

ULTRARAM

ULTRA-EFFICIENT MEMORY

AI-Native Memory Architecture

Enabling Energy-Efficient AI Infrastructure



James Ashforth-Pook

Co-founder & CEO - Quinas Technology Ltd.

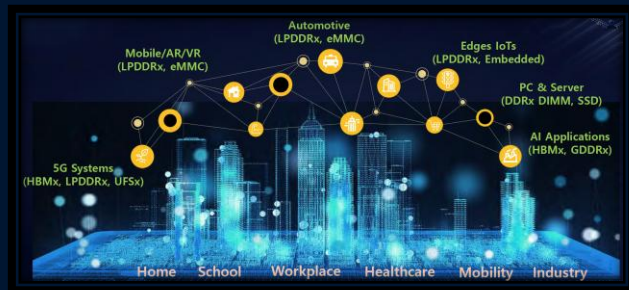
james@quinas.tech

www.quinas.tech

Explosive Data Growth Is Breaking the Memory Stack



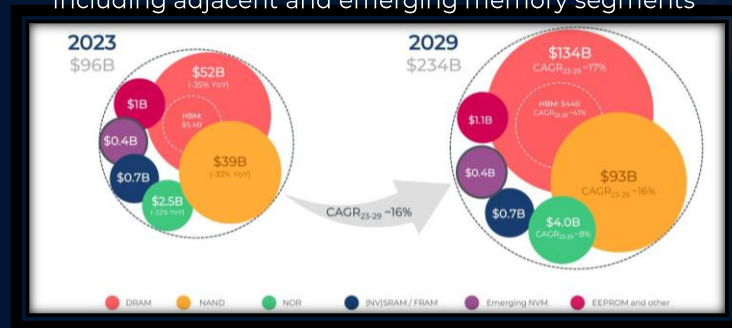
Explosive growth ahead:
AI to drive \$1T market by 2030



Memory powers every sector –
from edge to cloud

A \$234B+ Global Memory Market (and Growing)

Including adjacent and emerging memory segments



Representative memory market leaders

Source: Yole (public market estimates)



South Korea



USA



Taiwan



Japan

AI scaling is now constrained by memory energy - not compute

ULTRARAM™ is uniquely positioned to meet this need

Memory Scaling Has Structurally Failed

Why AI infrastructure urgently needs a new memory architecture

Digital Energy Pressure

- Data centers consume 2–3% of global electricity
Projected to double by 2030
- DRAM's constant power draw is a core inefficiency
- AI model training and inference are driving a step-change in memory-related energy demand

Industry at a Scaling Wall

- Moore's Law and Denard Scaling are breaking down
- DRAM: fast but volatile and power-hungry
- NAND: non-volatile but slow, with limited endurance

Performance vs. Power Trade-Off

AI, 5G/6G, and Edge workloads need memory that is:

- Fast, Non-volatile, and Energy-efficient
- No current memory tech achieves all three

Market Gap = Investor Opportunity

- \$250B+ memory market is ripe for disruption
- Legacy memory can't meet sustainability goals
- A new memory class is needed to unlock growth

ULTRARAM™ is designed to break the speed–power–volatility deadlock and redefine the future of memory.

ULTRARAM™: A Physics-Level Reset in Memory



1980

TOSHIBA



1966



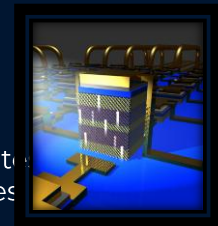
What is the ULTRARAM™ Breakthrough

- A revolutionary universal memory combining:
 - Non-volatility of flash.
 - Speed and endurance beyond DRAM.



Core Technology:

- Compound semiconductors replace silicon for enhanced performance. (InAs/GaSb-based quantum wells)
- Triple-barrier resonant tunneling (TBRT) eliminates refresh cycles, wear issues and drastically reduces power.



