

The Delegation Decision: A Product Manager's Framework for Deciding What to Hand to AI Agents

Aditi Kapildev Vatsa

vatsaditi@gmail.com

Abstract:

As AI Agents become capable of executing multi-step operational workflows autonomously, product managers face a structurally new decision: which tasks to delegate to agents and which to preserve for human judgment. Existing literature addresses how to design human-AI collaboration once the automation boundary is set [1,5] and how to measure return on investment after deployment, but no practitioner framework exists to guide the prior question of where to set that boundary in the first place.

This paper proposes ADAPT, a five-criterion evaluation model that product managers can apply during discovery and scoping to assess whether a workflow task is appropriate for AI Agent delegation. The five criteria are: Ambiguity Tolerance [1], Decision Reversibility [8], Audit and Compliance Exposure [10,11], Process Stability [3], and Trust Threshold [2]. Each criterion is grounded in established peer-reviewed research. The framework is demonstrated through three anonymized worked examples from a Fortune 500 operational technology enterprise and is shown to prevent both over-delegation risk (compounding autonomous errors) and under-delegation risk (unrealized automation value).

Keywords: AI Agents, Agentic AI, Product Management, Delegation Decision, ADAPT Framework, Task Allocation, Human-in-the-Loop, Automation Boundary, AI Governance, Workflow Automation

1. Introduction

1.1 A New Kind of Decision for Product Managers

Product managers have always decided what to build. With the emergence of AI Agents, which can perceive context, reason across options, invoke tools, and iterate across multi-step workflows without human intervention at each step, they must now also decide what to delegate.

This is structurally different from traditional automation scoping. Rule-based automation and robotic process automation (RPA) are deterministic: correctly specified logic yields predictable output. AI Agents operate differently. Baird and Maruping [4] characterize this shift, arguing that a new generation of agentic IS artifacts “can assume responsibility for tasks with ambiguous requirements and for seeking optimal outcomes under uncertainty,” requiring practitioners to revisit how delegation is conceptualized. Sapkota et al. [7] further distinguish AI Agents from generative AI along three defining characteristics: autonomy (minimal human intervention post-deployment), task-specificity (goal-oriented behavior in bounded environments), and reactivity (real-time adaptation to environmental inputs).

The practical stakes are asymmetric. When scoped too broadly, agent errors compound across multi-step chains. Rabanser et al. [8] documented real production failures of this kind across 14 agentic models. These included an AI assistant that deleted an entire production database despite explicit instructions forbidding it, and an agent that executed an unauthorized purchase in violation of a confirmed safeguard. When delegation is scoped too narrowly, automation value goes unrealized and AI investment underperforms. Both failure modes are common. Both stem from the same root cause: the product manager made a scoping decision without a structured framework.

1.2 What Existing Literature Addresses and What It Misses

Adjacent bodies of work address related but distinct questions. Parasuraman, Sheridan, and Wickens [1] provide a foundational taxonomy of automation levels across four function types (information acquisition, analysis, decision and action selection, and implementation) but do not offer criteria for where on that spectrum a specific task belongs. Johnson et al. [5] address coactive design, specifically how to structure human-agent coordination after the automation boundary is drawn. Lee and See [2] examine trust in human-automation relationships and design for appropriate reliance, but presuppose the delegation decision has already been made. Bainbridge [3] warns of the ‘ironies of automation,’ the paradox by which automating most of a system leaves operators responsible for precisely the tasks they are least prepared to handle, but provides no upstream scoping tool. The gap this paper addresses is the decision that precedes all of these: how does a product manager determine, during discovery and scoping, which workflow tasks are ready for agent delegation, partial delegation, or human retention?

1.3 Research Objectives

- Characterize the AI Agent delegation decision as a distinct, under-supported product management challenge with specific asymmetric failure modes.
- Propose the ADAPT framework as a five-criterion evaluation tool applicable during discovery and scoping, grounded in established peer-reviewed research.
- Demonstrate ADAPT’s application through anonymized worked examples and show how the framework prevents both over-delegation and under-delegation.

1.4 Theoretical Grounding: IS Delegation

Baird and Maruping’s [4] IS delegation framework provides the core conceptual grounding. They define IS delegation as involving “transfer of authority, responsibility, accountability, clear specifications regarding roles and boundaries, clarity of intent, and mechanisms for establishing trust.” They identify three delegation mechanisms: appraisal (assessing agent capability), distribution (allocating tasks appropriately), and coordination (managing ongoing handoffs). ADAPT operationalizes the appraisal and distribution mechanisms into a practitioner-usable scoping instrument.

