

A Real-World Application Based Evaluation of Serverless Architectures

Akhil Rana¹, Gayatri Barbade², A.L. Rane³

^{1,2,3} K.K. Wagh Institute of Engineering Education and Research, Nashik

Abstract—Serverless computing promises automatic scaling and cost reduction, yet the real-world performance of production serverless systems frequently diverges from vendor claims. This paper addresses that divergence by documenting the complete lifecycle of a CRUD-based inventory management system deployed on AWS — from architecture design through instrumented load testing and cost accounting. Four distinct test scenarios were executed under controlled conditions in the ap-south-1 region. Key measurements: a p50 response time of 161 ms under warm conditions, sustained below 250 ms even when concurrency increased twentyfold to 200 simultaneous users; a cold-start penalty averaging 1,450 ms on first invocation versus 180 ms on warm invocations; and a monthly cost projection 62% below a comparably sized EC2 instance. A significant finding concerns developer experience: a misconfigured IAM policy produced a silent failure requiring correlation of four separate log sources over three hours to diagnose — a concrete illustration of what this paper terms the operational paradox of serverless simplicity. Collectively, these results position serverless as the correct choice for traffic-variable, latency-tolerant workloads and the wrong choice where uniform response times are contractually required.

Index Terms—serverless computing, AWS Lambda, cold start latency, Function-as-a-Service, CRUD evaluation, cost analysis, cloud performance, DynamoDB.

I. INTRODUCTION

A. Motivation

Cloud infrastructure has undergone a sequence of shifts in who carries operational responsibility. Early cloud adoption kept most of that responsibility with the application team: virtual machines had to be provisioned, operating systems patched, and scaling policies written by hand. A subsequent tier of managed services offloaded some of that work, but left teams responsible for deciding how many instances to run.

The Function-as-a-Service model, introduced commercially by AWS Lambda in 2014, removes that question entirely — the developer writes a function body, and everything surrounding it is managed by the platform [1],[2].

This shift attracts engineering teams for two connected reasons. First, billing granularity becomes extreme: compute charges accumulate only while a function is actively executing, so quiet periods cost nothing. Second, scaling becomes reactive: the platform responds to incoming requests by spawning additional containers as needed, with no operator involvement [3]. For services with unpredictable or seasonal demand, these properties can convert a fixed infrastructure barrier into a variable cost that shrinks when usage shrinks.

What vendors communicate less clearly are the conditions under which these properties break down. When a function sits idle long enough, its container is reclaimed, and the next request must wait for a new container and runtime to be initialized. This pause — the cold start — sits between the user and the function body. Amazon has itself quantified a related effect in its own products: a 100 ms latency increase was associated with a 1% sales reduction [5]. For a transactional API where users expect responses within a few hundred milliseconds, a cold start exceeding one second is both perceptible and commercially consequential.

A second failure mode concerns fault isolation. Serverless architectures scatter failure evidence: each function invocation writes to its own log stream, API Gateway maintains a separate access log, and downstream service errors appear in still more

systems. Finding the root cause of a silent failure requires navigating all of these simultaneously [10].

B. Research Question

This study documents both phenomena — cold-start latency and operational complexity — quantitatively and qualitatively through a purpose-built CRUD application tested on AWS. The central research question is: across the four dimensions of scalability, cost, latency, and operational manageability, what does a production-grade serverless deployment on AWS actually deliver?

C. Contributions

- 1) End-to-end empirical evaluation of a production-representative serverless CRUD application on AWS, covering all four operation types under four distinct load scenarios.
- 2) Quantified cold-start analysis: mean 1,450 ms vs. 180 ms warm, with root-cause attribution to developer-controlled initialization patterns rather than platform behavior.
- 3) Full-stack cost breakdown across all AWS services — Lambda, DynamoDB, API Gateway, CloudFront, S3, Cognito — projecting 62% cost savings over provisioned EC2.
- 4) Structured adoption decision framework synthesizing experimental data and published benchmarks across five decision dimensions.

II. RELATED WORK

A. Cold Start Latency

The cold start is the defining performance liability of FaaS. Its duration is shaped by three factors under developer control: the runtime environment, deployment package size, and the volume of initialization code placed outside the handler function. Chowdhury et al. [4] decompose cold start into identifiable sub-phases and demonstrate that developer-authored initialization often accounts for the majority of delay. Laghari et al. [13] provided cross-provider measurements, testing AWS Lambda, Azure Functions, and Google Cloud Functions under identical Python workloads. Their data reveal that provider selection is itself a design decision when cold-start frequency matters. Table I situates published figures alongside this study's measurements.