Flash	DRAM
✗ High voltage switching (<20 V) ✗ Intrinsically slow P/E (10 μs) ✗ Low endurance (10^5) ✓ Non-volatile ✓ Non-destructive read ✓ Highly scalable	✓ Low voltage/energy switching (<2V) ✓ Fast operation (10 ns) ✓ High endurance (10^{16}) ✗ Volatile ✗ Destructive read ✗ Scaling challenges

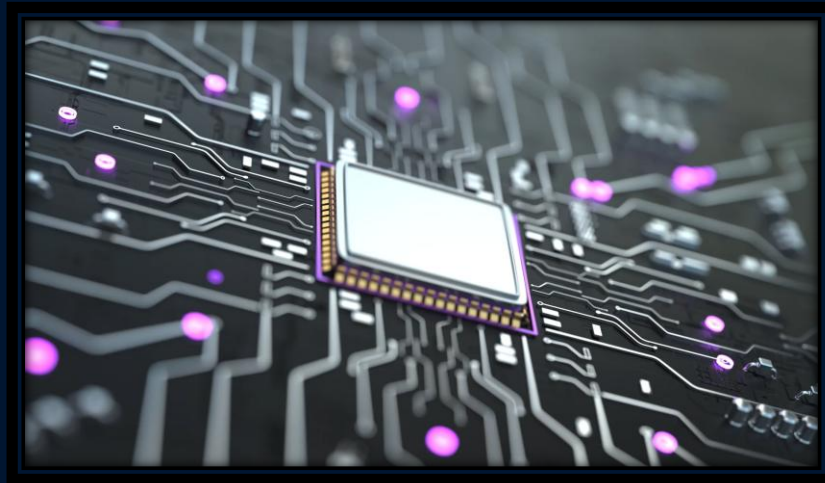
- Flash is non-volatile, but slow and wears out...
- DRAM is fast, but volatile and inefficient due to constant refreshing

ULTRARAM™ is not a hybrid — it is a new device enabled by different physics

ULTRARAM™ delivers DRAM-class speed with non-volatile persistence — enabled by novel device physics.

The Physics Behind AI-Native Memory

Persistent, high-bandwidth, ultra-low power memory enabled by novel quantum resonant tunnelling physics



<https://bit.ly/ultraram>

ULTRARAM
ULTRA-EFFICIENT MEMORY

Why the Industry Wants a Universal Memory

Simplified – Energy Efficient - Scalable



Collapse the Memory Hierarchy

One memory layer reducing reliance on SRAM, DRAM, Flash and HDD



Future Architectures

Enables in-memory, neuromorphic, and chiplet system designs



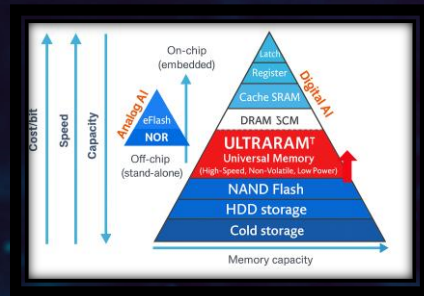
Performance & Energy

Low-latency, high-speed memory with dramatically lower energy per operation



Security

Enables secure architectures via controlled persistence, rapid zeroisation, and tamper-aware operation for trusted/encrypted compute..



Universal Memory enables a new class of fast, persistent system memory



Speed + Non-Volatility

Fast, persistent system memory enabling new compute models



Sustainability

Lower energy consumption and reduced system waste

ULTRARAM™ Is Already De-Risked

Patented Physics Breakthroughs:

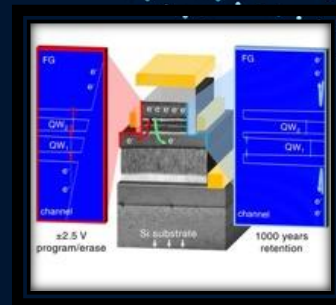
- Anchored by our proprietary Triple Barrier Quantum Resonant Tunneling (**TBRT**) technology, enabling step-change improvements in energy efficiency and performance.
- Protected by multiple global patent families (5 awarded and 8 pending patents in 4 patent families, with a fifth disclosure pending) across the US, UK, Europe, Japan, South Korea, China.
- More than 10 years of R&D

Advanced Materials and Design:

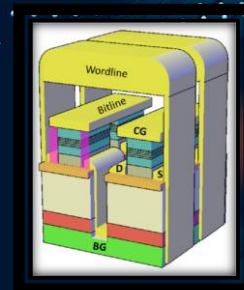
- Incorporates compound semiconductors like **InAs**, **GaSb**, and **AlSb** for unmatched endurance and efficiency.
- Scalable to silicon substrates, accelerating adoption.