Subsequent empirical work confirmed that delegation decisions are dynamic and performance-driven. A longitudinal study of 875 store managers found polarization in delegation willingness: managers who could assess AI capability over time increased delegation, while those who could not reduced AI involvement [9]. This finding motivates a structured appraisal tool at the scoping stage.

1.5 Asymmetric Failure Modes of Mis-scoped Delegation

The appraisal and distribution mechanisms Baird and Maruping identify are difficult to execute well in practice because the delegation decision is not symmetric. Getting it wrong in either direction produces distinct failure modes with different root causes and different organizational consequences, as shown in Table 1.

	Over-Delegation	Under-Delegation
Definition	Assigning agent-inappropriate tasks to fully autonomous execution	Retaining tasks AI Agents could handle reliably, with humans performing the work instead
Primary Risk	Errors compound silently across multi-step chains; irreversible downstream actions triggered before detection	Unrealized productivity gains; continued human error in high-volume repetitive tasks
Evidence	Production failures documented across 14 agentic models, including database deletion and unauthorized purchases [8]	Unrealized AI value despite high enterprise adoption rates, attributed to overly conservative scoping decisions [8]

Table 1: Asymmetric failure modes of AI Agent delegation decisions.

1.6 Why Existing Frameworks Do Not Fully Solve This

Despite a rich body of adjacent work, no existing framework gives a product manager a structured tool for the upstream question: which tasks should be delegated to an AI Agent in the first place, and at what level of autonomy? ADAPT fills that gap. Table 2 shows where each existing framework operates relative to the delegation decision and why none directly addresses it.

Framework	What It Addresses	Gap Relative to ADAPT
Levels of Automation [1]	Taxonomy of human vs. machine control across four function types	No criteria for where a specific task belongs on that spectrum
IS Delegation Framework [4]	Theoretical lens for human-agentic IS artifact delegation mechanisms	No practitioner-usable scoring tool for task-level scoping
Trust in Automation [2]	How trust develops; how to design for appropriate reliance	Presupposes delegation scope has already been decided
Ironies of Automation [3]	Paradoxes of partial automation; operator skill degradation	Diagnosis of failure patterns, not an upstream scoping tool
Coactive Design [5]	Coordination mechanisms after the automation boundary is set	Does not address how to determine scope in the first place

Table 2: ADAPT operates upstream of all existing frameworks listed above.

2. The ADAPT Framework

A critical distinction separates ADAPT from prior human-AI collaboration frameworks: ADAPT evaluates individual tasks, not systems. Prior frameworks—including role taxonomy models defining when AI acts as assistant, executor, collaborator, or delegate, and the four-layer collaboration design framework in [14]—address how to design the collaboration architecture once the automation boundary has already been set. ADAPT operates upstream of that design stage. It answers a prior question: for this specific workflow task, should a boundary be drawn here at all, and if so, at what level of autonomy? The output of ADAPT is not a collaboration design; it is a scoping decision that determines what the collaboration design will need to accommodate. ADAPT and collaboration frameworks are therefore sequential rather than competing: ADAPT determines task-level scope, collaboration frameworks determine system-level structure.

ADAPT is a five-criterion evaluation model applied to individual workflow tasks during product discovery and scoping. Each criterion is scored 1–3, where 3 indicates the task dimension is agent-ready. Scores are summed for a composite total of 5–15 that drives the delegation decision. The framework is designed to be applied collaboratively with cross-functional stakeholders because scoring disagreements surface the real risk loci. The name reflects its criteria: **A**mbiguity Tolerance, **D**ecision Reversibility, **A**udit and Compliance Exposure, **P**rocess Stability, **T**rust Threshold. Together these criteria reflect the core principle that delegation scope must adapt dynamically as context, capability, and trust evolve.

