

NVIDIA SK Group \$500 Billion AI Deal Explained (2026)

Published July 29, 2026 35 min read



NVIDIA SK Group \$500 Billion AI Deal Explained (2026)

Executive Summary

On July 24, 2026, NVIDIA and South Korea's SK Group announced plans for a comprehensive partnership representing more than \$500 billion in combined business, spanning [AI data center construction](#) and next-generation memory supply (Source: [nvidianews.nvidia.com](#)) (Source: [www.reuters.com](#)). The two sides signed letters of intent at the K-AI Summit in San Francisco alongside South Korean President Lee Jae Myung; the announcement does not publish the terms needed to determine whether, or to what extent, particular commitments are binding (Source: [www.datacenterdynamics.com](#)). NVIDIA CEO Jensen Huang framed the scale of the announcement directly, telling President Lee that "our two companies will enter into business partnerships that will represent over \$500 billion of business together" (Source: [m.koreaherald.com](#)).

The initiative has two main pillars. First, **SK Telecom** will build a **2-gigawatt** AI data center in South Korea using NVIDIA's **DSX** full-stack reference architecture and **Vera Rubin** accelerated computing, with the first facility scheduled to come online in 2027 (Source: [nvidianews.nvidia.com](#)). Second, **SK hynix** entered a long-term AI memory partnership under which NVIDIA secures supply of next-generation memory, principally **HBM4** (High Bandwidth Memory 4), while SK hynix and NVIDIA jointly optimize future memory designs (Source: [nvidianews.nvidia.com](#)). This memory agreement builds on an earlier multiyear technology partnership the two companies signed on June 7, 2026 (Source: [nvidianews.nvidia.com](#)), which SK hynix's newsroom describes as a "follow-up measure to solidify their previous long-term technical partnership" (Source: [news.skhynix.com](#)).

HBM4 is the technical center of gravity for the deal. Compared with the current production standard, HBM3e, which delivers "over 1.2 TB/s per stack through a 1024-bit wide interface and 16 independent channels," HBM4 doubles the interface to "2048 bits and 32 independent channels to deliver over 2.0 TB/s per stack, up to 3.3 TB/s in advanced configurations" (Source: [blogs.sw.siemens.com](#)). The standard was formalized by the JEDEC Solid State Technology Association as JESD270-4, with a later revision, JESD270-4A, published in December 2025 (Source: [www.jedec.org](#)). Market tracker TrendForce projects SK hynix will hold roughly a 50 percent share of global [HBM bit output](#) in 2026, down from 59 percent in 2025, as Samsung's share climbs to 28 percent (Source: [www.trendforce.com](#)).

The \$500 billion figure sits inside a much larger wave of Korea-US AI investment. South Korean presidential policy chief Kim Yong-beom disclosed that Korean and US technology firms had agreed to pursue \$950 billion in combined semiconductor, data center and physical AI partnerships, including SK Group's \$750 billion in long-term memory supply commitments and a separate \$200 billion Samsung-Broadcom memory and foundry agreement (Source: www.koreaherald.com) (Source: news.samsung.com). Parallel deals with Naver (a \$1 billion NVIDIA investment plus up to \$9 billion from Brookfield), Hyundai Motor Group, LG Group and Doosan Group round out a broader push to embed NVIDIA's compute stack across Korea's industrial base (Source: nvidianews.nvidia.com).

Investor reaction has been volatile rather than uniformly positive. SK hynix, which listed American Depositary Shares on the Nasdaq in July 2026 and disclosed in its SEC prospectus that it expected "net proceeds that we will receive in the offering will be approximately US\$28.0 billion" (Source: www.sec.gov), saw its shares "dropped 14.7%" on July 28 as part of a broader semiconductor selloff that Reuters attributed primarily to "concerns over AI infrastructure financing and intensifying competition from China," rather than to the NVIDIA-SK Group news alone (Source: www.reuters.com) (Source: www.reuters.com). For GPU buyers, the announcement is a directional signal that large-scale AI projects may place greater value on long-term memory-supply visibility. It does not establish that memory is the binding constraint for every GPU program through 2027, or that buyers without comparable agreements cannot obtain allocation.

Introduction and Background

The term "the NVIDIA SK Group \$500 billion AI deal" refers to a package of letters of intent that NVIDIA and South Korea's SK Group signed on July 24, 2026, at an AI summit in San Francisco (Source: nvidianews.nvidia.com). It is not a single contract with a fixed price tag. It is a bundle of commitments spanning **SK Telecom** (SK Group's telecommunications and cloud arm), **SK hynix** (the world's largest supplier of high bandwidth memory), and adjacent agreements involving **Naver**, **Brookfield**, **Samsung**, **Hyundai Motor Group**, **LG Group** and **Doosan Group**. Understanding why this deal matters requires separating three layers: the compute layer (data centers built on NVIDIA's Vera Rubin platform), the memory layer (SK hynix's HBM4 supply to NVIDIA), and the national policy layer (South Korea's effort to position itself as a global AI infrastructure hub).

The announcement did not emerge in isolation. It follows a decades-long technology relationship between SK Group and NVIDIA, and more immediately builds on a June 7, 2026 multiyear technology partnership between NVIDIA and SK hynix covering memory codevelopment for NVIDIA's Vera Rubin AI supercomputers, **Vera CPUs**, RTX Spark-powered PCs and Jetson Thor robotic computing platforms (Source: nvidianews.nvidia.com). It also follows an October 2025 commitment, made during the APEC CEO Summit in Gyeongju, in which NVIDIA pledged to prioritize the supply of 260,000 advanced AI GPUs to South Korean government agencies and companies over five years, described by NVIDIA's own blog as support for "more than a quarter-million NVIDIA GPUs" across sovereign clouds and industrial AI factories (Source: blogs.nvidia.com). The July 24 announcement is best understood as the point at which those earlier, smaller commitments were consolidated and scaled into a headline figure exceeding \$500 billion.

