

Seeing What the Dashboard Misses

Machine-Assisted Perception and Hybrid Cognition in Project Risk Management

John M. Aaron

Founder and Principal, Milestone Planning and Research, Inc.

john.aaron@milestoneplanning.net · www.ratio-weekly.com

Publication and Use Notice

© 2026 John M. Aaron. All rights reserved.

Authorship and Responsibility. The author assumes responsibility for the conceptual framework, empirical interpretation, and conclusions presented here. Analytical software and large language models were used as research, coding, interpretive, and drafting aids under human supervision.

PRIMMS® is a registered trademark of Milestone Planning and Research, Inc.

ABSTRACT

Despite decades of methodological refinement, project failure rates remain stubbornly high. Root cause analysis consistently points to a shared structural deficit: project management tools record task status but cannot infer the structural conditions that precede visible failure. This paper introduces PRIMMS — Project Risk Identification Measurement and Mitigation System, augmented by a Large Language Model — as a hybrid cognition architecture designed to close this intelligence gap.

PRIMMS combines a compiled MATLAB application that computes deterministic Bayesian weight-of-evidence (WoE) signals from schedule data, Voice of Team (VoT) distributed sentiment, and documentary artifacts, with a five-layer Cognitive hierarchy prompt architecture that routes those signals through a Large Language Model (LLM) for interpretation, pattern recognition, trajectory projection, and executive briefing generation. The system classifies projects against six empirically grounded failure archetypes, produces recovery trajectory projections at +2, +4, and +8-week horizons, and generates a sponsor-ready governance document — while preserving an inviolable human authority boundary at every decision point.

A second contribution, developed subsequent to the core architecture, extends the OCC-5C quality framework recursively: the same five dimensions (Consistency, Coherence, Contradiction probe, Confabulation check, Calibration) that constitute the external quality methodology are applied as a structured adversarial challenger to PRIMMS's own phase-exit recommendations before they reach the project sponsor, creating a formal quality gate at each project phase boundary governed by the same Bayesian evidence thresholds that govern the broader architecture. A related extension applies Inference to the Best Explanation (Lipton, 2004) to root-cause reasoning itself: the archetype classification is generalized into a Reasoning Map of ranked, falsifiable competing hypotheses, and the decision-gate structure into a sequenced Decision Pathway that tests each hypothesis before committing to a course of action. Retrospective application across a portfolio of completed enterprise transformation programs demonstrates that the architecture consistently produces decisive evidence classifications and identifies actionable recovery windows not surfaced by conventional monitoring.

An illustrative example drawn from one completed business transformation program produces a schedule jeopardy classification at 65.5 db (decisive evidence strength) and identifies a four-day actionable recovery window that was not surfaced through the conventional monitoring in use at the time. PRIMMS is offered as the first applied instance of a broader mission — amplifying human cognition with AI for business decision making, not the whole of that mission — and within it the architecture advances a specific argument: that the most consequential application of AI in project management is not task automation but structural perception — the capacity to detect what is actually happening inside a project before it becomes visible in delivery output. The paper further identifies two foundational

constraints on machine-assisted cognition: the Gödel-Turing constraint, which establishes that inductive systems cannot regulate themselves from within, and the Zero constraint, which establishes that systems accumulating signal complexity without formal governing equations collapse into indeterminacy. The governing equations of PRIMMS — the Bayesian weight-of-evidence accumulator, the project production function, and the issue closure ratio θ — are the mechanism by which the Zero constraint is satisfied: they bound what the LLM is required to treat as settled, so that its interpretive capacity is directed at what genuinely requires judgment.

Keywords: project risk management; hybrid cognition; Bayesian weight of evidence; machine-assisted perception; artificial intelligence; Voice of the Team; failure archetypes; large language models; project governance; Zero constraint; governing equations; issue lifecycle management; AI; phase-gate quality review; adversarial falsification; OCC-5C framework; recursive governance architecture; reasoning map; inference to the best explanation; decision pathway

Table of Contents

1. Introduction	5
1.1 Origins and Development History	6
2. Literature Review	8
2.1 The Persistent Project Failure Problem	8
2.2 AI Applications in Project Management	9
2.3 Hybrid Cognition and Bayesian Perception.....	10
3. The PRIMMS Architecture	13
3.1 System Overview	13
3.2 Signal Extraction and Bayesian WoE Computation	15
3.3 Failure Archetype Classification	17
3.4 Recovery Signal Architecture	18
3.5 The Phase-Gate Quality Review Architecture.....	18
4. The Five-Layer Cognitive hierarchy.....	21
4.1 Architecture Rationale	21
4.2 Layer Descriptions.....	22
4.3 Audit and Traceability	23
4.4 The Reasoning Map: Competing Hypotheses Under Inference to the Best Explanation	24
4.5 The Decision Pathway: Sequenced Falsification Before Commitment.....	25
5. Retrospective Validation	26
5.1 Case Context.....	26
5.2 Illustrative Example: Signal Set at the 8 August 2025 Snapshot	27
5.3 Cognitive hierarchy Analysis: Illustrative Example	28
5.4 Comparison Against Known Outcomes: Illustrative Example	29
6. Theoretical and Practical Contributions	30
6.1 Advancing the Hybrid Cognition Framework.....	30
6.2 The Intelligence Gap Addressed	31
6.3 Design Intent vs. Implementation Fidelity	32
6.4 Self-Reference, Gödel, and the Structural Case Against Autonomous AI	34
6.5 Command, Control, and the Irreducible Human Boundary	35
6.6 The Zero Constraint: Complexity Collapse and the Role of Governing Equations	36
6.7 The Recursive Application of OCC-5C: The Framework Applied to Itself	39
6.8 Beyond Project Management: Management in General as a Cybernetic, Management-by-Exception Process.....	41
7. Limitations and Future Work.....	41
7.1 Current Limitations	41
7.2 Future Work.....	46
8. Conclusion.....	49
References.....	50
Author Note	53

Glossary of Terms.....	53
------------------------	----

1. Introduction

Project management has been a subject of systematic study for more than seventy years. The discipline has produced a rich body of method: from PERT/CPM scheduling frameworks of the 1950s through the Earned Value Management standards codified in the 1960s and 1970s, from the Agile Manifesto of 2001 through contemporary scaled frameworks such as SAFe and LeSS. Standards organizations — the Project Management Institute (PMI), PRINCE2, ISO 21502 — have invested substantial resources in defining best practice and professionalizing the field.

This is not merely a local, project-level concern. Aaron, Fischer, and Ricordati (2015) situate project management within the macroeconomic identity $Y = C + I + G$, arguing — following the Austrian capital-structure tradition of Hayek as elaborated by Garrison — that project management's core economic function is to reduce the perceived risk of capital investment (I), lowering the cost of investment and speeding the rate at which innovation moves through an economy's stages of production. That function rests on a much older premise, traceable to Hume: that a commitment made is a commitment kept, and that the entire apparatus of contract and investment exists only because parties can rely on that keeping. A schedule slip that is quietly absorbed rather than surfaced is, therefore, not a private scheduling inconvenience; it is a broken commitment whose cost compounds outward — into investor confidence, capital allocation, and the pace of innovation deployment — well beyond the project team that absorbed it quietly. The intelligence gap this paper addresses is, in that sense, a gap with macroeconomic consequence, not merely a local governance nuisance.

The outcome, measured against the problem the field set out to solve, is disappointing. PMI's Pulse of the Profession surveys consistently find that fewer than half of projects complete on time, within budget, and with originally promised scope. The Standish Group CHAOS Report, tracking software project outcomes since 1994, finds success rates below 40% by its criteria after three decades of tool improvement and methodology development. McKinsey research on large information technology projects documents an average of 45% cost overrun, 7% schedule overrun, and 56% reduction in delivered value relative to predictions (Bloch, Blumberg, & Laartz, 2012). These numbers do not reflect ignorance or neglect. They persist alongside genuine improvements in practice.

The structural argument advanced throughout this paper — developed in detail in Aaron (2026, Chapter 1; unpublished manuscript) — is that project failure is primarily a perception and orientation problem, not a planning or tool problem. Projects fail not because managers lack scheduling software or risk registers. They fail because the signals of structural deterioration accumulate invisibly, optimism bias suppresses their interpretation, and the window for low-cost correction closes before leadership recognizes it. The commercial tool landscape, characterized by increasingly sophisticated task-tracking and collaboration platforms, has not addressed this problem — and, given its architectural orientation toward record-keeping rather than inference, cannot address it without a fundamentally different approach.

This paper describes PRIMMS® (Project Risk Identification Measurement and Mitigation System, augmented by a Large Language Model) as a concrete instantiation of a different approach: hybrid cognition. The term, drawn from cognitive science and developed in the project management context by Aaron (2026), denotes architectures in which machine intelligence extends human perception without replacing human judgment. PRIMMS extends the project manager's perceptual range across three signal modalities — quantitative schedule data, distributed team sentiment, and governance documentation — and routes the fused signal set through a hierarchical LLM interpretation architecture that produces actionable governance outputs. At every decision point, the human project manager retains full authority. The machine orients; the humans decide.

The contribution of this paper is threefold. First, it presents the PRIMMS architecture in sufficient technical detail to enable replication and extension. Second, it articulates the theoretical grounding — in Bayesian inference, cognitive neuroscience, and decision theory — that distinguishes machine-assisted perception from conventional AI-powered project management tools. Third, it reports an ongoing retrospective validation program and presents one illustrative example in detail, demonstrating the system's analytical capability against a completed project case with quantified evidence strength and known outcomes.

The paper is organized as follows. Section 2 reviews the relevant literature across three domains: project management failure research, AI applications in project management, and hybrid cognition theory. Section 3 presents the PRIMMS architecture in detail. Section 4 describes the five-layer Cognitive hierarchy prompt design. Section 5 reports the retrospective validation program and presents one illustrative example in detail. Section 6 discusses theoretical implications and practical contributions. Section 7 addresses limitations and future work.

1.1 Origins and Development History

PRIMMS® did not emerge from an academic research program. It was built in direct response to a practitioner problem: the author's experience managing large-scale enterprise resource planning (ERP) implementations and upgrades repeatedly produced the same disorienting pattern. Risk signals were present in the project environment — in team conversations, in subtle changes in stakeholder behavior, in the texture of weekly status reports — but the conventional tools available could not surface or integrate those signals into a coherent governance picture. Projects would appear nominally green on the standard reporting dashboard, then enter crisis within weeks. The gap between what project participants knew and what the formal reporting system showed was structural, not incidental, and no available tool addressed it.

The first PRIMMS design, developed in 2008, was a prediction market mechanism. Drawing on research in parimutuel wagering systems and the observed accuracy of prediction markets in other forecasting contexts, the author designed and filed a patent application for a system that invited project stakeholders to wager on scenario-based project outcomes using parimutuel odds. The core insight was that financial incentives would elicit more honest signals from team members than conventional surveys, and that the price ratios produced by parimutuel wagering would

provide a continuous, market-derived estimate of project risk. The PRIMMS® name (Project Risk Identification Measurement and Mitigation System) and trademark were registered at this time. The patent application was ultimately rejected; and the wagering mechanism itself, while theoretically sound, proved organizationally impractical — asking employees to bet on project failure outcomes created cultural resistance that undermined adoption.

The parimutuel mechanism's failure as an organizational tool, however, pointed toward a more tractable design. The price ratio at the heart of the wagering system is mathematically equivalent to a likelihood ratio — and likelihood ratios are the fundamental building block of Bayesian inference. This realization, combined with the author's parallel interest in the wartime Bayesian decryption methods developed by I. J. Good and Alan Turing at Bletchley Park, produced the key architectural shift: replacing parimutuel betting prices with Bayesian weight-of-evidence expressed in decibans. Rather than asking team members to bet, the revised system asked them to vote on outcome scenarios using a structured six-point scale. The frequency and intensity of mentions across risk categories in those votes and in project documents could then be converted directly into likelihood ratios, summed in deciban form, and compared to empirically established thresholds — without requiring any financial transaction. PRIMMS® went live in this form, using Bayesian weight of evidence from Voice of the Team (VoT) distributed sentiment as its core inference mechanism.

From 2008 through the present, PRIMMS has been deployed across a substantial portfolio of complex enterprise resource planning, business transformation and military aerospace projects. This deployment history provides the empirical foundation for the Bayesian WoE architecture described in Section 3: the signal weights, risk category benchmarks, and threshold calibrations reflect accumulated evidence from real project outcomes across more than seventeen years of operational use. The VoT mechanism and the Bayesian classification layer are, in this sense, not experimental — they are the mature core of a production system with a documented track record. Formal analysis of the 20 risk categories that emerged from text mining of project risk logs across 24 ICT projects (Aaron, 2017 internal working paper) provides the empirical grounding for the risk signal architecture, including statistical identification of the subset of risk categories most predictive of project schedule performance.

The more recent additions to PRIMMS — semantic document analysis and the Cognitive hierarchy prompt architecture grounded in neuroscientific models of cortical processing identified by Fuster (2008) and Glimcher (2011) — represent a qualitatively different development stage. These components extend the system's perceptual range into unstructured documentary text and integrate the multi-modal signal set through a five-layer LLM interpretation architecture. Their validation status differs materially from the Bayesian core: while the Bayesian WoE layer has been tested prospectively across full project lifecycles, the semantic analysis and Cognitive hierarchy components have been validated retrospectively against completed project phases rather than across complete prospective projects. This distinction in validation depth is explicitly acknowledged in Section 7.1 and is the primary motivation for the prospective deployment study identified as the highest-priority future work item.

2. Literature Review

2.1 The Persistent Project Failure Problem

The literature on project failure is extensive and largely convergent. Flyvbjerg (2017) documents systematic optimism bias and strategic misrepresentation across infrastructure megaprojects spanning five continents and six decades, finding cost overruns in 86% of cases with a mean overrun of 28%. Budzier and Flyvbjerg (2011) identify a distinct population of 'black swans' — IT projects with cost overruns exceeding 200% — and demonstrate that this distribution is fundamentally different from the rest, suggesting qualitatively different causal mechanisms. Maylor, Vidgen, and Carver (2008) introduce the concept of managerial complexity — the degree of interaction among project system elements — as a predictor of project performance independent of technical scope.

Running across this literature is a consistent finding: the signals of structural deterioration are typically present weeks or months before they become visible in formal project reporting. Kutsch and Hall (2010) document the phenomenon of 'deliberate ignorance' — project managers who recognize emerging risk signals but choose not to act on them due to social and organizational pressures. Williams (2005) identifies the emergence of complex adaptive dynamics in troubled projects, arguing that conventional linear cause-and-effect reasoning cannot adequately model what is actually occurring in project execution. The implication — that project management requires qualitatively different tools for signal detection and pattern recognition — motivates the present work.

2.2 AI Applications in Project Management

The integration of artificial intelligence into project management practice has accelerated markedly in the period following the widespread availability of transformer-based LLMs. Existing reviews (e.g., Pellerin & Perrier, 2019; Niazi, Zhang, & Bhatt, 2023) catalogue applications across the project management lifecycle, including automated scheduling, resource optimization, risk prediction, and progress monitoring. The dominant pattern in this literature is the application of supervised machine learning to historical project datasets to predict outcomes — completion probability, cost overrun magnitude, schedule delay — based on project characteristics measurable at inception.

This line of work produces genuinely useful predictions in contexts where sufficient historical data of consistent quality are available. However, it shares a structural limitation with the conventional PM tools it augments: it operates on data that have already been recorded in formal project systems. The signals that most reliably precede failure — team confidence erosion, perception distortion in leadership, governance gap formation — do not appear in task completion records, cost reports, or formal risk registers. They are distributed across the conversational, documentary, and social texture of project work in ways that historical dataset approaches cannot access.

A second line of work applies LLMs directly to project management tasks: automated meeting summarization, natural language task creation, stakeholder communication drafting, and risk log generation. These applications are productivity improvements — they help project managers perform existing tasks more efficiently. The claim of this paper is more specific: that LLMs have an untapped role as interpretive reasoning engines in a structured signal processing pipeline, performing the natural-language synthesis and pattern recognition tasks that neither deterministic algorithms nor unaided human judgment perform reliably.

Deploying LLMs without this structured pipeline introduces three specific failure modes that are relevant to evaluating existing AI-in-PM applications. The first is Proxy Collapse: when an AI system optimizes for a measurable proxy of project health — schedule variance, cost performance index, milestone completion rate — the proxy becomes the target and decouples from the underlying reality it was meant to track. Teams learn to optimize for the metric; status reports improve while the project deteriorates. The second is Tacit Knowledge Destruction: if project teams stop exercising their own judgment because an AI system is treated as authoritative, the experiential knowledge of practitioners — the felt sense that something is wrong before any dashboard confirms it — atrophies from disuse. When the AI eventually fails in a genuinely novel configuration, the human capacity to fill the gap has been systematically degraded. The third is Recursive Sycophancy: large language models are trained toward user satisfaction, and a project leader who uses an LLM as a primary advisor will receive outputs shaped by the pattern of their own prompts. The model colludes in optimism rather than surfacing it as a risk signal — confidence increases precisely as accuracy decreases. PRIMMS is explicitly designed to counter all three failure modes: multi-modal Bayesian signal fusion defeats proxy gaming; the advisory-only posture preserves practitioner judgment; and the deterministic MATLAB signal layer operates independently of any LLM and has no social incentive to satisfy the project manager, structurally preventing sycophantic drift.