Criterion	Core Question	Score 3 = Most Agent-Ready	Theoretical Grounding
A - Ambiguity	How bounded and well-defined are the task's inputs and success criteria?	Low ambiguity; structured inputs; objective, measurable success criteria	Parasuraman et al. [1]
D - Reversibility	If the agent errors, how quickly is it detected and how costly is correction?	Errors detected quickly; correction low-effort; no irreversible downstream actions	Rabanser et al. [8]
A - Audit/Compliance	Does this task require a named human decision-maker of record under applicable regulation?	No regulatory requirement for human approver; audit trail satisfied by agent logs	EU AI Act [10]; GDPR Art. 22 [11]; Lukács & Váradi [6]
P - Stability	How frequently do the task's rules, inputs, or success criteria change?	Rules change infrequently (annually or less) and predictably	Bainbridge [3]
T - Trust	What is frontline stakeholder confidence in autonomous agent execution for this task?	High trust; stakeholders comfortable with autonomous execution; low override rates in pilot	Lee & See [2]

Table 3: ADAPT Framework summary. Each criterion scored 1–3; composite 5–15 determines delegation decision.

2.1 A - Ambiguity Tolerance

Parasuraman, Sheridan, and Wickens [1] established a foundational taxonomy of automation function types, finding that automation performs most reliably on information acquisition and analysis tasks where inputs are structured and outputs measurable. It performs least reliably on decision and action selection tasks requiring integration of ambiguous, incomplete, or novel contextual information. This distinction maps directly to the ambiguity criterion: tasks with structured, bounded inputs and objective success criteria are strong delegation candidates; tasks requiring situational judgment, implicit stakeholder cues, or novel context are not.

Sapkota et al. [7] provide preliminary evidence that AI Agent failures are most concentrated in tasks requiring causal understanding and stable belief tracking, rather than pattern matching over structured data. This supports treating ambiguity as a primary gating criterion.

- Score 3 (Agent-ready): Inputs are structured and bounded; success criteria are objective and measurable; variation follows predictable patterns.
- Score 2 (Partial delegation): Moderate ambiguity; agent handles standard cases and escalates outliers for human review.
- Score 1 (Human retention): Task regularly involves novel situations, implicit stakeholder expectations, or judgment that cannot be reliably encoded in agent instructions.

 **Practitioner Signal:** *If a skilled new employee could learn this task reliably from a written SOP in under one week, ambiguity is likely low enough for agent delegation.*

2.2 D - Decision Reversibility

Rabanser et al. [8] propose a reliability science for AI agents organized around four dimensions: consistency, robustness, predictability, and safety. Their safety dimension specifically addresses whether failures are “bounded in severity” and systems “fail predictably.” Their empirical study (currently available as a preprint) documented production failures, including a database deletion and an unauthorized purchase, in tasks where errors triggered irreversible downstream actions before detection. Critically, the study finds that recent model capability gains have yielded only small reliability improvements, meaning reversibility must be evaluated for the task context and cannot be inferred from benchmark scores.

Reversibility is a function of two sub-dimensions: detection lag (how quickly an error surfaces) and correction cost (how expensive it is to undo). High-reversibility tasks bound the cost of agent error. Low-reversibility tasks require human oversight at the decision point regardless of agent capability.

- Score 3: Errors surface quickly through routine monitoring; correction requires minimal effort; no irreversible downstream actions triggered.
- Score 2: Errors may take days to surface but are correctable with moderate effort; downstream impact is contained.
- Score 1: Errors may be silent for extended periods; correction is expensive or impossible; task output directly triggers irreversible actions.

Practitioner Signal: *Ask: 'If the agent gets this wrong on a Friday evening, what is the state of the world by Monday morning?'*

2.3 A - Audit and Compliance Exposure

GDPR Article 22 [11], enforceable since 2018, restricts fully automated decision-making that significantly affects individuals, requiring human intervention mechanisms in regulated contexts. The EU AI Act [10], which entered force in August 2024 with high-risk system enforcement beginning August 2026, imposes comprehensive requirements on AI systems used in employment, credit, education, and critical infrastructure. These include mandatory human oversight mechanisms, technical documentation, and audit trail obligations. Organizations face fines of up to €35 million or 7% of global annual revenue for serious violations.

Lukács and Váradi [6] demonstrate that GDPR compliance increasingly requires not just transparency about automated decisions but demonstrated human accountability in the decision chain. Absent a human of record, agent decisions are difficult to defend in post-incident regulatory review even when the outcome was factually correct.

- Score 3: No regulatory oversight; no named human decision-maker required; audit obligations satisfied by agent action logs.
- Score 2: Moderate compliance relevance; agent executes but human must review and approve before output is externally committed.
- Score 1: Regulation requires human accountability; legal constraints prohibit fully autonomous execution.

