

NVIDIA DSX Due Diligence: Adoption and Contract Guide

Plan and build private AI compute with GPU Smith. We connect workload requirements with GPU hardware, networking, deployment and operating decisions so your infrastructure fits the work it must perform.

Published by GPU Smith · 2026-09-19 · nvidia dsx · dsx due diligence · ai factory · dsx os · dsx maxlps · gpu infrastructure · data center engineering · it ot integration · acceptance testing

Original article: <https://gpusmith.com/articles/en/nvidia-dsx-due-diligence-guide>



In this article

1. Executive Summary
 2. Introduction and Background
 3. Key Changes
 4. Architecture and Responsibility Map
 5. Implementation Considerations and Process Changes
 6. Procurement, Contract, and Acceptance
 7. Data Analysis and Evidence
 8. Implications and Future Directions
 9. Frequently Asked Questions (FAQs)
 10. Conclusion
-

Executive Summary

NVIDIA DSX is not a server, rack, or appliance. NVIDIA announced it on **May 31, 2026** as an AI factory-scale platform (nvidianews.nvidia.com) whose scope crosses generation-specific reference designs, simulation, operating software, power optimization, grid response, and data exchange. By contrast, DGX and NVL72 describe compute-system scope, while HGX is a GPU compute platform. A buyer therefore cannot make “DSX compliant” the complete specification. Due diligence must identify which of **DSX Reference Design, DSX Sim, DSX OS, DSX MaxLPS, DSX Flex, and DSX Exchange** is in scope, then assign every interface, dependency, license, support obligation, control privilege, and acceptance artifact.

The public stack is mixed in maturity and commercial model. NVIDIA calls DSX OS open-source and modular (docs.nvidia.com), but that platform label does not establish identical licenses or support for every dependency. **NVIDIA Config Manager** is documented as Developer Preview and not recommended for production (github.com); **DSX Exchange** has best-effort open-source support (github.com); and Run:ai commercial control-plane terms are subscription-based per GPU unless otherwise agreed (nvidia.com). Contracts need a dated component, version, license, entitlement, and support-owner matrix rather than a platform-level assurance.

Power claims require equally careful handling. NVIDIA markets **up to 40% more GPUs** in the same site-power envelope (docs.nvidia.com), while noting an assumed **power usage effectiveness (PUE) of 1.1** and that Vera Rubin testing remained in progress (docs.nvidia.com). Its September 2026 account of a Lambda test reports **19 nodes at an 85% power policy versus 16 at full power, 24% more cluster-wide token throughput, and 23% better performance per watt** (blogs.nvidia.com) (blogs.nvidia.com) (blogs.nvidia.com). Those are vendor-reported results, not a transferable project baseline. NVIDIA's own pilot guidance says the measured baseline, not a fixed table percentage, is the reference (docs.nvidia.com).

The investable scope is therefore an evidence system. Before capital approval, require a versioned bill of materials, topology, load schedule, calibrated simulation, uncertainty register, and generation-matched reference design. Before energization, approve trust zones, read and write privileges, interlocks, schemas, retention, rollback, and safe states. Before final acceptance, reproduce fixed-power throughput and service-level objectives (SLOs), security isolation, failure recovery, and grid-event behavior from system-level meters.

GPU Smith is an **adjacent independent engineering advisor**, not a DSX product vendor. Its published method centers written acceptance criteria and as-built records (gpusmith.com) and is relevant as a buyer-side discipline, not as a row in a product comparison.

Introduction and Background

The NVIDIA DSX due diligence question is not whether its individual ideas are useful. The question is whether a proposed private AI factory has a complete, contractible, and testable allocation of responsibility across those ideas. NVIDIA's public definition is explicit: **DSX is an AI factory-scale platform** (nvidia.com). That is a larger system boundary than a GPU, [HGX baseboard](#), [DGX system](#), [NVL72 rack](#), Kubernetes cluster, data hall, or utility interconnection considered alone.

The timing matters. The portfolio changed rapidly between the May announcement and this report's **September 19, 2026** cutoff. A public facilities walkthrough identifies **DSX Facilities Infrastructure Design Guide v2.0, dated August 19, 2026** (docs.nvidia.com), while Dynamic Power Software was still a **Developer Preview** on September 15 (docs.nvidia.com). There is no single public DSX-wide revision or price sheet. Every procurement baseline should consequently record the document revision, repository tag, license, entitlement, hardware generation, and support date separately.

This report uses “acceptance” narrowly: evidence that the delivered system meets buyer-supplied criteria under declared conditions. A digital twin is design evidence, not an operating result. A reference architecture is a controlled starting point, not the project's as-built design. A vendor performance claim is a hypothesis until the buyer reproduces it on the intended workload, service envelope, meters, firmware, and facility. NIST similarly treats verification, validation, and uncertainty quantification as lifecycle concerns for digital twins (nist.gov).

For an adjacent advisor such as GPU Smith, the appropriate role is buyer-side scope definition, vendor-quote verification, integration evidence, and acceptance planning. Its first-party description includes independent review of AI data-center and GPU-cloud investments (gpusmith.com), but it should not be represented as a competing DSX module or NVIDIA product.

