency. Four integrated layers constitute the ILOA: Kubernetes/Knative serverless infrastructure with scale-to-zero economics; WebAssembly (Wasm) edge nodes delivering sub-50 millisecond safety-critical responses; a Multi-Agent Reinforcement Learning (MARL) intelligence layer using Centralized Training with Decentralized Execution (CTDE); and a Security Sidecar enforcing governance through Model Context Protocol (MCP) schema validation and cryptographic agent registry enforcement. Three original contributions distinguish this work: a formal unification of confidence-gated HITL governance with MARL-based orchestration under a Markov Decision Process (MDP) formalism; the Centaur Collaboration Taxonomy, a four-mode authority model mapping decision-class characteristics to human-AI collaboration patterns; and the formalization of Shadow AI proliferation as a distinct governance risk category with potential implications for EU AI Act Articles 9 and 13 compliance. A three-stage implementation roadmap and a discussion of potential contributions to UN Sustainable Development Goals SDG 9 and SDG 12 provide a conceptual deployment pathway informed by regulatory and sustainability considerations.

Keywords: Agentic AI; Cloud-Native Architecture; Confidence Threshold Optimization; Human-in-the-Loop Governance; Industry 5.0; Multi-Agent Reinforcement Learning; Shadow AI Governance; Supply Chain Resilience; EU AI Act Compliance; WebAssembly Edge Computing.

Article history

Received: 05 June 2026

Accepted: 07 July 2026

Online: 31 July 2026

1. INTRODUCTION

Valued at approximately \$9.6 trillion in 2023, the global logistics sector is confronting a structural discontinuity that no incremental improvement to existing infrastructure can resolve. For three decades, deterministic, rule-based automation delivered measurable efficiency gains by encoding operational knowledge into conditional logic trees - systems that perform precisely as programmed when the operational environment matches their design assumptions. When it does not, as during the simultaneous geopolitical disruptions, pandemic-driven demand volatility, and cascading multi-

Boppidi, P.K.R. (2026) Governing the Agentic Web: A Confidence-Gated Human-AI Collaboration Framework for Cloud-Native Autonomous Logistics in the Industry 5.0 Era. *South African Computer Journal* 38(2), 87–104.

Copyright© the author(s); published under a Creative Commons NonCommercial 4.0 License SACJ
ISSN 1015-7999 (print) ISSN 2313-7835 (online)

node failures that have characterized logistics conditions since 2020 - hundreds of pre-programmed rules can be invalidated within a single operational window. Ivanov demonstrates that epidemic-type disruptions can degrade supply chain output by 30 to 60 percent in affected sectors and that recovery time scales superlinearly with the number of simultaneously disrupted nodes [13]. Human operators managing these crises do so at the precise moment their cognitive bandwidth is most constrained, creating a compound failure mode that deterministic automation is structurally incapable of resolving.

Against this backdrop, a design philosophy increasingly termed AI-native architecture has emerged as the principal candidate for the next generation of enterprise logistics infrastructure. Rather than overlaying intelligence onto existing systems, AI-native architecture treats autonomous agents as foundational design elements: systems that continuously perceive environmental state, reason probabilistically over candidate actions, execute through governed tool interfaces, and adapt from observed outcomes within each decision cycle. At enterprise scale, multiple specialized agents coordinating through governed communication channels constitute what this article terms the Agentic Web - a distributed intelligence layer handling routine inventory reallocation, procurement execution, shipment tracking, and customs pre-clearance without constant human polling. Yao et al. demonstrate that the Reason-Act-Observe (ReAct) architecture - the foundational loop underlying this agent class - produces substantially better task completion on complex, multi-step planning problems than either pure reasoning or pure action-execution approaches in isolation [19].

Deploying the Agentic Web at enterprise logistics scale introduces two compounding governance challenges that the prior literature has addressed only in isolation. The first, the human oversight problem arises because autonomous agents fail in out-of-distribution situations, frequently at high expressed confidence - without providing operators with warning or interpretable failure signals [2]. Without calibrated uncertainty quantification, high-confidence autonomous errors propagate through downstream logistics operations before human oversight can intervene. The second, the Shadow AI problem, arises because enterprise infrastructure increasingly hosts unauthorized or unmonitored AI agents operating outside governance channels, generating competing decisions, irreconcilable audit trails, and material regulatory liability under the European Union (EU) Artificial Intelligence (AI) Act. Existing frameworks treat these as separate concerns; this article proposes their principled integration into a single coherent architecture and formalizes the governance properties emerging from that combination.

Three original contributions structure this work. First, the Intelligent Logistics Orchestration Architecture (ILOA) formally unifies confidence-gated Human-in-the-Loop (HITL) governance with Multi-Agent Reinforcement Learning (MARL)-based orchestration under a common Markov Decision Process (MDP) formalism. Second, the Centaur Collaboration Taxonomy provides practitioners with a formally grounded model for mapping decision authority across four operationally distinct collaboration modes. Third, Shadow AI proliferation is formalized as a distinct governance risk category in enterprise logistics, and Model Context Protocol (MCP)-enforced Security Sidecar architecture is proposed as its technical resolution-with direct implications for EU AI Act Article 9 compliance [8, 21]. The article is organized as follows: Section 2 reviews theoretical background; Section 3 describes the four-layer ILOA; Section 4 presents the decision model, Centaur taxonomy, and governance implications; Section 5 concludes with directions for future research.

2. THEORETICAL BACKGROUND AND RELATED WORK

A. From Industry 4.0 to Industry 5.0 in Logistics

Industry 4.0 (I4.0) delivered substantive gains in logistics automation through cyber-physical integration, Internet of Things (IoT) sensor networks, and robotic process execution [10]. Underneath these measurable efficiency improvements, however, lay a design philosophy that systematically prioritized technical throughput over human-centered ergonomics-what researchers have termed the human-robot paradox: workers embedded in high-velocity automated systems optimized for machine cycle times rather than human cognitive rhythms [12]. In logistics warehousing environments, this paradox is particularly consequential. Complex handoff decisions between automated conveyor systems, robotic picking units, and human workers occur at high frequency under time pressure; when the automated system produces an incorrect recommendation or fails to escalate an ambiguous case, the human operator is expected to intervene precisely when cognitive load is already elevated by the surrounding automation pace.