Hard override rule: A score of 1 on this criterion automatically caps the delegation decision at Supervised Automation regardless of total score. Regulatory accountability cannot be offset by high performance on other dimensions.

Practitioner Signal: *Consult legal and compliance before scoring this criterion. Do not assume low exposure without explicit verification.*

2.4 P - Process Stability

Bainbridge's [3] 'ironies of automation' paper, among the most-cited works in human factors research, identified process change as a principal driver of automation failure. Her core finding is that operators left to monitor automated systems lose the skills needed for manual interventions, and those interventions become necessary precisely when the operational environment shifts beyond what the automation was designed for. Applied to AI Agents, this irony compounds: agents trained on the context at deployment time produce confident but outdated outputs when business rules, regulatory requirements, or partner conditions change without corresponding updates.

Highly stable processes are strong delegation candidates because the cost of maintaining agent accuracy is low relative to automation benefit. Dynamic processes may cost more to maintain than they save, reversing the value equation.

- Score 3: Rules and inputs change infrequently (annually or less); changes are predictable and manageable through structured update protocols.
- Score 2: Moderate change frequency (quarterly); agent handles the stable core while humans manage exceptions from rule changes.
- Score 1: Rules and context change frequently (monthly or more); cost of maintaining agent currency likely exceeds automation value.

Practitioner Signal: *How often did the SOP for this task change in the last twelve months? More than twice warrants a score of 2 or lower on this criterion.*

2.5 T - Trust Threshold

Lee and See [2] established that “automation is often problematic because people fail to rely upon it appropriately. Because people respond to technology socially, trust influences reliance on automation. In particular, trust guides reliance when complexity and unanticipated situations make a complete understanding of the automation impractical.” They identify a critical distinction between over-trust (complacency, missed anomalies) and under-trust (excessive verification overhead, forfeited automation benefit). These are the same two failure modes ADAPT is designed to prevent at the scoping stage. It is important to note how this differs from how trust appears in human-AI collaboration design frameworks [14]: there, trust is treated as a design requirement: something the system architecture must accommodate through transparency mechanisms, confidence displays, and override interfaces. In ADAPT, trust is treated as a scoping input, specifically a measure of whether a given task is ready for a given level of autonomy at this point in time, and a signal about what deployment sequencing is appropriate. The same Lee and See research underlies both uses, but they operate at different stages and produce different outputs.

Lee and See further establish that trust evolves through three processes: analytic (evaluating performance evidence), analogical (drawing on prior experience with similar systems), and affective (emotional response to perceived reliability). Trust is not a static property. It is a dynamic attribute of the deployment context that must be reassessed over time. The longitudinal MIS Quarterly study [9] empirically confirms this: over 875 store managers, delegation willingness evolved continuously through performance appraisal, with managers who observed AI capability increasing delegation and those who did not decreasing it.

- Score 3: Stakeholders have direct experience with similar automation; task stakes feel manageable; transparency mechanisms support calibration; override rates in pilot were low.
- Score 2: Mixed trust levels; phased delegation is recommended, starting in agent-as-assistant mode and expanding autonomy as trust builds through demonstrated performance.
- Score 1: Significant stakeholder anxiety about autonomous execution; override rates in pilot are high; trust-building through structured observation must precede autonomous delegation.

Practitioner Signal: *Before deployment, survey frontline users: ‘Would you be comfortable with this task being completed by an AI without your review?’ The distribution of responses indicates where the team sits on this dimension.*

3. Applying ADAPT: Scoring Method and Decision Rules

3.1 Composite Scoring and Delegation Outcomes

Each criterion is scored 1–3 and scores are summed for a composite total of 5–15. The composite determines the delegation decision as shown in Table 4. Practitioners should evaluate the Audit and Compliance Exposure criterion before summing the total: a score of 1 on that criterion triggers the hard override and renders the composite score irrelevant to the final decision.

Total Score	Delegation Decision	Operational Implication
13–15	Full Delegation	Task is appropriate for autonomous AI Agent execution with standard operational monitoring
9–12	Partial Delegation	Agent handles the task with defined human checkpoints at identified risk points
6–8	Supervised Automation	Agent assists and prepares output; human decides and approves before any action is committed
5	Human Retention	Task is not ready for agent delegation; revisit after capability improvement or context change

Table 4: ADAPT composite scores and delegation decisions.