South Korea's motivation is explicit in official statements. SK Group Chairman **Chey Tae-won** said that "by leveraging SK hynix's AI memory and SK Telecom's AI infrastructure capabilities, SK will collaborate with NVIDIA to build a world-class AI factory," describing the goal as helping Korea "become a global hub that drives AI innovation" rather than remain simply an adopter of AI technology built elsewhere (Source: nvidianews.nvidia.com). Jensen Huang, in turn, described South Korea as possessing "world-class networks and data centers, leadership in chip technology and vast industrial scale" (Source: nvidianews.nvidia.com). For NVIDIA, the deal is fundamentally about memory security: the company is the world's largest buyer of HBM, and the agreement gives it, in the words of NVIDIA enterprise vice president Raj Mirpuri, an opportunity to "codevelop... the next-generation SK Hynix AI memory" so it can "secure a stable supply of HBM memory" (Source: www.cnbc.com). SK Group, for its part, is one of South Korea's largest conglomerates, encompassing more than 175 affiliate companies and over 100,000 employees worldwide across artificial intelligence, semiconductors, energy and life sciences (Source: nvidianews.nvidia.com).

The rest of this report unpacks each component of the deal, defines the HBM4 memory technology at its core and how it differs from the current HBM3e standard, examines what the arrangement means practically for organizations buying GPU capacity, walks through the surrounding financial and market data, profiles the individual deals that make up the broader Korea-NVIDIA push, and considers the implications for the AI supply chain heading into 2027 and beyond.

Anatomy of the Deal: What the \$500 Billion Actually Covers

The \$500-billion-plus figure is a forward-looking estimate of "business together" over an unspecified multi-year horizon, not a signed, itemized contract value. Jensen Huang told reporters he did not provide details on "how the figure was calculated or over what period the partnerships would be carried out," a gap the Korea Herald noted directly in its coverage of the announcement (Source: m.koreaherald.com). That ambiguity matters for readers

trying to size the deal against other AI infrastructure announcements: the number should be read as a strategic target for combined compute purchases, memory purchases and infrastructure investment, rather than a cash commitment sitting on any single balance sheet.

SK Telecom: The 2-Gigawatt Vera Rubin DSX AI Factory

The compute half of the deal centers on SK Telecom, which will build a data center campus with **up to 2 gigawatts** of power capacity in South Korea. This buildout uses NVIDIA's **DSX** platform, described as "a full-stack reference architecture including software, hardware, and operations," combined with NVIDIA's next-generation **Vera Rubin** accelerated computing systems (Source: www.datacenterdynamics.com). The 2GW figure was first disclosed in a separate June 2026 SK Telecom announcement and was folded into the July 24 package as its flagship compute commitment (Source: www.datacenterdynamics.com). The first phase of the AI factory is slated to come online in 2027 (Source: nvidianews.nvidia.com), and will support what NVIDIA calls "sovereign, physical, agentic and enterprise AI services," initially serving South Korea before expanding to other regions in the Asia-Pacific (Source: nvidianews.nvidia.com).

A 2GW facility is exceptionally large by current data center standards: for context, NVIDIA's own APEC-era commitments to South Korea described individual industrial AI factories from Samsung, SK Group and Hyundai Motor Group at up to 50,000 GPUs each, with Naver deploying more than 60,000 GPUs (Source: blogs.nvidia.com). A 2GW campus implies a buildout an order of magnitude larger, consistent with hundreds of thousands of GPUs, which is why CNBC characterized the target power draw as indicating "a massive buildout with hundreds of thousands of graphics processing units" (Source: www.cnn.com).

SK hynix: The Long-Term AI Memory Partnership

The second pillar is a long-term AI memory partnership between NVIDIA and SK hynix, under which the two companies will "codevelop and optimize next-generation AI memory solutions, including HBM, to meet evolving infrastructure demands ranging from large language model training to agentic AI and physical AI" (Source: news.skhynix.com). Under the terms described by NVIDIA, this gives NVIDIA access to "a stable supply of next-generation AI memory," in exchange for SK hynix expanding "the foundation for growth" of its own business (Source: nvidianews.nvidia.com). SK Group Chairman Chey Tae-won framed the strategic logic in terms of intelligence production, saying that "in the AI era, competitiveness depends not just on how effectively AI is utilized, but on how much intelligence we can produce" (Source: news.skhynix.com).

This memory pillar did not appear out of nowhere. On June 7, 2026, NVIDIA and SK hynix had already announced "a multiyear technology partnership to advance next-generation memory for the global AI factory buildout and accelerate semiconductor design and manufacturing," under which SK hynix would use NVIDIA's CUDA-X libraries and PhysicsNeMo framework to speed chip design, and NVIDIA Omniverse and cuOpt to build autonomous fab digital twins (Source: nvidianews.nvidia.com). SK hynix's stock fell 7.68 percent on the Korea Exchange the trading day following that earlier announcement, closing at 1,911,000 Korean won (Source: www.rttnews.com), an early sign of the volatility that would recur around each subsequent Korea-NVIDIA milestone.

Beyond SK Group: Samsung, Naver, Hyundai, LG and Doosan

The July 24 summit produced a cluster of parallel agreements that, while not part of the SK Group deal itself, help explain its scale and context. Samsung Electronics and Broadcom separately signed a memorandum of understanding to expand collaboration "across memory and foundry technologies," an agreement the companies expect to be "estimated at more than \$200 billion across memory and foundry over the next five years through 2030," covering HBM supply for Broadcom's AI accelerators and Samsung's 2-nanometer foundry process (Source: news.samsung.com).

NVIDIA, Naver and Brookfield jointly announced plans to expand the NVIDIA DSX AI factory buildout at Naver's GAK Sejong data center "from 55 megawatts to 200 megawatts by 2028," with NVIDIA planning to invest \$1 billion into Naver Corp and Brookfield entering "a nonbinding term sheet to fund up to \$9 billion," with Naver covering the remainder (Source: nvidianews.nvidia.com) (Source: nvidianews.nvidia.com). Reuters independently confirmed this three-way arrangement, reporting that "Nvidia, Naver and Brookfield plan to expand Naver's AI data center in South Korea" (Source: www.reuters.com).