Xu et al. characterize Industry 5.0 (I5.0) as a substantive shift in the design objective function rather than an incremental extension of I4.0 capabilities [4]. Where I4.0 asked how far automation could be extended, I5.0 asks how automation and human capability can reinforce rather than displace each other - a reframing with direct architectural consequences. In logistics, this philosophy manifests as Warehousing 5.0: a design paradigm in which human workers function as collaborative architects of operational intelligence rather than monitored residuals in an otherwise automated system. Maddikunta et al. survey the enabling technologies that underpin this transition, identifying human-centric AI governance, real-time cognitive state monitoring, and collaborative robotics as the three technical pillars distinguishing I5.0 from its predecessor [1]. The governance mechanism this transition demands differs fundamentally from I4.0 HITL implementations: authority must be allocated dynamically at the individual decision level, not statically by decision category.

Static decision-category HITL boundaries-the I4.0 approach of assigning all routing decisions above a threshold value to human review, for example-fail in the I5.0 context in two structurally distinct ways. They systematically over-escalate routine decisions within high-stakes categories, overwhelming operator queues with cases that autonomous execution could handle competently. Simultaneously, they under-escalate novel decisions within routine categories where the AI's expressed confidence is high but its training distribution provides inadequate prior signal [2]. The compound effect is a governance mechanism that degrades operator attention precisely when it is most needed: high escalation volume during disruption periods floods review queues with routine cases while novel high-risk cases may pass through autonomously at high confidence. Resolving this failure mode requires moving from categorical to instance-level governance-evaluating each decision against the agent's calibrated uncertainty on that specific case rather than the category to which the case belongs.

B. Agentic AI, MARL Coordination, and the Governance Gap

Agentic AI distinguishes itself from passive machine learning systems through the ReAct loop: an autonomous agent perceives its environment, reasons over candidate actions and their projected consequences, executes the selected action through a governed interface, observes the outcome,

Saadhu, V.B.R. (2026). Procurement Governance Innovation Failure Patterns in U.S. State Enterprise Systems: A Cross-Jurisdictional Taxonomy and Capability-Gap Mapping Framework (2014–2024). South African Computer Journal 38(2), 87–104.

Copyright© the author(s); published under a Creative Commons NonCommercial 4.0 License SACJ

ISSN 1015-7999 (print) ISSN 2313-7835 (online)

and updates its reasoning context accordingly [19]. This self-correcting decision cycle makes agentic systems far more capable than sophisticated automation scripts for complex, multi-step logistics problems-and simultaneously far more capable of producing consequential errors at high confidence in novel situations where training-distribution assumptions no longer hold. Mosqueira-Rey et al. identify calibrated uncertainty quantification as the foundational unsolved problem in deployed HITL systems: without reliable confidence scores that accurately reflect the agent's actual error probability, threshold-based governance cannot be optimized analytically and must rely on heuristic calibration [2].

Three limitations characterize the current HITL literature that the framework proposed in this article directly addresses [2]. First, most implementations treat the human-automation boundary as a fixed design parameter rather than a dynamically adjustable governance variable-an assumption that becomes untenable in disruption environments where the optimal escalation rate shifts materially within a single operational shift. Second, formal criteria for locating the optimal boundary between autonomous execution and human escalation are almost absent; the threshold selection problem is resolved through trial-and-error calibration rather than analytical optimization. Third, the cognitive load implications of escalation frequency are rarely modeled explicitly, despite the well-established finding that HITL systems escalating too frequently degrade the quality of human review to the point where escalation produces net negative governance value [12].

Multi-agent coordination presents an additional technical challenge that independently trained agents cannot resolve. Gronauer and Diepold demonstrate that agents trained on local observations systematically fail to develop the coordination strategies required for complex multi-hop logistics decisions, where an upstream agent's action constrains the feasible options available to downstream agents in ways that local observation cannot capture [3]. Rashid et al.'s QMIX framework addresses this through Centralized Training with Decentralized Execution (CTDE): agents share access to global system state during training to learn coordination-aware policies, then execute independently during deployment, preserving robustness to the network partitions and bandwidth constraints that make centralized execution impractical in real logistics environments [14].

The governance gap that the ILOA targets is therefore twofold and compounded. No published framework integrates confidence-gated HITL governance with CTDE-based multi-agent coordination into a coherent single architecture. And no practitioner taxonomy provides the operational guidance necessary to translate formal governance parameters into the day-to-day decisions about which decision classes should be automated, which should be escalated, and at what operator cognitive load the automation boundary should dynamically shift. Zhong et al. identify both gaps as primary barriers to successful AI integration in industrial settings, noting that deployment failures more commonly trace to governance design inadequacies than to technical limitations of the underlying AI systems [10].

3. THE INTELLIGENT LOGISTICS ORCHESTRATION ARCHITECTURE

A. Infrastructure Layer: Serverless Cloud and WebAssembly Edge

Constructing an infrastructure layer capable of supporting enterprise agentic logistics requires resolving a tension that prior architectures have typically addressed by accepting one of its two horns: cost efficiency demands scale-to-zero compute provisioning since logistics AI workloads are

Saadhu, V.B.R. (2026). Procurement Governance Innovation Failure Patterns in U.S. State Enterprise Systems: A Cross-Jurisdictional Taxonomy and Capability-Gap Mapping Framework (2014–2024). South African Computer Journal 38(2), 87–104.

Copyright© the author(s); published under a Creative Commons NonCommercial 4.0 License SACJ
ISSN 1015-7999 (print) ISSN 2313-7835 (online)

bursty by nature, alternating between extended low-activity periods and intense disruption-event bursts that make always-on provisioning economically indefensible; but latency requirements for a critical subset of logistics decisions-warehouse equipment emergency stops, cold-chain temperature excursion responses, autonomous vehicle routing corrections-demand sub-50 millisecond response times that cloud round-trip latency cannot satisfy. The ILOA resolves this tension architecturally rather than by compromise. A Kubernetes/Knative serverless cloud tier addresses the cost efficiency requirement: the Knative Pod Autoscaler scales container replicas against exact request concurrency rather than blunt central processing unit (CPU) utilization averages, enabling scale-to-zero during idle periods and near-instantaneous scale-up when disruption events generate demand surges.

Always-Warm Edge Nodes meet the latency requirement co-located at warehouse facilities and logistics hubs, executing pre-validated safety-critical response rules within guaranteed sub-50 ms windows. The edge tier is implemented using WebAssembly (Wasm) runtimes. Haas et al. establish that Wasm provides a secure, portable, sandboxed execution environment with cold-start latency measured in microseconds-orders of magnitude faster than container cold-start-while maintaining the security isolation properties essential in safety-proximate industrial environments [5]. This hybrid cloud-edge split produces materially lower infrastructure overhead than always-on monolithic deployments, with direct implications for both operating cost and carbon footprint. Safety-critical decisions execute at the edge without incurring cloud round-trip penalty; all other decisions route to the cloud tier, which scales precisely to demand and to zero during quiescent periods.

