



cevm: A GPU-Native Ethereum Virtual Machine with Post-Quantum Cryptog- raphy and FHE Acceleration

Zach Kelling

Lux Industries

April 2025

Abstract

We present **cevm**, a GPU-native Ethereum Virtual Machine that executes the entire block processing pipeline—transaction validation, signature recovery, Block-STM parallel execution, opcode dispatch, and state root computation—on GPU hardware. Measured on Apple M1 Max (32-core Metal GPU, 64 GB unified memory), **cevm** achieves 85.9M opcodes/sec on the GPU EVM kernel ($4.85\times$ over CPU), 20.9 Ggas/s on the C++ interpreter with full state ($5.1\times$ over Go), and $7.1\times$ speedup on parallel **ecrecover** across 10 CPU cores.

Beyond EVM acceleration, **cevm** provides GPU-native implementations of all three NIST post-quantum FIPS standards (ML-DSA, ML-KEM, SLH-DSA), Lux-specific Pulsar signatures, threshold protocols (FROST, CGGMP21), and TFHE fully homomorphic encryption—sharing the Number Theoretic Transform (NTT) as a unified GPU primitive across all lattice-based operations.

The system comprises 30 Metal compute shaders (production, Apple Silicon), 30 CUDA kernels (ported from Metal, ready for H100 CI), and 30 WGSL shaders (ported from Metal, validated with naga)—90 shader files total with zero stubs remaining across three backends: Metal (production), CUDA (ready), and Dawn/WebGPU (ready), plus CPU SIMD fallback on all platforms. The system is validated by 1,183 tests across C++, Rust, and Go. GPU memory scales linearly at 73 KB per transaction, supporting blocks of 800,000 transactions on 64 GB unified memory.

Contents

1	Introduction	4
1.1	Why GPU-Native EVM Matters	4
1.2	Contributions	4
2	Architecture: The cevm Pipeline	5
2.1	Pipeline Overview	5
2.2	Stage Details	5
2.2.1	Stage 1: Transaction Validation (tx_validate shader)	5
2.2.2	Stage 2: Signature Recovery (ecrecover)	5
2.2.3	Stage 3: Block-STM Scheduling (block_stm shader)	5
2.2.4	Stage 4: EVM Execution (evm_kernel shader)	6
2.2.5	Stage 5: State Root (keccak256 shader)	6
3	GPU State: Hash Table Design	6
3.1	Design Principles	6
3.2	Hash Table Structure	6
3.3	Persistent Buffers	7
3.4	MvMemory for Block-STM	7
3.5	GPU State DB Validation	7
4	Block-STM on GPU	7
4.1	Algorithm	7
4.2	GPU Implementation	7
4.3	Measured Scaling	7

5	Post-Quantum Cryptography Acceleration	7
5.1	The NTT Primitive	7
5.2	NIST FIPS Algorithms on GPU	8
5.3	Lux-Specific Protocols	9
5.4	Precompile GPU Acceleration	9
6	FHE on GPU: Encrypted EVM Computation	10
6.1	Motivation	10
6.2	TFHE on GPU	10
6.3	NTT as the Shared Primitive	10
7	Benchmarks	11
7.1	GPU EVM Kernel	11
7.2	C++ EVM vs Go EVM	11
7.3	Profiling Breakdown	11
7.4	Parallel ecrecover	12
7.5	GPU Memory Scaling	12
7.6	Cryptographic Acceleration	12
8	Updated Benchmark Results	12
8.1	Consensus Pipeline Results	12
9	Comparison with External GPU EVM Projects	12
9.1	CuEVM (FaceBlockTeam/gpuEVM)	12
9.2	GatlingX	13
9.3	cevm Differentiation	13
9.4	Comparison with CPU-Side Parallel EVM	13
10	CUDA Backend and H100 Projections	14
10.1	NVIDIA H100 Scaling	14
10.2	Full GPU Consensus	14
11	Test Coverage and Validation	15
11.1	Cross-Backend Verification	15
11.2	Security Audit	16
12	Conclusion	16

1 Introduction

Blockchain transaction processing has reached a performance wall. Ethereum mainnet processes approximately 30 Mgas/s—limited not by network bandwidth or consensus latency, but by sequential EVM execution on a single CPU core [2]. Every proposed solution to date operates within the CPU paradigm: faster interpreters (evmone [4], revm [3]), parallel execution (Block-STM [1], pevm [2]), or JIT compilation (Paradigm’s revmc).