The same San Francisco summit produced further deals outside the memory and data center space. During his meeting with President Lee, Huang confirmed that "SK Hynix, Samsung and Nvidia are partnering to advance chip design and, of course, memory design," that NVIDIA was working with Hyundai Motor Group to "build autonomous Genesis vehicles as well as robotic systems," and that NVIDIA was collaborating with LG Group "on power-generation systems and robots for factories and homes" (Source: m.koreaherald.com) (Source: m.koreaherald.com). Separately, Doosan

Group expanded its own collaboration with NVIDIA "across physical AI, robotics and AI factory infrastructure," spanning Doosan Robotics, Doosan Bobcat, Doosan Enerbility and Doosan's electronics materials business, with Doosan Robotics developing an "Agentic Robot OS" using NVIDIA's Isaac Sim, Isaac Lab, Cosmos world models and Jetson Thor edge computing (Source: www.doosan.com).

Taken together, presidential policy chief Kim Yong-beom summarized the scale of the announcements by saying that "Korean companies and global big tech firms have agreed to pursue cooperation worth a combined \$950 billion, or 1,375 trillion won," of which SK Group's commitments alone total \$750 billion in long-term memory-chip supply partnerships with NVIDIA and other technology companies (Source: www.koreaherald.com). Kim was careful to characterize the nature of these commitments: "our companies have won long-term supply contracts, which will allow them to plan production with far greater stability," clarifying that the semiconductor cooperation consists of advance purchase orders rather than direct capital investment (Source: www.koreaherald.com). On the data center side, Korean and US firms separately agreed to pursue AI data center investment partnerships totaling "a combined power capacity of about 5 gigawatts and roughly 2 million graphics processing units," with Anthropic alone accounting for 1 gigawatt of that figure (Source: www.koreaherald.com).

Table 1 below summarizes the major components of the July 2026 Korea-NVIDIA announcement package and their disclosed scale.

COMPONENT	PARTIES	DISCLOSED SCALE	TIMELINE
AI factory compute	SK Telecom, NVIDIA	Up to 2 gigawatts, NVIDIA DSX and Vera Rubin (Source: www.datacenterdynamics.com)	First phase online 2027
Long-term memory supply	SK hynix, NVIDIA	Part of \$750 billion SK Group memory commitment (Source: www.koreaherald.com)	Multi-year, through 2030 window referenced across deals
Memory and foundry MOU	Samsung, Broadcom	Estimated over \$200 billion (Source: news.samsung.com)	Through 2030
AI factory expansion	Naver, NVIDIA, Brookfield	55 MW to 200 MW by 2028, \$1 billion NVIDIA investment, up to \$9 billion Brookfield financing (Source: nvidianews.nvidia.com)	By 2028
Total Korea-US semiconductor cooperation	Multiple Korean groups, US tech firms	\$950 billion combined (Source: www.koreaherald.com)	Over next five years

This table shows that the headline \$500 billion figure for NVIDIA and SK Group is one large component of an even larger \$950 billion wave of announced Korea-US technology cooperation, and that the individual pieces (AI factories, memory supply contracts and foundry MOUs) have different disclosed structures and execution conditions. Readers comparing this deal to other AI infrastructure announcements should be careful to note which of these structures underlies any given dollar figure, since a "\$500 billion partnership" of advance memory orders is a materially different commitment than \$500 billion in committed capital expenditure.

HBM4 versus HBM3e: The Memory Technology at the Center of the Deal

The following sections explain why HBM matters to this partnership and compare HBM4 with the current HBM3e generation.

Why HBM Supply Matters to AI Accelerator Output

The reason a memory supply deal can carry a headline figure larger than many GPU supply deals is that High Bandwidth Memory (HBM), a type of DRAM (Dynamic Random Access Memory) that stacks multiple memory dies vertically using Through-Silicon Vias (TSVs) to sit adjacent to a GPU on a shared silicon interposer, is an important input to AI accelerator output. Its availability should be assessed alongside GPU dies, advanced packaging, power, networking and vendor-specific allocation terms; this announcement does not establish HBM as the binding constraint for every accelerator program. SK hynix's own SEC filing describes HBM plainly as "a high-performance memory product that vertically interconnects multiple DRAM chips and increases data processing speed relative to traditional DRAM products" (Source: www.sec.gov). As the Siemens EDA (Electronic

Design Automation) blog explains, "without sufficient bandwidth, even the most advanced GPUs sit idle," meaning "modern AI system performance is increasingly memory bandwidth bound" (Source: blogs.sw.siemens.com). NVIDIA's own Vera Rubin platform illustrates the trend: TechInsights notes that the accompanying Vera CPU increases "CPU-to-GPU connectivity to 1.8 TB/s, enabling the GPU to offload KV cache memory to DRAM and reuse it later," a design response to the growing memory demands of large-context AI inference (Source: www.techinsights.com).

Technical Specifications Compared

HBM3e is the current production-grade standard. According to Siemens EDA, it is "the fifth-generation high bandwidth memory architecture, delivering over 1.2 TB/s per stack through a 1024-bit wide interface and 16 independent channels," is already deployed in NVIDIA's H100 and H200 systems and AMD's MI300 series, and is in volume production from SK hynix, Samsung and Micron (Source: blogs.sw.siemens.com).

HBM4 represents what Siemens describes as "not an incremental speed bump to HBM3e, but a redesign of the memory interface," doubling the interface width to 2048 bits and 32 independent channels, with pin speeds extending to 12.8 gigabits per second (Gb/s), and Samsung having demonstrated speeds of up to 13 Gb/s (Source: blogs.sw.siemens.com). Critically, HBM4 is "not backward compatible with HBM3 or HBM3e controllers," meaning any GPU program targeting HBM4, including NVIDIA's Vera Rubin platform, is effectively starting a new design cycle rather than reusing existing memory controller intellectual property (Source: blogs.sw.siemens.com). The standard itself was formalized by JEDEC, the global microelectronics standards body, as JESD270-4, with the underlying document describing an HBM4 DRAM that "is tightly coupled to the host compute die with a distributed interface," organized into independent channels that need not be synchronous with one another (Source: www.jedec.org); a revised version, JESD270-4A (version 1.1), was published in December 2025 (Source: www.jedec.org).