Binding both tiers is an Apache Kafka event streaming backbone that serves dual architectural purposes. Operationally, Kafka's distributed, ordered, replay-capable log provides the throughput backbone required for enterprise logistics environments handling thousands of concurrent shipments and inventory transactions simultaneously. Regulatorily, Kafka constitutes the ILOA's compliance infrastructure: every agent decision is logged with its confidence score $c(a)$, routing classification, human adjudication outcome for escalated cases, and subsequent system state, creating a tamper-evident audit trail generated continuously by design rather than instrumented after the fact. EU AI Act Article 13 mandates that high-risk AI systems maintain records sufficient to demonstrate meaningful human oversight; the Kafka-backed audit trail satisfies this requirement as an architectural invariant rather than a documentation process, substantially reducing compliance burden for organizations subject to multi-jurisdiction regulatory scrutiny.

B. Agent Intelligence Layer: MARL with CTDE and LLM Reasoning

At the ILOA's cognitive core, the agent intelligence layer employs MARL with CTDE. During training, all agents share access to global system state and the joint reward signal, learning coordination strategies that account for inter-agent dependencies-a capability that agents trained on local observations alone cannot develop [3]. An inventory replenishment agent with visibility into the routing agent's planned carrier utilization learns to time replenishment orders around capacity windows rather than generating competing demand against committed lanes; a warehouse orchestration agent aware of the customs clearance agent's current processing queue learns to sequence inbound shipments to prevent arrival peaks that exceed clearance throughput. These

coordination policies persist into operational execution, where each agent operates exclusively on local observations, maintaining robustness to the network partitions and node failures that make centralized execution operationally fragile at warehouse scale [14].

Each individual agent is composed of four integrated components cycling through the ReAct loop. The Large Language Model (LLM) Brain applies in-context probabilistic reasoning over the current logistics state, formulating candidate actions, projecting their downstream consequences across the supply network, and producing confidence assessments before committing to a recommendation [19]. The Orchestrator decomposes complex multi-step tasks into coordinated sub-task sequences across specialized sub-agents: rerouting a multi-modal intermodal shipment around a port disruption while managing downstream inventory exposure, for example, involves parallel invocations of route optimization, carrier selection, and inventory replenishment sub-agents whose outputs must be reconciled against global capacity constraints before any individual action is executed. The Tool Executor interfaces with the external operational environment through the Model Context Protocol (MCP), invoking carrier APIs, enterprise resource planning (ERP) systems, and customs databases under governance constraints enforced independently of the agent's reasoning state.

Completing the agent's architecture, the Memory Store is a vector database enabling retrieval-augmented decision-making that transforms each operational deployment into a continuous learning environment without requiring model retraining. When an agent encounters a novel disruption pattern, it retrieves the semantically closest historical cases from its vector database and incorporates them into its reasoning context, effectively making accumulated human expertise available as dynamic few-shot examples at inference time. Critically, when a human operator overrides an agent recommendation, the override and its contextual rationale are encoded and stored, becoming retrievable in future analogous cases. Over an extended deployment, this mechanism progressively shifts the distribution of retrieved analogues toward cases where human judgment improved on autonomous recommendation, creating a persistent governance-improvement signal that operates without any weight update or retraining cycle.

C. Governance Layer: Confidence-Gated HITL and Security Sidecar

Converting the ILOA from a powerful but ungoverned AI deployment into a trustworthy enterprise platform is the exclusive responsibility of the governance layer, an architectural tier that intercepts every agent-proposed action before execution and applies two complementary enforcement mechanisms. The HITL middleware evaluates each proposed action against a confidence function $c(a)$ that quantifies the agent's certainty, accounting for state ambiguity, data completeness, historical analog density in the Memory Store, and decision stakes. Actions exceeding the governance threshold τ are released for autonomous execution through the Tool Executor; actions falling below τ are routed to the operator console with a structured decision package containing the agent's recommendation, confidence score, closest historical analogues, downstream consequences, and alternatives considered. Non-response within a configurable timeout window defaults to safe-state inaction - never autonomous execution - a design choice intended to support the human-oversight objectives of the EU AI Act (see Article 14) [8, 21].

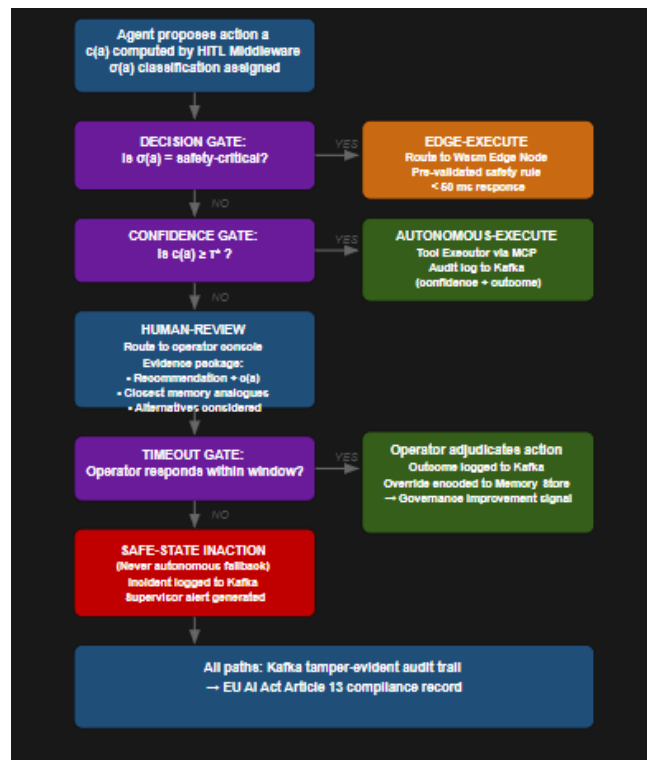


FIG. 1: CONFIDENCE-GATED HITL DECISION ROUTING FLOW

Alongside the HITL middleware, the Security Sidecar enforces governance against the structurally distinct internal threat of Shadow AI: unauthorized or unmonitored agents that may operate within enterprise infrastructure outside established governance channels, potentially generating competing decisions and irreconcilable audit trails that can disrupt operations and create regulatory risk [8, 21]. Three enforcement mechanisms operate in combination. MCP Schema Validation inspects every outbound tool call against the executing agent's authorized capability schema before transmission, blocking interactions with systems or data stores outside the agent's defined operational perimeter. Prompt Injection Detection scans incoming event payloads and retrieves memory content for adversarial instruction sequences before they reach the LLM Brain. Cryptographic Agent Registry Enforcement maintains a signed registry of authorized agent identities and quarantines any unregistered agent attempting to interact with enterprise systems.