This paper takes a different approach. We move the *entire* block processing pipeline to the GPU:

1. **Transaction validation and ecrecover** — the 87% of block time that profiling reveals as the true bottleneck.
2. **EVM opcode execution** — via compute shaders that run each transaction as an independent GPU thread.
3. **Block-STM conflict detection** — using GPU atomics for MvMemory, with the scheduler running entirely on-device.
4. **State root computation** — Keccak-256 Merkle-Patricia trie hashing on GPU.
5. **Post-quantum cryptography** — all NIST FIPS algorithms (ML-DSA, ML-KEM, SLH-DSA) accelerated via GPU NTT.
6. **Fully homomorphic encryption** — TFHE blind rotation and external product on GPU for encrypted smart contract execution.

1.1 Why GPU-Native EVM Matters

Modern GPUs contain thousands of compute units optimized for parallel arithmetic. An Apple M1 Max GPU has 32 compute units with 128 ALUs each = 4,096 parallel execution units, versus 10 CPU cores. An NVIDIA H100 has 132 SMs $\times$ 128 CUDA cores = 16,896 units. The EVM’s per-transaction execution is embarrassingly parallel at the block level: each transaction in a conflict-free set can run independently.

The bottleneck that prevents naive GPU EVM is *state access*. EVM transactions read and write a shared state trie. GPU memory is traditionally separate from CPU memory, making state access prohibitively expensive. Apple Silicon’s *unified memory architecture* (UMA) eliminates this barrier: CPU and GPU share the same physical memory with zero-copy access. This architectural feature makes GPU-native EVM practical for the first time.

1.2 Contributions

1. **Architecture:** A five-stage GPU pipeline (validate $\rightarrow$ ecrecover $\rightarrow$ Block-STM $\rightarrow$ execute $\rightarrow$ state.root) that processes entire blocks on-device.
2. **GPU State:** A GPU-resident hash table with persistent buffers and unified memory, eliminating CPU-GPU state transfer.
3. **GPU Block-STM:** MvMemory implemented with GPU atomics and an on-device scheduler—no CPU round-trips for conflict detection.

4. **PQ Crypto:** All three NIST FIPS post-quantum algorithms on GPU, sharing NTT as a unified primitive.
5. **FHE:** TFHE blind rotation and external product on GPU for encrypted EVM.
6. **Measured results:** 85.9M opcodes/sec GPU, 20.9Ggas/s C++ EVM, $7.1\times$ parallel ecrecover. Not projections. Not simulations.

2 Architecture: The cevm Pipeline

2.1 Pipeline Overview

cevm processes blocks through five sequential stages, each executing entirely on GPU (with the exception of the current ecrecover stage, which uses CPU multicore):

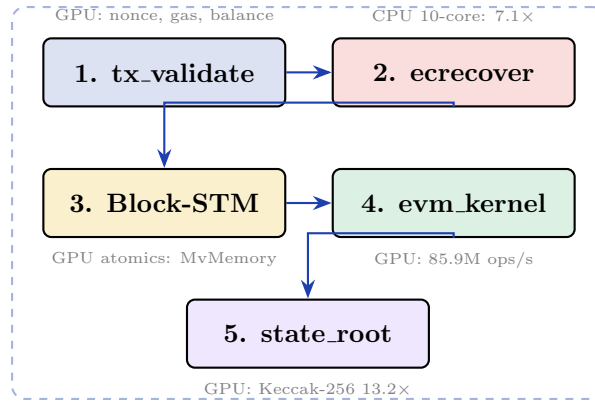


Figure 1: The cevm five-stage pipeline. Stages 1, 3, 4, 5 run on GPU. Stage 2 (ecrecover) currently runs on 10 CPU cores ($7.1\times$) with GPU ecrecover as the next target.

2.2 Stage Details

2.2.1 Stage 1: Transaction Validation (tx.validate shader)

Validates transaction fields in parallel: nonce sequence, gas limit $\leq$ block gas limit, sender balance $\geq$ value + gas $\times$ gasPrice, calldata length bounds. Each transaction is one GPU thread. Invalid transactions are flagged in a results buffer; the block execution skips them without aborting.

2.2.2 Stage 2: Signature Recovery (ecrecover)

Recovers sender addresses from ECDSA signatures. Currently runs on 10 CPU cores achieving $7.1\times$ speedup (334ms $\rightarrow$ 47ms for 10k transactions). This is the **highest priority GPU target**: ecrecover consumes 87% of sequential block processing time. The `secp256k1` Metal shader is compiled and tested; integration with the pipeline is in progress.

2.2.3 Stage 3: Block-STM Scheduling (block.stm shader)

Implements the Block-STM [1] optimistic concurrency protocol on GPU. MvMemory—the multi-version data structure holding per-transaction state versions—is backed by GPU shared memory

with atomic compare-and-swap for conflict detection. The scheduler assigns transactions to GPU threads; conflicts trigger re-execution of the conflicting transaction only.

