



TenQ: Work for a High-Rate Block Graph

The deterministic proof-of-work profile for TRAI DAG

Tenrai Technologies

tenrai-technologies.org

Litepaper v1.0 · September 2026

ABSTRACT

TENQ is the proof-of-work function designed for TRAI DAG's permissionless work layer. It converts the canonical, anchor-bound block header into a 256-bit work value and preserves the familiar nonce, extranonce and target interface used by miners and pools. The production profile combines a 32 MiB epoch cache, a 2 GiB shared dataset, an 8 KiB private scratchpad and four dependent integer programs over 1,024 rounds.

The same consensus function can be evaluated in two ways. A Full miner or verifier keeps the 2 GiB dataset resident; a Light evaluator reconstructs requested dataset nodes from the 32 MiB cache. Both return exactly the same result. The design therefore allows a memory–computation trade-off without creating a privileged verification rule.

TENQ is intended to make efficient use of commodity GPUs while forcing specialized miners to reproduce a broad mix of memory access, integer arithmetic, private state and control dependence. Hardware specialization and quantum search are therefore analyzed as cost models, while deterministic consensus behavior is specified exactly.

1 Role in TRAI DAG

TRAI DAG does not use proof of work to choose a longest chain. Miners produce anchored work blocks into a graph; stake-certified consensus later commits registration and execution. TENQ answers a narrower question: has this work proposal paid the required computational cost against the target fixed by its anchor? [1]

A canonical 323-byte preimage binds the chain identity, protocol version, work reference, parents, payload commitment, payload size and mass, reward destination, nonce and extranonce. The anchor selects the target and work profile. A block passes the low-level work predicate when

$$V = \text{TenQ}_A(P), \quad \text{BE256}(V) \leq T_A.$$

The target controls the probability of success, not the number of rounds performed by one evaluation.

Table 1. Production profile.

Resource	Value
Epoch cache	32 MiB
Materialized dataset	2 GiB
Attempt-local scratchpad	8 KiB
Logical lanes	32
Programs	4×256 instructions
Rounds	1,024
Dataset reconstruction	160 cache accesses per node
Work value	256-bit unsigned target value

Keeping nonce, extranonce and target difficulty is deliberate. Mining pools can allocate search space without changing the solved payload or reward destination, and existing anchor-based reward accounting remains coherent. A later useful-compute extension can reuse that settlement boundary, but it requires a separately specified job model, verification rule and activation path; it is not part of the v1 work profile.

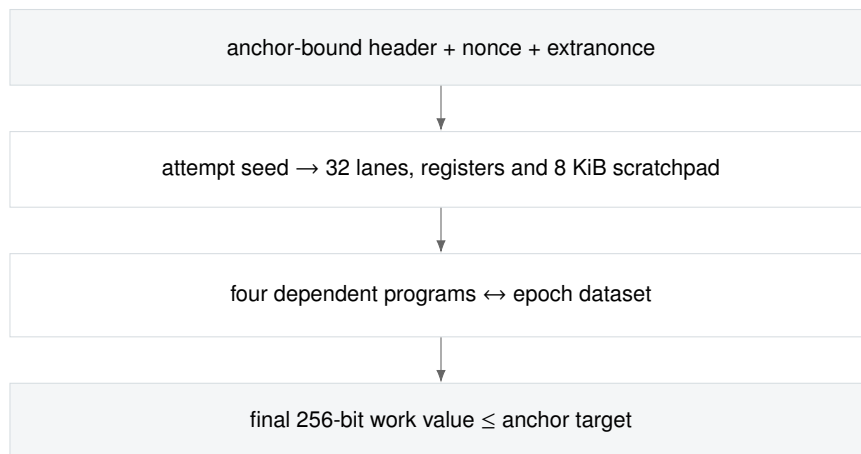


Figure 1. One TenQ work attempt. Ledger ordering remains a separate TRAI DAG consensus function.