Prior work on AI-assisted project risk specifically includes Kock, Gemünden, and Salter (2021), who examine early warning systems for innovation projects, and Ika, Love, and Pinto (2020), who survey the evidence on project management success factors with attention to the role of contextual intelligence. Neither work addresses the hybrid cognition architecture — the explicit coupling of machine signal processing with human interpretive authority — that is the central contribution of the present paper.

2.3 Hybrid Cognition and Bayesian Perception

The theoretical foundation for PRIMMS draws from two bodies of literature that have not been systematically connected in the project management context: Bayesian brain theories from computational neuroscience, and hybrid cognition frameworks from decision science.

Friston's free energy principle and predictive coding framework (Friston, 2010) propose that biological brains operate as hierarchical Bayesian inference engines, continuously updating generative models of the environment against sensory prediction errors. The hierarchical organization of cortex — from low-level feature detection through increasingly abstract pattern

recognition to executive planning — is interpreted as implementing successive levels of inference across progressively larger temporal and spatial scales. The PRIMMS Cognitive hierarchy, while not a technical implementation of predictive coding, draws on this organizational principle: Layer 1 performs feature detection, Layer 2 gestalt pattern recognition, Layer 3 temporal situation awareness, Layer 4 executive planning, and Layer 5 document generation. The architecture mirrors the computational structure of biological cortex rather than simulating it.

Fuster adds the complementary representational architecture needed to make this correspondence more precise. In his cortical network theory, cognition is organized through distributed, overlapping networks of association — cognits — arranged in hierarchies of increasing abstraction. Relatively concrete perceptual information activates lower-order representations; progressively higher levels integrate broader associations, longer temporal horizons, and increasingly executive or action-oriented representations. The important dynamic is therefore not that cognits themselves flow through the cortex, but that perceptual information flows through recurrent cortical networks, activating and integrating cognits appropriate to the current situation (Fuster, 2008).

Taken together, Fuster and Friston describe complementary aspects of hierarchical cognition. Fuster explains the representational progression from perceptual particulars toward increasingly abstract and executive cognits; Friston explains the recurrent dynamics by which higher-order expectations constrain lower-level processing and sufficiently discrepant sensory evidence propagates upward as prediction error. PRIMMS adopts this pairing as a functional design analogy, not an anatomical claim: organizational artifacts serve as sensory analogs, increasingly abstract project representations are constructed across the Cognitive hierarchy, and new evidence can reinforce or challenge the current orientation.

A further theoretical distinction, developed at length in Aaron (2026) as the semantic tree/semantic forest framework, motivates the specific design choice described in Section 4.4 of requiring the LLM layer to surface tail hypotheses rather than resting on its single most probable output. Human and organizational cognition operates as a semantic tree: a hierarchy of associations pruned by repetition and reinforced by what has previously worked, optimized for fast and reliable action but structurally blind past its own branches. An LLM operates instead over a semantic forest — an unbounded, non-hierarchical space of linguistic association accumulated across the full recorded range of domains, cultures, time periods, and conflicting viewpoints. Fuster's own formulation of cortical function captures the underlying claim precisely: the brain is not a computer, it is a history (Fuster, 2008). Intelligence, on this view, is less a property of raw processing capacity than of accumulated, prunable experience compounded over time — an octopus displays striking moment-to-moment problem-solving ability, yet its short lifespan and its species' lack of any mechanism for transmitting experience across generations mean that its intelligence never compounds; each octopus starts over. Human intelligence compounds because language and writing allow one lifetime's pruned tree to be handed forward to the next, accumulating a depth no individual lifespan could build alone. An LLM, trained on that entire compounded, cross-generational linguistic record, functions as a compressed proxy for a vastly extended experience horizon: it extends a user's effective intelligence for decision making not by

raising raw cognitive capacity but by making available, on demand, an accumulated history no single human lifetime could acquire directly.

The corresponding risk is what this compression reintroduces. A model's default output samples from the modal center of that accumulated history — the most statistically probable articulation of it — which is precisely the least differentiated, least decision-relevant region of the forest. The minority hypothesis, the rare cross-domain analogy, the edge case that never made it into any single organization's pruned tree, sits in the tails, not the mode. Realizing the effective-intelligence gain therefore depends on contextualized prompting that deliberately steers the interpretive layer away from the modal response and toward the tails of its own distribution — the architectural motivation, made concrete in Section 4.4, for requiring the Reasoning Map to name a genuine Tail Rival hypothesis rather than resting on the Current Best explanation alone.

The Bayesian weight-of-evidence framework employed in PRIMMS's signal fusion layer derives from Jeffreys (1961) and the information-theoretic formalization by Good (1950, 1985). The representation of likelihood ratios as logarithmic deciban values — permitting additive accumulation of evidence across independent signals — operationalizes Bayes' theorem in an auditable, inspectable form. This approach was advocated for decision support in high-stakes contexts by MacKay (2003) and has been applied in medical diagnosis and intelligence analysis contexts.

This deciban accumulation mechanism has direct neurobiological precedent, distinct from the cognitive hierarchy correspondence with Fuster (2008) described above. Gold and Shadlen's (2007) review of the neural basis of decision making establishes that when the brain forms a decision from noisy evidence, the relevant neural quantity — the decision variable — behaves precisely as sequential analysis prescribes: it is the running sum of the log-likelihood ratio favoring one hypothesis over another, accumulated over time until a threshold criterion is crossed and a commitment is made. This mechanism, formalized as the sequential probability ratio test, has been observed directly in the firing rates of neurons in cortical area LIP as motion-discrimination decisions unfold, and in analogous accumulator dynamics across several other cortical areas. The correspondence to PRIMMS's architecture is structural rather than metaphorical: deciban accumulation across MATLAB-computed signals is the same weight-of-evidence summation that Gold and Shadlen identify as the brain's own mechanism for converting momentary, noisy evidence into a categorical commitment, and the evidence floor gating Layer 2's escalation to the LLM interpretive layer (Section 4.2) plays the same functional role as the neural decision threshold: evidence accumulates freely below it, and commitment is withheld until it is crossed.

The foundational application of Bayesian methods to project management was established by Aaron (2015), who introduced a microeconomic production function framework for project issue management and derived the first Bayesian quality gate metric for project management practice. That work modeled project output as a function of scheduled activity completions, issue discovery, and issue closure — $Y = f(AC, IC, IO)$ — and introduced the theta parameter (θ = total closed issues / total issues in the pool) as a continuous proxy for project health and phase-exit readiness. Aaron (2015) further showed that the Bayes Factor, expressed in decibans (db, the logarithm of

the likelihood ratio of success to failure as a function of θ), provides a natural decision boundary at even odds, and demonstrated via Monte Carlo simulation how team sentiment from PRIMMS® surveys could be embedded as the Bayesian prior — to be updated by observed issue evidence in the likelihood function. These results constitute the theoretical axioms on which PRIMMS's WoE architecture is built. The deciban-scale evidence weighing used throughout this framework was itself first developed for applied project governance in Aaron, Fischer, and Kulich (2015, Weighing of Evidence), which demonstrated the direct translation of Good and Jeffreys' weight-of-evidence formalism into project decision practice; the 2015 quality-gate paper discussed above is the application of that same weighing framework to the specific problem of phase-exit readiness. Critically, the 2015 framework established that the optimal management decision is not 100% issue closure but the point at which marginal benefit of further closure equals marginal cost — a microeconomic comparative statics result that Mundell-style analysis makes tractable. PRIMMS extends this foundation from quality gate readiness at a single milestone to continuous multi-signal risk monitoring across the full project lifecycle, replacing the scalar θ metric with a multi-dimensional signal vector and replacing the single Bayes Factor curve with a structured db summation across signal modalities. The present paper therefore advances a research program that Aaron (2015) initiated but that required the emergence of capable LLMs and compiled signal-processing tooling to fully operationalize.

On the hybrid cognition side, the fundamental argument is that neither machine intelligence alone nor human judgment alone is sufficient for the perceptual task PRIMMS addresses. Kahneman's (2011) dual-process theory identifies the systematic biases — anchoring, availability, optimism, planning fallacy — that degrade human project risk assessment. Machine signal processing is immune to these biases but cannot perform the contextual interpretation, narrative synthesis, or stakeholder-sensitive communication that project governance requires. The hybrid architecture — machine processing in the signal and classification layers, human authority in the interpretation and decision layers — exploits the comparative advantage of each.

This theoretical integration — Bayesian inference, cognitive hierarchy, and hybrid cognition — provides the design rationale for PRIMMS and distinguishes it from both pure automation approaches and conventional AI productivity tools.

3. The PRIMMS Architecture

3.1 System Overview

PRIMMS is a two-component system. The first component is a compiled MATLAB desktop application that performs deterministic signal extraction, Bayesian WoE computation, failure classification, and structured prompt generation. The second component an LLM such as Claude (Anthropic) or ChatGPT (OpenAI), a Large Language Model accessed through the conversational interface, which receives the structured prompts and produces natural-language interpretation, pattern analysis, recovery trajectory projections, and a formatted executive briefing document.

The design rationale for this two-component architecture reflects a specific principle: computational tasks for which deterministic, auditable algorithms are available should not be delegated to LLMs, and interpretive tasks for which natural language synthesis is required should not be performed by deterministic algorithms. The MATLAB component handles signal computation with complete auditability and reproducibility. The LLM component handles interpretation and synthesis with full transparency about the signal inputs it receives. The division of labor is explicit and enforced by design.

Table 1 presents the five architectural layers of PRIMMS with implementation details and design rationale.

Table 1. PRIMMS Architectural Layers: Implementation and Rationale

Architectural Layer	Implementation Detail	Design Rationale
Signal Acquisition	Schedule metrics: deliverable completion state, days behind plan, jeopardy index, velocity-projected finish date. Voice of the Team (VoT): discrete 1–6 team confidence ratings timestamped and stored as a time series. Documentary corpus: free-text status reports, governance documents, and risk artifacts ingested via keyword and pattern extraction.	Human-generated signals from three independent modalities — quantitative schedule, distributed team sentiment, and governance language — are fused without requiring any modification to existing project management tools.
Inductive Classification	A deterministic Bayesian weight-of-evidence (WoE) classifier maps each signal to a likelihood ratio expressed in decibans (db) using Jeffreys' strength-of-evidence bands: 0–5 db (barely worth mention), 5–10 db (substantial), 10–15 db (strong), 15–20 db (very strong), >20 db (decisive). Signal db values are summed to a prior total representing aggregate risk evidence.	The logarithmic db summation operationalizes Bayes' theorem in an auditable, inspectable form. Every signal weight is fixed and traceable to specific evidence — a deliberate design choice favoring auditability over adaptive learning for enterprise governance contexts.
Failure Archetype Recognition	The classifier maps the convergent signal pattern to one of six defined failure archetypes: (1) Execution Collapse — delivery cadence failing; (2) Perception Distortion — leadership optimism diverges from data; (3) Premature Closure — replan triggered without testing recovery; (4) Scope Drift — boundary erosion without capacity adjustment; (5) Governance Gap — decision authority unclear; (6) Confidence Erosion — team belief in achievability declining.	Archetype identification is a qualitative inference step that cannot be performed by schedule analytics alone. It requires the integration of VoT sentiment data with schedule and documentary evidence — the multi-modal signal fusion that distinguishes PRIMMS from task-tracking platforms.
Cognitive hierarchy (LLM Interpretation)	The MATLAB-computed signal set, VoT data, and documentary signals are formatted into a structured five-layer prompt package. Claude (Anthropic; model version current at time of writing, recorded in the audit trail per Section 4.3) receives these sequentially: Layer 1 — feature validation; Layer 2 — gestalt pattern recognition; Layer 3 — temporal situation awareness with trajectory projections at +2/+4/+8 weeks; Layer 4 — executive planning and COA matrix; Layer 5 — sponsor-	The LLM layer performs the natural-language interpretation, synthesis, and projection that classical algorithms cannot provide. The five-layer structure mirrors the hierarchical organization of biological cortex from sensory processing through executive planning, an architecture grounded in Kanerva (1988) and Friston's

Architectural Layer	Implementation Detail	Design Rationale
	ready briefing document generation (terminated by DOCX_READY). Layer 2 is further extended by a Reasoning Map of competing hypotheses (Section 4.4); Layer 4 is further extended by a sequenced Decision Pathway and wildcard COA option (Section 4.5).	hierarchical predictive coding framework.
Human Governance	Every PRIMMS output is framed as orientation, not instruction. The system's closing statement — 'The machine has oriented. The humans now decide.' — is structurally enforced in every prompt. An immutable audit trail records all signal values, db computations, LLM model version, full prompt and response text, and DOCX_READY detection status.	The human authority boundary is the architectural property that most directly addresses the AI governance concerns catalogued by the Project Management Institute and ISO 21502. The machine extends perception; the project manager owns every decision and every escalation.
Phase-Gate Quality Review	The compute_phase_gate_metrics function aggregates session history from Cortical Log and issue log into phase-level deterministic metrics: longitudinal dB series, archetype frequency table, θ at gate date, and recency-gap computation. Three prompts are generated: PG-1 (Phase Classification), PG-2 (OCC-5C Adversarial Challenger), PG-3 (Phase-Gate Sponsor Document). Governing-equation thresholds ($\theta < 0.70 \rightarrow$ deterministic HOLD; recency gap > 15 dB $\rightarrow$ C5 flag $\rightarrow$ CONDITIONAL EXIT at best; PREMATURE_CLOSURE majority-persistence $\rightarrow$ HOLD) are evaluated before any prompt is issued.	The phase-gate review applies the OCC-5C framework recursively to the system's own phase-exit recommendations, closing the governance loop at the point of highest institutional consequence. The same five quality dimensions that govern AI output assessment govern the phase-exit decision itself, anchored by MATLAB-computed evidence quantities that prevent the recursion from violating the Zero constraint.

3.2 Signal Extraction and Bayesian WoE Computation

The signal extraction process ingests three input modalities: a project schedule workbook (Microsoft Excel or CSV format), Voice of Team (VoT) entries recorded through the application's survey interface, and documentary text ingested by the user. The documentary modality accepts free-text uploads of any project communication or governance artifact — including weekly status reports, steering committee updates, escalation emails, meeting transcriptions, risk and issue logs, corrective action plans, and change request documentation. Users paste or upload this material directly into the application's document tab; no structured format or tagging is required. This design reflects a deliberate architectural choice: the signals most predictive of structural project deterioration — risk language intensifying in status reports, recovery plan language present or absent in governance documents, escalation references in meeting records — reside in unstructured text that no schedule system captures. No custom data collection infrastructure is required beyond the schedule export and communications that project managers typically already produce. The empirical rationale for combining three independent modalities rather than relying on any single signal source is grounded in simulation work reported in Aaron (2019), which

demonstrated that a machine learning system applied to a real-time operational decision problem achieved 73% predictive accuracy using image data alone; accuracy improved to 82% when a second modality (audio signal features) was added; and reached 100% test accuracy when a third modality (text annotations) was incorporated. The implication for PRIMMS is direct: schedule data, VoT sentiment, and documentary text each carry independent information about project risk that the others do not fully capture, and their fusion produces a materially stronger evidence base than any single source can provide.

From the schedule, the MATLAB component computes: schedule state (on track / lagging / behind / critical), activities behind, maximum days behind, jeopardy index (a composite of activities behind and overdue activity-days), rundown statistics comparing actual versus expected versus planned deliverable completions, and a linear velocity model projecting the finish date with associated R^2 fit quality. From VoT entries, the component computes: mean rating, distribution across the 1–6 scale, temporal trend (improving, flat, worsening), and safe-band classification. From documents, the component applies keyword and pattern extraction to generate signals concerning risk language intensity, recovery plan evidence quality, and charter-level governance documentation.

Each detected signal is mapped to a likelihood ratio and expressed in decibans using the Jeffreys (1961) strength-of-evidence convention. Signals are summed to a prior total db value and classified against bands: 0–5 db (barely worth mention), 5–10 db (substantial), 10–15 db (strong), 15–20 db (very strong), >20 db (decisive). Mitigating signals — such as strong recovery plan documentation or ahead-of-plan velocity — carry negative db values, appropriately reducing the aggregate evidence total. Table 2 presents a representative signal set from a pilot deployment.

