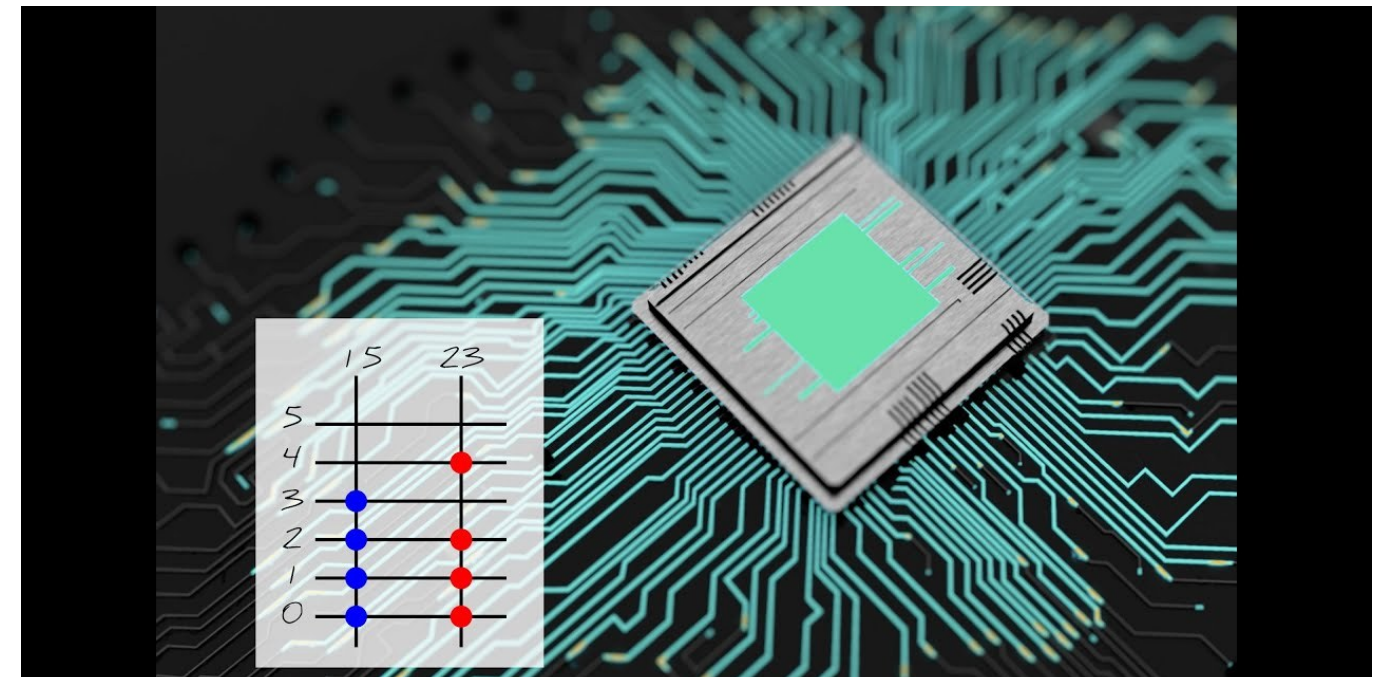


ARRAY



INVESTOR INVITATION

*A Compute-In-Memory Arithmetic core with
+85% efficiency and throughput gains
using standard cmos technology*

