

Building the Decision-Native Enterprise

An operating model for sustained, authorised action

ENTERPRISE AI

OPERATING MODEL

BOUNDED AUTHORITY

POLICY RENEWAL

ABSTRACT

Enterprise AI changes an organisation when its outputs become authorised actions. Sustaining that delegation requires both evidence that supports institutional confidence and standing authority that remains usable. This paper develops an operating response to two models of sustained delegation and renewable authority. The first shows how complementary controls and a delegation-evidence feedback can generate distinct entry and exit thresholds. The second shows when protecting policy-renewal effort, shared with case review, can sustain workloads that strict case priority cannot. Neither establishes the effectiveness of a particular enterprise architecture. This paper proposes one: five connected planes separating evidence, interpretation, policy and authority, execution, and assurance. Recurring decisions are managed through explicit, revocable authority envelopes, traceable policy compilation and a resourced renewal process. A worked procurement scenario, an implementation sequence and decision-level measures make the proposal operational. The contribution is a managerial synthesis of established decision-management, access-control and assurance practices, linked to the two models' distinct mechanisms. The architecture and its performance remain to be evaluated in organisations. Its objective is faster legitimate action within limits that accountable owners can explain, monitor, renew and withdraw.

KEYWORDS enterprise AI; decision systems; standing authority; policy renewal; operating model; human oversight.

1

From two models to an operating model

A model can produce a useful forecast without authorising a purchase, a draft response without permitting its release, or a risk signal without changing exposure. The practical unit of design in this paper is the recurring decision that commits the organisation to an action, allocation, classification, obligation or official position. Task productivity matters, but it does not by itself establish that decisions become faster, better or properly authorised.

The organisational adoption account in *Enterprise AI Adoption as an Organisational State Transition* [2] defines adoption through sustained, consequence-weighted delegation of consequential work without per-instance human ratification. Its minimal model combines complementary controls with a feedback between delegation, institutional evidence and confidence. Under the stated conditions, this produces alternative stable states and different entry and exit thresholds. The model is an existence proof. The prevalence of these dynamics in organisations remains an empirical question. Its managerial implication is to assess weak controls jointly and preserve usable evidence of outcomes through incidents and organisational change.

The queueing account in *The Authorisation Bottleneck* [3] examines a different mechanism. A standing mandate lets eligible cases bypass individual review while it is valid. When suspended, it must be renewed by a reviewer who also processes cases. At sufficiently high load, clearing cases first can prevent the renewal that would reduce future review work. Protecting renewal effort can then expand sustainable capacity. At lower load, renewal in otherwise idle time can be preferable, and immediate renewal can increase delay. The result depends on the eligible share, invalidation frequency, renewal effort and case-review capacity; it does not justify universal renewal priority.

The models are complementary, not a single derived theory. The adoption model's institutional evidence stock is not the renewal model's binary policy-validity state. The renewal model does not derive the adoption account's hysteresis, and neither proves that the operating design proposed here is necessary or sufficient for successful adoption. Together they motivate two management questions: what evidence sustains a defensible willingness to delegate, and what resources keep the resulting authority usable?

A **decision-native enterprise** treats recurring decisions, decision rights, executable policy and accountability evidence as managed operating assets. The term describes an architectural orientation, not a maturity rank. Strategic, contested, rights-affecting and irreversible decisions may remain human-led. The placement of human and machine authority should be explicit, justified and operable.

This proposal builds on established work. Human-AI allocation of decision authority is already an economic research problem [1]. Decision Model and Notation (DMN) formalises business decisions and rules [9]; attribute-based access control evaluates permissions against contextual attributes [6]; policy engines separate policy decisions from enforcement [10]. AI risk management already covers governance, measurement and continuing oversight [8]. Recent agent-governance

proposals also combine policy enforcement and accountability [7]. The contribution here is their operational integration around the decision, with evidence preservation and authority renewal made explicit. The five-plane decomposition is a design choice, not a newly proved architecture or a claim to have invented bounded delegation.

2 The decision-native operating architecture

Task automation produces or transforms an input. Workflow automation routes work through a process. Decision automation selects a permitted action and commits or triggers the organisational response. These categories can overlap within one system. Their boundary is institutional effect, not technical sophistication: a simple rule may approve a payment while an elaborate model remains advisory. If preparation falls from 40 to 10 minutes but approval still takes two days, task productivity improves much more than end-to-end decision time. These figures illustrate the distinction; they are not observed results.

The proposed architecture has five planes: evidence, interpretation, control, action, and assurance and learning. Policy and authority remain distinct tests within one control plane. Figure 1 shows the interfaces. The planes are functional responsibilities; they need not be separate products, teams or services.

