

TOROIDAL INFORMATION EXECUTION ENGINE

Executive Investor Brief

Dragrush AI | August 2026

Intellectual Property & Core Technology

The Toroidal Information Execution Engine (TIEE) is a patent-pending architecture developed by Dragrush AI. It eradicates computational waste by shifting from rigid, synchronous data processing to a fluid, asynchronous stream-based model.

The Singularity Gateway: An edge-level shield that mathematically annihilates malicious traffic and malformed vectors at absolute zero cost to the core CPU arrays.

The Quantum EV Logic Engine: A highly concurrent microservice layer that evaluates the expected value of every data packet using the wave-collapse scoring equation: $\text{Collapse}(\Psi) = \text{BaseEV} \times \omega_{\text{System}}$

The -6.666 Protocol: A core algorithmic function that instantly drops data yielding a negative execution score before it can consume internal memory or database connections.

Asynchronous Footprint Reduction: By transitioning from heavy processing threads to lightweight Goroutines, the engine shrinks the operational memory footprint from 2 MB down to 2 KB.

High-Yield Enterprise Use Cases

While the engine's asynchronous queuing principles can be fractally scaled down to optimize a local desktop computer, its primary commercial deployment targets high-throughput sectors:

AI Data Pipelines: Feeds massive, continuously cleansed datasets into vector databases and active machine learning models without risking data poisoning or stalling hardware.

High-Frequency FinTech Trading: Evaluates execution utility for millions of concurrent financial transactions with sub-millisecond response times.

Massive IoT Telemetry: Acts as a global shock absorber, capturing and batching continuous data streams from millions of edge sensors without crashing backend relational databases.

FinOps & Hardware ROI

The transition to a toroidal paradigm fundamentally alters the capital efficiency and physical constraints of massive data centers and AI factories.

Metric	Legacy Architecture	Toroidal Execution Engine
Thread Footprint	2 MB per thread	2 KB per Goroutine
System Throughput	10,000 requests/sec	200,000 requests/sec
GPU Utilization	~72% (Jagged)	99.4% (Sustained)
Compute Reclamation	Baseline Operation	36.8% Reclaimed

By mathematically destroying low-value payloads at the network edge, the system prevents "GPU starvation" and ensures that expensive tensor cores are only processing pure signal. This upstream filtration reclaims up to 40% of standard CPU overhead. Consequently, enterprises can immediately reallocate the reclaimed megawatts of power and cooling capacity to run 5% to 8% more GPU nodes within their exact same physical footprint.