Table 2 below summarizes the specification differences documented by Siemens EDA's semiconductor packaging engineering team.

SPECIFICATION	HBM3E	HBM4
Interface width	1024-bit	2048-bit (Source: blogs.sw.siemens.com)(https://blogs.sw.siemens.com/semiconductor-packaging/2026/04/24/hbm3e-hbm4-ic-design-guide/#:~:text=Interface%20width%20)
Independent channels	16	32 (Source: blogs.sw.siemens.com)(https://blogs.sw.siemens.com/semiconductor-packaging/2026/04/24/hbm3e-hbm4-ic-design-guide/#:~:text=Independent%20channels%20)
Bandwidth per stack	Over 1.2 TB/s, up to 1.33 TB/s	Over 2.0 TB/s, up to 3.3 TB/s in advanced configurations (Source: blogs.sw.siemens.com)(https://blogs.sw.siemens.com/semiconductor-packaging/2026/04/24/hbm3e-hbm4-ic-design-guide/#:~:text=Bandwidth%20per%20stack%20)
Capacity per stack	Up to 36GB	Up to 64GB (Source: blogs.sw.siemens.com)(https://blogs.sw.siemens.com/semiconductor-packaging/2026/04/24/hbm3e-hbm4-ic-design-guide/#:~:text=Capacity%20per%20stack%20)
Core voltage	1.1V	1.05V (Source: blogs.sw.siemens.com)(https://blogs.sw.siemens.com/semiconductor-packaging/2026/04/24/hbm3e-hbm4-ic-design-guide/#:~:text=Core%20voltage%20)
Controller compatibility	Backward compatible with HBM3	Not backward compatible with HBM3 or HBM3e (Source: blogs.sw.siemens.com)
Production status	In volume production	Ramping in 2026 across Samsung, SK hynix and Micron (Source: blogs.sw.siemens.com)(https://blogs.sw.siemens.com/semiconductor-packaging/2026/04/24/hbm3e-hbm4-ic-design-guide/#:~:text=Production%20status%20)

Reading this table alongside NVIDIA's own platform targets shows why the SK hynix deal matters operationally: NVIDIA's demands for Vera Rubin's HBM4 reportedly exceed the JEDEC baseline. TrendForce, citing Korean industry outlet Hankyung, reported that "NVIDIA is demanding HBM4 data rates exceeding 10Gb/s, well above the 8Gb/s standard set by JEDEC," with qualification tests being conducted in two tiers at 10 Gb/s and 11 Gb/s (Source: www.trendforce.com). NVIDIA's Vera Rubin GPU is expected to "pack 16 stacks for 576 GB" of HBM4 capacity, ahead of AMD's competing MI450 accelerator, which "tops at 432 GB" (Source: www.trendforce.com).

Supply Chain and Market Share

HBM4 supply is concentrated among three vendors, with SK hynix in the lead. TrendForce's projection has SK hynix leading "global HBM bit output with a 50% share" in 2026, a decline from 59 percent in 2025, while "Samsung's portion climbs from 20% to 28%," with Micron holding the remainder (Source: www.trendforce.com) (Source: www.trendforce.com). A separate TrendForce report on initial HBM4 allocation for NVIDIA's Vera Rubin specifically put SK hynix's share "in the mid-50% range, while Samsung holds the mid-20% range and Micron around 20%" (Source: www.trendforce.com). Notably, Samsung was first to actual shipment: TrendForce reported that "Samsung Electronics kicked off HBM4 shipments in February, while SK hynix has yet to announce deliveries," suggesting the competitive dynamic between the two Korean suppliers is not purely about volume share but also about qualification speed and data-rate performance (Source: www.trendforce.com). CNBC, citing TrendForce analyst Ellie Wang, described SK hynix's HBM leadership as having "positioned SK Hynix as one of the biggest beneficiaries of the rapid growth in AI infra" (Source: www.cnbc.com).

Implementation Guidance: What This Means for GPU Buyers and Enterprise AI Planners

Organizations planning AI infrastructure purchases in 2026 and 2027 should treat this deal as confirmation of several practical realities rather than as a signal that supply constraints are easing.

