



Validation That Travels: A CI/CD-Native Framework for Cloud-Ready Telecom DataOps Platforms

Sudhir Kumar Verma

Solution Architect, Ericsson

Abstract—This draft tackles a problem that has followed me across every telecom delivery: a mediation product has to be validated across lab, customer-like, cloud, and Kubernetes environments, and the usual failure is that deployment velocity quietly outruns the test evidence behind it. I propose making validation portable: declarative test inventories, configuration profiles per environment, synthetic CDR generators, and CI evidence reports, so the same suite and the same proof follows the product from an internal lab to a customer's Kubernetes cluster. I give the architecture, a domain risk matrix, an evaluation design with honest baselines, and an illustrative result set whose numbers are placeholders for logged measurements. The aim is a reliability-engineering method that ties software automation to evidence an outside reviewer could check, rather than to internal assertion.

Index Terms—CI/CD, Kubernetes, telecom DataOps, test automation, environment drift.

I. Introduction

Across every delivery pipeline I have owned, the same lesson repeats: reliability is designed in, not tested in. It is diagnosability, evidence quality, traceability, and the capacity to catch and understand a failure before it travels downstream properties of the architecture, not of the final regression run. This paper takes that stance into one specific problem: a mediation product has to be validated across lab, customer-like, cloud, and Kubernetes environments, and the usual failure is that deployment velocity quietly outruns the test evidence behind it.

Cloud-native mediation made an old problem sharper. When the same software runs in a lab, a customer staging cluster, and two cloud regions, configuration and environment drift becomes one of the most common and hardest to trace sources of failure. Portable tests are the obvious half of the answer; portable *evidence* is the half people forget.

As a solution architect and product owner for Ericsson Mediation, I built an automation framework and wired it into GitLab CI/CD across very different target environments. The lesson that stuck: the bug is rarely the code it is the gap between the environment you tested and the one you shipped to.

What I am proposing. I propose making validation portable: declarative test inventories, configuration profiles per environment, synthetic CDR generators, and CI evidence reports, so the same suite and the same proof follows the product from an internal lab to a customer's Kubernetes cluster.

To keep the claims testable, I organize everything around three questions an experiment could resolve:

RQ1. On the same workload, does the proposed method improve portability relative to an honest baseline (per-environment ad hoc suites, or the existing manual validation)?

- RQ2.** Does it achieve that without degrading detection quality or the quality of the evidence the method leaves behind?
- RQ3.** Which component is responsible for the improvement the full architecture, or one module that an ablation can isolate?

II. Background and Motivation

Cloud-native delivery solved real scaling problems for telecom data platforms and introduced a new dominant failure mode in their place: environment drift. When the same software runs in an internal lab, a customer-like staging cluster, and two or three cloud regions, the differences that matter are rarely in the code. They are in configuration, in platform behaviour, in the subtle ways one environment is not quite the other. A test suite that passes in the lab and a deployment that misbehaves in the field is, increasingly, the same bug wearing two costumes.

The conventional answer run more tests misses the point if those tests only ever run in one place. What is needed is validation that is portable across environments and, just as importantly, evidence that is portable with it, so that passing in one environment makes a checkable, environment-specific claim rather than a hopeful generalization. That is the gap between testing and trustworthy release, and it is the gap this framework is built to close.

I am careful to keep the public framing apart from my own evidence. Public sources establish the stakes; the contribution should stand on sanitized artifacts of my own test reports, CI logs, architecture notes, reviewer sign-off. The framework below is written so a synthetic or sanitized workload can demonstrate it without exposing anything proprietary.

III. Related Work and Context

I position this against established delivery and testing practice rather than a new field. The charging standards [1] fix the domain interfaces, OpenTelemetry [5] the observability vocabulary, and ISO/IEC/IEEE 29119 [4] the V&V language. My contribution is methodological: making validation and its evidence portable across environments.

Continuous delivery [3] and cloud-native operations [2] provide the practices this builds on, and the testing standards provide the V&V vocabulary. The specific contribution is portability of both tests and evidence across heterogeneous telecom environments, with environment fingerprinting so that drift defects are localizable. Where much CI/CD writing assumes a single target, telecom delivery is intrinsically multi-environment, and that is the condition this framework is shaped to.

Table 1. Domain risk and mitigation matrix.

