

ARGUS vLLM XPU Observatory: An Evidence-Aware Architecture, Prototype Evaluation, and Integration Roadmap

Aaron Singh

Abstract—This paper presents ARGUS vLLM XPU Observatory, a project in a twenty-project AI infrastructure portfolio. The current repository is an active prototype with implemented behavior, automated correctness tests, and explicit boundaries around unverified hardware or production claims. We describe the system model, current implementation, goals, and evaluation method, then propose five integrations that advance the project toward reproducible system-level validation. A local audit executed 3 project tests successfully. No accelerator, deployment, or performance conclusion is inferred unless a corresponding committed benchmark or profiler artifact exists.

Index Terms—ARGUS, AI infrastructure, reproducibility, prototype validation, systems evaluation, hardware–software co-design

I. INTRODUCTION

ARGUS vLLM XPU Observatory (P07) addresses a bounded problem within the portfolio’s track-02-llm-serving track. Its declared status is *Active Prototype* [1]. The project validates its central mechanism before adding device-specific acceleration, production orchestration, or real-world hardware. This ordering matters because optimization without a trusted reference can make incorrect behavior appear successful.

The project goals are to establish deterministic behavior, encode correctness as tests, create machine-readable evidence, identify limiting resources through measurement, and integrate results with adjacent portfolio systems without losing provenance. The repository contains one paper named exactly after its singular folder: `P07-argus-vllm-xpu-observatory.tex`.

II. TECHNICAL CONTEXT

Language-model serving couples execution with cache management, request scheduling, and latency objectives. A simulator can validate policy logic, but only a real model runtime and measured workload can support performance claims. The portfolio benchmark standard requires environment, workload, method, metric, artifact, reproduction, and limitation fields; unknown values remain explicitly unmeasured [3].

III. SYSTEM MODEL AND ARCHITECTURE

The prototype is organized around the following domain model:

$$\text{throughput} = \frac{\sum_i o_i}{t_{last} - t_{first}}, \quad L_{p95} = Q_{0.95}(L) \quad (1)$$

The equation is a design and test abstraction rather than a claimed empirical law. It supports invariants and expected-value checks while later implementations replace synthetic inputs with representative workloads or devices.

The software architecture contains an input/configuration layer, a deterministic core, validation and evidence output, and a local visualization. The principal inspected source artifacts are `python/observatory.py`. Unsupported real-world conditions are surfaced as limitations rather than silently simulated.

IV. DETAILED SCRIPT OPERATION AND RATIONALE

The observatory creates deterministic mock serving records and summarizes median latency, p95 latency, and serial token throughput. The script exists to validate the measurement schema and aggregation logic before connecting HTTP traffic, vLLM, a model, or Intel XPU telemetry.

The execution path is:

- 1) Generate a repeatable set of mock request records.
- 2) Validate that the record set is non-empty.
- 3) Sort latency samples and compute median and p95.
- 4) Calculate serial output-token throughput from total duration.
- 5) Write a benchmark-shaped JSONL record for later real adapters.

V. IMPLEMENTED PROTOTYPE

The metadata and source audit found these completed features [2]:

- Deterministic mock serving workload
- Median, p95, and serial throughput summary
- Unit tests and JSONL record

`python3 -m unittest discover -s tests -p 'test_*.py'` completed successfully with 3 tests. The inspected test artifacts are `tests/test_observatory.py`. This is evidence of local correctness for encoded cases, not production scale or hardware performance.

VI. TESTING METHODOLOGY AND OBSERVED RESULTS

Testing uses Python’s standard `unittest` discovery and exercises the public behavior of the reference implementation. The audit reran the suite from the project folder with `python3 -m unittest discover -s tests -p 'test_*.py'`. All 3 discovered tests passed. The result establishes correctness only for the encoded

TABLE I
IMPLEMENTATION ARTIFACTS AND WHY THEY EXIST

Artifact	Observed responsibility	Engineering rationale
python/observatory.py	simulate, summarize, main	Implements the inspectable, unit-tested project core.
scripts/reproduce.sh	Fixed test and demonstration entry point	Gives another developer one command for local reproduction.
PROJECT.yaml	and	Separates declared intent from evidence-backed implementation.
ANALYSIS.md		Makes outputs inspectable without upgrading simulation into a hardware claim.
streamlit_app.py		

local cases; it does not establish accelerator correctness, real-device behavior, production reliability, or benchmark completion.

