

TeraFab and the Compute Moat: Elon Musk’s Vision for Vertical AI Chip Manufacturing at Tesla

Douglas C. Youvan

doug@youvan.com

www.youvan.ai

January 29, 2026

In a bold shift that could redefine the boundaries between automotive engineering, robotics, and semiconductor manufacturing, Elon Musk has signaled Tesla’s intent to build its own vertically integrated semiconductor production infrastructure—codenamed “TeraFab.” This vision, unveiled during Tesla’s latest earnings discussions, aims to address a looming constraint in Musk’s AI-centric roadmap: the global shortage of high-throughput compute infrastructure. Tesla’s projected need for 100–200 billion chips per year—encompassing logic, memory, and packaging—positions TeraFab not just as a chip supply solution but as a strategic moat against dependency on external suppliers. This paper explores the technical, strategic, and geopolitical implications of Musk’s proposal, examining what it would mean for Tesla to enter a territory historically dominated by TSMC, Intel, and Samsung. By unpacking the motivations, feasibility, and industry context, we analyze whether TeraFab is a necessary evolution to support Tesla’s ambitions in autonomous driving, humanoid robotics, and AI inference, or a speculative overreach. The paper also investigates what this move reveals about the emerging race to control not just AI software or models, but the very infrastructure of intelligence production itself.

Keywords: Elon Musk, TeraFab, Tesla, AI chips, semiconductor manufacturing, compute moat, vertical integration, robotics, autonomy, packaging, AI hardware, logic chips, memory chips, foundry business, chip fab, infrastructure control.

A collaboration with GPT-4o and GPT-5.2-Thinking. CC4.0.

Outline:

1. Introduction: Tesla's New Frontier
 - 1.1 From automaker to AI infrastructure company
 - 1.2 The origin of the TeraFab concept
 - 1.3 Why chip production has become Tesla's bottleneck
2. The TeraFab Proposal in Musk's Words
 - 2.1 Summary of key statements from earnings calls and shareholder meetings
 - 2.2 Defining "logic, memory, and packaging all in one go"
 - 2.3 The staggering scale: 100–200 billion chips per year
3. Motivations Behind Tesla's Semiconductor Ambition
 - 3.1 Dependency risk and the fragility of global supply chains
 - 3.2 Autonomy and robotics: chip needs beyond current capacity
 - 3.3 Vertical integration as a compute moat strategy
4. Semiconductor Manufacturing: What It Takes
 - 4.1 Logic fabs vs memory fabs vs advanced packaging
 - 4.2 Cleanroom requirements, lithography, yield, and throughput
 - 4.3 The cost and complexity of building leading-edge fabs
5. Tesla's Position in the Global Foundry Landscape
 - 5.1 Current AI chip supply chain partners (Samsung, TSMC)
 - 5.2 Would TeraFab compete with or complement foundries?
 - 5.3 Intel Foundry Services and other potential collaborators
6. Strategic Implications of Full-Stack Compute Sovereignty
 - 6.1 What a vertically integrated Tesla chip stack would enable
 - 6.2 The platform strategy: chips for Tesla, xAI, and Optimus
 - 6.3 Ecosystem lock-in and chip-as-service futures
7. Risks and Constraints Facing the TeraFab Vision
 - 7.1 Technical risk: execution timelines and bleeding-edge challenges
 - 7.2 Capital intensity and opportunity cost for Tesla shareholders
 - 7.3 The temptation of overreach: when verticalization backfires
8. The Industry's Response and the Broader Landscape
 - 8.1 NVIDIA, Intel, AMD: how incumbents are watching the space
 - 8.2 Could Tesla's move accelerate geopolitical decoupling in semiconductors?
 - 8.3 Implications for AI cluster design and national chip policy

9. Philosophical Framing: Who Controls Intelligence Production?

9.1 AI as a physical infrastructure challenge

9.2 The inversion of Moore's Law: AI scaling as bottlenecked by atoms

9.3 Compute sovereignty as the new energy independence

10. Conclusion: TeraFab as Tesla's Final Form?

10.1 Recap of strategic drivers and operational challenges

10.2 Is this a vision of necessity or empire-building?

10.3 The future of AI is made of silicon—and ambition

11. References

11.1 Tesla earnings calls and shareholder transcripts

11.2 Semiconductor manufacturing literature and benchmarks

11.3 Industry coverage from Reuters, Bloomberg, Tom's Hardware, Wccftch

11.4 Expert commentary on AI hardware scaling and autonomy infrastructure

11.5 Strategic papers on vertical integration and foundry economics

Art

1. Introduction: Tesla's New Frontier

Tesla, once narrowly defined as a car company, is undergoing a radical metamorphosis. Under Elon Musk's increasingly expansive vision, Tesla now finds itself at the confluence of multiple high-growth, high-complexity sectors: robotics, artificial intelligence, energy infrastructure, and now—semiconductor manufacturing. Musk's recent articulation of the "TeraFab" concept signals not just a tactical pivot, but a strategic leap into one of the most capital-intensive and geopolitically sensitive domains in the world: the production of AI chips. In doing so, Tesla seeks to address a critical bottleneck in its broader roadmap—a looming shortfall in compute availability for AI applications ranging from autonomous driving to humanoid robotics. This section traces Tesla's evolution toward this new identity and explains why Musk now believes that the company must own the silicon stack to fully realize its ambitions.

1.1 From Automaker to AI Infrastructure Company

Tesla's transition from electric vehicle manufacturer to full-fledged AI infrastructure entity has been gradual but deliberate. The company's investment in neural net training, FSD (Full Self-Driving) systems, and the development of the Dojo supercomputer signaled early that Tesla was not content to outsource its most critical AI capabilities. The Dojo project itself—a custom-built training cluster optimized for vision-based deep learning—highlighted Tesla's increasing concern with owning the full AI pipeline. But while Dojo was a bold step, it was still downstream of chip production. TeraFab, by contrast, is upstream. It is a declaration that Tesla must control not just the software and training infrastructure, but the atomic-level substrate—transistors, interconnects, memory, and packaging—that enables intelligent behavior to manifest at scale.

This is not just a technical pivot; it is a redefinition of Tesla's business identity. While Dojo placed Tesla in competition with cloud AI providers like Google (TPUs) and NVIDIA (H100 clusters), TeraFab places Tesla in conceptual alignment with chip foundries and memory fabricators. The move signals that Tesla's aspirations in autonomy and robotics have reached a point where dependency on others is strategically unacceptable.

1.2 The Origin of the TeraFab Concept

The idea of "TeraFab" emerged publicly during Tesla's shareholder and earnings meetings in late 2025 and early 2026. In these forums, Musk openly speculated about the need to build a fab so large and so vertically integrated that it could combine logic, memory, and packaging under one roof. He argued that the global semiconductor supply chain—already under strain from geopolitical shocks, node transitions, and explosive AI demand—was not scaling fast enough to

meet Tesla's needs. Musk cited Tesla's internal AI chips, Optimus robots, and the compute demands of autonomous fleets as drivers of a future in which Tesla would require output on the order of 100–200 billion chips per year.