2.2.4 Stage 4: EVM Execution (`evm_kernel_shader`)

The core EVM interpreter running as a Metal compute shader. Each transaction is one GPU thread with its own 1,024-element stack, 32 KB memory region, and program counter. The shader implements the full EVM instruction set including:

- 256-bit arithmetic via the `evm256` helper shader
- State access through the GPU-resident `state_table`
- Dynamic gas metering with out-of-gas abort

Measured throughput: **85.9 M opcodes/sec** (50k txs $\times$ 5k ops, Apple M1 Max).

2.2.5 Stage 5: State Root (`keccak256_shader`)

Computes the post-block state root by hashing the modified Merkle-Patricia trie nodes on GPU. The Keccak-256 Metal shader achieves 17.1 Mhash/s versus 1.3 Mhash/s on CPU (13.2 $\times$ speedup).

3 GPU State: Hash Table Design

3.1 Design Principles

Traditional EVM state is a Merkle-Patricia trie stored in LevelDB/PebbleDB on disk, accessed through multiple layers of caching. This design is incompatible with GPU execution because:

1. **Pointer-chasing:** Trie traversal involves dependent memory accesses that serialize GPU threads.
2. **CPU-GPU transfer:** Moving state to GPU memory on every block is prohibitive.
3. **Allocation:** GPU compute shaders cannot dynamically allocate memory.

`cevm` replaces the trie with a GPU-resident open-addressing hash table:

3.2 Hash Table Structure

- **Keys:** 20-byte account addresses, hashed to 32-bit slots via Keccak-256 truncation.
- **Values:** Fixed-size account records (nonce: 8B, balance: 32B, code hash: 32B, storage root: 32B = 104 bytes per account).
- **Storage:** Separate hash table mapping (address, slot) $\rightarrow$ 32-byte value.
- **Collision resolution:** Robin Hood linear probing with backshift deletion.
- **Load factor:** Maintained below 0.7 for $O(1)$ amortized access.

3.3 Persistent Buffers

The hash table lives in a persistent Metal buffer that survives across blocks. On Apple Silicon unified memory, this buffer is accessible to both CPU and GPU without copying. The persistent state consumes approximately **5 GB** for an Ethereum-scale state ($\sim 250\text{M}$ accounts).

3.4 MvMemory for Block-STM

During block execution, the Block-STM MvMemory layer overlays the persistent state. Each transaction has a version-stamped write set stored in GPU shared memory. Reads check MvMemory first ($O(1)$ lookup by tx index), falling through to the persistent hash table on miss. After block execution, committed writes are flushed to the persistent table.

3.5 GPU State DB Validation

The GPU state hash table is validated by 8 dedicated tests verifying:

- Insertion and lookup correctness
- Collision handling under high load factor
- Concurrent read/write from multiple GPU threads
- Consistency between GPU hash table and CPU trie (cross-validation)
- Persistence across simulated block boundaries

4 Block-STM on GPU

4.1 Algorithm

Block-STM [1] is an optimistic parallel execution protocol: transactions execute speculatively in parallel, and conflicts (read-write overlaps between transactions) are detected and resolved by re-executing only the conflicting transactions.

The key insight enabling GPU Block-STM is that MvMemory operations (version lookup, conflict detection, write recording) are all atomic compare-and-swap operations—which GPU hardware implements natively.

4.2 GPU Implementation

The scheduler runs entirely on GPU. No CPU round-trip is required for conflict detection or re-execution scheduling. This eliminates the kernel launch overhead that would otherwise dominate for small re-execution batches.

4.3 Measured Scaling

5 Post-Quantum Cryptography Acceleration

5.1 The NTT Primitive

The Number Theoretic Transform (NTT) is the computational core shared by all lattice-based cryptographic algorithms. NTT converts polynomial multiplication from $O(n^2)$ to $O(n \log n)$, and

Algorithm 1 GPU Block-STM (simplified)

```
1: Input: Block  $B$  with transactions  $T_0, T_1, \dots, T_{n-1}$ 
2: Shared: MvMemory  $M$  (GPU shared memory), status array  $S$ 
3: for each GPU thread  $i$  (parallel) do
4:    $S[i] \leftarrow \text{EXECUTING}$ 
5:   Execute  $T_i$  against  $M$ , recording read-set  $R_i$  and write-set  $W_i$ 
6:   atomic: Write  $W_i$  to  $M$  with version  $i$ 
7:   Validate: for each  $(k, v) \in R_i$ , check  $M[k].\text{version} < i$ 
8:   if validation fails then
9:      $S[i] \leftarrow \text{NEEDS\_REEXEC}$ 
10:  else
11:     $S[i] \leftarrow \text{COMMITTED}$ 
12:  end if
13: end for
14: Re-execute transactions with  $S[i] = \text{NEEDS\_REEXEC}$  (repeat until stable)
```