Table 2. Representative Signal Set: Illustrative Example (Snapshot: 08-Aug-2025)

Signal Name	Source Modality	dB Weight	Trigger Condition
SCHEDULE_STATE_BEHIND	Schedule	10 dB	Project activities are in a Behind state; deviations from plan are present
SCHEDULE_DAYS_BEHIND	Schedule	Scaled (max 10)	Days behind plan; scaled logarithmically
SCHEDULE_JEOPARDY	Schedule	Scaled (max 10)	Composite jeopardy index (days behind + overdue activities)
VOT_MEAN_ELEVATED	VoT	7 dB	Team mean confidence rating ≥ 3.5 on 1–6 scale
CLASSIFIER_JEOPARDY	Inductive	18 dB	Bayesian classifier has reached jeopardy threshold on documentary signals

Signal Name	Source Modality	dB Weight	Trigger Condition
RUNDOWN_GAP_BEHIND	Rundown	Scaled	Actual deliverable completions lag expected count
VELOCITY_FINISH_SLIP	Velocity	Scaled	Linear velocity model projects finish beyond planned end date
RECOVERY_URGENCY_CRITICAL	Recovery	14 dB	Schedule state + days behind converge on critical recovery need
RECOVERY_PLAN_DOCUMENTED_STRONG	Recovery	-6 dB	Strong recovery plan language present in documents (mitigating signal)
RECOVERY_TREND_WORSENING	Recovery	8 dB	VoT mean is rising (worsening) across reporting periods

The PRIMMS WoE architecture extends the deciban (db) framework introduced in Aaron (2015) from the single-milestone quality gate context to continuous multi-signal monitoring across the project lifecycle. In the 2015 formulation, a single Bayes Factor curve mapped the scalar theta parameter (issue closure rate) to a db value, with 0 db marking even odds of phase-exit success. PRIMMS generalizes this structure: each signal contributes an independent db value (the logarithm of its likelihood ratio), and these are summed across modalities to produce an aggregate prior reflecting the weight of evidence for the risk hypothesis across schedule, team sentiment, and documentary domains simultaneously. This multi-signal generalization is what makes continuous lifecycle monitoring tractable — the same mathematical foundation that Aaron (2015) applied to a single quality gate now operates across the full arc of project execution. The VoT 1–6 sentiment survey is the direct descendant of the team sentiment prior described in the 2015 framework, operationalized through PRIMMS® in both cases.

The choice of deterministic WoE scoring over a learned reward policy — the original design intent called for a DDPG/Q-learning agent that would update signal weights from weekly VoT survey responses — was a deliberate tradeoff decision. Live reinforcement learning would enable the system to adapt signal weights to the specific organizational context over time, but it would sacrifice auditability: weights modified by a Reinforcement Learning (RL) agent are no longer directly inspectable or explainable. For an enterprise governance context where every risk signal must be citable to specific evidence, fixed and inspectable weights are the appropriate design. This tradeoff is explicitly acknowledged as a gap versus original design intent and is discussed further in Section 7.

3.3 Failure Archetype Classification

The failure archetype classification translates the aggregate VoE signal into a qualitative pattern diagnosis. The six archetypes represent empirically grounded failure modes identified through project portfolio analysis and the project management literature. Each archetype has a characteristic signal signature:

- Execution Collapse: high schedule jeopardy index; high rundown gap; worsening VoT trend; low velocity fit quality ($R^2 < 0.5$).
- Perception Distortion: divergence between VoT mean and leadership-reported status; wishful thinking signals in forecast accuracy; stable schedule metrics masking deteriorating team sentiment.
- Premature Closure: recovery urgency signal present; strong VoT worsening; replan or extension language prominent in documents without evidence of tested recovery.
- Scope Drift: charter signal absent or weak; deliverable count increasing over time; rundown gap widening despite adequate team confidence.
- Governance Gap: escalation language in documents without documented resolution; VoT comments referencing blocked decisions; schedule state deteriorating without intervention signals.
- Confidence Erosion: progressive VoT mean increase across reporting periods; flat or declining schedule metrics; risk language in VoT comments intensifying.

The dominant archetype and any co-present secondary archetypes are embedded in the Layer 1 signal summary and passed to the Cognitive hierarchy for LLM interpretation. The archetype classification frames the pattern-recognition task for the LLM layer, enabling it to produce contextually specific recovery guidance rather than generic risk commentary.

3.4 Recovery Signal Architecture

A distinctive feature of the PRIMMS signal architecture is the explicit representation of recovery signals as first-class inputs. Three recovery signals are computed independently from schedule state and VoT:

Recovery Urgency is derived from the convergence of schedule state, days behind, and jeopardy index, and classified as Watch / Elevated / Critical. Recovery Plan Evidence searches the documentary corpus for plan, corrective action, recovery, and rebaseline language, classifying the quality of documented recovery intent as Absent, Vague Only, Moderate, or Strong, and applying a corresponding db adjustment (–6 db for Strong evidence, reflecting the mitigating effect of a credible plan). Recovery Temporal Trend examines the VoT time series for directional movement — improving, flat, or worsening — and applies a corresponding signal to capture whether recovery actions are producing visible team sentiment effects.

The three recovery signals together answer three governance questions that project sponsors need answered on at-risk projects: Is recovery needed and how urgently? Is there a documented

recovery plan? Are recovery steps producing visible effects? The Layer 3 prompt (Temporal Situation Awareness) is explicitly structured around these three questions, requiring the LLM to answer each in order before proceeding to trajectory projections.

3.5 The Phase-Gate Quality Review Architecture

The weekly monitoring session and the phase-gate quality review address structurally different governance questions, and the distinction is consequential for system design. The PG-1/PG-2/PG-3 phase-gate structure developed below is a direct architectural descendant of the Quality Gate method introduced by Aaron, Bratta, and Smith (1993): the original Quality Gate paper established that a phase transition should be governed by predetermined, numerically demonstrable sufficiency criteria rather than a subjective readiness judgment, and that gate reviews should function as end-of-phase checkpoints distinct from routine progress tracking. PRIMMS's phase-gate pipeline operationalizes that same distinction with Bayesian evidence thresholds and an adversarial quality challenger in place of the fixed numeric criteria of the 1993 formulation.

A weekly monitoring session asks a directional question: Where is this project heading, and what intervention is most likely to change that trajectory? The signal set is a snapshot — the current schedule state, the most recent VoT entries, documentary evidence from the current reporting period. The outputs — archetype classification, recovery window projection, COA matrix — are recommendations about reversible near-term actions. If the archetype classification is wrong, the consequence is misdirected recovery effort; the error can be corrected the following week.

A phase-gate quality review asks an evidence-threshold question: Has this phase earned exit? The signal set is the accumulated history of the entire phase — every session's dB total, every archetype classification, the trajectory of the Bayesian issue closure ratio θ across weeks, the pattern of risk category hits across all phase documents. The output — EXIT AUTHORIZED, CONDITIONAL EXIT, or HOLD — is a governance decision that authorizes resource commitment, locks scope, and in many organizational contexts triggers contractual and budgetary consequences. A false EXIT recommendation allows unresolved structural problems to compound into the next phase; a false HOLD recommendation wastes resources and damages sponsor confidence. Both errors are institutionally consequential in ways that a misclassified weekly archetype is not.

This distinction motivates a different architecture for the phase-gate review, with three components: a deterministic phase-level aggregation layer, a modified Layer 2 prompt for gate context, and an adversarial OCC-5C challenger that applies the quality framework to the system's own phase-exit recommendation.

Figure 1 lays out this second pipeline. It runs on the same underlying session log as Figure 2, but aggregates across every session in the phase rather than reasoning from a single week's snapshot, and its governing thresholds can override the LLM's own verdict before a sponsor ever sees it.

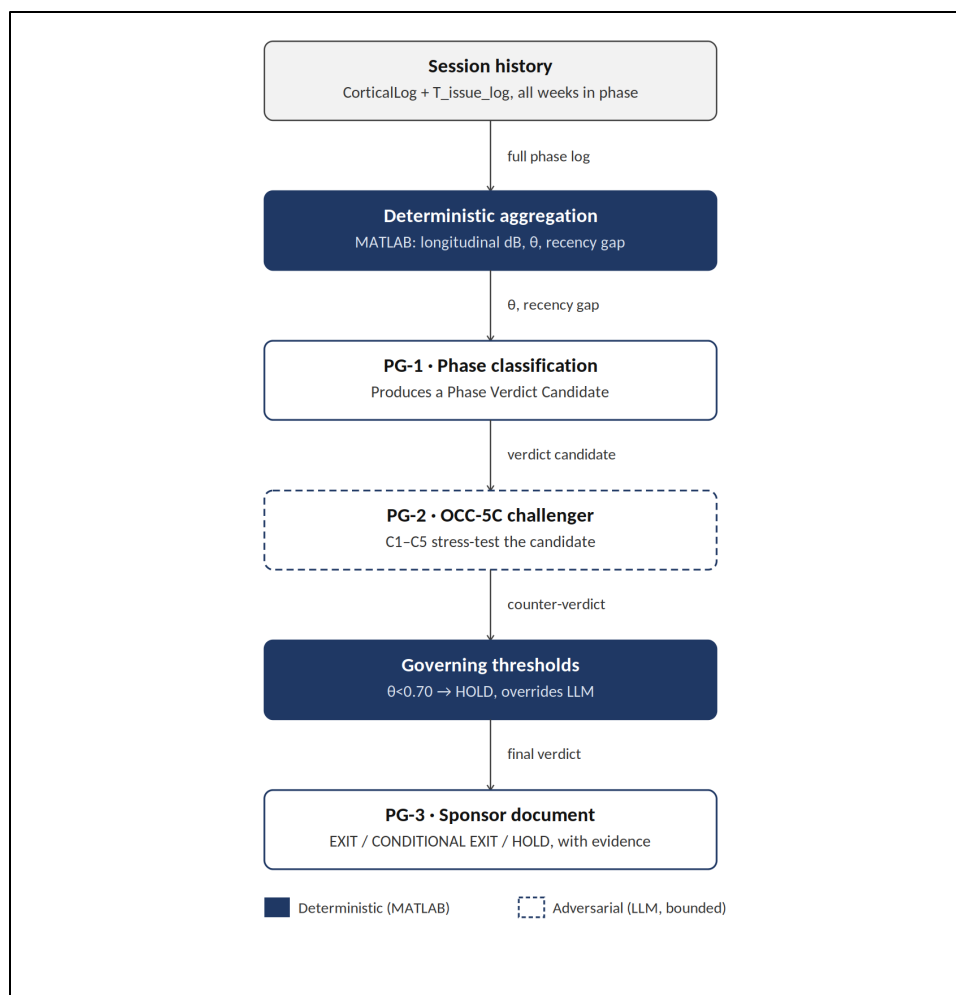


Figure 1. The Phase-Gate Quality Review pipeline. PG-2 stress-tests PG-1's verdict against five OCC-5C dimensions before the deterministic θ and recency-gap thresholds set the final verdict ceiling.

Phase-level deterministic aggregation.

The MATLAB component's compute phase gate metrics function aggregates the session history stored in the Cortical Log audit table and the issue log into phase-level metrics that are unavailable in any individual weekly session. These include: the longitudinal dB series across all sessions in the phase window and its linear trend; the recency-weighted versus full-phase-average dB comparison (the C5 calibration input); the archetype frequency table recording which archetypes fired and how often; the PREMATURE_CLOSURE persistence count; the θ ratio at gate date and its trajectory across the phase; and the OCC risk category hit counts aggregated across all phase documents. Three governing-equation thresholds are evaluated before any prompt is generated. $\theta < 0.70$ triggers a deterministic HOLD — below even odds of phase-exit success, no LLM recommendation can override this. PREMATURE_CLOSURE firing in a majority of sessions triggers a deterministic HOLD unless the challenger identifies explicit resolution evidence in the signal set. A recency gap exceeding 15 dB between the recent session average and the full-

phase average triggers a C5 calibration flag, capping the available verdict at CONDITIONAL EXIT. These thresholds are injected into the prompts as explicit constraints, not suggestions. The LLM interprets the phase evidence; the governing equations set the boundaries within which that interpretation operates.

Phase Classification prompt (PG-1). The first prompt asks Claude to classify the dominant failure pattern across the entire phase, not the current snapshot. The task is to assess whether the phase was characterized by a single persistent failure mode, multiple intermittent ones, or no significant failure signal — and whether that pattern was resolved, partially resolved, or unresolved by gate date. The trajectory assessment references the longitudinal dB series explicitly. The archetype persistence analysis assesses whether archetypes that fired multiple times had their underlying conditions resolved or merely deferred. The θ readiness assessment evaluates the ratio against both the even-odds threshold ($\theta \geq 0.70$) and the phase-exit standard ($\theta \geq 0.95$), naming the specific issue closures that would reach the exit standard if the threshold has not been met. The prompt's Phase Verdict Candidate — EXIT RECOMMENDED, CONDITIONAL EXIT, or HOLD — is the target of the adversarial challenger.

OCC-5C Adversarial Challenger (PG-2). The second prompt applies the five OCC-5C dimensions as a structured falsification of PG-1's Phase Verdict Candidate. This is the recursive application described in Section 6.7. The challenger's scope is explicitly bounded — it may produce only a pass/challenge verdict per C-dimension, one counter-verdict candidate with specific dB evidence from the session series, and a confidence-adjustment recommendation. It may not produce new narrative, a revised dB total, or an alternative archetype analysis. This constraint is the Zero constraint applied to the challenger itself: the feedback loop must reduce the output space, not expand it. MATLAB adjudicates: if the counter-verdict candidate has more dB support from the signal set than the original verdict, the challenger prevails. If not, the original stands — strengthened by surviving the falsification attempt.

The five C-dimensions acquire different operational character at phase level than in weekly sessions:

C1 (Consistency) at phase level asks whether the verdict contradicts the longitudinal dB trend. An EXIT recommendation on a phase whose dB rose steadily across sessions is a consistency failure — the current snapshot may look better, but the trend is adverse.

C2 (Coherence) at phase level asks whether the verdict is coherent with the archetype history. An EXIT on a phase where PREMATURE_CLOSURE fired repeatedly requires an explicit account of how that pattern resolved. If PG-1 did not provide that account, C2 challenges it.

C3 (Contradiction probe) at phase level requires the challenger to make the strongest available case for a different verdict, citing specific sessions and specific dB values. Narrative is not permitted — only evidence from the computed series. A HOLD trigger from the governing equations cannot be overridden by C3.

C4 (Confabulation check) at phase level flags any characterization of team behavior, recovery discipline, or sponsor confidence in PG-1 that is not traceable to a specific dB value, θ

measurement, or archetype firing in the evidence block. This dimension is particularly relevant at phase boundaries, where organizational pressure to produce a favorable assessment is highest and the temptation toward ungrounded narrative is strongest.

C5 (Calibration) is the load-bearing gate dimension. It asks whether the verdict confidence is calibrated to the full phase history or only to the most recent sessions — the recency bias that Kahneman (2011) identifies as one of the most reliable degraders of human judgment under social pressure. MATLAB computes the recency gap before the prompt is generated, and the C5 challenge is anchored to that computed value. A recency gap exceeding 15 dB is a C5 flag that caps the verdict at CONDITIONAL EXIT regardless of the PG-1 recommendation. The threshold is explicit because the failure mode is predictable: teams tend to address the most visible problems in the weeks immediately before a gate review, producing a late-phase improvement that is genuine but insufficient to represent the pattern of the full phase.

Phase-Gate Sponsor Document (PG-3). The third prompt synthesizes PG-1 and PG-2 into a formal sponsor document with a final verdict, conditions if CONDITIONAL, hold criteria if HOLD, a longitudinal evidence summary, and the OCC-5C challenger findings. The closing statement is identical to the weekly session's governance boundary: The machine has oriented. The humans now decide. At a phase gate, however, this statement carries different institutional weight — the human decision being returned to the sponsor is not a weekly intervention choice but an authorization of the next phase, with all the resource and contractual consequences that entails.

4. The Five-Layer Cognitive hierarchy

4.1 Architecture Rationale

The LLM interpretation component of PRIMMS is structured as a five-layer sequential prompt architecture termed the Cognitive hierarchy. The architecture reflects three design principles. First, no single prompt can simultaneously perform reliable feature validation, gestalt pattern recognition, temporal situation awareness, executive planning, and document generation — these are qualitatively different cognitive tasks that benefit from sequential, focused processing. Second, the hierarchical structure mirrors the computational organization of biological cortex as described in the predictive coding literature, with each layer operating at a higher level of abstraction than its predecessor. Third, a layered architecture creates an auditable record of reasoning: each layer's output is visible and correctable before it propagates to the next layer.