Evidence. Capture the event that makes a decision necessary, the relevant records and their provenance. Specify source, freshness, completeness and permitted use. Missing evidence differs from adverse evidence: an absent sanctions check is not a clear result. Preserve both the initial request time and the time at which the required evidence becomes complete. Otherwise a system can appear faster simply by delaying the measurement start.

Interpretation. Models classify, retrieve, forecast, summarise or propose. Their output should identify the claim, supporting evidence, model and prompt version, evaluated use, uncertainty where meaningful, and known limits. A self-reported confidence score grants no authority; any uncertainty threshold used for routing needs validation on the intended population. Deterministic checks may sit alongside probabilistic models, and either can be wrong or applied outside its intended scope.

Control. The policy test asks whether the proposed action is permitted, required or prohibited. The authority test asks whether this actor holds a current delegation for that action, population and context, within current exposure limits. Passing one test does not imply passing the other. A purchase can comply with policy but exceed the service's mandate; a spending delegation cannot legalise a prohibited purchase. Source policy, compiled logic and authority grants require separately identifiable versions and owners.

Action. Controlled interfaces commit the organisation: place the order, issue the refund or release the communication. They must enforce the control result at execution, not merely display it. Bind permission to the intended action and actor, reject stale grants, and reserve cumulative exposure atomically where concurrent actions could otherwise exceed a cap. Use idempotency, cancellation and compensating actions where feasible. A successful pre-check is insufficient if the state changes before commitment.

Assurance and learning. Connect the trigger, evidence references, interpretation, policy and authority versions, action, intervention and eventual outcome. Record blocked and escalated attempts as well as executed actions. The decision ledger is the linked record, not a requirement for blockchain or central storage of raw data. Protected source references, access controls and retention rules should permit reconstruction without unnecessary duplication of sensitive material. Append-only amendments and tamper evidence can support integrity; they cannot make false input true.

Assurance must distinguish a record from useful institutional evidence. The adoption model's stock includes audited outcomes, calibrated expectations and post-mortems, not merely a count of logged actions. Outcomes may be delayed, disputed or confounded; record that status and avoid treating absence of a complaint as success. Model retraining, policy revision and delegation expansion are separate governance events. Improving a model does not itself authorise a broader mandate.

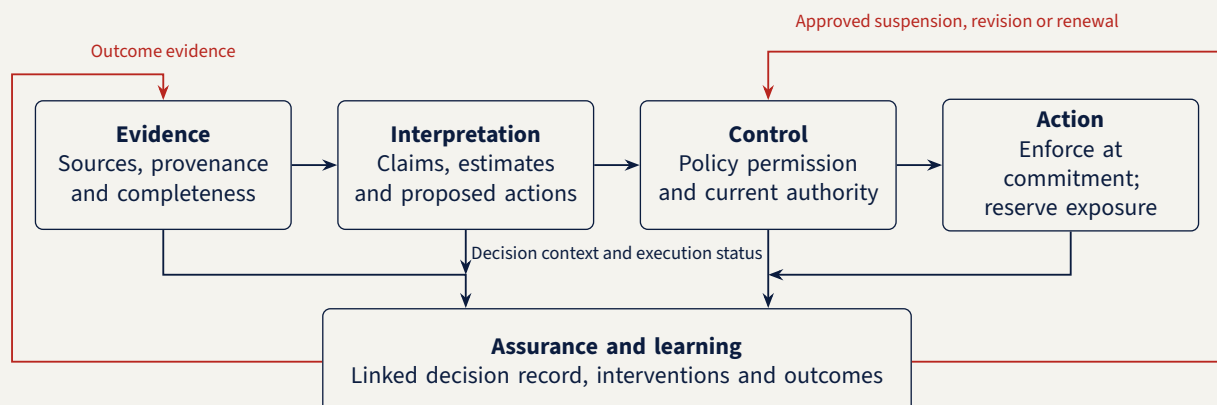


Figure 1. Five functional planes for bounded delegation. Policy and authority are separate tests within control. Every plane contributes to the linked decision record; assurance informs renewal and subsequent decisions. Feedback does not itself grant permission.

3 Authority must be bounded and maintained

An **authority envelope** specifies the machine-testable conditions under which a system may commit the organisation without further case approval. It implements an approved delegation; it cannot exhaust every legal obligation or convert an ambiguous source into determinate permission. Unresolved conditions require escalation or a safe block.

The envelope names the decision, permitted actions, affected population, required evidence, applicable policy, actor identity, operating context, per-action and aggregate limits, evaluated model range, expiry and suspension conditions. It also identifies the accountable outcome owner, renewal owner, intervention rights and recovery method. A procurement service might approve catalogue orders only for approved suppliers, unchanged terms, available budget and domestic delivery, below both individual and cumulative caps. It is not thereby authorised for procurement generally. Appendix B gives the operating specification.