TABLE I Cold Start Latency by Runtime —
Published Benchmarks vs. This Study

Runtime	Cold P50 (ms)	Cold P99 (ms)	Warm (ms)	C/W Ratio	Source
Python 3.9	250–400	600–900	5–15	~30×	Laghari et al. [13]
Node.js 18	150–300	400–700	5–12	~25×	Tanasa [15]
Java 11 (JVM)	900–1,800	2,000–2,600	8–20	~120×	IJARPR [16]
Go 1.21	120–200	300–450	4–10	~20×	Laghari et al. [13]
Rust (arm64)	16–80	150–200	3–8	~12×	Scanner.dev [17]
This Study (Python)	1,450 *	1,780	180	8.1×	Experiment

* Elevated due to DynamoDB client + Cognito JWT verifier initialized at module scope.

B. Cross-Provider Comparison

The SeBS benchmark suite [18] quantified network and I/O variability across AWS, Azure, and GCP. AWS exhibited the highest same-region ping latency (109 ms) compared to Azure (33 ms) and GCP (20 ms), contributing to overall response time independent of function initialization. Table II summarizes the cross-provider landscape.

TABLE II Cross-Provider Serverless Characteristics
(Published Data)

Metric	AWS Lambda	Azure Functions	GCP Cloud Fns	Source
Cold Start P50 (Python)	250–400 ms	800–1,200 ms	400–700 ms	[13]
Warm Latency P50	5–15 ms	10–25 ms	8–18 ms	[13]
Default Concurrency	1,000 exec.	200 exec.	1,000 exec.	[19]
Region Ping (same)	109 ms	33 ms	20 ms	SeBS [18]
Free Tier (req/mo)	1 M	1 M	2 M	Official Docs

C. Cost Modeling

Hassan et al. [12] synthesized seventeen independent cost studies and found that serverless billing beats provisioned alternatives decisively when peak traffic is at least two to three times average traffic, but approaches break-even when traffic is flat and sustained. Table III collects published cost-comparison findings.

TABLE III Serverless vs. Provisioned Infrastructure: Published Cost Findings

Study	Workload	Cost Saving	Condition	Notes
Hassan et al. [12]	General web API	50–70%	Variable traffic	SLR of 17 studies
IJISAE 2024 [20]	E-learning	56.8%	Simulated load	Also 71.8% faster
Hamza et al. [21]	Startup SaaS	~50%	EC2 migration	Qualitative
Besozzi et al. [22]	Per-unit price	~41%	Sustained load	Lambda pricier/unit
This Study	Inventory CRUD	62%	8-hour variable	AWS Cost Explorer

D. Developer Experience and Observability

De Oliveira et al. [10] surveyed practitioners and found configuration errors — particularly IAM misconfiguration and event routing mistakes — to be both the most common failure source and the most time-consuming to diagnose. Reza et al. [14] identify the permission surface as a distinct security concern: each function must carry only the permissions it strictly requires, and errors in either direction create security or reliability problems that surface only at runtime.

III. SYSTEM DESIGN AND METHODOLOGY

A. Application Architecture

The subject of the experiment is an inventory management system implementing the four standard CRUD operations on product records. The application was chosen because it combines authentication, stateful persistence, and network I/O in proportions representative of a real internal business tool. Four AWS services compose the backend: Lambda (one

function per operation, Python 3.9 runtime, 256 MB memory allocation); API Gateway (HTTP entry point, JWT enforcement); DynamoDB (on-demand capacity, single-table design); and Cognito (user pool, JWT issuance). A static frontend is distributed through CloudFront backed by S3. All resources were provisioned in the ap-south-1 (Mumbai) region.

B. Instrumentation

Three measurement systems ran simultaneously. CloudWatch collected function logs and invocation-level metrics. X-Ray captured distributed traces from API Gateway receipt through Lambda execution to DynamoDB acknowledgment, providing per-stage timing for each sampled request. Artillery.io, executed on a dedicated EC2 instance in the same region, generated load and recorded client-side p50 and p99 latency percentiles. Running the load generator in the same region eliminates internet path variability from measurements.