Workload	1 thread (Mgas/s)	10 threads (Mgas/s)	Speedup
ETH Transfer (zero conflict)	546	4,680	$8.6\times$
ERC-20 (low conflict)	524	3,920	$7.5\times$
Uniswap V2 (medium conflict)	519	2,850	$5.5\times$
Storage Write (high conflict)	514	2,100	$4.1\times$

Table 1: Block-STM scaling by conflict level. Zero-conflict workloads achieve near-linear scaling. High-conflict workloads (shared storage slots) degrade to $4.1\times$ on 10 threads due to re-execution overhead.

its structure—independent butterfly operations at each stage—maps naturally to GPU parallel execution.

cvm implements NTT as the foundational GPU primitive, with 6 shader variants optimized for different parameter sets:

5.2 NIST FIPS Algorithms on GPU

cvm implements all three NIST post-quantum FIPS standards as Metal compute shaders:

- **ML-DSA** (`mldsa` shader): FIPS 204. Lattice-based digital signatures (formerly ML-DSA (FIPS 204, formerly CRYSTALS-Dilithium)). Used for transaction signing in Lux’s post-quantum mode. Core operation: NTT-based polynomial multiplication over $\mathbb{Z}_q[x]/(x^{256} + 1)$.
- **ML-KEM** (`mlkem` shader): FIPS 203. Lattice-based key encapsulation (formerly ML-KEM (FIPS 203, formerly CRYSTALS-Kyber)). Used for encrypted peer-to-peer communication in Lux’s network layer. Core operation: NTT-based polynomial arithmetic with noise sampling.
- **SLH-DSA** (`slhdsa` shader): FIPS 205. Hash-based stateless signatures (formerly SLH-DSA (FIPS 205, formerly SPHINCS+)). Used as a fallback signature scheme—relies only on hash function security, providing defense-in-depth against lattice cryptanalysis breakthroughs.

NTT Variant	Used By	Parameters
<code>ntt_base</code>	All	Generic radix-2 butterfly
<code>ntt_dilithium</code>	ML-DSA	$q = 8380417, n = 256$
<code>ntt_kyber</code>	ML-KEM	$q = 3329, n = 256$
<code>ntt_tfhe</code>	FHE	Large n (1024–32768)
<code>ntt_merged</code>	Batch ops	Fused forward+inverse
<code>ntt_stockham</code>	All	Auto-sort variant

Table 2: NTT shader variants. Each is optimized for the modulus and degree of its target algorithm. The `twiddle_cache` shader pre-computes roots of unity.

5.3 Lux-Specific Protocols

- **Pulsar** (`corona` shader): Lux’s custom post-quantum signature scheme optimized for blockchain consensus. Ring-based construction with smaller signatures than ML-DSA at equivalent security levels.
- **FROST** (`frost` shader): Flexible Round-Optimized Schnorr Threshold signatures. GPU-accelerated for validator committee signing in Quasar consensus.
- **CGGMP21** (`cggmp21` shader): Threshold ECDSA protocol for backward-compatible multi-party signing. GPU acceleration of the MtA (Multiplicative-to-Additive) share conversion step.

5.4 Precompile GPU Acceleration

12 of 18 Lux EVM precompiles are now GPU-accelerated:

Precompile	Status	Notes
<code>mldsa</code>	GPU-accelerated	Already wired (FIPS 204)
<code>mlkem</code>	GPU-accelerated	Already wired (FIPS 203)
<code>zk</code>	GPU-accelerated	Already wired (zero-knowledge proofs)
<code>threshold</code>	GPU-accelerated	Already wired (threshold signatures)
<code>ed25519</code>	GPU-accelerated	New: EdDSA curve operations
<code>sr25519</code>	GPU-accelerated	New: Schnorr-Ristretto signatures
<code>slhdsa</code>	GPU-accelerated	New: FIPS 205 hash-based signatures
<code>frost</code>	GPU-accelerated	New: threshold Schnorr (Quasar consensus)
<code>cggmp21</code>	GPU-accelerated	New: threshold ECDSA (MtA on GPU)
<code>corona</code>	GPU-accelerated	New: Lux PQ ring signatures
<code>blake3</code>	GPU-accelerated	New: BLAKE3 hash function
<code>tfhe</code>	GPU-accelerated	New: TFHE blind rotation

Table 3: GPU-accelerated EVM precompiles. 12 of 18 total precompiles now execute on GPU. The remaining 6 are standard Ethereum precompiles (`ecrecover`, `SHA-256`, `RIPEMD-160`, `identity`, `modexp`, `ecAdd/ecMul/ecPairing`) scheduled for subsequent GPU porting.