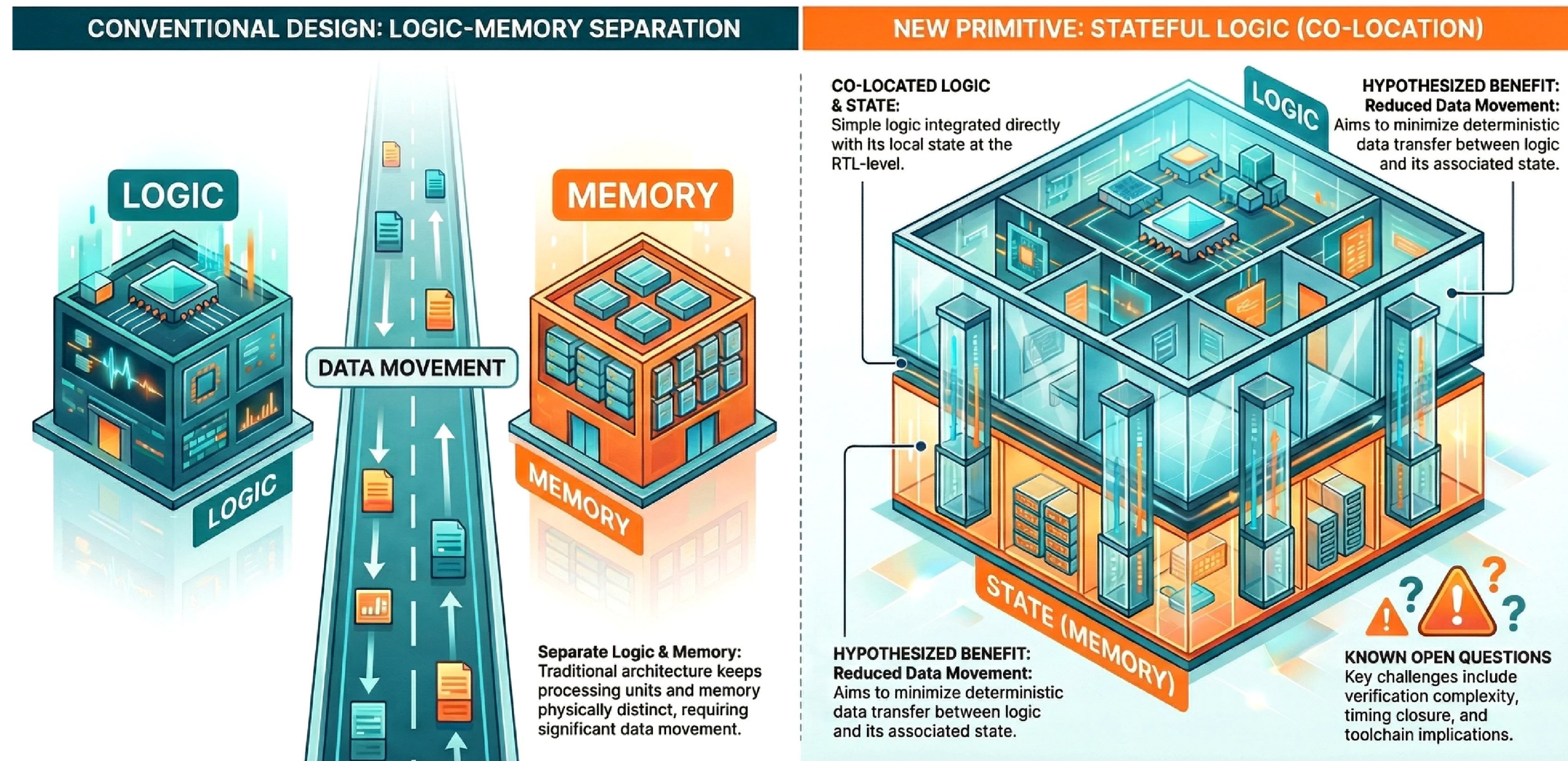


Revolutionizing Processor Arithmetic:

Next-Generation Computing Architectures

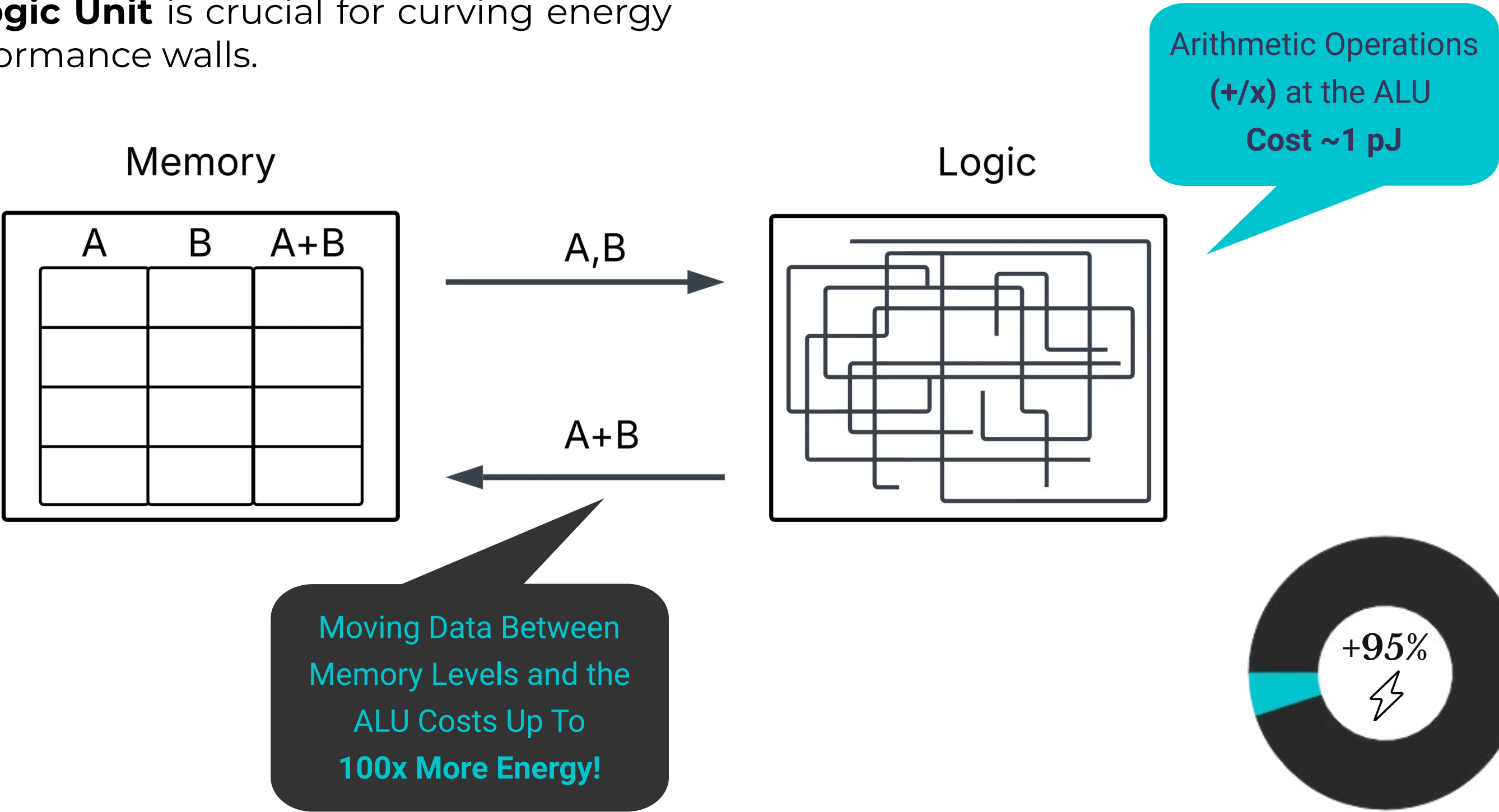
Our Compute-In-Memory architecture promises faster and more efficient arithmetic cores tailored for unparalleled processor efficiency and throughput by cutting out data transfer at the logical level. We solve the memory-wall problem at the mathematical level, not the physical (material) level, **drastically cutting energy and boosting performance**, to deliver next-generation processor cores that can achieve +85% increase in efficiency and performance using standard CMOS technology.

Stateful Logic: A New Architectural Primitive



Problem: Von Neumann Bottleneck

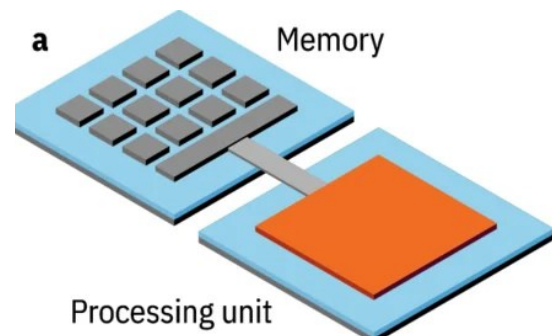
Solving the processors **bottleneck between Memory and Arithmetic Logic Unit** is crucial for curbing energy demands and performance walls.