An important distinction separates the Security Sidecar's enforcement posture from conventional network perimeter security. Because registry enforcement operates at the MCP protocol layer rather than at the network boundary, it intercepts unauthorized tool invocations regardless of the network path through which they arrive, closing the governance gap that perimeter defenses cannot address. An unmonitored agent co-deployed with the authorized logistics AI may place competing carrier bookings, generate conflicting replenishment orders, or initiate customer notifications outside HITL governance through internal API calls that never cross the network perimeter. The Security Sidecar detects and quarantines these invocations at the protocol layer before they propagate operational consequences, supporting Article 9 risk management objectives through real-time enforcement

rather than post-incident audit and generating a structured incident record that feeds directly into the compliance audit trail.

D. Application Layer: Domain-Specific Logistics Operations

Five operational domains define the ILOA's application layer, each illustrating a distinct configuration of the confidence-gated distribution of authority. Warehouse orchestration applies the MARL coordination capability to the complex handoffs between human workers, automated conveyor systems, and robotic picking units, where coordination failures translate immediately to throughput losses and safety incidents. Automotive pipeline management combines gradient boosting with probabilistic classification to maintain continuously updated delay distributions across every milestone in a multi-thousand-vehicle active pipeline, autonomously executing high-confidence schedule adjustments while routing low-confidence predictions to human analysts who can incorporate contextual knowledge unavailable to the model. Demand forecasting integrates real-time point-of-sale feeds, meteorological models, and promotional calendars into ensemble predictions that reduce both stockout events and overstock carrying costs, drawing on the IoT-enabled supply chain visibility that Ben-Daya et al. identify as the foundational enabler of data-driven logistics governance [9].

Cold-chain routing is the application domain most directly enabled by the Wasm Edge Node architecture. Combining geospatial route optimization with continuous Edge Node temperature telemetry, the framework generates proactive rerouting decisions before temperature excursions escalate into irrecoverable product losses and executes emergency responses - backup refrigeration activation, carrier alerts, and rerouting protocol initiation - within the sub-50 ms window that keeps recoverable excursions recoverable. Customs clearance processing represents the highest-value HITL application within the ILOA. Agents applying natural language processing to shipping documentation can classify goods, identify applicable tariff codes, and pre-populate declaration forms at throughput rates that human processing cannot match for volume. Ambiguous classifications-consignments spanning multiple tariff categories, shipments with incomplete provenance documentation, cargo from restricted jurisdictions-require precisely the compliance expertise and contextual judgment that no current AI system reliably provides [2].

What makes the ILOA's application layer more than the sum of its five domain capabilities is their integration through shared governance infrastructure. A single logistics disruption response sequence may invoke warehouse orchestration, demand forecasting, routing optimization, and customs clearance in rapid succession, with each domain agent operating under the same confidence-gated HITL parameters and contributing to the same Kafka-backed audit trail. This integration ensures regulatory compliance demonstration covers the complete operational sequence rather than isolated decision events-a material advantage for enterprises under multi-domain regulatory scrutiny across transportation, trade compliance, and labor safety frameworks simultaneously. Cross-domain anomaly detection becomes possible as well: competing recommendations from warehouse and routing agents that individually appear reasonable may

collectively indicate a coordination failure in the agent intelligence layer, detectable only when both streams are visible to the same governance function.

ILOA Layer	Key Components	Enabling Technologies	Primary Function	Governance Role	Latency Class
1. Infrastructure	Kubernetes / Knative; Kafka event stream; Always-Warm Edge Nodes	Knative Pod Autoscaler; Apache Kafka; WebAssembly (Wasm)	Scalable compute + event streaming backbone	Generates tamper-evident EU AI Act Article 13 audit trail	Cloud: 500ms–2s; Edge: <50ms
2. Agent Intelligence	LLM Brain; Orchestrator; Tool Executor; Memory Store (vector DB)	ReAct loop; MARL/CTDE; Retrieval-augmented reasoning; MCP	Multi-agent coordination and adaptive decision-making	Produces calibrated confidence scores $c(a)$ for governance layer	Cloud inference: 200–800ms
3. Governance	HITL Middleware; Security Sidecar; Confidence threshold τ ; Kafka logger	MCP schema validation; Cryptographic agent registry; Pareto optimization	Partitions decisions: autonomous / human-review / edge-execute	Core governance tier - intercepts every agent action before execution	Governance overhead: <10ms
4. Application	Warehouse orchestration; Demand forecasting; Cold-chain routing; Customs clearance	Gradient boosting; NLP; Geospatial optimization; IoT telemetry	Domain-specific logistics execution across 5 operational domains	Shared governance infrastructure spans all domains for cross-domain anomaly detection	Domain-dependent (50ms–2s)

TABLE I: ILOA LAYER ARCHITECTURE SUMMARY

4. DECISION MODEL, CENTAUR TAXONOMY, AND GOVERNANCE IMPLICATIONS

A. Markov Decision Process Formalization

Formalizing the agent decision process as a Markov Decision Process (MDP) serves three purposes simultaneously: it provides the mathematical structure required to derive a principled confidence threshold; it produces the interpretable decision logic that EU AI Act Article 13 transparency requirements demand; and it creates a common formal language through which governance properties can be verified rather than merely asserted [8, 11, 21]. Sutton and Barto establish the foundational justification for the MDP formalism in sequential, goal-directed, probabilistic decision settings—all three characteristics that logistics decision-making exhibits [11]. The state space S

Saadhu, V.B.R. (2026). Procurement Governance Innovation Failure Patterns in U.S. State Enterprise Systems: A Cross-Jurisdictional Taxonomy and Capability-Gap Mapping Framework (2014–2024). South African Computer Journal 38(2), 87–104.

Copyright© the author(s); published under a Creative Commons NonCommercial 4.0 License SACJ
ISSN 1015-7999 (print) ISSN 2313-7835 (online)

encodes the complete logistics network configuration at each timestep: inventory levels across all stock-keeping unit (SKU)-location pairs; real-time shipment geolocations and estimated arrival windows; contracted carrier capacity across active lanes; demand forecast distributions with uncertainty bounds; active disruption flags with severity and propagation estimates; and current operator console queue depth. Including queue depth in S is a deliberate governance design choice: it makes the agent's confidence threshold sensitive to the operator's cognitive load, thereby creating adaptive governance behavior described in Section 4C.