C. Test Scenarios

Scenario 1 — Warm Baseline: 10 requests per second for 15 minutes following a priming period establishing warm containers across all four function types. Scenario 2 — Cold Start: 20-minute idle, single request, record Init Duration from X-Ray; repeated 10 times. Scenario 3 — Scalability: ramp from 10 to 200 concurrent users over 20 minutes, hold peak for 10 minutes. Scenario 4 — Cost: 8-hour variable-traffic profile alternating 5 RPS and 100 RPS, extract AWS Cost Explorer billing breakdown.

D. Metrics

Response time was captured at p50 and p99 percentiles. Error rate was computed as the fraction of HTTP 5xx responses per scenario. Cold-start duration was read from the Init Duration field in X-Ray traces. Total cost was extracted from AWS Cost Explorer with service-level granularity, summing Lambda, DynamoDB, API Gateway, CloudFront, S3, and Cognito charges.

IV. EVALUATION RESULTS

A. Warm-State Baseline Performance

Under steady-state warm conditions (10 RPS, 15-minute window), all four operation types delivered

consistent and low-latency responses as shown in Table IV.

TABLE IV Baseline Performance — Warm State (10 RPS Steady)

Operation	p50 (ms)	p99 (ms)	RPS	Error Rate
GET (Read)	142	210	10.0	<0.01%
POST (Create)	168	248	10.0	<0.01%
PUT (Update)	175	261	10.0	<0.01%
DELETE	158	225	10.0	<0.01%
Overall Mean	161	236	10.0	<0.01%

Key Finding 1: All operations completed below 180 ms at the median. The p99 ceiling of 261 ms indicates low tail variance, consistent with DynamoDB single-digit millisecond access times adding predictably to function execution.

B. Scalability Under Concurrent Load

The scalability scenario ramped concurrency from 10 to 200 users — a twentyfold increase — and held at 200 for a further 10 minutes. Table V shows how response time and error rate evolved.

TABLE V Scalability Profile — Ramp 10 → 200 Concurrent Users

Concurrent Users	p50 (ms)	p99 (ms)	Effective RPS	Error Rate
10 (baseline)	161	236	10	<0.01%
25	172	254	24	<0.01%
50	185	278	49	0.02%
100	208	312	98	0.04%
150	229	348	146	0.07%
200 (peak)	247	389	195	0.09%

Key Finding 2: p50 at 200 concurrent users (247 ms) remained below 250 ms and the error rate peaked at 0.09% — well within the <0.1% threshold. Lambda's container pool expanded without any operator action. The p99 rise from 236 ms to 389 ms (65%) is notable for latency-SLA contexts.

C. Cold Start Measurements

TABLE VI Cold Start Trial Measurements (n = 10)

Trial	Duration (ms)	Trial	Duration (ms)
T1	1,410	T6	1,440
T2	1,480	T7	1,520
T3	1,390	T8	1,430
T4	1,510	T9	1,470
T5	1,460	T10	1,390
Mean	1,450	Std. Dev.	±42 ms

TABLE VII Cold Start Statistics vs. Published Python Benchmarks

Statistic	This Study	Typ. Python [16]	Typ. Node.js [15]	Typ. Java [16]
Mean cold (ms)	1,450	250–400	150–300	900–1,800
Warm mean (ms)	180	5–15	5–12	8–20
Cold/Warm ratio	8.1×	~30×	~25×	~120×
Std. deviation	±42 ms	±80 ms	±60 ms	±280 ms
P99 cold (ms)	1,780	600–900	400–700	2,000–2,600

Key Finding 3: The 1,450 ms mean cold start exceeds the typical Python range (250–400 ms). Code inspection identifies the cause: both the DynamoDB client and Cognito JWT verifier are instantiated at module scope, executing on every cold start. This is a developer design choice, not a platform characteristic.

D. Cost Analysis

Table VIII presents the cost breakdown from the 8-hour variable-traffic scenario, extrapolated to a monthly period and compared against the on-demand price of a t3.small EC2 instance in ap-south-1.

TABLE VIII Monthly Cost Projection — Service Breakdown

Service	Test Cost (\$)	% of Total	Monthly Est. (\$)	EC2 Equivalent (\$)	Saving
AWS Lambda	0.0031	51.7 %	0.93	—	—
Amazon DynamoDB	0.0011	18.3 %	0.33	—	—
API Gateway	0.0009	15.0 %	0.27	—	—
CloudFront / S3	0.0005	8.3 %	0.15	—	—
Cognito	0.0004	6.7 %	0.12	—	—
TOTAL	0.0060	100%	1.80	~4.75 *	62%

* *t3.small ap-south-1, on-demand, 100% utilisation, Q1 2025 pricing.*

Key Finding 4: A 62% monthly cost saving over provisioned infrastructure, consistent with the 50–70% range from Hassan et al. [12]. DynamoDB contributed 18.3% of total spend — invisible in a Lambda-only cost model.