Accordingly, this report treats **NVIDIA DSX architecture** as a responsibility system, maps the **NVIDIA DSX AI factory stack**, and separates a generation-specific **NVIDIA DSX reference architecture** from the delivered site. It provides an **NVIDIA DSX adoption checklist**, an **NVIDIA DSX vendor evaluation** method, measurable **NVIDIA DSX validation requirements**, and an **NVIDIA DSX contract requirements** schedule. NVIDIA DSX infrastructure consulting is discussed only as independent buyer-side engineering, never as a substitute for primary product evidence.

Key Changes

DSX changes the unit of procurement

DGX, **HGX**, and **NVL72** remain important equipment scopes, but they do not close the factory boundary. NVIDIA defines a Grace Blackwell rack-scale system around [GPUs connected through NVLink](#) (docs.nvidia.com); HGX brings together GPUs with related CPU, interconnect, networking, and software

technologies (nvidia.com). DSX reaches beyond those compute products into facility design, simulation, operation, and grid interaction. The buyer's work breakdown structure must expand accordingly.

- **Reference Design:** Establish the [generation-specific compute, network, storage, facility, and cluster baseline](#). Do not mix GB300 and Vera Rubin evidence because NVIDIA describes these architectures as generation-specific (docs.nvidia.com).
- **DSX Sim:** Model and validate infrastructure before and after physical deployment. Required inputs include geometry, equipment curves, power topology, cooling behavior, workload profiles, and calibration observations.
- **DSX OS:** Assemble scheduling, orchestration, cluster runtime, health, provisioning, security, and rack operations. “Open-source and modular” is an architecture description, not a universal warranty or support term.
- **DSX MaxLPS:** Coordinate facility design, power policy, and performance-per-watt techniques. NVIDIA says it connects facility and site design with Dynamic Power Software (docs.nvidia.com).
- **DSX Flex:** Adapt workload demand to grid conditions, including load shedding signals (nvidianews.nvidia.com).
- **DSX Exchange:** Move signals among power management, building management, cooling, grid interfaces, and compute schedulers (docs.nvidia.com).

These modules can be adopted together or selectively. Selection changes the integration burden, but it does not remove the underlying function. A site that declines DSX Exchange still needs a governed interface between operational technology (OT), facility systems, power policy, and workload control.

The DSX OS inventory creates a support-boundary problem

The operating stack contains both modular projects and commercial NVIDIA software. Its named components should be inventoried by capability, not compressed into one “DSX OS” line item.

- **Scheduling and grouping:** KAI Scheduler requires NVIDIA GPU Operator for GPU-requesting workloads (github.com); KAI consumes Karta gang instructions through a Karta-based pod grouper (github.com).
- **Inference orchestration:** Grove provides a declarative Kubernetes interface for inference workloads (github.com); Dynamo targets distributed, multi-node inference serving.
- **Cluster runtime:** AI Cluster Runtime records known-good combinations of drivers, operators, kernels, and system settings (github.com). Its security policy limits fixes to the latest released minor (github.com).
- **Topology and remediation:** Topograph requires Kubernetes **1.27 or later** and Helm **3.10+ or 4.x** (github.com); NVSentinel detects and remediates GPU faults on Kubernetes nodes (github.com).

- **Provisioning and isolation:** NVIDIA Infra Controller provides site-local, zero-trust bare-metal lifecycle management with data processing unit (DPU) enforced isolation (github.com). The site must still supply target firmware versions because the controller does not discover the newest release automatically (github.com).
- **Commercial control plane:** Run:ai is included with NVIDIA AI Enterprise (docs.nvidia.com); Infrastructure 8.2 documents Run:ai **2.26** for self-hosted and software-as-a-service delivery (docs.nvidia.com).

The conclusion is operational, not ideological: open source does not eliminate support risk, while commercial packaging does not prove interoperability. Each component needs an owner, supported version, maintenance window, security-response path, rollback plan, and exit mechanism.

Architecture and Responsibility Map

Table 1 maps each DSX layer to its principal lifecycle decision and minimum evidence. “Likely owner” is a recommended accountability assignment, not an NVIDIA warranty.

Layer	Lifecycle and function	Required inputs	Output or control	Likely accountable owner	Acceptance artifact
Reference Design	Concept and detailed design; generation-specific architecture	Workload SLO, GPU generation, utility limit, rack, network, storage, cooling and resilience requirements	Versioned topology and bill of materials	Chief architect with electrical, mechanical and network leads	Approved design basis, revision set, deviation register and itemized BOM
DSX Sim	Design through operations; digital twin and scenario analysis	Geometry, equipment curves, failure modes, workload traces and calibration observations	Predicted thermal, power, network and operational behavior	Modeling lead and discipline engineers	Model configuration, calibration record, uncertainty limits and scenario report
DSX OS	Commissioning and operations; schedule, serve, provision, observe and remediate	Supported-version matrix, identities, policies, images, firmware and cluster topology	Workload placement, lifecycle actions, health events and administrative state	Platform engineering and site reliability engineering	Software bill of materials, API schemas, isolation tests, recovery tests and runbooks
MaxLPS	Design and operations; optimize compute output within a fixed power budget	Validated equipment envelope, rack telemetry, workload priority and facility limit	GPU, rack or workload power policy	Platform lead with facilities energy manager	Baseline and candidate test using the same meter boundary, workload and SLO
Flex	Operations; respond to utility or grid events	Event schema, committed capacity, notice, duration, recovery ramp and workload flexibility	Curtailement, deferral, migration or restoration request	Energy manager with workload owner	Event trace, delivered reduction, SLO result, rollback and reconciliation report