Externally validated innovation:

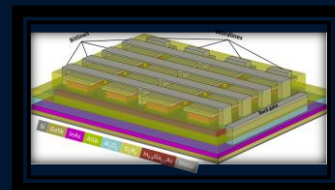
- **Flash Memory Summit (FMS) 2023 USA** - Most Innovative Startup
- **IC Taiwan Grand Challenge 2024** - Global Winner
- **WIPO Global Awards 2025 Geneva** - Spinout Intellectual Property



Memory Device Physics

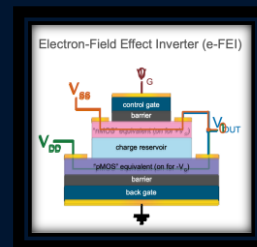


Scalable Bit Cell



Memory Array

16-bit RAM array design @ 4F2



Logic Device

Step-Change Performance at the Device Level

Performance Highlights:

- **Ultra-Low Switching Energy:** ~10 attojoules – per operation
 - (orders of magnitude lower than DRAM and NAND).
- **High Endurance:** Projected endurance $>>10^{15}$ cycles,
 - enabled by non-destructive, physics-based switching.
- **High Speed:** Sub-nanosecond to nanosecond-class switching,
 - competitive with DRAM.
- **Non-Volatility Data retention:** >10 years at elevated temperature,
 - with projected lifetime $>1,000$ years.
- **Wide Temperature Range:** Operational capability
 - from the lowest cryogenic temperatures up to 100°C .
- **Radiation Hardness:** Inherently radiation-tolerant device physics
 - for harsh environments (Testing TBA)

Parameter	DRAM	3D NAND	MRAM	ReRAM	PCRAM	ULTRARAM™	
Storage time	~60 ms	>10 years	>10 years	>10 years	>10 years	>10 years	●
Non-destructive read?	No	Yes	Yes	Yes	Yes	Yes	●
Switching energy	1 fJ	~10 fJ	~100 fJ	~1,000 fJ	>1 pJ	~10 aJ ^(*)	●
Switching voltage	<1 V	>10 V	<1.5 V	<3 V	1-3 V	≤ 2.5 V	●
Energy barrier	0.5 eV	1.6 eV	1.5 eV	1.4 eV	2.4 eV	2.1 eV	●
Cell size	$6F^2$	$\ll 4F^2$	$6F^2$	(4-12) F^2	(4-30)	$\leq 6F^2$	●
Switching time	10 ns	>10 ms	10-50 ns	10-100 ns	10-100 ns	~1 ns ^(*)	●
Endurance	10^6	10^5	10^{12}	Up to 10^{12}	Up to 10^{12}	$>10^{15}$ ^(*)	●

Energy Efficiency: Significantly reduced power consumption for memory-intensive and AI workloads.



