

NVIDIA DGX Spark Review: Price, Specs & Performance (2026)

Published July 10, 2026 38 min read



Executive Summary

The **NVIDIA DGX Spark** is a desktop "AI supercomputer" built around NVIDIA's **GB10 Grace Blackwell Superchip**, combining a 20-core Arm CPU (10 Cortex-X925 performance cores plus 10 Cortex-A725 efficiency cores) with a Blackwell-generation GPU carrying 6,144 CUDA cores and fifth-generation Tensor Cores (Source: [hothardware.com](https://www.theregister.com)) (Source: www.theregister.com). It ships in a 150mm by 150mm by 50.5mm chassis weighing 1.2 kilograms, headlined by **128 GB of coherent unified LPDDR5x memory** shared between CPU and GPU at up to 273 GB/s, and rated for **1 petaFLOP of FP4 AI performance** with sparsity (Source: www.theregister.com) (Source: docs.nvidia.com). It launched on October 15, 2025 at a Founders Edition MSRP of \$3,999 with a 4 TB SSD, then jumped to \$4,699 on February 23, 2026, an 18 percent increase NVIDIA attributed to "worldwide constraints in memory supply" (Source: [hothardware.com](https://www.theregister.com)) (Source: [wccftech.com](https://www.wccftech.com)). As of this writing the unit lists for \$4,699.00 on NVIDIA's own marketplace and \$4,679.00 on Amazon (Source: marketplace.nvidia.com) (Source: www.amazon.com).

Independent benchmarking from LMSYS Org found the DGX Spark running GPT-OSS 20B in Ollama at 2,053 tokens per second prefill and 49.7 tokens per second decode, roughly one-fourth the speed of an RTX Pro 6000 Blackwell Workstation Edition (10,108 tps / 215 tps) and slower than a single RTX 5090 (8,519 tps / 205 tps) (Source: www.lmsys.org). At smaller model sizes and higher batch counts the Spark performs far better: Llama 3.1 8B at batch 32 reached 368 tokens per second decode, and machine learning writer Sebastian Raschka found single-sample PyTorch inference "roughly on par with the 6-times more expensive H100 data center GPU" (Source: www.lmsys.org) (Source: sebastianraschka.com). The system's 273 GB/s of memory bandwidth, versus roughly 3.35 TB/s on an H100 or up to 819 GB/s on an Apple Mac Studio with M3 Ultra, is the binding constraint on large-model throughput, and reviewers consistently flag it as the device's central trade-off: 128 GB of capacity in exchange for bandwidth (Source: blog.exolabs.net) (Source: madcoolstuff.com).

Reception is polarized but broadly favorable for a narrow use case. Tom's Hardware's review concluded the Spark "is a well-rounded toolkit for local AI thanks to solid performance from its GB10 SoC, a spacious 128GB of RAM, and access to the proven CUDA stack. But it's a pricey platform if you don't intend to use its features to the fullest" (Source: www.tomshardware.com). Thermal and power-delivery issues have dogged early units: game programmer John Carmack publicly reported his unit capping power draw near 100W of the rated 240W and rebooting under sustained load, a pattern NVIDIA's own developer forum moderators confirmed as "a known, documented problem across the platform, not just your machine," addressed

partially by an April 2026 firmware update (Source: forums.developer.nvidia.com) (Source: forums.developer.nvidia.com). Software compatibility is also uneven: the GB10's consumer-class sm120/sm121 GPU architecture lacks the tcgen05 tensor-core feature found on datacenter Blackwell and NVIDIA's own Jetson AGX Thor, forcing many libraries onto older Ampere-class code paths (Source: www.reddit.com).

This report also resolves a recurring point of search confusion: "Apache Spark" (the open-source distributed-computing engine originally created at UC Berkeley) and "DGX Spark" (this desktop hardware) are unrelated products that merely share a marketing name. NVIDIA's actual GPU acceleration for Apache Spark comes through the separate **RAPIDS Accelerator for Apache Spark**, a plugin that runs on datacenter and cloud GPUs and has driven documented customer results such as AT&T's reported 70 percent reduction in AI pipeline cost and time (Source: www.nvidia.com). Overall, the DGX Spark occupies a specific niche as of July 2026: a CUDA-native, 128 GB unified-memory development box priced between Apple's Mac Studio and AMD's Ryzen AI Max+ 395 (Strix Halo) mini-PCs on one side and multi-GPU datacenter rigs on the other, best suited to prototyping, fine-tuning up to roughly 70 billion parameters, and privacy-sensitive local inference rather than sustained high-throughput production serving (Source: wccfttech.com) (Source: www.theregister.com).

Introduction and Background

NVIDIA first teased this machine at CES in January 2025 under the codename "Project DIGITS," positioning it as a way to bring "powerful local AI to the desktop with more memory than any consumer GPU can provide" (Source: hothardware.com). It was renamed **NVIDIA DGX Spark** and began shipping worldwide on October 15, 2025, with NVIDIA calling it "the world's smallest AI supercomputer" (Source: nvidianews.nvidia.com) (Source: nvidianews.nvidia.com). The launch carried symbolic weight: NVIDIA founder and CEO Jensen Huang personally delivered one of the first units to Elon Musk at SpaceX, echoing his 2016 hand-delivery of the original DGX-1 supercomputer to Musk's team at OpenAI, a machine Huang credits with helping "kickstart the AI revolution" that produced ChatGPT (Source: nvidianews.nvidia.com) (Source: www.theregister.com).

The Spark's reason for existing is a hardware gap that has become increasingly obvious as open-weight large language models (LLMs) have grown: most desktop and laptop GPUs, even flagship consumer cards, top out at 24 to 32 GB of video memory (VRAM), which is not enough to comfortably load, fine-tune, or run inference on the 30 billion-to-200 billion-parameter models that now define the open-source frontier. NVIDIA's own RTX 5090 "only" has 32GB of memory on board, and it's trivial to exhaust that pool with cutting-edge LLMs, according to Tom's Hardware's review (Source: www.tomshardware.com). Scaling up within NVIDIA's own product line means jumping to professional cards like the RTX Pro 6000 Blackwell Workstation Edition, which costs \$8,500 or more for 96GB of GDDR7, or renting the kind of multi-GPU cloud infrastructure that individual developers, small labs, and startups often cannot economically justify for exploratory work (Source: www.tomshardware.com).

