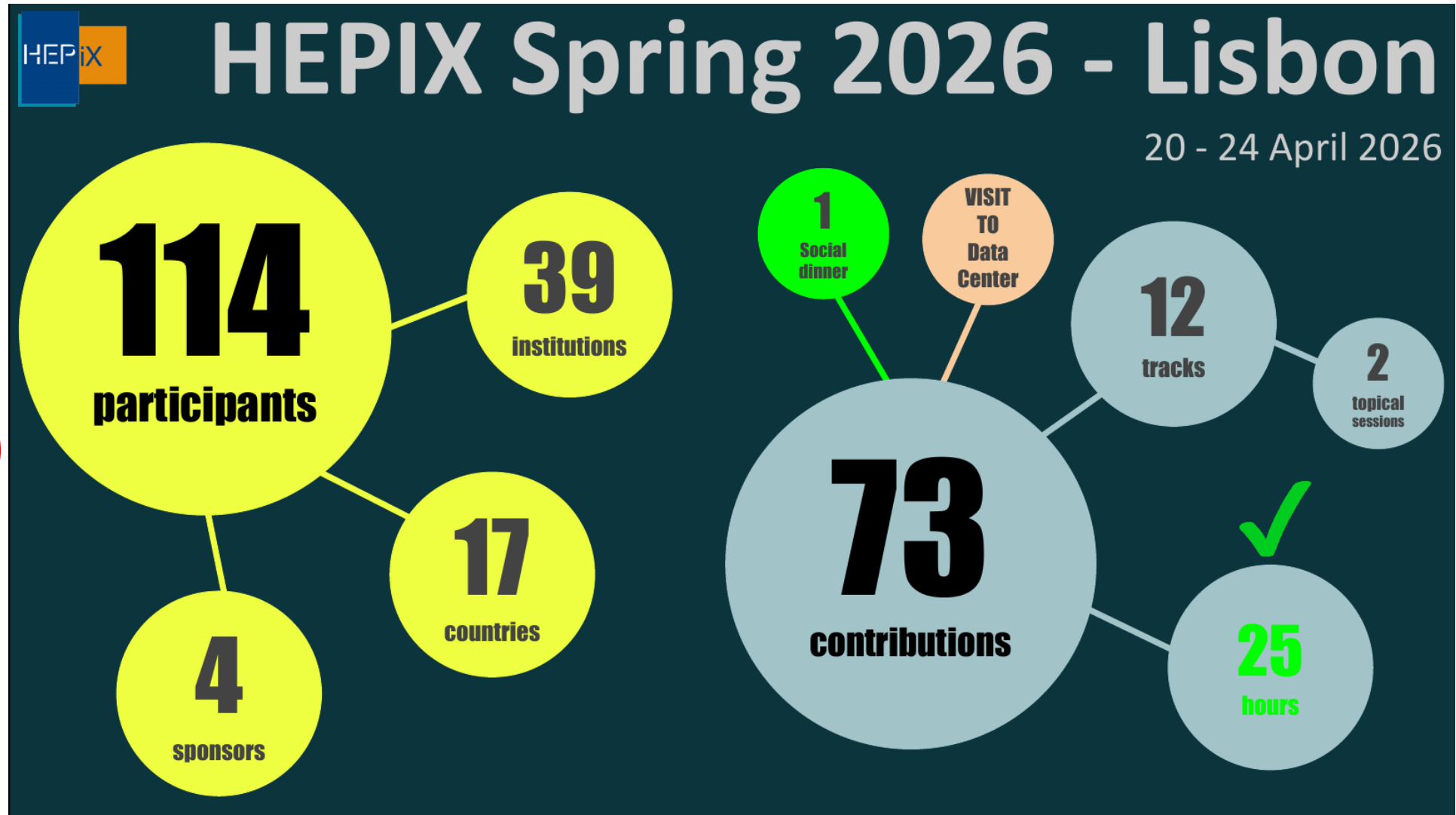


Centre de Calcul
de l'Institut National de Physique Nucléaire
et de Physique des Particules

