



The AI Infrastructure Race: Why Asia Represents the Next Frontier

Yesterday, Microsoft publicly revealed something remarkable—and deeply consequential. Despite adding two gigawatts of data center capacity last year—the entire output of the Hoover Dam—they still can’t meet demand. Microsoft’s CTO admitted that “even our most ambitious forecasts are just turning out to be insufficient on a regular basis.” Azure usage restrictions now extend through 2026. Internal Microsoft projects are being shut down to preserve capacity. Some customers are defecting to competitors. Amazon and Google report the same strain.

This isn’t a bubble. This is a shortage.

Why the Dot-Com Comparisons Are Dangerously Wrong

The dot-com comparisons fundamentally misread what’s happening. The skeptics see trillions flowing into AI infrastructure and immediately recall fiber optic overbuild in 1999. They’re looking at the wrong precedent. Those cycles built speculative capacity for demand that came too late; today, demand arrived first. Every hyperscaler is sold out.

Today, every watt of power and every chip that comes online is immediately consumed. Microsoft, Amazon, Google—they’re all turning away business. When trillion-dollar companies with unlimited capital can’t build fast enough, you’re witnessing something unprecedented. Compute, power, and bandwidth are the new scarce commodities.

Three realities explain why this moment is fundamentally different—and far earlier in its trajectory than most investors, even sophisticated ones, realize:

First, AI Capability Has Hit a Phase Change. The difference between what modern models can do today and what they could do two years ago is not incremental; it’s categorical. We use these advanced models (Claude, ChatGPT, Grok, Perplexity, and Gemini) extensively for complex programming and analysis every day—and the speed of improvement is staggering. Two years ago, ChatGPT was a clever chatbot. Today, these systems generate code, design chips, and conduct research beyond human foresight. Models are now *improving themselves*, finding optimizations their creators didn’t anticipate. This isn’t iteration—it’s exponential evolution. The jump from GPT-3 to GPT-5 was shocking; the next one will make it look primitive.

Second, Compute Is the Oxygen of Modern Intelligence. There are roughly three ways to advance AI: algorithmic breakthroughs, data quality improvements, and raw compute. Of these,

only compute delivers reliable, predictable progress. Double the compute, and capabilities improve proportionally—the scaling laws are as empirical as gravity. Algorithmic innovations are sporadic. Data quality faces diminishing returns. Compute scales predictably and relentlessly, delivering sustained capability gains. That’s why the hyperscalers’ capital expenditures look extreme: because they work. And this is why they are in an arms race. More GPUs and bigger data centers represent the surest path forward. It’s not speculation; it’s physics.

Third, Inference Demand Is Exploding and We’re In the Early Innings. Until recently, roughly 75% of AI power consumption went to training models, 25% to running them. That ratio is flipping—and not because training is slowing down (in fact, it is increasing exponentially) but because inference is skyrocketing (hyperbolic) as AI moves from labs into production at scale. Even cautious scenarios imply a tenfold increase in inference demand within just a few years and potentially far more as adoption ramps. As systems move toward continual learning—where models improve in real-time from deployment—the boundary between training and inference dissolves entirely, multiplying compute requirements yet again.

This isn’t irrational exuberance. If anything, AI’s transformative potential is *underestimated*, and perhaps massively so, because human intuition struggles with exponential growth and self-reinforcing loops. AI is unique among technologies: it improves itself. Every deployment creates data that trains the next generation, compounding intelligence in a closed feedback loop. The result is an exponential curve of capability, not a linear march.

Every previous technological revolution—electricity, railroads, and the internet—required decades to build infrastructure and train workforces. AI was born into a perfect environment, a world already wired for it: global cloud infrastructure, abundant compute, billions of connected devices, and because it speaks nearly every written human language—it’s immediately accessible to everyone. The result is near-instantaneous mass adoption driving transformative use cases, which feeds back into accelerating AI development. AI’s product is intelligence itself, and intelligence makes everything better, which means AI’s addressable market is limitless. We’re not climbing a mountain; we’re building a rocket ship while riding it.

Asia: The Engine Room of Intelligence

While the United States leads in model innovation and the brains and central nervous system of AI, Asia builds the machinery that makes it possible. The semiconductor, optical, and power supply chains that feed the world’s intelligence engines run through China, Taiwan, South Korea, and Japan.

For U.S. investors already long domestic AI leaders, exposure to Asia is not a geographical hedge—it’s the other half of the economic equation, providing equally critical AI infrastructure trading at traditional manufacturing multiples, and largely ignored by Western investors.

You cannot software your way out of needing physical semiconductors. And the semiconductor supply chain—from design to fabrication to assembly to testing—runs through Asia. This isn’t about choosing sides. It’s about recognizing that the moat isn’t just intellectual property; it’s decades of accumulated expertise, billions in installed capacity, and supply chains that cannot be

replicated quickly. Re-shoring advanced manufacturing to the U.S. will take at least a decade. In the meantime, if you want to build AI at scale, you depend on Asia.

Five Windows into the Asian AI Ecosystem

SK Hynix (South Korea) manufactures high-bandwidth memory (HBM), the specialized memory sitting directly on AI processors. They command over 50% market share in HBM3E, the latest generation. NVIDIA's most advanced chips depend entirely on SK Hynix and Samsung for HBM supply. There is no Western alternative. This is a bottleneck monopoly hiding in plain sight. As AI compute explodes, HBM demand scales with it—and SK Hynix controls the scarcest component in the entire stack.