NVIDIA's answer is a unified-memory architecture, an approach Apple pioneered commercially with its M-series systems on a chip (SoCs), which pair CPU and GPU on a shared, coherent memory pool rather than partitioning dedicated VRAM. AMD followed with its Ryzen AI Max+ 395 (marketed as "Strix Halo"), and the DGX Spark represents NVIDIA's own entry into this category, distinguished chiefly by full support for the CUDA software ecosystem that "remains the dominant software platform for AI development the world over," in Tom's Hardware's framing (Source: www.tomshardware.com). At its core sits the **GB10 Grace Blackwell Superchip**, a two-die package co-developed with MediaTek, fabricated on a TSMC 3-nanometer-class process and joined by NVIDIA's coherent, high-bandwidth NVLink-C2C interconnect (Source: www.theregister.com) (Source: www.tomshardware.com). This report evaluates the DGX Spark's architecture, price, and independently measured performance; explains its documented thermal and software-compatibility problems; disambiguates it from the separate "RAPIDS Accelerator for Apache Spark" GPU data-processing product; and situates it against its closest rivals as of July 2026.

Product Overview: Architecture and Specifications

The GB10 Superchip pairs a 20-core Arm CPU, split between 10 high-performance Cortex-X925 cores and 10 efficiency-focused Cortex-A725 cores, with a Blackwell-architecture GPU carrying 6,144 CUDA cores, fifth-generation Tensor Cores, and fourth-generation RT cores (Source: www.nvidia.com) (Source: hothardware.com). Unlike NVIDIA's larger Grace-CPU data center parts, which use Arm's Neoverse cores, GB10's CPU complex was co-engineered with MediaTek, giving the chip an unusual pedigree that blends mobile-derived Arm silicon with NVIDIA's Blackwell GPU microarchitecture (Source: www.theregister.com).

The headline feature is the memory system: **128 GB of LPDDR5x**, run across a 256-bit interface at 8,533 MT/s (roughly 4,266 MHz effective), delivering up to 273 GB/s of bandwidth shared coherently between CPU and GPU with no static partitioning required (Source: docs.nvidia.com) (Source: www.theregister.com). HotHardware's review highlights why this matters in practice: on a conventional x86 laptop with a discrete GPU, memory typically has to be manually split between CPU and GPU pools in the BIOS, so any workload that spills past the GPU's allocation has to be swapped in and out, whereas on the Spark's Arm-based unified architecture "all of the memory is available to either the CPU or the GPU part of the chip" (Source: hothardware.com). NVIDIA rates the GPU at up to 1 petaFLOP of FP4 (4-bit floating point) performance with sparsity, equivalent to roughly 1,000 TOPS (trillion operations per second) of inference throughput (Source: docs.nvidia.com). LMSYS Org's early-access review positions

that figure as "placing its AI capability roughly between that of an RTX 5070 and 5070 Ti" in raw compute terms, even though the Spark's memory capacity vastly exceeds either card (Source: www.lmsys.org). The Register's hands-on testing cautions that the theoretical 1-petaFLOP figure assumes both 4-bit precision and structured sparsity simultaneously, workloads that can actually exploit both are uncommon, so "the most you'll likely see from any GB10 systems is 500 dense teraFLOPS at FP4" (Source: www.theregister.com).

Storage comes as a user-replaceable, self-encrypting M.2 2242 NVMe SSD in 1 TB or 4 TB capacities; the Founders Edition reviewed by most outlets shipped with the 4 TB drive (Source: awesomeagents.ai) (Source: www.tomshardware.com). Networking is unusually generous for a device this size: a 10 GbE RJ-45 port, Wi-Fi 7, Bluetooth 5.4, and, most notably, two QSFP ports driven by an onboard **ConnectX-7 Smart NIC** capable of up to 200 Gbps, which lets two Spark units be cabled together directly to pool memory for larger models (Source: www.lmsys.org) (Source: hothardware.com). All I/O, including the four USB-C 3.2 ports at 20 Gbps (one dedicated to the included 240W power brick), a single HDMI 2.1a display output, and the Ethernet and QSFP jacks, is arranged on the rear panel of the compact, gold-toned Founders Edition chassis (Source: www.theregister.com). At 150mm by 150mm by 50.5mm and 1.2kg, HotHardware measured the footprint at "just under 6x6 inches for a footprint and only two inches in height (around 70 cubic inches, or just over 1 liter)" (Source: hothardware.com).

Table 1 below summarizes the core specifications drawn from NVIDIA's own product and documentation pages.

SPECIFICATION	VALUE
Superchip	NVIDIA GB10 Grace Blackwell, TSMC 3nm-class process (Source: www.nvidia.com) (Source: www.theregister.com)
CPU	20-core Arm: 10x Cortex-X925 + 10x Cortex-A725 (Source: www.theregister.com)
GPU	Blackwell architecture, 6,144 CUDA cores, 5th-gen Tensor Cores, 4th-gen RT Cores (Source: hothardware.com)
AI performance	Up to 1 PFLOP FP4 (sparse); ~1,000 TOPS inference (Source: docs.nvidia.com)
Memory	128 GB LPDDR5x, 256-bit bus, up to 273 GB/s (Source: www.theregister.com)
Storage	1 TB or 4 TB self-encrypting NVMe M.2 (Source: awesomeagents.ai)
Networking	ConnectX-7 (2x QSFP, up to 200 Gbps), 10 GbE, Wi-Fi 7, BT 5.4 (Source: www.lmsys.org)
I/O	4x USB-C (20 Gbps), 1x HDMI 2.1a (Source: hothardware.com)
Power	240W external PSU, GB10 TDP 140W (Source: www.theregister.com)
Dimensions / weight	150 x 150 x 50.5 mm, 1.2 kg (Source: www.theregister.com)
OS / CUDA	NVIDIA DGX OS (Ubuntu 24.04 LTS base), CUDA 13.0.2 (Source: www.tomshardware.com) (Source: awesomeagents.ai)
Model support	Inference up to 200B parameters (single unit); fine-tuning up to 70B (Source: docs.nvidia.com) (Source: www.theregister.com)
Price (Founders Edition)	\$4,699 (raised from \$3,999 in February 2026) (Source: marketplace.nvidia.com) (Source: wccftech.com)

Two DGX Spark units connected over ConnectX-7 can jointly address AI models with up to 405 billion parameters, and, as of a March 2026 NVIDIA developer-blog update, the platform now officially supports scaling to as many as four nodes, extending addressable model size to roughly 700 billion parameters for inference and enabling near-linear speedups on reinforcement-learning and fine-tuning workloads when inter-node communication is minimized (Source: www.nvidia.com) (Source: developer.nvidia.com). The Register notes that officially "Nvidia only officially supports two Sparks in a cluster," though it found nothing technically preventing larger ad hoc setups (Source: www.theregister.com). NVIDIA also sells the same GB10

platform through seven OEM partners: Acer, ASUS, Dell, Gigabyte, HP, Lenovo, and MSI each offer their own chassis, cooling, and storage variants of the underlying superchip, giving corporate and institutional buyers procurement flexibility beyond the Founders Edition (Source: www.nvidia.com) (Source: www.tomshardware.com).