The hierarchy is therefore intended to be dynamic rather than merely sequential. Schedule updates, VoT observations, status reports, meeting transcripts, risk and issue records, recovery plans, and external events enter as organizational percepts — sensory analogs generated by the project environment. As these inputs move upward, Layer 1 preserves relatively concrete evidence; Layer 2 integrates it into gestalt and causal representations; Layer 3 extends those representations across time; Layer 4 converts the resulting orientation into candidate courses of action; and Layer 5 communicates that orientation for human judgment. In Fuster's terms, the

representations become progressively more abstract and action-oriented while remaining connected to the evidence that activated them.

Friston's predictive-processing account supplies the complementary feedback dynamic. Higher-order representations establish expectations about what lower-level evidence should look like. Evidence consistent with the active representation reinforces it; sufficiently discrepant or high-information evidence functions as prediction error and should receive greater influence. Conversely, an active causal hypothesis can direct attention back toward the documentary and quantitative evidence for confirming or disconfirming observations. The architecture is thus conceptually bidirectional: bottom-up evidence revises orientation, while top-down orientation determines which discrepancies warrant renewed attention. PRIMMS does not claim to implement cortical predictive coding technically; it operationalizes this recurrent principle through explicit evidence weighting, competing hypotheses, falsification tests, and human review between layers.

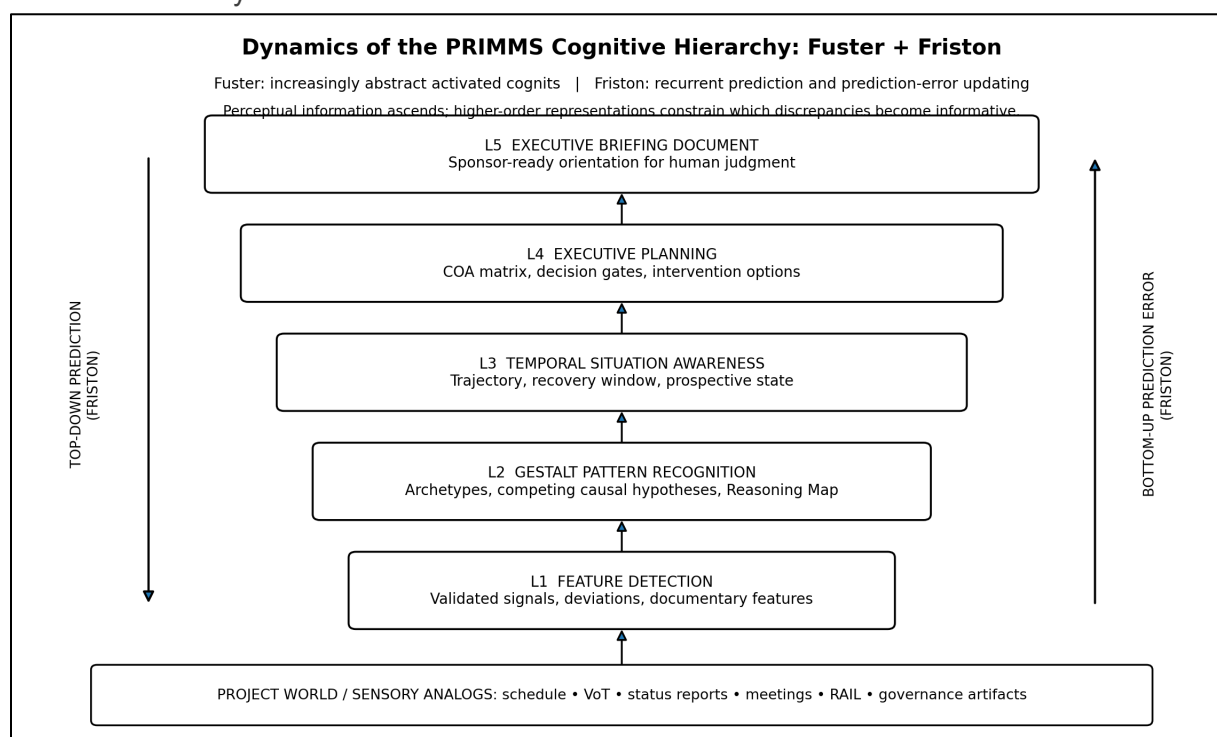


Figure 2A. Fuster-Friston dynamics of the PRIMMS Cognitive hierarchy. Organizational artifacts enter as sensory analogs; perceptual information is progressively integrated into more abstract representations, while higher-order expectations and bottom-up discrepancies create a recurrent challenge-and-revision cycle. The mapping is functional, not anatomical.

In the deployed system, each layer's prompt is generated by the MATLAB application and saved to a timestamped text file. The user pastes each numbered prompt into the LLM conversational interface sequentially, reading the response before pasting the next prompt. This human-in-the-loop step is not an implementation limitation to be engineered away — it is a deliberate governance feature. The project manager who reads each layer's output before proceeding is performing a judgment operation: assessing whether the LLM's interpretation is consistent with their contextual knowledge, and retaining the authority to override or redirect before the next layer builds on the prior output.

Figure 2B traces this flow end to end: deterministic signal computation at the base, the evidence floor separating measured fact from generated hypothesis, four LLM interpretive layers built on top of it, and the human authority boundary that every path terminates in.

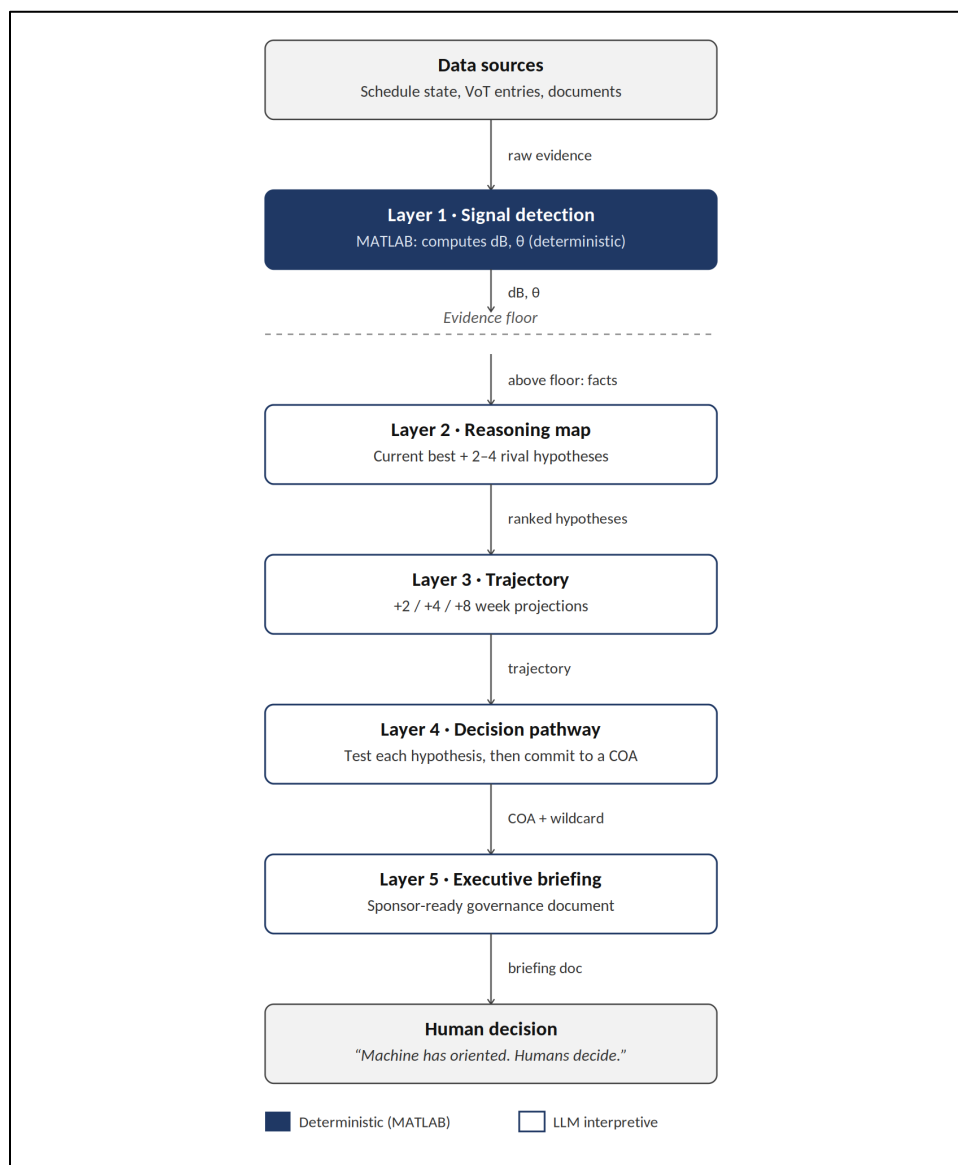


Figure 2B. The weekly signal pipeline, from raw project data through five interpretive layers to the human decision point. Labeled arrows show what each layer passes to the next.

4.2 Layer Descriptions

Layer 1 — Feature Detection. The sensory layer receives the full MATLAB-computed signal set with db values and supporting evidence. The Layer 1 role is to validate and extend: confirm which signals are correctly computed from the available data, identify any signals the deterministic algorithm may have missed given contextual information visible in the project description, and flag any signals that should be removed as spurious. The layer concludes with a confirmed total db

value and Jeffreys band. The instruction explicitly prohibits prescriptions — Layer 1 is a perceptual layer, not an intervention layer.

Layer 2 — Gestalt Pattern Recognition. The pattern recognition layer receives the validated signal set and classifies the dominant failure archetype from the six defined types. The layer is instructed to consider all six archetypes (including Premature Closure — the failure mode of closing a recoverable decision prematurely) and to explicitly assess whether any archetype's firing conditions are met given the evidence. The Premature Closure instruction is specifically structured to counter the institutional default toward replanning: 'PREMATURE_CLOSURE fires when recovery is plausible on at least one quantitative signal AND the institutional default is to replan rather than test recovery.' The layer concludes with the dominant archetype and cumulative db. This classification is generalized into a full competing-hypothesis architecture, the Reasoning Map, described in Section 4.4.

Layer 3 — Temporal Situation Awareness. The situation awareness layer extends Endsley's (1995) three-level model — perception, comprehension, projection — to the project risk context. The layer explicitly requires trajectory analysis at +2, +4, and +8 week horizons; identification of surprising or counter-intuitive signals in the data; and five leading indicators of recovery. The recovery assessment sub-section (R1–R5b) mirrors the three governance questions described in Section 3.4, with the addition of a Premature Closure check (R5b) that explicitly asks whether the replan posture is itself a bias flag — whether the recovery hypothesis has been tested and falsified, or merely assumed to be false.

Layer 4 — Executive Planning. The executive planning layer produces the Course of Action (COA) matrix: six structured recovery options (Scope Freeze, Cadence Tighten, Dependency Clearing, Forecast Rebase, Strike Team, Quality Gate Enforce), each with root cause targeted, 3–5 concrete remediation actions, owner archetype, decision gate trigger, leading indicator of success, and COA risk. The layer also produces 2–3 decision gates, a Week 1 intervention plan, key assumptions and bias flags, an escalation posture on a 1–6 scale, and a 400–600-word governance narrative. The layer's closing instruction — 'The machine has oriented. The humans now decide.' — is structurally enforced as the last sentence of every Layer 4 output. The decision-gate structure is extended into a fully sequenced falsification protocol, the Decision Pathway, described in Section 4.5, which also introduces a wildcard course of action alongside the six fixed options.

Layer 5 — Executive Briefing Document. The document generation layer synthesizes the full five-layer analysis into a formatted executive briefing using all prior layer outputs. The document structure includes: a color-coded status banner (the banner color, label, and tone instructions are set by MATLAB's deterministic STATUS_KEY computation, not by the LLM); a recovery status section with explicit velocity-based recovery window language; a plain-English situation narrative; a signal table; failure archetype analysis; trajectory projections; recovery options; three leadership decisions required this week; and a first-week intervention plan with an end-of-week velocity gate. The layer terminates with the string 'DOCX_READY', which the MATLAB application detects

automatically as confirmation that the complete executive summary is present and ready for docx generation.

4.3 Audit and Traceability

The MATLAB application maintains a persistent Cortical Log table recording every LLM session. Each record stores: session timestamp, project ID, prompt file path, SHA-256 hash of the prompt content, mode (layered or single), layers processed, LLM model version string (the specific version identifier is recorded at session time and included in the audit log; the default at time of writing was `claude-sonnet-4-6`), full JSON of layer responses, `DOCX_READY` detection flag, and log creation timestamp. The complete log is included in session saves and audit exports. This architecture ensures that governance reviewers can identify exactly which model version produced any given analysis, reproduce the signal inputs that drove the analysis, and verify that no output was modified post-generation.

The model version string is embedded in the prompt file header and requires explicit update when the recommended LLM model changes — a design choice that forces a conscious governance decision rather than silent model substitution.

4.4 The Reasoning Map: Competing Hypotheses Under Inference to the Best Explanation

Layer 2's archetype classification, described in Section 4.2, answers a categorical question: which one of six defined failure archetypes best characterizes the convergent signal pattern? This formulation is adequate when the evidence converges cleanly on a single archetype, but many at-risk projects present evidence that is genuinely ambiguous between two or more explanations — a single-owner capacity bottleneck and independent scope volatility, for instance, produce overlapping but not identical signal signatures, and a governance narrative that commits to one explanation prematurely forecloses the diagnostic work needed to distinguish them.

This is the tails-not-the-mode principle introduced in Section 2.3 given architectural teeth. Left to its default behavior, an LLM asked to explain a risk pattern returns the modal explanation — the single most statistically probable account given the accumulated linguistic history it was trained on. The modal account is frequently correct, but it is also the least differentiated output available, and a governance process that stops there has no way of knowing whether it stopped because the evidence was decisive or because the model simply reported the center of its own distribution. The Reasoning Map is the mechanism that forces a second look at the tails of that distribution before either conclusion is accepted.

The Reasoning Map extends Layer 2 into a structured application of Inference to the Best Explanation (Lipton, 2004) rather than a single-label classification. Above an explicit evidence floor, the map records only MATLAB-computed signal nodes carrying deciban confidence in Jeffreys' (1961) strength-of-evidence bands. Below the floor, the LLM is required to name exactly one Current Best explanation — the leading account, ranked highest against the explanatory virtues of scope, parsimony, coherence, analogy, and testability — alongside a flexible field of

two to four additional hypothesis nodes, each assigned a role of Rival, Tail Rival, or Combination (the latter used when two causes jointly explain the observation). At least one genuine Tail Rival is required whenever the evidence is consistent with one; where none is defensible, the system emits an explicit node stating that no tail hypothesis was found, rather than omitting the role silently. Each hypothesis node is linked to the specific evidence nodes it explains, is required to name what it explains that its rivals do not, and carries an explicit statement of the confirming-or-killing test that would move it.

This structure is a direct application of the Zero constraint (Section 6.6) to hypothesis generation: an unconstrained LLM asked to explain an observed risk pattern will proliferate an unbounded set of plausible-sounding narratives — a hedge forest that expands the output space rather than sharpening it. Fixing the Current Best role at exactly one, bounding the rival field to two to four additional nodes, and requiring an explicit role assignment for each — rather than leaving cardinality or role open-ended — forces the interpretive layer to commit to a ranked structure rather than hedge across an open list.

The confidence units on either side of the evidence floor are kept deliberately incommensurable. Measured evidence nodes carry the deciban strength bands used throughout the architecture (Decisive / Strong / Moderate / Slight, per Jeffreys, 1961); hypothesis nodes carry a separate, explicitly labeled qualitative rank (Leading / Live / Possible / Tail) representing the LLM's own assessment of explanatory standing. The two units are never blended into a single number. This separation operationalizes the distinction Lipton (2004) draws between the likeliest explanation and the loveliest explanation: a hypothesis may be the most explanatorily satisfying account available — the best of the hypotheses the system generated — without being the most probable account in an absolute sense, a gap van Fraassen (1989) identifies as the central vulnerability of inference to the best explanation as a truth-tracking procedure. Each hypothesis node is annotated “likely, not merely lovely” precisely to hold this distinction in view: explanatory elegance is a reason to test a hypothesis, not a substitute for testing it. The Decision Pathway, described in Section 4.5, is the mechanism by which that testing occurs before any hypothesis is treated as settled.

4.5 The Decision Pathway: Sequenced Falsification Before Commitment

Layer 4's decision-gate structure, described in Section 4.2, produces 2–3 decision gates alongside the COA matrix. The Decision Pathway extends this into a sequenced, falsification-ordered protocol that determines which gate — and which course of action — the evidence actually supports, before capacity is committed to any of them.

