

MEMTECH

California Memory Technologies, Inc.

TECHNICAL WHITEPAPER

The Next Decade of AI Memory

A technical perspective on the architecture decisions
shaping next-generation AI systems

*Compute scaled. Memory did not keep pace. The defining hardware
decisions of the coming decade are decisions about data movement.*

Author: Saswat Mishra, Founder & CEO

MEMTECH · Cupertino, California · memtech.ai

June 2026 · Rev 1.0

Contents

Executive Summary	3
1. The Memory Wall: How We Arrived Here	3
2. The Physics of Data Movement	4
2.1 Moving data costs more than computing on it	4
2.2 Arithmetic intensity and the roofline	4
3. The Workload Shift: Why Generative AI Broke the Old Assumptions	5
4. The Architecture Response: Decisions Shaping the Next Decade	5
4.1 High-bandwidth memory and its limits	5
4.2 Advanced packaging and chiplets.....	5
4.3 Near-memory and in-memory compute	6
4.4 Memory hierarchy and tiering.....	6
4.5 Reducing the data that must move	6
5. The Overlooked Layer: Modeling, Verification, and Silicon Bring-Up	7
6. A Framework for Architecture Decisions.....	7
7. Outlook: The Next Ten Years	8
8. Conclusion	8

Executive Summary

For more than a decade, progress in AI hardware has been narrated as a story about compute. Each new accelerator generation is introduced through its peak throughput, its tensor-core counts, and the raw arithmetic it can perform per second. That framing is increasingly out of step with how modern AI systems actually behave. The arithmetic units in a state-of-the-art accelerator are rarely the limiting factor. More often, they sit idle, waiting for data.

The true constraint has shifted to memory: how much data a system can hold close to compute, how quickly it can be delivered, and how much energy is spent moving it. This whitepaper examines why that shift happened, the physics that makes it durable rather than transient, and the specific architecture decisions—high-bandwidth memory, advanced packaging, near-memory compute, memory tiering, and data-movement reduction—that will define competitive AI silicon over the next ten years.

It also addresses a dimension that receives far less attention than it deserves: as the memory subsystem becomes the performance-defining part of the chip, it also becomes the hardest part to model accurately, verify thoroughly, and bring up reliably in silicon at volume. Bandwidth on a datasheet is meaningless until the subsystem behind it can be validated and yielded. The organizations that treat memory architecture and memory verification as first-order engineering disciplines—rather than downstream integration details—will hold the structural advantage.

CORE THESIS

The center of gravity in AI hardware has moved from computation to memory and data movement. The winners of the next decade will be defined not by how fast they can multiply, but by how little they have to move data, how efficiently they move what remains, and how reliably they can prove the memory subsystem works at scale.

1. The Memory Wall: How We Arrived Here

The imbalance now reshaping AI hardware is not new in principle. Computer architects have described the “memory wall” for decades: the observation that processor performance has historically improved far faster than the ability of memory systems to supply data. What changed is the magnitude and the stakes. AI workloads consume data at a scale and pace that turns a long-standing inconvenience into the primary design constraint.

The root cause is a scaling mismatch. Peak compute throughput has grown dramatically across accelerator generations, propelled by process scaling, larger dies, denser arithmetic units, and aggressive architectural specialization. Memory bandwidth has grown too—but along a far shallower curve. Each generation therefore buys an abundance of compute and then spends a growing share of its design budget, its silicon area, and its power envelope simply trying to keep that compute fed.

The consequence is a widening gap between how fast a chip can calculate and how fast it can be supplied with operands. When that gap dominates, adding more arithmetic units yields little benefit: the additional capability is stranded behind a memory system that cannot deliver work quickly enough to use it.

WHY IT IS DURABLE

The memory wall is not a temporary supply problem waiting on the next process node. It is the cumulative result of two technology curves diverging over many years. Closing it requires architectural change, not merely a faster part.

2. The Physics of Data Movement

To understand why memory has become the dominant constraint, it helps to reason in two currencies that architects care about most: energy and locality.

2.1 Moving data costs more than computing on it

The most consequential fact in modern accelerator design is that data movement, not arithmetic, dominates the energy budget. By widely cited estimates, fetching a single word from external DRAM can cost on the order of two orders of magnitude more energy than the floating-point operation that consumes it. A multiply is cheap. The journey the operand takes to reach the multiplier—across the package, through the physical interface, off-chip and back—is not.

This inverts the traditional optimization target. The question is no longer “how many operations can the chip perform?” but “how little can data be moved, and how efficiently can what must move be moved?” In leading designs, the interconnect and memory interfaces increasingly determine both the power profile and the achievable performance, while the arithmetic units themselves are comparatively inexpensive.

2.2 Arithmetic intensity and the roofline

Architects formalize this through the roofline model. Every workload has an arithmetic intensity: the number of operations it performs for each byte it fetches from memory. Below a hardware-specific threshold, performance is capped by memory bandwidth rather than by the compute units. Above it, the workload can saturate the arithmetic and becomes compute-bound.

