

The moat: we patented the hard part.

Patented scheduling

Intelligently routes each task to the fastest, most reliable nodes in real time.
US patent granted, EU pending.

Fault tolerant

Detects node failures and reroutes work automatically, so every job finishes.

Reliable by design

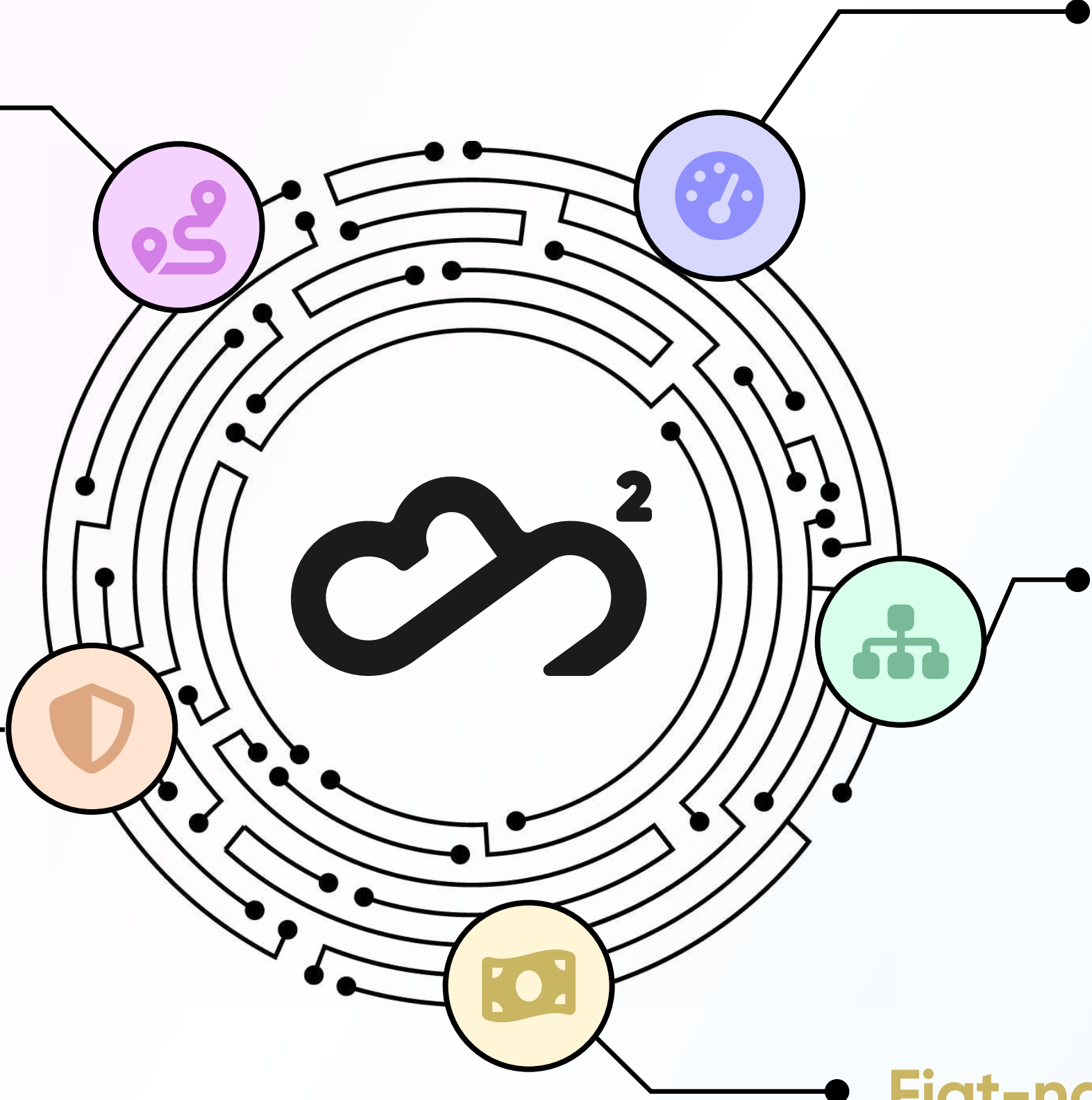
Turns an unreliable, heterogeneous fleet into SLA-backed, predictable compute.

Automatic parallelization

Splits any parallel job across the fleet and reassembles the results, no orchestration code.

Fiat-native

Real-money billing and payouts. Enterprises procure with confidence; workers cash out.



The market

CAGR: 44%

SOM: \$8.7B

Global GPU-as-a-Service (GPU rental) market

Source: Fortune Business Insights, 2026

SAM: \$20.6B

Public cloud end-user spending on AI inference and batch processing services.

Source: Deloitte, 2026

TAM: \$142.8B

Global spending on AI infrastructure

Source: Evolvance, 2026

Reliable like a hyperscaler. Cheap like a marketplace.