Failure mode	Sev.	Consequence	Mitigation
Lab-to-customer drift	High	False confidence before a rollout	Portable, declarative environment profiles
Manual validation effort	Medium	Slow, inconsistent release cycles	CI-triggered regression packs
Non-functional gaps	High	Performance surprises under real load	Load and platform tests in the same framework
Slow triage	Medium	Failures take too long to localize	Auto-tagged, environment-fingerprinted failure reports



Figure 1. Portable validation framework. The artifact that moves between environments is not just the test but the evidence it produces.

IV. Proposed Method

Everything below runs on stock tooling Python, Linux, CI/CD pipelines, structured logs, a SQL-style store and follows three moves I make on every project:

1. I separate what is tested (a declarative inventory) from where it runs (a config profile), so a new environment is a new profile, not a new test suite.
2. I drive everything from synthetic workloads calibrated to non-sensitive statistics, which keeps the framework shareable and the customer's data out of it.
3. I fingerprint each environment into the evidence report, so a failure that only appears in one cluster is immediately localizable.

Throughout, I keep the publishable method separate from any proprietary implementation: machine names, customer identifiers, exact parameters, and raw logs stay out, and a simulator or synthetic generator stands in for the confidential interface so the demonstrator can be shared.

IV.1 The method as pseudocode

Algorithm 1 The promotion gate I attach to each CI/CD stage.
Require: artifact a , environment profile π , gate policy G
Ensure: promote / hold decision with an evidence record
1: provision ephemeral env from π $\triangleright$ lab/pre-prod/cloud parity
2: for each gate $g \in G$ (functional, perf, integration, regression) do
3: $res_g \leftarrow \text{RUN}(g, a, \pi)$; attach logs to evidence
4: if res_g violates threshold then
5: return HOLD, classify($res_g \rightarrow$ code / config / env / data)
6: end if
7: end for
8: sign evidence bundle; return PROMOTE, evidence

Table 2. Notation used in Algorithm 1.

Symbol	Meaning
a	Build artifact under promotion.
π	Environment profile (lab / pre-prod / cloud).
G	Ordered gate policy.
res_g	Result of gate g .
$evidence$	Signed bundle of gate logs and verdicts.

IV.2 Design decisions and trade-offs

A few design choices carry most of the weight here, and each is a trade-off I made deliberately rather than a default I inherited:

- **Separate what from where.** The test inventory is environment-agnostic; the environment is a config profile. A new target becomes a new profile, not a forked suite, which is what keeps validation portable instead of duplicated.
- **Synthetic by default.** Driving validation from synthetic workloads calibrated to non-sensitive statistics keeps the framework shareable and the customer's real data out of it. Real traffic enters only as a final, sanitized confirmation.
- **Portable evidence, not just portable tests.** The artifact that travels is the fingerprinted evidence report, not only the test. That is what lets a customer see proof for their environment rather than a claim about a different one.

IV.3 A transparent system model

The scoring model is kept transparent on purpose; an opaque ranking that a release engineer cannot defend will not survive its first incident review. Each candidate x (a test paired with a target environment profile) is scored by a weighted sum of normalized risk features $\phi_k(x) \in [0, 1]$:

$$s(x) = \sum w_k \phi_k(x) \quad (\text{sum over } k = 1 \dots K), \quad \sum w_k = 1, \quad w_k \geq 0.$$

A simple, auditable decision rule then acts on that score against a budget B or a threshold τ chosen from the risk appetite of the environment: execute the top-ranked candidates in order of $s(x)$ until the cost budget B is exhausted, admitting a candidate when $s(x) \geq \tau$.

The weights w_k are not learned in a black box; they are set, reviewed, and re-set by engineers, which is what lets the method be argued about and trusted. Feature definitions ϕ_k are where the domain enters: environment sensitivity, past cross-environment divergence, customer impact of the covered function, and the cost of the run aligned one-to-one with the risk matrix in Table 1.

V. A Worked Example

Suppose a release passes every test in the internal lab and then misbehaves in a customer's Kubernetes cluster. The usual post-mortem discovers a configuration difference nobody tested. In my framework, the same declarative test inventory runs in both places, driven by a per-environment config profile, and each run is fingerprinted into its evidence report. The cross-environment diff shows the lab and the cluster diverging on exactly the cases the configuration difference touches.

The defect is therefore caught as a divergence in the evidence, before the rollout, rather than as an incident after it. Just as important, the evidence that the release was validated travels with it: when the customer asks what was tested, the answer is a fingerprinted report for their environment, not a claim about a different one.

