

# Designing Globally Scalable Distributed Review and Certification Systems

Rajendar Reddy Sama

**Abstract:** Certification backlogs in digital software marketplaces are produced by architecture, not by insufficient reviewer headcount. Position-ordered queues without urgency differentiation, manual reviewer assignment without skill-match logic, and static policy enforcement without feedback mechanisms cannot maintain SLA compliance as submission volume and product diversity scale. This paper presents a six-component distributed certification architecture whose controlling design element is a queue intelligence model: SLA-weighted, continuously updated priority scoring combined with automated reviewer allocation. Fault-tolerance patterns specific to certification failure modes - idempotent event handling, dead-letter queue management, circuit-isolated service dependencies - and a governance framework anchored by decision-sequence audit trails complete the specification. A staged five-step implementation pathway allows deployment without disrupting active review operations.

**Keywords:** distributed certification; review workflow; queue intelligence; SLA prioritization; event-driven microservices; multi-tenant platform; audit trail; software marketplace

## 1. INTRODUCTION

Certification backlogs in digital marketplaces are architectural failures wearing the appearance of staffing shortfalls. When the queue processes submissions in arrival order without differentiating urgency, reviewer competency, or business impact, no amount of additional reviewer capacity fully resolves the underlying problem - it delays it. Developer launch schedules are the pressure surface. A developer coordinating simultaneous platform availability, localization deployment, and a marketing campaign treats certification as a dependency that must hit its window. When it does not, the revenue loss and credibility damage trace directly back to the platform, not to the developer's planning [3].

Scaling amplifies every structural weakness in a flat-queue certification model. Multi-store environments accumulate policy divergence: reviewer teams in different regions apply the same policy differently, and without a shared decision record, no one knows it is happening until appeals surface the pattern. Category proliferation produces reviewer specialization gaps - a reviewer competent in one product type may accept assignments in adjacent categories without the depth required, and the assignment mechanism that allows this is operating without a skill model [1]. Geographic distribution creates handoff gaps during which submission queues grow but review throughput does not. Each of these problems is independent; none is resolved by adding compute capacity.

Queue intelligence - structured SLA-weighted priority scoring and automated reviewer allocation - is the architectural variable this paper identifies as decisive. Infrastructure scale is a genuine requirement. The argument is that scale without queue intelligence is insufficient, not that scale is unnecessary. This paper draws on the author's applied experience architecting DARSy, a globally distributed review and certification platform implemented across Wipro's enterprise marketplace operations in partnership with Microsoft, as a practical reference for the patterns described.

The remainder of this paper is structured as follows. Section II defines the requirements of global certification systems. Section III presents the six-component architecture. Section IV develops the queue intelligence model. Section V addresses scalability and reliability patterns. Section VI describes governance and continuous improvement. Section VII discusses limitations. Section VIII concludes.

*Independent Researcher, USA*

## 2. REQUIREMENTS OF GLOBAL CERTIFICATION SYSTEMS

Multi-tenancy, SLA differentiation, geographic continuity, audit completeness, and onboarding velocity - these five requirements do not coexist easily. They pull the architecture in competing directions, and the design choices made to satisfy one frequently constrain how another can be satisfied.

Multi-tenant policy configuration requires that stores and business units operating on shared infrastructure enforce their own intake criteria, review checklists, and approval policies without coupling those policies at the service level. Waseem et al.'s systematic mapping of microservices in DevOps establishes that tenant isolation at the service layer is the pattern that enables policy agility at scale [4]. The alternative, dedicated infrastructure per tenant, produces operational complexity that compounds with every new store onboarded.

SLA differentiation goes beyond committing to average turnaround times. Li et al.'s examination of microservices quality attributes identifies responsiveness as a first-class concern because latency in platform services produces downstream effects that cascade beyond the immediate interaction [8]. A flat queue with a published average turnaround time satisfies neither the deadline-critical developer nor the platform operator who needs to demonstrate SLA performance by tier.

Audit completeness is where the operational requirements of certification diverge most sharply from generic workflow management. Zhou et al.'s industry survey confirms that distributed platforms routinely lose the decision path between state transitions when audit is implemented as status updates rather than event sequences [9]. For certification, this loss is not merely an observability gap - it is a regulatory and appeal liability.

Table I summarizes the five core requirements, their principal design implications, and the architectural responses developed in subsequent sections.