V. DISCUSSION

A. Scalability: Where FaaS Delivers Unconditionally

A twentyfold traffic surge was absorbed with an error rate below 0.1% and p50 latency below 250 ms, with no configuration change or operational response required. For engineering teams accustomed to pre-scaling before anticipated load spikes, this removes an entire category of planned work. The residual p99 degradation lies within acceptable bounds for the inventory-management use case studied.

B. Cost Efficiency: Real but Workload-Dependent

The 62% cost saving arises because the test workload alternates between low and peak traffic; during quiet intervals the serverless bill approaches zero while an EC2 instance continues accumulating at its hourly rate. The Besozzi et al. [22] finding — that Lambda's per-unit compute price is 41% higher than equivalent EC2 capacity — clarifies why the advantage

disappears for flat, sustained workloads: at 100% utilisation, the higher per-unit price dominates and variable billing offers no benefit.

C. Cold Start: A Developer-Owned Problem

The measured cold-start time of 1,450 ms is 3.6 to 5.8 times higher than the typical published Python range (250–400 ms). The difference is entirely attributable to a coding pattern: placing the DynamoDB client and JWT verifier at module scope. This validates Chowdhury et al.'s [4] finding that developer-authored initialization code is often the dominant cold-start cost factor. Teams optimising for cold-start performance should audit their module-scope code before concluding the platform is the bottleneck.

D. The Operational Paradox

During development, a misconfigured IAM permission on the Lambda execution role caused DynamoDB write operations to fail silently: HTTP 200 OK was returned, but no record was persisted. Diagnosing the failure required correlating API Gateway access logs, Lambda execution logs, CloudWatch metrics, and an X-Ray trace — four independent sources. Total diagnostic time was approximately three hours. In a traditional single-process application, the same fault would have produced an exception in one log file, diagnosable in minutes. This corroborates de Oliveira et al. [10] and illustrates why observability infrastructure should be treated as a deployment prerequisite.

E. Limitations

1) Single region: all experiments were conducted in ap-south-1; cold-start and cost figures may differ across regions. 2) Single runtime: only Python 3.9 was tested; other runtimes would produce different cold-start profiles. 3) Workload representation: the CRUD inventory workload may not generalise to streaming, ML inference, or event-processing scenarios. 4) Traffic simulation: Artillery.io simulates concurrent HTTP clients; real user behaviour may differ in distribution.

VI. ADOPTION DECISION FRAMEWORK

Table IX integrates experimental findings and published benchmarks into a structured guide for teams evaluating serverless adoption.

TABLE IX Serverless Adoption Decision Framework

Dimension	Serverless Preferred	Serverless Avoid	Key Factor	Evidence
Traffic pattern	Variable / spiky	Flat and sustained	Billing structure	[12]; this study
Latency tolerance	>300 ms acceptable	Sub-100 ms required	Cold start risk	[4]; this study
Cost model	Pay-per-use wins	Break-even crossed	Traffic volume	[22]; this study
Team readiness	Observability in place	Debugging infrastructure absent	Fault isolation	[10]; this study
Security posture	IAM review defined	Per-fn perms unmanaged	Permission surface	[21]

VII. CONCLUSION

This paper reported an end-to-end empirical evaluation of AWS serverless architecture through the implementation and systematic instrumented testing of a production-representative CRUD application. Three principal contributions emerge. First, automatic scalability holds under real concurrent load: the system handled a twentyfold traffic surge with a sub-0.1% error rate and p50 latency below 250 ms, requiring no operator action. Second, cost efficiency is real but conditional: a 62% monthly saving over provisioned infrastructure was recorded for a variable-traffic workload, consistent with the broader literature, but this advantage depends on the traffic profile. Third, the cold-start penalty — $8.1\times$ the warm invocation time — is substantially a function of developer code organisation, not an immutable platform characteristic.

The operational paradox — effortless scaling achieved through the same architectural decomposition that makes failures difficult to diagnose — is the finding most relevant for adoption decisions. Teams that build observability infrastructure before their first production deployment will encounter this paradox on their own terms; teams that treat it as optional will encounter it through unexplained user-facing failures.