-  Low cost
-  Full managed stack
-  Fiat billing
-  QoS / SLA-backed

DePIN GPU



-  Low cost
-  Marketplace only
-  Crypto-native
-  No SLA

Hyperscaler



-  High cost
-  Full managed stack
-  Fiat billing
-  QoS / SLA-backed

Consumer GPU

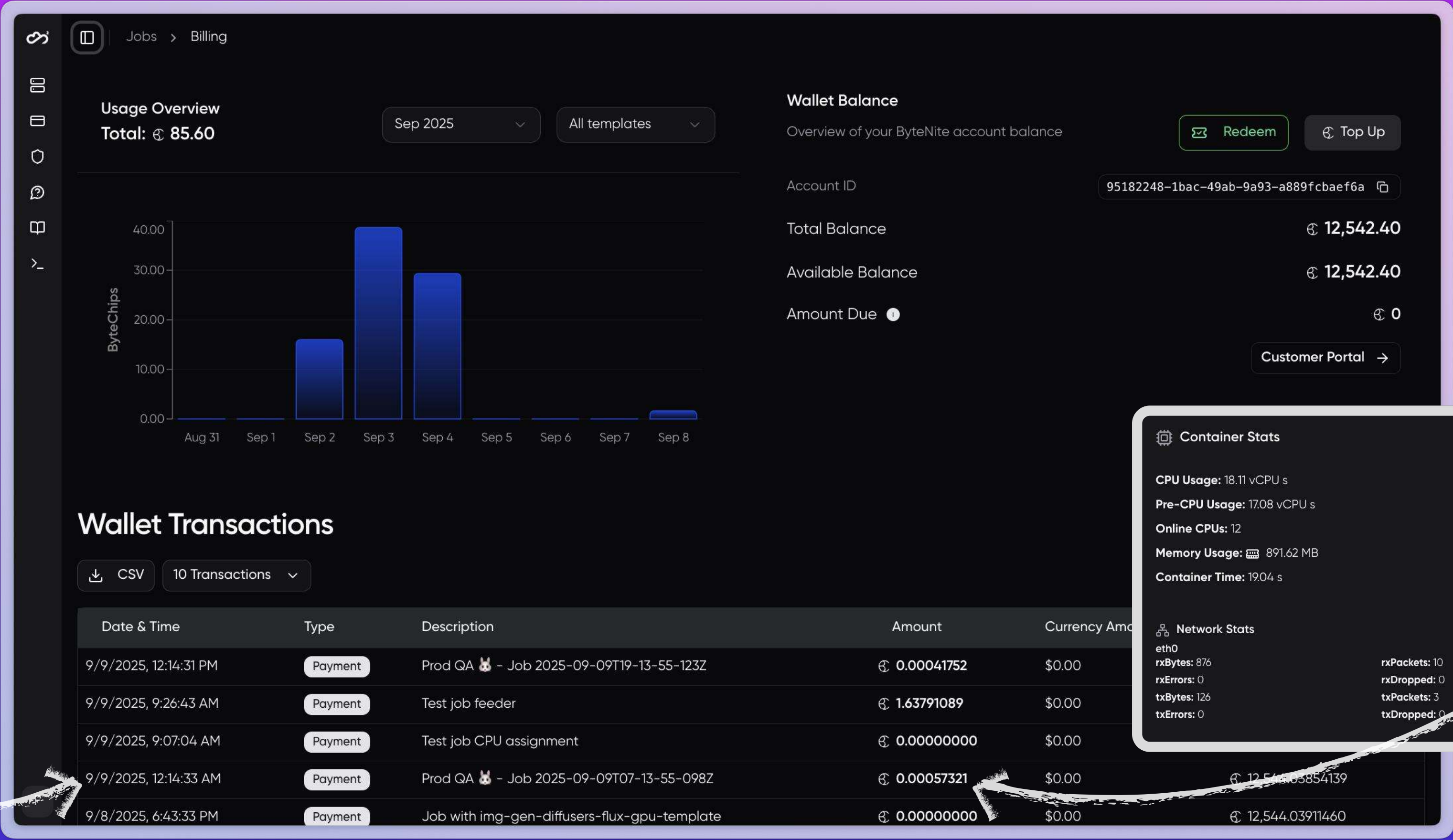


-  Low cost
-  Marketplace only
-  Fiat billing
-  No SLA

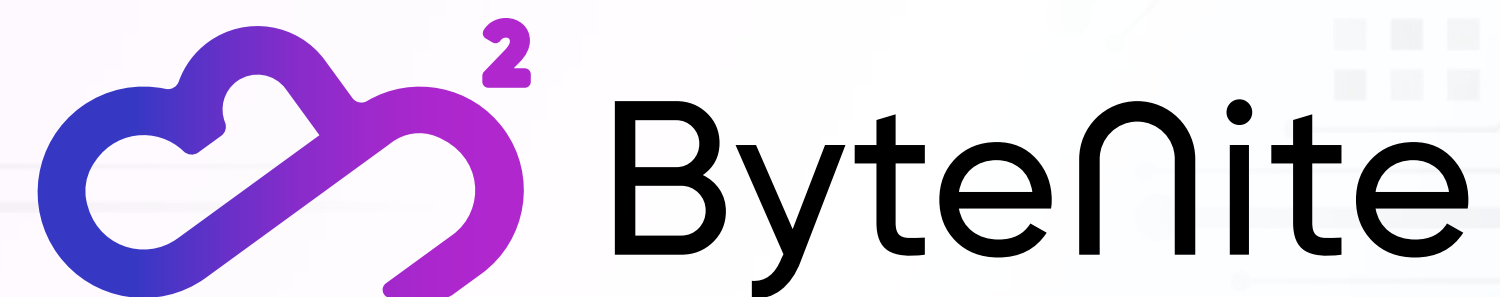
Customers pay per second for the CPU, GPU, RAM, and network they actually use.

Usage is tracked and itemized in real time, so customers can monitor spend and plan their budget.

Every job is a transaction: suppliers earn a real-money payout, and ByteNite keeps a margin.



Thank You



The distributed cloud for AI

LinkedIn



linkedin.com/company/bytenite

Website

www.bytenite.com

Documentation



docs.bytenite.com