Table 1. Global Certification System Requirements

| Requirement           | Design Implication                                     | Architecture Response                                                   |
|-----------------------|--------------------------------------------------------|-------------------------------------------------------------------------|
| Multi-tenancy         | Policy isolation per store/BU without duplication      | Tenant-scoped config in Submission Service; shared canonical schema     |
| SLA differentiation   | Tiered commitments enforced, not published             | Queue Engine: SLA-weighted priority scoring, continuous update          |
| Geographic continuity | Review throughput without fragmentation across regions | Geo-routing with unified event bus; shared audit state                  |
| Audit completeness    | Decision-sequence survivable across component failure  | Audit Service: immutable event log, actor+timestamp on every transition |
| Onboarding velocity   | New category or store live without engineering months  | Category schema as typed extension; policy config without redeploy      |

## ARCHITECTURE PATTERN FOR DISTRIBUTED CERTIFICATION

The six-component architecture does not emerge from generic microservices decomposition applied to certification. It emerges from the failure modes the requirements produce. Figure 1 illustrates the reference architecture, showing the six components, primary data flows, event publication paths, and the score feedback loop from the Reviewer Workbench to the Queue Engine.

API gateways serve the developer entry point: authentication, rate limiting, transport-level payload validation, and routing to the submission service. Concentrating these functions at a single entry point is an operational consistency decision as much as a security one. Abgaz et al.'s systematic review identifies boundary inconsistency as one of the most common sources of architectural drift in microservices migrations [5].

The submission service produces a canonical submission. The canonical schema - uniform representation across all stores and product categories - makes multi-tenant operation tractable. When the submission service isolates store-specific intake logic and emits a standardized event, every downstream service operates against a single, known data contract [7]. Category schema extensions are typed objects attached to the canonical core.

The queue engine is where the architecture's argument is implemented. Its function is to make priority decisions that are fair, transparent, and continuously current. The audit and notification services consume events from other components without modifying state - a separation that ensures audit completeness and developer communication are not operationally coupled to the components generating the events they record.

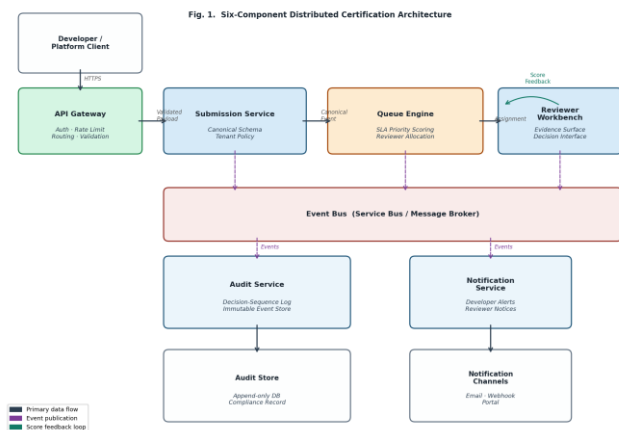


Fig. 1. Six-Component Distributed Certification Architecture

### 2.1. A. Submission Service Design

The canonical submission schema carries eight required fields: submission identifier, tenant identifier, product category, developer account reference, SLA tier, submission timestamp, artifact hash for idempotency, and applicable policy version. Category-specific metadata attaches as a typed extension, preserving schema stability for existing submissions when new categories are introduced.

### 2.2. B. Queue Engine as Architectural Center

Priority decisions are the queue engine's output. Every scheduling cycle, the engine updates priority scores across the pending submission population and services the assignment interface. The accuracy of those decisions - and the fairness of their aggregate effect across all developers and store tiers - is what makes queue intelligence substantively different from load-balancing mechanisms it superficially resembles [12].

## 3. QUEUE INTELLIGENCE AND SLA PRIORITIZATION

Priority score composition combines four factors. SLA deadline proximity carries the highest weight, increasing non-linearly as the remaining window narrows - a submission within 12 hours of its SLA boundary enters the critical escalation band regardless of other attributes. Submission age provides a floor weight: a submission that has queued for 96 hours triggers escalation regardless of other priority attributes, preventing indefinite deferral of low-urgency items. Store criticality applies a tier-based multiplier. Category complexity applies an inverse weight based on historical review duration distributions - a high-complexity category requiring 2.4 times the standard review duration must be weighted accordingly in SLA commitment calibration [13].

Score updates run continuously. A submission that was low-priority at intake may escalate to critical priority as its deadline approaches, without human intervention. This automatic update property is the mechanism that prevents SLA breaches from accumulating invisibly until they become visible failures.

Automated reviewer allocation solves the matching problem within the priority ordering the queue engine establishes. A reviewer's skill profile - the categories they are certified to review and the depth level of that certification - and their current availability state from workbench session status determine which submissions they are eligible to receive. Geographic routing preference minimizes cross-region handoffs that can introduce audit gap risk [10].