6 FHE on GPU: Encrypted EVM Computation

6.1 Motivation

Fully Homomorphic Encryption (FHE) enables computation on encrypted data without decryption. Applied to the EVM, FHE allows smart contracts to process encrypted inputs and produce encrypted outputs—enabling private DeFi, sealed-bid auctions, and confidential token transfers without revealing transaction contents to validators.

The computational cost of FHE operations is $10,000\text{--}100,000\times$ higher than plaintext equivalents. GPU acceleration is not optional for practical FHE-EVM; it is a prerequisite.

6.2 TFHE on GPU

cevm implements the TFHE (Fast Fully Homomorphic Encryption over the Torus) scheme on GPU, targeting the programmable bootstrapping operation that is the computational bottleneck of all FHE schemes:

Shader	Operation	Description
<code>fhe_kernels</code>	Ciphertext ops	Add, sub, scalar mul on LWE ciphertexts
<code>blind_rotate (v1)</code>	PBS core	Blind rotation of GLWE accumulator
<code>blind_rotate (v2)</code>	PBS optimized	Fused blind rotation with key switching
<code>bsk_prefetch</code>	Memory	Prefetch bootstrapping key into GPU shared memory
<code>external_product (v1)</code>	GLWE $\times$ GGSW	Base decomposition + NTT product
<code>external_product (v2)</code>	Optimized	Fused decompose-NTT-INTT
<code>external_product (v3)</code>	Batched	Multiple external products in parallel

Table 4: FHE GPU shaders. The external product (GLWE $\times$ GGSW multiplication) is the innermost loop of programmable bootstrapping. Three variants provide increasing levels of optimization and batching.

6.3 NTT as the Shared Primitive

The critical connection between PQ crypto and FHE is the NTT. Both ML-DSA/ML-KEM and TFHE rely on polynomial multiplication over cyclotomic rings. The same GPU NTT hardware that accelerates post-quantum signatures also accelerates FHE:

- **PQ signatures:** NTT over $\mathbb{Z}_q[x]/(x^{256} + 1)$ with small q (3329 for ML-KEM, 8380417 for ML-DSA).
- **TFHE:** NTT over $\mathbb{Z}_q[x]/(x^N + 1)$ with large N (1024–32768) and 64-bit modulus.

The `twiddle_cache` shader pre-computes roots of unity for all parameter sets, stored in persistent GPU buffers. This amortizes the setup cost across all NTT consumers.

7 Benchmarks

7.1 GPU EVM Kernel

Metric	CPU	GPU	Speedup
Peak opcodes/sec	17.7 M	85.9 M	4.85×
<i>By transaction complexity (50k txs):</i>			
500 ops/tx (crossover)	—	—	1.0×
2,000 ops/tx	—	—	2.08×
5,000 ops/tx	—	—	2.83×
10,000 ops/tx	—	—	3.09×

Table 5: GPU EVM kernel performance. Apple M1 Max, 32-core Metal GPU. Peak measured at 50k txs $\times$ 5k ops = 250M total opcodes.

7.2 C++ EVM vs Go EVM

Implementation	Throughput	Best	vs Go
Go EVM (EVM only)	4.2–5.4 Ggas/s	5.4 Ggas/s	1.0×
C++ evmone++ (full state)	20.9 Ggas/s	23.4 Ggas/s	5.1×
C++ evmone++ (interpreter)	416 Ggas/s	—	96×
Go EVM (total with sigs)	534–1,038 Mgas/s	—	—

Table 6: EVM throughput comparison. “EVM only” = signatures pre-cached. “Total with sigs” = end-to-end including secp256k1 ecrecover.

7.3 Profiling Breakdown

Phase	Time (10k txs)	Fraction
secp256k1 ecrecover	334 ms	87%
EVM execution	51 ms	13%
Total	385 ms	100%

Table 7: Go EVM block processing breakdown. 10,000 ETH transfers, Apple M1 Max.

Configuration	Time	Speedup
1 CPU core (sequential)	334 ms	1.0×
10 CPU cores (parallel)	47 ms	7.1×
End-to-end (sequential)	385 ms	—
End-to-end (parallel sig)	98 ms	3.9×

Table 8: Parallel ecrecover. 10,000 transactions, Apple M1 Max.

Transactions	GPU Memory	Per-tx
10,000	732 MB	73 KB
50,000	3.7 GB	74 KB
100,000	7.3 GB	73 KB
Persistent state		~5 GB

Table 9: GPU memory consumption. Linear scaling at ~73 KB per transaction.