2 Memory and execution

Epoch preparation begins with a public seed derived from chain identity, profile, revision and finalized epoch height. A 32 MiB cache is mixed through ordered recurrences. Dataset nodes are then derived independently from that cache, each through 160 data-dependent cache references. The Full path materializes all nodes into 2 GiB; the Light path reconstructs a node only when it is requested.

Every nonce attempt derives fresh registers, scratchpad and a first program key. Four phases then execute sequentially. Later program keys depend on earlier state, so the complete instruction stream is not merely an epoch-wide public program that can be screened before the attempt begins. Each round reads one dataset node and scratchpad state, performs fixed-width integer operations, updates lane-private scratchpad regions, mixes lane outputs and derives the next memory index and branch state.

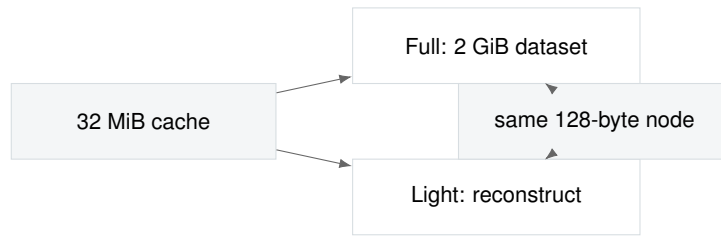


Figure 2. Memory residency changes cost, not validity.

The logical round order is part of consensus; the physical schedule is not. CUDA, OpenCL and CPU implementations may map the 32 logical lanes differently, but they must reproduce the same read-before-write transition and final reduction. This is what makes cross-platform evaluation deterministic rather than device-specific.

Why both shared and private memory matter

The 2 GiB dataset is reused across many attempts, so it is not 2 GiB of new work per nonce. The private scratchpad and registers, however, scale with concurrency. A specialized miner must therefore decide how much shared dataset to store, how much to recompute, and how much local state to provision per engine. The algorithm intentionally exposes this trade-off rather than relying on one large array as a proof of memory hardness.

3 Security model

Determinism

For a fixed preimage, authenticated anchor and epoch material, TENQ has one defined result. Cache generation, dataset reconstruction, program generation, round ordering and final reduction use exact integer semantics. Full and Light evaluation are mathematically equivalent because both call the same deterministic dataset-node function.

ASIC and FPGA resistance

A public function cannot prevent specialized hardware from implementing it. TENQ instead attempts to reduce the economic advantage of narrow specialization by combining irregular memory access, a broad integer instruction mix, branch dependence and private writable state. Full, Light and intermediate storage strategies all remain available to a miner. The proper comparison is therefore accepted proofs per unit of cost or energy after each implementation has been optimized [2, 3].

This also means the claim is falsifiable. If a future ASIC demonstrates a large, sustained cost advantage, the resistance objective has not been met merely because the ASIC still performs the nominal algorithm. The design should be judged by measured economics, not by the statement that a dataset is “large.”

Quantum search

Mining remains a public target search. Under a uniform-output model, one attempt succeeds with probability $p = (T_A + 1)/2^{256}$. Generic amplitude amplification reduces ideal query complexity from order $1/p$ to order $1/\sqrt{p}$ [4, 5]. Whether that becomes an economic mining advantage depends on the reversible cost of the complete TENQ oracle, including memory access, arithmetic, control state, uncomputation and fault tolerance.

A high block-production rate is not by itself a quantum deadline. TRAI DAG work is anchored to a finalized reference that remains eligible over a height-based registration window; a new sibling block does not automatically invalidate another miner’s anchor. Quantum analysis must therefore use the actual anchor lifetime rather than the average time between work blocks.

4 Performance and integration

The published production-profile CPU microbenchmark used a fully materialized 2 GiB dataset on an AMD EPYC 9V74 KVM guest with five visible virtual CPUs and a sustained quota of four CPU-seconds per second. Median Full-path throughput was 570.5 evaluations/s with one worker and 1,866.6 evaluations/s with five workers; the one-worker Light path measured 32.33 evaluations/s. Cache and dataset construction were measured separately at 2.221 s and 111.981 s [6].