The TeraFab concept arose not from idle speculation but from systemic risk analysis. Tesla's dependency on key partners—TSMC for advanced nodes, Samsung for memory, and third-party OSATs for packaging—created vulnerabilities in cost, availability, and time-to-market. Musk's answer is to collapse this disaggregated value chain into a single entity: Tesla itself. This would not only reduce exposure to supply shocks but also unlock aggressive optimization strategies between design, fabrication, and deployment. Whether Tesla will build such fabs from scratch, co-invest with existing players, or acquire existing capacity remains to be seen.

1.3 Why Chip Production Has Become Tesla's Bottleneck

The heart of Musk's concern is scale. As Tesla expands from vehicles to fleets of autonomous robotaxis and millions of general-purpose humanoid robots (via Optimus), its compute needs do not grow linearly—they grow exponentially. The sheer number of inference and training chips required across Tesla's roadmap places it in direct competition with hyperscalers, defense contractors, and national AI initiatives for scarce silicon. Moreover, the bottleneck is not just logic—it includes memory bandwidth, packaging density, and power efficiency, all of which require tight integration across previously siloed domains.

This bottleneck is particularly acute because Tesla is trying to deploy embodied AI, not just cloud-based services. In autonomous vehicles, chips must operate with tight energy budgets, harsh thermal environments, and low-latency constraints. Similarly, humanoid robots require edge inference capabilities with real-time responsiveness and embedded autonomy. General-purpose AI is shifting toward real-world deployment, and the chip supply chain is not ready. Musk's answer is simple: if no one will build it fast enough, Tesla must.

Thus, TeraFab emerges not as a vanity project, but as an operational necessity. It is the natural extension of Tesla's long-standing approach to vertical integration—applied this time not to motors or batteries, but to the silicon mind of the machine.

2. The TeraFab Proposal in Musk's Words

Elon Musk's proposal for TeraFab did not emerge from a single prepared announcement, but rather as a series of comments across Tesla's shareholder meetings and earnings calls. As with many of Musk's most ambitious initiatives, the vision has crystallized gradually, through offhand remarks, strategic warnings, and contextual responses to investor questions. Yet the throughline

is unmistakable: Musk believes that the current semiconductor supply chain is insufficient to meet Tesla’s AI-driven future, and that only a radical act of vertical integration can overcome the impending bottlenecks. This section compiles and interprets Musk’s statements, examining what he means when he says Tesla must integrate “logic, memory, and packaging all in one go,” and how the company could require an output of 100 to 200 billion chips per year.

2.1 Summary of Key Statements from Earnings Calls and Shareholder Meetings

During Tesla’s Q4 FY2025 earnings call and subsequent investor interactions in early 2026, Musk raised alarm bells about chip supply as a future limiting factor. He emphasized that the company’s ambitions in AI and robotics would eventually be constrained by how many chips it could acquire—not simply how many it could design or deploy. Musk warned that Tesla might need to build a “very big fab” in the near future, capable of producing chips at a scale far beyond current industry norms.

What set Musk’s remarks apart from previous chip-supply discussions was the integration of compute hardware into Tesla’s strategic core. No longer treating chips as outsourced components, Musk positioned them as fundamental infrastructure on par with Tesla’s batteries or vehicle platforms. His references to “TeraFab” evoked not just a single facility but a global-scale network of fabs—a vision akin to Tesla’s Gigafactory model, but adapted to the semiconductor domain.

Perhaps most telling was Musk’s direct implication that Tesla could not rely on external suppliers alone—not TSMC, not Samsung, and not Intel. While partnerships could still be valuable, the long-term future, in his words, would require Tesla to become compute-sovereign.

2.2 Defining “Logic, Memory, and Packaging All in One Go”

The phrase “logic, memory, and packaging all in one go” encapsulates Musk’s frustration with the fragmented nature of the chip supply chain. Typically, a chip’s lifecycle moves through several specialized domains:

- **Logic fabrication:** High-performance compute cores are built in foundries like TSMC or Samsung, using bleeding-edge nodes such as 5nm or 3nm.
- **Memory integration:** DRAM or HBM (High Bandwidth Memory) is sourced separately, often from vendors like Micron or SK Hynix.

- **Advanced packaging:** These components are combined through costly and complex packaging processes, often outsourced to OSATs (outsourced semiconductor assembly and test firms).

Musk envisions collapsing this disaggregated model into a unified operation. By doing so, Tesla could reduce latency between design and deployment, optimize for specific edge-AI needs, and avoid being caught in multi-vendor coordination failures. The proposal is not just about owning chip fabrication—it's about streamlining the entire AI hardware stack under one roof.

This integration is crucial for Tesla's targeted applications, such as Full Self-Driving (FSD) and Optimus robots, which require co-optimization between inference logic, low-latency memory, and heat-efficient packaging. For Tesla, the physical proximity and co-design of these domains could translate into speed, reliability, and cost advantages that no off-the-shelf solution can match.

2.3 The Staggering Scale: 100–200 Billion Chips Per Year

Among Musk's most headline-grabbing claims is that Tesla's future demand could reach between 100 and 200 billion chips annually. While he did not specify precisely what constitutes a "chip" in this count—full SoCs, memory dies, chiplets, or packaging units—the implication is clear: Tesla anticipates demand on a scale that dwarfs current production levels for any one company's internal use.

This projected output must be understood in light of Tesla's long-range vision. The company plans to deploy fleets of robotaxis, millions of Optimus humanoid robots, and possibly even its own AI cloud infrastructure via xAI. Each of these platforms would require multiple chips for training, inference, redundancy, and real-time sensing. When one multiplies the chip count per product by Musk's proposed deployment numbers—tens of millions of robots and autonomous vehicles—it becomes conceivable how chip demand might skyrocket.

However, the figure also serves a rhetorical purpose. By invoking a volume that borders on the fantastic, Musk is signaling to both the market and competitors that Tesla intends to control not only the application of AI, but also its material base. If these numbers are taken at face value, then Tesla's fab ambitions rival those of some of the world's largest chip manufacturers—not in revenue, but in raw throughput and vertical control. In this way, the 100–200 billion chip figure functions as a strategic declaration: Tesla is not waiting for the industry to catch up. It intends to build the infrastructure of intelligence itself.

3. Motivations Behind Tesla's Semiconductor Ambition

Tesla's push into semiconductor manufacturing is not a speculative moonshot but a calculated response to a growing set of interlocking pressures: global supply chain fragility, exploding compute demands in autonomy and robotics, and the intensifying race for AI infrastructure dominance. These pressures are reshaping Tesla's strategic posture—from a product-focused innovator to a vertically integrated infrastructure powerhouse. Musk's statements on TeraFab reveal that chip fabrication is no longer a supporting function for Tesla—it is becoming the very foundation of its next technological wave. This section explores the three main motivators driving Tesla toward full-stack semiconductor control.

3.1 Dependency Risk and the Fragility of Global Supply Chains

The COVID-19 pandemic, U.S.–China trade tensions, and semiconductor shortages have exposed the vulnerabilities of just-in-time supply chains in unprecedented ways. Tesla, like nearly every tech-adjacent company, faced manufacturing delays and production bottlenecks due to chip shortages—especially in the wake of increased demand for vehicle ECUs and AI accelerators. These events revealed a hard truth: even the most innovative company is hostage to the foundry capacity and geopolitical stability of its suppliers.