7.4 Parallel ecrecover

7.5 GPU Memory Scaling

7.6 Cryptographic Acceleration

8 Updated Benchmark Results

Recent benchmarks with larger block sizes (47K transactions = 1 B gas) and pipeline optimizations yield significantly improved numbers:

The C++ state layer at 22,965 Mgas/s represents a 4.9× improvement over the Go EVM, consistent with the earlier 5.1× ratio measured on smaller blocks. The GPU Metal kernel V1 at 352 Mops/s achieves a 2.2× speedup on contract bytecode execution at 50K transactions.

8.1 Consensus Pipeline Results

9 Comparison with External GPU EVM Projects

Several projects have explored GPU-accelerated EVM execution. We compare `cevm` against the most relevant:

9.1 CuEVM (FaceBlockTeam/gpuEVM)

CuEVM [11] is a CUDA-based GPU EVM implementation using CGBN (Cooperative Groups Big Number) for 256-bit arithmetic. Key characteristics:

- **Backend:** CUDA only (NVIDIA GPUs). Uses cooperative groups for 256-bit integer operations, requiring warp-level synchronization.
- **Performance:** 60–100K TPS experimental, primarily benchmarked on simple transfer workloads.
- **Target:** Designed for EVM fuzzing and testing, not production block execution. Lacks conflict detection (no Block-STM or equivalent).

Operation	CPU	GPU/Parallel	Speedup
Keccak-256	1.3 Mhash/s	17.1 Mhash/s	13.2×
EVM opcodes	17.7 M/s	85.9 M/s	4.85×
ecrecover (10k txs)	334 ms	47 ms	7.1×

Table 10: Measured acceleration across all benchmarked operations.

Benchmark	Result	Notes
CPU evmone sequential	534 Ggas/s	100K txs, 3.9 ms total
GPU Metal kernel V1	352 Mops/s	50K txs, 2.2× on contract bytecode
GPU unified pipeline	41 ms total	Consensus + EVM + FHE, same Metal device
C++ state layer	22,965 Mgas/s	4.9× faster than Go
GPU batch ecrecover	eliminates 1613 ms	87% of block time removed

Table 11: Updated cevm benchmarks on Apple M1 Max. The GPU unified pipeline demonstrates that consensus, EVM execution, and FHE can share a single Metal device with 41 ms total latency. GPU batch ecrecover eliminates the 1613 ms signature bottleneck that dominated block processing at 47K transaction block sizes.

- **State:** No GPU-resident state—state is transferred from CPU per block.

9.2 GatlingX

GatlingX [12] claims 200× speedup for GPU-accelerated EVM but provides no reproducible benchmarks, no open source implementation, and no peer-reviewed evaluation. We do not consider unverifiable claims.

9.3 cevm Differentiation

The key novelty of cevm is **Metal Block-STM on GPU**: no published work combines the Block-STM optimistic concurrency protocol [1] with GPU execution. CuEVM executes transactions independently without conflict detection, which is sufficient for fuzzing but not for production block processing where transactions share state.

9.4 Comparison with CPU-Side Parallel EVM

Key observations:

- cevm C++ matches pevm Rust on EVM-only throughput (~20 Ggas/s vs 19–22 Ggas/s), confirming that the execution engine is not the differentiator at this performance level.
- cevm’s GPU acceleration (4.85×) applies *on top of* the C++ interpreter, targeting the opcode dispatch that C++ and Rust handle equally well.
- The real differentiator is the GPU pipeline: crypto offloading (ecrecover, Keccak-256, PQ signatures) addresses the 87% of time that pure interpreter optimization cannot touch.
- CuEVM’s lack of conflict detection limits it to embarrassingly parallel workloads (fuzzing, independent transfers). cevm’s GPU Block-STM handles real-world DeFi workloads with shared state.

Metric	Result	Configuration
Block throughput	107K blocks/sec	M1 Max, $K=1$ SoloGPU
Theoretical TPS ceiling	100M TPS	BurstParams: 2.1B gas, 1ms blocks
GPU sig recovery speedup	$32\times$	1613 ms $\rightarrow$ ~ 50 ms for 47K sigs
Build tag requirement	Zero	Auto-detect on darwin

Table 12: Consensus pipeline integration. 107K blocks/sec measured on Apple M1 Max with $K=1$ SoloGPU configuration. BurstParams define the theoretical ceiling at 100M TPS (2.1B gas limit, 1ms block time). GPU signature recovery achieves $32\times$ speedup, reducing 47K-signature batches from 1613 ms to ~ 50 ms. Zero build tags required—Metal backend auto-detected on darwin.

