

The State and Future Direction of the Data Storage Research Community

Evidence from 2024–2026 and the Case for National Research Infrastructure

Jaime Cernuda¹ Suren Byna² Anthony Kougkas¹ Xian-He Sun¹

¹Illinois Institute of Technology ²The Ohio State University

Prepared within the StoreHub planning project
NSF Awards 2346504 and 2346505 (CISE Community Research Infrastructure)
<https://grc.iit.edu/research/projects/storehub>

July 2026

Abstract

We assessed the state of the data storage research community over the period 2024–2026 and projected its trajectory three to seven years forward. We combined venue program analysis, industry and market data, federal budget and award records verified against the NSF Awards API, an inventory of existing national research infrastructure, and the StoreHub community survey ($N=76$). Four conclusions follow from the evidence. First, artificial intelligence workloads have reorganized the research agenda of the storage field in under three years and have made storage performance a quantified, first-order component of national-scale computing systems. Second, the 2025–2026 memory and storage price surge, the sharpest on record, is widening the gap between the systems academics can afford and the systems whose behavior defines the research frontier. Third, the United States possesses a strong inventory of adjacent cyberinfrastructure but no facility dedicated to storage systems research: no testbed offers Compute Express Link (CXL) hardware, no successor replaced the PROBE cluster for destructive experimentation at scale, and production facilities allow researchers to run and measure workloads but not to modify the systems themselves. Fourth, community demand for such a facility is documented and specific. Based on this evidence, we conclude that a national storage research facility is warranted at the Mid-scale Research Infrastructure-1 class (approximately \$20–25M over five years), designed as a federation over existing infrastructure and distinguished from prior testbeds by an agentic-AI access layer that acquires real-system data continuously. We also state explicitly the conditions under which this conclusion should be revised.

Disclaimer. This material is based upon work supported by the National Science Foundation under Awards 2346504 and 2346505. Any opinions, findings, and conclusions or recommendations expressed in this material are those of the authors and do not necessarily reflect the views of the National Science Foundation, nor official positions of the authors' institutions.

Contents

Executive Summary	3
1 Introduction and Methodology	4
1.1 Purpose and Scope	4
1.2 Methodology and Evidentiary Standards	4

2	The Research Community in Transition, 2024–2026	4
2.1	Evolution of the Research Venues	5
2.2	The AI–Storage Research Interface	5
2.3	The Open-Source Storage Ecosystem	7
2.4	Industry Structure and Standards Formation	8
2.5	The Research Workforce Pipeline	8
3	Hardware and Economic Conditions	8
3.1	Quantified Storage Requirements of AI Workloads	9
3.2	Memory and Storage Supply Economics	9
3.3	Assessment of Emerging Storage Technologies	10
3.4	Cloud versus Owned Hardware	11
4	The Federal Funding and Infrastructure Landscape	11
4.1	NSF Programs, 2025–2026	12
4.2	Department of Energy and National Policy Context	12
4.3	The National Infrastructure Inventory	13
4.4	Costs of Comparable Facilities	14
5	The Infrastructure Gap	14
5.1	Testbeds at Realistic Scale	15
5.2	Access to Modern Hardware	16
5.3	Shared Traces and Datasets	16
5.4	Reproducibility Infrastructure	17
5.5	Community Benchmarks	17
5.6	Agentic-AI Integration	18
6	Assessment	18
6.1	The Case Against a National Center	18
6.2	The Case For	19
6.3	Verdict	19
7	Design of the Proposed Facility	19
7.1	Functions, in Priority Order	20
7.2	Integrate Rather than Build	21
7.3	Scale Points	21
7.4	Reference Configuration and Cost Estimate	21
7.5	Distinction from Prior-Generation Testbeds	22
8	Risks, Alternatives, and Conditions for Revision	23
8.1	Alternatives Considered	23
8.2	Principal Risks of Proceeding	24
8.3	Conditions for Revision	24
9	Conclusion	24
A	Consolidated Findings Register	26

Executive Summary

Between 2024 and 2026 the data storage research field underwent its fastest reorientation in decades. The share of AI-related papers at USENIX FAST, the field’s flagship venue, grew from roughly 14% in 2024 to more than 25% in 2025–2026, and AI themes won the conference’s best-paper award in both years [1, 2]. Inference state became a storage workload: key-value (KV) cache tiering across GPU memory, DRAM, flash, and object storage moved from a research topic to deployed infrastructure within twelve months of its first publication [3, 4]. Checkpointing for large-model training was codified as an industry benchmark with severe quantified requirements, including implied burst-write bandwidths of 3.6 TB/s on 100,000-accelerator clusters [5].

The same AI demand produced the sharpest memory and storage price surge on record. Server DRAM contract prices rose approximately 90% in a single quarter in early 2026; nearline hard drives sold out through calendar 2026; and enterprise flash roughly tripled in unit price [6–8]. Fixed-dollar academic equipment budgets lost an estimated 40–60% of their purchasing power for memory-rich systems between proposal and award, while hyperscaler long-term agreements absorbed supply ahead of small buyers [9].

Against this backdrop, the United States has no research facility dedicated to storage systems. The bare-metal testbeds the National Science Foundation funds (Chameleon, CloudLab) offer on the order of thirty purpose-built storage nodes between them and no CXL hardware; the production systems it funds (including the 400 PB all-flash storage subsystem of the Horizon leadership facility) allow researchers to measure but not to modify them; and the one dedicated systems-research cluster the country ever fielded at scale, PROBE, was retired around 2015 without a successor [10–17]. The StoreHub community survey ($N=76$) documents the demand side: 62% of respondents require administrator privileges that production and cloud systems categorically deny, 56% report research they cannot conduct for lack of specialized resources, and 98% state that early access to prototype storage hardware would be valuable [18].

This report weighs the case for and against new national infrastructure and concludes that a national storage research facility is warranted at the medium scale: approximately \$20–25M over five years, sized at the NSF Mid-scale Research Infrastructure-1 class and structured as a federation over existing infrastructure rather than a freestanding construction project. The proposed facility combines four functions: a dedicated, reconfigurable, destructible testbed of 200–400 storage-dense nodes with modern media; a national I/O observatory that trades instrumentation and benchmarking services to production facilities for traces and scale access; an agentic-AI access layer, generalizing patterns demonstrated by the NSF-funded IOWarp project, through which experiments self-document into community datasets and access extends to institutions without systems staff; and sustained community and workforce programs. An illustrative reference configuration (200 storage-dense nodes, roughly 54 PB across four tiers, 32 accelerators) prices at \$7.0M at mid-2025 component prices and \$9.6M at mid-2026 prices, quantifying both the capital requirement and the cost of the current supply shock. The report also states, explicitly, the evidence that would weaken or reverse this conclusion.

1 Introduction and Methodology

1.1 Purpose and Scope

This report was produced within the StoreHub planning project, funded under the NSF CISE Community Research Infrastructure (CIRC) program to assess whether the data storage research community requires new shared research infrastructure and, if so, of what kind [19]. It is written, however, as a standalone public assessment: the evidence, analysis, and conclusions are intended to be usable by any member of the community—researchers, program managers, facility operators, and policymakers—evaluating the state of storage research or the case for infrastructure investment, independently of the project that produced it. The report covers the period January 2024 through July 2026 as observed data and projects three to seven years forward (2029–2033). Each major section opens with a short overview of findings that can be read on its own; the body of the section supplies the evidence.

The report addresses four questions. What happened to the storage research field, its venues, its open-source ecosystem, its industry, and its workforce during 2024–2026 (Section 2)? What hardware and economic conditions must storage research now build on (Section 3)? What federal funding programs and existing infrastructure define the environment into which any new effort would launch (Section 4)? And does the combined evidence support a national center for storage research (Sections 5–8)?

1.2 Methodology and Evidentiary Standards

Claims in this report are dated and cited to primary or trade-press sources. Award numbers, amounts, and dates were verified against the NSF Awards API where possible. Federal budget figures distinguish presidential requests from congressional appropriations. Company-reported figures (valuations, revenue, performance multipliers) are identified as unaudited where they appear. Observed facts are separated from projections throughout; the consolidated findings register in Appendix A tags each finding accordingly. Known evidence gaps are stated where they occur rather than omitted.

Community demand data derive from the StoreHub Community Insight Survey ($N=76$), comprising a longitudinal instrument ($N=35$) fielded online and at HUG’24, SC’24, and SSDBM’25, and a live poll ($N=41$) conducted at the first StoreHub Community Workshop in June 2025. Respondents spanned fifteen universities and four national laboratories [18]. These are the community’s best available demand-side data; they are project-collected and are treated as survey evidence rather than as an independent census.

Two independent research passes over this landscape were conducted with different toolchains and reconciled. Where their quantitative results disagreed, the figure verifiable against a primary source was retained and the discrepancy is noted in the text; quantitative claims that could not be verified against primary sources were excluded.

2 The Research Community in Transition, 2024–2026

Overview of findings. *The storage research field reorganized around AI in under three years: AI-related work grew from roughly 14% to more than 25% of the flagship venue’s program and took its best-paper awards in 2025 and 2026, while inference state (KV caching) became a deployed storage workload within twelve months of first publication. Agent-mediated data access (MCP) became the de facto industry interface with no peer-reviewed storage semantics behind it yet. Open-source developer energy moved to AI-native data layers while the parallel file systems national facilities depend on are sustained by one or two vendors each. The workforce pipeline is rebalancing toward industry: record PhD production, but new enrollments down 15% and 61% of graduates leaving academia.*

2.1 Evolution of the Research Venues

The clearest quantitative signal comes from USENIX FAST. FAST ’24 accepted 22 papers from 123 submissions (17.9%), of which approximately three were AI-related, with a panel titled “Storage Systems in the LLM Era” foreshadowing the shift [20]. One year later, FAST ’25 accepted 36 papers from 167 submissions (21.6%) and ran two dedicated machine-learning sessions totaling nine papers, one quarter of the program, spanning vector search, LLM KV-cache storage, and GPU-native file systems [1]. The best-paper award went to Mooncake, a KV-cache-centric disaggregated serving architecture deployed in production for the Kimi LLM, the first time FAST’s top award recognized LLM-serving infrastructure [3]. By FAST ’26, the program carried dedicated “AI and LLMs” sessions covering LLM checkpointing, SSD-based LLM serving, and GPU checkpoint/restore, together with an “SSDs and CXL” session; one of two best papers concerned LLM-generated file systems [2]. The AI share of FAST thus moved from roughly 14% to more than 25% of accepted papers in two cycles, and AI themes took best-paper honors in both 2025 and 2026.

The pattern holds across venues. HotStorage ’24 gave its best-paper award to work on LLM-driven tuning of log-structured merge key-value stores [21]; HotStorage ’25 carried multiple CXL and memory-interconnect papers alongside vector-search and LLM position papers [22]. PDSW ’25 at SC25 featured LLM checkpointing, flexible-data-placement SSD write isolation, an object-storage roofline study, and closed with a panel on storage architectures for AI [23, 24]. SOSF ’24 included roughly seven storage and memory papers among a large LLM-training and serving cohort [25]. Two structural datapoints warrant attention. MSST, after its fiftieth-anniversary meeting in 2024, ran no research track at all in 2025 and returned in 2026 with a smaller sixteen-paper program, evidence of the fragility of mid-tier storage venues [26–28]. SSDBM renamed itself in 2025 from “Scientific and Statistical Database Management” to the International Conference on *Scalable Scientific Data Management*, a rebranding that reflects where the field’s energy has moved [29].

Across 2024–2026, the topics that rose were storage for AI training and inference (checkpointing, data loading, and KV-cache tiering, now a distinct subfield), vector and embedding storage, CXL and memory tiering, flexible-data-placement SSDs, and deployed-systems papers from hyperscalers. Persistent-memory research collapsed following Intel’s Optane cancellation; FAST ’26 carried no persistent-memory session, and the remaining papers repurpose or characterize orphaned hardware. Zoned-namespaces (ZNS) work plateaued as attention shifted to flexible data placement. Flash internals, caching, deduplication, erasure coding, and archival media (including two DNA-storage papers at FAST ’25) held steady as persistent niches. One evidence gap is noted: FAST ’26 official acceptance statistics conflict across public sources, and SC25’s storage-track composition could not be fully itemized.

2.2 The AI–Storage Research Interface

The community’s own composition reflects the transition: in the StoreHub survey, 81% of respondents report using AI or machine learning in their research and 47% identify AI applications (training data, checkpointing) among their primary research focus areas (Figure 1) [18].

Storage for AI. Checkpointing for large-model training became a first-class research topic driven by production failure rates. ByteDance’s ByteCheckpoint reported checkpoint stalls of intolerable minute-scale duration in production and demonstrated save-time improvements of up to $529\times$ [30]; Microsoft demonstrated just-in-time checkpointing at EuroSys ’24 [31]. MLCommons codified the workload in MLPerf Storage v2.0 (August 2025), whose failure model implies a checkpoint every 9.3 minutes on a 16,384-accelerator cluster to hold overhead below 5% [5]. The signature development of the period is the emergence of inference state as a storage workload: Mooncake and LMCache established multi-tier KV-cache storage across GPU memory, DRAM, NVMe, and object storage, and the pattern was productized within twelve months

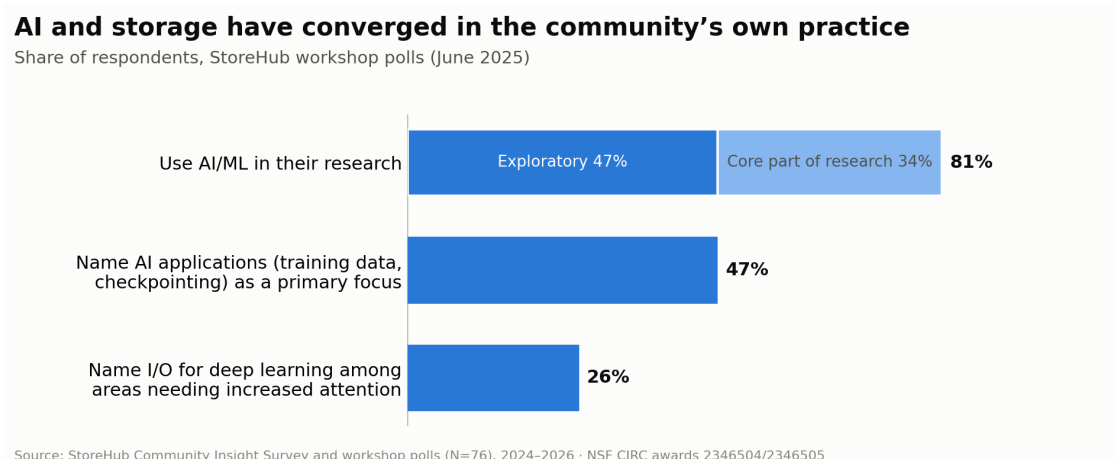


Figure 1: AI adoption among surveyed storage researchers [18].

by NVIDIA (NIXL and the Dynamo framework), WEKA, VAST Data, DDN, and Google Cloud [3, 4, 32].