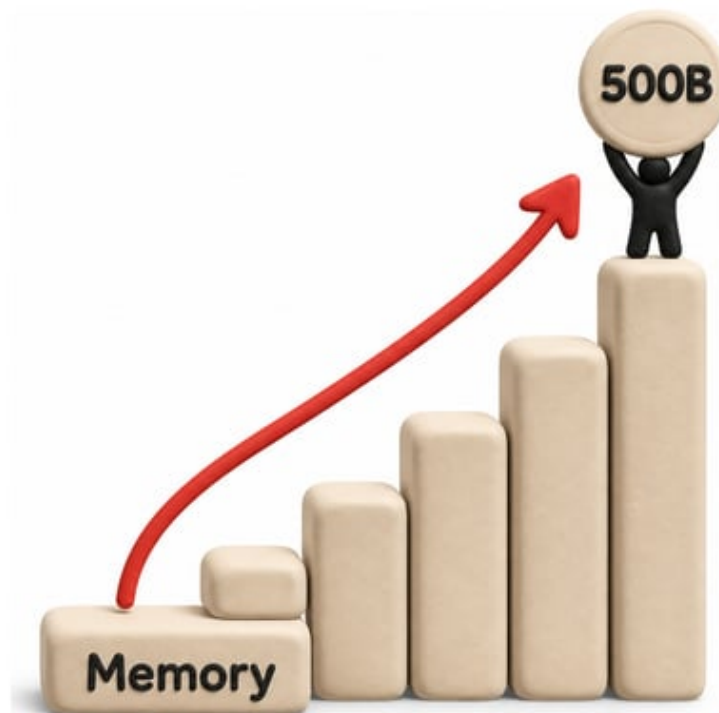
- **Long-term supply contracts are becoming the norm, not the exception.** South Korea's presidential policy chief characterized the underlying SK Group commitments explicitly as advance purchase orders: "these are advance contracts, advance orders for memory chips on that scale. This is not an investment" (Source: www.koreaherald.com). SK hynix's own leadership confirmed the same pattern extends well beyond NVIDIA: the company disclosed it had "concluded talks on around 10 such deals" with other major customers, typically running five years and including deposits as "financial safeguards" to ensure implementation (Source: www.reuters.com). Buyers without similar multi-year commitments in place risk being at the back of the allocation queue.
- **Memory, not GPU compute die availability, is the constraint to plan around.** CNBC's coverage of SK hynix's Nasdaq debut noted that memory vendors, including SK hynix, Micron and Samsung, are "implementing long-term contracts for memory, using their market power to lock in prices and orders years into the future," a shift from the historical practice of selling supply on a quarterly or annual basis (Source: www.cnbc.com).
- **HBM4 is not a drop-in replacement for HBM3e.** Because HBM4 is not backward compatible with HBM3e controllers, buyers planning platform transitions should expect vendors to require new system designs, new PHY (physical layer) intellectual property, and new qualification cycles rather than a simple firmware or memory-module swap (Source: blogs.sw.siemens.com).
- **Expect continued price and lead-time pressure through at least 2027.** SK hynix's own disclosed capital plans, up to \$720 billion in Korean facility expansion, including a \$390 billion fabrication cluster in Yongin, indicate the supply response to AI memory demand will take years to materialize (Source: www.cnbc.com). Industry veterans caution that expansion timelines are long: Counterpoint research director MS Hwang noted that "the earliest time frame that they can bring out manufactured wafers is end of 2027," even with accelerated construction schedules (Source: www.cnbc.com).
- **Diversify supplier relationships where possible.** With HBM4 output concentrated between SK hynix, Samsung and Micron, and NVIDIA already demonstrating a willingness to qualify multiple vendors at different data-rate tiers (Source: www.trendforce.com), buyers relying on a single GPU platform generation should track qualification announcements from all three memory vendors rather than assuming uniform availability.
- **Weigh the volatility risk priced into memory-linked equities against long-term demand.** Futurum Group CEO Daniel Newman cautioned that while memory suppliers may be a reasonable wager if AI demand persists, "this is how memory always acts in any megacycle or supercycle," warning that historically "it always crashes hard" when cycles turn (Source: www.cnbc.com).
- **Watch for signs that long-term contracts are capping near-term pricing upside.** Analysts covering SK hynix's own results said its adoption of long-term deals "could limit upside in memory prices," a factor that contributed directly to the company's second-quarter 2026 earnings falling short of forecasts despite record profit (Source: www.reuters.com).

Community reaction on technology forums has been more skeptical than official statements. One widely upvoted reader comment on a Tom's Hardware report on the deal characterized the letter-of-intent structure bluntly: it "really don't mean much until the money starts making an impact to profits," calling it "some gentleman agreement to try and stir up the market with a big number" (Source: www.tomshardware.com). Datacenter Dynamics' own reporting supports a cautious reading, noting plainly that "a letter of intent typically sets out initial terms for an agreement and is

usually non-binding" (Source: www.datacenterdynamics.com). Buyers should therefore treat the \$500 billion figure as a strong directional signal about where compute and memory supply is being allocated, rather than as evidence that specific capacity is contractually guaranteed to any given third party.

Data Analysis and Evidence

The scale of NVIDIA's own financial results provides context for why memory security commands a \$500 billion strategic commitment. NVIDIA reported record revenue of \$68.1 billion for the fourth quarter of fiscal 2026 (ended January 25, 2026), up 20 percent from the previous quarter and 73 percent year over year, with Data Center revenue reaching a record \$62.3 billion, up 75 percent year over year (Source: nvidianews.nvidia.com). Full fiscal 2026 revenue reached \$215.9 billion, up 65 percent from \$130.5 billion in fiscal 2025 (Source: nvidianews.nvidia.com). Momentum continued into the following quarter: NVIDIA reported "record revenue for the first quarter ended April 26, 2026, of \$81.6 billion, up 20% from the previous quarter and up 85% from a year ago," with Data Center revenue of \$75.2 billion, up 92 percent year over year (Source: nvidianews.nvidia.com).



On the supplier side, SK hynix's own regulatory disclosures show the scale of the memory boom underpinning the deal. In the prospectus it filed with the U.S. Securities and Exchange Commission ahead of its Nasdaq listing, SK hynix reported revenue of "97,147 billion (US\$63,765 million)" for full-year 2025, up from "66,193 billion in 2024" and "32,766 billion in 2023," alongside profit for 2025 of "42,948 billion (US\$28,190 million)" (Source: www.sec.gov) (Source: www.sec.gov). Revenue accelerated further into 2026: the filing shows first-quarter 2026 revenue of "52,576 billion (US\$34,510 million)," nearly triple the "17,639 billion in the first quarter of 2025" (Source: www.sec.gov). More than three-quarters of SK hynix's revenue now comes from RAM products, including HBM, according to CNBC's review of the company's business (Source: www.cnbc.com). To fund expansion, SK hynix's own filing states it expected "net proceeds that we will receive in the offering will be approximately US\$28.0 billion" (Source: www.sec.gov), and CNBC separately reported the company disclosed plans to spend "up to \$720 billion on expanding facilities to meet memory demand for AI" within South Korea, including a \$390 billion fabrication cluster in Yongin with four fabs now targeted for completion by 2033, an acceleration of more than a decade versus the prior schedule (Source: www.cnbc.com). The company also plans to spend approximately \$7.8 billion on new extreme ultraviolet (EUV) lithography machines by the end of 2027, equipment that costs up to \$400 million per unit and is manufactured exclusively by ASML of the Netherlands (Source: www.cnbc.com).