Pricing, Availability, and Total Cost of Ownership

The DGX Spark launched at a Founders Edition MSRP of \$3,999 for the 4 TB configuration on October 15, 2025 (Source: hothardware.com) (Source: www.theregister.com). On February 23, 2026, NVIDIA posted a forum announcement stating: "We have adjusted the MSRP of DGX Spark (Founders Edition) due to worldwide constraints in memory supply. DGX Spark continues to offer industry-leading AI performance in an ultra compact desktop form factor," raising the price to \$4,699, an increase of \$700, or about 18 percent (Source: wccfttech.com). Critically, NVIDIA confirmed the adjustment carried no hardware or configuration changes and would be honored globally, though it would not apply retroactively to orders already placed (Source: wccfttech.com). As of this writing, NVIDIA's marketplace lists the Founders Edition at \$4,699.00, bundled with a free 90-day NVIDIA AI Enterprise license and a complimentary Deep Learning Institute course valued at \$90 (Source: marketplace.nvidia.com). Amazon, fulfilled through Micro Center, currently lists the same unit at \$4,679.00 (Source: www.amazon.com). International pricing reflects the increase and local taxation; the Digit.in review of the Founders Edition reports an India retail price of Rs. 5,04,990, which it notes "immediately places it outside the realm of casual experimentation for most consumers" (Source: www.digit.in).

Context matters for evaluating that price. A comparable jump in VRAM within NVIDIA's own discrete-GPU lineup costs considerably more: the RTX Pro 6000 Blackwell Workstation Edition, with 96GB of GDDR7, runs \$8,500 and up, and HotHardware separately pegs "that datacenter GPU" at roughly \$10,000, before accounting for the cost of a compatible host workstation (Source: www.tomshardware.com) (Source: hothardware.com). Tom's Hardware, only half-joking, points to the Tinybox, a boutique eight-GPU AI workstation, as costing \$60,000 (Source: www.tomshardware.com). Against direct competitors, the DGX Spark now sits considerably above the market: AMD's Ryzen AI Max+ 395 ("Strix Halo") mini-PCs with 128 GB of unified memory range from \$2,499 (Beelink's GTR9 Pro) and \$2,599 (Framework's AI Max 300) to \$2,699 (GMKtec's EVO-X2) and \$3,500-plus for HP's Z2 Mini systems, and AMD's own branded Ryzen AI Halo box launched at \$3,999 with Windows 11 support (Source: wccfttech.com) (Source: www.tomshardware.com). On the cost-of-cloud side, Awesome Agents' review estimates the Spark "pays for itself after roughly 100-200 hours of cloud GPU time at H100 rates (\$2-4/hour)," implying a payback window of a few months for a daily-use ML engineer (Source: awesomeagents.ai). NVIDIA also sells the underlying GB10 platform through Acer, ASUS, Dell, Gigabyte, HP, Lenovo, and MSI, some of which, per Digit.in, are "priced much higher for no compute advantage over the vanilla DGX Spark," underscoring that the Founders Edition remains the reference price point for the category (Source: www.digit.in).

Performance and Benchmark Analysis

Independent benchmarking, rather than NVIDIA's own marketing figures, gives the clearest picture of where the DGX Spark excels and where it does not. LMSYS Org, the research group behind the Chatbot Arena leaderboard, ran an early-access unit against an RTX Pro 6000 Blackwell Workstation Edition, an RTX 5090, an RTX 5080, and Apple Mac hardware. Serving GPT-OSS 20B in MXFP4 through Ollama, the Spark achieved 2,053 tokens per second prefill and 49.7 tokens per second decode, versus 10,108 tps / 215 tps on the RTX Pro 6000 Blackwell, "roughly 4x faster," and 8,519 tps / 205 tps on a single RTX 5090, confirming that "the Spark's unified LPDDR5x memory bandwidth is the main limiting factor" (Source: www.lmsys.org) (Source: www.lmsys.org). On smaller, batchable models the picture flips: with SGLang serving Llama 3.1 8B in FP8 at batch size 1, the Spark hit 7,991 tokens per second prefill and 20.5 tokens per second decode, scaling to 7,949 tps / 368 tps decode at batch 32, which LMSYS called "excellent batching efficiency and strong throughput consistency" (Source: www.lmsys.org). Loading a 70-billion-parameter Llama 3.1 model in FP8 worked, at 803 tokens per second prefill but only 2.7 tokens per second decode, illustrating that decode throughput, which is memory-bandwidth-bound rather than compute-bound, is where the Spark's 273 GB/s ceiling bites hardest (Source: www.lmsys.org). Speculative decoding via the EAGLE3 algorithm in SGLang recovered some of that gap, delivering "up to a 2x speed-up in end-to-end inference throughput" across several tested models (Source: www.lmsys.org).

Machine learning author Sebastian Raschka's independent PyTorch benchmarks, run outside NVIDIA's own testing program, reached a similar conclusion from a different angle: for single-sample inference with a small, from-scratch 0.6-billion-parameter model, "the DGX Spark vastly outperforms the Mac Mini M4 Pro and is roughly on par with the 6-times more expensive H100 data center GPU, which is impressive," although "when it comes to batched runs, the H100 is the clear winner," which he attributes to the H100's superior memory bandwidth (Source: sebastianraschka.com) (Source: sebastianraschka.com). For standard chat-style Ollama inference, Raschka measured the Spark and a Mac Mini M4 Pro as roughly equivalent, both reaching "roughly 45 tok/sec when running gpt-oss-20B" (Source: sebastianraschka.com). Aggregator site Awesome Agents, synthesizing the LMSYS dataset alongside its own testing, summarized headline throughput as "Runs Llama 3.1 8B at 368 tokens per second (batch 32) and fits models up to 200B parameters in 128 GB unified memory," while cautioning that at 2.7 tokens per second decode, a loaded 70B model "is usable for batch processing or testing prompts, but not for interactive chat" (Source: awesomeagents.ai) (Source: awesomeagents.ai).