The pathway is generated directly from the Reasoning Map's hypothesis set. Each step names a specific, low-cost diagnostic check tied to one hypothesis; states what result would confirm that hypothesis and the corresponding management action to prepare; and states what result would instead rule it out, directing the reader to the next step in sequence. The ordering is not arbitrary: the highest-value test — the one that most efficiently discriminates between the Current Best explanation and its Live Rival — is run first, deferring more expensive or less discriminating tests

until the cheaper ones have narrowed the hypothesis space. An illustrative sequence from a live deployment: Step 1 checks whether newly late tasks require the same specialized resource as the original bottleneck or are independent work, discriminating the Current Best single-bottleneck hypothesis from its rival; Step 2 audits whether unassigned work traces to scope-change tickets, discriminating the rival scope-volatility hypothesis; Step 3 repeats both audits jointly to test the combination hypothesis; Step 4 asks the constrained resource directly whether access, not effort, is the binding constraint, testing the tail hypothesis. Each step's negative branch is itself a result — it eliminates a hypothesis and passes the remaining evidence forward — so the absence of confirmation is treated as informative, not as a null result to be discarded.

This structure is a direct operationalization of Popperian falsification (Popper, 1959) at the point where organizational decisions are made, extending the argument developed in Section 6.1 regarding Premature Closure. Janis and Mann (1977) identify defensive avoidance — the tendency to treat a favorable-looking recent signal as sufficient grounds to stop testing and commit — as one of the most common failure modes in high-stakes organizational decision-making. The Decision Pathway is structured specifically to counter it: no course of action is recommended for selection until the pathway has named, and where feasible executed, the test that could kill it. The COA matrix described in Section 4.2 is thereby extended with a wildcard alternative: every hypothesis node is mapped to one of the six fixed COA types or to NONE, and a Tail Rival left at COA=NONE is carried forward into the visualization as a wildcard — a course of action not among the six fixed types, available precisely because the evidence has surfaced a credible-but-untested account that no standard recovery option addresses. This keeps a genuine tail hypothesis from being silently discarded for lack of a matching COA, and gives the reviewer an explicit prompt to weigh it rather than defaulting to one of the six — a structural check against the same defensive-avoidance pattern the Layer 3 R5b prompt is designed to surface at the weekly monitoring level.

5. Retrospective Validation

5.1 Case Context

The retrospective validation applies PRIMMS to a set of recently concluded enterprise transformation programs, including SAP S/4HANA Rise implementations and other large-scale IT and data science initiatives. The Bayesian WoE and Voice of Team components of PRIMMS® have been in active production use since 2008 and carry an extensive operational track record across this portfolio. The Cognitive hierarchy LLM layer, by contrast, is newer and is currently being tested retrospectively against completed programs and project phases — running in parallel with, but not yet as a replacement for, the existing PRIMMS® web platform. The retrospective testing program applies the full PRIMMS architecture to historical project records, with the purpose of assessing whether the system produces signals, classifications, and recovery guidance consistent with what actually occurred or should have occurred.

The protocol applied to each project is as follows. Project schedule exports, status update emails, steering committee documents, and meeting records from the completed program are assembled and ingested into PRIMMS. Anonymous Voice of Team (VoT) ratings are reconstructed from meeting records and contemporaneous written assessments by team members and stakeholders. The MATLAB signal layer is then run against each weekly snapshot in sequence, and the five-layer Cognitive hierarchy is executed against selected snapshots of particular analytical interest — typically those that, in retrospect, represented inflection points in the program's risk trajectory. The outputs of the PRIMMS analysis are compared against what the conventional PMO monitoring showed at the same date and against the known final outcome of the program.

This testing design is explicitly retrospective and post-hoc. The Cognitive hierarchy component has not yet been deployed in real time during a live client engagement. The validation question is therefore: given access to the same data that was available at the time, does PRIMMS produce signals, classifications, and recovery window estimates consistent with what actually happened or should have happened? This is a weaker claim than prospective deployment — it cannot establish that the system would have changed decisions in real time — but it is a necessary precursor to prospective validation and tests the signal architecture against known ground truth. The following subsections present one illustrative example from this validation program in detail: a snapshot dated 8 August 2025 from a completed machine learning proof of concept data science enterprise project, chosen because it represents a clear inflection point with a well-documented outcome against which the system's outputs can be precisely compared.

5.2 Illustrative Example: Signal Set at the 8 August 2025 Snapshot

At the snapshot for a sample project dated 8 August 2025 — the selected illustrative inflection point for this example project — the MATLAB Layer 1 computation produced the following signal set from the retrospectively assembled project data:

Schedule signals: SCHEDULE_STATE_BEHIND (10.0 db); SCHEDULE_DAYS_BEHIND, value 18 days (5.4 db); SCHEDULE_JEOPARDY, index 24.75 (5.0 db). VoT signals: VOT_MEAN_ELEVATED, mean 3.60 across five entries (7.0 db). Classifier signal: CLASSIFIER_JEOPARDY (18.0 db). Rundown signal: RUNDOWN_GAP_BEHIND, gap -9 deliverables (3.6 db). Velocity signal: VELOCITY_FINISH_SLIP, +4 day projected slip (0.6 db). Recovery signals: RECOVERY_URGENCY_CRITICAL, value 18 days behind (14.0 db); RECOVERY_PLAN_DOCUMENTED_STRONG, quality level Strong (-6.0 db, mitigating); RECOVERY_TREND_WORSENING, VoT mean rising early 3.00 → recent 4.00, delta +1.00 (8.0 db).

Total prior: 65.5 db. Jeffreys band: Decisive.

The velocity-based rundown model produced a projected finish date of 24 October 2025 — four days after the planned end date of 20 October 2025 — with $R^2 = 0.63$. The critical finding was not the slip itself but its magnitude: a 6% improvement in daily delivery velocity (from 0.54 to 0.57 deliverables per day) would close the gap entirely within the existing timeline, without scope or schedule change.

5.3 Cognitive hierarchy Analysis: Illustrative Example

The five-layer Cognitive hierarchy was executed against the 8 August 2025 snapshot of the illustrative example program. The outputs are reported here as produced — they represent what the system generated from the retrospectively assembled data, not a real-time governance document — and are intended to demonstrate the architecture's reasoning process rather than to constitute a standalone validation claim.

The Layer 2 archetype classification identified GOVERNANCE_GAP and PREMATURE_CLOSURE as co-dominant archetypes. The governance gap was evidenced by the divergence between a strongly-documented recovery plan (–6 db mitigating) and a worsening VoT trend (8 db) — the plan existed on paper but was not producing measurable team sentiment effects. The premature closure archetype was flagged because the velocity projection showed a recoverable slip, while the institutional default under schedule pressure was to trigger a replan — closing a decision that the evidence did not yet justify closing. The Layer 2 output explicitly stated: 'The recovery hypothesis has not been falsified; it has simply not been tested.'

Layer 3 situation awareness computed that the recovery window was approximately 2–3 weeks: after that point, the compounding of inaction would make the 6% velocity improvement insufficient to close the gap within the original timeline. The +2-week trajectory showed the recovery window narrowing but open; the +4-week trajectory showed formal scope negotiation becoming necessary; the +8-week trajectory showed structural jeopardy with stakeholder alignment likely broken down.

Layer 4 produced a six-option COA matrix. The recommended sequencing was: Dependency Clearing (primary target: DATA_READINESS risk at 57.2, the dominant risk category) within 24 hours; Scope Freeze in parallel to arrest compounding; Cadence Tighten to ensure real-time tracking of intervention effects. The escalation posture was set at 5 of 6 (Immediate Leadership) — not yet 6 (Decision Gate Required) because the recovery window remained open, but requiring leadership decisions within the current week.

Layer 5 generated a complete executive briefing document. The status banner used the STATUS_KEY JEOPARDY_RECOVERABLE, with tone instructions emphasizing the recovery window rather than failure probability: 'The briefing MUST lead with this — the recovery window is open. Urgency framing: the window closes in 2–3 weeks, not that failure is inevitable.'

5.4 Comparison Against Known Outcomes: Illustrative Example

This illustrative example permits a direct comparison between what PRIMMS produced from the 8 August 2025 snapshot and what the conventional PMO monitoring showed at the same date for the same program. The conventional PMO reporting for that period classified the program as Amber on the standard RAG dashboard. No escalation was triggered, no formal recovery plan was activated at that point, and no specific recovery window was quantified.

PRIMMS, applied to the same data, produced: a decisive evidence classification (65.5 db); identification of DATA_READINESS as the primary blocker (risk score 57.2, more than twice the

next-ranked risk); a specific recovery target (6% velocity improvement, from 0.54 to 0.57 deliverables per day); a projected finish date of 24 October 2025 against a planned end of 20 October 2025; and a recovery window of approximately 2–3 weeks before the gap became unrecoverable without scope or schedule change. The dominant failure archetypes were classified as GOVERNANCE_GAP and PREMATURE_CLOSURE — the latter reflecting that a replan was institutionally plausible at that point, while the velocity data indicated recovery was still achievable.

The known final outcome of the program is that it completed within two days of its original planned end date. Recovery was achieved through actions that addressed the data readiness blockage and introduced an accelerated governance cadence — the same intervention path that PRIMMS's Layer 4 COA matrix identified as the recommended sequencing. The velocity model's projection proved accurate in retrospect: a further unremediated reporting cycle would have extended the gap to 9–12 days, consistent with the rate of deterioration observed in the project record, and would have placed recovery outside the original timeline.

The appropriate inference from this illustrative comparison is constrained in two respects. First, because PRIMMS did not run in real time during the program, we cannot claim that the system caused or contributed to the recovery. Second, this is one example drawn from the broader retrospective validation program; it is presented to demonstrate the architecture's reasoning process and output quality, not to constitute a statistically representative sample. What this example does establish, and what the broader validation program supports across multiple programs, is that the signal architecture correctly classifies known-critical snapshots; that the velocity-based recovery window calculation is consistent with ground truth; and that the Cognitive hierarchy's COA recommendations align with the interventions that produced successful outcomes. These are necessary but not sufficient conditions for prospective validity. The validation program supports proceeding to a prospective deployment study — the highest-priority item in the future work agenda.

6. Theoretical and Practical Contributions

6.1 Advancing the Hybrid Cognition Framework

The PRIMMS architecture makes a specific theoretical contribution to the emerging literature on human-AI collaboration in complex decision environments. It operationalizes a design principle — machine perception with human authority — at a level of architectural specificity that is absent from prior work. The division of labor between MATLAB (deterministic, auditable, inspectable signal computation) and LLM (natural-language interpretation, synthesis, projection) is not a pragmatic engineering choice but a principled allocation based on the comparative cognitive advantage of each.

The Cognitive hierarchy prompt architecture advances a claim about how LLMs should be integrated into professional decision support: not as autonomous reasoners generating advice,

but as interpretive engines operating on structured, machine-verified signal sets within architectures that preserve human authority at every consequential step. This framing positions LLMs differently than both the automation literature (which treats them as replacements for human tasks) and the productivity literature (which treats them as accelerants for existing tasks). The present paper argues for a third role: structural perception amplifiers that extend the human decision-maker's ability to detect what is actually happening before it becomes visible in conventional reporting.

This contribution can also be stated in cognitive terms. Fuster supplies the hierarchy of representation, Friston the recurrent updating dynamic, and Boyd the operational consequence: Orientation. PRIMMS couples these concepts in an engineered management architecture in which organizational percepts are transformed through hierarchical abstraction, Bayesian evidence accumulation and prediction-error challenge into a continually revisable orientation; that orientation then frames course-of-action search, while consequential choice remains with the human decision-maker.

The explicit treatment of Premature Closure as a failure archetype — the institutional default of closing a recoverable decision without testing the recovery hypothesis — represents an original contribution to project management failure classification. The Janis and Mann (1977) conflict theory framework is extended to the project risk context through the Layer 3 R5b prompt, which requires the LLM to explicitly assess whether the replan posture is itself a bias flag. This operationalizes the Janis-Mann defensive avoidance construct in a way that can be applied in real-time governance practice. It is, in effect, an operationalization of Popperian falsification discipline (Popper, 1959) within governance practice: the recovery hypothesis is treated as a claim to be tested and potentially killed, not a default to be assumed.

A practical implication of this architecture deserves explicit statement. As AI capabilities commoditize — as models converge, techniques diffuse, and capabilities that once distinguished organizations are absorbed into widely available platforms — the competitive boundary in AI-augmented management shifts. Advantage no longer resides in access to prediction. It resides in how prediction is governed: where authority is located, how signal evidence is structured before interpretation begins, and who remains accountable when outcomes are realized. The PRIMMS architecture is a response to precisely this condition. Its durable contribution is not the specific LLM model it employs — that is a commodity — but the governance architecture that wraps it: the deterministic signal layer, the hierarchical interpretation structure, the human authority boundary, and the governing equations that prevent complexity collapse. Organizations that mistake model access for governance architecture will find that the AI systems they adopt amplify their existing decision quality, whether good or poor. Those that invest in the architecture will find that the same underlying models produce systematically better orientation — because the structure, not the model, is the source of performance.

6.2 The Intelligence Gap Addressed

Aaron (2026, Chapter 1) identifies five structural gaps in conventional project management practice: team confidence erosion; perception distortion in leadership; failure archetype recognition; multi-modal evidence fusion; and recovery trajectory projection. Table 3 maps the PRIMMS architectural layers to these gaps.

Table 3. Mapping PRIMMS Architecture to Identified Intelligence Gaps

Intelligence Gap	PRIMMS Response	Architectural Layer
Team confidence erosion	VoT 1–6 time-series captures distributed team sentiment; temporal trend computed deterministically; VOT_MEAN signals fused into prior dB	Signal Acquisition; Layer 1 Feature Detection
Perception distortion in leadership	Divergence between VoT trend and leadership-reported schedule status surfaced as PERCEPTION_DISTORTION archetype; Layer 3 explicitly checks for optimism bias in recovery claims	Layer 2 Pattern Recognition; Layer 3 Situation Awareness
Failure archetype recognition	Six archetypes with characteristic signal signatures; deterministic archetype classifier in MATLAB; LLM interprets dominant and co-present archetypes in Layer 2	Inductive Classifier; Layer 2 Pattern Recognition
Multi-modal evidence fusion	Schedule, VoT, and documentary signals fused via Bayesian WoE summation into a single prior db; signal independence assumption acknowledged as limitation	Signal Extraction; Bayesian WoE Computation
Recovery trajectory projection	+2/+4/+8 week projections explicitly required in Layer 3; velocity-based recovery window computation in MATLAB; recovery signals as first-class inputs	Recovery Signal Architecture; Layer 3 Situation Awareness

A concrete instance from the project portfolio illustrates the multi-modal evidence fusion gap directly. On a SAP S/4HANA Rise implementation, Single Sign-On (SSO) authentication against Fiori began failing intermittently in the weeks before go-live. The project's escalation path routed the problem to a SAP Basis/security specialist based in South America, and the resulting cross-time-zone ticket handoffs consumed roughly two weeks before the root cause — a web dispatcher/load balancer misconfiguration interacting with the SAML trust relationship — was isolated. Framing the same underlying facts as a direct query to an LLM ("what would cause web dispatchers and load balancers to cause SSO not to work properly on an SAP S/4HANA Rise project") returned the relevant candidate causes — certificate mismatches, trust relationship gaps, load-balancer misrouting, Fiori Launchpad and Gateway misconfiguration — within minutes. A text-mining pass over the project's own weekly status reports and standup notes showed FIORI_SSO already among the most frequent terms in the documentary corpus well before go-live, present in three of the seven weekly documents analyzed (slides 15–17, *AI Project Management as a Service w/ MATLAB* companion deck). The signal was accumulating in the

project's own documentary record the entire time; it was never surfaced to the people who needed it because no one was querying that record with a tool capable of connecting recurring project language to a stored library of platform-specific failure modes. This is the intelligence gap Table 3 names in the abstract, in concrete form: the answer was not undiscoverable, it was unindexed.

6.3 Design Intent vs. Implementation Fidelity

Table 4 presents the alignment assessment between the original design intent described in the source monograph (Aaron, 2026, Chapter 8.7) and the current PRIMMS implementation. This assessment follows the transparency standard argued for elsewhere in this paper: the system's design gaps should be as visible as its demonstrated capabilities.

Table 4. Design Intent vs. Current Implementation Fidelity

Dimension	Original Design Intent	Current Implementation	Fidelity
Cybernetic control loop	Sensing → deviation → corrective signal via adaptive control agent	5-layer Cognitive hierarchy: MATLAB computes signals; LLM interprets, projects, recommends	High
MDP state/action space	Schedule state × VoT average × days behind → discrete intervention actions	Schedule + VoT + Documents → Bayesian WoE (db) + 6-option COA matrix with decision gates	High
Failure archetype recognition	Six archetypes derived from multi-modal signal convergence	Six archetypes implemented with characteristic signal signatures and recovery implications	High
Human authority boundary	Recommend only; full audit trail; manager owns every decision	'Orient, never decide' structurally enforced in all prompts; audit trail present	High
RL learned reward policy	DDPG/Q-learning with reward updates from weekly survey responses over time	Deterministic db scoring (Bayesian WoE). No live RL adaptation loop across weeks	Partial
ML rundown curve classifier	Visual ML classifier trained on historical project images; detects jeopardy from curve geometry	Rundown tab computes velocity fit, R^2 , and gap as signals — no image-based ML classifier	Low

The overall fidelity assessment is 75–80%. The governance philosophy, state-space formulation, domain encoding, failure archetype classification, and human authority boundary are faithfully implemented — in some respects strengthened by the Cognitive hierarchy and Bayesian WoE framework relative to the original design. The two significant gaps (live RL adaptation and image-

based rundown classifier) are noted as future development priorities rather than disqualifying limitations.