A usable envelope needs a lifecycle. Source-policy change, supplier deterioration, model change or evidence failure may suspend some or all of its scope. The control plane stops new automatic commitments; assurance opens a renewal task with the reason, affected scope, required evidence, owner and deadline. Renewal restores permission only after the relevant owner approves the revalidated conditions. Suspension must not silently become expiry followed by automatic reactivation.

The renewal model shows why renewal cannot always be deferred until the case queue is empty [3]. In its one-reviewer model, let candidate arrivals have rate λ , full-effort case review rate μ , full-effort renewal rate ν , invalidation rate δ , and eligible share $x \in [0, 1]$. The four rates are positive and use the same time unit. Under the admissible controls and stationary stability definition in [3, Appendix B], the supremum stabilisable candidate-arrival rate is:

$$\lambda_{\text{cap}} = \max \left\{ \mu, \frac{\mu}{1 - x + \delta/\nu} \right\}. \quad (1)$$

Renewal expands capacity exactly when $\nu x > \delta$. This is a result of the companion model, not an organisational estimate. It assumes Poisson arrivals; mutually independent exponential case effort, renewal effort and valid-policy lifetimes; independent eligibility; immediately observed invalidation; divisible reviewer effort with preemptive resume; and admissible standing authority and manual fallback. The control set includes declining renewal and reviewing every case. Renewal changes routing for future arrivals; already queued cases remain for review. Capacity is a strict-load feasibility boundary, not a target service level. A stable queue may still impose unacceptable delay. Appendix G explains the resource balance and the protected-share condition; the stability proofs are in [3, Appendix C].

The operating implication is to measure case and renewal demand before allocating scarce review effort. Where the assumptions are informative and workload warrants it, protect renewal effort, potentially above a finite backlog trigger. At low load, useful case service should not be interrupted by a universal renewal-first rule. Shared reviewers may maintain several policies, and case difficulty may change after automation; the single-policy formula supplies no ready-made enterprise staffing rule. Simulation, observed workload and service constraints are needed for those settings.

Restoring permission is also different from correcting an unsafe policy. If an incident calls the policy's substance into question, renewal requires resolving that concern. A valid flag certifies neither accuracy nor lawfulness. If manual fallback is not independently admissible, suspension must lead to a safe hold or another approved process. It cannot route prohibited actions to a human for nominal approval.

4 Implementation and evaluation

The first implementation should be narrow enough to stop safely and important enough to expose the operating issues. Select recurring decisions with stable policy, meaningful waiting time, observable outcomes and a credible recovery path. Do not select solely for the proportion that could be automated. The decision inventory should include both case work and the effort needed to keep its governing authority current.

Diagnose before redesign. Establish end-to-end lead time from initial request, decision-ready time, stage-level waiting, case-review effort and renewal effort. For each approval, identify the risk controlled and the evidence or judgement added. Remove an approval only when its contribution is demonstrably redundant or replaced. Map the dependencies shared by several mandates; a single policy change can suspend many envelopes at once.

Compile and test. Translate determinate policy into versioned controls with traceability to the authoritative source. Resolve precedence and mark ambiguity explicitly. Test boundaries, conflicting rules, missing evidence, concurrent exposure, stale permissions and revocation. Shadow evaluation should distinguish model error from policy ambiguity, incomplete data and inconsistent human practice. Human agreement is not ground truth. Appendix C describes compilation and human oversight.

Pilot and maintain. Begin with a small eligible population or exposure cap. Exercise suspension, recovery and the renewal queue under representative load. Measure the reviewer budget consumed by cases, renewal, monitoring and appeals. Expansion should follow evidence rather than a calendar. Appendix D provides a 90-day planning sequence, not a validated delivery duration or a guarantee of readiness.

The procurement scenario in Appendix E illustrates the entire cycle. A supplier-status change suspends automatic ordering. Individually reviewed purchases proceed only through a separately permissible fallback. A named owner revalidates the mandate using protected effort where justified; renewed authority applies to future eligible requests. Restoring the mandate does not itself clear the existing queue or reverse commitments already made.

Evaluate the redesign against a credible comparator. Where feasible, compare improved case assistance, case-first renewal and protected renewal at the same total reviewer budget. Assignment should respect shared reviewers because cases compete for their time. Fix candidate populations upstream of treatment-dependent triage, eligibility definitions, outcome windows and acceptable quality margins in advance. Use a common assistance configuration to isolate renewal allocation; a comparison of entire operating designs instead estimates their combined effect. Use risk-adjusted lead times, useful completions, harmful approvals, missed exceptions, appeals and full operating costs. Track unfinished cases and delayed outcomes; do not exclude them because they make results less favourable.