Market reaction to the deal itself, and to the broader AI infrastructure buildout, has been sharp, occasionally negative, and multi-causal. On July 28, 2026, four days after the SK Group announcement, "SK Hynix closed 14.65% lower, while Samsung Electronics lost more than 13%" in a broad semiconductor selloff that also affected Micron, Seagate, Western Digital and Sandisk in the United States, plus Tokyo Electron, Advantest, SoftBank

Group and Kioxia in Japan (Source: www.cnbc.com). Reuters measured the same session slightly differently, reporting SK Hynix "dropped 14.7%" while Samsung Electronics "closed 13.4% lower, notching their worst one-day fall in almost two decades" (Source: www.reuters.com). Seeking Alpha attributed part of the move directly to the NVIDIA-SK Group news, describing how "shares of memory and AI-related stocks were largely in the red on Monday after a \$500B AI infrastructure and chip collaboration between Nvidia and South Korea's SK Group, raising potential investor concerns over high capital spending in this space and returns on" investment (Source: seekingalpha.com). Reuters offered a more multi-causal reading of the sharper Tuesday slide, attributing it to "concerns over AI infrastructure financing and intensifying competition from China," including reports that Chinese firms were developing domestic deep ultraviolet lithography equipment, a rival stock listing from Chinese memory maker CXMT, and a separate report that NVIDIA could provide a roughly \$250 billion financing backstop for an OpenAI data center project (Source: www.reuters.com) (Source: www.reuters.com). CNBC's own analysis cited an additional factor, a Standard Chartered view that memory prices could peak in 2027 (Source: www.cnbc.com). Acadian Asset Management's Owen Lamont summarized the underlying uncertainty facing investors: "no one has any idea how this AI process is going to affect our economy, and so I think it's going to be rocky no matter what" (Source: www.cnbc.com). Standard Chartered's Sundeep Gantori offered a more constructive framing of the same volatility, arguing "the market opportunity remains sufficiently large for multiple players to benefit and coexist" (Source: www.cnbc.com).

Table 3 below places NVIDIA's recent quarterly results alongside SK hynix's disclosed financial position to illustrate the scale mismatch between quarterly AI revenue and multi-year memory infrastructure spending.

METRIC	VALUE	PERIOD
NVIDIA Q4 FY2026 total revenue	\$68.1 billion, up 73% year over year (Source: nvidianews.nvidia.com)	Quarter ended Jan 25, 2026
NVIDIA Q4 FY2026 Data Center revenue	\$62.3 billion, up 75% year over year (Source: nvidianews.nvidia.com)	Quarter ended Jan 25, 2026
NVIDIA Q1 FY2027 total revenue	\$81.6 billion, up 85% year over year (Source: nvidianews.nvidia.com)	Quarter ended Apr 26, 2026
SK hynix 2025 annual revenue (SEC filing)	US\$63,765 million (Source: www.sec.gov)	Full year 2025
SK hynix Q1 2026 revenue (SEC filing)	US\$34,510 million (Source: www.sec.gov)	Quarter ended Mar 31, 2026
SK hynix Korea facility capex	Up to \$720 billion (Source: www.cnbc.com)	Multi-year, through 2033
SK hynix Nasdaq ADS net proceeds (SEC estimate)	Approximately US\$28.0 billion (Source: www.sec.gov)	July 2026 listing

The pattern in this data is that NVIDIA's Data Center revenue is compounding at 75 to 92 percent year over year even before the SK Group deal's compute buildout comes online in 2027, while SK hynix's own multi-year capital commitments (\$720 billion domestically) already exceed the disclosed \$500 billion NVIDIA-SK Group figure on a standalone basis. This suggests the \$500 billion partnership figure, while large in absolute terms, represents an incremental layer of demand visibility for SK hynix rather than the entirety of its AI-driven investment case, a distinction that is easy to lose when headline figures are compared without their underlying time horizons and structures.

Case Studies and Real-World Examples

SK Telecom's 2-Gigawatt Vera Rubin DSX AI Factory

SK Telecom's planned facility is the clearest embodiment of the compute side of the deal. The project uses NVIDIA's DSX platform, "a full-stack reference architecture including software, hardware, and operations," and will run NVIDIA Vera Rubin systems powered specifically by SK hynix HBM4 memory (Source: www.datacenterdynamics.com). NVIDIA describes the goal as providing compute for "AI training, inference, and agentic

workloads for companies and organizations in South Korea, later expanding to other regions" (Source: www.datacenterdynamics.com). This is a live, dated infrastructure commitment (first phase online in 2027) rather than a hypothetical scenario, and it is the single largest disclosed AI factory in the July 2026 announcement package.

SK hynix's Nasdaq Listing and Capacity Race

SK hynix's decision to list ADSs on the Nasdaq in July 2026 is a directly observable real-world case of a chip supplier restructuring its capital base specifically to fund AI memory expansion. Its SEC prospectus described the offering as covering "177,900,000 ADSs," each representing one-tenth of a common share (Source: www.sec.gov), with cornerstone investors including Baillie Gifford Overseas, Coatue Management and Situational Awareness Partners indicating interest in purchasing "up to an aggregate of US\$7 billion" of the shares offered (Source: www.sec.gov). Counterpoint's MS Hwang described the on-the-ground effect in Korea plainly: "everybody is coming," with hotels near SK hynix's facilities "fully booked" as "cloud companies and chipmakers... are all lining up to sign a long-term contract" (Source: www.cnbc.com) (Source: www.cnbc.com). SK hynix is also expanding in the United States, building a \$4 billion advanced packaging plant in West Lafayette, Indiana, scheduled for completion in 2028, and expects to receive up to \$458 million in CHIPS and Science Act funding plus up to \$570 million in Commerce Department loans (Source: www.cnbc.com) (Source: www.cnbc.com).

SK hynix's Second-Quarter 2026 Earnings and the Limits of Long-Term Deals

SK hynix's own quarterly results, reported July 29, 2026, add important nuance to how durable the long-term supply commitments described above may prove. Reuters reported that the company's "quarterly operating profit soared more than sixfold to a record high," yet this "fell short of lofty investor expectations, heightening market concerns about slower AI spending by big tech firms" (Source: www.reuters.com). Specifically, operating profit reached "60.5 trillion won for the April-June period, compared with 9.2 trillion won a year earlier and short of a 64 trillion won forecast by LSEG SmartEstimate," while "quarterly revenue rose 257% to 79.3 trillion won, below a 84 trillion won estimate" (Source: www.reuters.com). Shares fell a further "closed down 9.6%" on the earnings day even as the company reaffirmed demand strength (Source: www.reuters.com).