AI for storage. The application of machine learning to storage-system internals bifurcated during the period. The learned-index line of work, which since 2018 had promised to replace classical index structures with learned models, entered a consolidation phase: the 2024–2025 literature at SIGMOD and PVLDB is dominated by skeptical evaluation, robustness analysis, and benchmarking of existing designs under updates and shifting distributions, rather than by new index families—the pattern of a subfield maturing from promise into engineering discipline. The frontier energy migrated instead to language-model-based system administration and tuning. Modern storage engines expose configuration spaces of hundreds of interacting parameters governing compaction, flushing, caching, and resource sharing, and navigating those spaces has historically required scarce human expertise. A line of work beginning with the HotStorage ’24 best paper, which asked whether modern LLMs can tune and configure log-structured merge key-value stores [21], progressed within eighteen months to full-cycle frameworks: ELMo-Tune-V2 combines LLM-driven workload characterization, iterative configuration refinement, and real-time adjustment for RocksDB, reporting throughput improvements of up to approximately $14\times$ over default configurations on standard benchmarks and 26–34% end-to-end gains for applications built above the store [33], with successor systems extending agent-driven tuning across heterogeneous storage engines. The research question this line poses—whether the operational knowledge of storage systems can be captured and applied by language-model agents—connects directly to the agentic-interface developments discussed below, and answering it requires exactly the instrumented, reconfigurable systems that Section 5 finds scarce.

Vector and semantic storage. Vector search moved from specialist libraries into mainstream data-systems venues and cloud platforms during the period, and its storage dimension became explicit. On the research side, disk-resident vector indexing emerged as a distinct problem class: Starling (SIGMOD ’24; arXiv:2401.02116) demonstrated I/O-optimized graph-index layouts serving 33 million vectors from a 10 GB disk segment at sub-millisecond latency, reporting $43.9\times$ throughput gains over prior disk-based designs, and vendor systems presented billion-vector production deployments at VLDB ’25. The capacity arithmetic explains why this is a storage problem: one billion 1,024-dimension single-precision embeddings occupy roughly 4 TB before indexing, and IBM Research’s 2025 demonstration of a 100-billion-vector index carried a 153 TiB footprint [34]. Commoditization followed within the same window: AWS moved S3 Vectors from preview to general availability at up to two billion vectors per index, claiming a 90% cost reduction

relative to dedicated vector databases [35], and each major cloud added comparable primitives. Practitioner experience converged on a corrective consensus—most production retrieval systems index fewer than ten million documents, and vector search is becoming a feature of general-purpose data systems rather than a standalone database category. For storage research, the durable questions are accordingly less the index structures than the substrate beneath them: embeddings constitute a new persistent data type whose tiering across memory and flash, whose update and re-embedding semantics, and whose index-maintenance economics at scale remain unsolved, and these are precisely the questions that require realistic hardware to study.

Agent-native storage. Anthropic’s Model Context Protocol (MCP), announced November 25, 2024, became the de facto agent-to-data interface within one year: it was adopted by OpenAI, Google, Microsoft, and AWS, and donated to the Linux Foundation’s Agentic AI Foundation on December 9, 2025 [36, 37]. Storage products now ship MCP servers. As of mid-2026, however, no landmark peer-reviewed storage-venue paper defines agent-native storage semantics: consistency, provenance, access control, and context-aware data reduction for agent-mediated access remain open. The earliest funded systems effort in this space is IOWarp (NSF CSSI Frameworks, awards 2411318/2411319, approximately \$5M, 2024–2029), which couples a content assimilation and transfer architecture with an extensive MCP deployment: more than sixteen scientific MCP servers and 150 tools spanning HDF5, ADIOS, Slurm, Darshan, and workflow systems [38–40]. Its verified mechanisms include system-monitoring MCP servers that expose live DRAM, NVMe, GPU, and parallel-file-system telemetry to agents; a Darshan MCP server giving agents access to real I/O traces; and provenance capture and replay for agentic workflows. These constitute a working model for automatic acquisition of real-system data by agents rather than through manual benchmarking. Two caveats apply: IOWarp’s published performance claims are self-reported, and its peer-reviewed agentic footprint remains thin.

2.3 The Open-Source Storage Ecosystem

The most-starred open-source storage project effectively left open source in approximately thirteen months. MinIO removed administrative features from its AGPL community edition (March–June 2025), shipped a final community release in October 2025, entered maintenance mode in December 2025, and archived its repository, read-only at 61,100 stars, on April 25, 2026 [41–43]. Redistribution was immediate: permissively licensed alternatives (RustFS, SeaweedFS, Garage) and Ceph RGW absorbed migrations, and practitioners now price single-vendor open-source storage as a sustainability risk.

The parallel file systems underpinning U.S. research computing are healthy but institutionally captive. Ceph sustains an annual major-release cadence under the Ceph Foundation with IBM as the dominant employer of core developers [44]. Lustre ships on an approximately thirteen-month cadence with DDN’s Whamcloud as the principal development organization [45]. DAOS survived Intel’s exit through divestment to HPE under the Linux Foundation’s DAOS Foundation and runs in production on Aurora at 230 PB, yet its public community is small, roughly 950 GitHub stars [46]. Developer energy concentrates instead in AI-native data layers: every major vector database out-stars every distributed file system, and 2024–2026 venture capital flowed to AI data-path software (LanceDB \$30M, Tigris \$25M) rather than general-purpose file systems [47, 48]. Object-storage semantics are displacing POSIX as the default substrate for new systems, from embedded storage engines whose only disk is an object store to language runtimes shipping S3 clients as built-ins [49, 50].

2.4 Industry Structure and Standards Formation

AI repriced the storage-vendor market between January 2025 and April 2026. DDN accepted \$300M from Blackstone at a \$5B valuation [51]; VAST Data closed an approximately \$1B round at a \$30B post-money valuation in April 2026, with NVIDIA participating [52, 53]. NVIDIA’s AI Data Platform reference design and certification program (March 2025) enrolled effectively every major storage vendor, which now design to NVIDIA certification tests rather than POSIX benchmarks; its NIXL library provides a unified interface for moving KV-cache data across the memory-storage hierarchy [4, 54]. The vendor-neutral counterweight is SNIA’s Storage.AI project (August 2025), founded by fifteen companies, conspicuously without NVIDIA, whose named problem areas (KV-cache placement, GPU-initiated I/O, data-services offload) map directly onto open research questions [55].

Frontier AI laboratories treat storage as a competitive in-house discipline. Meta is the exception that publishes: its engineering literature documents dedicated checkpoint storage at 240 PB with 2 TB/s sustained throughput and, in July 2026, checkpoint-burst congestion control and cache-hit-rate figures for inference serving [56, 57]. OpenAI publishes nothing on storage but maintains a full compute-storage engineering organization [58]. The structural implications for academia are two: no university has access to systems at the 10,000–100,000-GPU scale where the field’s characteristic failure modes occur, and no I/O traces from frontier training runs have ever been published. The community’s most productive feedback loops operate precisely where open data exists, notably MLPerf Storage and the Backblaze drive-reliability dataset [59].

2.5 The Research Workforce Pipeline

CRA Taulbee data released in June 2026 show a record 1,909 computer science PhDs awarded in academic year 2024–25 alongside a 15% decline in new doctoral enrollments and a record 61% of employed new PhDs entering industry [60]. Taulbee does not break out storage as a subfield; storage-specific pipeline claims are therefore inferential, and the absence of such tracking is itself a measurement gap inexpensive for NSF to close. Demand is barbell-shaped: growing on the AI-infrastructure side, where frontier laboratories offer new PhDs six- and seven-figure packages and observers document academic brain-drain concern [61], and shrinking on the traditional side. The NSF-sponsored community visioning workshop called in 2018 for focused educational and training activity in storage systems [62]; the 2024–2026 evidence indicates that the concern has materialized, with storage expertise accumulating in AI laboratories and vendors rather than universities.

3 Hardware and Economic Conditions

Overview of findings. *The storage demands of AI systems are now quantified and severe: checkpoint bursts approaching 3.6 TB/s at frontier scale, roughly 40 GB of cache state per long-context inference request, and minimum per-accelerator bandwidths now defined by industry benchmarks. The 2025–2026 memory and storage price surge is the sharpest on record—server DRAM up roughly 90% in one quarter, enterprise flash near \$500/TB, nearline drives sold out through 2026—and it disproportionately harms academic buyers, who lost 40–60% of purchasing power. Among emerging technologies, CXL memory expansion is real and in production while computational storage collapsed as a category; no NSF testbed offers CXL hardware. Cloud resources complement but cannot substitute for owned hardware: device-level, kernel-level, and fault-injection research is infeasible on rented infrastructure at any price.*

3.1 Quantified Storage Requirements of AI Workloads

Meta’s Llama 3 publication provides the clearest published evidence of scale. The 405-billion-parameter pre-training run occupied 16,384 H100 GPUs for 54 days and suffered 466 job interruptions, 419 of them unexpected, roughly one failure every three hours, with 58.7% attributable to GPUs and high-bandwidth memory [56]. At such failure rates, frequent checkpointing is the only path to acceptable machine utilization; Meta reports bursty checkpoint writes that saturate a dedicated 240 PB storage deployment delivering 2 TB/s sustained and 7 TB/s peak throughput. MLPerf Storage v2.0 converted this operational reality into a benchmark, defining checkpoint sizes of approximately thirteen bytes per parameter (15 TB at one trillion parameters) and implying, for a trillion-parameter model on 100,000 accelerators at 5% overhead, roughly 3.6 TB/s of burst write bandwidth and more than 14 PB of checkpoint writes per day [5]. On the ingest side, MLPerf Storage v1.0 (September 2024), the first broad architecture-neutral measurement of AI training I/O, implies roughly 2.7 GB/s of read bandwidth per H100 for the most demanding workloads; the v2.0 round drew 26 organizations and more than 200 results [59, 63]. At the frontier, DeepSeek’s open-sourced Fire-Flyer File System reports 6.6 TiB/s of aggregate read throughput from 180 storage nodes; these figures are self-reported, with open-source code but no independent replication [64].

Inference moved the pressure to memory tiering. A single 128K-token context on a 70-billion-parameter model requires roughly 40 GB of KV cache; four concurrent long-context requests exceed the size of the quantized model weights. KV cache therefore spills from GPU memory into DRAM, local flash, and networked storage, making storage systems a direct determinant of inference throughput. Mooncake demonstrated 59–498% gains in effective request capacity under latency service-level objectives, in production across thousands of nodes serving more than 100 billion tokens per day [3]; LMCACHE reports 3–10× reductions in time-to-first-token from GPU-DRAM-disk tiering [32]; NVIDIA’s Dynamo framework makes the hierarchy explicit from HBM through object storage [65]. Retrieval-augmented generation adds parallel capacity pressure: one billion 1024-dimension floating-point embeddings occupy roughly 4 TB before indexing, and IBM Research demonstrated a 100-billion-vector index with a 153 TiB footprint [34]. On the capacity side, hard-drive exabyte shipments reached approximately 1,621 EB in calendar 2025, up 22% year over year, with nearline drives constituting roughly 85–90% of exabytes shipped [66].

3.2 Memory and Storage Supply Economics

As late as November 2024, TrendForce forecast DRAM price declines through 2025 [67]. AI server demand invalidated the forecast within months. Observed contract-price movements: server DRAM rose 43–48% quarter-over-quarter in 4Q25 [68]; conventional DRAM rose 90–95% quarter-over-quarter in 1Q26, a record [6]; 2Q26 added a further 58–63% [69]. At the module level, a 64 GB DDR5 RDIMM moved from roughly \$255 (3Q25) to more than \$900 (1Q26) [70, 71]. The structural driver is a capacity reallocation rather than a demand blip: high-bandwidth memory consumes roughly three times the wafer area per bit of DDR5, HBM’s share of DRAM wafer starts is rising toward 30% by end-2027, and SK hynix and Micron report their entire 2026 HBM output sold out [72, 73]. TrendForce guides 2027 RDIMM bit-supply growth below server demand growth, implying the shortage persists into 2027; this is a forecast, not observed data [74].

NAND and hard drives followed. After 2023 losses, major memory suppliers cut investment or output sharply: Micron’s fiscal 2023 net capital expenditures fell from \$11.98B to \$7.01B, a reduction of more than 40%; SK hynix cut 2023 investment by at least 50%; and Western Digital reduced flash wafer starts by 30% [75–77], leaving no slack when AI enterprise-SSD demand arrived: 1Q26 NAND contract prices rose 55–60% quarter-over-quarter, with a further 70–75% in 2Q26 [69]. Nearline hard-drive lead times stretched beyond 52 weeks, with Western Digital and Seagate reporting capacity effectively sold out through calendar 2026 under long-term agreements extending to 2027–2028 [7, 78]. Enterprise SSD pricing moved from roughly \$80–150/TB in mid-2025 to the \$500/TB class for high-capacity drives by 1Q26, with SSD dollars-

Technology	Verdict (mid-2026)	Anchoring evidence
CXL direct-attach memory	Shipping; production	Azure M-series preview (Nov. 2025); Meta production ASIC (2026) [79, 80]
CXL pooling / fabrics	Credible roadmap; pilots only	SC25 demonstrations; CXL 4.0 fabrics target 2027+
Memory-tiering software	Mature; in production fleets	Mainline Linux; Google fleet-wide since 2016 [81, 82]
Computational storage (category)	Stalled; survives as embedded feature	SmartSSD withdrawn; NGD folded; LANL niche persists [83, 84]
MRAM	Shipping; permanently niche	Everspin ~\$55M/yr revenue [85]
SCM successors (XL-Flash)	Credible roadmap; no volume	10M-IOPS drive sampling 2H26 [86]
DNA storage	Research stage	Write cost estimates $\geq$ \$100M/TB [87, 88]
Glass/ceramic archival	Credible roadmap; pre-product	Cerabyte 1 PB/rack pilot 2025/26 [89]

Table 1: Emerging storage technology assessments as of mid-2026.

per-terabyte exceeding twenty times hard-drive levels [8].