Retour HEPiX printemps @Lisbonne

20-24 avril 2025

- **Volonté de resserrer les liens avec WLCG**
 - **Un représentant WLCG au board**
 - **Une ou des sessions thématiques (lien avec WLCG) à chaque workshop**
 - @Lisbonne : « WLCG OTF mid-long term evolution of Facilities »
 - Discussion débutée à Lanzhou (automne 2025)
- **Création d'une nouvelle session (track) sur l'IA**
 - **Pour répondre à un nombre de talks croissants sur cette thématique depuis 3~4 HEPiX**
 - **« Applied AI in computing center infrastructure »**
- **Élection du co-chair asiatique (mandat de 4 ans, renouvelable une fois)**
 - **Tomoaki Nakamura s'est proposé de nouveau, et a été ré-élu.**
- **Actuellement, co-chair européen : Pepe Flix, co-chair nord-américain : Ofer Rind**



- Site Reports [11]
- Session thématique : évolution des sites à moyen-long terme (1/2) [7]
- Session thématique : évolution des technologies (basé sur le HEPiX WG TechWatch) (2/2) [8]
- Computing and batch system [5]
- Applied AI in computing center infrastructure [4]
- Storage [6]
- Software [7]
- Security and networking [6]

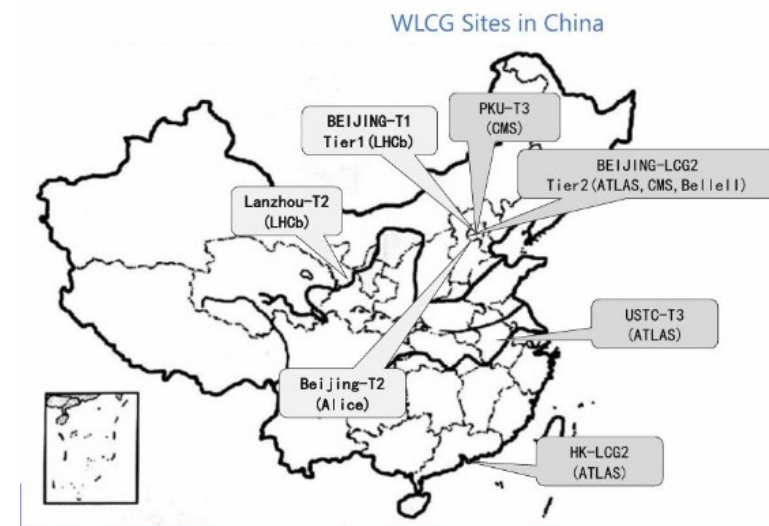
Site Reports

Site Reports I



- **11 site reports** (total 3h40) – sharing their developments, issues and solutions
- Same **usual tools deployed** at sites: HTCondor/Slurm, dCache (grid), Ceph/Lustre, OpenStack, k8s, NextCloud
- More **AI tools (LLM) at sites**: anomaly detection, help to write code...
 - Usually various LLM are provided: Mistral, Qwen3, gpt-oss-120b, ... (models are local)
- **New Computing Room at RAL**
 - ... utilising waste heat, part of de-carbonisation plan.
- **Also new datacentre in Guimarães, PT, managed by LIP**
- **CERN: new director general and new head of IT**
 - Security: password policy change (based on NIST 800-63b): length (15 characters), not complexity
- **IHEP**: growing number of (Tier-2) sites