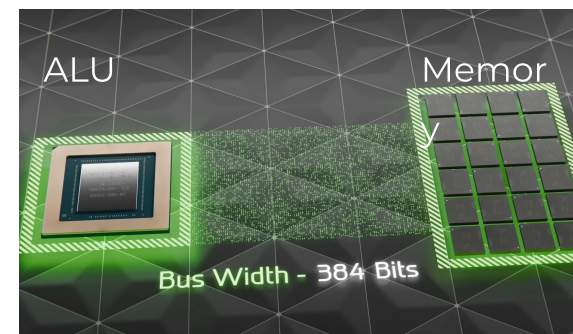
Causes for Von Neumann Bottleneck

The **Von Neumann Bottleneck** makes up for **60-90% of the time latency and energy consumption** in modern processors, due to four main reasons:

The **peripheral distance travelled between units is huge** compared to the distances travelled inside the subunits.



We have a **bandwidth problem** when trying to increase bit capacity too much.



Energy costs are disproportionate.

~1 pJ to Add/Multiply
vs.
~100 pJ to Move Data



A single operation requires information to migrate between the memory and the ALU **back and forth several times.**



C o m p u t e - I n - M e m o r y

- Transformative improvements are achieved by performing computations directly in the memory. **Up to 90% Time & Energy efficiency in some processors.**
- CIM has been implemented in some **niche applications**, but requires **expensive R&D** (often this means new materials, new machines, new processes, new generation fabs, etc.) and/or **complicated designs.**



Research Budgets into CIM Over the Last Decade (Estimates Based on Public Information)

IBM

\$400 Million USD

intel

\$1 Billion USD

SAMSUNG

\$500 Million USD

Traction

2020

Theoretical background.
Mathematical Proofs.

2024

- IEEE International Conference, Publication, and Peer-Review
- Independent IP Valuation at \$8 M USD
- Third-Party Market Report

2026

Completed VHDL
Simulations.
Ready for FPGA
Emulations and
Physical Design.

2022

International (PCT) patent
application filed in key
global jurisdictions.