Fine-tuning tells a more favorable story. The Register fine-tuned Meta's 3-billion-parameter Llama 3.2 model on roughly one million tokens of training data using the Spark's "125 teraFLOPS of dense BF16 performance" and completed the job "in just over a minute and a half," compared with just under 30 seconds on a 48 GB RTX 6000 Ada workstation card that had cost roughly twice as much a year earlier, though The Register notes the RTX 6000 Ada's advantage narrows or disappears once model size or context length exceeds its 48 GB capacity (Source: www.theregister.com) (Source: www.theregister.com). An RTX 3090 Ti, with only 24 GB of GDDR6X, could not complete the same fine-tuning run at all, triggering a CUDA out-of-memory error (Source: www.theregister.com). Awesome Agents reports peak fine-tuning throughput using the Unsloth library of 53,658 tokens per second for Llama 3.1 8B LoRA and 5,079 tokens per second for Llama 3.3 70B QLoRA, noting Unsloth alone "delivers 2.5x speed-ups over standard Hugging Face transformers on the Spark" (Source: awesomeagents.ai). NVIDIA's own March 2026 technical blog post on multi-node scaling reports that four Spark nodes running the Nanochat fine-tuning workload reach roughly 74,600 tokens per second of aggregate throughput, up from about 18,400 tokens per second on a single unit, a near-4x speedup, and that reinforcement-learning throughput in NVIDIA Isaac Lab scales from 630 frames per second on one node to 2,520 frames per second on four (Source: developer.nvidia.com) (Source: developer.nvidia.com). For a single agent workload sharded across nodes with tensor parallelism, NVIDIA measured time-to-first-token dropping from 33,415 milliseconds on one node to 21,384 ms on two and 15,552 ms on four, running Llama 3.3 70B Instruct in NVFP4 with a 32K-token input (Source: developer.nvidia.com).

Table 2 consolidates the independently reported throughput figures across model sizes and frameworks, illustrating the sharp performance cliff between models that comfortably fit the Spark's bandwidth profile and those that merely fit its memory capacity.

MODEL / WORKLOAD	FRAMEWORK	BATCH	PREFILL (TOK/S)	DECODE (TOK/S)	SOURCE
Llama 3.1 8B (FP8)	SGLang	1	7,991	20.5	(Source: www.lmsys.org)
Llama 3.1 8B (FP8)	SGLang	32	7,949	368	(Source: www.lmsys.org)
GPT-OSS 20B (MXFP4)	Ollama	1	2,053	49.7	(Source: www.lmsys.org)
DeepSeek-R1 14B (FP8)	SGLang	8	2,074	83.5	(Source: www.lmsys.org)
Llama 3.1 70B (FP8)	SGLang	1	803	2.7	(Source: www.lmsys.org)
Nemotron 3 Super 120B (NVFP4)	TensorRT-LLM	1 (128K ctx)	2,855	18	(Source: developer.nvidia.com)
RTX Pro 6000 Blackwell (reference)	Ollama	1	10,108	215	(Source: www.lmsys.org)
RTX 5090 (reference)	Ollama	1	8,519	205	(Source: www.lmsys.org)

The pattern across every independent benchmark is consistent: the DGX Spark is bandwidth-constrained, not compute-constrained. As Awesome Agents summarizes, "273 GB/s shared between CPU and GPU sounds like a lot until you're serving a 70-billion-parameter model from it. For context, an H100 has 3.35 TB/s of HBM3 bandwidth, more than 12 times as much" (Source: awesomeagents.ai). MadCoolStuff's review, which benchmarked the Spark against both an RTX 5090 and a 256 GB Mac Studio with M3 Ultra, put it bluntly: "The 1 PFLOPS FP4 sparse number is the marketing peak. We did not see it," while noting per-token latency on long-context 70B inference was "slower than a quantized fit on a 5090, and noticeably slower than an M3 Ultra running its 800 GB/s pool" (Source: madcoolstuff.com) (Source: madcoolstuff.com). One creative workaround comes from EXO Labs, which paired a DGX Spark with an Apple Mac Studio M3 Ultra over a network link, running the compute-bound prefill phase on the Spark (measured at roughly 100 TFLOPs of FP16 performance) and the memory-bandwidth-bound decode phase on the Mac Studio (512 GB of memory at 819 GB/s). On an 8,192-token prompt with Llama 3.1 8B, the combined system finished in 2.32 seconds versus 4.34 seconds on the Spark alone and 6.42 seconds on the Mac Studio alone, "delivering 2.8x overall speedup compared to M3 Ultra alone" (Source: blog.exolabs.net).

Thermal Behavior, Reliability, and Software Compatibility

Thermal management has been the single most contentious topic in DGX Spark coverage since launch. Reports first surfaced publicly in October 2025 when programmer John Carmack, well known for his work on Doom and Quake, flagged that his unit was capping power draw at roughly 100W despite the rated 240W power supply and 140W GB10 chip TDP, and was experiencing spontaneous reboots under sustained load. A moderator on

NVIDIA's own developer forum later confirmed the pattern in a support thread, writing: "Since launch, the DGX Spark has had widely reported problems with its power delivery subsystem. John Carmack publicly noted back in October 2025 that his unit was capping at around 100W instead of the rated 240W and was experiencing spontaneous reboots under sustained load" (Source: forums.developer.nvidia.com), adding that "this is a known, documented problem across the platform, not just your machine" (Source: forums.developer.nvidia.com). Awesome Agents' review, drawing on multiple owner reports, describes CPU temperatures "hitting 95C during sustained training runs" with systems that "would throttle, spontaneously reboot, or shut down after 20-30 minutes of continuous load" (Source: awesomeagents.ai) (Source: awesomeagents.ai). NVIDIA responded with firmware updates through 2026; the forum moderator specifically noted, "there was a firmware update pushed in late April 2026 that specifically addressed the USB Power Delivery Controller and Embedded Controller stability," recommending owners run `fwupdmgm upgrade` before pursuing a hardware RMA (return merchandise authorization) (Source: forums.developer.nvidia.com). Separately, HP's community forum reported that a "major firmware update for the DGX Spark" substantially reduced idle power consumption on ConnectX-7-equipped systems, indicating NVIDIA has continued iterating on power management well past launch.