- Lanzhou **2025**
- Beijing-T2 Alice **2024**



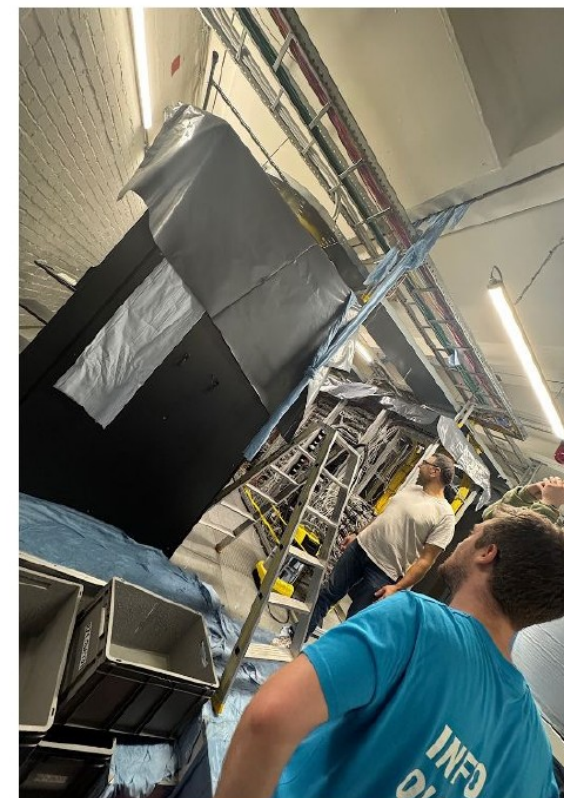
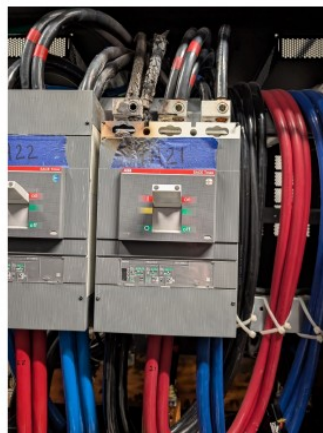
Site Reports II



- **DESY: “Fun” with water!**
 - Heavy rain made the ceiling to come down with a lot of water from the drain pipes
- **Fire at GSI (February ‘26)**
 - Fortunately nobody injured
 - FAIR not directly damaged, but UNILAC was
- **Brookhaven power overload**



J. Flix - Lisbon Workshop Wrap-up





Medium to long term evolution of facilities

(Follow up to the Lanzhou topical session)

Topical Chapters and Facilitators (Editors)



- Introduction - TCB Chairs
- Chapter 1 - Experimental Computing Requirements
- • Chapter 2 - Facility Evolution - A. Dewhurst, J. Flix, A. Melo, A. De Salvo
- Chapter 3 - Data Management - K. Ellis, J. Elmsheuser, plus contributions from FTS experts
- Chapter 4 - Network Infrastructure and Management - S. McKee, E. Martelli
- Chapter 5 - Workflow Management - M. Mascheroni
- Chapter 6 - Security and Authentication & Authorization Infrastructure - M. Litmaath, J.C. Luna
- Chapter 7 - Services - M. Litmaath, J. Andreeva, P. Paparrigopoulos
- Chapter 8 - New Architectures and Infrastructures - D. Benjamin, I. Osborne, O. Smirnova
- Chapter 9 - Sustainability - D. Britton (our contact with the [WLCG Sustainability Forum](#))
- Conclusions - TCB Chairs
- References

*Aiming for about 5 pages for each chapter - We don't need a lot of text, we need good text!
Each chapter will list several key high-level milestones, risks & opportunities, etc.*

Follow-up on mid-long term evolution of facilities 1/2

Topical session (HEPIX + OTF)



HEPiX and [WLCG Open Technical Forum](#) (OTF#10) joint session

- **Facilities** are one of the **building blocks** of WLCG
 - Facility **evolution** is driven by the various **communities'** requirements,
 - and vice-versa - community requirements **change** depending on the **opportunities** that Facilities enable
- HL-LHC challenges: **WLCG Technical Roadmap**

Follow-up on mid-long term evolution of facilities (Topical Session)



Topical session jointly organized with WLCG OTF that aims to explore and converge on likely directions for WLCG facility evolution, informed by ongoing developments in storage, workflow management, and resource provisioning. By bringing together cross-site and cross-experiment perspectives on CPU, GPU, disk, tape, and network evolution, the goal is to identify a small set of concrete milestones and a 3–5 year timeline for deploying and validating production-ready infrastructure in advance of Run 4, feeding directly into the WLCG Technical Roadmap.