SK hynix President Song Hyun-jong linked the earnings call directly to the kind of long-term contract structure underpinning the NVIDIA deal, saying "major customers are still requesting more memory supply" and that the company was "seeking more long-term supply agreements to better manage chip price volatility" (Source: www.reuters.com). The company disclosed it "has concluded talks on around 10 such deals" beyond the NVIDIA agreement, typically running five years and including "financial safeguards such as deposits to ensure contract implementation" (Source: www.reuters.com). BNK Investment and Securities analyst Lee Min-hee offered a note of caution that applies directly to the wider Korea-NVIDIA push: "there are concerns that tech firms will take a breather in infrastructure spending" (Source: www.reuters.com). This episode illustrates a tension at the heart of the \$500 billion NVIDIA-SK Group announcement: the same long-term contract structure that gives suppliers demand visibility can also cap near-term pricing upside, a tradeoff SK hynix's own management acknowledged contributed directly to its earnings miss.

Naver, NVIDIA and Brookfield's GAK Sejong Data Center

The Naver expansion is a concrete example of how the broader Korea-NVIDIA push extends beyond SK Group. NVIDIA, Naver and Brookfield plan to grow the DSX AI factory buildout at Naver's GAK Sejong hyperscale data center in Sejong, South Korea, "from 55 megawatts to 200 megawatts by 2028," with Naver founder and chairman Haejin Lee describing the arrangement as having "propelled NAVER's vision for the AI factory business into a robust execution phase" (Source: nvidianews.nvidia.com). NVIDIA's planned \$1 billion investment is explicitly conditional: it is "subject to customary closing conditions and NAVER finalizing at least \$9 billion of committed financing for the project, separate from NVIDIA's planned investment" (Source: nvidianews.nvidia.com), illustrating that even nominally "signed" elements of this announcement wave carry real execution risk and financing contingencies.

Samsung and Broadcom's Parallel \$200 Billion Memory and Foundry MOU

Announced the same day and at the same San Francisco summit, the Samsung-Broadcom memorandum of understanding is a useful comparison case because it shows a different structure for a similarly sized headline number. Samsung's Vice Chairman and CEO of the Device Solutions Division, Young Hyun Jun, said the expanded Broadcom collaboration would help "deliver greater value to customers while advancing the AI infrastructure of the future," and Broadcom's Charlie Kawwas said the goal was to "continue to deliver technologies that power the next generation of AI infrastructure" by "combining Samsung's memory and foundry expertise with Broadcom's AI and connectivity leadership" (Source:

news.samsung.com). Unlike the NVIDIA-SK Group announcement, this MOU explicitly covers Samsung's 2-nanometer foundry process for Broadcom's chip designs in addition to HBM memory supply, extending to "advanced packaging technologies built on Samsung's 2nm process, including 2.3D and 2.5D integration" (Source: news.samsung.com).

The July 28, 2026 Semiconductor Selloff

The market's reaction to the wave of Korea-NVIDIA announcements is itself a case study in how AI infrastructure deals now move global equity markets. Within days of the SK Group announcement, "SK Hynix closed 14.65% lower" alongside double-digit percentage declines at Samsung Electronics, Samsung SDI, LG Innotek and Kioxia, while the VanEck Semiconductor ETF, ticker SMH, extended prior losses in US trading (Source: www.cnbc.com). Reuters reported that "together, the two companies account for nearly half of the benchmark KOSPI index, which closed down 10.8%, marking its biggest one-day decline since the early days of the U.S.-Iran conflict in March" (Source: www.reuters.com). Acadian Asset Management's Owen Lamont warned that "leveraged exchange-traded products could be adding to market swings," noting that "the entire ecosystem of levered ETFs in Korea, also in Hong Kong and in the United States, are possibly adding volatility and magnifying market fluctuations" (Source: www.cnbc.com). This episode underscores that a memory supply agreement of this scale carries financial market consequences well beyond the two primary signatories.

Implications and Future Directions

The NVIDIA-SK Group deal signals at least three structural shifts likely to persist into 2027 and beyond. First, AI infrastructure agreements are increasingly national in scope rather than purely corporate: South Korea's presidential office is now an active participant and public communicator of deal terms, framing the SK Group and Samsung agreements as components of a coordinated \$950 billion national strategy built around "three megaprojects" spanning semiconductors, AI data centers and physical AI (Source: www.koreaherald.com). Buyers and competitors should expect similar government-brokered mega-deals from other jurisdictions competing for sovereign AI infrastructure status.

Second, the memory supply chain is consolidating around long-term, pre-negotiated contracts rather than spot allocation, a shift memory vendors are actively encouraging because, as CNBC noted, such agreements "typically require customers to provide longer-term demand visibility," allowing suppliers like SK hynix "to plan its spending with more confidence" (Source: www.cnbc.com). This favors large, well-capitalized buyers who can commit to multi-year volumes, and structurally disadvantages smaller AI infrastructure operators who cannot offer similar demand guarantees. SK hynix's own management confirmed this dynamic is already reshaping how the company sells memory: it has "concluded talks on around 10 such deals" beyond the NVIDIA agreement alone (Source: www.reuters.com).

Third, the HBM4 transition itself is a technology inflection point that will reshape competitive positioning among GPU vendors. Because NVIDIA is reportedly demanding HBM4 data rates above the JEDEC baseline specification (Source: www.trendforce.com), memory suppliers that fail to hit these custom performance tiers risk losing allocation to Hynix specifically even while continuing to sell standard-specification HBM4 to other customers. This dynamic already appears to be playing out: Samsung passed NVIDIA's qualification tests first, while SK hynix, despite its overall HBM leadership, was described as "still optimizing its product to pass the 11Gb/s test" as of TrendForce's March 2026 reporting (Source: www.trendforce.com).