Transparency cannot be optional. When developers ask why their submission was deprioritized, the platform must provide an explainable score breakdown. Platforms that cannot explain their priority decisions generate appeals that consume reviewer capacity - exactly the capacity the priority model was intended to protect [8]. Unexplained priority signals are trust liabilities.

Table 2. SLA Priority Scoring Factors

| Scoring Factor         | Description                                       | Weight Basis              | Update Frequency         |
|------------------------|---------------------------------------------------|---------------------------|--------------------------|
| SLA Deadline Proximity | Inverse remaining window; linear escalation < 12h | Highest non-0.40)         | (≥Every scheduling cycle |
| Submission Age         | Elapsed time; escalation at 96h                   | Floor weight (0.15)       | Every scheduling cycle   |
| Store Criticality      | Tier-based multiplier (Tier 1–3)                  | Mid weight (0.25)         | On assignment change     |
| Category Complexity    | Inverse of historical review duration             | Calibration weight (0.20) | Rolling 30-day window    |

## 4. SCALABILITY AND RELIABILITY DESIGN PATTERNS

Scaling each component to its own demand profile is more efficient than uniform platform scaling. Intake services see volume spikes largely independent of review session activity; queue services experience reprioritization workload that scales with active submission count; reviewer workbench services scale with concurrent session activity. Shi et al.'s study confirms that per-component scaling policies calibrated to actual demand profiles outperform uniform scaling for heterogeneous-load platforms [10].

Four reliability patterns address failure modes specific to certification workflows. Idempotent event handling is the first - and the one most frequently underspecified. A submission event delivered twice through retry must not produce duplicate review assignments or duplicate audit records. Each event carries a compound idempotency key from submission identifier and event type; the receiving service checks prior processing before acting [9].

Dead-letter queues handle submissions exceeding the retry threshold without successful processing. Rather than blocking the main queue, unprocessable events route to an isolated channel for manual inspection. Circuit breakers around notification and audit services prevent cascade propagation during review volume spikes. When either service returns errors above threshold, the circuit opens and the calling component buffers outbound events. Xu et al. establish that failure isolation at component boundaries is a prerequisite for maintaining throughput SLAs when individual service tiers degrade [11].

Reconciliation jobs provide the safety net for state inconsistency: a reviewer whose session terminates unexpectedly may leave a submission assigned-but-inactive, which the job corrects without requiring manual intervention [14].

**Table 3. Reliability Patterns and Recovery Paths**

| Pattern             | Failure Scenario                                           | Mechanism                                                                | Recovery Path                                            |
|---------------------|------------------------------------------------------------|--------------------------------------------------------------------------|----------------------------------------------------------|
| Idempotent Handling | Event Duplicate delivery on retry                          | eventCompound idempotency (submissionId +check eventType) state mutation | Prior-keyprocessing                                      |
| Dead-Letter Queue   | Metadata violation; unresolved loop                        | Route to isolated retryDLQ after retries                                 | Manual inspection; schema correction; re-injection       |
| Circuit Breaker     | Notification/audit service degradation during volume spike | Open above threshold; outbound                                           | circuitHalf-open errorprobe; auto-recovery; event replay |
| Reconciliation Job  | Assigned-but-inactive submission (session drop)            | Scheduled for assignment state clear                                     | scanRe-queue stale submission; assignment record         |

## 5. GOVERNANCE, ANALYTICS, AND CONTINUOUS IMPROVEMENT

The certification platform generates operational data that can improve its own performance - but only if governance mechanisms are designed to act on it. Four dashboard metrics warrant specific attention because they inform qualitatively different corrective actions.

Review volume by category identifies demand trend shifts that precede staffing and SLA calibration decisions. Average turnaround by store disaggregates platform-level SLA performance into per-tenant performance, revealing whether latency is systemic or concentrated. SLA compliance rate measured weekly is the leading indicator for capacity planning - declining compliance trends precede breach events by a measurable interval. Policy confusion rate - the proportion of rejections successfully appealed - identifies policies generating inconsistent reviewer application [16].

Feedback loops from developers and reviewers produce inputs dashboard metrics cannot. Taibi et al. note that reviewer-identified policy ambiguity is a systemic quality problem when left unaddressed - inconsistency compounds as reviewer populations grow [15]. Audit trail design determines the quality of every governance activity that depends on it. The decision-sequence model - every state transition with actor identity, timestamp, and evidence reference - compresses investigation time and provides the documentary record enterprise marketplace environments increasingly require [18].