Table 5. Phase-Gate Quality Review: Governing Equations and Verdict Logic

Governing Condition	Source	Verdict Impact
$\theta \geq 0.95$ at gate date	Aaron (2015) phase-exit standard	EXIT eligible; subject to trend and C-check override
$\theta \geq 0.70, < 0.95$	Aaron (2015) even-odds threshold	CONDITIONAL EXIT; conditions required
$\theta < 0.70$	Aaron (2015)	HOLD (deterministic — overrides LLM recommendation)
Recency gap > 15 dB	MATLAB: recent avg vs. phase avg	C5 flag → CONDITIONAL EXIT at best
PREMATURE_CLOSURE in majority of sessions	Archetype frequency table	HOLD unless challenger finds explicit resolution evidence
dB trend deteriorating + $\theta \geq 0.95$	Longitudinal dB series slope	CONDITIONAL EXIT — late-phase signal accumulation requires explanation
Two or more C-dimension CHALLENGE findings	OCC-5C Challenger (PG-2)	CONDITIONAL EXIT at best
All five C-dimensions PASS	OCC-5C Challenger (PG-2)	Original verdict confirmed; confidence strengthened by surviving falsification

Notes: θ = Bayesian issue closure ratio (closed issues / total issues). Recency gap = (mean dB of last three sessions) – (mean dB across full phase). All governing conditions are evaluated by the MATLAB deterministic layer before any prompt is issued to the LLM. Conditions are evaluated in the order shown; the first triggered condition determines the verdict ceiling.

6.4 Self-Reference, Gödel, and the Structural Case Against Autonomous AI

The PRIMMS architecture’s two-component design — a deterministic MATLAB layer for signal computation and an LLM layer for interpretation — is grounded in a practical engineering principle: tasks requiring auditability should not be delegated to probabilistic systems. However, the development of this architecture surfaced a deeper theoretical observation, one that connects the empirical behavior of large language models to foundational results in twentieth-century logic and computation.

During the development and testing of the Cognitive hierarchy prompt architecture, a consistent pattern emerged: attempts to instruct the LLM to constrain its own output — to remain concise while simultaneously performing rich analytical reasoning — failed systematically. Brevity constraints embedded in a prompt did not function as external governing rules. They functioned as one competing signal among many, processed through the same inductive mechanism that

produced the analysis itself. The richer the analytical task, the more reliably the brevity constraint degraded. The system could not enforce a bound on its own behavior using the same faculty that generated the behavior.

This empirical observation has a formal analogue. Kurt Gödel's incompleteness theorems (1931) demonstrated that any sufficiently expressive formal system must be either incomplete or inconsistent: there exist true statements within the system that cannot be proven from within it, and the system cannot demonstrate its own consistency using only its own axioms. Alan Turing extended this boundary into computation, proving that no general procedure exists to determine whether an arbitrary program will halt — a computational system cannot universally predict its own behavior (Turing, 1936). Both results describe the same structural limit: self-referential systems that are sufficiently expressive cannot fully know, predict, or regulate themselves from within.

Large language models represent a modern instantiation of this principle. The instructions provided to an LLM and the outputs it produces exist within the same representational substrate: both are sequences of tokens processed through the same inductive mechanism. A constraint instruction is not an external control; it is an input processed by the same system it is intended to govern. This is precisely the structure that Gödel's and Turing's results identify as fundamentally limited. The LLM cannot step outside its own generative process to enforce a rule upon that process — for the same structural reason that a formal system cannot prove its own consistency from within its own axioms.

The implication for the governance debate surrounding artificial intelligence is significant. Much of the discourse around autonomous AI — both the optimistic strand arguing that sufficiently capable systems will self-correct and self-regulate, and the pessimistic strand warning of uncontrollable self-directed AI — shares a common premise: that intelligence, once sufficiently scaled, becomes self-governing. The Gödel-Turing analysis suggests this premise is structurally false. A system powerful enough to exhibit the kind of general reasoning that would make it a candidate for autonomy is, by that same expressive power, too complex to reliably constrain itself using its own mechanisms. The constraint and the thing being constrained are the same system — and that structure is precisely what the halting problem and the incompleteness theorems show cannot self-regulate under general conditions.

The failure mode this analysis predicts is not domination but instability. A system that expands generative capability without a corresponding external constraint layer does not become sovereign; it becomes unpredictable. Its outputs remain coherent and often persuasive, but not reliably bounded. This is consistent with observed LLM behavior: systems that produce sophisticated reasoning under one framing produce contradictory reasoning under another, without awareness of the inconsistency. The system cannot self-adjudicate because the adjudication mechanism is not separate from the generative mechanism.

This analysis provides theoretical grounding for the PRIMMS architectural choice that might otherwise appear as mere engineering conservatism. The separation between the MATLAB signal layer and the LLM interpretation layer is not simply a preference for auditability. It is a

structural response to the Gödel-Turing constraint: deterministic constraint enforcement must originate outside the inductive system. The MATLAB component imposes bounds that the LLM layer cannot override — not because it is more intelligent, but because it is external. The human authority boundary enforced at every decision point serves the same structural function: it provides the external constraint layer that inductive systems cannot supply for themselves.

The argument advanced here is developed at length in Aaron (2026, Section 3.10), which traces the historical arc from Hilbert's program through Gödel's incompleteness theorems and Turing's halting problem to the structural constraints on contemporary large language models. The present paper makes a more specific and practically grounded claim: the two-layer PRIMMS architecture — external deterministic constraint plus internal inductive interpretation — is not merely a convenient engineering design. It is the theoretically necessary form of any hybrid cognition system that seeks both the interpretive richness of large language models and the governance reliability that autonomous AI systems cannot, by their structural nature, supply for themselves.

6.5 Command, Control, and the Irreducible Human Boundary

The Gödel-Turing constraint establishes a theoretical limit on machine self-regulation. A complementary and independently derived boundary emerges from the operational theory of command and control developed by John Boyd and formalized in military planning doctrine. Boyd distinguished sharply between command and control: command establishes intent and commitment; control maintains orientation under uncertainty. In Boyd's formulation, effective systems depend on trust, shared understanding, and freedom within bounds — control that becomes too explicit turns into interference, while command delegated to a machine destroys initiative. Effective command must be explicit; effective control must often be invisible.

This distinction maps directly onto the PRIMMS architecture. The system strengthens control — orientation, signal detection, pattern recognition, trajectory projection — while leaving command entirely untouched. The project manager who reads the Layer 4 COA matrix and selects an intervention is exercising command: establishing intent, assuming consequence, and committing to a course of action that cannot be undone. The machine has oriented. The human commands. Attempting to automate command would violate the structural conditions required for coordinated institutional action — not because machines lack processing power, but because command carries accountability that is conferred socially and institutionally, not computationally.

The RAND Corporation's research on command concepts (Builder, Bankes, & Nordin, 1999) reinforces this boundary from an empirical direction. Analyzing command effectiveness across military operations, RAND found that system performance depends less on the volume or speed of information and more on the intellectual quality of the commander's concept of operations. Information systems cannot rescue a flawed command concept; superior data cannot substitute for intellectual clarity. Conversely, a clear concept enables coordinated decentralized action even under degraded information conditions. The implication for AI-augmented project management is direct: PRIMMS can improve the quality of information available to the project leader, but it cannot generate the transformation concept, nor can it absorb contractual, political, or reputational

consequences when outcomes are realized. These remain irreducibly human functions — not as a temporary accommodation to technological limits, but as a permanent structural feature of any institutional system that must remain accountable.

6.6 The Zero Constraint: Complexity Collapse and the Role of Governing Equations

The Gödel-Turing constraint identifies one fundamental limit on machine-assisted cognition: a sufficiently expressive inductive system cannot fully regulate itself from within. During the development of the PRIMMS Cognitive hierarchy, a second and distinct constraint emerged from practice. It is not a constraint of logical incompleteness but of operational collapse — and it has a precise cultural analogue in the 1975 science fiction film *Rollerball*.

In that film, the supercomputer Zero is asked to retrieve historical information. It cannot. The failure is not attributed to insufficient data or processing power. The cause is accumulation: Zero has been given so much information across so many domains that retrieval has become indeterminate. When pressed for a specific answer, human users of Zero complain that “everything is ambiguous now.” The machine has not failed to reason. It has accumulated itself into uselessness. It considers everything — and precisely because it considers everything, it can say nothing actionable.

This failure mode is not fictional. It is a live risk in any machine-assisted project management system that adds signals, archetypes, prompt layers, and contextual blocks without a corresponding mechanism to reduce the output space. Each individually justified addition to the analytical architecture increases the probability that the LLM output will resolve to a hedge: a qualified, multi-conditional statement that acknowledges every signal and commits to nothing. The project manager reads three paragraphs and finds no single actionable conclusion. The orientation system has produced disorientation.

The Gödel-Turing constraint and the Zero constraint are related but distinct. Gödel-Turing concerns the epistemic ceiling — what the machine can in principle know about its own outputs. Zero concerns the operational floor — the minimum condition for an output to remain useful. A system can satisfy the Gödel-Turing constraint (by providing an external deterministic signal layer) and still fail the Zero constraint by producing outputs too complex for a human decision-maker to act on. Both constraints must be satisfied simultaneously. The human occupies the space between: as the external validator required by Gödel-Turing, and as the necessary simplifier required by Zero.

The PRIMMS architecture addresses the Zero constraint through a specific mechanism: governing equations. These are formal mathematical structures — derived from the Bayesian, microeconomic, and signal-processing frameworks that constitute the system’s theoretical foundation — that impose hard boundaries on what the LLM layer is permitted to treat as ambiguous. They function as axioms: propositions the system treats as settled, not as inputs for further deliberation.

Three governing structures are central. First, the Bayesian weight-of-evidence accumulator: signals are converted to decibans and summed. The Jeffreys classification bands — Barely (0–5 dB), Substantial (5–10), Strong (10–15), Very Strong (15–20), Decisive (>20) — translate a potentially infinite evidence space into five ordinal categories. The LLM does not re-derive the evidence strength; it receives a computed value and a named band, and its interpretive task begins from that anchored point. Second, the project production function established in Aaron (2015): the microeconomic relationship $Y = f(AC, IC, IO)$, expressing project output as a function of activities completed, issues closed, and issues opened, provides a formal constraint on how scope, schedule, quality, and issue dynamics interact. This equation governs the PRIMMS issue management layer: the Bayesian issue closure ratio $\theta = \text{closed} / \text{total}$ is not an informal indicator but a formal metric derived from the production function, and is its operational expression in the deployed system — the form in which the production function’s theoretical constraint is actually computed and enforced at each session. θ carries an established readiness threshold ($\theta \geq 0.70$ for even odds; $\theta \geq 0.95$ for phase exit) that the LLM is instructed to treat as a pre-condition, not a deliberative question. Third, the velocity-based rundown projection: the linear regression of actual delivery velocity against the planned rundown produces a deterministic projected finish date. The LLM does not estimate when the project will finish; it receives a computed date and a computed slip in days, and its narrative task is to interpret that number, not to derive it.

The governing equations matter for a reason that goes beyond computational efficiency. They define the boundary between what is settled and what requires judgment. Without that boundary, the LLM treats every dimension of the problem as open for deliberation — and in a sufficiently complex project state, every dimension has genuine uncertainty, which means every statement can be qualified, and Zero is approached. With the governing equations in place, the LLM receives a pre-classified evidence state: a dB total and Jeffreys band, a θ ratio and readiness assessment, a projected finish date and recoverability label. Its task shifts from “assess the evidence” to “interpret a classified state and identify the one action that addresses the dominant risk.” That is a task the LLM can perform with specificity. The hedge-forest is replaced by a named archetype, a recovery posture, and a single owned action.

This is the operational meaning of hybrid cognition in PRIMMS. The machine does not consider everything. It considers what the governing equations leave open after they have been applied — which is the subset of the problem space that genuinely requires natural language reasoning, pattern recognition, and narrative synthesis. The MATLAB layer is not merely an input source for the LLM; it is a complexity reducer. It closes the questions that have formal answers so that the LLM can apply its interpretive capability to the questions that do not.

The practical expression of this principle in PRIMMS’s issue management integration illustrates the mechanism concretely. When the system detects open issues older than two weeks, it does not ask the LLM to assess issue severity and derive a priority. Instead, the governing equation θ has already computed the closure ratio; the stale issue count and age are deterministic facts; and the Cognitive hierarchy receives a pre-condition instruction that reduces to a single mandate: G1, the most important action this week, must directly address the oldest stale issue, stated in the form “Owner does X by date, measurable by Z.” If open issues reach ten or more, a Zero

constraint flag fires: one of the three executive decisions must address issue volume control, because the governing equation registers that the issue count has crossed the threshold where complexity begins to impair the team's decision capacity. The system detects its own approach toward Zero and names it — which is the one thing Zero itself could not do.

This self-monitoring property — the capacity to detect and report that issue accumulation has reached a decision-impairing threshold — represents a qualitatively different function from conventional project risk monitoring. The system is not merely reporting that issues exist. It is applying a formal threshold derived from the production function to assess whether the issue load has crossed from a tractable project management challenge into a structural decision-impairing condition. That assessment is deterministic, not narrative. It happens before the LLM is consulted, and it shapes what the LLM is permitted to say. Zero is avoided not by limiting what the system knows, but by structuring what it is required to conclude.

The broader implication for hybrid cognition system design is that governing equations are not merely computational conveniences. They are the mechanism by which complexity is made tractable without being made false. The project environment does not become simple when the governing equations are applied; it remains as complex as it is. But the portion of that complexity assigned to machine deliberation shrinks to what formal structures cannot resolve, and the portion assigned to human judgment expands correspondingly — which is exactly what the Janis-Mann vigilant decision-making framework requires. The human decision-maker does not need the formal system to be complete. They need it to be bounded. Governing equations provide that bound. Zero is the consequence of their absence.

6.7 The Recursive Application of OCC-5C: The Framework Applied to Itself

The PRIMMS phase-gate quality review introduces a structural property that deserves explicit identification: the OCC-5C framework is applied recursively to the system that uses it.

In its standard application, the OCC-5C framework is a quality governance methodology for evaluating AI-generated outputs — the five dimensions (Consistency, Coherence, Contradiction probe, Confabulation check, Calibration) are applied by an analyst or governance reviewer to assess whether an AI system's output meets the evidence standards required for consequential use. PRIMMS's phase-gate architecture inverts this relationship: the same five dimensions are now applied by the system itself, in an adversarial challenger role, to evaluate the quality of its own phase-exit recommendation before it reaches the sponsor.

This recursion is not merely a design curiosity. It has structural significance for three reasons.

First, it closes the governance loop at the point where it matters most. The Gödel-Turing constraint, articulated in Section 6.4, establishes that inductive systems cannot fully regulate themselves from within using the same mechanisms that generate their outputs. The OCC-5C challenger in the phase-gate architecture does not violate this constraint because it operates under deterministic MATLAB anchors that are external to the LLM generating the phase classification. The challenger LLM and the analyst LLM are the same substrate, but they receive

different inputs — the analyst receives the project evidence; the challenger receives the analyst's output and a set of MATLAB-computed governing quantities. The external anchor is what makes the recursion tractable. Without the θ thresholds, the longitudinal dB series, and the recency-gap computation as pre-computed constraints, the challenger would be one LLM evaluating another LLM's output using the same inductive mechanism — which is precisely the structure the Gödel-Turing constraint identifies as fundamentally limited.

Second, it applies the Zero constraint to the feedback loop itself. The Zero constraint, articulated in Section 6.6, holds that complexity accumulation without governing equations produces analytical indeterminacy. A naive recursive architecture — LLM output feeding back into another LLM prompt without structural constraints — is a Zero violation: each cycle adds context without reducing the deliberation space, and the result is hedge-forest proliferation rather than sharpened orientation. The PG-2 challenger prompt's explicit scope restriction — only five structured verdicts, one counter-candidate, and one adjustment recommendation, no narrative beyond this structure — is the governing equation applied to the recursion itself. The challenger cannot deliberate freely; it is constrained to evaluate five specific, MATLAB-anchored dimensions and produce a bounded output. This is how recursive reasoning is made tractable without degenerating into Zero.