The principal hypothesis is that, at a fixed reviewer budget and within a prespecified decision class, the proposed design increases useful, authorised completions per unit time while meeting prespecified quality, cost and service thresholds. A more specific renewal hypothesis is that protected effort improves authority availability and prevents sustained backlog growth where independently estimated conditions predict it. Evidence against the design includes a quality breach, costs exceeding the prespecified limit, or failure to improve useful completion in an adequately powered comparison. Failure of the predicted availability and queue response challenges the renewal mechanism in that application. An inconclusive study does not establish either success or failure. Short pilots cannot establish long-run stability or rare-event safety. Absence of a statistically significant quality difference is not evidence of equivalence.

5 Scope and conclusion

The design is not an automation mandate. Missing evidence, unstable policy, unreliable eligibility classification, unobservable invalidation, delayed outcomes and inadmissible fallback can invalidate the proposed operating comparison. Contested values, serious irreversible consequences, rare decisions and duties reserved to people may favour human-led deliberation. The architecture can still improve evidence and ownership without delegating the final choice.

Human involvement is useful when it supplies information, judgement, authority, legitimacy or a safeguard. It is not effective merely because a person clicks approve. Automation-bias research warns against equating nominal oversight with independent scrutiny [5, 12]. Article 14 of the EU AI Act specifies human-oversight measures for high-risk systems within its scope, including enabling oversight personnel to understand limitations, recognise over-reliance, interpret outputs, intervene and stop operation [4]. This is not a general permission to automate every other decision. Applicable duties and commencement provisions require assessment for the actual use.

The Robodebt Royal Commission's recommendations on automated decision-making emphasise lawful foundations, review, transparency and independent monitoring [11]. More efficient execution cannot repair defective authority or inaccessible recourse. Nor should an intact evidence record preserve confidence in a substantively unsafe practice: records must support challenge and withdrawal as well as continued use.

The adoption account explains how delegation and outcome evidence can sustain an organisational state. The renewal account explains how shared review capacity can maintain the authority on which delegation depends. This paper proposes the operating response: separate interpretation from permission, implement bounded authority, resource its renewal, preserve usable outcome evidence and measure decision performance alongside harm and cost. It contributes a practical synthesis and an evaluation design. Its effectiveness remains to be demonstrated in organisations.

A

Contribution boundary and design dependencies

The principal audience is executives, enterprise architects and governance owners responsible for consequential AI use. This is an operating-model paper with a worked scenario, not a new mathematical theory, a field evaluation or an implementation standard. The architecture can be implemented using existing decision, identity, policy and workflow systems. Procurement of a new platform is not a precondition.

Existing foundation	What is inherited	What this paper adds
Human-AI decision authority [1]	Allocation of decision rights between people and AI is already a research object.	Operating ownership and interfaces for a recurring decision class.
DMN [9]	Explicit business decisions, decision requirements and rules.	Connection to current mandates, evidence and renewal responsibilities.
Attribute-based access control [6] and policy engines [10]	Contextual permission and separation of policy evaluation from enforcement.	Business consequence, aggregate exposure and accountable delegation around the action.
NIST AI RMF [8] and agent governance [7]	Continuing governance, monitoring and policy control are established.	A managerial synthesis connecting the companion papers' evidence and renewal mechanisms.
Companion models [2, 3]	Conditional transition dynamics and shared-resource capacity results.	Testable operating hypotheses, implementation sequence and decision-level measures.

The argument is conditional: an admissible decision class and observable operating conditions support bounded delegation; explicit controls and ownership make that delegation operable; resourced renewal keeps it available; evaluated outcomes may produce usable institutional evidence. None of these steps guarantees the next. An impeccably logged action may have no observable outcome. A formally valid policy may encode a bad decision rule. Faster renewal may restore a mandate that should instead be retired.

The five planes need not map one-to-one to the adoption model's illustrative control dimensions. Verification, permissioning, observability, process design and staff capability cut across several planes. Likewise, the renewal model's eligible share is a routing parameter for one decision class, not the adoption model's realised consequence-weighted delegation measure. Maintaining these distinctions prevents architectural labels from being mistaken for measurements of the formal models.

B

Authority-envelope specification and runtime behaviour

Each envelope should have a stable identifier, version, accountable outcome owner and explicit institutional approval. A human-readable statement should express the same scope as its executable implementation. Approval must identify the evidence supporting delegation and the conditions that would invalidate it.