2025

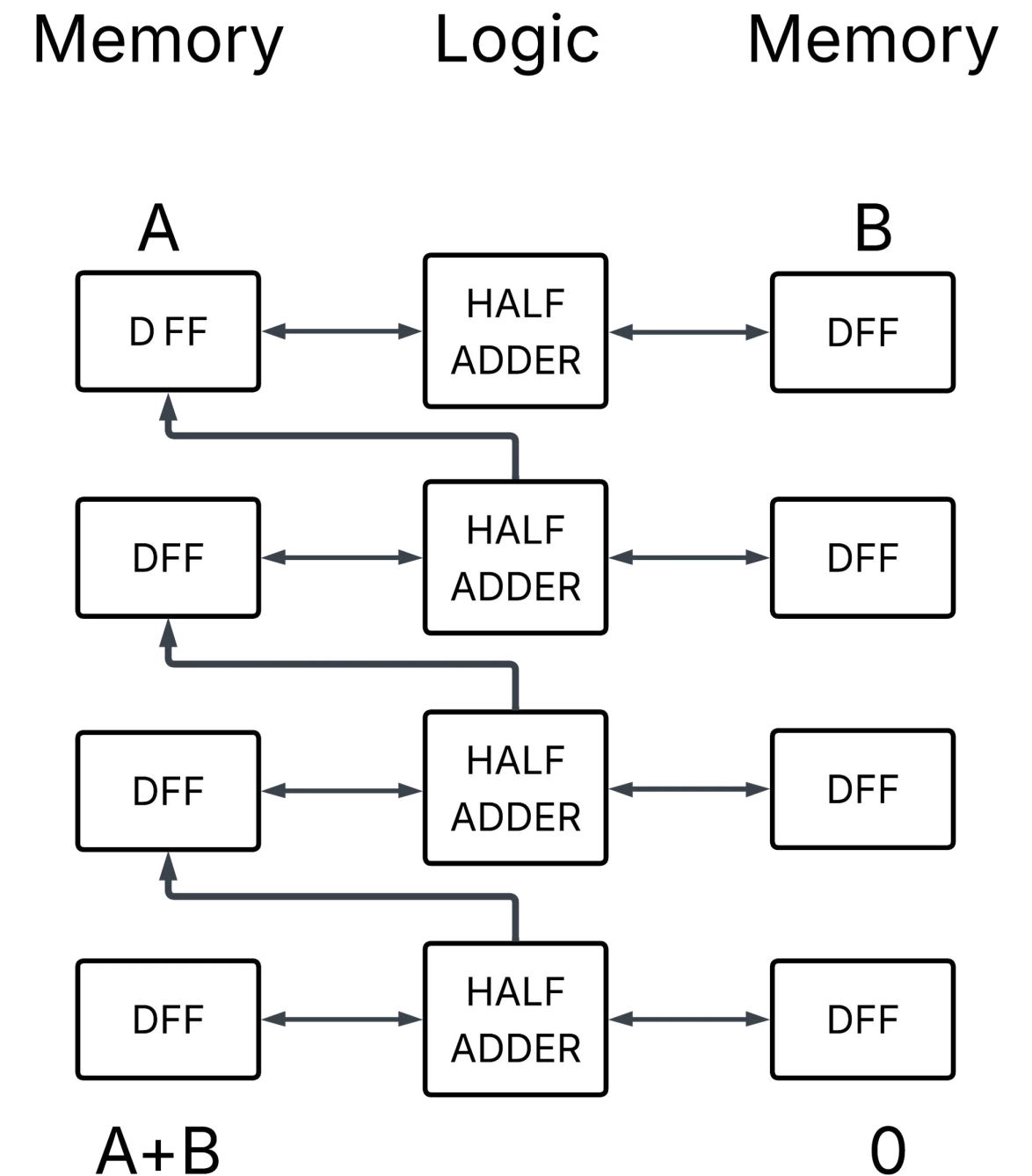
Logical Synthesis and Analysis for In-Memory Matrix
Multiplication (mathematical operation behind AI
and Computer Graphics) and Bitcoin Mining Engine.

O u r T e c h n o l o g y

We embed the arithmetic logic directly into On-Chip Flip Flop arrays, dramatically reducing data movement between memory and arithmetic cores. The innovation is in the mathematical algorithm which redefines addition, and that can be executed with a regular grid of ordinary components (DFFs+AND/XOR gates). **Our SLFA is mathematically more efficient than other adders.**

This architecture **lowers energy per operation while increasing throughput**—particularly in hashing-dominant workloads (mining), vector/matrix operations (AI, graphics), and signal processing. It is an architectural shift that meaningfully expands the performance-per-watt achievable in ASIC designs.

- Low-Power, Low-Cost Design
- Linearly Scalable Bit Capacity
- In-Situ Processing (Super-Fast and Energy-Efficient)
- Based on Standard CMOS Memory and Transistor Technology



Comparison Table

	Traditional Adder Architecture		Non Von Neumann Architectures			
	Ripple Carry-Over (1st Gen)	Parallel (Carry-Look Ahead, Carry-Save, Kogge-Stone, etc.)	Near-Memory Computing	Compute-In-Memory Transistor Arrays (FeFETs, MRAM, ReRAM, etc.)	ASIC (On-Chip Pipelines)	Simple and Linear Fast Adder (Bit-Level Embedded Logic & Memory)
Competitors	Low-Power Systems Texas Instruments	Traditional Processors Intel, Nvidia, IBM, AMD, etc.	NeuroBlade, Nvidia, MythicAI	IBM, Intel, Samsung, SK Hynix, MythicAI	BitMain, Etched, Nvidia, Cerebras, NeuroBlade	Our IP
Energy (Efficiency)	◐	○	●	●	●	●
Time (Latency)	○	◐	●	●	●	●
Manufacturing (Costs and Requirements)	●	◐	◐	○	◐	●
Technology (R&D)	●	●	◐	○	◐	●
Applicability (Use Cases)	◐	●	◐	◐	◐	●
Scalability (Bit Length)	●	◐	◐	◐	◐	●

ARRAY

One Patented Architecture. Multiple Industries. Multiple Revenue Streams.