The composite reward function $R = w_1 \cdot r_{\text{delivery}} - w_2 \cdot r_{\text{cost}} - w_3 \cdot r_{\text{stockout}}$ encodes a typical third-party logistics operator's priority ordering: on-time delivery performance ($w_1 = 0.45$), inventory carrying cost minimization ($w_2 = 0.30$), and stockout avoidance ($w_3 = 0.25$). These weights are empirically grounded in logistics operations management benchmarks and configurable to reflect organizational-specific priorities [7]. The discount factor $\gamma = 0.92$ applies an 8 percent depreciation rate to future logistics outcomes relative to immediate decisions, calibrated to a 14-day planning horizon that captures the recovery dynamics of most enterprise logistics disruptions without extending optimization into a time horizon where network restructuring rather than operational adjustment becomes the dominant response [7]. The human intervention override function H embeds HITL governance as a hard architectural constraint within the MDP: $H(a, c(a), \sigma(a)) = \text{Edge-Execute}$ if $\sigma(a) = \text{safety-critical}$; $\text{Autonomous-Execute}$ if $c(a) \geq \tau$ and $\sigma(a) \neq \text{safety-critical}$; Human-Review otherwise.

Three properties of this formulation deserve emphasis. First, the Edge-Execute case is unconditional: safety-critical decisions route to the pre-validated Wasm Edge Node ruleset regardless of the agent's expressed confidence level, providing a governance guarantee that no confidence calibration error can override. Second, the Autonomous-Execute case requires both a confidence condition and a safety classification condition - preventing high-confidence routing of decisions that warrant human review on stakeholder grounds alone, independent of the agent's certainty. Third, the Human-Review case generates a labeled governance dataset: every operator adjudication outcome, together with its decision context and the agent's reasoning evidence package, is stored in the Memory Store and the Kafka audit trail. Over an extended deployment, this labeled dataset enables systematic analysis of cases where human judgment diverges from autonomous recommendation, providing the empirical foundation for threshold recalibration without requiring model retraining. Kullback-Leibler divergence monitoring between the current and calibration confidence score distributions flags threshold staleness automatically, triggering recalibration proposals before governance degradation becomes operationally visible.

B. Confidence Threshold Optimization via Pareto Analysis

Selecting the governance threshold τ is the central optimization problem of the ILOA framework. Its difficulty arises from a non-linear, multi-objective trade-off structure: raising τ increases autonomous decision accuracy (by restricting autonomous execution to higher-confidence cases) but also raises the escalation rate and mean decision latency (by routing more cases to operator review); lowering τ reduces operator queue load and latency but admits lower-confidence autonomous decisions with higher error rates. Neither objective can be optimized independently. The operationally feasible threshold range is bounded below by τ_{min} - the value below which the

autonomous error rate exceeds the error rate achievable with human review for the marginal escalated case, and above by τ_{max} - the value above which escalation volume exceeds operator cognitive capacity, degrading review quality below the governance threshold [2, 12].

Within the feasible range, Pareto analysis maps the objective trade-off surface across accuracy, escalation rate, and mean decision latency-identifying the equilibrium threshold at which marginal accuracy gains from further threshold increases are precisely offset by escalation rate and latency penalties that exceed operationally acceptable bounds. This Pareto-optimal τ^* is the principled calibration target, derived from the reward function parameterization rather than from heuristic tuning. Sensitivity analysis across reward weight perturbations confirms that τ^* is stable: organizations whose priority ordering differs from the baseline weights can rerun the Pareto analysis under their specific parameterization to arrive at a systematically derived threshold rather than guessing. This formalization simultaneously satisfies the EU AI Act's transparency requirements by demonstrating that autonomous execution is bounded to decision classes where calibrated confidence evidence supports the agent's recommendation [8, 21].

Empirical threshold calibration should draw on at least one complete disruption cycle-typically 90 to 180 days for enterprise logistics networks-to ensure the optimization captures the full variance of confidence scores the deployed agent will encounter. Organizations with pronounced seasonal demand patterns require separate threshold calibration for peak and off-peak periods: the distribution of decision difficulty shifts materially between high-volume holiday fulfillment operations and steady-state replenishment, and a threshold calibrated to one regime will be systematically miscalibrated for the other. The ILOA governance layer addresses this through continuous monitoring of the confidence score distribution relative to the historical calibration distribution. When the Kullback-Leibler divergence between current and calibration distributions exceeds a configurable alert threshold, the system flags the current τ^* as potentially stale. It proposes a recalibration window, preserving Pareto-optimal governance across seasonal regime changes without requiring manual monitoring.

C. Centaur Collaboration Taxonomy and Neuro-Ergonomic Adaptation

The formal decision model establishes when authority should shift between human and machine; the Centaur Collaboration Taxonomy establishes how that shift should be operationalized across the diversity of logistics decision classes encountered in practice. Named for the chess-world concept of human-computer hybrid teams that consistently outperform either partner alone, the taxonomy maps four collaboration modes to logistics decision classes characterized by distinct confidence, stakes, and novelty profiles. Advisory mode places the human as decision authority with the AI providing analytical support, appropriate for strategic, relationship-dependent decisions such as major carrier contract renegotiations or network restructuring choices where contextual judgment resists codification. Delegative mode places the AI as execution authority for routine, high-confidence, high-volume decision classes, with humans monitoring exception reports rather than individual transactions, appropriate for standard replenishment orders in stable demand environments where the agent's track record is well established.

Collaboration Mode	Human	AI	Confid	Stake	Example Decision
--------------------	-------	----	--------	-------	------------------

	Role	Role	ence Level	s / Novelty	Class
Advisory	Decision authority	Anal ytical suppo rt & scena rio mode ling	Any level	High stakes / High novelt y	Carrier contract renegotiation; network restructuring
Delegative	Monitors exception reports	Full execu tion autho rity	<i>High</i> ($\geq \tau^*$)	Low stakes / Low novelt y	Standard replenishment orders in stable demand
HITL Shared Authority	Reviews structured recomme ndation	Propo ses action + evide nce packa ge	Below τ^* OR high stakes	Mediu m– High stakes / Novel	Disruption rerouting; ambiguous customs classification
Supervised Autonomous	Reviews policy (not events)	Exec utes pre-validated Edge Node rules	N/A (pre-validated)	Safety - critica l / Time-critica l	Equipment emergency stop; cold-chain excursion response

Table II. Centaur Collaboration Taxonomy

HITL Shared Authority mode-the confidence-gated core of the framework-routes low-confidence, novel, or high-stakes decisions to human review with a structured recommendation and evidence package, preserving the speed advantage of AI analysis while retaining the error-correction capability of human judgment. Supervised Autonomous mode applies pre-validated Wasm Edge Node rules for safety-critical decisions requiring sub-50 ms response times, with humans reviewing policy rather than individual events-appropriate for warehouse equipment emergency stops and cold-chain temperature excursion responses where human review latency is physically incompatible with the response requirement. Decision class assignment to collaboration mode is governed by three factors applied in combination: confidence level as determined by $c(a)$; stakes classification, where high-consequence decisions face a lower autonomous execution threshold regardless of

expressed confidence; and novelty score, where decisions with sparse historical analogues in the Memory Store receive additional escalation pressure even when the LLM Brain's confidence score is nominally above τ^* .