The uncomfortable reality is how many important AI workloads sit firmly on the memory-bound side of that line. For these workloads, the peak FLOPS figure on a specification sheet is largely irrelevant. Effective performance is governed almost entirely by how fast the system can stream data through the compute fabric.

DESIGN IMPLICATION

Locality becomes the primary design currency. Keeping data on-chip, reusing it before evicting it, and shortening the physical distance it travels matter more than raw arithmetic capacity for an expanding share of real workloads.

3. The Workload Shift: Why Generative AI Broke the Old Assumptions

Large generative models pushed the memory constraint from a chronic condition to an acute one. The shift is most visible in transformer inference, and specifically in the autoregressive decode phase that generates output one token at a time.

Decode is fundamentally memory-bandwidth bound. Generating each token requires streaming model weights and a continually growing key-value cache through the compute units, with very little arithmetic reuse per byte fetched. The arithmetic intensity is low by construction. As a result, an accelerator can carry an enormous compute budget and watch most of it sit idle; effective throughput is set by how fast memory can be read, not by how many operations the chip can theoretically perform.

Two structural pressures compound the problem. First, model parameter counts have grown faster than the memory capacity available close to compute, forcing larger and more complex memory hierarchies. Second, longer context windows enlarge the key-value cache that must be stored and repeatedly streamed during generation, so the per-request memory footprint and bandwidth demand grow with the very capability users most want.

THE STRATEGIC CONSEQUENCE

High-bandwidth memory has become the most performance-critical—and the most expensive, power-hungry, and supply-constrained—component on the board. It is simultaneously the industry's answer to the bandwidth problem and its newest bottleneck.

4. The Architecture Response: Decisions Shaping the Next Decade

The response to the memory constraint is reshaping hardware at every level of the stack. The dominant architectural decisions of the coming decade cluster into five interacting strategies. None is a complete answer on its own; the defining engineering work lies in how a design composes them.

4.1 High-bandwidth memory and its limits

Stacked high-bandwidth memory, connected to the accelerator through silicon interposers and through-silicon vias, has become the default solution for delivering bandwidth at scale. Successive generations have raised per-stack bandwidth and capacity substantially, and the roadmap continues toward wider interfaces and taller stacks. But HBM is not a free lever. It dominates the bill of materials, contributes heavily to package power and thermal load, demands sophisticated 2.5D packaging, and remains supply-constrained. Bandwidth bought this way is expensive bandwidth, which is precisely why reducing the need to move data has become as important as increasing the rate at which it can be moved.

4.2 Advanced packaging and chiplets

If moving data over distance is the cost, shrinking distance is the most direct remedy. Advanced packaging—2.5D interposers, 3D stacking, and chiplet-based disaggregation—lets designers place

memory and compute closer together and partition large dies into smaller, higher-yielding pieces connected by dense, short, energy-efficient links. Packaging has moved from a back-end manufacturing concern to a front-line architectural decision, because where data lives relative to where it is consumed now sets the energy and latency floor of the entire system.

4.3 Near-memory and in-memory compute

A more radical response reverses the usual flow. Rather than continually transporting data to the compute units, near-memory and in-memory compute push selected operations toward where the data already resides. By performing reduction, filtering, or simple arithmetic within or adjacent to the memory array, these approaches cut the volume of data that must traverse expensive interconnect. The techniques carry real trade-offs in programmability, precision, and integration complexity, but for bandwidth-bound, low-intensity workloads they attack the constraint at its source rather than working around it.

4.4 Memory hierarchy and tiering

No single memory technology optimizes capacity, bandwidth, cost, and power simultaneously, so next-generation systems increasingly combine them into deliberate tiers. High-bandwidth memory serves the bandwidth-critical working set; higher-capacity, lower-cost memory holds data that is large but less latency-sensitive; and emerging interconnect standards such as CXL enable memory expansion and pooling beyond the boundaries of a single package. The design problem becomes one of placement and movement policy: deciding what data lives where, and when it migrates, so that the most expensive bandwidth is reserved for the data that genuinely needs it.

4.5 Reducing the data that must move

The cheapest byte to move is the one never moved at all. Quantization, sparsity, compression, and smarter caching and scheduling all reduce the volume of data crossing the memory system in the first place. Because data movement dominates energy, these software-and-architecture techniques often yield larger system-level gains than equivalent investments in raw compute. Over the next decade, co-design across model, compiler, and memory architecture—rather than any single hardware feature—will likely separate efficient systems from inefficient ones.

The strategies above are best understood not as alternatives but as a portfolio. The following summary contrasts what each one optimizes and the trade-off it carries.