More strategically, Tesla's reliance on external partners such as TSMC (for logic), Samsung (for memory), and third-party OSATs (for packaging and testing) creates a supply stack that is highly fragmented and outside of its operational control. Any disruption—be it natural disaster, political instability, or vendor prioritization—could derail Tesla's aggressive timelines for FSD, Optimus, or its AI training efforts. Musk's proposal to internalize chip production reflects a shift toward supply chain sovereignty. It is an attempt to insulate Tesla from external chokepoints and to guarantee not just availability, but priority, customization, and security across its AI hardware needs.

3.2 Autonomy and Robotics: Chip Needs Beyond Current Capacity

Tesla's ambitions extend well beyond electric vehicles. The company's FSD platform already relies on custom-designed inference chips, and its Dojo supercomputer aims to train AI models with Tesla's proprietary data at immense scale. The next phase—autonomous robotaxis and humanoid robots (Optimus)—requires further magnitudes of compute, deployed not only in the cloud but at the edge.

Unlike general-purpose AI workloads that can be batched and delayed, autonomous systems must operate in real-time with fail-safe reliability. This creates intense pressure on performance-

per-watt, thermal management, and latency—all of which are strongly influenced by the co-design of logic, memory, and packaging. Moreover, these systems are meant to proliferate in the millions or even tens of millions. Every unit shipped increases the cumulative chip burden—not just for inference but also for onboard control, vision processing, and sensory integration.

The global chip industry, even with its accelerating expansion, is not optimized for one company's exponential edge compute needs. Tesla's forecasted demand curve may far exceed what any single vendor can or will reserve. In this light, TeraFab is not about disruption for disruption's sake—it is about building compute infrastructure that doesn't currently exist, in the exact form factor and throughput that Tesla requires.

3.3 Vertical Integration as a Compute Moat Strategy

Tesla's history is deeply rooted in vertical integration. From battery cell production (with Panasonic and later in-house) to custom drive units and software-defined vehicle platforms, the company has repeatedly demonstrated that owning more of its stack leads to better margins, tighter integration, and faster innovation cycles. TeraFab represents the application of this logic to the semiconductor domain, with a new twist: the compute stack is now the ultimate moat.

In an AI-driven economy, controlling access to scalable, low-latency compute is akin to owning the railroads in the age of industrialization. If Tesla can build and operate a vertically integrated chip pipeline—design, logic, memory, packaging, testing, and deployment—it would no longer be at the mercy of NVIDIA's supply cadence or TSMC's node availability. Instead, it would wield asymmetric leverage: its competitors would depend on public supply chains, while Tesla would operate on a private highway.

This control also enables optimization loops across the stack. Tesla could tightly align neural network architecture, chip design, and real-world deployment feedback—leading to gains in performance and cost-efficiency that competitors cannot easily replicate. Just as Apple's chip design efforts led to a leap in mobile efficiency and performance via the M-series processors, Tesla seeks to unlock a similar effect—only this time, in AI for embodied systems.

In this sense, TeraFab is not merely a manufacturing initiative. It is a sovereignty project, a control stack, and a long-term moat rolled into one.

4. Semiconductor Manufacturing: What It Takes

For Tesla to execute Elon Musk's TeraFab vision, it must enter one of the most technically demanding and capital-intensive industries in the world. Semiconductor manufacturing,

especially at the leading edge, is not a monolithic process but a layered ecosystem involving logic processing, memory fabrication, and advanced packaging—each with its own specialized infrastructure, materials, and bottlenecks. What Musk proposes is not merely “building a fab,” but collapsing this entire value chain into a vertically integrated AI chip factory capable of handling the full stack—from transistor layout to multi-die packaging. This section provides an overview of what this entails, covering the distinctions among different types of fabs, the extreme cleanroom and lithography requirements involved, and the staggering costs that make semiconductor manufacturing one of the most capital-intensive endeavors on Earth.

4.1 Logic Fabs vs Memory Fabs vs Advanced Packaging

The term “fab” often obscures the diversity of processes involved in building a modern semiconductor. A logic fab is optimized for creating compute-intensive components—CPUs, GPUs, neural engines—at leading-edge nodes such as 5nm, 3nm, or even the experimental 2nm class. These require extreme ultraviolet (EUV) lithography, multi-patterning techniques, and highly precise doping and etching operations. TSMC, Intel, and Samsung are among the few entities with experience in this domain, and even they struggle with consistent yields and ramp-up timelines.

Memory fabs, on the other hand, are specialized for producing DRAM, NAND, or high-bandwidth memory (HBM). The manufacturing constraints are different: while density is still key, the lithography demands are slightly more relaxed, and the focus shifts toward stacking layers vertically and ensuring endurance. Samsung and SK hynix dominate this space, with enormous economies of scale and deeply refined processes.

Advanced packaging refers to the post-fabrication assembly of multiple logic and memory dies into a single functional unit. It is increasingly critical for high-performance AI chips. Techniques such as 2.5D interposers, chiplet-based architectures, and through-silicon vias (TSVs) have emerged to enhance bandwidth, reduce latency, and lower power consumption. OSAT providers like ASE and Amkor have specialized in these processes, but integrating them tightly with in-house logic and memory presents a formidable challenge.

What Musk proposes—combining all three under one manufacturing ecosystem—is virtually unprecedented at commercial scale. Each element represents a distinct industrial vertical. Merging them implies not just capital investment, but radical organizational synthesis.

4.2 Cleanroom Requirements, Lithography, Yield, and Throughput

At the core of semiconductor fabrication is the cleanroom—a hyper-controlled environment that filters out particles far smaller than a human hair. For leading-edge nodes, even a single airborne particle can destroy a wafer die worth thousands of dollars. Cleanrooms are typically rated by how many particles of 0.5 microns or larger are permitted per cubic meter. For example, a Class 1 cleanroom allows only one such particle per cubic foot of air. Maintaining these standards across tens of thousands of square meters is a non-trivial feat of engineering and logistics.

Lithography is the heart of transistor miniaturization. EUV lithography—required for nodes at 7nm and below—relies on extreme ultraviolet light at 13.5nm wavelength. It requires precision optics (such as those made by ASML), multi-million-dollar machines, and photoresists that are still evolving to handle EUV’s quirks. Any attempt by Tesla to produce its own logic chips at the bleeding edge would require either partnerships with EUV-capable suppliers or licensing and integrating ASML’s entire process chain—a significant feat, even for nations, let alone companies.

Yield refers to the percentage of usable dies on a wafer. Early runs at advanced nodes can produce yields as low as 30%, rising to 90% or higher only after extensive process tuning. Poor yield erodes margins, delays schedules, and inflates unit cost. Throughput, meanwhile, is constrained by wafer processing time, lithography exposure cycles, and tool availability. The high-volume production Musk envisions—billions of chips per month—requires yields and throughput approaching the limits of physics and logistics.