For research budgets the arithmetic is direct. One terabyte of server DRAM cost approximately \$4,100 in 3Q25 and more than \$14,400 by 1Q26. A cluster specified at proposal time in 2024 may be 40–60% over budget at award time in 2026; industry reporting describes memory price increases above 200% reshaping HPC storage architectures, with supply partners advising planning for 10–20% monthly increases through end-2026 [9]. Allocation priority compounds the problem: hyperscaler long-term agreements absorb supply first, so universities face both higher prices and 36–52-week lead times, and buyers are downshifting module sizes, which causes academic testbeds to diverge in configuration from the production systems whose behavior researchers study. No public NSF or DOE statement on memory-price procurement impact was found as of July 2026; this is a data gap rather than evidence of absence.

3.3 Assessment of Emerging Storage Technologies

The 2024–2026 window resolved much of the uncertainty surrounding emerging storage technologies. Technologies became real when the price shock supplied an economic forcing function and no application changes were required; they stalled when they demanded ecosystem changes or were orders of magnitude off on cost. Table 1 summarizes the assessment.

CXL merits elaboration because it defines a concrete infrastructure gap. After a genuine trough in expectations [90], the DRAM price crisis revived CXL in its simplest form. Microsoft enabled CXL-attached memory in Azure M-series virtual-machine preview in November 2025, the first announced cloud deployment [79]; Meta disclosed in June 2026 a custom CXL 2.0 ASIC in production that recycles decommissioned DDR4 modules into DDR5-only servers, motivated by the price spike rather than by the disaggregation vision that drove the original enthusiasm [80]. Sobering context remains: roughly two-thirds of servers shipped in early 2025 were CXL-capable but nearly none had CXL enabled, and multi-host pooling remains at the demonstration stage [91]. CXL modules are largely allocated to hyperscalers, with open-market availability through 2026 essentially limited to evaluation samples. The salient fact for NSF: neither CloudLab nor Chameleon lists dedicated CXL hardware as of mid-2026, and most academic CXL papers run on NUMA emulation [10]. This is a concrete, addressable infrastructure gap. Conversely, memory-tiering software is the quiet success of the period, an area in which academic work has a track record of landing in the mainline Linux kernel and which arguably offers the highest leverage per research dollar [81]. The collapse of computational storage as a marketed category is a cautionary precedent for proposals that invoke

industry waves [83].

3.4 Cloud versus Owned Hardware

The case for shifting academic storage research onto rented infrastructure is strongest on affordability and scale. Owned hardware now costs 1.5–3.5 times its 2024 price with year-long lead times, while cloud pricing has so far absorbed the shock. Five-year total-cost models find hardware to be only about 35% of the true cost of an owned cluster once power, cooling, and four to six staff FTEs are counted [92], and few storage groups can fund a systems administrator from a typical award. The access infrastructure exists and was recently renewed: CloudBank 2.0, a \$20M five-year award announced in April 2025, brokers access to five commercial clouds for roughly 500 projects annually [93], and the NAIRR pilot has supported more than 600 projects since January 2024 [94]. Cloud also offers the one thing academia cannot own: the ability to rent a thousand-node experiment for hours, the configuration in which emergent behaviors such as metadata storms and tail latency appear.

The case against is economic and epistemic. For sustained, data-heavy workloads, the exact profile of storage research artifacts under continuous benchmarking, owned hardware wins: S3-class object storage costs approximately \$282 per terabyte-year [95], which owned bulk capacity undercuts within one to two years even at inflated 2026 prices, and long-running endurance, aging, and trace studies are precisely the workloads per-hour metering punishes [96]. Egress compounds the cost: at typical rates of \$0.05–0.09 per gigabyte, routine export of a 50 TB working set costs thousands of dollars per round trip, and researcher egress waivers are capped [97]. The epistemic objection is decisive for a subset of the field: cloud NVMe exposes no firmware, no zoned namespaces, no host-managed SMR, no CXL fabric control, and no power-fault injection. The published literature reflects this, with ZNS and CXL-SSD papers running on emulators for lack of open hardware. Multi-tenant noise undermines reproducible baselines, and renting concedes the research agenda to the abstractions vendors choose to expose, at precisely the moment the field’s open problems live below those abstractions.

The evidence therefore supports a hybrid posture rather than substitution: cloud, through CloudBank and NAIRR, for scale experiments and work above the block and object interfaces; owned or testbed hardware for anything touching device internals, kernel stacks, new interconnects, failure semantics, or reproducible baselines. Two named holes remain as of mid-2026: no NSF testbed offers CXL hardware, and nothing replaced PRObE, the retired thousand-node facility for large-scale destructive storage experimentation [11, 12]. The NSF-sponsored community vision report explicitly recommended shared storage-research infrastructure [62]; no federal report from 2024–2026 recommends replacing owned testbeds with commercial cloud.

4 The Federal Funding and Infrastructure Landscape

Overview of findings. *The NSF program created specifically for community research infrastructure (CIRC) is archived. Five programs remained active through 2025–2026, issuing solicitations or making awards: CSSI, Mid-scale RI-1, ACSS, IDSS, and the NAIRR ecosystem. Across the national policy record of the period—NAIRR, the AI Action Plan, the Genesis Mission, the CREATE AI Act—storage appears as supporting language but never as a funded priority in its own right. The country holds a strong inventory of adjacent infrastructure (Chameleon, CloudLab, FABRIC, ACCESS, DOE facilities, OSN) but no facility dedicated to storage research, and the one dedicated systems testbed it ever fielded (PRObE) was retired without a successor. The costs of comparable facilities are well documented: operations run roughly twice acquisition over a facility’s life, and storage constitutes 8–10% of leadership-system cost.*

4.1 NSF Programs, 2025–2026

The environment into which any new storage-research infrastructure would launch is defined first by the fate of the program purpose-built to fund it. The CISE Community Research Infrastructure program (CIRC, solicitation NSF 23-589) was designed for community infrastructure of exactly this kind, with award classes from \$50K planning grants to \$2–5M Grand awards and an anticipated \$24M annual budget [98]. Its design carries two templates any center proposal should internalize: Medium and Grand awards must devote 20–25% of budget to community outreach, and Grand awards gate their final two years on a site visit and an approved sustainability plan. As of July 2026, however, the CIRC program page lists the program as archived; the scheduled September 2025 deadline did not run, and no successor solicitation has been posted. Among the 2024 CIRC cohort was a single storage-focused award: the StoreHub planning grant (awards 2346504/2346505, Illinois Institute of Technology and The Ohio State University, \$100K total, 2024–2026), within which this report was produced [19]. CIRC planning awards were designed to lead to CIRC Medium or Grand submissions; that path currently has no open solicitation to receive it, while the alternative path CIRC itself anticipated, planning toward Mid-scale Research Infrastructure, remains open. The situation illustrates a general condition: community-infrastructure planning efforts across CISE, not only in storage, currently lack the follow-on program they were designed to feed into.

Five programs remained active through 2025–2026, issuing solicitations or making awards. CSSI (NSF 22-632) funds software institutes at up to \$5M and demonstrably funds storage and I/O work at that scale, as the IOWarp Frameworks award shows; its December 2025 cycle ran on schedule, a meaningful signal of continuity [38, 99]. Mid-scale RI-1 (NSF 24-598) funds implementation awards of \$4M to just under \$20M on a biennial cycle, with the next preliminary-proposal deadline on September 1, 2026, making it the nearest-term large-infrastructure opportunity in the current program calendar [100]. Mid-scale RI-2 (\$20M–\$100M) anticipates no new awards in FY2026, with its next solicitation unreleased [101]. ACSS (NSF 24-583) Category II explicitly funds innovative prototypes and testbeds with novel technologies [102], and IDSS (NSF 25-544) is the Office of Advanced Cyberinfrastructure’s first data-infrastructure-specific solicitation of the period [103]. In ACCESS, storage appears only as allocatable resources rather than as a service track, and no ACCESS resource is a storage-research platform [104, 105]. The NAIRR pilot reached more than 700 projects across all fifty states by June 2026; NSF leadership stated that NAIRR is transitioning to a permanent program with funding doubling in the fiscal year, and a \$35M operations center award is pending [106, 107]. Storage in NAIRR is subsidiary: file systems and object stores attached to compute, and dataset hosting, not a resource category in its own right. The AI Institutes provide the scale model rather than the funding path: roughly 27–29 institutes at up to \$20M over five years each demonstrate that NSF will sustain \$4M-per-year thematic centers, and equally that storage and systems research has not yet won one [108, 109].

The budget context explains the volatility. The May 2025 presidential request sought a cut of approximately 56% to NSF [110]; Congress instead enacted \$8.75B (a 3.4% reduction) in January 2026, but the agency that emerged had abolished its 37 divisions, terminated roughly 1,750 grants, and lost most of its rotator positions, and in April 2026 the National Science Board was dismissed while the FY2027 request again proposes deep cuts [111–113]. The practical consequence for infrastructure planning is that program continuity can no longer be assumed; durable efforts will be those designed to draw on multiple complementary programs and to remain viable across solicitation gaps of the kind that affected CIRC.

4.2 Department of Energy and National Policy Context

The Department of Energy’s storage momentum is on the procurement side. NERSC-10 (“Doudna,” delivery late 2026) pairs IBM Storage Scale with a VAST Data quality-of-service storage system offering schedulable performance guarantees, the first VAST deployment in DOE HPC [114, 115]. OLCF’s Discovery (2028) will

ship the first factory-built storage product embedding open-source DAOS [116], and Argonne hosts NVIDIA and Oracle partnership systems announced in October 2025 [117]. The installed base already includes the largest storage systems ever fielded: Aurora’s 230 PB DAOS at roughly 31 TB/s, Frontier’s approximately 700 PB Orion Lustre system, and Perlmutter’s 36 PB all-flash scratch [118–120]. On the research side, post-Exascale-Computing-Project software stewardship keeps HDF5, ADIOS2, and Darshan alive within the NGSST/CASS structure, with no publicly verifiable budget line [121, 122]; ASCR’s single largest storage-relevant research call of the period was a \$35M laboratory announcement for data management, storage, and visualization in January 2025 [123]; and ASCR’s January 2022 workshop on scientific data management set an intellectual agenda closely aligned with the academic community’s direction [124]. Congress funded ASCR at \$1.12B for FY2026, a 9.8% increase, even as the research mix tilts toward AI [111, 125].

Across the entire 2024–2026 national compute-policy corpus, storage appears consistently as supporting language but never as a funded priority in its own right. The July 2025 AI Action Plan and its executive orders frame infrastructure as a datacenter-buildout and export problem [126]; the Genesis Mission executive order (November 2025) directs DOE to build a platform spanning the seventeen national laboratories and required a ninety-day inventory of federal computing, storage, and networking resources [127]; the CREATE AI Act, which would give NAIRR a statutory basis, passed the House Science Committee in June 2026 and was reintroduced in the Senate in April 2026 [128, 129]; the Senate’s American Science Acceleration Project makes data its first pillar [130]. The policy spotlight remains GPUs, datacenters, and power [131, 132]. The compute-divide literature supplies the demand-side corollary: academia has exited compute-intensive AI, with one survey finding 87.8% of researchers sharing roughly 9.7% of GPUs [133, 134]. The storage version of that argument, that academics cannot observe or reproduce the I/O behavior of frontier systems, is implied throughout the policy debate but rarely stated; this report states it explicitly.

4.3 The National Infrastructure Inventory

An unusually strong inventory of adjacent infrastructure exists, which means the design question for a storage center is not what to build from scratch but which specific gap none of the existing facilities covers.

Chameleon Cloud is the closest existing platform: hundreds of bare-metal reconfigurable nodes, roughly thirty dedicated storage nodes, storage-hierarchy nodes mixing device classes, composable hardware since August 2024, and documented storage-experiment patterns, funded through approximately 2028 by a \$12M Phase 4 award [13, 14, 135, 136]. Its limits are equally clear: storage nodes are a small fraction of the testbed, devices are single-terabyte class, no at-scale parallel file system can be stood up, and contention is heavy. CloudLab offers more than 1,800 bare-metal nodes with local storage on every node, likewise funded through 2028 (Phase IV, \$12M) [137, 138], with storage hardware incidental rather than purpose-built. FABRIC, a Mid-scale RI-1 network testbed of approximately \$21.8M, carries attachable NVMe as a slice component and is the natural wide-area data-movement layer for any distributed storage facility [139]. ACCESS resource providers are production systems, valuable as I/O-workload observatories and closed to systems modification; their storage highlights include Stampede3’s 13 PB all-flash VAST deployment and Bridges-2’s extension through 2027 [140, 141]. The most significant recent addition is Horizon, the \$457M NSF Leadership-Class Computing Facility, whose 400 PB all-solid-state storage subsystem makes NSF’s flagship the first U.S. academic system with DOE-class storage scale; production began in spring 2026 [15–17]. The Open Storage Network operates S3 pods at seventeen sites, grown from an unusually small federal seed of roughly \$1.8M; it is a data-hosting service rather than a research platform, S3-only, with no low-level pod access [105, 142]. DOE user facilities operate the largest storage systems in existence, accessible to academics through allocation programs strictly as production assets: researchers may run workloads on them and measure their performance, but may not alter the system software, reconfigure the hardware, or run experiments that risk disruption. The exception that proves a model exists is OLCF’s ACE testbed program, which explicitly sandboxes emerging compute, network, and storage architectures outside

Anchor	Scale	Notes
OSN federal seed (2018)	~\$1.8M	National S3 network bootstrap; hosting, not research [142]
CIRC Grand	\$2–5M / 5 yr	Program archived; caps remain the community-infrastructure reference [98]
CSSI Frameworks (IOWarp)	~\$5M / 5 yr	Software institute; no hardware center [38]
CloudLab, per phase	\$10–12M / 4–6 yr	Four phases $\approx$ \$42.8M total [138]
Chameleon Phase 4	\$12M / 4 yr	More than \$42M cumulative [136]
Mid-scale RI-1	\$4M–<\$20M	Next preliminary deadline Sept. 1, 2026 [100]
AI Institute	~\$20M / 5 yr	Proven center template [108]
FABRIC (RI-1)	~\$21.8M	Plus international extension [139]
Mid-scale RI-2	\$20M–<\$100M	No FY2026 awards anticipated [101]
LCCF Horizon	\$457M	Includes 400 PB storage subsystem [15]
NERSC operations	~\$130M / yr	FY2024 level [145]

Table 2: Costs of comparable national research infrastructure (amounts as obligated or announced).

production constraints [143]. Internationally, France’s Grid’5000, with reservable extra disks on roughly 800 bare-metal servers, is the closest analog to a storage-capable research testbed [144].