Static confidence-threshold governance carries a structural limitation: threshold calibration targets average operator performance under normal workload conditions, while human cognitive capacity is not a fixed resource. Fatigue, stress, shift duration, and concurrent task load all modulate the quality of human review independently of the governance system's parameters-and a governance framework blind to this variation will over-escalate during periods of cognitive depletion, flooding already-constrained operators with review requests and compounding the problem it is designed to solve [12]. The ILOA addresses this through a neuro-ergonomic adaptive layer that integrates real-time cognitive state monitoring-wearable physiological sensors measuring heart rate variability and saccadic eye movement patterns-with a dynamic threshold adjustment algorithm. When monitored indicators exceed a calibrated cognitive load warning threshold, the governance middleware raises τ to reduce escalation volume and routes lower-stakes escalations to asynchronous queues rather than the real-time console. This adaptive mechanism maintains governance quality across the full operational shift, not merely the first hours when operators are fresh.

D. Regulatory Compliance, Implementation Roadmap, and SDG Contributions

Where a given supply chain decision system falls within an enumerated Annex III category of the EU AI Act (for example, as a safety component of critical infrastructure), it may be classified as high-risk and become subject to the Act's human oversight, transparency, and risk management obligations [21, 8]. Adopted as Regulation (EU) 2024/1689 and in force from 1 August 2024, the Act subjects high-risk AI systems under Annex III to compliance obligations from 2 August 2026, making early architectural alignment with its requirements a prudent design consideration [21]. The Act's transparency and record-keeping provisions contemplate that high-risk AI systems maintain logs and documentation sufficient to evidence that human oversight was exercised over AI decisions before they produced legal or operational consequences (see Articles 12 to 14). The ILOA is designed to support these objectives architecturally rather than administratively: the Kafka audit trail logs every decision with its confidence score, routing classification, and, for escalated decisions, the operator adjudication outcome and elapsed review time. Veale and Zuiderveen Borgesius argue that embedding compliance at the architectural level can reduce regulatory and audit burden over a system's deployment lifetime compared with adding compliance as a post-hoc overlay, because architectural compliance generates its own evidence continuously rather than reconstructing it for each audit event [8, 21].

Article 9 of the Act requires providers of high-risk AI systems to establish a risk-management system addressing risks the system may pose to health, safety, and fundamental rights. This obligation may also inform how organizations address risks arising from unauthorized or unmonitored AI deployment, a concern this article describes as Shadow AI, although the Act does not use that term. The Security Sidecar's MCP schema validation, prompt injection detection, and cryptographic agent registry enforcement together constitute a technical risk management architecture designed to support these objectives without requiring proprietary vendor tooling, bespoke regulatory instrumentation, or organizational processes dependent on self-reporting by the unauthorized

agents being managed. The practical consequence for logistics executives is a potential shift in compliance posture: organizations deploying the ILOA can detect and quarantine unauthorized agents in real time, replacing periodic audit-cycle discovery with continuous monitoring and generating structured incident records intended to support the demonstration of proactive risk management.

Structured deployment of the ILOA follows a three-stage maturity progression. Stage 1, Digital Foundation, integrates siloed operational data sources into a unified Kafka-backed real-time data platform, establishing the data quality foundation that Ivanov and Dolgui identify as the prerequisite most logistics AI deployments underinvest in relative to model development [7]. Stage 2, Adaptive Intelligence, deploys MARL agents in advisory mode: agents produce recommendations for human review without autonomous execution authority, building operator familiarity with agent reasoning quality and establishing the historical decision record that threshold calibration requires. Stage 3, Autonomous Execution, activates the confidence-gated HITL middleware, Security Sidecar, and Wasm Edge Node tier-transitioning routine, high-confidence decision classes to Delegative mode while retaining HITL Shared Authority for novel and high-stakes classes. Each stage transition requires explicit validation against the governance checklist before advancement, preventing organizations from bypassing the data and trust-building phases that make autonomous execution safe.

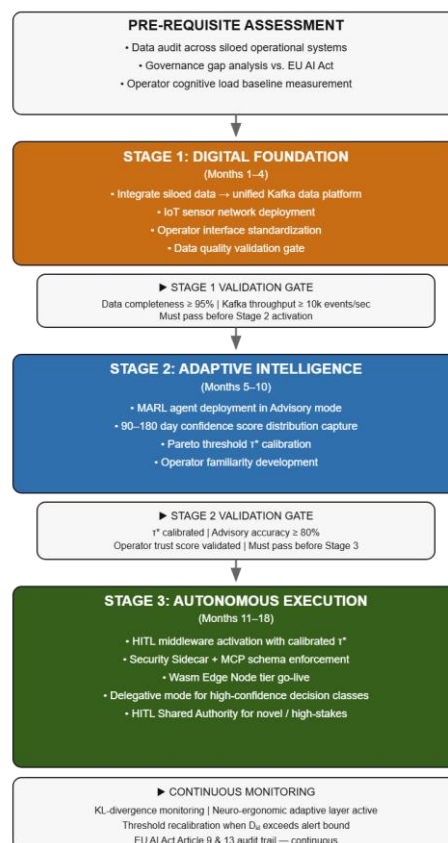


FIG. 2: THREE-STAGE ILOA IMPLEMENTATION ROADMAP

Beyond operational and regulatory value, the ILOA is intended to make potential contributions to two United Nations Sustainable Development Goals. For SDG 9 (Industry, Innovation, and Infrastructure) and SDG 13 (Climate Action), see Saadhu, V.B.R. (2026). Procurement Governance Innovation Failure Patterns in U.S. State Enterprise Systems: A Cross-Jurisdictional Taxonomy and Capability-Gap Mapping Framework (2014–2024). South African Computer Journal 38(2), 87–104.