Tesla has proven adept at scaling hardware factories for vehicles and batteries, but semiconductors operate at a different order of magnitude. Just a single mask set for a modern logic chip can cost \$10–20 million, and tuning yield requires a feedback loop between design, test, and fab operations that typically takes years to refine.

4.3 The Cost and Complexity of Building Leading-Edge Fabs

Building a modern semiconductor fab is one of the most expensive and regulated industrial undertakings in the world. A state-of-the-art logic fab—such as those being built by TSMC in Arizona or Intel in Ohio—can cost anywhere from \$15 billion to \$25 billion for just one phase of operations. These costs include not only the facility itself but also ultra-pure water systems, extreme power redundancy, climate stabilization, chemical delivery infrastructure, and tens of thousands of machines calibrated to atomic tolerances.

The complexity is not limited to construction. A fab requires a long-term ecosystem: thousands of skilled engineers, semiconductor toolchain vendors, safety certification processes, and continuous supply of rare materials like photoresists, noble gases, and silicon-on-insulator substrates. If Tesla intends to build even one TeraFab—let alone a network—this will require a multi-decade investment and coordination across dozens of industrial partners.

Moreover, regulatory approval and environmental impact studies are nontrivial. The water and energy demands of a chip fab are staggering. Leading-edge fabs can consume over 10 million gallons of ultra-pure water per day and require gigawatt-scale power supplies. Building these facilities in locations already stressed by climate or grid limitations introduces logistical and political risk.

Despite these barriers, Musk has a history of challenging industrial norms and succeeding where others predicted failure. The cost and complexity of TeraFab do not necessarily rule it out—they simply raise the stakes. What Tesla is contemplating is not a manufacturing plant. It is a geopolitical-grade infrastructure asset, one that may sit at the crossroads of global power, national policy, and technological sovereignty.

5. Tesla's Position in the Global Foundry Landscape

Tesla's proposal to create a vertically integrated semiconductor manufacturing operation—TeraFab—places it in a complex and evolving global ecosystem of chip suppliers and foundries. While Tesla has historically relied on external partners for fabrication and memory, the company's ambition to consolidate logic, memory, and packaging internally would fundamentally shift its relationship to these players. Whether TeraFab would serve as a complementary strategy or a competitive one depends not only on Tesla's execution but on the willingness of the existing semiconductor establishment to evolve alongside it. This section explores Tesla's current chip supply chain relationships, analyzes whether TeraFab implies direct competition with legacy foundries, and assesses potential for partnership—especially with Intel Foundry Services and other emerging options.

5.1 Current AI Chip Supply Chain Partners (Samsung, TSMC)

Tesla's current AI hardware pipeline is built upon a network of established semiconductor partners. The company has designed several custom chips, including those powering its Full Self-Driving (FSD) platform and the Dojo supercomputer, but has not fabricated them in-house. Instead, Tesla relies on major foundry players to produce its designs at scale.

TSMC (Taiwan Semiconductor Manufacturing Company) is believed to be the primary logic fab partner for Tesla's AI inference chips, particularly at advanced nodes like 7nm and 5nm. TSMC offers the kind of manufacturing precision, throughput, and node maturity required to realize Tesla's custom silicon. Their deep experience in high-performance chip fabrication makes them the default choice for logic-heavy applications.

Meanwhile, Samsung has reportedly signed a significant deal with Tesla—reportedly worth \$16.5 billion—to provide advanced memory chips and may also participate in logic fabrication under contract. Samsung's position as both a DRAM/NAND supplier and a logic fab (with leading-edge EUV capabilities) makes it a strategic dual-use partner. This relationship appears particularly valuable in Tesla's attempts to co-optimize logic and memory for edge-AI applications.

Despite these partnerships, Musk's public statements suggest that Tesla sees reliance on even its most trusted partners as a risk. The desire to internalize production is not necessarily a reflection on TSMC or Samsung's capabilities, but rather on Tesla's belief that it must outpace industry norms in both capacity and customization.

5.2 Would TeraFab Compete With or Complement Foundries?

At first glance, TeraFab appears to position Tesla as a direct competitor to foundries like TSMC and Samsung. If Tesla succeeds in building a high-throughput, vertically integrated chip operation, it could theoretically fabricate not just internal silicon but chips for external customers—thus entering the foundry business in the conventional sense. This would pit Tesla against companies whose business models revolve around contract chip manufacturing.

However, a closer analysis suggests a more nuanced relationship. Tesla's immediate motivation is not to sell wafers to third parties but to guarantee internal supply. TeraFab's primary function would be to serve Tesla's AI roadmap—supporting in-house products like FSD computers, Dojo nodes, Optimus robots, and xAI inference systems. In this sense, Tesla may resemble Apple, which designs its own chips but relies on partners for fabrication, with the critical difference that Tesla wants to bring that fabrication under its own roof.

In scenarios where capacity constraints persist or geopolitical disruptions continue, Tesla's TeraFab could actually relieve pressure on global foundries by absorbing its own demand. This would free up capacity for other customers and reduce Tesla's prioritization conflicts with foundry partners. In such a case, TeraFab would act more as a strategic buffer than a head-to-head competitor.

Still, should Tesla succeed in creating surplus capacity, the possibility of offering foundry services cannot be discounted. In doing so, Tesla could challenge Intel’s emerging foundry ambitions and even alter pricing dynamics across specific chip verticals.

5.3 Intel Foundry Services and Other Potential Collaborators

Among Tesla’s potential partners in realizing the TeraFab vision, Intel Foundry Services (IFS) stands out. Musk has publicly acknowledged discussions with Intel regarding possible collaboration, particularly in scenarios where Intel might co-invest or support fab construction with Tesla. Intel’s renewed push into contract manufacturing—IFS aims to challenge TSMC by offering U.S.-based fab capacity—aligns with Tesla’s desire for sovereignty over chip production, especially within domestic or geopolitically aligned territories.

A Tesla–Intel partnership would benefit both parties: Tesla would gain access to Intel’s manufacturing infrastructure and process nodes (including advanced packaging capabilities through Intel’s Foveros and EMIB technologies), while Intel would secure a marquee customer with extreme throughput demand. Moreover, such a partnership might be favored by U.S. policymakers eager to localize semiconductor production and reduce reliance on East Asian supply chains.

Other potential collaborators include GlobalFoundries, which specializes in mature nodes and automotive-grade silicon; Applied Materials and ASML, which provide lithography and process tools; and memory-centric companies like Micron, which might see benefit in cross-licensing arrangements or advanced integration projects.

Ultimately, whether Tesla chooses to go it alone, partner selectively, or hybridize internal and external capacity will depend on execution risk, capital allocation, and regulatory strategy. The one constant is clear: TeraFab places Tesla in the gravitational field of the global foundry landscape—not merely as a customer, but as a gravitational force in its own right.

6. Strategic Implications of Full-Stack Compute Sovereignty

If Tesla succeeds in building and operationalizing TeraFab, it would not merely improve supply chain resilience—it would fundamentally shift the dynamics of power in AI, robotics, and the semiconductor industry. Full-stack compute sovereignty—where design, fabrication, memory, packaging, deployment, and feedback all occur within a single vertically integrated system—transforms Tesla from a consumer of compute into a manufacturer of intelligence infrastructure. This model enables a new class of control, optimization, and monetization that competitors, including traditional automakers and cloud providers, may struggle to match. The implications

go well beyond autonomy or manufacturing efficiency—they touch on how intelligence is deployed, controlled, and scaled across domains.