1

ARITHMETIC CORE

Our fundamental asset

ARRAY architecture

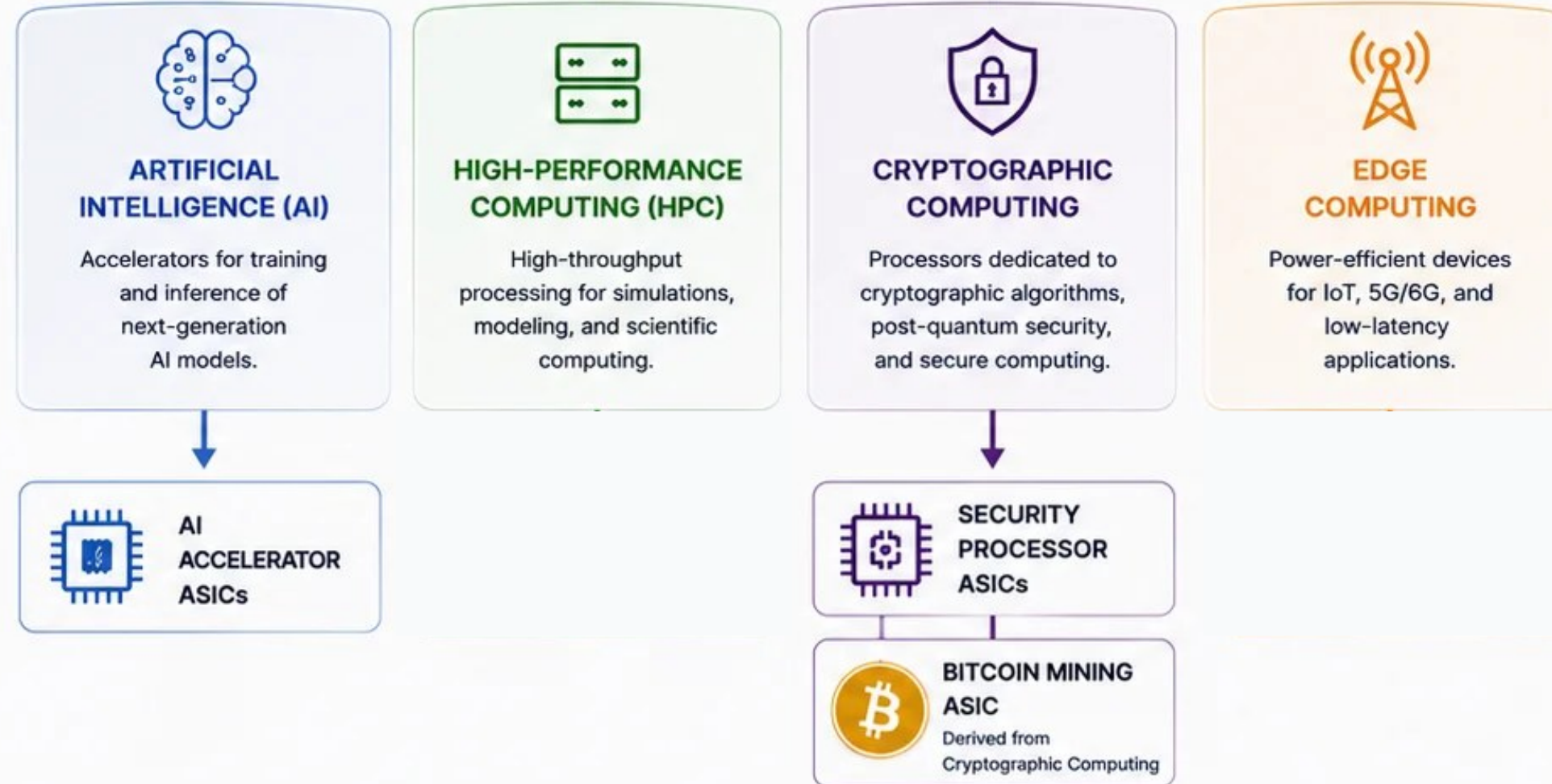


- Patented Arithmetic Architecture
- Superior Efficiency & Performance
- Lower Power, Higher Throughput
- Scalable Across Multiple Workloads

2

INDUSTRIES WE ENABLE

A licensable architecture across high-impact markets



3

FIRST COMMERCIAL IMPLEMENTATION

Products derived from our architecture

OUR COMMERCIALIZATION MODEL & REVENUE STREAMS



UPFRONT LICENSE FEES
One-time license fees for the right to use our patented architecture.



