

COMPUTING INFRASTRUCTURE BREAKOUT– wrap up

Where does HEP computing go from here?

CLARIPHY AI Collaboration Meeting • Madison, WI • Sept 14–16, 2026

Session Organizers

Daniel Díaz (SDSC / NRP / CMS) • Ian Fisk (Flatiron Institute) • Rob Gardner (UChicago / ATLAS) • Oliver Gutsche (FERMILAB)

Today's computing model: distributed, pledged, tiered

The WLCG tier-1 / tier-2 / tier-3 structure is the resource mode. Institutions pledge dedicated hardware, sized for steady-state simulation, reconstruction, and analysis at LHC data rates.

- Built around dedicated, owned hardware pledges per institution
- Sized for predictable, **batch-style** CPU-intensive work
- Worked well when compute and storage costs were falling year over year

Tier-0

CERN — raw data, first-pass processing

Tier-1

National centers — storage, reprocessing

Tier-2

Institutional — simulation, analysis

Tier-3

Local / campus — user analysis

That model is now colliding with hardware prices

3–5×

rise in DRAM prices since early 2025: the AI memory shortage, expected to persist to 2030

35%

of system bill of materials is now memory, up from 15–18% a quarter earlier

~2 yr

backorder on enterprise disk; 2026 supply was booked out before February

The tier model assumed ~170 institutions could each keep buying the same hardware every year, and that falling \$/unit would turn flat money into growing capacity. Both halves of that assumption broke at once.

And the scale gap is widening, not closing

2026 AI capital spend: four hyperscalers vs. DOE's entire Genesis call

\$760B



Microsoft, Alphabet, Meta, Amazon combined capex

\$293M



DOE Genesis Mission FY26 funding call · a ratio of ~2,600 : 1

20–40×

More CPU and disk that ATLAS and CMS need at HL-LHC under current computing models. Flat budgets deliver roughly 2× over the same period — a shortfall of 10–20×.

HL-LHC starts June 2030. We are in Long Shutdown 3 right now — this is the window where the computing model gets decided.

Both agencies are consolidating — at the same time

DOE[†]

- HEP research line: \$431M → \$283M in the FY27 request (−34%)
- Genesis Mission: 20 HEP awards from 381 applications (5.2%)
- Awards are large, multi-institution, industry-partnered
- No new major construction starts expected for 6–7 years

NSF[‡]

- ~6,100 new grants in FY26 — lowest since the early 1980s
- 46% below the 2021–2024 average
- Cutting the number of solicitations by half or more
- Explicitly “focusing on longer awards with larger funding amounts”

The one P5 computing recommendation DOE is funding (4f: shared software, sustained computing R&D) is being delivered “in support of the DOE-wide Genesis Mission.” Computing is the surviving line.

[†][HENP for DPF July 2026](#)

[‡]<https://www.nature.com/articles/d41586-026-02574-6>

The demand is changing shape, not just size

The tier model was built for one workload: CPU batch, highly parallel, latency-tolerant. Three different shapes now sit on top of it, and each wants a different machine and a different operating model.

CPU batch

Simulation, reconstruction, analysis
trains

Queue for hours. Scales out. Latency does
not matter.

Pledged cores

GPU training

Foundation models, surrogate
simulation

Dense multi-GPU nodes, scheduled,
memory-bound.

Hard to pledge

GPU inference

Reconstruction, triggering, tagging

Bursty, latency-bound, wants to be a
service.

Service-shaped

Interactive / agentic

LLM access, agentic workflows, coding

Always on, sub-second, metered per
token.

Budget-shaped

Discussion Topic: Industry partnerships

Both agencies have made partnership a condition of the money, not a bonus. The community has decades of precedent for doing this well — the question is whether we scale it.

It is written into the awards

Every HEP Genesis selection is multi-institution. AMD, Google, NVIDIA, AWS, SuperMicro and GlobalFoundries appear across the list — as named partners, not vendors.

It is how NAIRR exists at all

~\$100M of in-kind contribution from 28+ companies: NVIDIA \$30M, Microsoft \$20M in Azure credits, plus OpenAI, Google, AWS, Meta, IBM, Intel and Anthropic.

We already have a working model

CERN openlab has run public-private R&D since 1983 — currently Oracle, Siemens, NVIDIA, Intel, PaaS and others, with a defined governance and cost-sharing structure.

11

Discussion Topic: Joint procurements

Joint hardware procurement

- Multiple institutions/experiments co-propose for shared hardware on a single award.
- Distributed across sites and made available to the wider community.
- Aggregate demand could support larger, better-optimized procurements and avoid every institution independently buying specialized hardware.
- The attraction is that it matches what both agencies say they want to fund: fewer, larger, multi-institution awards.