6.1 What a Vertically Integrated Tesla Chip Stack Would Enable

A fully internal Tesla chip stack would create seamless interoperability across hardware, software, and deployment environments. By owning the chip lifecycle from architecture to fabrication to in-field application, Tesla could achieve optimizations that are simply not available to companies operating within fragmented supply chains. Benefits include:

- **Real-time feedback loops** between deployed chips (e.g., in robotaxis or robots) and the training clusters used to improve them.
- **Co-optimization of silicon and neural network architectures**, enabling specialized inference units tailored to specific tasks like perception, planning, or manipulation.
- **Power efficiency and thermal management breakthroughs** by tailoring chip design to Tesla’s known constraints in automotive and humanoid form factors.
- **Accelerated deployment cycles**, reducing the lag between model improvement and hardware rollout across global fleets.

Importantly, a vertically integrated stack reduces latency not just in data but in iteration—giving Tesla the rare ability to innovate faster in a domain (semiconductors) typically known for slow cycles and vendor dependencies.

6.2 The Platform Strategy: Chips for Tesla, xAI, and Optimus

TeraFab does not exist in isolation—it would support a trifecta of interlinked Tesla initiatives:

- **Autonomous Vehicles:** The FSD platform requires inference chips optimized for edge performance, vision processing, and real-time safety validation. Tesla’s current chips are custom-designed but externally fabricated. TeraFab would allow the next-generation designs to be built with tailored manufacturing constraints in mind.
- **Dojo and xAI:** Tesla’s in-house supercomputer, Dojo, and Musk’s separate AI company, xAI, both rely on high-throughput training hardware. The demands here differ: large-scale, high-bandwidth, low-latency interconnects for distributed training. A vertically integrated Tesla chip program could feed both inference and training pipelines with shared architectural DNA.

- **Optimus (Humanoid Robot):** Arguably the most demanding application, Optimus will require edge inference, actuation control, safety computation, and ultra-efficient power draw—across millions of units if Musk’s projections are realized. No existing chip platform satisfies all these constraints simultaneously. TeraFab offers the possibility of custom chiplets or packages for robot brains, aligned precisely with Optimus’s evolving form factor.

Together, these three platforms (cars, cloud, and humanoid robots) suggest Tesla is not simply pursuing chip sovereignty for supply security—it is building a **platform architecture**, where chips are foundational assets in a full-stack deployment ecosystem that stretches from data center to sidewalk.

6.3 Ecosystem Lock-In and Chip-as-Service Futures

Perhaps the most transformative implication of compute sovereignty is the possibility of ecosystem lock-in and service monetization. If Tesla controls the chip stack, it can:

- Tie functionality to silicon: Just as Apple ties features to its M-series chips, Tesla could bind FSD capabilities, Dojo access, or robot functionality to hardware authentication, discouraging third-party servicing or unauthorized use.
- Deploy chips as a service: Through fleet updates or subscription models, Tesla could deliver neural model upgrades tied to specific chip architectures, effectively leasing performance improvements to customers over time.
- Create proprietary toolchains and developer ecosystems: Owning the chip design opens the door for Tesla to offer APIs, simulation environments, and training tools tailored to its own silicon—deepening customer and developer dependence.

This path leads to chip-as-infrastructure—where Tesla’s silicon becomes the gateway to its entire AI product ecosystem. In such a scenario, TeraFab is not simply about chipmaking. It is about controlling the interface between artificial intelligence and the real world, mediated through hardware Tesla both owns and governs.

In this light, the strategic logic of TeraFab becomes clear: by controlling its chips, Tesla controls its destiny—not just in cars, but in a world increasingly run on embodied intelligence.

7. Risks and Constraints Facing the TeraFab Vision

While TeraFab presents a compelling vision for Tesla's future as a vertically integrated AI and semiconductor powerhouse, it is not without immense challenges. From technical feasibility to financial burden to strategic overreach, the path Musk envisions is strewn with risks that could threaten execution or divert the company from its core strengths. These risks are not speculative—they are embedded in the history of every prior attempt to internalize chip manufacturing by non-traditional players. Understanding these constraints is critical to assessing whether TeraFab is a disruptive breakthrough or an expensive misstep. This section outlines the three main categories of risk: technological, financial, and strategic.

7.1 Technical Risk: Execution Timelines and Bleeding-Edge Challenges

Semiconductor fabrication at the cutting edge is one of the most technologically complex processes ever developed by humanity. Even experienced players like Intel and Samsung, with decades of R&D and global supply networks, have struggled with node transitions, yield optimization, and lithography scaling. For Tesla—a company with no historical fabrication experience—to succeed in building a fab capable of producing high-performance AI logic, HBM memory, and advanced chiplet packaging is an undertaking with extremely high execution risk.

Among the key challenges:

- **Time-to-ramp:** Building a fab is not the same as operating one. Even after construction, it can take years to tune processes, establish repeatable yields, and achieve high wafer throughput. A delayed fab offers no short-term relief from chip shortages and may become obsolete if AI hardware architectures evolve during its ramp.
- **EUV and process complexity:** If Tesla targets 5nm, 3nm, or lower nodes, it must either partner with ASML and other equipment makers or build extraordinary in-house photonics and lithography competence. These steps are non-trivial and require rare expertise.
- **Process knowledge:** Foundries like TSMC employ tens of thousands of engineers, each operating within mature feedback systems, defect-tracking protocols, and atomic-level process controls. Even hiring top talent, Tesla would face a multi-year learning curve to internalize this tacit knowledge.

Failure to meet technical milestones would not only delay Tesla's product roadmap—it could also erode credibility with investors and partners, especially if capital is redirected from more mature initiatives.

7.2 Capital Intensity and Opportunity Cost for Tesla Shareholders

The financial commitment required to build even a single high-end semiconductor fab is immense. Estimates for TSMC's and Intel's latest fabs range from \$15 billion to \$25 billion *per site*, and this does not include the cost of multiple redundant lines, R&D, yield stabilization, and packaging infrastructure. A full TeraFab network—especially one with integrated logic, memory, and packaging—could require \$50–100 billion over a decade.

For Tesla shareholders, this raises a critical question: is this the best use of capital?

- **Dilution of focus:** Investors may worry that instead of focusing on vehicle margins, robotaxi deployments, or energy storage, Tesla is taking on fab risks that traditionally belong to companies like Intel or Samsung.
- **Return on investment:** Even if successful, chip manufacturing has very different financial characteristics than software or services. It is cyclical, margin-sensitive, and vulnerable to global demand fluctuations. A fab may not deliver Tesla-like returns even in the best-case scenario.
- **External incentives:** Governments like the U.S., through the CHIPS Act, may offer subsidies or tax credits for domestic manufacturing. But whether these incentives offset the sheer scale of the required investment remains uncertain.