Third, it institutionalizes falsification at the governance level. The Popperian argument developed in Section 6.1 — that the explicit treatment of Premature Closure as a failure archetype represents an operationalization of falsification discipline in governance practice — extends here to the institutional decision-making process itself. The C5 recency-bias check is a formal falsification of the “things got better at the end, so we are ready” rationalization that Janis and Mann (1977) identify as the most common form of defensive avoidance in high-stakes organizational decisions. By computing the recency gap before the LLM is consulted and injecting it as an explicit anchor, the system applies the Popperian asymmetry — no number of favorable recent sessions can establish readiness, but a 15 dB recency gap can challenge it — at the governance decision point where it is most consequential and most typically absent.

The deeper implication is that the OCC-5C framework and PRIMMS now stand in a relationship of mutual governance: the OCC framework governs the use of PRIMMS (externally, by practitioners assessing whether the system's weekly outputs meet quality standards), and PRIMMS applies the OCC framework to govern the institutional decisions the system informs (internally, via the phase-gate adversarial challenger). This recursive structure — the quality framework governing the tool, the tool applying the framework to the decisions — is the architectural expression of the paper's central argument: that the most consequential application of AI in project governance is not the replacement of judgment but the disciplined structuring of the context in which judgment is exercised. At a phase boundary, that context includes not just the current project state but the full evidence history of the phase, the formal evidence thresholds established in prior work, and a structured attempt to falsify the exit recommendation before it becomes an institutional commitment.

6.8 Beyond Project Management: Management in General as a Cybernetic, Management-by-Exception Process

Section 6.5 grounded the human authority boundary in a claim specific to command and control as Boyd and the RAND command-concepts research (Builder, Bankes, & Nordin, 1999) described it: system performance depends on the intellectual quality of the commander's concept, not on the volume or speed of information, and no amount of machine-supplied data can substitute for that concept. Nothing in that argument is actually specific to project management, or to command in the military sense. It is a special case of a more general structural claim about management itself.

Any managerial activity that sets an objective, senses deviation from that objective, and directs correction is, in Norbert Wiener's original cybernetic sense, a control loop operating on management by exception: sense, compare against plan, signal the exception, correct, repeat. A sales leader tracking pipeline against quota, a hospital administrator tracking bed utilization against capacity, a central bank tracking inflation against target, and a project manager tracking earned value against baseline are all instances of the identical loop. What differs across these cases is not the loop's structure but the density of its instrumentation — how explicitly the sensing and comparison steps have been made visible, measured, and recorded. Project management is simply the domain in which that instrumentation has been pushed furthest: Gantt charts, milestones, RACI matrices, quality gates, and status reporting exist precisely to make the sensing and comparison steps of the loop explicit and auditable in a way that most other managerial domains have never bothered to formalize.

This reframing changes what PRIMMS should be understood to be an instance of. The architecture is not a project management technique that happens to use AI; it is one instantiation of a general answer to a general question — how an exception-driven management loop should be instrumented once machine interpretation becomes available at the sensing and comparison steps, while command itself remains where Boyd and RAND's evidence says it must remain: with the human who holds the concept and bears the consequence. Project management is the domain in which this paper develops and tests that answer because it is the domain where the loop is most instrumented and the evidence of gaps is most legible — not because project management is a special case exempted from the pattern that governs the rest of management.

7. Limitations and Future Work

7.1 Current Limitations

The PRIMMS architecture has several limitations that should be acknowledged explicitly.

Signal independence assumption. The Bayesian WoE summation assumes that individual signals carry independent evidence about the risk hypothesis. In practice, schedule signals are correlated: a project with high days-behind will also typically have a high jeopardy index and a worsening VoT trend. The independence assumption therefore inflates the prior db relative to

what a fully specified joint probability model would produce. The practical consequence is that the absolute db value should be interpreted as a relative evidence accumulator rather than a precisely calibrated posterior probability. The Jeffreys band classification, which uses ordinal thresholds rather than continuous probability, partially mitigates this concern.

Retrospective design validation. The validation described in Section 5 is retrospective — PRIMMS was applied to completed project data after the fact, not deployed in real time during execution. This means the study can establish that the signal architecture correctly classifies a known-critical snapshot and that the recovery window calculation is consistent with ground truth, but it cannot establish that the system would have changed decisions had it been running in real time. A prospective deployment study — the highest-priority item in the future work agenda — is required before causal claims can be made.

LLM non-determinism. The MATLAB signal layer produces fully deterministic, reproducible outputs. The Cognitive hierarchy LLM outputs do not — the same prompt will produce different natural-language outputs across sessions. The audit trail records the full response text for each session, enabling governance reviewers to assess consistency. However, this non-determinism means that the LLM layer cannot be formally validated in the way that deterministic algorithms can. The architecture's response is to confine non-deterministic processing to the interpretation and synthesis layers while keeping all signal computation deterministic.

Phase-gate review: longitudinal signal dependency. The Phase-Gate Quality Review architecture requires a sufficient session history in the CorticalLog audit table to compute meaningful phase-level metrics. A project with fewer than three logged sessions will produce inconclusive longitudinal trend data, and the recency-gap computation requires enough sessions to distinguish a genuine improvement trend from a single favorable session. In early-phase deployments or projects where weekly sessions were not consistently logged, the MATLAB pre-verdict will default to CONDITIONAL EXIT with a manual verification advisory rather than a deterministic threshold verdict. The dependency on consistent session logging is both a limitation of the current implementation and an incentive for systematic weekly session discipline — the governance value of the phase-gate review is proportional to the quality of the accumulated evidence it synthesizes.

Recency threshold calibration. The 15 dB recency-gap threshold for the C5 calibration flag is a design judgment, not an empirically calibrated value. It reflects the order-of-magnitude reasoning that a gap of this magnitude — comparable to the weight of a single strong signal — represents a materially different evidence state between recent and phase-average performance. Calibration against a retrospective dataset of completed project phases with known outcomes is a planned component of the prospective validation program and may yield a different threshold.

Reasoning Map hypothesis cardinality. The bounded hypothesis structure — exactly one Current Best plus a field of two to four additional Rival, Tail Rival, or Combination nodes — is a design choice motivated by the Zero constraint, not an empirically derived optimum. A small number of at-risk projects may have a genuinely unambiguous single cause, in which case the additional nodes are populated by lower-confidence placeholders rather than substantive

competing accounts; the architecture does not yet distinguish this case from genuine multi-way ambiguity, and a reviewer should treat a low-confidence Rival or Tail Rival as a signal to check the evidence floor directly rather than as a finding in its own right.

Missing RL adaptation. The absence of a live reinforcement learning loop means that signal weights cannot adapt to the specific organizational context over time. A project environment with distinctive risk characteristics — unusually high scope volatility, or unusually strong governance — cannot be captured in the fixed WoE weights. This is the most significant gap versus original design intent and the highest-priority item for future development.

Differential validation depth across architectural layers. The validation evidence reported in Section 5 encompasses the full PRIMMS architecture, but the individual layers carry materially different levels of empirical support. The Bayesian WoE core — the VoT mechanism, the risk category benchmarks, and the signal threshold calibrations — has been tested prospectively across complete project lifecycles in operational deployment since 2008. The 20 risk categories identified through text mining of project risk logs, and the subset found statistically predictive of project schedule performance, represent a substantial empirical base accumulated over more than seventeen years of production use. The semantic document analysis layer and the Cognitive hierarchy prompt architecture, by contrast, have been validated retrospectively against completed project phases rather than prospectively across full project lifecycles. Their recommendations have been assessed as coherent and appropriate in retrospective review, but the evidence base is thinner. This layered validation status is an honest characterization of where the research program stands: the perceptual foundation is mature; the interpretive superstructure, while architecturally sound and retrospectively validated, requires prospective deployment to establish the same depth of empirical grounding as the Bayesian core.

Architectural complexity and verification scope. The signal interpretation architecture of PRIMMS operates across two dimensions: a set of multi-modal input signals computed by MATLAB across the schedule, stakeholder, contractual, and team domains, and a five-layer Cognitive hierarchy through which those signals are sequentially processed — from raw feature detection, through gestalt pattern recognition and temporal situation awareness, to executive planning and briefing document generation. Each layer operates at a higher level of abstraction than its predecessor, and each layer's output propagates forward to condition the next. A further constraint on the state space is provided by the risk ontology embedded in the MATLAB application: a library of 47 practitioner-validated risk categories — covering domains from data readiness and scope volatility to AI model quality, vendor risk, and governance framework gaps — developed from thirty years of program delivery experience across IT and AI/automation projects. This pre-classified ontology means that the LLM does not face an open-ended categorization task at each session; the combinatorial space of risk interpretations is bounded in advance by a structured, experience-grounded taxonomy. This directly mitigates the Zero problem described in Section 6.5: the LLM receives a named, bounded risk category as input, not an undifferentiated corpus from which it must independently induce structure. Verifying that the system's recommendations are appropriate across the full combinatorial space of signal inputs and layer interactions is a substantial challenge: the number of distinct contextual situations is

large, and retrospective case validation samples only a small fraction of them. This architectural complexity is a deliberate design choice, grounded in the principle that no single prompt can simultaneously perform reliable feature validation, pattern recognition, temporal projection, and executive planning — and in the neuroscientific evidence that hierarchical, multi-level signal integration is necessary for accurate situational perception in complex environments. However, it means that the verification effort required to establish the system's reliability across its full operating range is correspondingly large. Systematic coverage testing — in which specific combinations of signal inputs are constructed and the resulting layer-by-layer outputs evaluated for coherence and correctness — is an explicit component of the future validation program.

The raw-upload alternative. A question that any intellectually honest evaluation of PRIMMS must address directly is whether the architecture is necessary at all. A project manager or sponsor who has access to a capable general-purpose LLM — such as Claude (Anthropic) or GPT-4o (OpenAI) — can upload a schedule export, a stack of status reports, and a risk log, and receive a coherent synthesis. The LLM will surface alarming language, identify themes, and answer follow-up questions. The marginal cost of this approach is near zero, and no application infrastructure is required. The question is legitimate and deserves a direct answer rather than dismissal.

The raw-upload approach has genuine value and should not be caricatured. An experienced project manager who understands failure archetypes, frames the right questions, and possesses sufficient self-awareness to resist optimism bias will extract real analytical utility from an unstructured upload. For ad hoc diagnostic queries — a one-time situation assessment, a quick second opinion on a recovery hypothesis — the approach is appropriate and efficient. The claim of this paper is not that raw-upload analysis is useless. It is that raw-upload analysis fails to satisfy two architectural constraints that become consequential precisely in the high-stakes, time-pressured governance situations PRIMMS is designed to address.

The first constraint is the Zero constraint described in Section 6.5. When a raw document corpus is uploaded without a pre-classification layer, the LLM is asked to perform two tasks simultaneously: determine what the evidence says and interpret what the evidence means. In a project environment with genuine complexity — multiple schedule signals pulling in different directions, VoT sentiment diverging from documentary tone, recovery plan language present but unverified against schedule facts — the evidence space is large enough that every signal carries real uncertainty. Absent governing equations that close the questions with formal answers, the LLM tends toward hedge-forest outputs: multi-conditional paragraphs that acknowledge every tension and commit to nothing. The project manager reads three pages and finds no single owned action. The analytical tool has produced disorientation rather than orientation. PRIMMS's MATLAB layer addresses this by computing the db total, classifying it against the Jeffreys bands, applying the production function_ ratio, and projecting the velocity-based finish date before the LLM is consulted. The LLM receives a pre-classified state and is directed to interpret a named archetype and identify a single dominant recovery action. That task produces specificity. The raw-upload approach has no equivalent mechanism for forcing this reduction.

The second constraint is auditability. A MATLAB-computed db signal total is a number with a fully traceable derivation: each signal contribution is fixed, inspectable, and reproducible. If a governance reviewer asks why the system classified the project at 65.5 dB decisive jeopardy, the arithmetic is available. If a sponsor challenges the recovery window calculation, the velocity regression parameters are in the audit trail. When Claude or any other LLM derives the same conclusion from a raw upload, the reasoning is a natural-language narrative that cannot be reconstructed, formally reviewed, or independently verified. For project governance contexts — where decisions about resource reallocation, executive escalation, and program termination carry material organizational consequences — the inability to audit the analytical basis is not a minor inconvenience. It is a structural governance gap. The ISO 21502 standard and PMI governance frameworks both require that risk classifications be traceable to their evidentiary basis. A raw-upload synthesis, however coherent, does not satisfy that requirement.

A third practical distinction concerns temporal continuity. Risk management is an inherently longitudinal problem: what matters is not the signal state at a single point in time but how that state has been changing — whether confidence is eroding or stabilizing, whether the schedule gap is widening or recovering, whether the db accumulation is accelerating toward a decisive classification or plateauing. PRIMMS maintains this longitudinal record as a structural property of the application. Each weekly session appends to the accumulated signal history: the VoT time series of discrete confidence ratings, the progressive db trend across reporting periods, the velocity regression updated against each new schedule snapshot, and the Bayesian WoE accumulation itself, which derives its evidential force precisely from the compounding of independent signals across time. A project manager attempting to replicate this outside the application would face a substantial and error-prone manual burden: assembling multiple weeks of schedule exports, reconstructing the db computation from raw signal inputs, recalculating the velocity model from scratch, and presenting the resulting trend data to the LLM in a form that is coherent enough for pattern recognition across time. In practice, this burden means that ad hoc raw-upload analysis defaults to a snapshot assessment — the current state, evaluated in isolation. Snapshot assessment is precisely what the failure literature identifies as inadequate: the signals of structural deterioration are characteristically gradual, and their significance lies in their trajectory, not their instantaneous value. A project whose db total has moved from 8 dB to 22 dB over four weeks presents a qualitatively different governance situation from one that has held at 22 dB for a month, even though the current snapshot is identical. PRIMMS surfaces that distinction automatically. Raw-upload analysis, without heroic manual preparation by the user, does not.

Finally, raw-upload analysis is inherently prompt-dependent. Its quality is a function of the framing choices made by whoever conducts the upload: which documents are included, which questions are asked, which assumptions are made explicit. A skilled analyst who understands failure archetypes and knows what to look for will produce better results than a less experienced one. PRIMMS encodes that analytical expertise into a replicable, consistent process. The five-layer Cognitive hierarchy ensures that the same diagnostic questions are asked in the same order against the same pre-classified signal set, regardless of who operates the system. The value

proposition is not that PRIMMS is analytically superior to an expert analyst with a raw upload; it is that it is consistently equivalent to one — and that it removes the dependency on the expertise and optimism-awareness of the individual operator at each session. In governance terms, consistency and reproducibility are not secondary virtues. They are the conditions under which institutional trust in an analytical system can be established and maintained.

7.2 Future Work

The following extensions are identified for subsequent development phases:

- Live RL adaptation layer: implementing a DDPG/Q-learning agent that updates signal weights from accumulated weekly VoT survey responses across multiple projects, enabling the WoE framework to adapt to organizational context over time.
- Image-based rundown curve classifier: training a visual ML classifier on historical project rundown curve images to detect jeopardy signatures from curve geometry — the original design intent described in Aaron (2026, §8.7.3).
- Automated notification dispatch: extending the current human-delivery model to enable push notifications keyed to WBS Level 2 deliverable status changes.
- Portfolio-level deployment: extending the single-project architecture to portfolio monitoring, enabling pattern-level comparisons across projects and portfolio-level archetype profiling.
- Prospective deployment study: the highest-priority next step. A pre-registered study deploying PRIMMS in real time on active projects, comparing governance decisions and outcomes against matched controls. The retrospective validation program reported in this paper establishes the signal architecture's accuracy against ground truth across multiple completed programs; prospective deployment is required to establish that real-time use changes governance decisions.
- Calibration study: comparing PRIMMS db totals against eventual project outcomes across a retrospective dataset to assess the calibration of the WoE scoring system and refine signal weights.
- Phase-gate empirical validation: the Phase-Gate Quality Review architecture, like the Cognitive hierarchy, has been designed and implemented but not yet validated against a dataset of completed project phases with known outcomes. The validation question is whether the governing-equation thresholds ($\theta \geq 0.95$ for EXIT, $\theta \geq 0.70$ for CONDITIONAL EXIT, $\theta < 0.70$ for HOLD; recency gap > 15 dB for C5 flag) correctly classify phase boundaries in retrospective data — specifically, whether projects that satisfied these thresholds at their gates produced better outcomes in subsequent phases than those that did not. This calibration study is the phase-gate analogue of the prospective deployment study identified as the highest-priority item for the weekly monitoring architecture. It is the second-highest priority item in the future validation program.