Digit.in's independent thermal testing, conducted in June 2026 with instrumented temperature probes, found the top surface reaching around 41.4 degrees Celsius under load and the rear exhaust climbing to around 52.6 degrees Celsius, with peak wall power draw around 220W, results that are meaningfully cooler than the failure reports circulating on NVIDIA's own forums, suggesting firmware revisions and unit-to-unit manufacturing variance both play a role (Source: www.digit.in) (Source: www.digit.in). LMSYS's own testing, run closer to launch, reported the opposite experience, finding the Spark "maintains sustained throughput across high-intensity tests without thermal throttling," which it credited to NVIDIA's "metal-foam cooling design" (Source: www.lmsys.org). This report presents both findings honestly: thermal behavior appears to vary considerably across individual units, workload types, and firmware versions, and prospective buyers running long, sustained fine-tuning jobs should budget for the possibility of throttling until they confirm their own unit's behavior on current firmware.

Software compatibility is the second major friction point. The GB10's GPU uses a consumer-oriented compute capability, referred to as sm120 or sm121 in different community reports, that differs architecturally from both NVIDIA's datacenter Blackwell parts (sm100) and its own Jetson AGX Thor robotics module, which shares the "Blackwell" branding but includes the tcgen05 tensor-core feature the Spark's GPU lacks. In a widely discussed Reddit thread, an NVIDIA support representative reportedly explained the omission by noting the chip "not has tcgen05 like jetson Thor or GB200, due die space with RT Cores and DLSS algorithm," a design trade-off that frustrated developers expecting full datacenter-class tensor core behavior in a device positioned for AI development (Source: www.reddit.com). The practical consequence, per the same thread, is that "sm80-class kernels can execute on DGX Spark because Tensor Core behavior is very similar," meaning many CUDA libraries fall back to six-year-old Ampere-generation code paths rather than exploiting native Blackwell optimizations (Source: www.reddit.com). The same thread and others report more mundane compatibility issues, including basic HDMI display output failures with certain monitors, a problem the original poster says was reported by "myself and servethehome" (Source: www.reddit.com). The Register encountered a related class of bug during its testing, reporting that on more than one occasion "the GPU robbed enough memory from the system to crash Firefox, or worse, lock up the system," a side effect of applications not yet optimized for GB10's unified memory model (Source: www.theregister.com). Awesome Agents also flags a diagnostic quirk in which the standard `nvidia-smi` monitoring tool "reports 'Memory-Usage: Not Supported' on UMA," confusing users accustomed to discrete-GPU memory reporting (Source: awesomeagents.ai). On the positive side, several reviewers highlight that DGX Spark ships with a working software stack out of the box: Docker with GPU passthrough, Ollama, CUDA 13.0.2, and JupyterLab are pre-installed, and Sebastian Raschka's tests found the Spark's `torch.compile` support notably more mature than Apple's MPS backend, at least prior to a PyTorch 2.9 fix that resolved most of the Mac-side compilation errors he had previously encountered (Source: awesomeagents.ai) (Source: sebastianraschka.com).

Apache Spark GPU Acceleration: A Separate "Spark" Ecosystem

Because "DGX Spark" and "Apache Spark" share a name, search queries about GPU-accelerated Apache Spark and about the NVIDIA DGX Spark hardware frequently intersect, even though the two are functionally unrelated products from different origins. Apache Spark is an open-source, distributed, big-data processing engine, originally created at UC Berkeley and now maintained under the Apache Software Foundation, used for SQL analytics, extract-transform-load (ETL) pipelines, and machine learning at enterprise data scale, typically across large clusters of servers rather than a single desktop. NVIDIA's GPU acceleration path for this software is a separate product entirely: the **RAPIDS Accelerator for Apache Spark**, a plugin NVIDIA describes as leveraging GPUs "to accelerate processing by combining the power of the RAPIDS cuDF library and the scale of the Spark distributed computing framework," designed so that "you can run your existing Apache Spark applications on GPUs with no code change by launching Spark with the RAPIDS Accelerator for Apache Spark plugin jar and enabling a single configuration setting" (Source: docs.nvidia.com) (Source: docs.nvidia.com). The plugin works by replacing supported SQL and DataFrame operations with GPU-accelerated equivalents, automatically falling back to the CPU implementation for operations it does not yet support, and it is deployed against datacenter and cloud GPUs, on platforms such as AWS EMR and Databricks, rather than desktop hardware like the DGX Spark reviewed above (Source: docs.nvidia.com).

NVIDIA's marketing for the RAPIDS Accelerator for Apache Spark emphasizes three benefits: "Faster Execution Time," "Reduced Infrastructure Costs," and "Quick Time to Value," positioning the plugin as a way to "do more with less: Spark on NVIDIA GPUs completes jobs faster with less hardware when compared to CPUs, saving time as well as on-premises capital costs or operational costs in the cloud" (Source: www.nvidia.com). Customer case studies cited on NVIDIA's page include telecommunications company AT&T, which "reduced both the cost and time of their AI pipeline by 70 percent" using GPU-accelerated Spark for data preparation and feature engineering (Source: www.nvidia.com), and Adobe, whose Senior Director of Machine Learning William Yan reported: "We're seeing significantly faster performance with NVIDIA-accelerated Spark 3 compared to running Spark on CPUs. With these game-changing GPU performance gains, entirely new possibilities open up for enhancing AI-driven features in our full suite of Adobe Experience Cloud apps" (Source: www.nvidia.com). Databricks, whose Matei Zaharia is Apache Spark's original creator and now the company's chief technologist, is quoted describing the collaboration as leading to "faster data pipelines, model training and scoring, that directly translate to more breakthroughs and insights for our community of data engineers and data scientists" (Source: www.nvidia.com). None of these figures relate to the DGX Spark desktop hardware discussed elsewhere in this report; readers researching GPU-accelerated Apache Spark for enterprise data engineering should evaluate the RAPIDS Accelerator on Databricks or AWS EMR documentation directly, while readers researching the desktop AI development box should disregard Apache Spark benchmark figures entirely, as they describe an unrelated architecture and deployment model.

Reception and Community Perspectives

Sentiment note: this section reflects public reviewer opinion, developer-forum commentary, and social sentiment rather than laboratory-verified fact, and is presented as such.