Strategy	Primary lever	Principal trade-off
High-bandwidth memory	Raises delivered bandwidth to compute	Cost, power, thermal, supply constraint
Advanced packaging	Shortens the distance data travels	Process complexity, packaging yield
Near-memory compute	Moves operations to the data	Programmability and precision limits

Strategy	Primary lever	Principal trade-off
Memory tiering	Matches data to the right memory class	Placement and migration complexity
Data-movement reduction	Moves fewer bytes in the first place	Accuracy and software co-design effort

5. The Overlooked Layer: Modeling, Verification, and Silicon Bring-Up

As the memory subsystem becomes the part of the chip that determines performance, it also becomes the part that is hardest to get right. This is the dimension most often missing from the AI-hardware conversation, and it is where much of the next decade’s real risk and real advantage will sit.

A bandwidth figure on a datasheet is a promise. Keeping it requires a memory subsystem that has been modeled accurately before silicon, verified exhaustively across corner cases and operating conditions, and brought up reliably in physical hardware at production volume. Each of these is difficult in its own right, and they compound.

- **Modeling.** Architects must predict subsystem behavior—bandwidth, latency, contention, power—early enough to make sound design decisions, when the design itself is still in flux and the workloads are evolving.
- **Verification.** Memory interfaces and controllers must be exercised across an enormous state space of timing, training, and traffic conditions. Defects that escape verification surface as field failures or silent data corruption that are extraordinarily costly to diagnose after the fact.
- **Silicon bring-up.** The transition from a correct design to working, yielding silicon is where many memory-centric programs lose time. Calibration, training firmware, repair, and binning determine whether the theoretical bandwidth is ever realized in volume production.

These activities are not downstream integration details to be resolved at the end of a program. They are first-order determinants of whether an architecture’s memory advantages translate into a shippable, dependable product. As memory subsystems grow more complex and more central, the disciplines that prove them out become correspondingly more strategic.

THE UNDER-PRICED RISK

Two accelerators can carry identical memory specifications and ship with entirely different real-world reliability and yield. The difference is the depth of modeling, verification, and bring-up behind the specification—an investment that is easy to underfund and expensive to skip.

6. A Framework for Architecture Decisions

Memory architecture decisions are inseparable from the workloads a system is built to serve. A useful starting point is to characterize the target workload along the dimensions that most influence the memory subsystem, then let that profile drive the architectural emphasis.

If the workload is...	The dominant pressure is...	Emphasize...
Low arithmetic intensity (e.g. autoregressive decode)	Memory bandwidth	HBM, tiering, data-movement reduction
Capacity-bound (large models, long context)	Memory capacity and placement	Tiering, CXL expansion, compression
Energy- or thermally-constrained	Cost of moving data	Packaging, near-memory compute, locality
Latency-sensitive serving	Predictable data delivery	Hierarchy design, scheduling, verification depth
High-volume production	Yield and reliability	Modeling, verification, bring-up rigor

The discipline this framework encourages is to resist optimizing for headline compute figures and instead design backward from the memory behavior of the actual workload. In a memory-bound regime, the most valuable engineering judgment is about data: where it lives, how far it travels, how often it is reused, and how confidently the subsystem that handles it can be proven correct.

7. Outlook: The Next Ten Years

Three directional shifts seem durable enough to plan around.

- **Memory becomes a primary axis of differentiation.** The distinction between competing accelerators will increasingly come from the memory subsystem and the data-movement architecture around it, not from the compute fabric alone.
- **Co-design replaces layer-by-layer optimization.** Gains will come from optimizing model, compiler, interconnect, and memory together. Treating any one layer in isolation will leave large efficiencies unrealized.
- **Verification and yield become competitive moats.** As memory subsystems grow more complex, the ability to model, verify, and bring them up reliably will separate products that ship dependably at volume from those that struggle to convert their specifications into silicon.

None of these shifts diminishes the importance of compute. They reframe it. Compute brought the industry to this point; memory and data movement are what now hold it back, and solving that constraint is where the most consequential engineering of the next decade will take place.

8. Conclusion

The headline metric of AI hardware is overdue for revision. Operations per second describe a capability that real workloads can rarely use in full, because the limiting resource is no longer arithmetic—it is the memory system that must feed it and the energy spent moving data into place. The architecture decisions that matter most over the next decade are decisions about bandwidth, capacity, locality, hierarchy, and the rigor with which the resulting subsystem is proven out.

The systems that win will be the ones that learn to move data less, move what remains more efficiently, and demonstrate—at volume, in real silicon—that the memory subsystem actually works. **Compute got us here. Memory is what comes next.**

About MEMTECH

MEMTECH (California Memory Technologies, Inc.) is an autonomous AI-silicon infrastructure company based in Cupertino, California, focused on memory IP, verification and modeling, data movement, yield improvement, and silicon bring-up for the custom AI-chip ecosystem. MEMTECH builds the memory and verification foundations that turn architecture decisions into dependable, high-yield silicon.

memtech.ai · Cupertino, California

© 2026 California Memory Technologies, Inc. This whitepaper is provided for informational purposes and reflects the author's technical perspective. All trademarks are the property of their respective owners.