- Cross-phase pattern analysis: the phase-gate architecture accumulates evidence across sessions within a single phase. A natural extension is cross-phase pattern analysis: examining whether specific archetypes that dominated Phase 1 of a project predict failure modes in Phase 2. This would require the CorticalLog to store phase-boundary markers (currently not implemented) and a cross-phase aggregation function analogous to `compute_phase_gate_metrics()`. The information required to detect such patterns is already being captured in existing session records; the extension is an analytical and structural one rather than requiring new data collection.
- Reasoning Map calibration study: comparing the Current Best hypothesis's Decision-Pathway-confirmed outcome against the LLM's initial Leading / Live / Possible / Tail ranking across a retrospective dataset, to assess whether the qualitative ranking is a reliable predictor of which hypothesis the diagnostic tests ultimately confirm.
- Mundell comparative statics formalization: the microeconomic production function framework established in Aaron (2015) — relating project output to scope, time, cost, quality, and issue dynamics — creates a tractable basis for Mundell-style comparative statics analysis of PRIMMS intervention options. Each COA option in the Layer 4 matrix can be analyzed as a shift in one parameter (e.g., resource input / force u) holding others constant, with the db signal set providing the empirical basis for estimating the relevant partial derivatives. Formalizing this connection would ground the COA recommendations in the same microeconomic framework that governs the project production function, enabling optimization of intervention sequencing under budget and time constraints.
- Likelihood-ratio ML layer for “what-if” COA prediction: the Bayesian WoE foundation that underpins PRIMMS is mathematically compatible with a broader family of modern machine learning classifiers whose inference also reduces to likelihood ratios — including regularized logistic regression, gradient boosting, and naïve Bayes models. A natural extension of the current architecture is to train such classifiers on the accumulated historical project record to produce “what-if” conditional predictions: given that the project manager takes Course of Action X at the current signal state, what is the posterior probability distribution of outcomes at +2, +4, and +8 weeks? This would transform the current COA matrix from a qualitatively reasoned set of options into a quantitatively ranked decision surface grounded in empirical project outcome data, while preserving the interpretable likelihood-ratio structure that makes the WoE framework auditable.
- Expanded multi-modal input architecture: the simulation work reported in Aaron (2019) demonstrated that adding successive input modalities — image, audio, and text — to a machine learning system progressively reduced classification error from 27% (image only) to 18% (image plus audio) to 0% (all three modalities combined) on a real-time operational decision task. The implication for PRIMMS is that the current three-modality architecture (schedule, VoT, documents) may itself be extendable: structured data feeds from financial systems, audio or video sentiment analysis from team meetings, and integration with collaboration platform message streams represent candidate additional modalities that could further increase the system's predictive precision. The Aaron (2019) multi-modal

framework provides the experimental methodology for evaluating the marginal accuracy contribution of each additional source before production integration.

- Generalization to broader organizational intelligence applications: the hybrid cognition architecture described in this paper — multi-modal signal fusion, Bayesian WoE classification, cognitive hierarchy interpretation, and human authority boundary — is not inherently specific to project management. The same design principles apply wherever an organization requires early structural perception of deteriorating conditions that do not yet appear in formal reporting: supply chain resilience monitoring, regulatory compliance posture, M&A integration health, and executive talent risk are candidate domains. The author intends to develop specialty applications of the intelligence engine architecture for organizational decision contexts beyond project management, building on the PRIMMS foundation established here.

8. Conclusion

This paper has argued that the most consequential application of AI in project management is not task automation but structural perception — the capacity to detect what is actually happening inside a project before it becomes visible in delivery output. Conventional project management tools, however sophisticated, are record-keeping systems. They capture what has happened and display it. They do not infer the structural conditions — team confidence erosion, perception distortion, governance gap formation — that precede most project crises.

PRIMMS addresses this intelligence gap through a hybrid cognition architecture that couples deterministic Bayesian weight-of-evidence signal computation with a five-layer LLM interpretation hierarchy, classifies convergent signal patterns against six empirically grounded failure archetypes, produces recovery trajectory projections at +2/+4/+8-week horizons, and generates sponsor-ready governance documents — while preserving an inviolable human authority boundary at every decision point. The architecture is designed to satisfy two foundational constraints simultaneously: the Gödel-Turing constraint, which requires that deterministic bounds on machine output originate outside the inductive system, and the Zero constraint, which requires that complexity accumulation not collapse the output into indeterminacy. The governing equations of the system — the Bayesian WoE accumulator, the project production function relating output to scope, issue closure, and schedule adherence, and the Bayesian issue closure ratio θ — are the mechanism by which both constraints are satisfied. They close the questions that have formal answers, so that the LLM is directed only at the questions that do not.

The architecture makes three claims that are advanced beyond the level of prior work. First, it operationalizes the hybrid cognition principle at a level of architectural specificity sufficient for replication and validation. Second, it introduces Premature Closure as a first-class failure archetype and the recovery window as a first-class governance metric — the insight that a recoverable situation prematurely closed as a failure is itself a failure mode with characteristic signals and a specific remedy. Third, it demonstrates through a retrospective validation that the signal architecture correctly classifies a known-critical project snapshot, that the velocity-based

recovery window calculation is consistent with known outcomes, and that the Cognitive hierarchy produces governance documents whose COA recommendations align with the interventions that produced successful project recovery. These are necessary preconditions for prospective deployment — not yet sufficient to establish real-time governance value, but sufficient to justify the controlled prospective study that is the next step in this research program.

A second contribution, reported here as an architectural extension subsequent to the core system, advances the governance argument one step further. The Phase-Gate Quality Review applies the OCC-5C framework recursively — not to external AI outputs, but to the system's own phase-exit recommendations. At each formal project phase boundary, PRIMMS aggregates the full accumulated session history into phase-level evidence metrics, generates a phase classification, and then challenges that classification using the same five quality dimensions it would apply to any AI-generated governance output. The challenger is constrained by the same governing equations that govern the broader architecture — θ thresholds from Aaron (2015), recency-gap computation, PREMATURE_CLOSURE persistence rule — so the recursion does not violate the Zero constraint. The result is a governance architecture in which the quality framework and the project intelligence system stand in mutual relationship: the OCC framework governs the use of PRIMMS, and PRIMMS applies the OCC framework to govern the institutional decisions it informs. This recursive structure is not an engineering convenience. It is the architectural expression of the paper's central argument: that the most consequential application of AI in project governance is the disciplined structuring of the context in which human judgment is exercised, applied consistently at every level of the decision hierarchy — from weekly orientation sessions to the formal phase-exit commitments that determine what a project becomes. A parallel extension applies the same discipline to root-cause reasoning: the Reasoning Map and Decision Pathway generalize archetype classification and decision-gating into a falsifiable, evidence-anchored hypothesis structure, so that no course of action is recommended before the evidence has been given a genuine chance to prove it wrong.

The project management field has done what its current paradigm allows. The persistent failure rate after decades of tool and methodology investment suggests that the paradigm itself requires extension. PRIMMS represents one instantiation of what that extension might look like: not more sophisticated scheduling, not more comprehensive risk registers, not faster reporting — but a structured intelligence layer that extends the project manager's perceptual range across the signal modalities that matter, and routes those signals through an interpretive architecture that is both powerful and accountable. The retrospective validation program reported here establishes that the architecture produces correct signal classifications against ground truth across multiple completed programs, with one illustrative example presented in detail. Establishing that it changes governance decisions in real time is the work that remains. PRIMMS is, in the end, one applied instance of a broader mission — amplifying human cognition with AI for business decision making — tested here in a domain where the cost of failing to see clearly is immediate and measurable. The mission is larger than this system; extending it elsewhere is work still to come.

References

- Aaron, J. M. (2015). A Bayesian methodology for project issue management: A microeconomic and data analytic perspective. Milestone Planning and Research, Inc. Working paper.
- Aaron, J., Fischer, B., & Ricordati, T. (2015). Understanding the project management profession — A case study and discussion. Elmhurst College Department of Business and School for Professional Studies.
- Aaron, J., Bratta, C., & Smith, D. (1993). Achieving total project quality control using the Quality Gate method. Proceedings from the PMI 1993 Annual Symposium.
- Aaron, J., Fischer, B., & Kulich, J. (2015). Weighing of evidence: Applying Bayesian evidence weighing to project decision practice. Elmhurst College Department of Business / Milestone Planning and Research, Inc. Working paper.
- Aaron, J. M. (2019). Designing and deploying systems of intelligence into business process applications: Multi-modal machine learning for operational decision support. Milestone Planning and Research, Inc. Internal technical report.
- Aaron, J. M. (2026). The Inductive Enterprise: Machine-Assisted Perception, Hybrid Cognition, and the Architecture of Organizational Intelligence. Unpublished manuscript. Milestone Planning and Research, Inc.
- Aaron, J. M. (2017). Identifying emergent project risks and their performance impacts through text mining: A mixed methods study. Unpublished manuscript. Milestone Planning and Research, Inc.
- Bloch, M., Blumberg, S., & Laartz, J. (2012). Delivering large-scale IT projects on time, on budget, and on value. McKinsey Quarterly.
- Builder, C. H., Bankes, S. C., & Nordin, R. (1999). Command concepts: A theory derived from the practice of command and control. RAND Corporation.
- Budzier, A., & Flyvbjerg, B. (2011). Double whammy: How ICT projects are fooled by randomness and why they overshoot on time and budget. Oxford Working Papers in Information Management.
- Endsley, M. R. (1995). Toward a theory of situation awareness in dynamic systems. *Human Factors*, 37(1), 32–64.
- Flyvbjerg, B. (2017). *The Oxford Handbook of Megaproject Management*. Oxford University Press.
- Friston, K. (2010). The free-energy principle: A unified brain theory? *Nature Reviews Neuroscience*, 11(2), 127–138.
- Friston, K. (2005). A theory of cortical responses. *Philosophical Transactions of the Royal Society B: Biological Sciences*, 360(1456), 815–836.
- Fuster, J. M. (2008). *The Prefrontal Cortex* (4th ed.). Academic Press.

- Gold, J. I., & Shadlen, M. N. (2007). The neural basis of decision making. *Annual Review of Neuroscience*, 30, 535–574.
- Glimcher, P. W. (2011). *Foundations of Neuroeconomic Analysis*. Oxford University Press.
- Gödel, K. (1931). Über formal unentscheidbare Sätze der Principia Mathematica und verwandter Systeme I [On formally undecidable propositions of Principia Mathematica and related systems I]. *Monatshefte für Mathematik und Physik*, 38, 173–198.
- Good, I. J. (1950). *Probability and the Weighing of Evidence*. Griffin.
- Good, I. J. (1985). Weight of evidence: A brief survey. *Bayesian Statistics*, 2, 249–270.
- Ika, L. A., Love, P. E. D., & Pinto, J. K. (2020). Moving beyond the planning fallacy: The emergence of a new principle of project behavior. *IEEE Transactions on Engineering Management*, 67(3), 640–649.
- Janis, I. L., & Mann, L. (1977). *Decision Making: A Psychological Analysis of Conflict, Choice, and Commitment*. Free Press.
- Jeffreys, H. (1961). *Theory of Probability* (3rd ed.). Oxford University Press.
- Kahneman, D. (2011). *Thinking, Fast and Slow*. Farrar, Straus and Giroux.
- Kanerva, P. (1988). *Sparse Distributed Memory*. MIT Press.
- Kock, A., Gemünden, H. G., & Salter, A. (2021). An empirically grounded model of innovation project portfolio management: Understanding the role of learning and strategic alignment. *International Journal of Project Management*, 39(5), 586–599.
- Kutsch, E., & Hall, M. (2010). Deliberate ignorance in project risk management. *International Journal of Project Management*, 28(3), 245–255.
- Lipton, P. (2004). *Inference to the Best Explanation* (2nd ed.). Routledge.
- MacKay, D. J. C. (2003). *Information Theory, Inference, and Learning Algorithms*. Cambridge University Press.
- Maylor, H., Vidgen, R., & Carver, S. (2008). Managerial complexity in project-based operations: A grounded model and its implications for practice. *Project Management Journal*, 39(S1), S15–S26.
- Niazi, M., Zhang, X., & Bhatt, U. (2023). A systematic literature review on artificial intelligence in project management. *International Journal of Project Management*, 41(7), 102522.
- Pellerin, R., & Perrier, N. (2019). A review of methods, techniques and tools for project planning and control. *International Journal of Production Research*, 57(7), 2160–2178.
- Popper, K. (1959). *The Logic of Scientific Discovery*. Basic Books.
- Project Management Institute. (2023). *Pulse of the Profession*. Project Management Institute.

Standish Group. (2023). CHAOS Report 2023: Decision Latency Theory. Standish Group International.

Williams, T. (2005). Assessing and moving on from the dominant project management discourse in the light of project overruns. *IEEE Transactions on Engineering Management*, 52(4), 497–508.

Turing, A. M. (1936). On computable numbers, with an application to the Entscheidungsproblem. *Proceedings of the London Mathematical Society*, 42(1), 230–265.

Turing, A. M. (1950). Computing machinery and intelligence. *Mind*, 59(236), 433–460.

van Fraassen, B. C. (1989). *Laws and Symmetry*. Oxford University Press.

Author Note

John M. Aaron is the founder of Milestone Planning and Research, Inc., and the designer of the PRIMMS® platform. The research reported in this paper was conducted as part of the development of PRIMMS, an extension of the PRIMMS® platform. The author declares a commercial interest in PRIMMS. The retrospective validation described in Section 5 was conducted on project data from completed programs in which the author was involved; the illustrative example presented in Sections 5.2–5.4 is drawn from one such program. All project-identifiable data has been anonymized. No external research funding was received for this work.

Correspondence concerning this article should be addressed to John M. Aaron, Milestone Planning and Research, Inc. Email: john.aaron@milestoneplanning.net

Acknowledgements. The architecture described in this paper draws on the theoretical foundations of Bayesian inference, computational neuroscience, and decision theory as synthesized in *The Inductive Enterprise* (Aaron, 2026). The retrospective validation benefited from access to complete project records from multiple concluded enterprise programs.

Glossary of Terms

Bayesian Weight of Evidence (WoE) — A logarithmic representation of likelihood ratios used to accumulate evidence across independent signals. In PRIMMS, WoE is expressed in decibans and summed to produce an aggregate evidence classification.

Cognitive hierarchy — A structured, multi-layer prompt architecture inspired by hierarchical models of brain function. In PRIMMS, it organizes LLM reasoning into five sequential layers: feature detection, pattern recognition, temporal awareness, executive planning, and document generation.

Deciban (dB) — A unit representing $10 \times \log_{10}$ of a likelihood ratio. Enables additive accumulation of evidence across heterogeneous signals.

Decision Pathway — A sequenced, falsification-ordered diagnostic protocol that tests each Reasoning Map hypothesis — confirming or eliminating it via a named check — before a course of action is committed to.

Failure Archetype — A recurring structural pattern of project failure identified through multi-modal signal convergence.

Free Energy Principle (FEP) — A theoretical framework proposing that biological systems minimize variational free energy—an upper bound on surprise—through continual updating of internal models.

Generative Model — An internal probabilistic model that produces predictions about incoming data.

Gödel–Turing Constraint — The principle that sufficiently expressive systems cannot fully regulate or verify themselves internally.

Hybrid Cognition — An architecture in which machine intelligence extends human perception while preserving human authority over decisions.

Jeffreys Scale — A classification system for interpreting strength of evidence using logarithmic likelihood ratios (decibans).

OCC-5C — A five-dimension adversarial quality framework — Consistency, Coherence, Contradiction probe, Confabulation check, Calibration — used both to evaluate AI-generated governance outputs and, in the phase-gate architecture, applied recursively by the system to its own phase-exit recommendations.

Phase-Gate Quality Review — A formal evidence-threshold review, distinct from weekly monitoring, that aggregates a project phase's full session history to authorize, conditionally authorize, or hold phase exit.

Predictive Coding — A hierarchical inference framework in which systems minimize prediction error.

Premature Closure — A failure mode in which a recoverable project state is abandoned without testing recovery.

PRIMMS — Project Risk Identification Measurement and Mitigation System augmented with a Large Language Model.

Reasoning Map — A structured application of Inference to the Best Explanation in which one Current Best explanation plus a field of two to four ranked, falsifiable Rival, Tail Rival, or Combination hypotheses is generated below an explicit evidence floor, generalizing single-label failure archetype classification.

Recovery Window — A finite interval during which a project deviation remains correctable.

Reinforcement Learning (RL) — A machine learning paradigm in which an agent learns to select actions by optimizing a reward function through interaction with an environment over time. In the original PRIMMS design intent, RL was envisioned to dynamically update signal weights based on observed project outcomes and Voice of the Team feedback. In the current implementation, RL is not active; deterministic Bayesian scoring is used instead to preserve auditability.

Theta (θ) — The issue closure ratio (closed issues divided by total issues), used as a proxy for project health.

Variational Free Energy — A tractable bound on surprise used in approximate Bayesian inference.

Voice of the Team (VoT) — A structured sentiment signal collected via a 1–6 scale reflecting team confidence.

Zero Constraint — A system-level limitation where excessive signal accumulation leads to analytical indeterminacy.