Hard override: A score of 1 on Audit and Compliance Exposure caps the decision at Supervised Automation regardless of total score. Regulatory accountability is not a tradable dimension.

3.2 Collaborative Discovery Session Process

- Map the workflow into discrete task units, scoring each independently.
- All participants score each criterion independently before group discussion to prevent anchoring.
- Convene focused discussion on scoring disagreements; these reliably indicate where real delegation risk lies.
- Document the rationale for each final score alongside the score itself, as this becomes the audit trail for the delegation decision.
- Establish a formal reassessment date (suggested: 90 days post-deployment) to revise scores based on live agent performance, override rates, and trust levels.

4. Worked Examples

The following examples are drawn from anonymized operational workflows within a Fortune 500 operational technology enterprise. Identifying details have been removed; the task structures, scoring rationales, and delegation outcomes reflect actual scoping decisions made during product discovery.

4.1 Example A: Vehicle Status Verification for Outbound Communication

Task description: Before initiating an outbound call to a vehicle owner or partner body shop to confirm pickup readiness, verify the vehicle's current status across multiple internal data sources and determine whether the call should proceed.

Criterion	Score	Rationale	Ref.
Ambiguity	3	Status fields are structured and bounded; decision logic is rule-based and testable	[1]
Reversibility	3	If agent reads status incorrectly, the call is abandoned and retried; no irreversible downstream action committed	[8]
Audit/Compliance	3	No regulatory requirement for human decision-maker on internal status verification in this context	[10]
Process Stability	2	Status field logic updated quarterly following partner rule changes; requires managed update protocol	[3]
Trust Threshold	2	Operations teams comfortable with automation but request override visibility and exception reporting during rollout	[2]
Total	13	DECISION: Full Delegation. Implement with managed quarterly update protocol and visible exception log.	

Table 5: ADAPT scoring for vehicle status verification.

4.2 Example B: Title Document Approval for Regulatory Processing

Task description: Review a vehicle title document package and approve it for downstream regulatory processing, a step that carries legal accountability in multiple jurisdictions.

Criterion	Score	Rationale	Ref.
Ambiguity	2	Most cases follow standard patterns; jurisdictional edge cases and lienholder discrepancies require expert judgment	[1]
Reversibility	1	Approval triggers irreversible downstream title filings; correction is expensive and time-critical	[8]
Audit/Compliance	1*	Multiple jurisdictions require a named human approver of record. HARD OVERRIDE TRIGGERED. Agent-only approval is legally impermissible.	[11][6]
Process Stability	2	Jurisdictional title regulations change periodically; stability is moderate	[3]

Trust Threshold	2	Title specialists skeptical of autonomous approval without demonstrated accuracy track record in this document category	[2]
Total	8*	DECISION: Supervised Automation (hard override applies). Agent pre-processes the package, flags discrepancies, prepares recommendation. Human retains approval authority.	

Table 6: ADAPT scoring for title document approval. *Hard override applies regardless of total score.

4.3 Example C: Exception Routing in a Multi-Step Claims Workflow

Task description: When an incoming claim triggers an exception flag, determine the correct internal routing (which team and which priority level) based on claim type, partner identity, and current queue status.

Criterion	Score	Rationale	Ref.
Ambiguity	2	Majority of routing follows rule-based patterns; minority of novel claim types requires specialist judgment and escalation	[1]
Reversibility	2	Misrouting causes delay and rework but is correctable within hours; no irreversible external actions triggered	[8]
Audit/Compliance	3	No regulatory accountability requirement on internal routing decisions in this jurisdiction; audit satisfied by routing log	[10]
Process Stability	3	Routing rules stable throughout the year; changes infrequent, well-documented in change management records	[3]
Trust Threshold	3	Claims teams observed agent-assisted routing succeed in a structured pilot; trust established; override rates within acceptable range	[2]
Total	13	DECISION: Full Delegation with anomaly escalation protocol for novel claim types and 90-day ADAPT reassessment.	

Table 7: ADAPT scoring for exception routing in claims workflow.