NV $>1k$ yrs



Super
Fast



High
endurance



Low
disturb



Ultra-
efficient

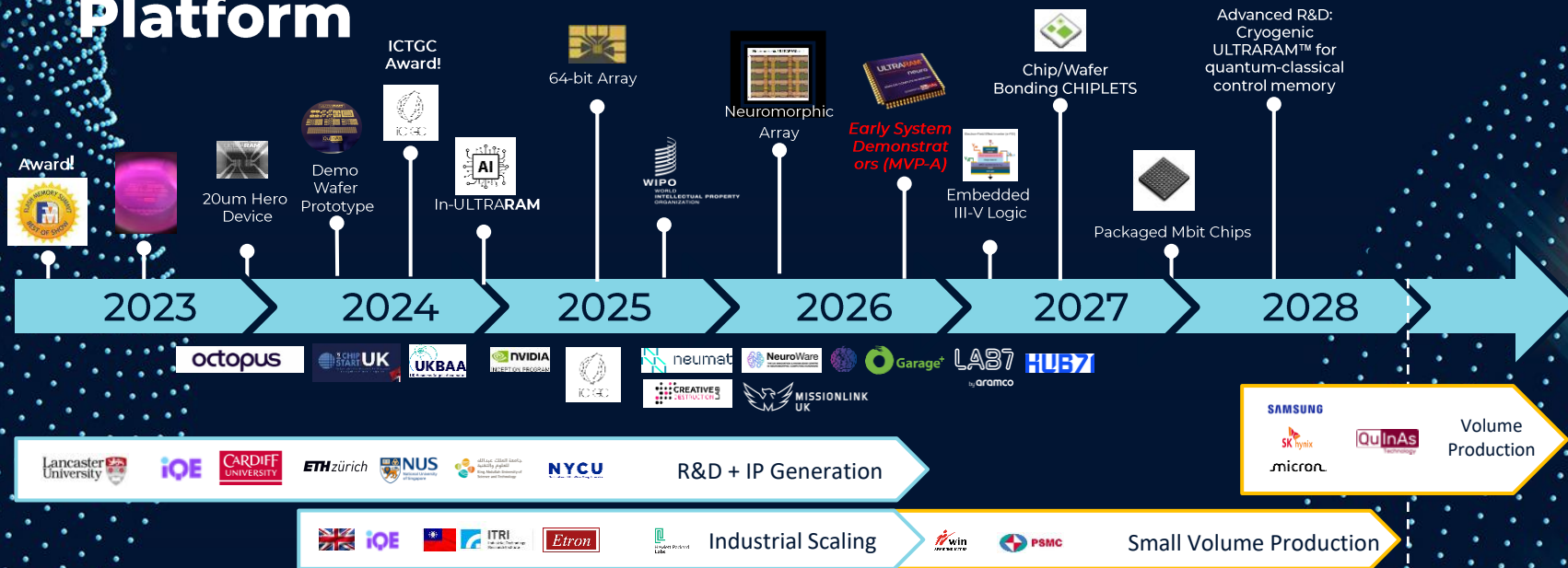


Wide temp
range



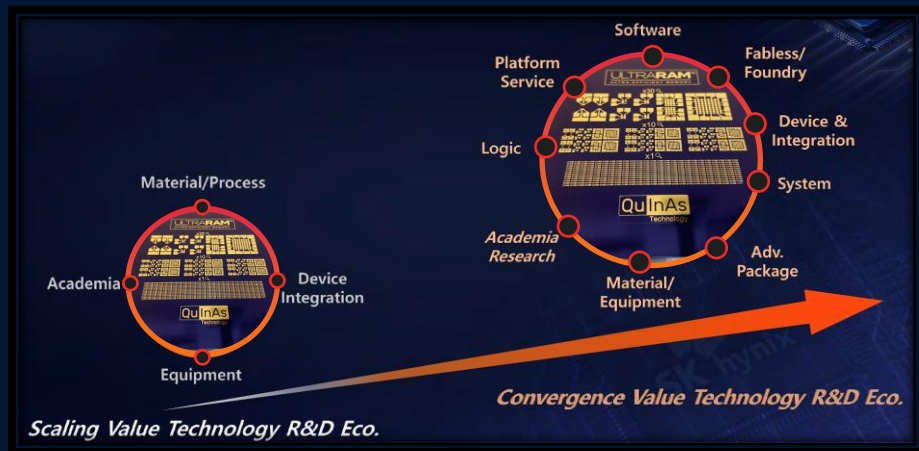
RadHard

From Breakthrough Device to AI-Native Memory Platform



- **2026–2027: Platform Validation & Initial Commercialisation**
Pilot-scale arrays, MVP-A system demonstrators, early adopter deployments, and small-volume production.
- **2028+: Platform Scaling & Market Expansion**
Packaged memory products, advanced chiplet integration, and volume adoption across telecoms, datacentres, and IoT.

Execution Through a Convergent Semiconductor Ecosystem



- End-to-end coverage: materials, epitaxy, devices, packaging, and system integration
- Parallel progress: physics, manufacturing, and applications developed concurrently
- De-risked scale-up: early alignment with industrial fabrication and packaging partners

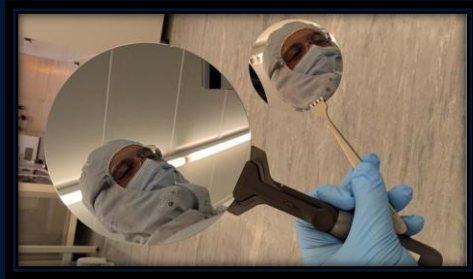
ULTRARAM™ development is embedded across materials, devices, systems, and manufacturing — reducing execution and scale-up risk.