Layer	Lifecycle and function	Required inputs	Output or control	Likely accountable owner	Acceptance artifact
Exchange	Integration and operations; broker factory signals	Endpoint inventory, schemas, timestamps, identity, access policy and retention	Normalized telemetry and authorized control messages	Integration architect and OT security owner	Interface control document, conformance tests, audit trail and outage test

The table exposes the central due-diligence point: an input without a named supplier and quality criterion is a hidden buyer obligation. A generated output without a consuming owner is shelfware. A control without an interlock, bounded authority, and rollback is an operational hazard.

IT and OT boundary

OT must be treated as a safety and availability domain, not merely another observability source. NIST says OT security has unique performance, reliability, and safety requirements (csrc.nist.gov) and describes fail-to-a-known-state design for unexpected conditions (nvlpubs.nist.gov). A safety instrumented system can remain independent from other control systems; DSX software should not bypass that separation.

The minimum register should cover:

- **Compute telemetry:** GPU power, clocks, temperature, errors, job state, useful throughput, latency and quality result.
- **Server management:** Baseboard management controller (BMC), firmware, boot, inventory, health and power controls. Redfish offers a schema-based RESTful model (dmtof.org) and a multivendor out-of-band interface (dmtof.org).
- **Network state:** Fabric health, port counters, congestion, topology, access control, management reachability and time synchronization.
- **Rack and facility state:** Power distribution, cooling distribution, temperatures, flow, pressure, valve and fan state, uninterruptible power supply status, and building alarms.
- **Grid state:** Event identifier, requested megawatts, start, duration, ramp, baseline rule, telemetry acknowledgement and restoration instruction.
- **Control authority:** Read-only, advisory, bounded write, emergency override, manual approval, and prohibited actions for every endpoint.
- **Failure behavior:** Stale-data timeout, loss of Exchange, partial network partition, meter disagreement, rejected actuation, scheduler loss and recovery ramp.

BACnet is designed for monitoring and control of heating, ventilation, air conditioning, refrigeration, and other building systems (data.ashrae.org). OPC UA covers client, server, and user authentication (opcfoundation.org) and can correlate client and server audit logs (opcfoundation.org). Protocol support alone is insufficient. The interface control document must state exact object paths, units, schema versions, update rates, quality flags, time source, retention, and write constraints.

Interoperability claims should be tested across the selected equipment because ASHRAE characterizes BACnet as a vendor-independent approach to interoperable equipment and controls ([ashrae.org](https://www.ashrae.org)). The test should demonstrate discovery, units, writable ranges, alarm propagation, authentication, clock behavior, and loss-of-link handling rather than only a successful connection.

Identity and isolation should follow zero-trust principles. NIST rejects implicit trust based only on location or ownership (csrc.nist.gov) and treats subject and device authentication and authorization as discrete functions (csrc.nist.gov). Kubernetes allows pod-to-pod communication by default (kubernetes.io), so namespace separation alone is not tenant isolation. Its own guidance identifies RBAC, quotas, and network policies as essential controls (kubernetes.io). Require role-based access control (RBAC), quotas, network policies, node and device isolation, secrets management, image provenance, and negative tests across tenants.

The security owner should also make authorization and telemetry governance explicit. SNMPv3 defines an access-control subsystem that provides authorization services ([rfc-editor.org](https://www.rfc-editor.org)), while IETF telemetry guidance calls out storage encryption, access methods, and retention practices ([ietf.org](https://www.ietf.org)). CISA's cross-sector goals address IT and OT in one framework ([cisa.gov](https://www.cisa.gov)), which is a useful review structure even though project-specific controls still require an engineering decision.

Implementation Considerations and Process Changes

Evidence before capital approval

A responsible DSX adoption checklist begins before a purchase order. The design authority should freeze a traceable baseline containing:

- **Workload envelope:** Models, precision, context length, request distribution, batch policy, training or inference mix, latency percentiles, quality threshold, availability, recovery objective, and growth cases.
- **Facility envelope:** Utility and generator capacity, redundancy, voltage, rack density, cooling temperatures and flows, water constraints, floor loading, fire strategy, maintenance states, and expansion phases.
- **Generation boundary:** Exact GPU, CPU, switch, DPU, rack, firmware, driver, CUDA, operator, Kubernetes, storage, and serving-stack versions.
- **Topology and BOM:** Port-level network design, power path, cooling path, management network, spares, cables, optics, meters, sensors, and replaceable units.
- **Simulation basis:** Source and date of each input, calibration observation, parameter uncertainty, scenario matrix, predicted value, later measured value, and disposition.
- **Decision thresholds:** Buyer-supplied tolerances for temperature, flow, pressure, energy, loss, queue time, latency, availability, recovery, and model-output quality.