RECURRING ROYALTIES
Ongoing royalties based on the volume of chips or systems shipped.



MANUFACTURING PARTICIPATION
Strategic participation in manufacturing yields (1–3% in-kind allocation of ASICs).



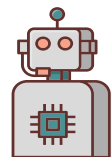
STRATEGIC MINING OPERATIONS
Selective deployment of proprietary mining infrastructure to capture additional value.

Total Addressable Market

Artificial Intelligence

\$200BUSD/
year in MP

Matrix multiplication is the core operation of neural networks. It directly affects the training speed and scalability of AI models.



Graphics and Gaming

\$100BUSD/
year in MP

Transformations like rotation, scaling, and translation in 3D graphics rely on matrix operations.



Blockchain and Bitcoin Mining

\$60BUSD/
year in MP

Matrix multiplication is fundamental to many cryptographic algorithms used for secure key exchange, encryption, and decryption.



Annual Spending on Microprocessors:

\$540 Billion USD

Digital Signal Processing

\$40BUSD/
year in MP

DSP for audio, image, and video data relies on matrix multiplication for tasks like imaging, filtering, coordinate transformations, and data compression.



Data Analysis and Big Data

\$60BUSD/
year in MP

Principal Component Analysis (PCA) and ML models leverage matrix multiplication to analyze correlations and patterns in massive datasets.



Cryptography and Security

\$30BUSD/
year in MP

Matrix multiplication is fundamental to cryptographic algorithms used for securing sensitive data, especially in real-time applications.



There is a wide range of industries to which our IP is applicable, but we believe the Bitcoin Mining and AI industries are our first go-to strategy because of the huge benefits our IP will have in mining architectures and ROI projections.

Business Model

Our business model is centered on licensing our patented processor architecture to semiconductor companies and ASIC developers through upfront licensing fees, and royalties.

Bitcoin mining and AI Accelerators represent our first commercial applications because of the market opportunities.

Our architecture addresses a foundational inefficiency in the computing stack by **reducing data movement**, unlocking significant energy savings across data centers and compute-intensive applications.