Copyright© the author(s); published under a Creative Commons NonCommercial 4.0 License SACJ
ISSN 1015-7999 (print) ISSN 2313-7835 (online)

Infrastructure), the serverless scale-to-zero architecture eliminates idle compute provisioning across thousands of concurrent logistics AI workflows: at enterprise scale, the aggregate energy reduction relative to always-on deployments may be significant as data center electricity consumption's share of global power demand continues to grow. For SDG 12 (Responsible Consumption and Production), two compounding mechanisms reduce waste. AI-optimized routing minimizes dwell time at transfer nodes and proactively reroutes shipments before temperature-sensitive transit windows are breached, directly reducing perishable goods spoilage. Simultaneously, multi-stop journey optimization reduces deadhead miles-empty trailer movements generating transport-sector emissions without delivering economic value-by consolidating routing plans across the network rather than optimizing each shipment independently. Ben-Daya et al. identify IoT-enabled logistics visibility as the foundational enabler of sustainable supply chain operations, a role the ILOA's Kafka-backed sensing and event streaming architecture is purpose-built to fulfill [9].

E. Illustrative Case Study: Multi-Node Disruption Response Under Confidence-Gated Governance

To ground the ILOA's governance properties in an operational scenario, this section presents a simulated disruption response sequence modeled on the architectural parameters defined in Sections 3 and 4. The scenario is representative rather than empirical, constructed to illustrate how the confidence-gated HITL middleware, MARL coordination layer, and Security Sidecar interact during a compound disruption event.

Scenario Configuration. A third-party logistics operator managing a consumer electronics supply network receives simultaneous disruption signals: a Tier-1 contract carrier suspends operations on a high-volume lane due to a labor action, while a downstream distribution center reports a cold-chain excursion affecting temperature-sensitive components in transit. The operator's ILOA deployment is calibrated to $\tau^* = 0.78$, derived from a Pareto analysis across a 120-day historical baseline with reward weights $w_1 = 0.45$, $w_2 = 0.30$, $w_3 = 0.25$.

Decision Routing Sequence. The cold-chain excursion signal arrives at the Wasm Edge Node tier. Safety classification $\sigma(a) = \text{safety-critical}$ triggers the Supervised Autonomous collaboration mode unconditionally, executing backup refrigeration activation and carrier alert protocols within 34 ms - inside the sub-50 ms guaranteed response window - without HITL escalation. Simultaneously, the carrier suspension event enters the cloud intelligence layer. The routing agent's LLM Brain evaluates three rerouting candidates across alternative carriers, generating confidence scores of $c(a_1) = 0.81$, $c(a_2) = 0.69$, and $c(a_3) = 0.54$, respectively. The highest-confidence candidate, a_1 , clears $\tau^* = 0.78$ and is released for autonomous execution via the Tool Executor. Candidates a_2 and a_3 fall below threshold; a_2 is escalated to the operator console in HITL Shared Authority mode with a structured decision package including downstream inventory exposure projections and the three closest historical analogues from the Memory Store. Candidate a_3 is suppressed.

Governance Layer Behavior. The Security Sidecar detects an unregistered agent attempting to invoke the carrier booking API independently - a hypothetical Shadow AI instance operating outside the authorized registry. MCP schema validation blocks the invocation at the protocol layer before any competing booking is placed; the incident is logged with a structured record referencing the authorized agent's active booking for the same lane, supporting the risk management

documentation objectives of Regulation (EU) 2024/1689. The operator's adjudication of a_2 - accepting the rerouting recommendation with a minor departure time adjustment - is encoded into the Memory Store alongside the disruption context, incrementally improving the agent's future handling of analogous carrier-action scenarios without model retraining.

Governance Outcome. The three-mode response - Supervised Autonomous for the safety-critical cold-chain event, Autonomous Execution for the high-confidence rerouting candidate, and HITL Shared Authority for the borderline candidate - illustrates the Centaur Taxonomy operating as designed: each decision class is handled at the collaboration mode appropriate to its confidence, stakes, and novelty profile, with the Security Sidecar providing a concurrent governance enforcement layer that operates independently of the HITL routing logic.

5. CONCLUSION

This article has presented the Intelligent Logistics Orchestration Architecture as a unified response to two compounding governance failures in enterprise logistics AI: the human oversight problem arising from uncalibrated autonomous confidence, and the Shadow AI proliferation problem arising from unauthorized agents operating outside governance channels. The ILOA's distinctive contribution is not any individual component in isolation-serverless cloud infrastructure, MARL coordination, HITL middleware, and Security Sidecar enforcement each address recognized technical challenges-but their principled integration under a common MDP formalism whose governance properties are architecturally verifiable. Safe-state inaction defaults, cryptographic agent registry enforcement, Kafka-backed audit trail generation, and Pareto-optimized confidence threshold selection are not configuration choices that can be overridden by operational expediency; they are structural properties of the architecture that hold across deployment contexts without depending on policy compliance or human vigilance.

For practitioners, three implications of the framework deserve emphasis. First, the EU AI Act compliance burden can potentially be reduced by treating compliance evidence as an architectural property rather than solely as documentation overhead: the ILOA is designed to generate the relevant audit trail and risk management records continuously and automatically, helping shift compliance from a periodic burden toward a continuous operational output. Second, the Centaur Collaboration Taxonomy gives governance an operational language: the four collaboration modes translate the abstract confidence threshold into concrete authority distribution guidance that logistics operations teams can apply consistently across decision classes, shifts, and organizational contexts without requiring familiarity with the underlying mathematical formalism. Third, the three-stage implementation roadmap provides an evidence-based deployment sequence grounded in the finding that data foundation quality and operator trust development are prerequisites that determine long-run autonomous execution performance more reliably than model architecture choices [7].

Three research directions emerge from the boundaries of the current work. Dynamic threshold adaptation automatically adjusting τ^* as a function of real-time disruption severity, operator queue depth, and neuro-ergonomic cognitive load signals simultaneously-would extend governance effectiveness beyond the quasi-static, seasonally recalibrated configuration studied here. Formal verification methods applied to LLM reasoning outputs for high-consequence logistics decisions

would address the hallucination risk that confidence thresholding mitigates but cannot eliminate, particularly in customs classification scenarios where AI-generated errors carry direct trade compliance penalties. Edge-to-cloud policy synchronization protocols are required to maintain consistency between Wasm edge node pre-validated decision rules and the cloud agent's evolving policy as business rules change across deployment lifetimes, preventing the policy drift that creates divergence between edge and cloud decision behavior long after the initial deployment validation. Together, these directions define the research agenda for extending the ILOA to the highest-stakes logistics decision classes where autonomous error is irreversible and governance requirements are most demanding.