The gap, stated plainly: the United States once had a dedicated systems-research testbed at scale, PRObE, approximately 2010–2015, serving some 400 users on repurposed national-laboratory clusters [11, 12], and has had nothing comparable since. Between small reconfigurable testbeds and untouchable production giants lies a specific unoccupied niche: dedicated, destructible, reconfigurable storage systems at hundreds-of-nodes and tens-of-petabytes scale, with modern media (NVMe, CXL, flexible-data-placement SSDs) and fault-injection capability. Every facility above is a potential partner precisely because none of them is this.

4.4 Costs of Comparable Facilities

The costs of comparable facilities, ordered by scale, are summarized in Table 2. Two regularities emerge from these figures. First, acquisition is roughly one third of lifetime cost: NSF’s \$10M ACSS-class systems consistently obligate \$30–38M over their lives once operations and expansion supplements are counted. A realistic facility budget therefore carries approximately twice its acquisition cost in operations and staffing across five years, a ratio this report adopts as a design principle in Section 7. Second, the storage fraction of leadership systems runs 8–10% (Frontier’s Orion contract exceeded \$50M of a roughly \$600M system), a useful sanity check on subsystem estimates.

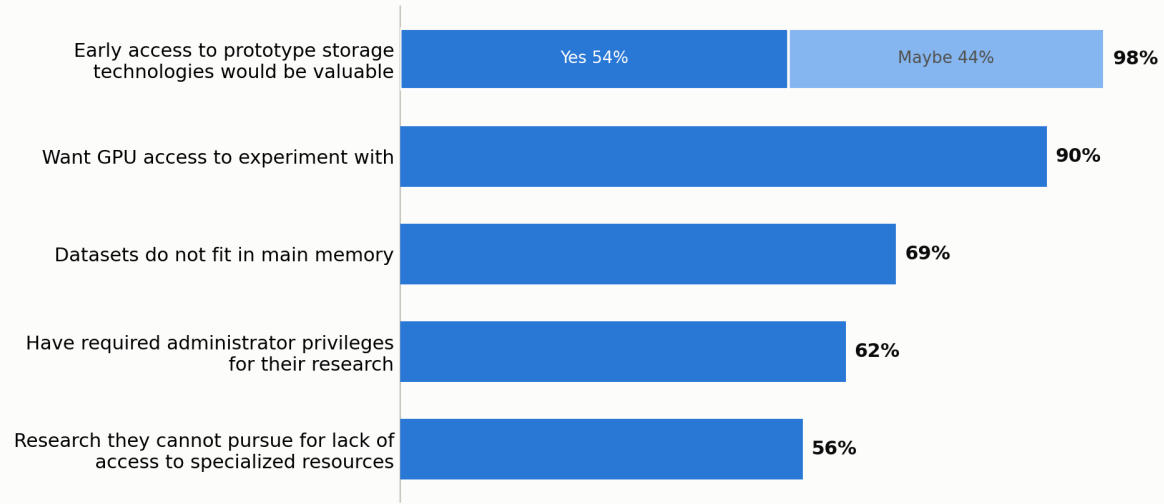
A hardware-cost caveat applies to this window. The NAND shortage described in Section 3 roughly doubles to triples the capital cost of a flash-heavy testbed proposed in 2026 versus 2025 planning assumptions [8]. Proposals written now should hedge NAND pricing explicitly, weight capacity toward hard-drive and tape tiers, and pursue donated or loaner media programs for scarce device classes.

5 The Infrastructure Gap

Overview of findings. *Three independent lines of evidence—the supply of facilities, the community’s surveyed demand, and the location of the field’s research phenomena—converge on the same gap: no dedicated, modifiable, modern-media storage testbed exists at meaningful scale. Access to modern hardware (CXL, programmable devices) and realistic-scale experimentation are the strongest-evidenced needs; shared traces are a real gap whose remedy is unproven; reproducibility is documented but under-measured. One commonly claimed gap does not survive scrutiny:*

The community’s demand signals are strong and specific

Share of respondents, StoreHub Community Insight Survey and workshop polls



Source: StoreHub Community Insight Survey and workshop polls (N=76), 2024–2026 · NSF CIRC awards 2346504/2346505

Figure 2: Infrastructure needs reported by the community in the StoreHub Community Insight Survey and workshop polls (N=76, 2024–2025) [18].

community benchmarks exist and are healthy—what is missing is infrastructure to run them on. Agent-native storage is an emerging gap in which NSF already funds the software pattern (IOWarp) but not the hardware substrate.

The question before NSF is not whether storage research matters; Sections 2 and 3 establish that AI made storage a quantified, first-order component of the most valuable computing systems in the world. The question is whether the community lacks shared infrastructure that a national effort would supply. Candidate gaps are examined individually below; some are strongly evidenced, some partially, and one commonly asserted gap is not supported by the evidence.

5.1 Testbeds at Realistic Scale

Supply. No U.S. testbed is dedicated to storage research. Chameleon and CloudLab offer roughly thirty purpose-built storage nodes between them, single-terabyte-class devices, and no capacity to stand up an at-scale parallel file system; FABRIC treats storage as a slice component; ACCESS and DOE production systems are measurable but not modifiable; OSN is S3 hosting only; and PRObE was retired around 2015 without a successor (Section 4).

Demand. In the StoreHub survey (Figure 2), 62% of respondents require root or administrator privileges to conduct their research, precisely what production HPC, ACCESS, NAIRR, and commercial clouds categorically deny; 56% report research they cannot conduct for lack of access to specialized computing or storage resources; 58% name customizable resource allocation as an essential user service; and respondents rated current infrastructure adequacy 3.4 on a five-point scale against importance ratings of 4.1–4.3 for storage performance and scalability [18]. The root-access requirement replicated across both survey cohorts independently. Dataset scale compounds the access problem: 69% of respondents report primary datasets that exceed main memory, with the distribution peaking in the 10–100 TB range (Figure 3).

Phenomenology. The behavior the field most needs to study occurs at scales academics cannot reach:

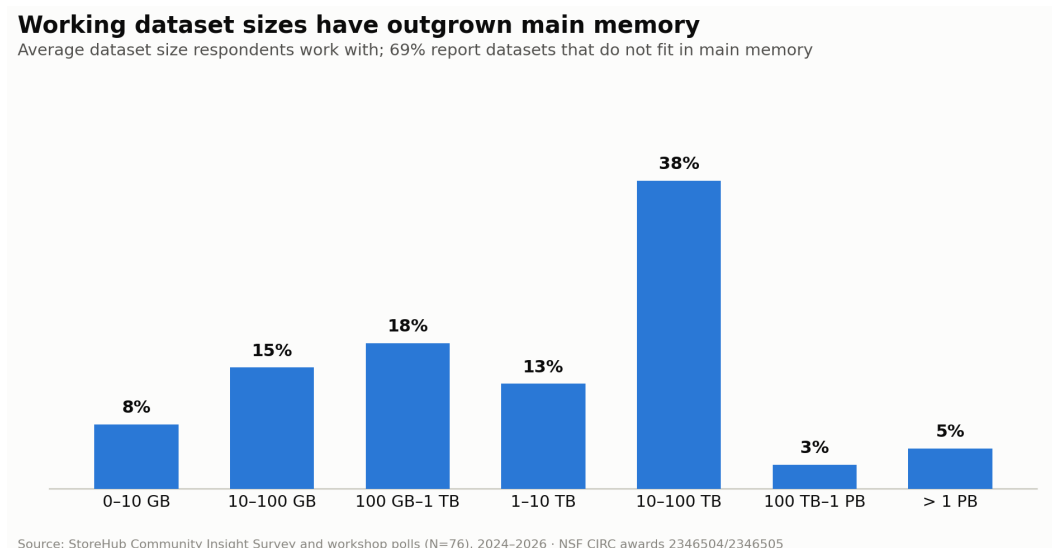


Figure 3: Distribution of primary research dataset sizes reported by survey respondents [18].

failure-driven checkpointing on 16,384-GPU clusters with a failure every three hours [56], burst congestion that hyperscaler engineering blogs describe and no university can reproduce [57], and KV-cache tiering on fleets serving more than 100 billion tokens per day [3]. The compute-divide literature documents the general problem [133]; its storage corollary has no advocate in the current policy debate.

5.2 Access to Modern Hardware

Neither CloudLab nor Chameleon lists dedicated CXL hardware as of mid-2026, and most academic CXL and ZNS work runs on emulation for lack of open hardware (Section 3). CXL modules are largely hyperscaler-allocated, and the price shock deepened the wall: 40–60% purchasing-power loss and 36–52-week lead times for memory-rich academic systems [9]. Demand matches supply’s failure: 68% of surveyed researchers identify CXL SSDs and 68% computational or software-defined storage as essential to their research, 46% open-channel SSDs, and 98% state that early access to prototype storage hardware would be valuable (Figure 4) [18].

The mismatch is structural: the device classes at the center of the field’s open research questions—and of SNIA Storage.AI’s published problem statements [55]—are the ones academic researchers currently cannot buy, rent, or borrow.

5.3 Shared Traces and Datasets

The gap itself is solid. SNIA’s IOTTA trace repository, which has served more than 860 papers, is aging and sparsely fed; no I/O traces from frontier training runs have ever been published; and Meta, the sole hyperscaler publishing storage architecture systematically, describes phenomena that exist nowhere in the open literature as data [57]. The field’s most productive feedback loops run exactly where open data exists, notably MLPerf Storage and the Backblaze drive-reliability dataset. What the evidence does not establish is the remedy’s supply side: nothing demonstrates that hyperscalers or vendors would donate traces to a national repository if one existed. Donation incentives are a design problem a center must solve, not a solved input; the service-for-access exchange proposed in Section 7 is the candidate mechanism.

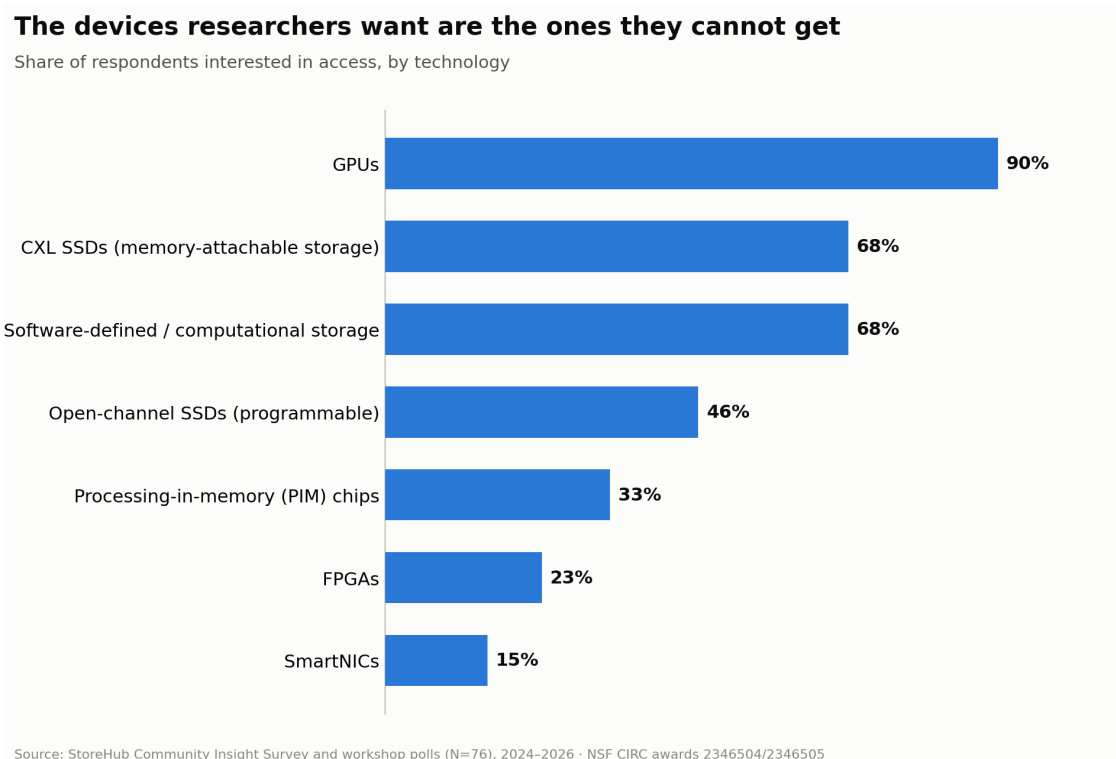


Figure 4: Technology interest among surveyed storage researchers [18].

5.4 Reproducibility Infrastructure

The mechanism is documented: multi-tenant noise on shared systems undermines reproducible baselines (Section 3), and the deployment complexity of modern storage stacks on shared academic clusters is itself a documented barrier [146]. The demand signal exists but is soft: 39% of surveyed researchers report difficulty participating in reproducibility initiatives owing to hardware constraints and background noise [18]. What is missing is systematic measurement of the reproducibility failure in storage specifically, for example artifact-evaluation failure rates at the major venues. The gap is real but under-measured, and quantifying it would be an inexpensive, high-value study for NSF to commission, alongside the storage-workforce tracking gap noted in Section 2.

5.5 Community Benchmarks

In the interest of balance, one commonly asserted gap is not supported by the evidence. Benchmarks exist and are healthy: MLPerf Storage v2.0 drew 26 organizations and more than 200 results and codified check-pointing [5, 59]; the IO500 continues; DLIO covers scientific deep-learning I/O [147]. The actual gap is that academics largely cannot run these benchmarks at representative scale or contribute results from systems they control, which is a testbed-and-access problem, not a benchmark-definition problem. A proposal that promises new community benchmarks as a headline deliverable would solve the wrong problem; one that promises a place to run and extend the existing ones would not.

5.6 Agentic-AI Integration

MCP became the de facto agent-to-data interface within a year of release, and storage products ship MCP servers, yet no landmark peer-reviewed storage-venue paper defines agent-native storage semantics as of mid-2026 (Section 2). The projection that agent-mediated data access becomes a mainstream storage-research area within two to four years implies that the defining abstractions are being set now. The only funded infrastructure occupying this space is NSF’s own IOWarp, whose verified mechanisms, live-telemetry MCP servers, trace access, and provenance capture for agentic workflows, constitute a working pattern for the automatic acquisition of real-system data by agents [38, 39]. The gap, precisely stated, is that the software pattern exists and is federally funded while the hardware substrate and community scale for it do not.