Professional reviewer sentiment converges on a similar theme: the DGX Spark is a well-engineered, CUDA-native niche product rather than a mass-market machine. Tom's Hardware's formal verdict states that the DGX Spark "is a well-rounded toolkit for local AI thanks to solid performance from its GB10 SoC, a spacious 128GB of RAM, and access to the proven CUDA stack. But it's a pricey platform if you don't intend to use its features to the fullest," awarding it credit for efficiency, ease of use, and "unique ConnectX 7 connectivity options for local clustering" while flagging that it "doesn't run Windows (yet)" (Source: www.tomshardware.com) (Source: www.tomshardware.com). LMSYS Org, an academic research group rather than a consumer outlet, called it "a gorgeous, well-engineered mini supercomputer that trades raw power for accessibility, efficiency, and elegance," concluding it is "not built to compete head-to-head with full-sized Blackwell or Ada-Lovelace GPUs" but is nonetheless well suited to "model prototyping and experimentation," "lightweight on-device inference," and "research on memory-coherent GPU architectures" (Source: www.lmsys.org) (Source: www.lmsys.org). Awesome Agents' aggregate score of 7.8 out of 10 summarizes the trade-off directly: "The DGX Spark is the best local AI development box you can buy for under \$10K, and it's not close," while cautioning that "it isn't a datacenter replacement" (Source: awesomeagents.ai). MadCoolStuff scored the unit 7 out of 10, summarizing its identity as "a development box, not a throughput box," recommending it to "solo researchers and small teams who need a development box for 70B-class work without leasing cloud hours" while advising throughput-focused single-model users to buy an RTX 5090 tower instead (Source: madcoolstuff.com) (Source: madcoolstuff.com).

Community sentiment on developer forums and Reddit skews more critical, largely centered on the gap between marketing claims and delivered performance. One widely upvoted Reddit post in r/LocalLLaMA, titled a "PSA," describes the poster's decision to return their unit after concluding NVIDIA had shipped "handheld gaming scraps" repurposed to compete with Apple and AMD's Strix Halo, based largely on the sm120/sm121 tensor core limitations discussed above (Source: www.reddit.com). Commenters on that same thread were split: one wrote that after an initial rough setup, "a few weeks later when the second one arrived, I could 'apt-get' what I needed... I would say that I'm getting a broader range of capabilities more smoothly by having CUDA available than I did on my Mac," while another argued that for the price, "you were right to have those expectations" of full compatibility (Source: www.reddit.com) (Source: www.reddit.com). A third commenter framed the trade-off in relative terms: after the February 2026 price increase, the Spark went from "nearly double the price" of AMD's Strix Halo to "more like 60-70% more expensive," while still preferring Strix Halo "even at the same price" for that user's specific workload, underscoring how sensitive community sentiment is to the pricing swings documented earlier in this report (Source: www.reddit.com). Sebastian Raschka's more measured personal verdict, from a research and prototyping perspective rather than a consumer purchasing one, is that "if you don't expect miracles or full A100/H100-level performance, the DGX Spark is a nice machine for local inference and small-scale fine-tuning at home," praising it as quiet enough for office use, "even under full load it's barely audible" (Source: sebastianraschka.com) (Source: sebastianraschka.com). Ollama, one of the software partners NVIDIA worked with at launch, described the collaboration in its own blog post as ensuring the platform "runs fast and efficiently out-of-the-box," reflecting the software ecosystem's broadly cooperative posture toward the hardware even where individual users report friction (Source: ollama.com).

Data Analysis and Evidence

Stepping back from individual anecdotes, several quantitative threads recur across the fact base gathered for this report. First, on pricing volatility: the \$700 increase NVIDIA imposed in February 2026 represents an 18 percent jump on the original \$3,999 Founders Edition price within roughly four months of general availability, a swing NVIDIA attributed entirely to "memory supply constraints" rather than any hardware change (Source:

wccfttech.com). Wccfttech's reporting places the increase in broader industry context, noting that AMD's competing Strix Halo-based systems "have seen prices for those being bumped up, too, due to DRAM constraints that have engulfed the tech market," indicating the Spark's price hike reflects a market-wide memory supply shock rather than an NVIDIA-specific pricing decision (Source: wccfttech.com).

Second, on the memory-bandwidth-versus-capacity trade-off that defines the product category: DGX Spark's 273 GB/s sits between AMD's Strix Halo platform (roughly comparable bandwidth in the same unified-memory category) and Apple's M3 Ultra Mac Studio, whose 512 GB configuration runs at 819 GB/s, nearly three times the Spark's rate, according to EXO Labs' measurements (Source: blog.exolabs.net). Against that, the Spark carries roughly 4x the raw FP16 compute of the M3 Ultra by EXO's estimate (approximately 100 TFLOPs versus approximately 26 TFLOPs), which is why disaggregated prefill-decode architectures like EXO's combined setup can outperform either machine running alone (Source: blog.exolabs.net). Against a full datacenter H100, the gap is far starker: 273 GB/s versus roughly 3.35 TB/s, a bandwidth deficit of more than 12x that fully explains why the Spark tracks the H100 closely on single-sample latency but falls far behind on batched throughput (Source: awesomeagents.ai).

Third, on multi-node scaling economics: NVIDIA's own data shows scaling efficiency depends heavily on workload communication pattern. Reinforcement-learning workloads in Isaac Lab, where nodes largely compute independently and synchronize only at step boundaries, scaled from 630 frames per second on one node to 1,241 on two and 2,520 on four, essentially linear (Source: developer.nvidia.com). By contrast, LLM inference, which requires continuous layer-by-layer synchronization across nodes, scales sub-linearly: token generation throughput scaled by only about 3x moving from one to four concurrent agent tasks even though the workload itself was 4x larger, because "as additional nodes are added, this overhead becomes increasingly dominant, limiting scaling efficiency" (Source: developer.nvidia.com). NVIDIA's Distributed Data Parallel (DDP) fine-tuning benchmark for Qwen3 4B showed cleaner scaling, with throughput moving from 2.03 samples per second on one node to 6.12 samples per second on three nodes, a 3x improvement matching the 3x increase in node count (Source: developer.nvidia.com).

Fourth, on reliability incident volume: a single thread on NVIDIA's own developer forum titled "DGX Spark Performance Degradation - GPU Power Draw Issue" accumulated 69 replies and roughly 4,400 views as of this report's research window, and multiple parallel threads document the same `MODS-020000600139` PowerStress diagnostic failure code, suggesting the power-delivery issue affected a non-trivial share of the early production run rather than a handful of isolated units (Source: forums.developer.nvidia.com).

