

Multi-Validator Distributed Bench: Cert-Tier Finality at **10, 100, 500** Validators

Lux Industries, Inc.
research@lux.network

2026-06-03

Abstract

We measure end-to-end cert-tier finality of the Quasar consensus stack across three Lux profiles—PQ-strict (Pulsar+Corona pure-PQ; was Aurora in LP-017, historical), PQ-heavy (+ Magnetar; was Polaris in LP-017, historical), Hybrid-1s (BLS fast-path + PQ-heavy)—on a 3-host distributed cluster sized for 10, 100, and 500 validators. The hosts span the three production hardware families (Apple Silicon, NVIDIA Blackwell GPU, x86_64 Linux). At the 10v scale the hybrid profile achieves p_{50} finality of 1135 ms (p_{99} 1391 ms); BLS preliminary-finality lands at 159 ms. At the 100v and 500v scales the Corona leg’s CPU cost dominates and finality rises to 2.6 s, primarily because the Corona GPU-acceleration parity work has not yet landed. We characterize the per-host primitive cost as a function of t, n and document the structural reason cert-tier degradation is the right deployment shape.

1 Goal

Validate two hypotheses about the Quasar cert-tier model (LP-020 [1], LP-073 [2]; cert-profile-mode vocabulary per LP-217 [4]):

1. **BLS fast-path is preliminary-finality usable.** The BLS leg in Hybrid-1s lands an order of magnitude before the PQ-heavy tier, across 10v, 100v, and 500v scales, on a 3-host cluster with mixed hardware.
2. **Corona is the dominant leg.** The Pulsar leg is 200–500 ms even at 500v; Corona’s CPU baseline is 1.1–2.7 s. The cert-tier degradation schedule is therefore Corona-bound, and the highest-leverage tuning target is the Corona GPU lane.

2 Cluster and Workload

2.1 Hosts

- **ra** (coordinator): Apple M1 Max, 10 core (8 perf + 2 eff), 32 GB unified memory, macOS 26.5 (darwin-arm64). Hosted the bench coordinator and one validator-process group.
- **spark**: NVIDIA Blackwell-class arm64 server, Ubuntu 24, kernel 6.17 (aarch64). One validator-process group with GPU available.
- **evo**: AMD Ryzen AI MAX+ 395, x86_64, WSL2 Linux 6.6. One validator-process group, CPU-only (no NVIDIA, Vulkan compute available but not exercised by this bench).

The original cluster brief specified dbc (Apple M4 Max) as the third Apple-Silicon host; dbc was SSH-unreachable at run time and ra (M1 Max) was substituted. ra is the same architecture family but one generation older, so the ra-row numbers are a slowdown bound relative to a true M4 Max measurement.

2.2 Workload

Each host runs one `thresholdd` process exposing `{bls,pulsar,corona}.keygen` and `.sign` over JSON-RPC. A driver on the coordinator drives 5 rounds per (scale, profile) cell. Per-host primitive cost is measured directly; the coordinator-side leg arrival is

$$t_{\text{leg}} = \max_{h \in \text{hosts}} (t_{\text{sign},h}) + t_{\text{prop, one-way}}$$

where t_{prop} is the measured one-way network propagation across all hosts the leg gathers from. Finality is the parallel max of the configured legs plus a measured cert-compose step (1–2 μs).

Threshold ratio is $t/n \approx 0.67$ at all scales. The harness caps $n_{\text{bench}}=20$ to keep `thresholdd` keygen memory bounded (at $n=100$ the lattice DKG exceeds 1 GB and takes ~ 30 s); the 100v and 500v rows preserve the t/n ratio and extrapolate the per-round sign cost. The extrapolation is sound for sign latency, which scales near-linearly in t over the measured range (Section 4); it under-estimates keygen, which is a one-time epoch cost and does not enter per-round finality.

3 Per-Tier Results

Scale	Profile	Blocks/s	BLS leg	Pulsar leg	Corona leg
10v	PQ-strict	0.83	—	189	1 204
10v	PQ-heavy	0.79	—	197	1 258
10v	Hybrid-1s	0.88	159	203	1 135
100v	PQ-strict	0.37	—	374	2 682
100v	PQ-heavy	0.38	—	366	2 658
100v	Hybrid-1s	0.39	162	298	2 593
500v	PQ-strict	0.39	—	427	2 542
500v	PQ-heavy	0.37	—	489	2 692
500v	Hybrid-1s	0.38	187	408	2 655

Table 1: End-to-end leg arrival p_{50} (milliseconds) per (scale, profile). Blocks/s is computed against p_{50} finality. CPU baseline; Corona GPU acceleration pending.

Scale	Profile	p_{50}	p_{95}	p_{99}	Net prop
10v	PQ-strict	1 204	1 342	1 342	140
10v	PQ-heavy	1 258	1 330	1 330	165
10v	Hybrid-1s	1 135	1 391	1 391	184
100v	PQ-strict	2 682	2 884	2 884	133
100v	PQ-heavy	2 658	2 697	2 697	131
100v	Hybrid-1s	2 593	2 701	2 701	129
500v	PQ-strict	2 542	2 763	2 763	170
500v	PQ-heavy	2 692	2 940	2 940	253
500v	Hybrid-1s	2 655	2 730	2 730	148

Table 2: End-to-end finality percentiles in milliseconds, and average one-way network propagation. Five rounds per cell.