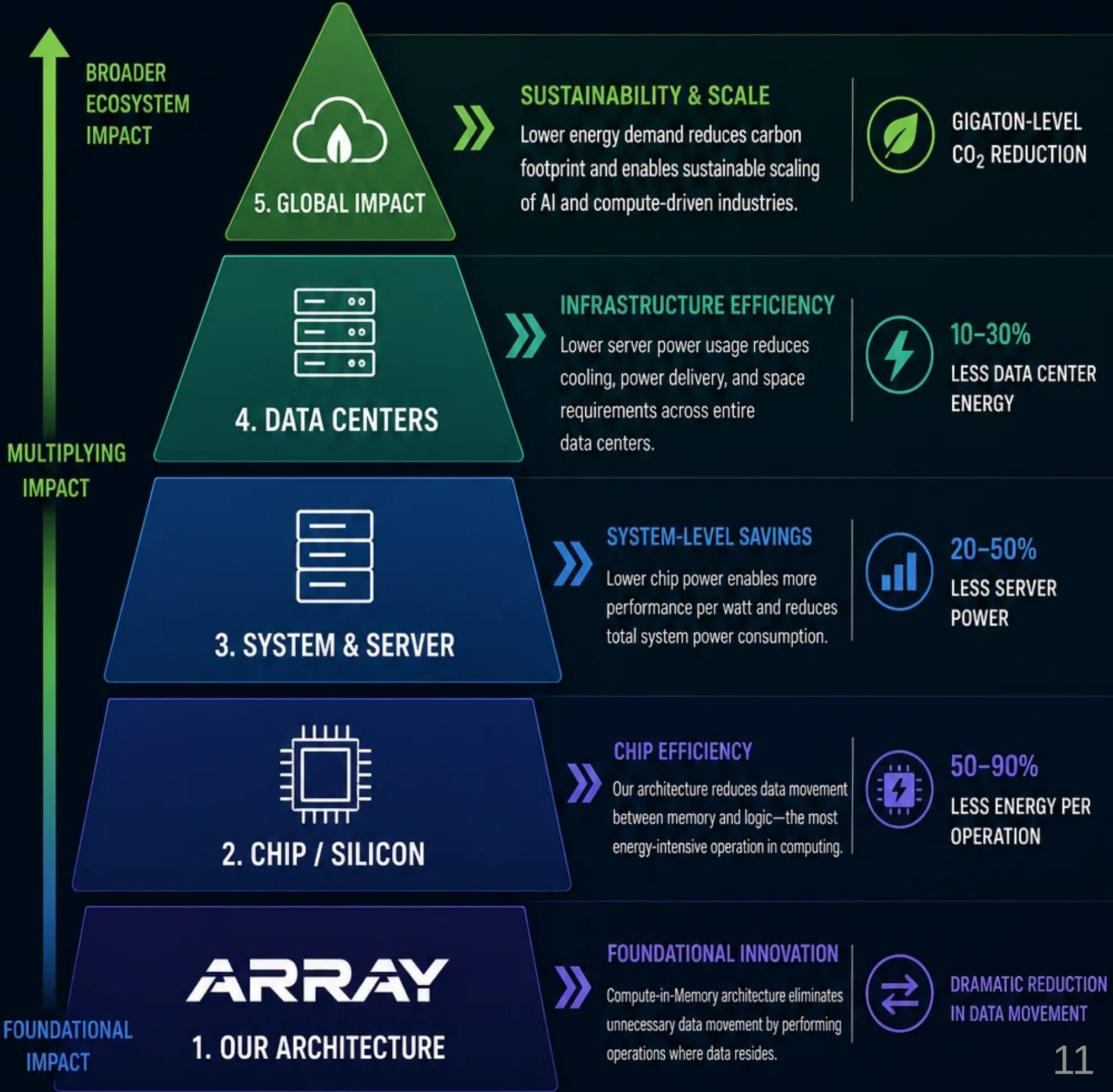
 **THE COMPOUNDING BENEFIT**

Small savings at the silicon level compound into massive impact at global scale.

=

 **TRANSFORMING THE ENERGY ECONOMICS OF COMPUTING**

LESS DATA MOVEMENT. **LOWER ENERGY.** GREATER IMPACT.



P r o c e s s t o M a r k e t



**Revenue from Commercial License Agreements
(6 Months On)**



L e t ' s P a r t n e r

We are launching a dedicated venture to commercialize our next-generation arithmetic core in Bitcoin Mining Engines and AI Accelerators, establishing a new efficiency standard and creating a strategic partnerships.

- Our core IP has received a **Third-Party valuation by Intra-Com Group GmbH** (Germany-based IP Firm) of **\$8 Million USD**. This valuation will increase during 2026 as key jurisdictions (China, Canada, India, Japan, S. Korea, Singapore, UK, USA) grant the international application.
- We are seeking **\$6.6 Million USD** of investment to develop a licensable asset and co-design the next-generation of high performance processors.
 - **\$500k Pre-Seed (6 Months)**
 - Emulate and benchmark the arithmetic core on FPGAs.
 - Fully de-risk the technology with performance and efficiency data.
 - **\$6.1 Million Seed (10 Months)**
 - Deliver a complete, licensable, Compute-In-Memory Engine
 - Ready for integration into commercial ASIC designs: AI, Mining, Cryptography, etc.

THANK YOU!



www.arrayarchitecture.com



Juan Pablo Ramirez

Architect, Founder & CEO

juan.ramirez@arrayarchitecture.com



Pablo César Vázquez Estrella

Co-Founder & Chief Financial Officer

pablo.vazquez@arrayarchitecture.com

www.arrayarchitecture.com

[ARRAY ARCHITECTURE, CORP.]