6 Assessment

Overview of findings. *The case against a new center rests on the funded existing testbeds, the least favorable hardware-price environment on record, a difficult budget climate, the possibility that industry supplies the needed infrastructure on its own, and a permanent ceiling on the scale any academic facility can reach. The case in favor rests on a gap documented independently through facility supply, community demand, and the location of the field’s research phenomena; on a narrow period (2025–2028) during which academic work can still influence the abstractions of AI-era storage; on the observation that the costs of inaction grow with time; on the inability of every alternative to support research below the block and object interfaces; and on the existence of proven precedents at every scale point. The conclusion: a national storage research facility is warranted at the medium scale—approximately \$20–25M over five years, Mid-scale RI-1 class, federated rather than freestanding—with a defined growth path as milestones are met. A monolithic “small NERSC” new start is not supported in the current climate.*

6.1 The Case Against a National Center

Stated at full strength, five arguments counsel against new infrastructure. First, the existing base is funded and partially underused: Chameleon and CloudLab hold fresh \$12M Phase-4 awards through approximately 2028, Chameleon added composable hardware in 2024, and CloudBank 2.0 and NAIRR already serve scale experiments and AI workloads; a marginal dollar spent extending these arguably buys more than a first dollar spent on a new organization. Second, the timing for capital acquisition is the worst in the industry’s recorded history: flash roughly tripled in price, DRAM rose 90% in a quarter, and the shortage is projected to persist into 2027–2028, so a flash-heavy testbed bought in 2026–2027 costs two to three times its 2025 price. Third, the budget environment is unfavorable to new starts: the agency emerged from FY2026 reorganized, with CIRC archived, Mid-scale RI-2 paused, and deep FY2027 reductions proposed, and no national policy instrument treats storage as a funded priority, so storage-specific infrastructure cannot rely on a dedicated policy mandate for support. Fourth, industry may deliver the substrate anyway: NVIDIA’s certification regime and SNIA Storage.AI are actively standardizing the AI-storage interface, MLPerf Storage codifies the workloads, and vendors have incentives to seed universities with hardware once the shortage clears. Fifth, the scale objection cuts against the proposal itself: if the interesting failure modes live at 10,000–100,000-GPU scale, a \$20M center cannot reproduce them either, and the center risks becoming a 2015-era testbed for 2030 problems while partnership instrumentation on Horizon- and Doudna-class systems might capture the phenomena more cheaply.

6.2 The Case For

Five arguments support proceeding. First, the gap is real, specific, and evidenced from three independent directions, supply, demand, and phenomenology (Section 5), and converges on the niche Section 4 identifies: dedicated, destructible, reconfigurable storage systems at hundreds-of-nodes and tens-of-petabytes scale with modern media and fault injection. Every existing facility is a potential partner precisely because none of them occupies that niche. Second, the window argument: the interface layer of AI-era storage is being decided in 2025–2028, in the contest between NVIDIA’s stack and SNIA Storage.AI and in the pre-paradigm phase of agent-native storage semantics, and academic influence on those abstractions requires that academic researchers be working on real systems while the abstractions are still open to change. Infrastructure that becomes available after the standards have settled can only document the outcome; infrastructure available while they are being formed can influence them. Third, the counterfactual is deteriorating rather than static: price pressure is pushing academic hardware configurations further from production reality, the talent flow is one-directional (61% of new PhDs to industry, enrollment down 15%), and waiting compounds the divide. Fourth, the economics of the alternatives fail exactly where the gap is: cloud cannot host device-level, kernel-level, fault-injection, or reproducible-baseline research at any price, sustained data-heavy workloads are what cloud metering punishes, and the 62% root-access requirement is categorically unmeetable on production or multi-tenant systems. Fifth, the model is proven at every scale point: PRObE demonstrated the dedicated destructible testbed, Chameleon and CloudLab demonstrate \$12M-per-phase community operation with more than a thousand publications, the AI Institutes demonstrate NSF’s willingness to sustain \$20M five-year thematic centers, OSN shows a national storage federation bootstrapping on \$2M, and IOWarp demonstrates the agentic software layer. Nothing here requires an unprecedented organizational invention, only an unprecedented combination.

6.3 Verdict

The evidence supports a national storage-research facility at the medium scale, built as a federation; it does not support a monolithic smaller-NERSC new start in the current climate. Three considerations settle the matter between the two cases. First, the strongest points against, price shock, budget volatility, and the scale ceiling, are arguments about sizing, timing, and design rather than about whether the gap exists, and none of them is addressed by inaction, the costs of which grow with time. Second, the scale objection is answered by design rather than dollars: the facility’s response to frontier-scale phenomenology is instrumentation partnerships with the systems that have it, Horizon’s 400 PB subsystem, Doudna’s quality-of-service storage, and OLCF’s ACE testbed, all of which for the first time operate storage novel enough to need research partners, while the facility owns only what must be owned, the destructible root-access middle scale. Third, on sizing, the two independent research passes behind this report converged on the recommendation while diverging on scale, one supporting \$30–80M over five years outright, the other, grounded in the documented costs of comparable facilities (Table 2), supporting \$20–25M as the scale that closes the PRObE gap. Given the current pause in Mid-scale RI-2 and the proposed FY2027 reductions, a staged path carries materially less programmatic risk than a single larger commitment.

The recommendation is therefore a national AI-and-storage research facility at approximately \$20–25M over five years, sized at the Mid-scale RI-1 class and complemented by software-institute support and by NAIRR and ACCESS allocation integration, structured as a multi-institution federation, with a defined growth path toward larger-scale national infrastructure as operational milestones are met.

7 Design of the Proposed Facility

Overview of findings. *Four functions, in priority order: a 200–400-node destructible testbed with root access and modern media (CXL, FDP, GPUDirect paths, fault injection); a national I/O observatory that trades instrumentation services to production facilities for traces and access; an agentic access layer through which every experiment self-documents into community datasets; and community and workforce programs at roughly 20% of budget. The facility builds only these; identity, allocation, data movement, and scale access are all integrated from existing systems. An illustrative reference configuration prices acquisition at \$7.0M (2025 prices) to \$9.6M (2026 prices), within a five-year total of \$22–27M; the facility’s distinguishing capability (the emerging-device laboratory) is among its least expensive subsystems.*

7.1 Functions, in Priority Order

F1: The dedicated testbed. The irreplaceable core is a facility of roughly 200–400 storage-dense nodes and 10–30 PB across tiers offering what no existing facility does: root and kernel access; raw and programmable devices, including flexible-data-placement SSDs, zoned namespaces, open-channel devices where obtainable, and computational-storage samples; a CXL pool, expansion first and pooling as silicon matures; GPUDirect and GPU-initiated I/O paths with sufficient accelerators to exercise them; reconfigurable fabrics (NVMe-over-Fabrics, RDMA, programmable switches); power-fault and cluster-scale fault injection, the PRObE function; and reservation-based isolation for reproducible baselines. Hardware strategy must reflect the price environment: capacity weighted toward hard-drive and tape tiers, modest owned flash, vendor loaner and donation programs for scarce media, and explicit NAND-price contingencies.

F2: The national I/O observatory. The most cost-effective of the four functions is instrumentation, trace capture, and curation operated as a service to the new national systems, Horizon, Doudna, Discovery, and the ACCESS resource providers, in exchange for traces and scale access. This exchange is the mechanism that converts the unproven trace-donation remedy of Section 5 into a transaction both sides want: facilities obtain instrumentation and benchmarking expertise they do not currently staff, and the community obtains the traces it has never had. The observatory would also revive the IOTTA function under modern governance and anchor continuous MLPerf-Storage-based benchmarking.

F3: The agentic access layer. Every facility capability, allocation, experiment launch, telemetry, traces, and datasets, is exposed as MCP endpoints from the first day of operation, generalizing the pattern IOWarp has already demonstrated under NSF funding rather than inventing a new one [38, 39]. Three consequences distinguish the design from earlier testbeds. Experiments self-document: every run’s traces, configurations, and outcomes are born machine-readable and repository-ready, addressing the trace gap structurally. The access barrier drops: a graduate student at an institution without systems staff can run a storage experiment against real hardware through a conversational agent, which is the substantive answer to NAIRR’s democratization mandate applied to systems research. And the facility becomes the natural venue for developing and evaluating the agent-native storage semantics—consistency, provenance, and access-control models for agent-mediated data access—that the peer-reviewed literature does not yet provide, on real systems rather than in simulation.

F4: Community and workforce programs. The CIRC template of 20–25% of budget devoted to community programs is adopted deliberately: workshops (55% of surveyed researchers prefer workshops and conferences for engagement [18]), summer schools, artifact-evaluation hosting for the major venues, and the storage-workforce tracking that no existing instrument provides.

7.2 Integrate Rather than Build

The facility builds three things: the storage-dense destructible hardware of F1, the observatory tooling and governance of F2, and the agentic layer of F3 atop IOWarp. Everything else is integration with the inventory of Section 4 and with proven federation precedents: identity through InCommon and CILogon [148]; allocation through the existing ACCESS processes and through NAIRR resource-provider status when that pathway opens; wide-area data movement through FABRIC and the Open Science Data Federation’s cache-origin model [139, 149]; edge service deployment following the SLATE pattern [150]; bare-metal lifecycle mechanics in partnership with Chameleon and CloudLab rather than by reinvention; scale access through OLCF ACE-style negotiated agreements and the service-for-access exchanges with Horizon and Doudna [143]; and cloud burst through CloudBank for the experiments cloud serves well [93]. The facility is deliberately conventional in its organization and governance, reusing proven federation models throughout; its novelty is concentrated where the gap analysis indicates it must be, in the hardware it provides and in the interfaces through which that hardware is used.

7.3 Scale Points

Three scale points bracket the design, anchored to Table 2. A small facility, \$5–12M over four to five years, the CIRC-Grand and CloudLab-phase class, funds a coordination hub with F2 and F3 in full and F1 as a modest 50–100-node testbed, wholly dependent on partnerships for scale; it is viable as a fallback and does not close the PRObE gap. The recommended medium facility, \$20–25M over five years, the Mid-scale RI-1 and AI-Institute class, funds a purpose-built few-hundred-node, 10–30 PB facility with all four functions, staff, and national allocation, budgeted with operations at approximately twice acquisition, roughly \$7–8M of hardware and \$14–16M of operations, staffing, and community programs, with explicit NAND hedges. A large facility, \$30–100M, the Mid-scale RI-2 class, would add thousand-node fault injection and co-located archival research; it represents a natural growth path once a medium-scale facility has established a multi-year operating record.

7.4 Reference Configuration and Cost Estimate

The costs of comparable facilities (Section 4) bound the plausible budget; this section provides a concrete estimate within those bounds. Table 3 states the unit prices on which the estimate rests, at the two price points the 2025–2026 market defines; Table 4 applies them to an illustrative reference configuration for the recommended medium-scale facility. The configuration is illustrative rather than final—a proposal-stage design would be refined with vendors and partners—but the arithmetic is committed: totals follow from the cited unit prices, and the estimate should be read with an uncertainty of roughly $\pm 30\%$.

Three observations follow from the tables. First, the same configuration costs roughly \$2.6M more at 2026 prices than at 2025 prices, with the flash tier alone accounting for \$1.5M of the difference—a quantified statement of the procurement-timing risk, and the basis for the phased-acquisition strategy: procure nodes, fabric, HDD, and tape early; back-load the bulk of the flash tier into the supply easing projected for 2027–2028 [74]; and pursue vendor loaner and donation programs for the scarcest device classes (CXL modules, FDP and computational-storage samples), where research value is highest per dollar and market availability is lowest. Second, the emerging-device laboratory—the capability no other facility offers—is among the least expensive subsystems, at roughly 5% of acquisition; the facility’s differentiation does not depend on bulk capacity spending. Third, the five-year budget assembles consistently with the anchors of Section 4: acquisition of \$7.5–9.5M depending on procurement timing; operations and staffing at eight to ten loaded FTEs, \$2.0–2.5M per year, or \$10–12.5M over five years; and community programs at the CIRC-template share, \$4–5M; for a total of \$22–27M, consistent with the recommended \$20–25M class once donation offsets and partner cost-sharing are counted.

Component	Mid-2025	Mid-2026	Basis
Enterprise NVMe flash, per TB	\$80–150	~\$500	Contract/street pricing [8]
Nearline HDD, per TB	>\$6	>\$15	Market-research vendor ASP [151]
Server DRAM (DDR5 RDIMM), per TB	~\$4,100	~\$14,400	Module contract data [70, 71]
LTO tape, per TB (media, compressed)	~\$2	~\$2.10	Analysis and current vendor price [152, 153]
Storage server, dual-socket base	\$12–18K	\$12–18K	Vendor list estimate
400 Gb/s fabric, per port (switch + NIC)	\$1.5–3K	\$1.5–3K	Vendor list estimate
GPU node, 8 accelerators (H100 class)	\$250–350K	\$250–350K	Market estimate
Operations staffing, per loaded FTE-year	\$190–250K		Derived from TCO models [92]

Table 3: Unit-price basis for the facility cost estimate. The 2025/2026 divergence in memory and flash reflects the supply shock documented in Section 3.

Subsystem (illustrative configuration)	At 2026 prices	At 2025 prices
200 storage-dense nodes (dual-socket, 384 GB DRAM each)	\$4.1M	\$3.3M
Experimental flash tier, 4 PB NVMe	\$2.0M	\$0.5M
Capacity tier, 20 PB nearline HDD	~\$0.30M	\$0.12M
Archival tier, 30 PB tape (library + media)	\$0.35M	\$0.35M
GPU partition, 4 nodes $\times$ 8 accelerators	\$1.2M	\$1.2M
CXL pool and emerging-device laboratory (FDP, ZNS, computational-storage, programmable devices)	\$0.5M	\$0.4M
Fabric: 400 Gb/s RDMA/NVMe-oF plus programmable (P4) switching	\$0.8M	\$0.8M
Fault-injection and power instrumentation, spares	\$0.4M	\$0.35M
Acquisition total	~\$9.6M	~\$7.0M

Table 4: Illustrative reference configuration for the medium-scale facility (~ 200 nodes, ~ 54 PB across four tiers, 32 accelerators). Totals computed from Table 3; $\pm 30\%$.

7.5 Distinction from Prior-Generation Testbeds

Testbeds of the PRObE generation provided researchers with remote machine access and documentation, and left everything else—experiment configuration, measurement, data collection, and publication of results—to each individual research group. The facility proposed here differs in that data collection is built into the infrastructure itself. Both researchers and software agents interact with the facility through MCP interfaces that expose live system telemetry; every experiment automatically records its configuration, traces, and outcomes into the observatory’s repository; standard benchmarks run continuously in the background rather than only when individual groups choose to run them; and the facility’s own operational record—component failures, device wear, thermal behavior, congestion under checkpoint-style burst loads—is itself curated and published as a community dataset from the beginning of operations.