**Table 4. Key Operational Governance Metrics**

| Metric                   | Measurement Approach                                    | Governance Use                                             |
|--------------------------|---------------------------------------------------------|------------------------------------------------------------|
| Review Volume Category   | bySubmission event count per category per week          | Staffing and SLA calibration planning                      |
| Average Turnaround Store | Time from submission event to decision event per tenant | Identify systemic vs. concentrated latency                 |
| SLA Compliance (weekly)  | RateDecisions committed window / totalcapacity          | Leading indicator for adjustment                           |
| Policy Confusion Rate    | Appeals upheld / totalrejections per version            | Direct policy revision effort to high-ambiguity provisions |

## 6. DISCUSSION AND LIMITATIONS

Priority weight calibration is the sharpest operational risk in queue intelligence deployment. An algorithm that prioritizes by SLA tier will, if poorly configured, systematically advance high-tier stores at the expense of developers who cannot afford high-tier enrollment. Maximum deferral thresholds - guaranteeing no submission waits beyond a defined boundary regardless of priority score - are not optional. They are the fairness floor on which queue intelligence legitimacy rests [19].

Metadata quality is the dependency most consistently underestimated before deployment. Priority scoring depends on accurate SLA tier assignment at intake; reviewer allocation depends on complete skill profile records; governance analytics depend on consistent category classification across tenants. Schema validation at the submission service boundary must be enforced before queue intelligence is activated, not alongside it. Activating queue intelligence on a platform with inconsistent metadata produces a model that is confidently wrong rather than visibly uncertain.

The Azure-native service mapping in this paper is both a practical strength and a portability boundary. Rahman et al.'s systematic review of API management confirms that gateway abstraction patterns generalize across platforms, but configuration semantics and operational characteristics differ enough to require re-specification, not just substitution [20].

Three research directions remain open: fairness measurement in priority-weighted review queues; intelligent reviewer assignment incorporating expertise development trajectories; and adaptive policy maintenance driven by decision outcome analysis.

## 7. CONCLUSION

The central claim of this paper is not that certification platforms need better infrastructure. It is that they need queue intelligence. Infrastructure scale without priority management produces throughput; queue intelligence produces throughput with SLA compliance and developer equity. The six-component architecture presented here specifies the queue engine as the controlling design element, and every other component is specified in terms of what it must provide for queue intelligence to operate correctly.

Staged adoption makes the architecture accessible to organizations whose current certification infrastructure is not designed for replacement. Submission model standardization and canonical schema adoption provide immediate operational discipline. Event-driven communication adoption enables reliable state transitions. Queue intelligence activation creates the

platform's operational differentiation. Governance analytics and continuous policy refinement complete the maturity arc.

Certification platforms serve as trust infrastructure for the software ecosystems they govern. A flat queue treats a deadline-critical developer submission identically to a routine patch update. Queue intelligence restores the differentiation the business context demands.