VI. Evaluation Design

A credible evaluation can ride on real cleared data, sanitized logs, or a calibrated synthetic workload whichever path clears review first. Whichever path, the proposed method is compared against at least one honest baseline on the same workload, and every figure reports sample size, conditions, and whether the data are real, sanitized, or synthetic. Table 3 lists the metrics and, more importantly, what each is meant to reveal.

Table 3. Evaluation metrics: what each one is meant to reveal.

Family	Measures	Why it matters
Portability	Environments covered by one suite	Whether the evidence actually travels
Drift detection	Cross-environment result divergence caught	The specific failure this framework targets
Velocity	Release lead time vs. evidence completeness	Whether speed and proof move together

Concretely, I would run the evaluation as a fixed protocol so that anyone could repeat it:

1. Fix a workload (real-approved, sanitized, or synthetic-calibrated) and freeze it for the whole comparison.
2. Run the baseline and the proposed method under identical conditions, changing only the method under test.
3. Repeat each condition at least three times to expose variance, not just a single lucky or unlucky run.
4. Record the metrics of Table 3 with mean, variance, and a confidence interval, plus the environment fingerprint.

- Run the ablation of Table 5 to attribute any gain to a specific component rather than to the architecture as a whole.

VII. Illustrative Results and Reading Them Honestly

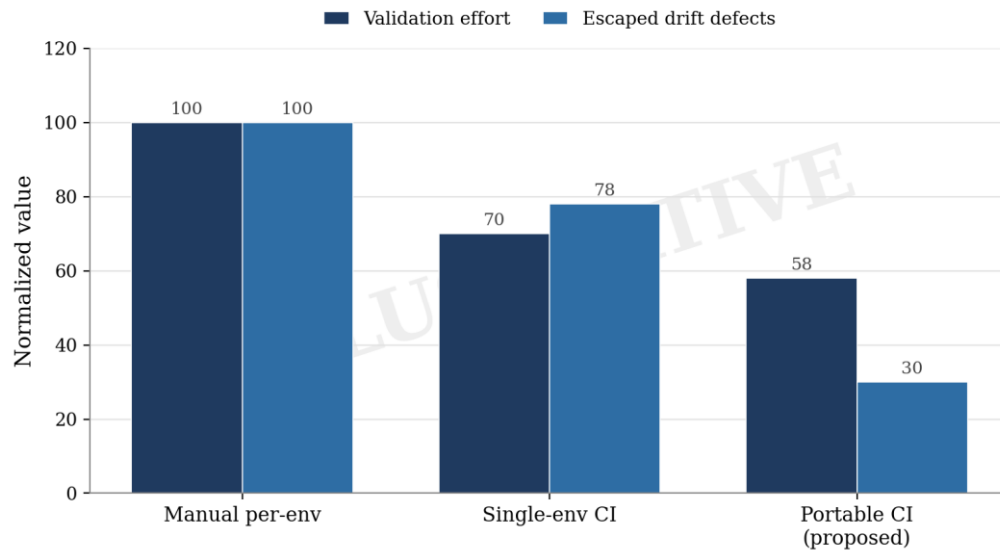


Figure 2.: escaped environment-drift defects by validation approach.

Figure 2 is the target shape: a measurable gain in portability, drift detection, or velocity, with the other two held steady rather than quietly spent. I want to be plain that these are illustrative. A mature version will report means, variance, confidence intervals, and a clear workload description, and for any external-recognition purpose it should carry independent review rather than internal language alone. The schema in Table 4 is what makes those numbers reproducible.

Table 4. Minimal publishable event schema used across the evaluation.

Field	Definition	Type	Purpose
event_id	Unique event / test-execution id	String/UUID	Traceability and replay
timestamp	Event or observation time	ISO time	Sequencing and drift
source	Originating subsystem or stage	Categorical	Failure localization
risk_class	Mapped operational risk	Low/Med/High/Crit	Prioritization
expected	Defined pass/fail/alert	Categorical	Validation
observed	Measured result	Cat./numeric	Comparison
evidence	Pointer to log/report/ticket	URI/hash	Auditability

VIII. Reproducibility Plan

I would start with a lightweight reference prototype, not a production integration: read a bounded input stream, normalize it to the canonical schema, apply the checks, and emit an evidence bundle (input summary, rules run,

pass/fail, timing, anomalies, trace ids, reviewer notes). Proprietary interfaces are swapped for stand-ins a synthetic CDR generator for customer traffic, disposable container profiles for the environments, a REST mock for the northbound interface so a small, GitHub-safe demonstrator can reproduce the results without any employer-owned data.