This design also addresses the specific limitation that characterizes each class of existing facility. Production systems such as the DOE machines and Horizon run the large-scale scientific and AI workloads in which the phenomena of research interest actually occur, but researchers cannot modify them: they grant

no administrative privileges, permit no kernel or firmware changes, and cannot host destructive experiments. Existing testbeds such as Chameleon and CloudLab grant exactly that kind of low-level control, but they do not carry production-scale workloads, so the phenomena of interest do not occur on them, and they provide no mechanism for systematically collecting and publishing the data an experiment generates. The proposed facility is designed to supply all three properties at once: it is modifiable in the way a testbed is, it observes production-scale storage behavior through its instrumentation partnerships with the systems that have it, and it collects and publishes the resulting data as a matter of infrastructure rather than of individual effort.

8 Risks, Alternatives, and Conditions for Revision

Overview of findings. *Every alternative to a dedicated facility—extending general-purpose testbeds, cloud credits, distributed vendor donations, reliance on DOE initiatives—fails specifically where the gap is, though each is valuable as a component of the design. The principal risks of proceeding are procurement timing, program volatility, utilization, agentic-layer immaturity, and the scale ceiling; each has a stated mitigation. The conclusion is falsifiable: the report lists six concrete developments that would weaken or reverse it, and three that would strengthen it.*

8.1 Alternatives Considered

Extend Chameleon and CloudLab. Storage-dense racks and CXL hardware could be added through supplements plausibly costing \$2–5M. This is the fastest path to the modern-hardware gap and should be pursued regardless of the center decision. It cannot, however, reach tens-of-petabytes scale or host destructive fault injection on shared general-purpose testbeds, provides neither observatory nor agentic layer, and leaves storage a minority tenant of facilities with other missions, the arrangement whose thirty-node ceiling defines the status quo.

Cloud credits at scale. Expanding CloudBank-style allocations is the right answer for scale-out work above the block and object interfaces and the wrong answer below them: no firmware, no kernel control, no fault injection, no reproducible baselines, no CXL access (Section 3). The sustained, data-heavy profile of storage research is cloud metering’s worst case, and egress economics compound it.

The distributed-instrument model. Vendor-donated hardware pods at many campuses, on the OSN pattern, are politically attractive and worth using for media diversity. The record argues against them as the primary model: donations concentrate at elite institutions, deployed hardware is isolated from the broader community, results do not reproduce across heterogeneous sites, and no root-access norm exists. OSN itself became a hosting service rather than a research platform precisely because nothing funded the research function [142].

Emerging federal collaboration pathways. A final alternative holds that Department of Energy initiatives will meet the need. The evidence supports treating them as collaboration opportunities rather than substitutes: the Genesis Mission inventories federal storage resources but funds AI-for-science applications; DOE’s storage momentum is on the procurement side rather than in research access (Section 4); and ACE-style testbed access, while valuable, is allocation-gated, non-destructive, and shaped by laboratory missions. These same developments make DOE facilities natural partners for the instrumentation and access exchanges described in Section 7. Each alternative examined here is, in fact, a component of that design; none, alone or in combination, occupies the niche.

8.2 Principal Risks of Proceeding

Five risks are named with their mitigations. Procurement timing: buying flash into a shortage projected to persist to 2027–2028; mitigated by tiering, loaner programs, and back-loading flash purchases into the projected easing, and bounded by the fact that device diversity rather than bulk flash carries most of the research value per dollar. Program volatility: individual funding programs can pause, as CIRC and Mid-scale RI-2 did; the mitigation is a deliberately multi-program structure, combining software-institute support, mid-scale infrastructure funding, NAIRR and ACCESS allocation integration, and DOE partnerships, treated as part of the facility’s design rather than as a contingency. Utilization: a facility the community does not use; the $N=76$ survey is encouraging but small, so first-year milestones should include independent demand validation at larger N . Agentic-layer immaturity: MCP is twenty months old and agent-native semantics unpublished; mitigated by building on funded, running IOWarp code, and bounded by a graceful failure mode, since the layer degrades to conventional access. The scale ceiling: a few hundred nodes never equals 100,000 GPUs; this is permanently true, is answered by the observatory function, and should be stated plainly in any proposal rather than overclaimed.

8.3 Conditions for Revision

The verdict is falsifiable, and program officers should hold it to that standard. It weakens or reverses if any of the following occurs: CIRC returns with materially larger caps, or Chameleon and CloudLab Phase-5 plans, visible by roughly 2027, pivot decisively to storage-dense, CXL-equipped, fault-injection-capable configurations, in which case the extension alternative deserves first claim; commercial clouds begin exposing device-level control, programmable flash translation layers, zoned or flexible-placement access, fault injection, and single-tenant isolation, at academic price points; a larger- N independent survey fails to replicate the StoreHub survey findings; the memory and flash shortage extends materially past 2028, making any owned-hardware strategy untenable and reducing the design to its observatory and agentic subset; the service-for-access exchanges fail in practice, that is, if Horizon-, Doudna-, and ACCESS-class operators decline instrumentation partnerships within the facility’s first two years; or vendor and consortium programs begin offering root-access, modern-media systems to universities at scale. Conversely, the verdict strengthens if the CREATE AI Act passes and NAIRR opens its anticipated resource-provider pathway, if artifact-evaluation data quantifies the reproducibility gap, or if subsequent workforce data confirms continued pipeline erosion.

9 Conclusion

Three years of evidence tell a consistent story. The storage research field reorganized itself around AI faster than in any prior transition in its history; the storage demands of frontier AI systems are now precisely quantified and severe; and the systems that exhibit those demands are inaccessible to the academic researchers expected to study them. The economics of 2025–2026 widened that gap: record price surges, hyperscaler-first allocation, and configuration divergence between what universities can buy and what production runs. The national infrastructure inventory is strong everywhere except the one place the field needs it, and the community has said, in its own survey, precisely what it lacks: privileged access, modern devices, isolation, and scale.

The remedy this report recommends is deliberately conservative in form and specific in function: a federated, medium-scale national facility, approximately \$20–25M over five years through Mid-scale RI-1, that owns only what nothing else provides—a destructible, root-access testbed built on modern media—and obtains everything else through partnerships in which the facility supplies instrumentation and benchmarking services to existing systems in exchange for data and access. Its distinguishing investment, an agentic

access layer through which the facility continuously acquires and publishes real-system data, means that the facility's value to the community grows over time as its published datasets accumulate, and aligns the effort with national priorities in artificial-intelligence infrastructure. Absent action, the interface standards, the agent-native semantics, and the trace norms of AI-era storage will be written entirely inside companies, and national infrastructure built later can only replicate, not shape, the field it was meant to serve. The window in which the alternative remains available is measured in years, not decades, and it is open now.

A Consolidated Findings Register

Findings are tagged [FACT] (observed and cited) or [PROJECTION] (extrapolated from observed trajectories).

1. [FACT] AI-related papers grew from roughly 14% of the FAST '24 program to 25% or more at FAST '25 and FAST '26, with AI themes taking best-paper awards in both 2025 and 2026; persistent-memory research collapsed after Optane's cancellation, and MSST ran no research track in 2025 [1–3, 26–28].
2. [FACT] Inference state became a storage workload: Mooncake and LMCache established KV-cache tiering across GPU memory, DRAM, flash, and object storage, productized within twelve months by NVIDIA, WEKA, VAST, DDN, and Google Cloud [3, 4, 32].
3. [FACT] Checkpointing is a quantified first-class workload: Meta's Llama-3 run suffered unexpected failures every three hours on 16,384 GPUs, and MLPerf Storage v2.0 implies approximately 3.6 TB/s of burst write bandwidth at 100,000 accelerators [5, 56].
4. [FACT] The 2025–2026 price surge is the sharpest on record: server DRAM rose approximately 90% quarter-over-quarter in 1Q26, nearline hard drives sold out through calendar 2026, enterprise flash reached the \$500/TB class, and academic purchasing power for memory-rich systems fell 40–60% between proposal and award [6–9].
5. [FACT] No U.S. facility is dedicated to storage research: Chameleon and CloudLab offer roughly thirty purpose-built storage nodes combined and no CXL hardware; production systems are measurable but not modifiable; PROBE has had no successor since approximately 2015 [10–13].
6. [FACT] Community demand is documented: in the StoreHub survey ($N=76$), 62% require root access, 56% report research they cannot conduct, 68% identify CXL and computational storage as essential, 98% value prototype-hardware access, and infrastructure adequacy rates 3.4 of 5 [18].
7. [FACT] Across the 2024–2026 policy corpus, NAIRR, the AI Action Plan, the Genesis Mission, and the CREATE AI Act, storage appears as supporting language but never as a funded priority in its own right; the NSF program created for this purpose (CIRC) is archived, while Mid-scale RI-1, CSSI, ACSS, and IDSS remain active programs [98, 100, 127, 128].
8. [FACT] MCP became the de facto agent-to-data interface within one year of release, no peer-reviewed storage-venue paper yet defines agent-native storage semantics, and NSF's IOWarp is the earliest funded infrastructure in the space, with verified mechanisms for agents to acquire live telemetry, traces, and provenance from real systems [36–38].
9. [PROJECTION] KV-cache and model-state storage consolidates into a standard tier by approximately 2028–2030 through either NVIDIA's stack or SNIA Storage.AI; the interface contest, observable in 2025–2026, defines the period in which academic and NSF influence on the abstractions remains feasible [54, 55].
10. [PROJECTION] Agent-native storage semantics become a mainstream research area within two to four years; the defining abstractions are being set now, and infrastructure influence requires arrival before approximately 2028.
11. [PROJECTION] A credible facility comes in three sizes, \$5–12M, \$20–25M, and \$30–100M; the medium scale closes the PROBE gap, and an illustrative reference configuration prices its acquisition at \$7.0–9.6M depending on procurement timing (Section 7.4). Given demonstrated program volatility, a multi-program funding structure is more robust than dependence on any single solicitation [100, 101].
12. [PROJECTION] On the do-nothing path the divide compounds: academic hardware diverges further from production configurations under price pressure, the workforce pipeline (61% of new PhDs to industry,

enrollment down 15%) thins as storage becomes a named national bottleneck, and the standards of AI-era storage are written entirely inside companies [9, 60].

References

- [1] USENIX Association. FAST '25 technical sessions. <https://www.usenix.org/conference/fast25/technical-sessions>, 2025. Accessed July 2026.
- [2] USENIX Association. FAST '26 technical sessions. <https://www.usenix.org/conference/fast26/technical-sessions>, 2026. Accessed July 2026.
- [3] Ruoyu Qin, Zheming Li, Weiran He, Jialei Cui, Feng Ren, Mingxing Zhang, Yongwei Wu, Weimin Zheng, and Xinran Xu. Mooncake: Trading more storage for less computation — a KVCache-centric architecture for serving LLM chatbot. In *Proceedings of the 23rd USENIX Conference on File and Storage Technologies (FAST '25)*, 2025. Best Paper Award. Preprint: <https://arxiv.org/abs/2407.00079>.
- [4] NVIDIA. Introducing NVIDIA Dynamo, a low-latency distributed inference framework. <https://developer.nvidia.com/blog/introducing-nvidia-dynamo-a-low-latency-distributed-inference-framework-for-scaling-reasoning-ai-models/>, March 2025.
- [5] MLCommons. MLPerf storage v2.0: Benchmarking checkpointing for large-scale training. <https://mlcommons.org/2025/08/storage-2-checkpointing/>, August 2025.
- [6] TrendForce. 1q26 conventional DRAM prices rise 90–95% qoq. <https://www.trendforce.com/press-center/news/20260202-12911.html>, February 2026.
- [7] The Register. AI blamed again as hard drives sell out. https://www.theregister.com/2026/02/20/ai_blamed_again_as_hard_drives_sell_out/, February 2026.
- [8] Tom Coughlin. SSD storage capacity prices are over 20 times HDD storage capacity prices. Forbes, <https://www.forbes.com/sites/tomcoughlin/2026/04/16/ssd-storage-capacity-prices-are-over-20-times-hdd-storage-capacity-prices/>, April 2026.
- [9] HPCwire. How the memory shortage is impacting AI and HPC projects. <https://www.hpcwire.com/2026/01/15/how-the-memory-shortage-is-impacting-ai-and-hpc-projects/>, January 2026.
- [10] CloudLab. CloudLab hardware. <https://www.cloudlab.us/hardware.php>, 2026. Accessed July 2026.
- [11] Parallel Data Laboratory, Carnegie Mellon University. PRObE: Parallel reconfigurable observational environment. <https://www.pdl.cmu.edu/PRObE/index.shtml>, 2015.
- [12] Garth Gibson, Gary Grider, Andree Jacobson, and Wyatt Lloyd. PRObE: A thousand-node experimental cluster for computer systems research. *USENIX ;login.*, 38(3), 2013. https://www.pdl.cmu.edu/PDL-FTP/HECStorage/07_gibson_036-039_abs.shtml.
- [13] Chameleon Cloud. Chameleon hardware description. <https://www.chameleoncloud.org/about/hardware-description/>, 2026. Accessed July 2026.
- [14] Chameleon Cloud. Maximizing storage research on chameleon. <https://chameleoncloud.org/blog/2022/03/29/maximizing-storage-research-on-chameleon/>, March 2022.
- [15] HPCwire. NSF announces \$457 million leadership-class computing facility at UT austin. <https://www.hpcwire.com/2024/07/29/nsf-announces-457-million-leadership-class-computing-facility-at-ut-austin/>, July 2024.