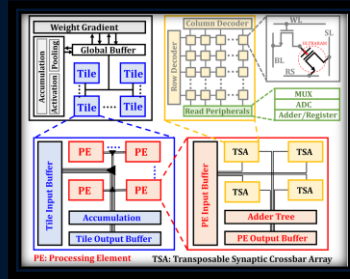
Full-Stack Partnerships for ULTRARAM™ Scale-Up

Covers the full stack from materials → AI System Validation → manufacturing

Materials & Epitaxy



AI System Validation



Manufacturing Readiness



Industrial Epitaxy & Materials Platform

- Transition from MBE → industrial MOCVD (IQE)
- ULTRARAM™ demonstrated on 6" GaSb wafers
- Compatible with III-V manufacturing workflows

Device & System-Level Validation

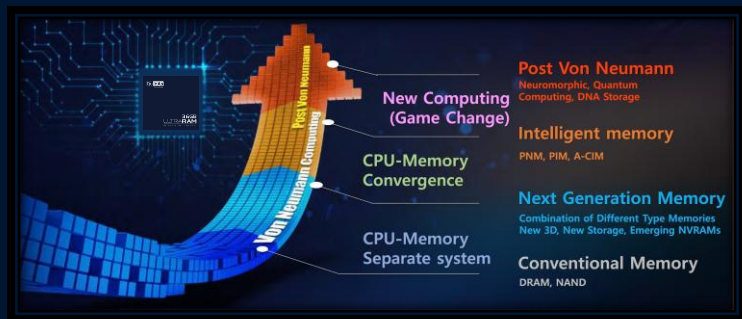
- Physics-based ULTRARAM™ models (IIT Roorkee)
- Integrated into Synopsys design flows
- System-level validation via NeuroSim (Georgia Tech)

Fabrication Readiness

- Atomic Layer Etching (ALE) quantum-well precision
- Oxford Instruments + LayTec process control
- Enables repeatable, scalable fabrication

De-risks ULTRARAM™ from lab-scale innovation to foundry-scale semiconductor production

ULTRARAM Is Emerging as a Strategic Memory Platform for AI Infrastructure



- The memory wall has shifted from bandwidth to energy, latency, and data movement.
- HBM solves bandwidth but exposes system-level power and cost limits.
- Persistent, compute-adjacent memory is required to collapse the memory hierarchy and reduce data movement.

Strategic Markets & Commercial Engagement

AI Infrastructure / AI Accelerators

- Microsoft | Meta
- NAVER | Rebellions | AxeleraAI | Tenstorrent | Fractile

Agentic & Edge AI

- €12M EU Agentic Embodied Edge AI Consortium
- Persistent memory for robotics and embodied AI

Memory Ecosystem

- Kioxia | SK Hynix | Taiwan semiconductor ecosystem

Defence & Secure Compute

- Raytheon | BAE Systems | Leonardo | Northrop
- NATO DIANA ecosystem

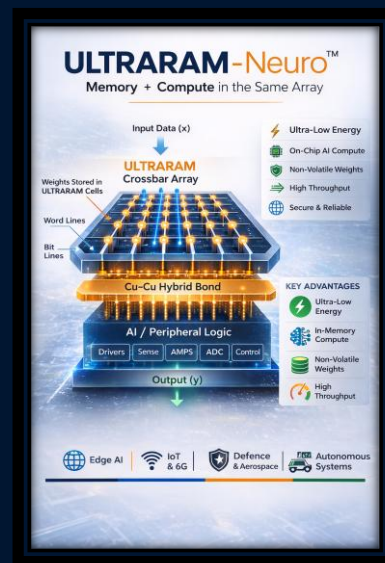
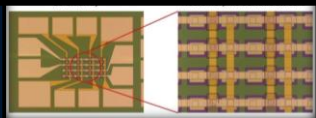
ULTRARAM™ is attracting early interest across AI infrastructure, accelerators, robotics and defence applications..

First Physics-Validated ULTRARAM™ Compute-in-Memory Demonstration MVP-A

MVP-A Validates:

- * ULTRARAM™ device physics integrated into compute architecture
- * Analog MAC operations using multi-level memory states
- * System-level performance via NeuroSim (Georgia Tech)
- * Measurable energy & density advantages vs CMOS systems

Literature	Die Layout	Type	Efficiency (TOPS/W)	Density (TOPS/mm ²)	Array Size	ADC Type
BM LABS		RRAM-CMOS	30 (8 levels)	0.29	128X128X2	TDC
ULTRARAM		Floating Gate	80 (32 Levels)	~0.29 (6F2)	Development	TDC
NVIDIA H100		Full-CMOS	9.36	0.04 (400F2)	-	-
Google TPU v1		Full-CMOS	2.3	0.06	-	SAR



De-risks ULTRARAM™ as a viable compute-in-memory platform for AI hardware..