Specification	Minimum operating content
Decision and mandate	Named choice, permitted actions, affected population, source authority, accountable owner and approved actor identity.
Evidence and policy	Required sources, freshness, completeness, policy version, precedence and explicit prohibitions.
Scope and exposure	Jurisdiction, business unit, channel, per-action cap, cumulative cap, aggregation window and shared dependencies.
Evaluated operation	Model and configuration versions, supported population, uncertainty handling and evidence of performance.
Intervention and recovery	Step-up route, safe hold, appeal route, stop authority, cancellation boundary and compensating action.
Lifecycle	Effective time, expiry, suspension triggers, renewal evidence, renewal owner and approval record.

The enforcement path should bind the permission decision to the actual actor, action and relevant state. Identity verification must not be confused with business authority. Where several services share a cumulative cap, checking the current balance separately in each service is insufficient: reserve exposure in a common, concurrency-safe operation before commitment, and reconcile failures. Retries should not create duplicate commitments or consume inconsistent reservations.

Revocation exercises should distinguish work not yet authorised, authorised but not committed, and already committed. New commitments must stop within the specified operational bound. In-flight work requires cancellation or containment

where possible; completed actions require whatever remedy is available. An emergency stop does not erase external effects. If a control service is unavailable, enter the predefined safe state rather than treating its silence as permission.

The evidence record should preserve the original decision context and link corrections instead of silently replacing history. It should identify policy and mandate versions, the permission result, actual execution status and unresolved outcomes. Raw customer data may remain in systems of record if protected references and retention arrangements support later reconstruction. Test portability and reconstruction after vendor changes, not merely export of an unreadable log.

C

Policy compilation and meaningful human oversight

Source and decompose. Identify the controlling instrument, jurisdiction, owner, effective date and precedence. Separate definitions, permissions, duties, prohibitions, thresholds, evidence requirements, exceptions and appeal rights. Written and performed policy may disagree; that disagreement needs an accountable resolution.

Formalise and mark ambiguity. Express determinate provisions in decision tables or declarative policy. DMN supplies an established notation [9]. A phrase such as “reasonable” or “in the customer’s best interests” cannot safely become a numerical threshold merely because software needs one. Unresolved interpretation is an explicit escalation condition. Language models can assist extraction and comparison but cannot supply institutional approval for a semantic choice.

Test and approve. Replay representative and adverse cases, test rule boundaries and contradictions, and review differential impacts where relevant. Expert agreement helps identify candidate interpretations but does not establish lawfulness or fairness. Obtain approval from the policy owner and the relevant operational, legal or risk authority. Version the rule, its source mapping and the tests together.

Deploy, monitor and renew. Separate policy evaluation from enforcement, as established policy engines do [10]. Detect source changes and invalidate dependent controls and grants within an explicit bound. Create renewal work rather than leaving stale logic active. Renewal may confirm the old scope, narrow it, replace it or retire it. Model changes, policy changes and mandate changes each need their own approval path.

Three control modes are useful. Straight-through execution operates within the envelope, with monitoring and risk-based sampling. Step-up provides a qualified reviewer with the evidence, applicable rule, reason for escalation, available actions and deadline. Human-led work leaves framing and judgement with a person while AI assists research or preparation. These are functional modes, not a scale on which more autonomy is necessarily better.

Decision characteristics	Starting control design	Qualification
Stable rule, bounded consequence, prompt recovery and reliable evidence	Straight-through execution with exposure caps and sampling	Aggregate exposure and disparate impact can make small cases material.
Repeated operational choice with measurable outcomes and recoverability	Narrow envelope; step-up for specified anomalies	Classify exceptions independently of the model’s asserted confidence.
Material financial, customer, employee or regulatory consequence	Human decision or dual control supported by structured evidence	Involvement must add competence, authority or a safeguard.
Contested, rights-affecting, strategic or difficult-to-reverse choice	Human-led deliberation unless an applicable regime and governance design support another arrangement	This is a conservative design default, not a universal legal classification.

Sampling must cover apparently eligible cases, not only escalations. Otherwise missed exceptions remain invisible. Reviewers need enough time, competence, evidence and authority to disagree [5, 12]. The organisation must staff both difficult cases and mandate renewal; reducing routine work can increase the average difficulty of the remaining queue. Monitor that change rather than assuming a fixed saving in headcount.

D

Decision inventory and a 90-day implementation sequence

Decision debt is a portfolio label for recurring work that still demands bespoke handling despite an apparently stable, defensible policy or delegation pattern. It is an analogy, not an accounting liability or a newly established economic quantity. Repeated expert choices are a candidate for examination, not proof that the choices should be encoded. Redesign may automate a decision, simplify an approval while retaining a person, or confirm that apparent repetition conceals genuine ambiguity.