Table 2. Production-profile reference measurements.

Path	Workers	Median eval/s
Full	1	570.5
Full	2	1,151.0
Full	4	1,893.7
Full	5	1,866.6
Light	1	32.33

For work-block rate B and average count u of direct parent proofs not already cached, a basic verifier-demand model is $B(1 + u)$. At 100 blocks/s this is 100 evaluations/s when all parents are already known and 1,700 evaluations/s if all 16 permitted parents are new for every block. Against the measured 1,866.6 evaluations/s reference rate, those cases consume 5.36% and 91.07% of the measured proof capacity respectively. The calculation is a work-component budget, not a whole-node throughput claim.

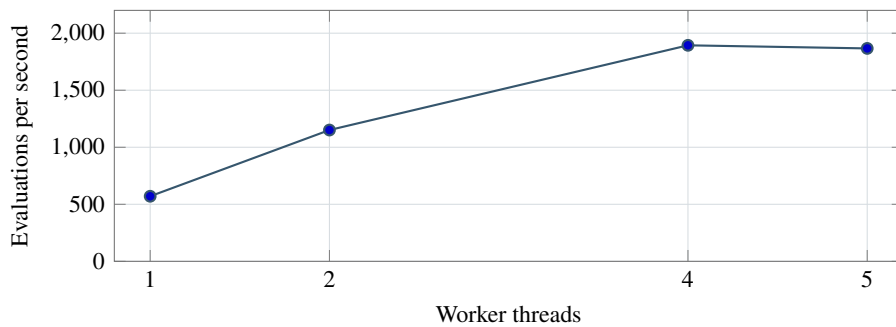


Figure 3. Measured Full-path CPU evaluation rate. The fifth worker shares a four-CPU quota.

In deployment, the work verifier is only one component. Nodes also verify post-quantum signatures, execute transactions, maintain state, relay blocks and retain enough old epoch material to validate still-eligible anchors and parents. Profile changes therefore require coordinated activation; reducing rounds or reconstruction work locally would define a different proof-of-work function.

The current TRAIDAG v1 profile keeps the compute settlement lane disabled. A future useful-compute system may reuse the surrounding anchor and settlement architecture, but it requires an independent job format, verification protocol and activated settlement rules. Those mechanisms are outside the v1 TenQ validity predicate.

5 Summary

TenQ combines a public epoch cache and dataset with nonce-specific writable state and dependent integer programs. Full and Light implementations evaluate the same function, making verification deterministic while allowing a clear memory–computation trade-off. The production profile retains the ordinary target-search interface needed by existing miner and pool workflows and is designed so that many independent proofs can be checked in parallel.

Deterministic evaluation and Full/Light equivalence follow from the specification. ASIC/FPGA resistance remains an economic property to be measured against optimized hardware, while quantum search admits the generic amplitude-amplification model and depends in practice on the reversible cost of the complete work predicate. The resulting security claims are therefore explicit about which properties are established by construction and which require empirical hardware evidence.

References

- [1] Tenrai Technologies and contributors. *TRAIDAG: Certified Ordering over a Proof-of-Work Block Graph*. Technical Whitepaper v1.2.3, September 2026.
- [2] tevador and RandomX contributors. *RandomX design*. Project design document and reference implementation.
- [3] IfDefElse and contributors. *ProgPoW: Programmatic Proof-of-Work*. Design and reference implementation.
- [4] L. K. Grover. A fast quantum mechanical algorithm for database search. *STOC*, 1996.
- [5] G. Brassard, P. Høyer, M. Mosca and A. Tapp. Quantum Amplitude Amplification and Estimation. 2002.
- [6] Tenrai Technologies. *TenQ v1.0.0 Benchmark Records*. September 2026. Companion source-release data.