NIST says credibility assessment should include verification, validation, and uncertainty quantification ([nist.gov](https://www.nist.gov)). Calibration should compare model outputs with physical observations at declared operating points. Validation should use observations that were not merely fitted. Sensitivity analysis should identify which

uncertain inputs can change the procurement decision. No simulation should become a warranty or acceptance criterion unless its intended use, error bound, and measurement method are explicit.

The validation plan should state the intended purpose because NIST frames verification and validation as showing that a twin meets its intended purpose and design goals (nist.gov). A model accepted for concept screening may not be accepted for protective settings or contractual capacity.

Partner assets should enter that same evidence chain. Vertiv describes its contribution as SimReady power and cooling assets, validated interfaces, and repeatable infrastructure building blocks (vertiv.com). Buyers should still record the asset revision, source curves, applicable equipment configuration, calibration observation, and permitted use.

Full-stack adoption versus selected modules

Full-stack adoption may simplify integration when a buyer accepts NVIDIA's supported combinations and operating cadence. Selective adoption may preserve an existing scheduler, Kubernetes distribution, observability system, BMC/DCIM/BMS, identity provider, security information and event management platform, network automation, or storage control plane. Neither approach is automatically lower risk.

Score each capability on a five-year worksheet:

- **License and support:** Subscription metric, included components, environment count, nonproduction rights, response targets, renewal, audit terms, and end-of-support.
- **Integration:** Adapters, schemas, testing, vendor coordination, data mapping, time synchronization, identity, secrets, monitoring, and automation.
- **Operations:** Staffing, on-call coverage, training, test environments, spares, maintenance windows, incident ownership, and escalation.
- **Change:** Release cadence, compatibility window, regression suite, rollback, firmware qualification, and planned downtime.
- **Portability:** Export formats, configuration ownership, telemetry retention, API stability, replacement path, and migration labor.
- **Disconnected operation:** Local registries, mirrored packages, offline updates, certificate lifecycle, entitlement behavior, vulnerability data import, and recovery without external dependencies.

OpenTelemetry is vendor and tool agnostic (opentelemetry.io), but an open telemetry format does not ensure every operational action is portable. Prometheus defaults to **15 days** of local sample retention if no size or time limit is set (prometheus.io); a contract that needs forensic, capacity, or utility reconciliation history must specify longer durable retention elsewhere.

Procurement, Contract, and Acceptance

DSX vendor evaluation should issue one responsibility matrix and one requirements traceability matrix across facility, hardware, software, network, and operations. ISO/IEC/IEEE 15288 is useful framing because it structures information exchange among acquirers, suppliers, and other stakeholders ([iso.org](https://www.iso.org)) across conception, development, production, use, support, and retirement ([iso.org](https://www.iso.org)).

Table 2 converts that lifecycle into an RFP and contract schedule. Pass values are deliberately left as buyer inputs where no universal threshold exists.

Deliverable	Contract content	Factory or site test	Pass criterion and artifact	Responsible party
Design baseline	Reference-design generation and revision, BOM, topology, calculations, approved deviations	Document and constructability review	Buyer-approved version, no unresolved critical interfaces; signed design package	Design authority and discipline leads
Software and support matrix	Component, version, license, source, entitlement, support owner, severity response, maintenance and end date	Verify deployed inventory against schedule	Exact match or approved deviation; machine-readable inventory and contracts	Platform vendor and procurement
API and schema package	OpenAPI or equivalent definitions, event schemas, units, timestamps, compatibility and deprecation	Positive, negative, load, stale-data and downgrade tests	Buyer-set error rate and compatibility window; conformance report	Integration vendor
Security evidence	Threat model, trust zones, identities, RBAC, network policy, software bill of materials (SBOM), signing, vulnerabilities and audit integration	Tenant breakout, privilege, image, secret, logging and restore tests	No unauthorized path; complete evidence pack and remediation record	Security owner and suppliers
Facility and controls	Single lines, sequences, point list, calibration, interlocks, manual override and safe state	Integrated systems test under normal, maintenance and fault states	Buyer-set electrical, thermal and recovery limits; calibrated trace	Commissioning authority and controls vendors
Performance and power	Workload, SLO, measurement boundary, repetitions, warm-up, duration and uncertainty	Baseline and candidate run under the same energy budget	Buyer-set useful jobs or tokens, latency, quality and availability; raw telemetry	Workload owner and independent test lead
Training and handover	Role-based curriculum, runbooks, spares, escalation, backup, rollback and exit procedure	Operator scenarios and restoration exercise	Named staff demonstrate procedures; attendance, results and as-built manual	Prime contractor and operations owner

The table makes the “support boundary” inspectable. NIST’s Secure Software Development Framework recommends putting core requirements in acquisition documents and contracts (nvlpubs.nist.gov), and it calls for defined criteria for software security checks (nvlpubs.nist.gov). Provider agreements should allocate responsibility and specify conformance attestations (nvlpubs.nist.gov). NTIA’s SBOM minimum elements

include supplier, component name, version, unique identifier, and dependency relationship ([ntia.gov](https://www.ntia.gov)). NTIA also emphasizes automatic generation and machine-readable formats for scale ([ntia.gov](https://www.ntia.gov)). Those fields should be joined to runtime inventory, entitlement, and support owner, not delivered as an orphan file.