The inventory records the trigger, available actions, outcome owner, decision-maker, approval chain, policy, required evidence, candidate volume, eligibility, consequence, reversibility, exposure, stage-level latency, overrides, appeals, execution interface and recovery mechanism. Add policy dependencies, suspension events, renewal effort, renewal waiting time,

renewal owner and the reviewer pool shared with cases. Prioritisation should consider setup and continuing maintenance costs, not only current waiting time.

Days 1-30: diagnose and select. Choose a value stream such as purchase request to commitment or complaint to remedy. Establish request-to-outcome and decision-ready timestamps. Measure median and tail latency, waiting versus active work, unresolved cases, case-review demand and renewal demand. Use a stated follow-up horizon or a justified time-to-event method; completed cases alone can understate waiting. Identify the purpose of each approval. Select a bounded candidate and name the outcome, policy, risk, renewal and technical owners. If evidence or lawful authority is missing, resolve that before selecting an automation target.

Days 31-60: compile and test. Resolve source precedence, compile determinate rules, specify the envelope and define suspension and renewal routes. Test historical cases, rare events, stale grants, concurrent limits, missing records and inaccessible reviewers. Run in shadow mode and investigate disagreement against an independently approved decision standard. Estimate eligible share and renewal burden separately from model accuracy. A narrower eligible population can be the correct result.

Days 61-90: pilot bounded authority. Launch only if evidence and approvals warrant it. Begin with limited scope or exposure, review outcomes and exception age, and exercise stopping and recovery under load. Record actual renewal allocation and compare it with the proposed schedule. Test that suspended authority cannot execute and that renewal restores only its approved scope. Keep service and harm thresholds visible beside throughput.

Continuing operation. Revalidate model, policy and mandate separately. Widen only after representative performance, acceptable outcome quality and demonstrated intervention capacity. Narrow or suspend when source policy changes, outcomes deteriorate, eligibility becomes unreliable or the intervention path fails. Ninety days is a planning scaffold; long outcome windows, rare harms or complex approval requirements may make it insufficient.

E Worked procurement scenario

This is a constructed operating example, not a deployment result. An employee requests standard equipment from an approved supplier. The price is within the catalogue range, budget is available and delivery is domestic, but the request exceeds the requester's personal approval limit. Conventional processing may reassemble the same evidence across procurement, finance, legal and a director. Whether those approvals are redundant must be established, not assumed.

The decision is whether the organisation may commit this purchase under its current policy and delegation. Evidence includes the catalogue item, supplier status, budget, contract version, delivery jurisdiction, requester identity, cost centre, conflicts and aggregate exposure. A model may identify an unusual description or suspected order splitting. It supplies a signal, not permission. Missing evidence leaves the case incomplete while end-to-end waiting continues to be counted.

Policy checks category eligibility, approved terms, budget use, segregation of duties and jurisdiction. Authority checks the procurement service's mandate, current scope and expiry, individual and aggregate limits, evaluated model range and intervention conditions. The service's approved mandate can differ from the employee's personal limit; neither overrides a policy prohibition. Ambiguous classification or non-standard terms step up to the designated owner.

At commitment, the service rechecks material state, reserves exposure and creates the purchase order through an idempotent interface. If cancellation is assumed before dispatch, verify that the contract and interface actually support it. The record links the evidence, model, policy result, mandate and transaction identifier. Monitor delivery, price variance, disputes, cancellations, supplier incidents and concentration, including cases that remain unresolved.

Now suppose a supplier-status change suspends this envelope. New automatic commitments stop. A separately admissible manual process can assess individual purchases; if the change prohibits purchases altogether, they remain blocked. The renewal task identifies the changed condition and the evidence required to restore or narrow the mandate. Where renewal and case review draw on the same constrained staff, backlog pressure is not a sufficient reason to defer renewal indefinitely.

The renewal owner checks the supplier evidence, obtains any required policy approval, retests the affected scope and records the renewed grant. Future eligible requests may then proceed. Existing queued requests remain under their assigned review process; changing that routing would require a separate authorised design and falls outside the renewal model's baseline. Already issued orders require monitoring or remediation independently of renewal.

Procurement owns the operating outcome, finance the budget controls, legal the interpretation of relevant terms and duties, risk the exposure criteria, and technology the runtime. A named renewal owner coordinates mandate restoration. Internal audit independently assesses control operation. Shorter lead time and lower manual effort count as benefits only alongside acceptable error, fraud, supplier, concentration and recourse outcomes.