- [16] Blocks & Files. VAST Data, Dell, Versity, and Spectra Logic are shining storage stars on TACC’s horizon. <https://blocksandfiles.com/2025/11/18/vast-data-dell-versity-and-spectra-logic-are-shining-storage-stars-on-taccs-horizon/>, November 2025.
- [17] Texas Advanced Computing Center. Proposals for allocations on NSF LCCF horizon to begin april 15. <https://lccf.tacc.utexas.edu/construction-updates/2026/proposals-for-allocations-on-nsf-lccf-horizon-to-begin-april-15/>, 2026.
- [18] StoreHub Project Team. StoreHub community insight survey: Results of the 2024–2025 community instruments. Project data report, nsf awards 2346504/2346505, Illinois Institute of Technology and The Ohio State University, 2025. $N=76$ across two instruments. <https://grc.iit.edu/research/projects/storehub>.
- [19] National Science Foundation. Awards 2346504/2346505: StoreHub: A community infrastructure for shaping the future of data storage research. https://www.nsf.gov/awardsearch/showAward?AWD_ID=2346504; https://www.nsf.gov/awardsearch/showAward?AWD_ID=2346505, 2024. Verified via NSF Awards API, July 2026.
- [20] USENIX Association. FAST ’24 technical sessions. <https://www.usenix.org/conference/fast24/technical-sessions>, 2024. Accessed July 2026.
- [21] ACM HotStorage. Hotstorage ’24 program. <https://www.hotstorage.org/2024/program.html>, 2024.
- [22] ACM HotStorage. Hotstorage ’25 accepted papers. <https://www.hotstorage.org/2025/accepted.html>, 2025.
- [23] PDSW Workshop. The 10th international parallel data systems workshop (PDSW ’25), held with SC25. <https://pds.w.org/pds25/index.shtml>, 2025.
- [24] Meng Tang, Zhaobin Zhu, Luanzheng Guo, James G. Bandy, Tim Carlson, Sarah Neuwirth, Anthony Kougkas, Xian-He Sun, and Nathan R. Tallent. Quantifying AWS S3 I/O performance boundaries using the roofline model. In *Proceedings of the 10th International Parallel Data Systems Workshop (PDSW ’25), held with SC25*, 2025. DOI: 10.1145/3731599.3767513.
- [25] ACM SIGOPS. SOSP 2024 accepted papers. <https://sigops.org/s/conferences/sosp/2024/accepted.html>, 2024.
- [26] MSST. MSST 2024: 50th Anniversary Conference. <https://www.msstconference.org/2024/>, 2024.
- [27] MSST. MSST 2025 Conference Program. <https://www.msstconference.org/2025conference/>, 2025.
- [28] MSST Conference. MSST 2026 research program. <https://www.msstconference.org/2026-research-program/>, 2026.
- [29] SSDBM. International conference on scalable scientific data management (SSDBM 2025). <https://ssdbm.org/2025/>, 2025.
- [30] Borui Wan, Mingji Han, Yiyao Sheng, Yanghua Peng, Haibin Lin, Mofan Zhang, Zhichao Lai, Menghan Yu, Junda Zhang, Zuquan Song, Xin Liu, and Chuan Wu. ByteCheckpoint: A unified checkpointing system for large foundation model development. arXiv:2407.20143; subsequently published at NSDI ’25, 2024. <https://arxiv.org/abs/2407.20143>.

- [31] Tanmaey Gupta, Sanjeev Krishnan, Rituraj Kumar, Abhishek Vijeev, Bhargav Gulavani, Nipun Kwatra, Ramachandran Ramjee, and Muthian Sivathanu. Just-in-time checkpointing: Low cost error recovery from deep learning training failures. In *Proceedings of the 19th European Conference on Computer Systems (EuroSys '24)*, 2024. DOI: 10.1145/3627703.3650085.
- [32] Yuhan Liu, Yihua Cheng, Jiayi Yao, Yuwei An, Xiaokun Chen, Shaoting Feng, Yuyang Huang, Samuel Shen, Rui Zhang, Kuntai Du, and Junchen Jiang. LMCache: An efficient KV cache layer for enterprise-scale LLM inference. arXiv:2510.09665, October 2025. <https://arxiv.org/abs/2510.09665>.
- [33] Viraj Thakkar, Qi Lin, Kenanya Keandra Adriel Prasetyo, Raden Haryosatyo Wisjnnunandono, Achmad Imam Kistijantoro, Reza Fuad Rachmadi, and Zhichao Cao. ELMo-Tune-V2: LLM-assisted full-cycle auto-tuning to optimize LSM-based key-value stores. arXiv:2502.17606, 2025. <https://arxiv.org/abs/2502.17606>.
- [34] IBM Research. A 100-billion-vector storage index for AI. <https://research.ibm.com/blog/cas-100-billion-vector-storage-ai>, 2025.
- [35] Amazon Web Services. Amazon S3 vectors now generally available with increased scale and performance. <https://aws.amazon.com/blogs/aws/amazon-s3-vectors-now-generally-available-with-increased-scale-and-performance/>, December 2025.
- [36] Anthropic. Donating the Model Context Protocol and establishing the Agentic AI Foundation. <https://www.anthropic.com/news/donating-the-model-context-protocol-and-establishing-of-the-agentic-ai-foundation>, December 2025.
- [37] Linux Foundation. Linux foundation announces the formation of the Agentic AI Foundation. <https://www.linuxfoundation.org/press/linux-foundation-announces-the-formation-of-the-agentic-ai-foundation>, December 2025.
- [38] National Science Foundation. Award 2411318: Collaborative research: Frameworks: IOWarp: Bending the I/O fabric for advancing AI-infused scientific workflows. https://www.nsf.gov/awardsearch/showAward?AWD_ID=2411318, 2024. Verified via NSF Awards API, July 2026.
- [39] IOWarp Project. IOWarp MCP servers for scientific computing. <https://github.com/iowarp/iowarp-mcps>, 2025.
- [40] Gnosis Research Center. IOWarp project overview. <https://grc.iit.edu/research/projects/iowarp/>, 2025.
- [41] Blocks & Files. MinIO removes management features from basic community edition object storage code. <https://blocksandfiles.com/2025/06/19/minio-removes-management-features-from-basic-community-edition-object-storage-code/>, June 2025.
- [42] InfoQ. MinIO enters maintenance mode: S3-compatible alternatives. <https://www.infoq.com/news/2025/12/minio-s3-api-alternatives/>, December 2025.
- [43] MinIO, Inc. minio/minio repository (archived read-only april 25, 2026). <https://github.com/minio/minio>, 2026.
- [44] Ceph Foundation. v20.2.0 tentacle released. <https://ceph.io/en/news/blog/2025/v20-2-0-tentacle-released/>, November 2025.
- [45] OpenSFS. Lustre 2.17.0 released. <https://www.opensfs.org/lustre-2-17-0-released/>, December 2025.

- [46] Blocks & Files. DAOS’ post-Optane resurrection. <https://blocksandfiles.com/2025/04/15/daos-post-optane-resurrection/>, April 2025.
- [47] LanceDB. LanceDB series a funding announcement. <https://www.lancedb.com/blog/series-a-funding>, June 2025.
- [48] SiliconANGLE. Tigris data raises \$25m for AI-optimized cloud storage service. <https://siliconangle.com/2025/10/09/tigris-data-raises-25m-ai-optimized-cloud-storage-service/>, October 2025.
- [49] SlateDB Project. SlateDB: An embedded storage engine built on object storage. <https://slatedb.io/>, 2024.
- [50] Bun. Bun 1.2 documentation: S3 object storage. <https://bun.com/docs/runtime/s3>, 2025.
- [51] Blackstone. Blackstone invests \$300 million at a \$5 billion valuation in DDN. <https://www.blackstone.com/news/press/blackstone-invests-300-million-at-a-5-billion-valuation-in-ddn-ai-and-data-intelligence-solutions-leader-to-fuel-further-rapid-growth/>, January 2025.
- [52] VAST Data. VAST series f financing at \$30 billion valuation. <https://www.vastdata.com/press-releases/vast-series-f-financing-at-30-billion-valuation>, April 2026.
- [53] CNBC. Nvidia backs AI company VAST Data. <https://www.cnbc.com/2026/04/22/nvidia-backs-ai-company-vast-data.html>, April 2026.
- [54] NVIDIA. NVIDIA and storage industry leaders unveil new class of enterprise infrastructure for the age of AI. <https://nvidianews.nvidia.com/news/nvidia-and-storage-industry-leaders-unveil-new-class-of-enterprise-infrastructure-for-the-age-of-ai>, March 2025.
- [55] SNIA. SNIA announces Storage.AI. https://www.snia.org/news_events/newsroom/announces-storageai, August 2025.
- [56] Llama Team, AI @ Meta. The llama 3 herd of models. arXiv:2407.21783, July 2024. <https://arxiv.org/abs/2407.21783>.
- [57] Meta Engineering. Meta’s AI storage blueprint at scale. <https://engineering.fb.com/2026/07/01/data-infrastructure/metas-ai-storage-blueprint-at-scale/>, July 2026.
- [58] OpenAI. Careers: Software engineer, compute–storage. <https://openai.com/careers/software-engineer-compute-storage-san-francisco/>, 2026. Accessed July 2026.
- [59] MLCommons. MLPerf storage v2.0 benchmark results. <https://mlcommons.org/2025/08/mlperf-storage-v2-0-results/>, August 2025.
- [60] Computing Research Association. New CRA Taulbee survey findings show record degree production alongside a cooling enrollment pipeline. <https://cra.org/crn/2026/06/cra-update-new-cra-aulbee-survey-findings-show-record-degree-production-alongside-a-cooling-enrollment-pipeline/>, June 2026.
- [61] Fortune. AI companies court AI PhDs with huge pay packages, raising fears of an academic brain drain. <https://fortune.com/2025/06/25/ai-companies-court-ai-phds-with-huge-pay-packages-raising-fears-of-an-academic-brain-drain/>, June 2025.
- [62] George Amvrosiadis, Ali R. Butt, Vasily Tarasov, Erez Zadok, and Ming Zhao. Data storage research vision 2025: Report on NSF visioning workshop. Technical report, National Science Foundation, 2019. Workshop co-chairs, with contributions from workshop participants. <https://par.nsf.gov/servlets/purl/10086429>.

- [63] MLCommons. MLPerf storage v1.0 benchmark results. <https://mlcommons.org/2024/09/mlperf-storage-v1-0-benchmark-results/>, September 2024.
- [64] DeepSeek-AI. Fire-flyer file system (3FS). <https://github.com/deepseek-ai/3FS>, February 2025.
- [65] NVIDIA. How to reduce KV cache bottlenecks with NVIDIA Dynamo. <https://developer.nvidia.com/blog/how-to-reduce-kv-cache-bottlenecks-with-nvidia-dynamo/>, September 2025.
- [66] Tom Coughlin. C4Q 2025 and 2025 hard disk drive industry update. Forbes, <https://www.forbes.com/sites/tomcoughlin/2026/02/02/c4q-2025-and-2025-hard-disk-drive-industry-update/>, February 2026.
- [67] TrendForce. DRAM price forecast for 2025. <https://www.trendforce.com/presscenter/news/20241118-12365.html>, November 2024.
- [68] TrendForce. 4q25 server DRAM contract prices rise 43–48% qoq. <https://www.trendforce.com/presscenter/news/20250924-12733.html>, September 2025.
- [69] TrendForce. 2q26 memory contract price update. <https://www.trendforce.com/presscenter/news/20260331-12995.html>, March 2026.
- [70] Network World. Server memory prices could double by 2026 as AI demand strains supply. <https://www.networkworld.com/article/4093752/server-memory-prices-could-double-by-2026-as-ai-demand-and-strains-supply.html>, 2025. Counterpoint Research module-price data.
- [71] Kenanga Research. 2QCY26 Strategy: Market Strategy. <https://www.kenanga.com.my/wp-content/uploads/2026/04/Market-Strategy-260401-2QCY26-Strategy-Kenanga.pdf>, 2026.
- [72] Ellie Wang. HBM wafer allocation outlook, TrendForce presentation at FMS 2025. https://files.futurememorystorage.com/proceedings/2025/20250805_BMKT-102-1_Ellie-Wang.pdf, August 2025.
- [73] Micron Technology. Fiscal Q1 2026 earnings materials. <https://investors.micron.com/static-files/088991c5-a249-4f66-a0a6-258d9b66f3f9>, December 2025.
- [74] TrendForce. 3q26 memory price forecast and 2027 supply outlook. <https://www.trendforce.com/presscenter/news/20260709-13140.html>, July 2026.
- [75] Micron Technology. Micron technology, inc. reports results for the fourth quarter and full year of fiscal 2023. <https://investors.micron.com/news-releases/news-release-details/micron-technology-inc-reports-results-fourth-quarter-and-full-6>, September 2023.
- [76] SK hynix. SK hynix Reports Second Quarter 2023 Financial Results. <https://news.skhynix.com/sk-hynix-reports-second-quarter-2023-financial-results/>, 2023.
- [77] Western Digital. Western digital reports fiscal second quarter 2023 financial results. <https://investor.wdc.com/news-releases/news-release-details/western-digital-reports-fiscal-second-quarter-2023-financial>, 2023.
- [78] heise online. WD and Seagate confirm: Hard drives for 2026 sold out. <https://www.heise.de/en/news/WD-and-Seagate-confirm-Hard-drives-for-2026-sold-out-11178917.html>, February 2026.
- [79] Astera Labs. Leo CXL smart memory controllers on Microsoft Azure M-series virtual machines. <https://www.asteralabs.com/news/astera-labs-leo-cxl-smart-memory-controllers-on-microsoft-azure-m-series-virtual-machines-overcome-the-memory-wall/>, November 2025.

- [80] Tom’s Hardware. Meta fights soaring hardware costs by reusing old DDR4 server memory via a custom CXL 2.0 chip. <https://www.tomshardware.com/pc-components/dram/meta-fights-soaring-hardware-costs-by-reusing-old-ddr4-server-memory-in-new-ddr5-only-servers-custom-cxl-2-0-chip-marries-legacy-ddr4-2400-with-cutting-edge-ddr5-6400>, June 2026.
- [81] Andres Lagar-Cavilla, Junwhan Ahn, Suleiman Souhlal, Neha Agarwal, Radoslaw Burny, Shakeel Butt, Jichuan Chang, Ashwin Chaugule, Nan Deng, Junaid Shahid, Greg Thelen, Kamil Adam Yurtsever, Yu Zhao, and Parthasarathy Ranganathan. Software-defined far memory in warehouse-scale computers. In *Proceedings of the 24th International Conference on Architectural Support for Programming Languages and Operating Systems (ASPLOS ’19)*, 2019. <https://research.google/pubs/software-defined-far-memory-in-warehouse-scale-computers/>.
- [82] LWN.net. Weighted memory interleaving. <https://lwn.net/Articles/948037/>, 2023.
- [83] TechRadar Pro. Samsung and AMD made a revolutionary SSD together — then it was left to wither in the shadows and nobody knows exactly why. <https://www.techradar.com/pro/samsung-and-amd-made-a-revolutionary-ssd-together-then-it-was-left-to-wither-in-the-shadows-and-nobody-knows-exactly-why>, 2025.
- [84] Eideticom. Los alamos labs deploys advanced computational storage. <https://www.eideticom.com/blog/los-alamos-labs-deploys-advanced-computational-storage-eideticom-news>, 2022.
- [85] Everspin Technologies. Q3 2025 financial results (form 8-k exhibit). <https://www.sec.gov/Archives/edgar/data/1438423/000162828025049548/mram-20250930xexx991.htm>, 2025.
- [86] Tom’s Hardware. Kioxia preps 10-million-IOPS XL-Flash SSD with peer-to-peer GPU connectivity. <https://www.tomshardware.com/pc-components/ssds/kioxia-works-with-nvidia-to-prep-xl-flash-ssd-thats-3x-faster-than-any-ssd-available-10-million-iops-drive-has-peer-to-peer-gpu-connectivity-for-ai-servers>, 2025.
- [87] Blocks & Files. DNA data storage: When the physics work but the economics don’t. <https://www.blocksandfiles.com/data-protection/2026/01/16/dna-data-storage-when-the-physics-work-but-the-economics-dont/4090361>, January 2026.
- [88] National Academies of Sciences, Engineering, and Medicine. *Rapid Expert Consultation on Archival Data Storage Technologies for the Intelligence Community*. The National Academies Press, 2024. <https://doi.org/10.17226/27445>.
- [89] Blocks & Files. Cerabyte roadmap. <https://blocksandfiles.com/2025/07/02/cerabyte-roadmap/>, July 2025.
- [90] SemiAnalysis. CXL is dead in the AI era. <https://newsletter.semianalysis.com/p/cxl-is-dead-in-the-ai-era>, 2023.
- [91] ServerMall. CXL in 2026: Memory expansion and pooling. <https://servermall.com/blog/cxl-in-2026-memory-expansion-and-pooling/>, 2026. Analyst enablement estimates.
- [92] Introl. GPU infrastructure TCO: A 5-year cost model. <https://introl.com/blog/gpu-infrastructure-tco-5-year-cost-model>, 2025.
- [93] San Diego Supercomputer Center. CloudBank 2.0: \$20m NSF award. <https://www.sdsc.edu/news/2025/PR20250409-CloudBank.html>, April 2025.