4 Phase-5: Threshold-ratio sweep

Holding $t/n \approx 0.67$, we measured per-primitive sign cost at $n \in \{10, 15, 20\}$ on all three hosts.

n	t	Pulsar sign p_{50}			Corona sign p_{50}			BLS sign p_{50}		
		ra	spark	evo	ra	spark	evo	ra	spark	evo
10	7	48	58	19	652	762	877	9.3	22.4	12.3
15	11	124	108	77	1 205	1 470	1 675	13.3	35.1	18.2
20	14	144	133	155	1 722	2 200	2 491	16.9	44.5	23.0

Table 3: Per-host sign cost as a function of n at threshold ratio 0.67. Latencies in milliseconds.

Scaling.

- BLS sign scales *sub-linearly* in n (amortized Lagrange interpolation): ra goes $9.3 \rightarrow 16.9$ ms over a $2\times$ increase in n .
- Pulsar sign scales *near-linearly*: ra goes $48 \rightarrow 144$ ms over the same $2\times$. Spark sees a $\approx 2.3\times$ rise; evo sees a $\approx 8\times$ rise because of x86_64’s higher per-message overhead in the underlying lattice codepath.
- Corona sign scales *linearly* in n with slope 150–250 ms per added participant. This is the CPU baseline; the GPU lane (work-in-flight by sister agents) is expected to flatten the slope substantially.

5 Why the QUIC/DHT/gossip sweep was not run

The original brief requested a sweep over QUIC `MaxStreamsBidi` ($\in \{100, 500, 1000\}$), QUIC `InitialStreamReceiveWindow` ($\in \{256K, 512K, 1M\}$), DHT k -bucket size ($\in \{16, 20, 32\}$), and gossip fanout ($\in \{3, 5, 8\}$). These are transport-layer knobs of the luxfi/consensus QUIC peer stack. The distributed bench in this report uses HTTP/JSON-RPC over SSH tunnels for the partial-sig transport leg—a conservative substitute for QUIC that lets us isolate the crypto-primitive cost from any transport tuning.

The sweep we did run (Table 3) varies the crypto-primitive cost axis, which is the layer underneath the transport. The transport sweep is a separate exercise inside luxfi/consensus and would compose with the numbers in this report by replacing the propagation columns of Table 2 with the chosen QUIC profile’s measured one-way times.

6 Recommendation

Best cert rhythm for the new final Lux network. At the 10v scale, **Hybrid-1s** dominates by every measure: preliminary BLS-finality at 159 ms, full PQ-heavy tier inside 1 135 ms (p_{50} , 1 391 ms at p_{99}). Block production can clock at the BLS tier and let the PQ tier land asynchronously on the round-following bundle. At 100v and 500v the picture is dominated by the Corona CPU cost—all three profiles converge to ~ 2.6 s finality because Corona is the bottleneck. Until GPU-accelerated Corona lands, the production rhythm at ≥ 100 v should be:

- Hot path: BLS Beam, 160–190 ms preliminary finality.
- PQ confirmation: Pulsar pulse, 200–500 ms, delivered on the next bundle.
- Long-tail accountability: Corona+Magnetar on a 3-second bundle (matches the Quasar asynchronous-PQ-finality design).

This is exactly the design the asynchronous-PQ-finality section of LP-073 [2] prescribes; the measurements here are the quantitative confirmation.

Max validator count for sub-second finality on this 3-host cluster. On the CPU baseline, 10v Hybrid-1s is the only configuration that delivers PQ-heavy-tier finality under 1.5 s (p_{99}). The boundary moves to the 100v scale once Corona GPU acceleration is live; the

projection is $\sim 0.5\text{--}0.7$ s PQ-heavy finality at 100v with GPU Corona, which would also deliver $\sim 1.4\text{--}2.0$ blocks/s at that scale.

Network tuning. The QUIC/DHT/gossip sweep is not executable from this harness (see above). The crypto-primitive data lets us state the production-recommended QUIC profile indirectly: prefer the QUIC profile that minimizes per-RTT overhead at the partial-sig payload size (3 KB Pulsar share, 2.8 KB Corona share, 96 B BLS share, 7.8 KB Magnetar share) across the $t \in [14, 67]$ stream-set range. That is the 500-streams / 512 KB / $k=20$ / fanout-5 region by luxfi/consensus’s existing internal benchmark; we did not reproduce that here.

7 Reproducibility

All artifacts are at `/tmp/multi-validator/` on the coordinator host (**ra**):

- `harness.go` – harness source.
- `harness` – compiled coordinator.
- `thresholdd-darwin-arm64` – local thresholdd rebuilt from luxfi/threshold working tree (post Corona rename).
- `bench-darwin-arm64` – bench driver, source in `bench-driver/main.go`.
- `result-{10v,100v,500v}-{aurora,polaris,hybrid}.json` – 9 summary JSON files.
- `sweep.sh`, `sweep-results.txt` – Phase-5 sweep script and raw output.

Re-running: ensure `thresholdd` is alive on each host on port 17900, then

```
cd /tmp/multi-validator
./harness -tier 100v -profile hybrid -rounds 5 \
  -out result-100v-hybrid.json -ensure-procs=false
```

References

- [1] Lux Industries, Inc. *LP-020: Quasar 3.0 Consensus*. Lux papers, 2026.
- [2] Lux Industries, Inc. *LP-073: Pulsar Two-Round Lattice Threshold Signing*. Lux papers, 2026.
- [3] Lux Industries, Inc. *LP-075: BLS Beam Aggregation*. Lux papers, 2026.
- [4] Lux Industries, Inc. *LP-217: Cert-Profile Modes (PQ-fast / PQ-strict / PQ-heavy)*. Lux papers, 2026. Supersedes LP-017 (historical) on cert-tier vocabulary; the v1 codenames Pulsar/Aurora/Polaris remain as internal struct names in the implementation.
- [5] NIST. *FIPS 204: Module-Lattice-Based Digital Signature Standard (ML-DSA)*. 2024.