Discussion Topic: Inference-as-a-service (e.g., SONIC)

The problem it solves

Experiment workflows are CPU-bound and latency-tolerant; the ML inside them wants a GPU. Rather than putting a GPU in every worker node, the job calls out to a shared inference server.

How it works

SONIC is the client pattern inside experiment software; SuperSONIC is the Kubernetes-based server infrastructure — Triton, load balancing, autoscaling — so the CPU-to-GPU ratio can be tuned independently of the job mix.

Who uses it

CMS for ParticleNet, DeepMET, DeepTau and Part; ATLAS for ExaTriXX and Traccc; IceCube for CNN event classifiers. Deployed at Purdue (Geddes, Anvil), NRP Nautilus, and the UChicago ATLAS Analysis Facility.

Why it matters here

It is a template for treating expensive hardware as a shared service rather than a per-site purchase — which is exactly the shape a metered, allocation-based model would need.

Source: SuperSONIC documentation, fastmachinelearning.org.

More Discussion Ideas

- Where are the most promising synergies across communities?
- What concrete collaborative opportunities exist near-term?
- Are there achievable demonstrator projects to build momentum?
- What resources, infrastructure, or org. support is needed?
- What workforce / training / community needs should we address?
- Inference-as-a-service, sustaining infra, industry partnerships?

Discussion Topic: Industry partnerships

Framing: "Forward deployed scientists" — science still gets done at universities and labs; industry's motive is revenue, not discovery

Leverage: find what makes us useful — benchmarks/evaluations, and data for training (trade data access for compute?)

Advocacy: use HEP's congressional advocacy as the model for making the case to industry

Governance — open questions:

- How do we stay in charge of which scientific questions get asked?
- What do we share — calibration data, metadata, papers, technotes?
- At what point don't they need us anymore?

What we want back: money, tokens — and a public return when the data and expertise are publicly funded

Loose ends: joint procurement under Genesis is a political problem; startups (not just frontier labs) as expertise-exchange partners

Link to our report: FDP × CMS Open Data — *owner-set access policy under Genesis* — is the concrete mechanism for staying in charge

Discussion Topic: Inference-as-a-service (e.g., SONIC)

laaS + shared procurement for GPU work; users want to deploy their own models, not just consume hosted ones

Training is the blocker, not inference — would a hand-off service (submit data → get model back) unblock people?

Tension: hand-off is efficient, but people want to learn to do it themselves — culture and training matter as much as capacity

Allocation management and equity are unsolved

Discovery gap: no way to find what laaS exists or how to get access

Expect federation, not central provisioning — many Foundation Models with different data I/O, some services hosted by experiments

Open: is there will for a community-wide physics foundation model? Are we pushing people to agentic workflows, or will a good enough product pull them?

Discussion Topic: Joint procurements

There is precedent for buying together

- CMS and ATLAS jointly procured silicon sensors — in principle the better model
- CERN negotiated with Dell; anyone HEP-affiliated could buy at CERN prices; years ago, US LHC Tier2s common portal

Extend the same model to computing

- Aggregate demand across campuses, not just across experiments
- Federated identity is the analogy — if we can align there, why not on purchasing?
- Funding agencies are the right actors to drive this

Why it is hard

- Procurement cycles differ experiment to experiment
- Vendor sales teams and territories
- Institutional purchasing requirements — RFQs, competitive bid rules

Bottom line: the technical case is easy; many organizational pieces have to align at once, and agency-level action is the lever

Recommendations for CLARIPHY

- Inference-as-a-service: Do we need to invest into big foundations models or do we enable the model of small individual models?
- Industry partnership: We should use our HEP advocacy to congress as a model for advocacy to industry
- Joint procurement is a bureaucratic nightmare. Better the funding agencies negotiate hardware prices with vendors for large tenders and everyone gets the same prices (if that is possible)

Computing Infrastructure — addendum

Synergies — Open data is the fertile ground — cross-experiment, cross-framework (incl. theory); metadata drives software selection · shared access to AmSC / NAIRR services and lab GPUs (ORNL, ANL) with allocation models that fit project timescales · a common platform for agentic-tooling development

Near-term opportunity — Fermi Data Platform × CMS Open Data — ~3 PB; under Genesis the data owners set access policy (open and closed); joint demand case to DOE

Demonstrator — IRI-pathfinder-style cross-facility workflow (precedent: AI + Harvester on Perlmutter, job manager to ALCF/ORNL/NERSC via IRI) — needs a compute + latency requirements manifest, containers for the software environment, federated ID for experiment groups

Resources needed — Resource needs are time-dependent (CPU → GPU training → CPU): capacity at some times, capabilities at others · shared development allocations for collaborative projects (US ATLAS student NERSC model) · effort to port experiment software environments · network bandwidth as a first-class resource for IaaS

Workforce — Training on writing allocation proposals for training and inference services · equitable access to tokens across the community