Tokyo Electron (Japan) manufactures the equipment that even Intel and TSMC depend on. They're the "ASML of Japan"—irreplaceable, with expertise that China desperately wants but cannot replicate. As both the US and China race to build domestic chip capacity, they both must buy from Tokyo Electron. Heads or tails, they win.

Cambricon Technologies (China) is China's leading indigenous AI chip designer and the embodiment of semiconductor sovereignty. While currently years behind NVIDIA in raw performance, Cambricon represents something more important than today's specifications: China's demonstrated ability to develop advanced AI silicon despite export restrictions. When Huawei unexpectedly produced 7nm chips, they revealed capabilities intelligence agencies didn't know existed. China may be 3-5 years behind today, but they were 10 years behind just 3 years ago. That acceleration is geometric, not linear. When a nation with unlimited resources and political will commits to catching up, betting against them is expensive.

Hon Hai Precision/Foxconn (Taiwan) is where AI hardware physically manifests. While attention focuses on chip design and software, someone must manufacture and assemble the millions of servers filling data centers globally. Foxconn's manufacturing scale, supply chain integration, and execution capabilities are irreplaceable. Beyond assembly, they're leveraging AI and robotics internally—smart manufacturing initiatives have increased revenue per employee by over 80%—while also building robots for other companies, positioning them at the intersection of AI infrastructure demand and supply. This manufacturing dominance has largely left the US and won't return. The AI buildout isn't just about designing better chips—it's about producing them at unprecedented scale and speed.

Zhongji Innolight (China) builds the optical interconnects—the photonic nervous system of intelligence—allowing AI processors to communicate at light speed. In modern AI training, thousands of GPUs work in parallel, continuously exchanging massive data volumes. Optical transceivers are the only technology capable of handling these bandwidth requirements. The physics is inescapable: more AI chips means exponentially more interconnects. Innolight is a leading supplier in an exploding market with physics-driven demand that cannot be wished away. As compute scales, optical demand scales quadratically. There is no U.S. equivalent pure-play.

These are just a few examples of many compelling opportunities in our investment portfolio, and they represent themes you simply cannot access through U.S. public markets: bottleneck components, manufacturing dominance, sovereign technology development, and physics-constrained infrastructure.

The China Card: Two Structural Advantages

First, China's drive for chip independence has moved from aspiration to execution, and they have committed immense political will, national resources, and effectively unlimited capital to building a domestic semiconductor industry **with an intensity that dwarfs the Manhattan Project.** Companies like NAURA Technology and AMEC aren't just government-subsidized; they're national priorities. They reflect a national commitment measured in trillions of yuan. Progress once thought a decade away is arriving years early. Political will and capital intensity at this scale create enduring investment optionality. This is a multi-decade mission with the full weight of the state behind every advance. The gap is narrowing faster than most expect.

Second, China possesses massive energy capacity advantages. The U.S. data center developers face multi-year interconnection queues and transformer delivery times exceeding two years. China faces none of these bottlenecks. They've built power generation and grid capacity aggressively while America's infrastructure ages and bureaucracy slows everything. Compute requires power. China has it; America is scrambling for it. Energy is the ultimate limit on compute, and China has structural advantages that will matter enormously in the coming race.

The Generational Barbell (U.S. and Asia)

These dynamics set the stage for a dual opportunity—one driven by innovation, the other by industrial scale. We strongly recommend exposure to both sides of this epochal transformation and history's greatest technological race: participation in perhaps the most revolutionary advancement society will experience—achieving intelligence that makes everything before it look infantile—and the highest-stakes corporate and geopolitical competition ever waged. On one side: American innovators build intelligence itself. On the other: the Asian industrial ecosystem manufactures the materials, memory, optics, and energy that make it real. Each amplifies the other; together they form the full stack of the AI economy. The companies building the infrastructure for this metamorphosis are disproportionately located in Asia. They're trading at fractions of their US customers' valuations. And they're essential to both sides of the U.S.-China AI race.

The dot-com comparisons aren't just wrong—they're wrong in the opposite direction. That era bet on demand that didn't materialize. Today, real demand overwhelms every attempt to meet it. This isn't a bubble waiting to pop. If anything, the profit potential and economic impact of AI is massively ***underestimated*** because humans struggle to comprehend hypervelocity growth and recursive learning systems.

The winners of the AI age won't just be those who build the best models. They'll also be the ones building the backbone—the power, cooling, memory, and machinery that every AI system depends on. Many of those companies are in Asia. Most investors haven't noticed yet.

We have. And we invite you to participate in both sides of a generational opportunity. While the world watches NVIDIA's stock price, the real opportunities are hiding in plain sight across the Pacific—irreplaceable companies powering a race neither side can afford to lose. **We believe this moment will be remembered as both the greatest technological leap in history and the most capital-intensive industrial race ever fought.**

Welcome to the Era of Superintelligence.

Bradford Stanley, CFA

Chief Investment Officer of The Stanley-Laman Group. Designer of SLG's proprietary Ad-Star® System. Brad received a BS in Computer Science from Carnegie Mellon University and is a Chartered Financial Analyst Charterholder

Disclosure: The preceding represents the opinions of The Stanley-Laman Group, Ltd., a Registered Investment Advisor, and is not intended to be used as investment recommendations. All strategies outlined and the views expressed herein offer the risk of loss of principal and are not suitable for all investors. Investors are advised to consult with qualified investment professionals relative to their individual circumstance and objectives.