Acceptance criteria belong in the statement of work, not in a post-installation negotiation. NASA's acquisition guidance says final criteria must be in the contract statement of work ([nasa.gov](https://www.nasa.gov)) and that results must be captured in a test report or acceptance data package ([nasa.gov](https://www.nasa.gov)). ASHRAE Standard 202 likewise provides a framework for design documents, specifications, procedures, documentation, and reports ([ashrae.org](https://www.ashrae.org)).

ASHRAE defines Standard 202 as the minimum acceptable commissioning process for new buildings and systems ([ashrae.org](https://www.ashrae.org)). DOE guidance places the post-installation measurement and verification report before final acceptance ([energy.gov](https://www.energy.gov)). Together, these support a contract sequence in which measured evidence precedes, rather than follows, payment and handover approval.

Pilot and evidence pack

The pilot should move through five gates:

1. **Design release:** Approved requirements, responsibility matrix, versioned BOM, topology, interfaces, simulation basis, risks, and test plan.
2. **Factory acceptance test (FAT):** Inventory, firmware, configuration, cabling, functional controls, security baseline, telemetry completeness, failure injection, and supplier corrections.
3. **Site acceptance test (SAT):** Installed power and cooling, network paths, identity, time, facility integration, failover, workload and power tests under site conditions.
4. **Operational proving period:** Declared workload mix across normal operations, maintenance, faults, tenant changes, curtailment, recovery, and software updates.
5. **Handover:** As-built drawings, configurations, source and binaries, licenses, SBOM, schemas, raw test data, calibration certificates, training results, open items, warranties, and exit export.

ASHRAE's AI data-center framework places integrated-system testing at commissioning **Level 5** ([ashrae.org](https://www.ashrae.org)). This is where DSX's cross-domain promise should be proven: not just that each subsystem starts, but that the factory respects thermal, power, performance, safety, and recovery constraints when signals and dependencies interact.

Handover is also an operating control. Federal commissioning guidance describes a Systems Manual as the information needed to understand, operate, and maintain the delivered systems ([wbdg.org](https://www.wbdg.org)). ASHRAE Standard 202 includes training requirements for continued successful system performance ([ashrae.org](https://www.ashrae.org)). The contract should therefore make completion of training scenarios and delivery of the Systems Manual conditions of acceptance, not optional closeout tasks.

Data Analysis and Evidence

The correct quantitative unit is useful work within a declared resource and service envelope. For inference, that might be accepted tokens or requests meeting latency and quality thresholds per megawatt-hour. For training, it might be successful steps, samples, or completed jobs at a defined numerical and convergence target. GPU count alone is not output.

Define, for both baseline and candidate:

Useful efficiency = accepted work completed / measured facility energy

The meter boundary must be identical. The US Department of Energy defines PUE as total annual facility energy divided by annual IT equipment energy ([energy.gov](https://www.energy.gov)). For a pilot lasting hours, report facility and IT energy directly rather than treating an annualized PUE assumption as measured project evidence. MLCommons likewise requires system-level power measurement (github.com) and says power and performance should come from the same benchmark run (github.com).

Table 3 is the claim-evidence ledger that should be maintained through procurement and commissioning.

Claim or observation	Published conditions	Project applicability	Required reproduction and approval status
Up to 40% more GPUs (docs.nvidia.com)	NVIDIA illustration; same site-power envelope; PUE 1.1 and Vera Rubin testing stated as in progress (docs.nvidia.com)	Not a base-case capacity multiplier	Obtain dated model, cooling, workload, SLO, firmware and meter conditions; reproduce locally; unapproved until then
Lambda: 24% more token throughput (blogs.nvidia.com)	Vendor-reported comparison within the same facility budget	Useful directional evidence, but missing a public raw dataset and complete workload/SLO disclosure	Re-run buyer workload with raw telemetry, repetitions, uncertainty, latency and quality checks; vendor-reported only
GB200 pilot envelopes	NVIDIA runbook: 375 kW for GB200 NVL72 and 405 kW for GB300 NVL72 (docs.nvidia.com)	Starting test values, not universal facility limits	Replace with OEM and project-approved envelopes; verify calibration and protection settings
Peer-reviewed flexible-load result	Pre-DSX 256-GPU demonstration (nature.com); 25% reduction for three peak hours (nature.com)	Supports feasibility of flexible AI demand, not DSX product validation	Reproduce site event, workload, recovery and SLO conditions; contextual evidence only
Project observed result	Buyer-declared baseline, candidate, energy boundary, workload, SLO, versions and uncertainty	Directly relevant only to the approved configuration	Sign raw-data hash, test report, deviations and approval; becomes acceptance evidence

The ledger prevents three analytical errors. First, density is not throughput. Second, a lower GPU power cap may allow more racks yet still change latency, queueing, errors, or useful output. Third, site power, IT power, provisioned power, and nameplate power are not interchangeable.