Case Studies and Real-World Examples

SpaceX and the Symbolic Handoff

On October 13, 2025, NVIDIA CEO Jensen Huang personally delivered one of the first DGX Spark units to Elon Musk at SpaceX's Starbase, Texas facility, an event NVIDIA's press release frames as bookending a nearly decade-long arc: "In 2016, we built DGX-1 to give AI researchers their own supercomputer. I hand-delivered the first system to Elon at a small startup called OpenAI, and from it came ChatGPT, kickstarting the AI revolution" (Source: nvidianews.nvidia.com). The Register notes the DGX Spark's gold-cladded chassis design deliberately echoes the original DGX-1's aesthetic (Source: www.theregister.com).

University of Wisconsin-Madison: IceCube Neutrino Observatory, South Pole

Researchers at the University of Wisconsin-Madison's IceCube Neutrino Observatory deployed a DGX Spark at the South Pole to run AI models analyzing subatomic particle data from the facility's under-ice sensor array. Benedikt Riedel, computing director at the Wisconsin IceCube Particle Astrophysics Center, explained the deployment constraints: "There's no hardware store in the South Pole, which is technically a desert, with relative humidity under 5% and an elevation of 10,000 feet, meaning very limited power. DGX Spark allows us to deploy AI in a compartmentalized and easy fashion, at low cost and in such an extremely remote environment, to run AI analyses locally on our neutrino observation data" (Source: blogs.nvidia.com).

NYU Global AI Frontier Lab: Clinical Radiology Report Evaluation

At New York University's Global AI Frontier Lab, the ICARE (Interpretable and Clinically-Grounded Agent-Based Report Evaluation) project runs end-to-end on a DGX Spark, using collaborating AI agents to evaluate how closely AI-generated radiology reports align with expert sources without sending medical imaging data to the cloud. Lucius Bynum, a faculty fellow at the NYU Center for Data Science, said: "Being able to run powerful LLMs locally on the DGX Spark has completely changed my workflow. I have been able to focus my efforts on quickly iterating and improving the research tool I'm developing" (Source: blogs.nvidia.com). NYU professor Kyunghyun Cho separately described the broader lab-wide impact: "DGX Spark allows

us to access peta-scale computing on our desktop. This new way to conduct AI research and development enables us to rapidly prototype and experiment with advanced AI algorithms and models, even for privacy- and security-sensitive applications, such as healthcare" (Source: nvidianews.nvidia.com).

Harvard Kempner Institute: Genetic Drivers of Epilepsy

At Harvard's Kempner Institute for the Study of Natural and Artificial Intelligence, neuroscientists led by Institute Co-Director Bernardo Sabatini are using a DGX Spark to analyze roughly 6,000 genetic mutations in excitatory and inhibitory neurons, building protein-structure and neuronal-function prediction maps to guide which variants merit further wet-lab testing for links to epilepsy. NVIDIA's account of the project describes the Spark as functioning as "a bridge between benchtop and cluster-scale computing," letting researchers validate workflows and timing on a single unit before scaling successful pipelines to large GPU clusters for massive protein screens (Source: blogs.nvidia.com).

Stanford University: Biomni Biological Agent Prototyping

Stanford researchers use DGX Spark to prototype complete training and evaluation pipelines for their Biomni biological-agent workflows locally before scaling to large GPU clusters, enabling a tight, iterative development loop. NVIDIA reports the Stanford team measured "performance similar to big cloud GPU instances, about 80 tokens per second on a 120 billion-parameter gpt-oss model at MXFP4 via Ollama, while keeping the entire workload on a desktop" (Source: blogs.nvidia.com).

University of Delaware and ISTA: Campus-Scale and Compact Training Deployments

The University of Delaware, through an ASUS Ascent GX10 (a GB10-based OEM system), deployed its first unit to support research spanning sports analytics and coastal science; professor Sunita Chandrasekaran, director of the university's First State AI Institute, called the arrival "transformative for research," noting it lets teams "run large AI models directly on campus instead of relying on costly cloud resources" (Source: blogs.nvidia.com). At the Institute of Science and Technology Austria (ISTA), researchers use an HP ZGX Nano AI Station, also GB10-based, to train and fine-tune LLMs directly on a desktop; their open-source LLMQ software "enables working with models of up to 7 billion parameters," relying on the platform's 128 GB of unified memory to keep the entire model and its training data resident without the complex memory management typically required on consumer GPUs (Source: blogs.nvidia.com).

Implications and Future Directions

The DGX Spark's trajectory since October 2025 illustrates a broader pattern in the AI hardware market: unified-memory desktop systems are becoming a distinct product category, sitting between consumer GPUs and datacenter accelerators, and NVIDIA is no longer alone in defining it. AMD's Ryzen AI Max+ 395 platform, sold under names like the Ryzen AI Halo, directly targets the same 128 GB unified-memory niche at a lower price with native Windows 11 support, and Tom's Hardware's own coverage frames the competitive dynamic explicitly, cross-linking to a review titled "AMD Ryzen AI Halo review: AMD builds a DGX Spark of its own" (Source: www.tomshardware.com). NVIDIA has responded in kind: the same coverage cross-links to news that NVIDIA unveiled "RTX Spark" at Computex 2026, a related but distinct Superchip platform bringing a similar Arm-CPU-plus-Blackwell-GPU-plus-128GB-unified-memory formula to mainstream Windows laptops and desktops, and that Microsoft has separately debuted a "Surface RTX Spark Dev Box" on the same underlying platform, suggesting the DGX Spark's core architecture will proliferate well beyond its original niche developer-kit form factor over the next one to two years (Source: www.tomshardware.com) (Source: www.tomshardware.com).

The February 2026 price increase also signals a structural risk for this entire product category: because unified-memory AI PCs depend on large pools of the same LPDDR5x DRAM that mobile devices, other AI PCs, and now data-center memory buyers are all competing for, they are unusually exposed to memory-market volatility in a way that discrete-GPU systems, which use dedicated GDDR or HBM supply chains, are not. Tom's Hardware's broader industry reporting on "the rise of local agentic computing" facing "a brutal reality: rising DRAM prices" suggests this is not a one-off event specific to the Spark but a sector-wide dynamic likely to recur as more vendors chase the same unified-memory approach (Source: wccftech.com).