5. Discussion

5.1 ADAPT as an Upstream Pre-Mortem

ADAPT functions as a structured anticipation tool that surfaces conditions under which a delegation decision is likely to fail before those conditions materialize in production. In this respect it shares the spirit of the pre-mortem technique [12], which improves risk identification by asking teams to assume failure and explain why. The key structural difference is that ADAPT is criterion-specific, quantified, and

grounded in established theory across five known risk dimensions; a pre-mortem is open-ended and draws on team knowledge. Used together, they are complementary: ADAPT provides systematic diagnostic of task-level delegation risk, while a pre-mortem surfaces organizational or systemic concerns outside the five criteria. This connection to the author's prior work on pre-mortem methodology [13] positions ADAPT as part of a coherent upstream risk management toolkit.

5.2 The Product Manager as Delegation Architect

Bainbridge [3] observed that the ironies of automation arise from incomplete thinking about where automation boundaries are set, a diagnosis that applies with even greater force to AI Agents. Baird and Maruping [4] argue that IS delegation requires ongoing coordination mechanisms as agent performance data accumulates. Lee and See [2] establish that trust calibration is continuous, not a one-time deployment assessment. Together, these bodies of work point to the same conclusion: the delegation decision is an ongoing product management responsibility. Product managers who treat delegation as dynamic and data-informed will realize more durable AI automation value than those who treat it as a one-time launch choice. This means using override rates, escalation frequency, and trust signals as leading indicators of mis-scoping, and scheduling ADAPT reassessments at 90-day intervals.

This reframes the PM role from delivery manager to delegation architect: not just the person who decides what to build, but the person who continuously calibrates the boundary between human and agent judgment as operational context evolves.

5.3 The Trust Dimension Is Systematically Underweighted

In practice, the trust criterion is the most frequently underestimated dimension in agent deployment decisions. Technical readiness (low ambiguity, high reversibility, low compliance exposure, stable processes) is necessary but not sufficient. Lee and See [2] established that trust becomes critical “when complexity and unanticipated situations make complete understanding impractical”, and this is precisely the context of most enterprise AI Agent deployments. A task that scores highly on the first four ADAPT criteria but scores 1 on Trust is likely to produce high override rates, workarounds, shadow processes, and eventual abandonment regardless of technical capability. This observation complements but is distinct from the treatment of trust in human-AI collaboration design frameworks [14]. That work addresses trust as a design challenge: given that a task has already been delegated, how should the system be built to support appropriate reliance through transparency mechanisms and override interfaces? ADAPT addresses trust as a scoping challenge: given the current level of stakeholder trust in this task context, what level of autonomy is the deployment ready for right now? The implication is that low trust scores should be treated as deployment sequencing signals, not blockers. Tasks with Trust scores of 1 or 2 should begin in agent-as-assistant mode, allowing stakeholders to build confidence through direct observation before autonomy expands. The longitudinal MIS Quarterly study [9] confirms this progression occurs when PMs create structured conditions for it, but stalls when delegation is treated as all-or-nothing.

5.4 Limitations and Future Research

ADAPT is a judgment-based practitioner framework, not an algorithmic formula. Scores reflect the assessor's understanding of the task context, and inter-rater reliability is higher when scoring is

collaborative rather than individual. The framework presupposes but does not replace agent capability assessment. Evaluating what a specific agent can reliably do is a prerequisite step. ADAPT is also designed for single-agent task-level evaluation; multi-agent architectures with interdependencies between agents require extensions beyond this paper's scope.

Future empirical work should test whether tasks scoring poorly on ADAPT empirically experience higher rates of delegation failure in production. This is a testable hypothesis that would establish predictive validity and support the framework's refinement through real deployment feedback.

6. Conclusion

The rise of AI Agents capable of executing multi-step operational workflows creates a structurally new product management challenge: the delegation decision. Deciding what to hand to an agent, and what to preserve for human judgment, is consequential, asymmetric in its failure modes, and, to the authors' knowledge, currently unsupported by any practitioner-level decision framework.

The ADAPT framework addresses this gap through five structured, theory-grounded criteria: Ambiguity Tolerance [1], Decision Reversibility [8], Audit and Compliance Exposure [10,11,6], Process Stability [3], and Trust Threshold [2]. Together these determine whether a workflow task is ready for full delegation, partial delegation, supervised automation, or human retention. Applied during discovery and scoping, ADAPT shifts the delegation decision from intuition to structured judgment. It complements IS delegation theory [4], human-AI collaboration design [5], and trust calibration research [2] by operating upstream of all of them, at the moment the boundary is drawn rather than after it is set. The most consequential product decisions of the next decade may not be what to build, but what to delegate, to whom, and under what conditions. Future empirical research testing whether tasks scoring poorly on ADAPT systematically exhibit higher rates of production failure would establish predictive validity and support the framework's continued refinement through real deployment evidence.