The opportunity cost of committing capital to TeraFab must be weighed against all the other frontiers Tesla is trying to lead—robotics, autonomy, training infrastructure, and vehicle production scaling. A misallocation could hinder all of them.

7.3 The Temptation of Overreach: When Verticalization Backfires

Tesla's success has long been tied to its aggressive vertical integration strategy—from building its own battery cells and drivetrain components to owning retail and service networks. However, vertical integration is not universally beneficial. When applied too broadly or without sufficient organizational expertise, it can lead to complexity that chokes innovation, slows execution, and increases systemic fragility.

Historical analogs offer cautionary tales:

- **Intel's IDM bottlenecks:** Intel's integrated device model (design + fab) eventually became a liability, as the company struggled to keep up with TSMC's pace while maintaining backward compatibility with its product stack.

- **Apple’s selective verticalization:** Apple excels at chip design (e.g., M1/M2), but outsources fabrication to TSMC—a deliberate choice to avoid being dragged into fab economics and technical risk.
- **General Motors’ failed semiconductor aspirations:** Decades ago, GM attempted to develop its own microelectronics division. The venture faltered due to scale mismatches, internal friction, and a lack of sustained focus.

For Tesla, the risk lies in assuming that excellence in EVs and AI design translates directly to excellence in nanoscale process control and atomic-level manufacturing. If the company overestimates its ability to master these domains, it could end up with a fragmented, costly, and underutilized fabrication operation—a monument to ambition, but a drag on performance.

In short, while vertical integration has served Tesla well, there are limits to what even a visionary company can internalize without risk of strategic overreach. TeraFab may be one such frontier where wisdom lies in partnership, not autonomy alone.

8. The Industry’s Response and the Broader Landscape

Tesla’s TeraFab ambition has reverberated across the semiconductor, AI, and geopolitical landscapes. While Elon Musk’s ability to dominate headlines is not new, the implications of a major industrial player building its own chip megafab—especially one encompassing logic, memory, and packaging—has caused both intrigue and unease among incumbent chipmakers, AI cluster builders, and national security strategists. In many ways, TeraFab is more than a company initiative; it is a provocation to an entire industry structure built on specialization, interdependence, and decades of supply chain optimization. This section examines how key players are reacting, what geopolitical consequences may emerge, and how Tesla’s move could reshape national and corporate strategies for AI deployment at scale.

8.1 NVIDIA, Intel, AMD: How Incumbents Are Watching the Space

For incumbent chipmakers, Tesla’s TeraFab is both a threat and a validation.

- **NVIDIA** likely sees TeraFab as a challenge to its long-standing dominance in AI infrastructure. While Tesla already designs its own AI inference chips (e.g., for Full Self-Driving) and custom training chips (e.g., Dojo), those efforts have been niche relative to the broader demand for NVIDIA’s H100s and future B-series architectures. However, if Tesla’s TeraFab scales to support both internal and external AI workloads—especially via xAI—this could represent the beginning of an alternative AI stack, bypassing NVIDIA’s

hardware altogether. The implication is that embodied AI (cars, robots) may not need general-purpose accelerators when a custom vertical stack is available.

- **Intel**, by contrast, may view TeraFab as a potential ally. Intel Foundry Services (IFS) has struggled to win large anchor customers in its effort to compete with TSMC and Samsung. Musk's public acknowledgment of potential collaboration with Intel suggests an avenue for mutual benefit. Intel could provide early node access, fabrication guidance, or packaging infrastructure in exchange for long-term volume commitments. However, should Tesla eventually internalize everything, Intel would be left with another hyperscaler that ultimately decided to go solo.
- **AMD** remains peripheral to Tesla's current architecture needs but may view the TeraFab announcement as confirmation that bespoke silicon is the new strategic norm. As AMD pushes further into AI acceleration (through Instinct MI series and Xilinx IP), it may need to consider partnerships with vertical integrators or risk being excluded from closed-loop ecosystems like Tesla's.

More broadly, all three incumbents must now contend with the reality that the most ambitious AI companies may not be customers in the future—they may be competitors with fabs.

8.2 Could Tesla's Move Accelerate Geopolitical Decoupling in Semiconductors?

The global semiconductor industry is already fracturing under geopolitical pressure. U.S.–China tensions, export controls on advanced nodes, and national security concerns have made chip manufacturing a matter of statecraft. Tesla's TeraFab concept could accelerate this decoupling by reinforcing a new norm: critical compute must be produced within national or company borders.

- **U.S. policymakers** may welcome Tesla's fab ambitions as a way to reduce dependence on TSMC and East Asian supply chains. The CHIPS and Science Act, which offers tens of billions in incentives for domestic manufacturing, could support TeraFab if sited on U.S. soil. A Tesla chip plant producing logic, memory, and packaging would be a crown jewel in Washington's goal of re-shoring semiconductor sovereignty.
- **China**, which is aggressively investing in self-sufficiency via SMIC and other state-supported efforts, may interpret Tesla's verticalization as justification for its own. If Tesla walls off its silicon stack, other major technology companies—especially those facing export controls—may feel pressure to accelerate internal R&D for sovereign chip production.

- **Europe, Japan, and India** may view Tesla’s model as a template: blend industrial ambition, AI verticalization, and chip sovereignty into a single national project. TeraFab’s success could signal that the age of ultra-globalized chip dependency is ending—and that national champions must arise.

Tesla’s entry into fabrication therefore carries more than technical significance; it is a soft-power move that challenges the assumptions of shared global infrastructure and invites emulation.

8.3 Implications for AI Cluster Design and National Chip Policy

TeraFab also threatens to disrupt the traditional logic of AI cluster development. Today’s supercomputers—whether used by OpenAI, Meta, or Google—are constructed from off-the-shelf components: NVIDIA GPUs, AMD CPUs, HBM memory, and third-party interconnects. The assumption is that AI scale = buying more standardized parts.

Tesla’s vision flips this: AI infrastructure should be designed holistically, with the chip, model, deployment context, and product lifecycle co-optimized. This has several implications:

- **AI cluster design** may become less modular and more monolithic. Rather than scaling with commodity hardware, future AI systems—especially embodied ones—may scale with **vertically integrated pods** of co-designed compute, trained on in-house data, and deployed directly into physical fleets.
- **Training/inference separation** may blur. If Tesla produces its own training clusters (Dojo) and inference hardware (robot/vehicle chips), it could pioneer architectures optimized for **full-cycle feedback** between the real world and the model update loop—something no current cloud provider offers natively.
- **National chip policy** may need to adapt to companies, not just countries, as sovereign compute actors. TeraFab would not be a state project, but its consequences could rival those of Intel’s Ohio plant or TSMC’s Arizona facility. Governments may need to consider new categories of support, export control, and security oversight tailored to private AI fabs with global reach.

In total, TeraFab is not just a Tesla story—it is a reconfiguration of the landscape in which AI is built, regulated, and deployed. It marks a shift from data-centric AI dominance to matter-centric AI sovereignty, where the company that controls atoms controls intelligence.