On the software side, NVIDIA's continued firmware updates through the first half of 2026, addressing both idle power consumption and USB Power Delivery Controller stability, along with the March 2026 expansion of officially supported clustering from two to four nodes, indicate the platform is still maturing well past its initial ship date rather than arriving feature-complete (Source: developer.nvidia.com). NVIDIA's introduction of Tile IR and cuTile Python, intended to give developers kernel-level portability from DGX Spark development environments to full Blackwell data-center GPUs "with minimal code changes," points toward the company's long-term strategy of using the Spark as an on-ramp into its broader CUDA and Blackwell

ecosystem rather than as a standalone product category (Source: developer.nvidia.com). For organizations evaluating GPU-accelerated data engineering rather than desktop AI development, the fact that RAPIDS Accelerator for Apache Spark requires "no code change" to existing Spark applications suggests NVIDIA is investing in lowering the barrier to GPU adoption for existing enterprise Spark workloads on cloud and datacenter infrastructure, a parallel but organizationally distinct effort from the DGX hardware line (Source: docs.nvidia.com).

Frequently Asked Questions (FAQs)

What is the NVIDIA DGX Spark? It is a desktop "AI supercomputer" built around NVIDIA's GB10 Grace Blackwell Superchip, combining a 20-core Arm CPU with a Blackwell GPU and 128 GB of coherent unified memory in a 150mm by 150mm by 50.5mm chassis, designed for local AI prototyping, inference, and fine-tuning (Source: www.nvidia.com).

How much does the NVIDIA DGX Spark cost? The Founders Edition launched at \$3,999 in October 2025 and now lists for \$4,699 on NVIDIA's marketplace after a February 2026 price increase attributed to memory supply constraints, with retailers like Amazon pricing it around \$4,679 (Source: marketplace.nvidia.com) (Source: www.amazon.com).

What are the NVIDIA DGX Spark's specs? A GB10 Grace Blackwell Superchip with a 20-core Arm CPU and Blackwell GPU (6,144 CUDA cores), 128 GB of LPDDR5x unified memory at up to 273 GB/s, 1 TB or 4 TB self-encrypting NVMe storage, ConnectX-7 networking up to 200 Gbps, and a 240W external power supply (Source: www.theregister.com).

How fast is the DGX Spark in benchmarks? Performance varies dramatically by model size: it reached 368 tokens per second decode on Llama 3.1 8B at batch 32, but only 2.7 tokens per second decode on a loaded Llama 3.1 70B model, because memory bandwidth, not compute, is the binding constraint at larger sizes (Source: www.lmsys.org).

Is Apache Spark related to NVIDIA DGX Spark? No. Apache Spark is an open-source distributed data-processing engine unrelated to this hardware. NVIDIA's GPU acceleration for Apache Spark comes through the separate RAPIDS Accelerator for Apache Spark plugin, which runs on datacenter and cloud GPUs rather than the DGX Spark desktop device (Source: docs.nvidia.com).

Does the DGX Spark run Windows? No; it ships with NVIDIA DGX OS, a customized build of Ubuntu 24.04 LTS. NVIDIA's newer, related "RTX Spark" platform targets Windows machines, but the DGX Spark itself is Linux-only as of this writing (Source: www.tomshardware.com).

Does the DGX Spark have thermal or reliability problems? Multiple independent sources, including NVIDIA's own developer forum moderators, confirm a documented pattern of power-delivery and thermal issues on some units, involving power capping near 100W and reboots under sustained load, though firmware updates through 2026 have addressed at least some cases (Source: forums.developer.nvidia.com).

How does the DGX Spark compare to a Mac Studio or AMD Ryzen AI Max+ 395? The Spark offers native CUDA support that neither competitor has, but the Mac Studio (M3 Ultra) offers up to 819 GB/s of memory bandwidth versus the Spark's 273 GB/s, and AMD's Ryzen AI Max+ 395 platforms generally undercut the Spark's price by \$1,000 to \$2,000 at comparable 128 GB memory capacity (Source: blog.exolabs.net) (Source: wccfttech.com).

Conclusion

The NVIDIA DGX Spark is a genuinely novel product: a CUDA-native desktop machine with more unified memory than any consumer GPU, capable of loading models up to 200 billion parameters and fine-tuning models up to roughly 70 billion parameters locally, without the multi-thousand-dollar step up to professional datacenter cards. Independent benchmarking from LMSYS Org, Sebastian Raschka, The Register, and multiple aggregator reviews consistently confirms that its 273 GB/s of memory bandwidth, not its 1-petaFLOP FP4 compute headline, is the real determinant of performance, and that the machine shines on smaller, batchable models and fine-tuning workloads while struggling with sustained, high-throughput inference on the largest models it can technically load. Its price rose 18 percent within four months of launch due to industry-wide memory shortages, its thermal and power-delivery behavior has drawn public criticism from prominent developers, and its consumer-class Blackwell GPU architecture creates real software compatibility friction relative to NVIDIA's own datacenter and Jetson Thor products. None of this appears to be marketing overreach so much as the ordinary growing pains of a genuinely new hardware category competing against Apple's Mac Studio and AMD's Ryzen AI Max+ 395 on one side and NVIDIA's own datacenter GPUs on the other. For developers, researchers, and small teams who specifically need full CUDA compatibility, 128 GB of coherent memory, and privacy-preserving local execution, and who understand its bandwidth-limited profile going in, the evidence gathered here supports the recurring reviewer conclusion that it is the strongest option in its specific category, even if it is not, and was never positioned to be, a datacenter replacement. Separately, and importantly for search accuracy, none of this reflects on GPU-accelerated Apache Spark data engineering, which is an unrelated NVIDIA product line, the RAPIDS Accelerator for Apache Spark, evaluated on entirely different hardware and by an entirely different set of enterprise customers.

Tags: nvidia dgx spark, dgx spark review, dgx spark price, dgx spark specs, dgx spark benchmarks, grace blackwell gb10, unified memory ai pc, local llm hardware, rapids accelerator for apache spark, ai supercomputer

DISCLAIMER

This document is provided for informational purposes only. No representations or warranties are made regarding the accuracy, completeness, or reliability of its contents. Any use of this information is at your own risk. GPUSmith shall not be liable for any damages arising from the use of this document. This content may include material generated with assistance from artificial intelligence tools, which may contain errors or inaccuracies. Readers should verify critical information independently. All product names, trademarks, and registered trademarks mentioned are property of their respective owners and are used for identification purposes only. Use of these names does not imply endorsement. This document does not constitute professional or legal advice. For specific guidance related to your needs, please consult qualified professionals.