No numerical performance result is promoted by this test run. Where scripts emit JSON or JSONL, those outputs remain raw or simulation-specific until a reviewed summary includes hardware, software, workload, warm-up, repetition, correctness threshold, Git revision, and limitations.

VII. CLAIM BOUNDARIES AND RISKS

The project records these unverified or excluded claims:

- No vLLM, model, HTTP service, or XPU is measured

The main risk is confusing synthetic or modeled behavior with deployed-system behavior. Other risks include incomplete workloads, platform-dependent timing, missing failure injection, and interfaces not yet exercised across device boundaries. Performance claims require a reviewed record meeting the portfolio standard.

VIII. EVALUATION PLAN

Evaluation proceeds through correctness tests, deterministic reproduction with Git and environment metadata, repeated benchmarks reporting latency/throughput/memory/error metrics, and a named profiler capture tied to exact hardware and source revision. Success requires reference equivalence within a documented tolerance, preservation of safety and resource invariants, clear failures, and evidence reproducible from a clean environment.

IX. GOALS, MILESTONES, AND SUCCESS CRITERIA

The project goals are staged so that correctness precedes performance and integration. A goal is complete only when its proof artifact is committed or otherwise reviewable; prose or a simulated number alone is insufficient.

X. FUTURE WORK AND INTEGRATIONS

The five project-specific next steps are:

- 1) Real HTTP workload
- 2) vLLM adapter
- 3) Intel XPU metrics
- 4) Add benchmark-standard metadata validation and regression thresholds.
- 5) Publish serving summaries to SENTINEL and POLARIS.

The early items complete declared evidence; the later items connect downstream portfolio consumers. Each integration should add tests and a reviewable artifact such as JSONL evidence, a report, profiler capture, deployment manifest, trace, or labeled data set.

XI. CONCLUSION

ARGUS vLLM XPU Observatory is an evidence-aware active prototype: its implemented behavior and tests are real, while unbuilt hardware, deployment, and performance goals remain labeled. Completing the five integrations in dependency order will advance it from a learning artifact toward a credible portfolio component.

REFERENCES

- [1] Aaron Singh, “ARGUS vLLM XPU Observatory PROJECT.yaml,” local portfolio repository.
- [2] Aaron Singh, “ARGUS vLLM XPU Observatory: README, ANALYSIS, source, and test artifacts,” local portfolio repository.
- [3] Aaron Singh, “AI Infrastructure Portfolio Benchmark Standard,” local portfolio repository.

TABLE II
AUDITED TEST MATRIX

Test artifact and case	Behavior being checked	Observed result
tests/test_observatory.py:test_repeat	Repeat.	Pass (local)
tests/test_observatory.py:test_summary	Summary.	Pass (local)
tests/test_observatory.py:test_empty	Empty.	Pass (local)

TABLE III
DECLARED STATUS VERSUS AUDITED EVIDENCE

Audit field	Finding	Evidence source
Declared status	Active Prototype	PROJECT.yaml
Evidence-backed status	Active local prototype; 3 tests passed	Source plus local unittest run
Accepted measured results	None recorded in measured_results	PROJECT.yaml
Mismatch / claim boundary	No vLLM, model, HTTP service, or XPU is measured	PROJECT.yaml and ANALYSIS.md
Next proof required	Real HTTP workload; vLLM adapter	Planned features

TABLE IV
PROJECT GOALS AND REQUIRED PROOF

ID	Goal	Completion evidence
G1	Real HTTP workload	Passing tests and a reviewed source artifact
G2	vLLM adapter	Machine-readable result with reproduction metadata
G3	Intel XPU metrics	Reference-equivalence or domain-correctness report
G4	Connect real HTTP and vLLM workloads	Named profiler, deployment, or integration artifact
G5	Detect reproducible serving regressions on Intel XPU	Dashboard/report link preserving provenance and limitations