9. Philosophical Framing: Who Controls Intelligence Production?

Beneath the technical and strategic layers of the TeraFab vision lies a deeper philosophical question: *Who should control the production of intelligence in the physical world?* Musk's proposed verticalization of AI chip fabrication is not just a business decision—it is a move that challenges traditional notions of computation as a fluid, virtual, and infinitely scalable resource. It forces a reorientation toward the *material constraints* of AI, and by extension, toward the *politics, ethics, and ownership* of intelligence itself. This section explores three interlocking philosophical lenses: AI as physical infrastructure, the reversal of Moore's Law as a practical limit on abstraction, and compute sovereignty as the new frontier of autonomy and control.

9.1 AI as a Physical Infrastructure Challenge

For decades, discussions of artificial intelligence have been dominated by abstractions—algorithms, data, software layers, neural architectures. But as AI systems become embedded in vehicles, robots, and sensor-driven environments, a shift is taking place. Intelligence is no longer just software running on general-purpose hardware—it is becoming *a physical infrastructure problem*.

In Tesla's case, this infrastructure spans from training clusters (Dojo) to real-time inference engines deployed in edge environments (vehicles, robots). These systems must meet physical-world constraints: latency, energy, heat dissipation, reliability, and real-time responsiveness. You cannot deploy large-scale embodied AI without also solving for the physical logistics of heat sinks, power converters, interconnects, and silicon yields.

TeraFab reframes AI as something you *build*, not merely program. In this framing, intelligence is not just trained—it is manufactured. It lives not only in weights and activations, but in the lattice of silicon and copper that must be etched, layered, and cooled. This brings AI into the realm of factories, minerals, and national infrastructure—shifting power away from the cloud and into the foundry.

9.2 The Inversion of Moore's Law: AI Scaling as Bottlenecked by Atoms

Moore's Law, the decades-long trend of exponential growth in transistor density, once gave software developers the illusion of infinite scaling. Every few years, hardware improvements would "catch up" with algorithmic demands. That illusion has now broken down. AI's appetite for compute—especially in large model training and edge inference—is outpacing hardware improvements, leading to a new reality: *software is bottlenecked by matter*.

This inversion of Moore’s Law forces a reckoning with the finiteness of physical resources. Chip yields depend on atomic precision. Wafer capacity is limited by rare materials and fragile logistics. Packaging density runs into thermal ceilings. Even power distribution becomes a constraint at scale. As AI ambitions soar, it is no longer enough to write better code—you must build better matter.

TeraFab is Musk’s answer to this inversion. It internalizes the material basis of AI, allowing Tesla to scale not just by refining algorithms but by controlling the substrate of intelligence. It is a return to constraint-aware engineering, where the frontier is no longer cloud elasticity but wafer availability, cleanroom purity, and node transitions. In this view, whoever best controls *atoms at scale* will win the next decade of AI deployment.

9.3 Compute Sovereignty as the New Energy Independence

In the 20th century, national power was shaped by control over oil, gas, and energy infrastructure. In the 21st, compute sovereignty may play an equivalent role. Nations and corporations alike are beginning to treat access to compute—not just in abstract cycles but in physical silicon—as a critical resource, with strategic implications for economic growth, military power, and social stability.

Just as energy independence was pursued to escape geopolitical vulnerability, compute independence is emerging as a hedge against export controls, supply chain fragility, and the monopolization of AI platforms. Musk’s TeraFab fits directly into this frame. By internalizing chip production, Tesla gains the ability to deploy and iterate its AI systems without being held hostage by TSMC’s capacity, Samsung’s priorities, or NVIDIA’s pricing.

This sovereignty is not just about security—it is about leverage. Tesla could, in theory, deploy compute not only to its own fleet and robots, but to *licensed partners*, enforcing alignment through silicon itself. Chip design becomes a tool of governance; fab ownership becomes a gatekeeping layer on AI deployment. In this light, TeraFab is not merely an industrial ambition. It is a political statement: *the future of intelligence must be sovereign, embodied, and made—by us*.

In sum, the TeraFab vision invites us to reframe artificial intelligence not as ethereal, virtual, or disembodied, but as a grounded phenomenon—rooted in factories, forged in atoms, and governed by those who control its production. The question is no longer “who writes the algorithm?” It is: *who builds the substrate where the algorithm lives?*

10. Conclusion: TeraFab as Tesla's Final Form?

Tesla's TeraFab initiative represents a radical inflection point—not just in the company's trajectory, but in the evolution of artificial intelligence as a material and geopolitical force. It fuses chip manufacturing, AI deployment, and industrial sovereignty into a single, coherent, and intensely ambitious strategy. As with many of Musk's transformative ventures—from reusable rockets to electric cars to planetary colonization—TeraFab poses a provocative question: *Is this the logical culmination of Tesla's mission, or the first step in a new empire of silicon intelligence?* This final section distills the strategic logic, evaluates the deeper intent, and considers the broader implications of Tesla's possible transformation into the world's most vertically integrated AI company.

10.1 Recap of Strategic Drivers and Operational Challenges

At its core, TeraFab is a response to three strategic pressures:

- **Dependency risk** from external foundries, packaging firms, and memory suppliers.
- **Exponential demand** for high-performance, edge-deployable AI compute—driven by Tesla's FSD systems, humanoid robots, and training clusters.
- **The need for speed**, flexibility, and architectural co-optimization between AI models and hardware in a feedback-rich deployment environment.

By bringing chip fabrication in-house, Tesla seeks to eliminate chokepoints, accelerate innovation, and carve a sustainable moat in a future where intelligence is embedded in everything from vehicles to infrastructure to robots.

Yet this path is riddled with operational challenges. Leading-edge semiconductor fabrication is among the most complex, resource-intensive, and execution-sensitive undertakings in modern industry. Tesla must master wafer yields, lithography, thermal management, cleanroom logistics, and packaging precision—while contending with the capital drag and opportunity cost of building fabs in parallel with scaling global production of vehicles and robots. No company has ever attempted such verticalization in AI hardware at this scale—success would be historic, failure could be deeply disruptive.

10.2 Is This a Vision of Necessity or Empire-Building?

TeraFab walks a fine line between pragmatic response and grand ambition. On one hand, it arises from genuine constraints: global chip shortages, vendor prioritization conflicts, and the

growing divergence between general-purpose GPUs and Tesla’s bespoke AI workloads. Musk has repeatedly emphasized that Tesla’s roadmap could *fail* if chip supply does not scale with need. In this light, TeraFab is *defensive infrastructure*—a necessity born from systemic fragility.

On the other hand, the scope of Musk’s proposal suggests more than defensive thinking. A fab capable of producing 100–200 billion chips per year is not just for internal use—it is a sovereign-scale platform, potentially capable of serving other AI players, licensing hardware, or dictating the standards of embodied AI. Combined with Tesla’s autonomous fleets, xAI’s model development, and Optimus’s real-world deployment, TeraFab would anchor a vertically unified *intelligence empire*—one in which Tesla controls the substrate of thought, from transistor to behavior.

The tension, then, is between existential need and imperial reach. TeraFab may be both: a survival strategy that, once achieved, enables strategic dominance far beyond its original scope.