Three directions merit follow-on work: (1) a controlled measurement of how moving SDK initialization inside the handler body affects both cold-start and warm-start latency; (2) extension of the cost model to sustained high-volume traffic, specifically locating the break-even point for this application's request shape; and (3) a randomised comparison of time-to-diagnosis between equivalent serverless and container-based deployments using fault injection.

ACKNOWLEDGMENT

The authors thank Prof. A.L. Rane, Department of Computer Engineering, K.K. Wagh Institute of Engineering Education and Research, Nashik, for his guidance throughout this work. The authors also acknowledge the use of the AWS Free Tier and AWS Academy resources during the experimental phase of this study.

REFERENCES

- [1] Amazon Web Services, “Serverless Application Lens – AWS Well-Architected Framework,” AWS White Paper, 2023.
- [2] [2] M. Taibi and D. A. Tamburri, “A comparative study of microservice and serverless architectures,” *J. Syst. Softw.*, vol. 182, p. 111073, 2021.
- [3] [3] G. Casale et al., “The future of serverless computing: A research agenda,” in *Proc. 1st Int. Workshop Serverless Syst., Appl. Paradigms*, pp. 1–6, ACM, 2023.
- [4] [4] A. Chowdhury, S. Paul, M. Rahman, and M. J. Zohir, “Serverless computing: A study of cold start problem,” in *Proc. 2022 ICISSET*, pp. 426–431, IEEE, 2022.
- [5] [5] G. Linden, “Make data useful,” presented at Stanford University CS345, Nov. 2006. [Online]. Available: <https://glinden.blogspot.com>
- [6] [6] A. S. Gill and R. S. Sandhu, “A review of serverless computing,” *J. Cloud Comput.*, vol. 12, no. 1, p. 32, 2023.
- [7] [7] S. Alam et al., “Opportunities and challenges of serverless computing: A survey,” *Security Privacy*, vol. 5, no. 2, p. e202, 2022.
- [8] [8] S. R. Chowdhury and L. K. Shar, “Observability in the serverless era,” *ACM Trans. Softw. Eng. Methodol.*, vol. 32, no. 4, 2023.

- [9] [9] N. M. B. E. Hassan et al., “When is serverless a cheaper option?” *IEEE Access*, vol. 10, pp. 70914–70931, 2022.
- [10] [10] N. R. T. A. de Oliveira, A. C. Hora, and M. T. Valente, “An empirical study on the challenges of serverless platform users,” in *Proc. 44th ICSE: SEIP*, pp. 333–344, 2022.
- [11] [11] N. C. R. de S. Filho et al., “Performance and cost evaluation of serverless FaaS platforms,” *J. Cloud Comput.*, vol. 12, no. 1, p. 108, 2023.
- [12] [12] M. M. Hassan et al., “Performance and cost of serverless computing: A systematic literature review,” *Cluster Comput.*, vol. 25, no. 6, pp. 4053–4074, 2022.
- [13] [13] A. A. Laghari, J. He, and G. Gui, “A comparative analysis of cold start latency across cloud providers,” *IEEE Access*, vol. 10, pp. 67401–67412, 2022.
- [14] [14] S. M. P. Reza et al., “Security challenges in serverless computing,” in *Proc. 2022 IEEE DASC*, pp. 1–8, IEEE, 2022.
- [15] [15] A. Tanasa, “Size is (almost) all that matters for optimizing AWS Lambda cold starts,” *Medium*, Sept. 2023.
- [16] [16] “Serverless computing performance analysis for real-time applications,” *IJARPR*, vol. 2, no. 12, 2024.
- [17] [17] Scanner.dev, “Serverless speed: Rust vs. Go, Java, and Python in AWS Lambda,” Sept. 2025.
- [18] [18] P. Copik et al., “SeBS: A serverless benchmark suite for FaaS computing,” *ACM Middleware*, 2021.
- [19] [19] Amazon Web Services, “Lambda quotas,” *AWS Documentation*, 2025.
- [20] [20] P. Rao and A. Menon, “Benchmarking serverless efficiency for e-learning platforms,” *Int. J. Intell. Syst. Appl. Eng.*, vol. 12, no. 23s, 2024.
- [21] [21] M. Hamza, M. A. Akbar, and R. Capilla, “Understanding cost dynamics of serverless computing,” in *Proc. ICSOB 2023*, Springer LNBIP, vol. 500, 2024.
- [22] [22] V. Besozzi et al., “Demystifying serverless costs on public platforms,” *arXiv:2506.01283*, 2025.