Outperforming DDR5 in Power, Performance & Potential – MVP-D

System-level persistent memory that reduces HBM power, capacity, and cost pressure in AI data centres

Why ULTRARAM™ Stands Out:

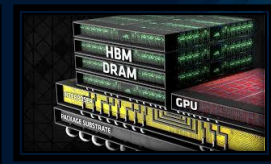
Combines DRAM-class speed with non-volatile persistence to enable HBM-efficient system architectures.

Sustainable: Eliminates refresh and reduces memory energy per inference and training step — directly addressing AI data-centre power and cooling limits. Reducing AI infrastructure energy intensity at system level.

Scalable: Designed for seamless integration into existing global memory manufacturing and packaging systems.

Reliable: Long retention, high endurance, and radiation resistance for mission-critical use cases.

Future Proof: Positioned for future computing paradigms, including AI, neuromorphic and the secure quantum-era systems.



Validated by ETH Zurich SAFARI Group

Modelling predicts disruptive gains in efficiency and scalability.

Early Adopters, Beachhead & Expansion Markets

PRIMARY BEACHHEAD — MVP-A First, MVP-D Follow-on

MVP-A — AI & Edge Computing (Neuromorphic / In-Memory Compute)

- Collapses memory and compute for AI-native architectures
- Exploits speed, non-volatility, and ultra-low energy.
- First MVP to validate the platform and unlock scale



MVP-D — Digital Persistent Memory (DRAM-Class)

- Addresses digital memory systems requiring speed + persistence
- Natural progression toward DRAM-adjacent systems
- Leverages the same ULTRARAM™ device platform

MVP-A is the first commercial platform, enabled by the same ULTRARAM™ device foundation that later scales into system memory

MVP-D



EARLY ADOPTERS — Validation & De-Risking

Space, Defence & Quantum

- Radiation-tolerant, non-volatile memory for extreme environments
- Early adoption driven by physics advantages, not cost
- Initial deployments, qualification, and market credibility

EXPANSION MARKETS — Post-MVP-A Scale

- Data Centres — Persistent system memory and energy savings
- Telecoms (5G/6G) — Low-latency, energy-efficient infrastructure
- Automotive / Industrial / IoT — Reliability and endurance and temperature resilience



Team

Founder team



James Ashforth-Pook
CEO

38+ years in semiconductors and high-growth global startups / scale-ups. Commercialisation & global partnerships



Prof Manus Hayne

CSO
inventor of ULTRARAM™ (device physics and architecture) 35+ years experience.



Dr Peter Hodgson
CTO

Expertise in semiconductor materials with 15+ years experience in III-V devices process integration & manufacturability



Research Team



Dr Serdar Tekin
Postdoc (arrays)



Kacper Burczyk
PhD* (scaling)



Bianca Giuroiu
PhD (logic)

Global Tech Advisor Network



Dr J. Iwan Davies
Grp Tech Dir, IQE PLC



Prof Avirup Dasgupta
Indian Institute of Tech Roorkee



Prof Kwon Kee-Won
Sung Kyun Kwan University



Prof Dr Onur Mutlu
Full Professor, Dept. Information Tech and Elec. Eng.



Global Semiconductor & AI Ecosystem Momentum

- **£5m+ non-dilutive** funding secured through EU and UK semiconductor innovation programmes across pre-and post-spinout development
- Led consortium formation and programme delivery across major UK semiconductor innovation initiatives alongside IQE, Cardiff University and leading UK semiconductor partners
- First strategic investor secured in South Korea (2025), alongside engagement with Korean AI infrastructure leaders including NAVER
- Malta-based EU subsidiary established as Quinas' European semiconductor hub, enabling engagement with Chips JU, IPCEI and IMEC, participation in a €12M RIA programme on Agentic Embodied Edge AI, and a €500k Malta VC co-investment commitment
- Active collaborations spanning the UK, EU/CH, India, KSA, Taiwan, South Korea and Singapore



This momentum underpins Quinas' transition from technology validation to globally scalable AI memory infrastructure.