NVIDIA's current public evidence should be read precisely. The MaxLPS pilot runbook gives **375 kW for GB200 NVL72** and **405 kW for GB300 NVL72** baseline envelopes (docs.nvidia.com) and requires at least **one hour** with ramp-up and cool-down removed (docs.nvidia.com). Those instructions are a useful minimum

structure, but a project may require longer steady-state windows and repeated operating days.

NVIDIA's September report identifies Lambda's result as the first MaxLPS validation and a **five-rack, 19-node cluster** (blogs.nvidia.com). No fetched Lambda-authored technical report, raw dataset, or independent audit established the complete workload, duration, latency distribution, error rate, software versions, repetitions, or uncertainty. It should remain a vendor-attributed result.

Flexible load has broader supporting evidence. A peer-reviewed field demonstration used a **256-GPU cluster** (nature.com) and reduced power by **25% for three hours** during peak demand (nature.com). That predates DSX and validates the general capability, not DSX Flex. For any live grid program, define the event baseline first because FERC guidance notes that baselines enable measurement of demand-response load reduction (ferc.gov).

Other public programs reinforce the need to test response speed and service quality together. EPRI Europe reports a live demonstration reducing AI data-center load by **30% to 40% within seconds** (epri.com), but its fetched summary does not expose the complete baseline, workload, duration, or SLO measurements. The figure is useful for scenario selection, not for guaranteeing DSX performance.

Implications and Future Directions

DSX's most important implication is organizational. The factory's useful output depends on interactions that usually belong to different teams: utility and energy management, electrical and mechanical engineering, BMS and DCIM controls, server and network operations, Kubernetes and AI platform engineering, security, finance, procurement, and workload owners. A single prime contractor may coordinate them, but the buyer remains accountable for requirements that the contract does not allocate.

As of September 2026, adoption should be gated by evidence maturity:

- **Adopt now:** Generation-matched reference design methods, versioned topology and BOM discipline, calibrated simulation, explicit interface registers, component inventory, system-level metering, and workload-based acceptance.
- **Pilot under controls:** Dynamic power allocation, DSX Exchange write paths, cross-rack power policy, and grid-responsive scheduling. Limit authority, establish interlocks, keep manual recovery, and collect raw data.
- **Contract before dependence:** Commercial support, delivery dates, repository-to-supported-product mapping, security response, compatibility windows, API and schema stability, air-gap behavior, data retention, and exit rights.
- **Do not assume:** That "open source" means every component is production-supported, that a compatible hardware list is certification, or that a digital twin prediction is an acceptance result. NVIDIA's Rack Management Service documentation itself warns that a hardware list does not establish certification, support, or continuing compatibility (docs.nvidia.com).

The public software will continue to move. KAI's public long-term-support policy is **one year from release** (github.com); NVSentinel **1.0.0** moved from Experimental to Beta/Stable (github.com). Buyers need a controlled qualification lane that can lag upstream releases, reproduce regressions, and preserve a recoverable production baseline.

Future value will depend less on the number of named modules and more on whether interfaces become stable, observable, and substitutable. OpenAPI defines a language-agnostic interface contract for HTTP APIs (openapis.org); equivalent rigor is needed for events, telemetry semantics, facility points, and grid instructions. A portable AI factory is one whose owner can export state, replace a component, replay tests, and prove continued conformance.

Energy governance should persist beyond commissioning. ISO 50001 provides a framework for establishing, maintaining, and improving an energy management system (iso.org). That favors a recurring measurement and review process over a one-time headline benchmark.

Operational observability should remain replaceable as well as durable. OpenTelemetry's vendor- and tool-agnostic positioning (opentelemetry.io) can support that goal, but only if the buyer retains schemas, collectors, configuration, and historical exports.

Frequently Asked Questions (FAQs)

Is NVIDIA DSX the same as DGX?

No. DGX is a compute-system family; DSX is a factory-scale platform covering design, simulation, operating software, power optimization, grid flexibility, and signal exchange. A DSX-aligned factory may contain DGX or NVL systems, but purchasing those systems does not complete the DSX integration or acceptance scope.

Which NVIDIA DSX modules does a private AI factory need?

There is no universal bundle. Every factory needs the underlying functions, but it may use DSX products, existing enterprise tools, partner systems, or custom integration. The decision should be made per capability using interoperability, support, operational skills, air-gap needs, lifecycle cost, and exit cost. A selective deployment still requires an interface owner and acceptance evidence for omitted DSX modules' replacement functions.

What are the main DSX MaxLPS requirements?

The buyer needs generation-specific hardware and facility envelopes, calibrated power telemetry, a controlled workload, priority and SLO policy, safe control bounds, identical baseline and candidate measurement boundaries, and rollback. NVIDIA's public runbook changes rack load and DPS enforcement while holding other experimental conditions constant (docs.nvidia.com). Exact production eligibility and support must be confirmed in writing while DPS remains documented as Developer Preview.