10.3 The Future of AI Is Made of Silicon—and Ambition

In closing, TeraFab reminds us that the future of artificial intelligence is not solely written in code or dreamed in algorithmic architectures. It is *etched into matter*, constrained by physics, governed by logistics, and shaped by those bold enough to make their own tools when the world will not provide them. Musk’s proposal collapses the abstraction of AI into the rawness of foundries, wafers, and engineered electrons.

Whether TeraFab becomes a functional facility, a global fab network, or merely a forcing function that pushes partners and governments to accelerate their own chip roadmaps, it has already done something important: it has reframed compute as destiny. It has reminded us that intelligence does not emerge from nothing—it is built, layer by layer, in factories, under microscopes, inside cleanrooms, by people and systems who understand not just the software of thought, but the hardware of power.

In this framing, TeraFab is not the end of Tesla’s evolution. It may be the beginning of something larger: a world where *those who control the silicon control the AI—and those who control the AI, shape the world*.

11. References

11.1 Tesla earnings calls and shareholder transcripts

- Tesla, Inc. Tesla Releases Fourth Quarter and Full Year 2025 Financial Results (Investor Relations press release). January 28, 2026.
- Tesla, Inc. Q4 2025 Quarterly Update Deck (PDF). January 28, 2026.
- Tesla, Inc. Q4 2025 and Full Year 2025 Financial Results and Q&A Webcast (Investor Relations landing page). January 2026.
- Tesla, Inc. Tesla Fourth Quarter 2025 Production, Deliveries & Deployments (Investor Relations press release). January 2, 2026.
- Tesla, Inc. (TSLA) Q4 2025 Earnings Call Transcript. Seeking Alpha (SA Transcripts). January 28, 2026.
- Tesla (TSLA) Q4 2025 Earnings Call Transcript. The Motley Fool. January 28, 2026.
- Earnings call transcript: Tesla Q4 2025 beats forecasts; stock edges up. Investing.com. January 28–29, 2026.
- Tesla 2025 Annual Shareholder Meeting (public livestream video). YouTube. November 2025.
- Tesla’s 2025 Annual Shareholder Meeting (Transcript). SingjuPost. November 7, 2025.
- Tesla, Inc. Form 8-K (Annual Meeting voting outcomes and related disclosures). U.S. SEC EDGAR. November 6, 2025.

11.2 Semiconductor manufacturing literature and benchmarks

- International Roadmap for Devices and Systems (IRDS). IRDS Roadmap (latest available chapters and summaries).
- ASML. EUV lithography technology overview and system context (official materials).
- JEDEC Solid State Technology Association. High Bandwidth Memory (HBM) standard documentation (e.g., HBM3 family).
- TSMC. Advanced packaging overview (including CoWoS and related offerings).
- Mahajan, R. Advanced Packaging Architectures for Heterogeneous Integration (industry technical presentation). 2021.

- Boston Consulting Group (BCG). Emerging Resilience in the Semiconductor Supply Chain. May 8, 2024.
- McKinsey & Company. A roadmap for US semiconductor fab construction (construction, workforce, schedule, and cost constraints). January 27, 2023.
- Tesla Q4 2025 Update Deck sections on in-house inference silicon (AI4/AI5/AI6), memory scaling, and deployment timelines (for Tesla-specific chip targets). January 28, 2026.

11.3 Industry coverage from Reuters, Bloomberg, Tom's Hardware, Wccftch

- Ghosh, S. (Reuters). Musk plans Tesla mega AI chip fab, mulls potential Intel partnership. November 6–7, 2025.
- Reuters. Tesla invests \$2 billion in Musk's xAI and reiterates Cybercab production starts this year (includes chip-plant / memory-bottleneck remarks). January 28, 2026.
- Ludlow, E.; Hull, D. (Bloomberg News). Musk Says Tesla Needs to Build "TeraFab" to Manufacture Chips. January 28–29, 2026.
- Ludlow, E.; Hull, D. (Bloomberg Law / Bloomberg News). Musk Says Tesla Needs to Build Its Own Factory to Make Chips. January 29, 2026.
- Carlson, K. (Bloomberg News). Tesla Plots \$20 Billion Splurge to Support Musk's AI Future (earnings context for capex and AI infrastructure). January 28, 2026.
- Tom's Hardware. Elon Musk wants foundry partners to build 100–200 billion AI chips per year (scale claim and feasibility framing). November 2025.
- Zuhair, M. (Wccftch). Elon Musk's Mega "TeraFab" Project Will Combine Memory, Chips, and Packaging in One Facility. January 29, 2026.
- Wccftch. Elon Musk claims Tesla will build a 2nm "TeraFab" (follow-on reporting and technical skepticism). January 7, 2026.
- Business Insider. Tesla Q4 earnings recap (AI pivot, capex, and "TeraFab" mention in broader narrative). January 2026.
- Investor's Business Daily. Tesla Q4 2025 earnings coverage and capex/AI manufacturing notes. January 2026.

11.4 Expert commentary on AI hardware scaling and autonomy infrastructure

- Stanford HAI. AI Index Report (latest edition and underlying trend tables on compute, deployment, and ecosystem scaling).

- International Energy Agency (IEA). Energy and AI (report PDF). 2025/2026 edition.
- International Energy Agency (IEA). AI is set to drive surging electricity demand from data centres (news release). April 10, 2025.
- International Energy Agency (IEA). Energy demand from AI (report section / analysis page).
- Reuters. Nvidia CEO comments on the difficulty of building leading-edge manufacturing capacity (contextualizing “extremely hard” foundry capability). November 7, 2025.
- European Commission energy portal. In focus: Data centres – an energy-hungry challenge (IEA-aligned demand figures and context). November 17, 2025.
- Patterson, A. (EE Times). Rapidus, Siemens join 2-nm alliance targeting 2027 production (industry timeline and ecosystem coordination). June 26, 2025.
- IBM Newsroom. Rapidus and IBM expand collaboration to chiplet packaging technology for 2nm generation semiconductors. June 3, 2024.

11.5 Strategic papers on vertical integration and foundry economics

- Coase, R. H. The Nature of the Firm. *Economica*, 1937.
- Williamson, O. E. Markets and Hierarchies. 1975.
- Teece, D. J. Profiting from Technological Innovation: Implications for Integration, Collaboration, Licensing and Public Policy. *Research Policy*, 1986.
- Stigler, G. J. The Division of Labor Is Limited by the Extent of the Market. *Journal of Political Economy*, 1951.
- Kuan, J.; West, J. Technical, Organizational and Business Model Modularity: Evidence from Fabless Semiconductor Design. 2021.
- Kim, K. Economics of R&D specialization in the semiconductor industry (doctoral dissertation / academic study). 2018.
- BCG. Emerging Resilience in the Semiconductor Supply Chain (capex projections and geographic rebalancing relevant to “build vs buy” decisions). May 8, 2024.
- McKinsey & Company. Semiconductor fab construction challenges in the United States (execution constraints that shape vertical-integration feasibility). January 27, 2023.

The author has published ~ 5,000 AI-assisted papers at:
<https://www.researchgate.net/profile/Douglas-Youvan/research> . Google Scholar has indexed these preprints, and if you want a particular reference, the AI mode of a Google search (Gemini) has efficiently catalogued these preprints.