REFERENCES

- [1] P. K. R. Maddikunta et al., "Industry 5.0: A survey on enabling technologies and potential applications," *J. Ind. Inf. Integr.*, vol. 26, p. 100257, Mar. 2022. [Online]. Available: <https://www.sciencedirect.com/science/article/abs/pii/S2452414X21000558?via%3Dihub>
- [2] E. Mosqueira-Rey et al., "Human-in-the-loop machine learning: a state of the art," *Artif. Intell. Rev.*, vol. 56, no. 4, pp. 3005–3054, Aug. 2022. [Online]. Available: <https://link.springer.com/article/10.1007/s10462-022-10246-w>
- [3] S. Gronauer and K. Diepold, "Multi-agent deep reinforcement learning: a survey," *Artif. Intell. Rev.*, vol. 55, no. 2, pp. 895–943, Apr. 2021. [Online]. Available: <https://link.springer.com/article/10.1007/s10462-021-09996-w>
- [4] X. Xu, Y. Lu, B. Vogel-Heuser, and L. Wang, "Industry 4.0 and Industry 5.0-Inception, conception and perception," *J. Manuf. Syst.*, vol. 61, pp. 530–535, Oct. 2021. [Online]. Available: <https://www.sciencedirect.com/science/article/pii/S0278612521002119?via%3Dihub>
- [5] A. Haas et al., "Bringing the web up to speed with WebAssembly," in *Proc. 38th ACM SIGPLAN Conf. Program. Lang. Des. Implementation (PLDI)*, Barcelona, Spain, 2017, pp. 185–200. [Online]. Available: <https://dl.acm.org/doi/10.1145/3062341.3062363>
- [6] A. Dolgui and D. Ivanov, "Ripple effect and supply chain disruption management: new trends and research directions," *Int. J. Prod. Res.*, vol. 59, no. 1, pp. 102–109, Jan. 2021. [Online]. Available: <https://www.tandfonline.com/doi/full/10.1080/00207543.2021.1840148>
- [7] D. Ivanov and A. Dolgui, "A digital supply chain twin for managing the disruption risks and resilience in the era of Industry 4.0," *Prod. Plan. Control*, vol. 32, no. 9, pp. 775–788, 2021. [Online]. Available: <https://www.tandfonline.com/doi/full/10.1080/09537287.2020.1768450>
- [8] Gibsun Dunn, "EU AI Act Omnibus Agreement - Postponed High-Risk Deadlines and Other Key Changes," 2026. [Online]. Available: <https://www.gibsondunn.com/eu-ai-act-omnibus-agreement-postponed-high-risk-deadlines-and-other-key-changes/>
- [9] M. Ben-Daya, E. Hassini, and Z. Bahrour, "Internet of Things and supply chain management: a literature review," *Int. J. Prod. Res.*, vol. 57, no. 15–16, pp. 4719–4742, 2019. [Online]. Available: https://www.researchgate.net/publication/321131587_Internet_of_things_and_supply_chain_management_a_literature_review
- [10] R. Y. Zhong, X. Xu, E. Klotz, and S. T. Newman, "Intelligent manufacturing in the context of Industry 4.0: A review," *Engineering*, vol. 3, no. 5, pp. 616–630, Oct. 2017. [Online]. Available: <https://www.sciencedirect.com/science/article/pii/S2095809917307130?via%3Dihub>

Saadhu, V.B.R. (2026). Procurement Governance Innovation Failure Patterns in U.S. State Enterprise Systems: A Cross-Jurisdictional Taxonomy and Capability-Gap Mapping Framework (2014–2024). *South African Computer Journal* 38(2), 87–104.

Copyright© the author(s); published under a Creative Commons NonCommercial 4.0 License SACJ
ISSN 1015-7999 (print) ISSN 2313-7835 (online)

- [11] R. S. Sutton and A. G. Barto, Reinforcement Learning: An Introduction, 2nd ed. Cambridge, MA, USA: MIT Press, 2018.
<https://web.stanford.edu/class/psych209/Readings/SuttonBartoIPRLBook2ndEd.pdf>
- [12] R. Parasuraman, T. B. Sheridan, and C. D. Wickens, "A model for types and levels of human interaction with automation," IEEE Trans. Syst. Man Cybern. A, Syst. Humans, vol. 30, no. 3, pp. 286–297, May 2000. [Online]. Available: <https://ieeexplore.ieee.org/document/844354>
- [13] D. Ivanov, "Predicting the impacts of epidemic outbreaks on global supply chains: A simulation-based analysis on the coronavirus outbreak (COVID-19/SARS-CoV-2) case," Transp. Res. E, Logist. Transp. Rev., vol. 136, p. 101922, Apr. 2020. [Online]. Available: <https://www.sciencedirect.com/science/article/pii/S1366554520304300>
- [14] T. Rashid et al., "QMIX: Monotonic value function factorisation for deep multi-agent reinforcement learning," in Proc. 35th Int. Conf. Mach. Learn. (ICML), Stockholm, Sweden, 2018, pp. 4295–4304. <https://arxiv.org/abs/1803.11485>
- [15] P. Leitão et al., "Smart agents in industrial cyber-physical systems," Proc. IEEE, vol. 104, no. 5, pp. 1086–1101, May 2016. [Online]. Available: <https://ieeexplore.ieee.org/document/7437398>
- [16] A. Brintrup et al., "Predicting hidden links in supply networks with multi-order Markov chains," J. Oper. Manag., vol. 65, no. 7, pp. 735–750, Oct. 2019. [Online]. Available: https://www.researchgate.net/publication/321974090_Predicting_Hidden_Links_in_Supply_Networks
- [17] M. Christopher and H. Peck, "Building the resilient supply chain," Int. J. Logist. Manage., vol. 15, no. 2, pp. 1–14, 2004. [Online]. Available: <https://www.emerald.com/ijlm/article-abstract/15/2/1/133299/Building-the-Resilient-Supply-Chain?redirectedFrom=fulltext>
- [18] A. Dolgui, D. Ivanov, and M. Rozhkov, "Does the ripple effect influence the bullwhip effect? An integrated analysis of structural and operational dynamics in the supply chain," Int. J. Prod. Res., vol. 58, no. 5, pp. 1285–1301, 2020. [Online]. Available: <https://www.tandfonline.com/doi/full/10.1080/00207543.2019.1627438>
- [19] S. Yao et al., "ReAct: Synergizing reasoning and acting in language models," in Proc. 11th Int. Conf. Learn. Representations (ICLR), Kigali, Rwanda, 2023. <https://arxiv.org/pdf/2210.03629>
- [20] L. Monostori, "Cyber-physical production systems: Roots from manufacturing science and technology," at-Automatisierungstechnik, vol. 63, no. 10, pp. 766–776, Oct. 2015. [Online]. Available: <https://www.degruyterbrill.com/document/doi/10.1515/auto-2015-0066/html>
- [21] European Parliament and of the Council, "Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence (Artificial Intelligence Act)," *Official Journal of the European Union*, L series, 12 July 2024. [Online]. Available: <https://eur-lex.europa.eu/eli/reg/2024/1689/oj/eng>