How should buyers validate NVIDIA's 40% claim?

Do not apply the headline percentage as a capacity multiplier. Obtain the dated test configuration, workload, cooling and PUE basis, firmware and software versions, power boundary, SLO results, duration, repetitions, and uncertainty. Then reproduce useful work within the same measured megawatt-hours and SLO. Record both positive and negative results in the claim-evidence ledger.

What belongs in an NVIDIA DSX contract?

Include the design revision, generation-specific BOM, version and support matrix, licensing, API and schema commitments, OT point list, security model and SBOM, simulation calibration, FAT and SAT procedures, measurable pass criteria, raw-data ownership, training, incident response, upgrade cadence, warranty allocation, data portability, and exit assistance. DOE measurement guidance specifically calls for documentation of every assumption and data source ([energy.gov](https://www.energy.gov)).

For HTTP interfaces, an OpenAPI definition can provide a language-agnostic contract (openapis.org), but the RFP must still cover event streams, non-HTTP protocols, compatibility, and operational semantics.

Can DSX support multi-tenancy and air-gapped deployments?

DSX OS is positioned for multi-tenant infrastructure, but implementation evidence is still required. Test RBAC, quotas, network policy, device assignment, data separation, noisy-neighbor behavior, audit linkage, and administrative boundaries. Kubernetes itself identifies RBAC, quotas, and network policy as essential shared-cluster controls (kubernetes.io). For an air gap, test installation, entitlement, local registries, offline updates, vulnerability feeds, certificate rotation, recovery, and ongoing operation with external routes physically or logically unavailable. No public platform label substitutes for those tests.

Conclusion

NVIDIA DSX is best understood as a factory-scale architecture and operating model, not a product-shaped shortcut around systems engineering. Its six layers can organize a private AI factory program, but they also expose new buyer obligations across facilities, compute, networking, Kubernetes, security, power, and grid operations.

The sound adoption decision is capability-by-capability. Freeze the hardware generation and document set. Inventory every DSX OS component with its version, license, maturity, and support owner. Treat Exchange, MaxLPS, and Flex as controlled IT/OT integrations with explicit authority, safe states, and rollback. Calibrate simulations, preserve uncertainty, and convert selected outputs into buyer-owned requirements only when measurement methods are defined.

Most importantly, purchase evidence rather than adjectives. The contract should contain the reference baseline, interface definitions, raw-data rights, security artifacts, FAT and SAT tests, training, operational proving period, and exit package. This lifecycle view is consistent with ISO/IEC/IEEE 15288 spanning conception through support and retirement (iso.org). NVIDIA's published density and efficiency figures are useful hypotheses and pilot-design inputs. They become project facts only after useful output, energy, and service quality are measured together on the delivered configuration.

That evidence hierarchy also keeps commercial choices reversible. A buyer can adopt the modules that satisfy current requirements, preserve existing systems where they perform better, and revisit the boundary as software and hardware mature. The resulting decision is a controlled engineering commitment, not an endorsement of an undifferentiated stack.

Source links

Links cited in this article, in order of first appearance. Full addresses are provided for readers using extracted text.

1. <https://nvidianews.nvidia.com/news/dsx-infrastructure-ai-factory>
2. <https://docs.nvidia.com/dsx>
3. <https://github.com/dsx-ai-factory/nv-config-manager>
4. <https://github.com/dsx-ai-factory/dsx-exchange/blob/main/SUPPORT.md>
5. <https://www.nvidia.com/en-us/agreements/cloud-services/runai-control-plane/>
6. <https://docs.nvidia.com/dsx/maxlps/overview>
7. <https://blogs.nvidia.com/blog/from-megawatts-to-tokens-how-nvidia-maximizes-ai-factory-production/>
8. <https://docs.nvidia.com/datacenter/dps/versions/latest/guides/runbooks/maxlps-simple-mode-pilot.html>
9. <https://gpusmith.com/>
10. <https://www.nvidia.com/en-us/data-center/products/dsx/>
11. <https://gpusmith.com/articles/en/dgx-vs-hgx-vs-nvl72-vs-mgx>
12. <https://docs.nvidia.com/dsx/facilities-infra/reference-design-overview>
13. <https://docs.nvidia.com/datacenter/dps/versions/latest/>
14. <https://www.nist.gov/publications/credibility-consideration-digital-twins-manufacturing>
15. <https://gpusmith.com/articles/en/nvlink-5-vs-nvlink-6>
16. <https://docs.nvidia.com/dgx/dgxgb200-user-guide/>
17. <https://www.nvidia.com/en-us/data-center/hgx/>
18. <https://gpusmith.com/articles/en/how-to-build-an-ai-datacenter>
19. <https://docs.nvidia.com/dsx-exchange/architecture>
20. <https://github.com/kai-scheduler/KAI-Scheduler>
21. <https://github.com/dsx-ai-factory/workload-map>
22. <https://github.com/ai-dynamo/grove>
23. <https://github.com/NVIDIA/aicr>
24. <https://github.com/NVIDIA/aicr/security>
25. <https://github.com/dsx-ai-factory/topograph>
26. <https://github.com/NVIDIA/NVSentinel>
27. <https://github.com/dsx-ai-factory/infra-controller>
28. <https://github.com/dsx-ai-factory/infra-controller/blob/main/docs/operations/firmware-updates.md>
29. <https://docs.nvidia.com/ai-enterprise/planning-resource/licensing-guide/latest/overview.html>
30. <https://docs.nvidia.com/ai-enterprise/release-8/latest/index.html>