- [94] NAIRR Pilot. About the national AI research resource pilot. <https://nairrpilot.org/about>, 2026. Accessed July 2026.
- [95] Amazon Web Services. Amazon S3 pricing. <https://aws.amazon.com/s3/pricing/>, 2026. Accessed July 2026.
- [96] Sarah Wang and Martin Casado. The cost of cloud: A trillion-dollar paradox. Andreessen Horowitz, <https://a16z.com/the-cost-of-cloud-a-trillion-dollar-paradox/>, May 2021.
- [97] Amazon Web Services. Data egress waiver available for eligible researchers and institutions. <https://aws.amazon.com/blogs/publicsector/data-egress-waiver-available-for-eligible-researchers-and-institutions/>, 2021.
- [98] National Science Foundation. CISE community research infrastructure (CIRC), solicitation NSF 23-589. <https://www.nsf.gov/funding/opportunities/circ-community-infrastructure-research-computer-information/12810/nsf23-589/solicitation>, 2023. Program status listed as Archived, July 2026.
- [99] National Science Foundation. Cyberinfrastructure for sustained scientific innovation (CSSI), solicitation NSF 22-632. <https://www.nsf.gov/funding/opportunities/cssi-cyberinfrastructure-sustained-scientific-innovation/nsf22-632/solicitation>, 2022.
- [100] National Science Foundation. Mid-scale research infrastructure-1, solicitation NSF 24-598. <https://www.nsf.gov/funding/opportunities/mid-scale-ri-1-mid-scale-research-infrastructure-1/nsf24-598/solicitation>, 2024.
- [101] National Science Foundation. Mid-scale research infrastructure-2 program update. <https://www.nsf.gov/funding/opportunities/mid-scale-ri-2-mid-scale-research-infrastructure-2/updates/118947>, 2025.
- [102] National Science Foundation. Advanced computing systems & services (ACSS), solicitation NSF 24-583. <https://www.nsf.gov/funding/opportunities/advanced-computing-systems-services-adapting-rapid-evolution/nsf24-583/solicitation>, 2024.
- [103] National Science Foundation. Integrated data systems and services (IDSS), solicitation NSF 25-544. <https://www.nsf.gov/funding/opportunities/idss-integrated-data-systems-services/nsf25-544/solicitation>, 2025.
- [104] National Science Foundation. NSF ACCESS awardees will advance innovations. <https://www.nsf.gov/news/nsf-access-awardees-will-advance-innovations>, 2022.
- [105] ACCESS. Open storage network resource description. <https://allocations.access-ci.org/resources/osn.access-ci.org>, 2026. Accessed July 2026.
- [106] HPCwire. NSF’s antypas reflects on successes of NAIRR pilot at TPC26. <https://www.hpcwire.com/2026/06/04/nsfs-antypas-reflects-on-successes-of-nairr-pilot-at-tpc26/>, June 2026.
- [107] National Science Foundation. NAIRR operations center, solicitation NSF 25-546. <https://www.nsf.gov/funding/opportunities/nairr-oc-foundations-operating-national-artificial-intelligence/nsf25-546/solicitation>, 2025.
- [108] National Science Foundation. National artificial intelligence research institutes. <https://www.nsf.gov/focus-areas/ai/institutes>, 2026. Accessed July 2026.

- [109] National Science Foundation. NSF announces \$100 million investment in national artificial intelligence research institutes. <https://www.nsf.gov/news/nsf-announces-100-million-investment-national-artificial>, July 2025.
- [110] Computing Research Association. President releases devastating NSF budget request. <https://cra.org/govaffairs/blog/2025/06/president-releases-devastating-nsf-budget-request-proposes-to-turn-back-the-clock-more-than-20-years-on-nsf-funding-join-the-discussion-and-take-action/>, June 2025.
- [111] Computing Research Association. FY26 CJS and E&W minibus analysis. <https://cra.org/govaffairs/blog/2026/01/fy26-cjs-ew-minibus/>, January 2026.
- [112] Science. Exclusive: NSF faces radical shake-up as officials abolish its 37 divisions. <https://www.science.org/content/article/exclusive-nsf-faces-radical-shake-officials-abolish-its-37-divisions>, 2025.
- [113] Computing Research Association. National science board dismissed; FY2027 request analysis. <https://cra.org/govaffairs/blog/2026/04/nsb-fired-2026/>, April 2026.
- [114] NERSC. Doudna storage solutions. <https://www.nersc.gov/news-and-events/news/doudna-storage-solutions>, July 2025.
- [115] Blocks & Files. VAST Data cracks into HPC with doudna supercomputer win. <https://www.blocksandfiles.com/ai-ml/2025/07/04/vast-data-cracks-into-hpc-with-doudna-supercomputer-win/1591929>, July 2025.
- [116] Oak Ridge Leadership Computing Facility. ORNL, AMD, and HPE to deliver DOE’s newest AI supercomputers: Discovery and lux. <https://www.olcf.ornl.gov/2025/10/27/ornl-amd-and-hpe-to-deliver-does-newest-ai-supercomputers-discovery-and-lux/>, October 2025.
- [117] NVIDIA. NVIDIA, Oracle, and US department of energy AI supercomputers for scientific discovery. <https://nvidianews.nvidia.com/news/nvidia-oracle-us-department-of-energy-ai-supercomputer-scientific-discovery>, October 2025.
- [118] Argonne Leadership Computing Facility. DAOS overview: Aurora data management. <https://docs.alcf.anl.gov/aurora/data-management/daos/daos-overview/>, 2026. Accessed July 2026.
- [119] Oak Ridge Leadership Computing Facility. Orion file system. <https://www.olcf.ornl.gov/olcf-resources/data-visualization-resources/orion/>, 2026. Accessed July 2026.
- [120] NERSC. Perlmutter scratch file system. <https://docs.nersc.gov/filesystems/perlmutter-scratch/>, 2026. Accessed July 2026.
- [121] PESO Project. 2025 PESO project report. <https://pesoproject.org/files/2025PESOProjectReport.pdf>, January 2026.
- [122] CASS. Darshan: CASS software catalog. <https://cass.community/software/darshan.html>, 2026. Accessed July 2026.
- [123] insideHPC. DOE ASCR: \$35m available for HPC data management. <https://insidehpc.com/2025/01/doe-ascr-35m-available-for-hpc-data-management-proposal-deadline-may-13/>, January 2025.
- [124] DOE ASCR. Workshop on management and storage of scientific data. <https://doi.org/10.2172/1845707>, January 2022.

- [125] Computing Research Association. DOE office of science FY2027 budget request analysis. <https://cra.org/govaffairs/blog/2026/05/doe-sc-fy2027-pbr/>, May 2026.
- [126] The White House. Executive order: Accelerating federal permitting of data center infrastructure. <https://www.whitehouse.gov/presidential-actions/2025/07/accelerating-federal-permitting-of-data-center-infrastructure/>, July 2025.
- [127] The White House. Executive order: Launching the genesis mission. <https://www.whitehouse.gov/presidential-actions/2025/11/launching-the-genesis-mission/>, November 2025.
- [128] U.S. House Committee on Science, Space, and Technology. H.R. 2385, CREATE AI act. <https://science.house.gov/2026/6/h-r-2385-create-ai-act>, June 2026.
- [129] Office of Senator Todd Young. Young, colleagues introduce bill to advance AI innovation, reliability for americans. <https://www.young.senate.gov/newsroom/press-releases/young-colleagues-introduce-bill-to-advance-ai-innovation-reliability-for-americans/>, 2026.
- [130] Office of Senator Martin Heinrich. American science acceleration project (ASAP). <https://www.heinrich.senate.gov/asap>, 2025.
- [131] GovTech. Two years into NAIRR pilot, shared infrastructure boosts AI innovation. <https://www.govtech.com/education/higher-ed/2-years-into-nairr-pilot-shared-infrastructure-boosts-ai-innovation>, 2025.
- [132] Center for Security and Emerging Technology. The NAIRR pilot: Estimating compute. <https://cset.georgetown.edu/article/the-nairr-pilot-estimating-compute/>, 2024.
- [133] Tamay Besiroglu, Sage Andrus Bergerson, Amelia Michael, Lennart Heim, Xueyun Luo, and Neil Thompson. The compute divide in machine learning: A threat to academic contribution and scrutiny? arXiv:2401.02452, January 2024. <https://arxiv.org/abs/2401.02452>.
- [134] Ji-Ung Lee, Haritz Puerto, Betty van Aken, Yuki Arase, Jessica Zosa Forde, Leon Derczynski, Andreas Rücklé, Iryna Gurevych, Roy Schwartz, Emma Strubell, and Jesse Dodge. Surveying (dis)parities and concerns of compute hungry NLP research. arXiv:2306.16900, 2023. <https://arxiv.org/abs/2306.16900>.
- [135] Chameleon Cloud. Composable hardware on chameleon now. <https://chameleoncloud.org/blog/2024/08/19/composable-hardware-on-chameleon-now/>, August 2024.
- [136] University of Chicago CS. Chameleon testbed secures \$12 million in funding for phase 4. <https://cs.uchicago.edu/news/chameleon-testbed-secures-12-million-in-funding-for-phase-4-expanding-frontiers-in-computer-science-research/>, August 2024.
- [137] CloudLab. CloudLab hardware documentation. <https://docs.cloudlab.us/hardware.html>, 2026. Accessed July 2026.
- [138] National Science Foundation. Award 2431419: CloudLab phase IV. https://www.nsf.gov/awardsearch/showAward?AWD_ID=2431419, 2024. Verified via NSF Awards API, July 2026.
- [139] FABRIC Testbed. FABRIC public project record. <https://portal.fabric-testbed.net/experiments/public-projects/990d8a8b-7e50-4d13-a3be-0f133ffa8653>, 2026. Accessed August 2026.
- [140] Texas Advanced Computing Center. Stampede3 user guide. <https://docs.tacc.utexas.edu/hpc/stampede3/>, 2024.

- [141] Pittsburgh Supercomputing Center. NSF grants extend bridges-2 and neocortex at PSC. <https://www.psc.edu/nsf-grants-extend-bridges-2-and-neocortex-at-psc/>, May 2026.
- [142] San Diego Supercomputer Center. Open storage network expansion. https://www.sdsc.edu/news/2024/PR20241022_OSN.html, October 2024.
- [143] Oak Ridge Leadership Computing Facility. ACE testbed documentation. https://docs.olcf.ornl.gov/ace_testbed/index.html, 2026. Accessed July 2026.
- [144] Grid'5000. Grid'5000 testbed. <https://www.grid5000.fr/>, 2026. Accessed July 2026.
- [145] DOE Office of Science. ASCR FY2025 congressional budget request. https://science.osti.gov/-/media/budget/pdf/sc-budget-request-to-congress/2025/Advanced-Scientific-Computing-Research-3_15_24---Final.pdf, 2024.
- [146] Jaime Cernuda, Luke Logan, Noah Lewis, Suren Byna, Xian-He Sun, and Anthony Kougkas. Jarvis: Towards a shared, user-friendly, and reproducible I/O infrastructure. Work-in-progress presentation, 9th International Parallel Data Systems Workshop (PDSW '24), held with SC24, 2024. <https://pdsw.org/pdsw24/index.shtml>.
- [147] Hariharan Devarajan, Huihuo Zheng, Anthony Kougkas, Xian-He Sun, and Venkatram Vishwanath. DLIO: A data-centric benchmark for scientific deep learning applications. In *Proceedings of the 21st IEEE/ACM International Symposium on Cluster, Cloud and Internet Computing (CCGrid '21)*, 2021.
- [148] CILogon. CILogon: An integrated identity and access management platform for science. <https://www.cilogon.org/>, 2026. Accessed July 2026.
- [149] OSG Consortium. Open science data federation (OSDF). <https://osg-htc.org/services/osdf>, 2023.
- [150] SLATE Project. SLATE: Services layer at the edge. <https://slateci.io/>, 2020.
- [151] John Monroe. Ministering our dataverse: The need for new technologies. Further Market Research; Designing Storage Architectures for Digital Collections, Library of Congress, 2026. https://digitalpreservation.gov/meetings/DSA2026/0107_monroe_Library%20of%20Congress_DSA_09%20March%202026%20V11.83-Final.pdf.
- [152] Georg Lauhoff and Sassan Shahidi. Data storage trends: NAND, HDD and tape storage. Designing Storage Architectures for Digital Collections, Library of Congress, 2025. https://digitalpreservation.gov/meetings/DSA2025/010201_lauhoff_LoC2025_IBM_G_Lauhoff.pdf.
- [153] Hewlett Packard Enterprise. HPE LTO-9 Ultrium 45TB RW Data Cartridge. <https://buy.hpe.com/us/en/storage/storage-media/tape-media/hpe-lto%E2%80%99-ultrium-45tb-rw-data-cartridge/p/q2079a>, 2026. Accessed August 2026.