F Measures, accountabilities and failure tests

Report by decision class and value stream. An unweighted total that mixes routine refunds with acquisitions has little management meaning. The adoption model's consequence-weighted delegation share is a research measure with specified scope, weights and time windows, not a universal executive score or a target to maximise. The renewal model's eligible share is different: eligibility is not realised delegation when authority is suspended.

Measure	Definition and use	Guardrail
End-to-end and decision-ready lead time	Request to action; separately, complete evidence to action	Keep both clocks; report tails and unfinished cases.
Straight-through decision rate	Share of a fixed eligible arrival cohort executed without per-instance human approval by a stated follow-up horizon	Retain unfinished cases in the denominator; report follow-up, mandate availability and candidate-wide rate.
Authority availability	Usable mandate time divided by observation time; separately, share of candidate arrivals encountering usable authority	Separate time and arrival weighting; report partial scope; do not infer decision quality.
Renewal burden and waiting	Active renewal effort, waiting before renewal, suspension-to-restoration time and effort allocation while suspended	Calendar delay is not active effort; identify shared reviewers.
Backlog and exception age	Cases awaiting or receiving review; age by consequence and route	Distinguish growth, service delay and unresolved renewal work.
Escalation yield and missed exceptions	Share of escalations materially changed by review; audited errors in apparently eligible cases	High yield alone can conceal under-escalation.
Decision outcomes and recourse	Useful completion, harmful approval, valid rejection, appeal and remedy by class	Prespecify windows and acceptable margins; track unresolved outcomes.
Control coverage and exposure	Current tested rules, active mandates and cumulative commitments	Limits are ceilings, not targets; stale logic is not coverage.
Total operating cost	Setup, case work, renewal, monitoring, appeals and remediation	Compare like populations and equal available resource budgets.

The board and CEO approve material delegation principles and demand evidence on exposure, recourse and duties retained by people. The COO owns cross-functional decision flow and reviewer capacity, including renewal. The CIO and architects provide identity, versioned policy and authority services, controlled interfaces, observability and records. Legal interprets applicable sources and duties; risk sets exposure, sampling, incident and acceptance criteria. These responsibilities must fit the organisation's actual mandates rather than be inferred from job titles.

The business outcome owner is accountable for performance and resources; the policy owner approves rule meaning; the renewal owner is responsible for restoring, narrowing or retiring suspended authority. AI governance connects models to the decisions they influence. Frontline experts supply boundary cases. Internal audit tests independently rather than owning the design it must assess. Several functions contribute, but one named outcome owner prevents shared work from becoming unowned work.

Failure tests should include invisible invalidation, stale policy, concurrent cap breaches, shared-source suspension of many envelopes, misleading confidence, missed exceptions, overloaded reviewers and inaccessible appeals. Repeated renewals may indicate unstable source policy or an uneconomic eligible scope. A portable ledger may preserve inaccurate history. A well-staffed exception queue may still leave renewal starved. These observations should narrow the design or challenge its application rather than be explained away as insufficient organisational maturity.

G Capacity accounting and protected renewal

This appendix makes the resource interpretation of equation (1) explicit. It restates the accounting and constant-allocation results of *The Authorisation Bottleneck* [3, Appendices B–C]; it introduces no new capacity theorem. The result concerns authorisation capacity for one decision class, not the capacity of downstream execution or the quality of the resulting decisions.

Let v be the stationary fraction of time the mandate is valid, r the fraction of total available reviewer time spent renewing, and c the fraction actually spent on cases. Balance of invalidations and renewals gives $\delta v = vr$. Renewal uses only invalid time, so $r \leq 1 - v$ and hence $0 \leq v \leq v/(v + \delta)$. Independent Poisson arrivals encounter the time-average valid fraction. Manual arrivals therefore have long-run rate $\lambda(1 - xv)$, giving $\mu c = \lambda(1 - xv)$. The effort budget $c + r \leq 1$ implies

$$\lambda \leq F(v) := \frac{\mu(1 - \delta v/v)}{1 - xv}, \quad 0 \leq v \leq \frac{v}{v + \delta}. \quad (\text{G.1})$$

The denominator is positive because $\delta > 0$. Moreover,

$$F'(v) = \frac{\mu(x - \delta/v)}{(1 - xv)^2}.$$

The largest bound is therefore attained at $v = 0$ when $vx \leq \delta$, or at $v = v/(v + \delta)$ when $vx > \delta$, producing equation (1). This is an upper-bound argument. Achievability for every strict load below the frontier requires the companion paper's separate stability proof: permanent manual review supplies the first endpoint and immediate renewal while suspended supplies the second. No conclusion here relies on operation at equality.