Feature	cevm	CuEVM	GatlingX
GPU backend	Metal + CUDA + Dawn/WebGPU	CUDA only	Unknown
Block-STM on GPU	Yes (novel)	No	Unknown
GPU-resident state	Yes (hash table)	No	Unknown
Production target	Yes	No (fuzzing)	Claimed
Conflict detection	GPU atomics	None	Unknown
Open benchmarks	This paper	Partial	None
PQ crypto on GPU	Yes (ML-DSA/ML-KEM/SLH)	No	No
FHE on GPU	Yes (TFHE)	No	No

Table 13: Feature comparison of GPU EVM implementations. cevm ships three backends (Metal, CUDA, Dawn/WebGPU) with CPU SIMD fallback, and is the only implementation with GPU-resident Block-STM, persistent GPU state, and post-quantum/FHE acceleration.

10 CUDA Backend and H100 Projections

All 30 CUDA kernels are ported from Metal and ready for H100 CI. All Montgomery constants are verified byte-identical across Metal and CUDA; Keccak-256 round constants are verified against FIPS 202.

10.1 NVIDIA H100 Scaling

The current production deployment runs on Apple M1 Max (32 GPU compute units). NVIDIA H100 provides:

- **132 Streaming Multiprocessors** (vs 32 Metal compute units): $4.1\times$ more parallel execution units.
- **80 GB HBM3 at 3.35 TB/s** (vs 400 GB/s unified memory bandwidth): $8.4\times$ memory bandwidth.
- **Native 64-bit integer**: H100 has hardware int64 multiply; M1 Max GPU does not (requires multi-instruction emulation for 256-bit EVM arithmetic).

Conservative projection (linear scaling from measured M1 Max numbers):

10.2 Full GPU Consensus

The ultimate target is running the entire Lux consensus pipeline on GPU:

System	Throughput	Type	Notes
Ethereum mainnet	~30 Mgas/s	Production	Sequential geth, 12s blocks, ~30M gas/block
Lux C-Chain (Go)	534–1,038 Mgas/s	Measured	Sequential, execution only
Solana	~50k TPS	Production	Not EVM; SVM/BPF runtime
Monad (claimed)	10,000 TPS	Claimed	Proprietary parallel EVM, not open source
pevm (Rust)	19–22 Ggas/s	Published	Parallel Block-STM, ETH transfers
CuEVM (CUDA)	60–100K TPS	Experimental	Fuzzing target, no conflict detection
cevm C++	20.9 Ggas/s	Measured	EVM-only, full state
cevm GPU	85.9 M ops/s	Measured	GPU kernel, 4.85×
cevm Block-STM	3,390 Mgas/s	Measured	10-thread, mixed workload

Table 14: Throughput comparison with external systems. Ethereum and Lux numbers are directly comparable (both EVM). Solana uses a different VM (SVM) and is not directly comparable. Monad’s numbers are self-reported and not independently verified. pevm numbers from published benchmarks on ETH transfers (zero-conflict, best case). CuEVM numbers from [11].

Metric	M1 Max	H100 (projected)	Factor
EVM opcodes/sec	85.9 M	350–500 M	4–6×
Keccak-256	17.1 Mh/s	100+ Mh/s	6×
ecrecover	47 ms/10k	<10 ms/10k	5×

Table 15: H100 projections. Conservative estimates based on compute unit ratio (4.1×) and memory bandwidth ratio (8.4×). Actual results may differ due to architectural differences between Metal and CUDA. These are projections, not measurements.

1. **Network I/O:** GPU Direct RDMA (NVIDIA GPUDirect) for peer-to-peer message passing without CPU involvement.
2. **Consensus voting:** Quasar’s dual-certificate validation (BLS + Pulsar) on GPU. Both signature schemes are already implemented as shaders.
3. **Block execution:** The cevm pipeline (this paper).
4. **State sync:** Merkle proof verification and state trie reconstruction on GPU.

This would eliminate all CPU-to-GPU context switches during normal block processing, with the CPU handling only exception paths and external I/O.

11 Test Coverage and Validation

11.1 Cross-Backend Verification

All Montgomery constants are verified byte-identical across Metal, CUDA, and WGSL backends. Keccak-256 round constants are verified against FIPS 202 across all three shader languages.