Table 5. Planned ablations and the effect each is expected to expose.

Ablation	What is removed	Expected effect
Remove configuration profiles	One hard-coded environment	The suite stops travelling; false failures elsewhere
Remove environment fingerprint	No provenance on results	Drift defects cannot be localized
Remove synthetic CDR generator	Depend on production data	Evaluation blocked by confidentiality
Remove CI evidence report	Results live in job logs	Proof does not follow the deployment
Remove baseline comparison	Report proposed only	Weakens the empirical argument

Table 6. From this draft to a submission: the work that remains.

Step	Minimum action	Artifact
Dataset	Generate or sanitize 500–5,000 test executions across ≥ 3 environment profiles	CSV/JSON dataset with schema notes
Baseline	Document the existing manual or rule-based process	Baseline runtime and error report
Prototype	Build the proposed pipeline without proprietary code	GitHub-safe demonstrator
Evaluation	Run ≥ 3 repeated trials per workload	Metrics with mean, variance, CI
Review	Have one domain expert review assumptions and threats	Reviewer comments and revision log

IX. Deployment and Integration

This integrates as a first-class part of the delivery pipeline rather than a test suite run on the side. The declarative inventory and per-environment profiles live in version control next to the product; a pipeline stage materializes the right profile for the target, runs the suite, and publishes the fingerprinted evidence bundle as a release artifact. A promotion to the next environment is gated on that bundle.

The synthetic generators are the part worth investing in early, because they decouple validation from access to customer data and let the same workloads run anywhere. I would still keep one sanitized real-traffic profile per major customer as a final confirmation step, since synthetic data can only approximate the real distribution it stands in for.

X. Why This Sits in My Line of Work

It is my kind of problem: lived pipeline experience recast as a method that can be tested without touching anything proprietary. My record runs along the software/operations seam automation, diagnostics, integration, cross-component debugging and the paper's best version is a small, tightly bounded experiment. The background justifies asking; the data must justify concluding.

XI. Threats to Validity

Synthetic workloads can only approximate a real customer's traffic mix; a framework that looks portable on generated data must still be re-validated against at least one sanitized real workload before strong claims are made. More broadly, the results here are modeled with illustrative data, so the work should not be presented as empirical until a real or faithfully simulated dataset is documented; generalization across environments with different failure costs needs at least two workload profiles; and confidentiality means I publish methods, metrics, and anonymized patterns while keeping proprietary detail out. It is worth stating plainly what would falsify the central claim: if, on logged data with honest baselines, the method failed to improve any of portability, drift detection, or velocity without a compensating loss elsewhere, the contribution would not hold, and I would rather discover that in an ablation than paper over it. Research has to be falsifiable to count, and I have tried to leave this method exposed to exactly that risk.

XII. Broader Applicability and Future Work

Portable validation applies wherever one product must be proven across many heterogeneous deployments not only telecom, but any platform sold into varied customer infrastructure. The environment-fingerprinted evidence bundle is the most broadly useful idea, because it converts 'it passed our tests' into 'here is proof for your environment', which is a stronger and more honest claim in any multi-deployment setting.

Future work must close the synthetic-versus-real gap. A framework that looks portable on generated workloads has to be re-validated against at least one sanitized real-traffic profile, and the divergence between the two measured rather than assumed away. I would also quantify how much release lead time the portability actually returns, since velocity and evidence are supposed to move together here, not trade off.

XIII. Conclusion and Next Steps

The framework is ready for its evidence. Next: generate the non-confidential workload, implement baseline and prototype, run each condition at least three times with variance reported, complete the ablation, and get an outside reviewer to try to break it. That, and only that, converts this from architecture into result.

References

- [1] 3GPP TS 32.240, *Charging architecture and principles*, 2020.
- [2] Cloud Native Computing Foundation, "Kubernetes documentation," 2019. [Online]. Available: <https://kubernetes.io/docs/>
- [3] J. Humble and D. Farley, *Continuous Delivery*. Addison-Wesley, 2010.
- [4] ISO/IEC/IEEE 29119-1:2021, *Software and Systems Engineering Software Testing*, 2020.
- [5] OpenTelemetry, "Documentation," 2020. [Online]. Available: <https://opentelemetry.io/docs/>