Looking further ahead, TechInsights projects that CPU-to-GPU connectivity demand, driven by inference workloads that increasingly rely on large key-value cache memory, will be "sustained at a 22% CAGR (compound annual growth rate) from 25 to 30," while broader GPU accelerator demand is expected to grow at an 18 percent CAGR over a similar window (Source: www.techinsights.com). If these projections hold, the SK hynix agreement announced in July 2026 will likely be only one of several similarly structured long-term memory deals NVIDIA signs with its major suppliers over the remainder of the decade, and the Samsung-Broadcom \$200 billion MOU signed the same day suggests other GPU and AI accelerator vendors are already following the same playbook of trading long-term purchase commitments for memory-supply security (Source: news.samsung.com).

Finally, the sharp equity market reaction of late July 2026 suggests investors remain unconvinced that current AI capital expenditure levels will generate commensurate returns, even as the underlying compute and memory demand data continues to point upward. Standard Chartered's Sundeep Gantori framed this tension as a matter of valuation rather than fundamentals, arguing that "what matters is risk-reward and at current valuations, risk-reward has improved" following the selloff (Source: www.cnbc.com). Whether that view or the more skeptical "letter of intent" reading of the deal proves correct will likely become clear as SK Telecom's 2027 AI factory milestone and NVIDIA's subsequent quarterly Data Center revenue disclosures arrive over the following several quarters.

Frequently Asked Questions (FAQs)

What is the NVIDIA SK Group \$500 billion AI deal? It is a package of letters of intent signed July 24, 2026, under which SK Telecom will build a 2-gigawatt AI data center using NVIDIA's DSX platform and Vera Rubin computing, and SK hynix will supply NVIDIA with next-generation HBM memory under a long-term partnership, together representing more than \$500 billion in projected business (Source: nvidianews.nvidia.com).

Is the \$500 billion figure a binding financial commitment? No. It was disclosed through signed letters of intent, which Datacenter Dynamics describes as agreements that "typically set out initial terms" and are "usually non-binding" (Source: www.datacenterdynamics.com), and Jensen Huang himself did not disclose how the total was calculated or over what period it would be realized (Source: m.koreaherald.com).

What is HBM4 and how is it different from HBM3e? HBM4 is the sixth-generation High Bandwidth Memory standard, doubling HBM3e's interface width from 1024 bits to 2048 bits and its channel count from 16 to 32, while increasing bandwidth from about 1.2 TB/s per stack to over 2.0 TB/s, up to 3.3 TB/s in advanced configurations (Source: blogs.sw.siemens.com). It is standardized by JEDEC as JESD270-4 and is not backward compatible with HBM3e memory controllers (Source: blogs.sw.siemens.com).

Who supplies HBM4 to NVIDIA? SK hynix, Samsung and Micron are all qualified or in the process of qualifying HBM4 for NVIDIA's Vera Rubin platform, with SK hynix projected to hold roughly a 50 percent bit-output share of the overall HBM market in 2026 and Samsung around 28 percent (Source: www.trendforce.com).

What does the deal mean for enterprises buying NVIDIA GPUs? It reinforces that GPU buyers should expect memory, not GPU compute silicon, to remain the primary supply bottleneck through at least 2027, and that securing long-term supply agreements, similar in structure to the advance purchase orders described by South Korea's presidential office (Source: www.koreaherald.com), is becoming the standard mechanism for guaranteeing allocation rather than relying on spot-market purchasing.

Is this deal part of a larger South Korean AI investment push? Yes. It sits within a disclosed \$950 billion package of Korea-US semiconductor and AI infrastructure partnerships, which also includes a separate \$200 billion Samsung-Broadcom memory and foundry MOU, a Naver-NVIDIA-Brookfield data center expansion, and collaborations with Hyundai Motor Group, LG Group and Doosan Group (Source: www.koreaherald.com).

Conclusion

The NVIDIA SK Group deal announced on July 24, 2026 combines a 2-gigawatt AI data center buildout led by SK Telecom with a long-term next-generation-memory partnership with SK hynix, and it forms one component of a broader wave of Korea-US AI infrastructure agreements. The headline \$500 billion figure describes projected business rather than a fixed, binding contract value, a distinction that matters when comparing AI infrastructure announcements. At its technical core, the deal is a bet on HBM4 and on coordinated supply planning between NVIDIA and SK hynix. For organizations planning GPU procurement, the practical lesson is to evaluate memory supply alongside GPU availability, packaging, system integration, power, networking and project scale. Long-term supply arrangements may be useful for large, predictable deployments, but this announcement does not establish a universal bottleneck or allocation rule for the wider GPU market ([NVIDIA partnership announcement](#)). Whether the deal's projected scale materializes will depend on execution and subsequent disclosures.

Tags: nvidia sk group deal, sk hynix hbm4, nvidia ai partnership, hbm4 vs hbm3e, vera rubin, sk telecom ai factory, nvidia gpu supply chain, hbm memory market

DISCLAIMER

The information contained in this document is provided for educational and informational purposes only. We make no representations or warranties of any kind, express or implied, about the completeness, accuracy, reliability, suitability, or availability of the information contained herein.

Any reliance you place on such information is strictly at your own risk. In no event will [GPUSmith](#) or its representatives be liable for any loss or damage including without limitation, indirect or consequential loss or damage, or any loss or damage whatsoever arising from the use of information presented in this document.

This document may contain content generated with the assistance of artificial intelligence technologies. AI-generated content may contain errors, omissions, or inaccuracies. Readers are advised to independently verify any critical information before acting upon it.

All product names, logos, brands, trademarks, and registered trademarks mentioned in this document are the property of their respective owners. All company, product, and service names used in this document are for identification purposes only. Use of these names, logos, trademarks, and brands does not imply endorsement by the respective trademark holders.

This document does not constitute professional or legal advice. For specific guidance related to your business needs, please consult with appropriate qualified professionals.

© 2026 [GPUSmith](#). All rights reserved.