For a constant share $h \in (0, 1]$ reserved for renewal whenever suspended, with remaining effort available for cases and no reassignment of idle case effort, the exact stability condition is

$$\lambda < C(h) := \frac{\mu[\delta + h(\nu - \delta)]}{\delta + h\nu(1 - x)}. \quad (\text{G.2})$$

In the expansion regime $\nu x > \delta$, at a load $\mu < \lambda < \lambda_{\text{cap}}$, this is equivalent to

$$h > h_{\min} := \frac{\delta(\lambda - \mu)}{\mu(\nu - \delta) - \lambda\nu(1 - x)}. \quad (\text{G.3})$$

The denominator is positive and $0 < h_{\min} < 1$ in this range. This is a strict threshold for the specified allocation class, not a universal minimum for all scheduling rules. The long-run fraction of total available reviewer time spent renewing is $h\delta/(\delta + \nu h)$, which differs from the share h reserved while suspended. At lower load, renewing only when idle can preserve useful case service; a blanket renewal-first recommendation does not follow.

Estimates must use active effort for μ and ν , observed valid exposure for δ , and a fixed eligibility definition for x . Calendar renewal delay includes waiting and cannot stand in for renewal effort. Retain unfinished exposure when estimating rates. Parameter uncertainty, monitoring work and service targets require additional slack; several interacting mandates or heterogeneous reviewers require a different model. In particular, the capacity calculation does not establish that the five-plane architecture is necessary or sufficient to realise these allocations.

References

- [1] Susan C. Athey, Kevin A. Bryan, and Joshua S. Gans. The Allocation of Decision Authority to Human and Artificial Intelligence. *AEA Papers and Proceedings*, 110:80–84, 2020. doi:10.1257/pandp.20201034.
- [2] Decision Superintelligence Labs. *Enterprise AI Adoption as an Organisational State Transition*. Research paper, August 2026. URL <https://www.dsilabs.ai/research/enterprise-ai-adoption-state-transition>. Dated 1 August 2026.
- [3] Decision Superintelligence Labs. *The Authorisation Bottleneck*. Research paper, 10 August 2026. URL <https://www.dsilabs.ai/research/the-authorisation-bottleneck>.
- [4] European Union. Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024. *Official Journal of the European Union*, L, 2024/1689, July 2024. URL <https://eur-lex.europa.eu/eli/reg/2024/1689/oj>. Article 14: Human oversight.
- [5] Kate Goddard, Abdul Roudsari, and Jeremy C. Wyatt. Automation bias: a systematic review of frequency, effect mediators, and mitigators. *Journal of the American Medical Informatics Association*, 19(1):121–127, 2012. doi:10.1136/amiainl-2011-000089.
- [6] Vincent C. Hu, David Ferraiolo, Rick Kuhn, Adam Schnitzer, Kenneth Sandlin, Robert Miller, and Karen Scarfone. *Guide to Attribute Based Access Control (ABAC) Definition and Considerations*. Special Publication 800-162, National Institute of Standards and Technology, January 2014. doi:10.6028/NIST.SP.800-162. Includes updates as of 2 August 2019.
- [7] Ken Huang, Kyriakos Rock Lambros, Jerry Huang, Yasir Mehmood, Hammad Atta, Joshua Beck, Vineeth Sai Narajala, Muhammad Zeeshan Baig, Muhammad Aziz Ul Haq, Nadeem Shahzad, and Bhavya Gupta. AAGATE: A NIST AI RMF-Aligned Governance Platform for Agentic AI. arXiv:2510.25863v2, November 2025. URL <https://arxiv.org/abs/2510.25863v2>. Preprint; version consulted revised 3 November 2025.
- [8] National Institute of Standards and Technology. *Artificial Intelligence Risk Management Framework (AI RMF 1.0)*. NIST AI 100-1, January 2023. doi:10.6028/NIST.AI.100-1.
- [9] Object Management Group. *Decision Model and Notation*, version 1.5, August 2024. URL <https://www.omg.org/spec/DMN/1.5>. Formal specification.
- [10] Open Policy Agent. *Open Policy Agent Documentation*. URL <https://www.openpolicyagent.org/docs>. Accessed 22 September 2026.
- [11] Royal Commission into the Robodebt Scheme. *Report*. Commonwealth of Australia, July 2023. URL <https://robodebt.royalcommission.gov.au/publications/report>. Chapter 17, recommendations 17.1 and 17.2; corrected edition, 11 July 2023.
- [12] Linda J. Skitka, Kathleen L. Mosier, and Mark Burdick. Does automation bias decision-making? *International Journal of Human-Computer Studies*, 51(5):991–1006, 1999. doi:10.1006/ijhc.1999.0252.