31. <https://csrc.nist.gov/pubs/sp/800/82/r3/final>
32. <https://nvlpubs.nist.gov/nistpubs/SpecialPublications/NIST.SP.800-82r3.pdf>
33. https://www.dmtf.org/sites/default/files/standards/documents/DSP0266_1.19.0.html
34. <https://data.ashrae.org/bacnet/>
35. <https://reference.opcfoundation.org/specs/OPC-10000-1/4.4.1>
36. <https://csrc.nist.gov/pubs/sp/800/207/final>
37. <https://kubernetes.io/docs/concepts/security/multi-tenancy/>
38. <https://www.rfc-editor.org/rfc/rfc3411.html>
39. <https://datatracker.ietf.org/doc/html/rfc9232>
40. <https://www.cisa.gov/cybersecurity-performance-goals>
41. <https://www.vertiv.com/en-anz/about/news-and-events/news-releases/2026/vertiv-brings-converged-physical-infrastructure-to-nvidia-vera-rubin-dsx-ai-factories/>
42. <https://opentelemetry.io/docs/what-is-opentelemetry/>
43. <https://prometheus.io/docs/prometheus/latest/storage/>
44. <https://www.iso.org/standard/81702.html>
45. <https://nvlpubs.nist.gov/nistpubs/specialpublications/nist.sp.800-218.pdf>
46. https://www.ntia.gov/files/ntia/publications/sbom_minimum_elements_report.pdf
47. <https://swehb.nasa.gov/spaces/7150/pages/16450634/SWE-034%2B-%2BAcceptance%2BCriteria>
48. <https://www.ashrae.org/technical-resources/bookstore/commissioning>
49. <https://www.ashrae.org/technical-resources/standards-and-guidelines/titles-purposes-and-scopes>
50. <https://www.energy.gov/cmei/femp/measurement-and-verification-activities-required-energy-savings-performance-contract>
51. <https://www.ashrae.org/technical-resources/ai-data-center-framework/commissioning-performance-validation>
52. <https://legacy.wbdg.org/building-commissioning/commissioning-documents>
53. <https://www.energy.gov/cmei/femp/cooling-water-efficiency-opportunities-federal-data-centers>
54. https://github.com/mlcommons/inference_policies/blob/master/power_measurement.adoc
55. <https://www.nature.com/articles/s41560-025-01927-1>
56. <https://gpusmith.com/articles/en/gpu-rack-power-requirements-planning-guide>
57. <https://ferc.gov/sites/default/files/2020-04/napdr-mv.pdf>
58. <https://europe.epri.com/press-releases/dcflex-momentum-builds-across-europe>
59. <https://docs.nvidia.com/rms/documentation/reference/hardware-compatibility-list/>
60. <https://github.com/kai-scheduler/KAI-Scheduler/blob/main/SUPPORT.md>
61. https://github.com/NVIDIA/NV Sentinel/blob/main/RELEASE_NOTES_v1.0.0.md
62. <https://spec.openapis.org/oas/v3.2.0.html>
63. <https://www.iso.org/standard/69426.html>

THE PUBLISHER

About GPU Smith

GPU Smith provides independent engineering and research for private AI infrastructure. We help technical buyers and operators reason about GPU workload sizing, hardware selection, procurement, deployment and inference operations, from the compute system to the networking, power and cooling requirements around it.

Start with the workload

Useful infrastructure decisions begin with the models, data, throughput, latency and operating constraints a team actually has. GPU Smith helps connect those requirements with system design and deployment choices. The goal is a privately controlled AI environment whose capacity and operational demands are understood before a purchasing decision.

Hardware and supplier research

Our public reference library covers GPUs, complete systems and networking components. The server-vendor directory and manufacturing research help teams investigate suppliers, compare options and follow sources behind technical claims. We distinguish manufacturer specifications from measured benchmarks and advertised capabilities from tested configurations. Reference pages are research resources, not a live inventory listing or a binding equipment quote.

Deployment and operations

GPU Smith's engineering scope includes infrastructure integration, networking, power, cooling and the practical operation of inference workloads. Our published articles explain the tradeoffs behind deployment, procurement and ongoing operations so teams can ask better questions and document decisions.

Work with GPU Smith

Explore the [hardware library](#), [GPU references](#), [networking references](#) and [vendor research](#). [Contact GPU Smith](#) with your workloads and deployment constraints to discuss a project and confirm current scope, pricing and availability.

Inclusion of a manufacturer or product in our research does not imply a partnership, certification or endorsement.

Website: <https://gpusmith.com>

This article is educational. Verify consequential medical, legal, financial, regulatory and technical decisions with appropriate professionals and the cited primary sources. Company services are described separately from the article's research findings.