## 8. REFERENCES

- [1] E. A. Brewer, "CAP twelve years later: How the 'rules' have changed," *IEEE Computer*, vol. 45, no. 2, pp. 23–29, Feb. 2012. <https://ieeexplore.ieee.org/document/6133253>
- [2] A. Balalaie, A. Heydarnoori, and P. Jamshidi, "Microservices architecture enables DevOps: Migration to a cloud-native architecture," *IEEE Software*, vol. 33, no. 3, pp. 42–52, May–Jun. 2016. <https://ieeexplore.ieee.org/document/7436659>
- [3] P. Jamshidi, C. Pahl, N. C. Mendonça, J. Lewis, and S. Tilkov, "Microservices: The journey so far and challenges ahead," *IEEE Software*, vol. 35, no. 3, pp. 24–35, May–Jun. 2018. <https://ieeexplore.ieee.org/document/8354433>
- [4] M. Waseem, P. Liang, and M. Shahin, "A systematic mapping study on microservices architecture in DevOps," *J. Syst. Softw.*, vol. 170, art. no. 110798, Dec. 2020. <https://www.sciencedirect.com/science/article/abs/pii/S0164121220302053>
- [5] Y. Abgaz et al., "Decomposition of monolith applications into microservices architectures: A systematic review," *IEEE Trans. Softw. Eng.*, vol. 49, no. 8, pp. 4213–4242, Aug. 2023. <https://ieeexplore.ieee.org/document/10160171>
- [6] M. Usman, S. Ferlin, A. Brunstrom, and J. Taheri, "A survey on observability of distributed edge & container-based microservices," *IEEE Access*, vol. 10, pp. 86904–86919, 2022. <https://ieeexplore.ieee.org/document/9837035>
- [7] A. Rahmatulloh et al., "Event-driven architecture to improve performance and scalability in microservices-based systems," in *Proc. IEEE ICADEIS*, 2022. <https://ieeexplore.ieee.org/document/10037390>
- [8] S. Li et al., "Understanding and addressing quality attributes of microservices architecture: A systematic literature review," *Inf. Softw. Technol.*, vol. 131, art. no. 106449, Mar. 2021. <https://www.sciencedirect.com/science/article/abs/pii/S0950584920301993>
- [9] X. Zhou et al., "A fault analysis and debugging of microservice systems: Industry survey, benchmark system, and empirical study," *IEEE Trans. Softw. Eng.*, vol. 49, no. 4, pp. 1887–1908, Apr. 2023. <https://ieeexplore.ieee.org/document/8580420>
- [10] T. Shi, H. Ma, G. Chen, and S. Hartmann, "Auto-scaling containerized applications in geo-distributed clouds," *IEEE Trans. Serv. Comput.*, vol. 16, no. 6, pp. 4261–4274, Nov.–Dec. 2023. <https://ieeexplore.ieee.org/document/10255283>
- [11] J. Xu et al., "CoScal: Multifaceted scaling of microservices with reinforcement learning," *IEEE Trans. Netw. Serv. Manag.*, vol. 19, no. 4, pp. 3995–4009, Dec. 2022. <https://ieeexplore.ieee.org/document/9904920>
- [12] Angelos-Christos Maroudis, Theodoros Theodoropoulos, John Violos, Aris Leivadeas, and Konstantinos Tserpes. 2024. Leveraging Graph Neural Networks for SLA Violation Prediction in Cloud Computing. *IEEE Trans. on Netw. and Serv. Manag.* 21, 1 (Feb. 2024), 605–620. <https://doi.org/10.1109/TNSM.2023.3292392>
- [13] D. -D. Vu, M. -N. Tran and Y. Kim, "Predictive Hybrid Autoscaling for Containerized Applications," in *IEEE Access*, vol. 10, pp. 109768–109778, 2022, <https://ieeexplore.ieee.org/document/9919851>
- [14] J. Bogner, J. Fritzsch, S. Wagner and A. Zimmermann, "Microservices in Industry: Insights into Technologies, Characteristics, and Software Quality," 2019 IEEE International Conference on Software Architecture Companion (ICSA-C), Hamburg, Germany, 2019, pp. 187–195, <https://ieeexplore.ieee.org/document/8712375>
- [15] D. Taibi, V. Lenarduzzi, and C. Pahl, "Microservices anti-patterns: A taxonomy," in *Microservices Science and Engineering*, Springer, 2020, pp. 111–128. <https://arxiv.org/pdf/1908.04101>
- [16] Yingying Wang, Sarah Bornais, and Julia Rubin. 2024. Microservice Decomposition Techniques: An Independent Tool Comparison. In *Proceedings of the 39th IEEE/ACM International Conference on Automated Software Engineering (ASE '24)*. Association for Computing Machinery, New York, NY, USA, 1295–1307. <https://dl.acm.org/doi/10.1145/3691620.3695504>
- [17] Guadalupe Ortiz et al., "A microservice architecture for real-time IoT data processing: A reusable Web of Things approach for smart ports," *ScienceDirect*, vol. 81, 2022. <https://www.sciencedirect.com/science/article/pii/S0920548921000994#article>
- [18] R. Laigner, Y. Zhou, M. A. V. Salles, Y. Liu, and M. Kalinowski, "Data management in microservices: State of the practice, challenges, and research directions," *Proc. VLDB Endow.*, vol. 14, no. 13, pp. 3348–3361, 2021. <https://vldb.org/pvldb/vol14/p3348-laigner.pdf>
- [19] P. Francisco et al., "Microservices testing: A systematic literature review," *ScienceDirect*, vol. 188, 2025. <https://www.sciencedirect.com/science/article/abs/pii/S0950584925002095>
- [20] Nneka Ochuba et al., "Systematic Review of API Gateway Patterns for Scalable and Secure Application Architecture," *Journal of Frontiers in Multidisciplinary Research* 02(01):94-100, 2021. [https://www.researchgate.net/publication/393361386\\_Systematic\\_Review\\_of\\_API\\_Gateway\\_Patterns\\_for\\_Scalable\\_and\\_Secure\\_Application\\_Architecture](https://www.researchgate.net/publication/393361386_Systematic_Review_of_API_Gateway_Patterns_for_Scalable_and_Secure_Application_Architecture)