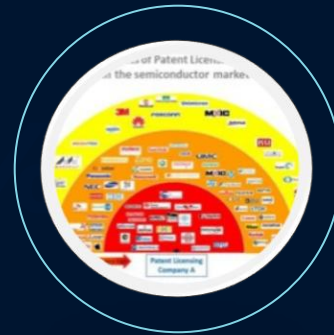
Business Model - Fabless, Licensing-Led Memory Platform



Innovation



Industrialise



License



Scale

Near-Term Commercialisation Pathway

Evaluation Partners → Design Partners → MVP-A System Demonstrators → Licensing & NRE Revenue → Chiplets & Packaged Devices → Royalties

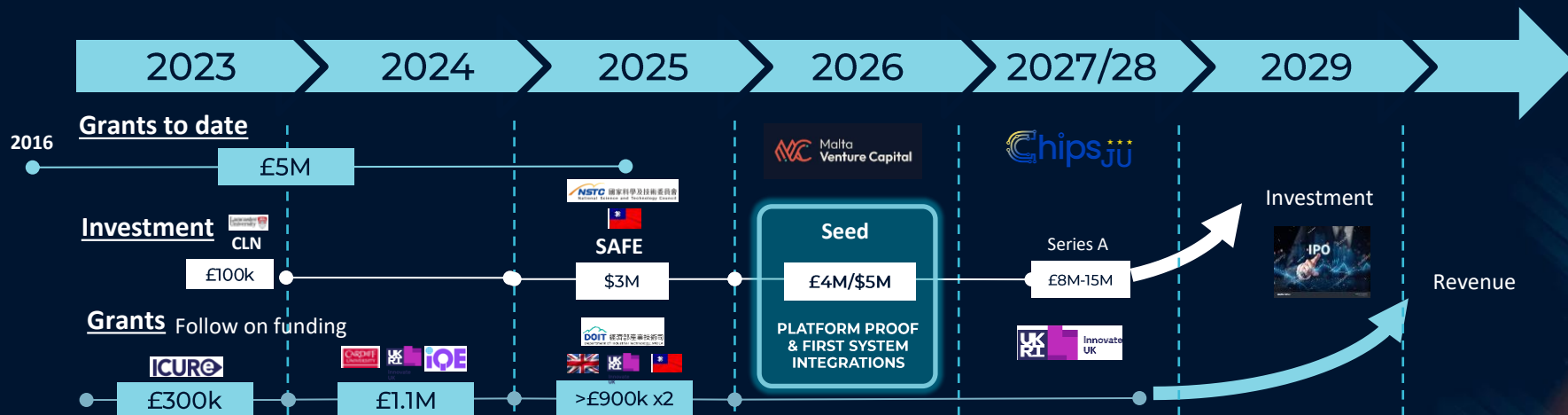
Revenue Streams:

- IP Licensing & Royalties
- Strategic co-development
- AI memory chiplets & packaged devices
- Defence & AI infrastructure applications

Scaling Strategy:

- Partner-led manufacturing & packaging
- Semiconductor ecosystem integration
- Capital-efficient scale-up through global foundry and packaging partners

Funding Roadmap



Seed Execution Focus (18–24 Months:

- TRL-5→6 completion
- MVP-A system demos
- First manufacturing integrations (fabless).

What this round unlocks:

- MVP-A system demos with lead partners
- Foundry & package alignment for MVP-D
- Advanced packaging & chiplet integration

Seed capital converts deep, publicly funded validation into platform revenue and scale

Delivers:

- A partner-ready ULTRARAM™ MVP-A for neuromorphic and in-memory AI compute
- System-level proof of ultra-low-energy AI workloads
- A manufacturable path to scale with industry partners

18-24 month execution → Series A readiness

ULTRARAM™
ULTRA-EFFICIENT MEMORY



Seed Round: Turning Breakthrough into an AI-Native Memory Platform



HQ UK

- IP
- Commercial
- Partnerships

Great Portland St. London W1W 5PF



R&D Labs UK

- Device Physics
- Pathfinding & Arrays

LU Physics, Lancaster LA1 4YB



Taiwan

- Pilot Production
- Foundry
- Adv. Packaging

Taipei Arena 105037



Malta EU

Valletta Europe



- EU Incentives
- ChipsJU
- CMOS Design

Contact: ULTRARAM™ overview



Quinas operates a distributed, fabless model optimised for capital efficiency, speed, and global scale.