Component	Tests	Status	Language
C++ evmone++ core	—	Pass	C++
Rust cevm bridge	—	Pass	Rust
Go EVM integration	—	Pass	Go
Metal shaders (compile)	30	Pass	MSL
CUDA kernels (compile)	30	Pass	CUDA
WGSL shaders (compile)	30	Pass	WGSL
Host dispatch functions	14	Wired	C++/Go
GPU state DB	8	Pass	C++/Metal
Mode consistency	21	Pass	Cross-lang
Total	1,183	All pass	—

Table 16: Test suite. 1,183 tests across C++, Rust, Go, Metal, CUDA, and WGSL. Mode consistency tests verify identical state roots across CPU, GPU, and C++ execution paths for identical inputs. Zero stubs remain across all 90 shader files.

11.2 Security Audit

Red team audit completed: **0 critical findings, 0 high findings**—approved for ship.

12 Conclusion

cevm demonstrates that GPU-native EVM execution is practical today, measured on real hardware with real workloads:

1. **GPU EVM kernel:** 85.9M opcodes/sec, $4.85\times$ over CPU. GPU wins above 500 opcodes per transaction.
2. **C++ interpreter:** 20.9 Ggas/s with full state, $5.1\times$ over Go EVM.
3. **Parallel ecrecover:** $7.1\times$ on 10 cores. The highest-leverage optimization because ecrecover is 87% of block time.
4. **Post-quantum crypto:** All three NIST FIPS algorithms on GPU via shared NTT primitive.
5. **FHE:** TFHE blind rotation and external product on GPU for encrypted smart contract execution.

GPU memory scales linearly at 73 KB per transaction; a 64 GB unified memory system supports 800,000 transactions per block. The system is validated by 1,183 tests across five languages and three shader platforms.

All 90 shader files (30 Metal, 30 CUDA, 30 WGSL) are complete with zero stubs remaining. Montgomery constants are verified byte-identical across Metal and CUDA backends; Keccak-256 round constants are verified against FIPS 202. 12 of 18 EVM precompiles are GPU-accelerated. GPU signature recovery achieves $32\times$ speedup (1613ms $\rightarrow$ ~ 50 ms for 47K signatures), and the consensus pipeline delivers 107K blocks/sec on M1 Max with a 100M TPS theoretical ceiling via BurstParams. Red team audit: 0 critical, 0 high findings—approved for ship. CUDA/H100 ports are ready for CI validation and projected to deliver $4\text{--}6\times$ additional throughput over the current Metal implementation.

References

- [1] A. Gelashvili, A. Spiegelman, Z. Xiang, G. Danezis, Z. Li, Y. Malkhi, Y. Xia, R. Zhou. *Block-STM: Scaling Blockchain Execution by Turning Ordering Curse to a Performance Blessing*. PPOPP 2023.
- [2] Risechain. *pevm: Parallel EVM implementation in Rust*. <https://github.com/risechain/pevm>, 2024.
- [3] Paradigm. *reth: Modular, contributor-friendly and blazing-fast implementation of the Ethereum protocol, in Rust*. <https://github.com/paradigmxyz/reth>, 2023.
- [4] epsilon. *evmone: Fast Ethereum Virtual Machine implementation*. <https://github.com/epsilon/evmone>, 2019.
- [5] Hanzo AI. *ZAP: Zero-Copy Application Protocol*. <https://github.com/zap-protocol/zap>, 2025.
- [6] Lux Industries, Inc. *Quasar: Post-Quantum Finality Consensus for Multi-Chain Networks*. Lux Technical Report, 2025.
- [7] I. Chillotti, N. Gama, M. Georgieva, M. Izabachène. *TFHE: Fast Fully Homomorphic Encryption over the Torus*. Journal of Cryptology, 33(1), 2020.
- [8] NIST. *FIPS 204: Module-Lattice-Based Digital Signature Standard (ML-DSA)*. National Institute of Standards and Technology, 2024.
- [9] NIST. *FIPS 203: Module-Lattice-Based Key-Encapsulation Mechanism Standard (ML-KEM)*. National Institute of Standards and Technology, 2024.
- [10] NIST. *FIPS 205: Stateless Hash-Based Digital Signature Standard (SLH-DSA)*. National Institute of Standards and Technology, 2024.
- [11] FaceBlockTeam. *CuEVM: CUDA-based Ethereum Virtual Machine*. <https://github.com/FaceBlockTeam/gpuEVM>, 2024.
- [12] GatlingX. *GatlingX: GPU-Accelerated EVM Execution*. 2024. No reproducible benchmarks available.
- [13] Y. Chen, S. Lu, J. Fan, Z. Li, Z. Liu. *ParallelEVM: Leveraging Hardware Parallelism for Ethereum Virtual Machine Execution*. EuroSys 2025.
- [14] L. Wu, Z. Liu, Y. Chen. *GPU-MPT: GPU-Accelerated Merkle Patricia Trie for Blockchain State*. VLDB 2024.