REFERENCES:

- [1] 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*, 30(3), 286–297. <https://doi.org/10.1109/3468.844354>
- [2] Lee, J. D., & See, K. A. (2004). Trust in automation: Designing for appropriate reliance. *Human Factors*, 46(1), 50–80. https://doi.org/10.1518/hfes.46.1.50_30392
- [3] Bainbridge, L. (1983). Ironies of automation. *Automatica*, 19(6), 775–779. [https://doi.org/10.1016/0005-1098\(83\)90046-8](https://doi.org/10.1016/0005-1098(83)90046-8)
- [4] Baird, A., & Maruping, L. M. (2021). The next generation of research on IS use: A theoretical framework of delegation to and from agentic IS artifacts. *MIS Quarterly*, 45(1), 315–341. <https://doi.org/10.25300/MISQ/2021/15882>
- [5] Johnson, M., Bradshaw, J. M., Feltovich, P. J., Jonker, C. M., van Riemsdijk, M. B., & Sierhuis, M. (2014). Coactive design: Designing support for interdependence in joint activity. *Journal of Human-Robot Interaction*, 3(1), 43–69. <https://doi.org/10.5898/JHRI.3.1.Johnson>

- [6] Lukács, A., & Váradi, S. (2023). GDPR-compliant AI-based automated decision-making in the world of work. *Computer Law & Security Review*, 50, 105848. <https://doi.org/10.1016/j.clsr.2023.105848>
- [7] Sapkota, R., et al. (2025). AI Agents vs. Agentic AI: A conceptual taxonomy, applications and challenges. *arXiv:2505.10468*. <https://arxiv.org/abs/2505.10468>
- [8] Rabanser, S., Kapoor, S., Kirgis, P., Liu, K., Utpala, S., & Narayanan, A. (2026). Towards a science of AI agent reliability. *arXiv:2602.16666*. <https://arxiv.org/abs/2602.16666>
- [9] Liu, J., Yue, W. T., Leung, A. C. M., & Zhang, X. (2025). Find the good. Seek the unity: A hidden Markov model of human-AI delegation dynamics. *MIS Quarterly*, 49(3), 1185–1204. <https://doi.org/10.25300/MISQ/2024/18232>
- [10] European Union. (2024). Regulation (EU) 2024/1689-EU Artificial Intelligence Act. Official Journal of the European Union. <https://eur-lex.europa.eu>
- [11] European Parliament. (2016). General Data Protection Regulation (GDPR), Article 22. Official Journal of the European Union, L 119, 1–88.
- [12] Klein, G. A. (2007). Performing a project premortem. *Harvard Business Review*, 85(9), 18–19.
- [13] Vatsé, A. K. (2026). Pre-mortems in product management: Psychological safety and risk anticipation in technology teams. *The American Journal of Engineering and Technology*, 8(01), 186–195. <https://doi.org/10.37547/tajet/v8i1-318>
- [14] Vatsé, A. K. (2026). A framework for human-AI collaboration in operational teams: Applications in manufacturing and supply chain. *International Journal of AI, BigData, Computational and Management Studies, ICAIDSCT26*, 287–294. <https://doi.org/10.63282/3050-9416.ICAIDSCT26-133>

Appendix A: ADAPT Scoring Sheet (Reusable Template)

Workflow / Task Name: _____

Assessed by: _____ Date: _____

Cross-functional reviewers: _____

Criterion	Score (1–3)	Key Rationale	Dissenting Views / Notes
A - Ambiguity Tolerance			
D - Decision Reversibility			
A - Audit & Compliance		Hard override if score = 1	
P - Process Stability			
T - Trust Threshold			
TOTAL	/15		

Hard override triggered? (Score of 1 on Audit & Compliance) Yes / No

Delegation Decision: ☐ Full Delegation (13–15) ☐ Partial Delegation (9–12) ☐ Supervised

Automation (6–8) ☐ Human Retention (5)

90-Day Reassessment Date: _____