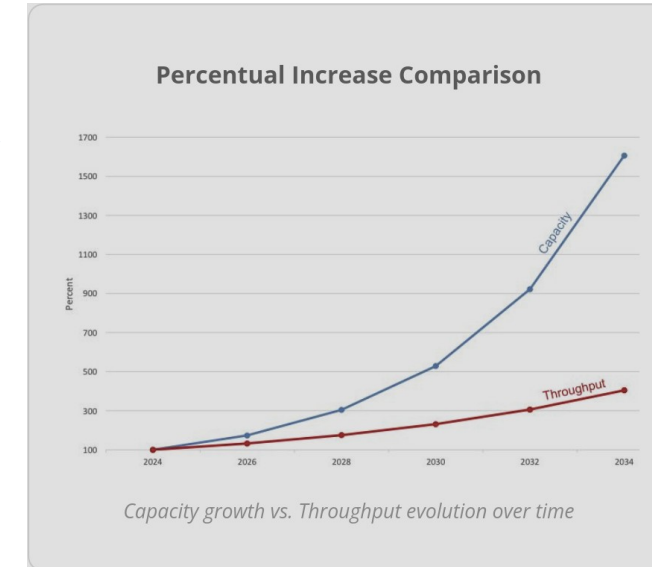
Follow-up on mid-long term evolution of facilities 2/2

Topical session (HEPIX + OTF)



Topical Presentations and Discussions (see the [live notes](#) for details, executive summary):

- **Operational Effort at Sites**
 - Goal is to understand the highest time consuming areas (and what can be optimized) ⇒ plan to make a **survey on operational effort** by next HEPiX
- **Evolution of tape (archival) storage incl. performance requirements**
 - Summary of [OTF#8](#) “tape jamboree”
- **Storage Capacity vs. Throughput**
 - Capacity is increasing significantly quicker than throughput both for tape and disk
 - Flash storage may help, but no unified approach (caching, tiered, or stand-alone)
- **Facilities Evolution: U.S. CMS and IHEP (China)**
 - U.S. CMS future processing purchases at the Tier-2 sites, storage at Fermilab
 - IHEP: recent network upgrades and new configurations, planning GPU purchases
- **Services Deployments with K8s - panel discussion**
 - Important to choose the “right tool for the job”
 - Good for workforce development, but a steep learning curve

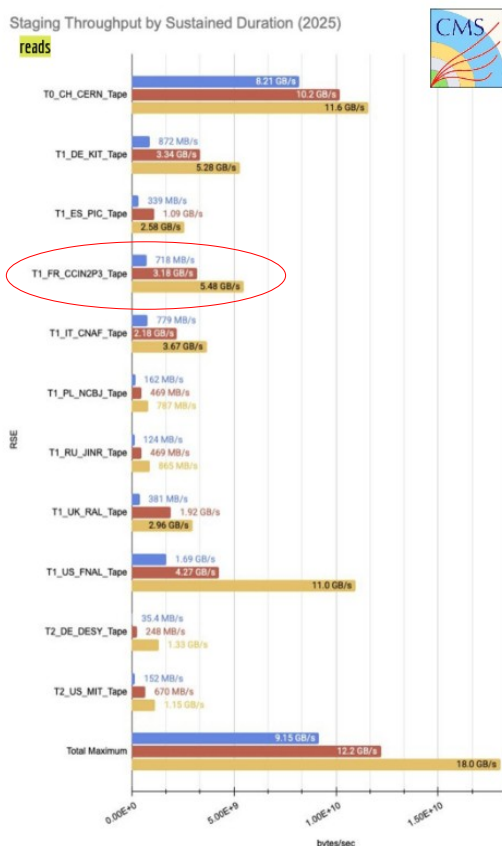


More robots required to maintain the performances (similar pattern for disks)

CMS Performance Tape Tests

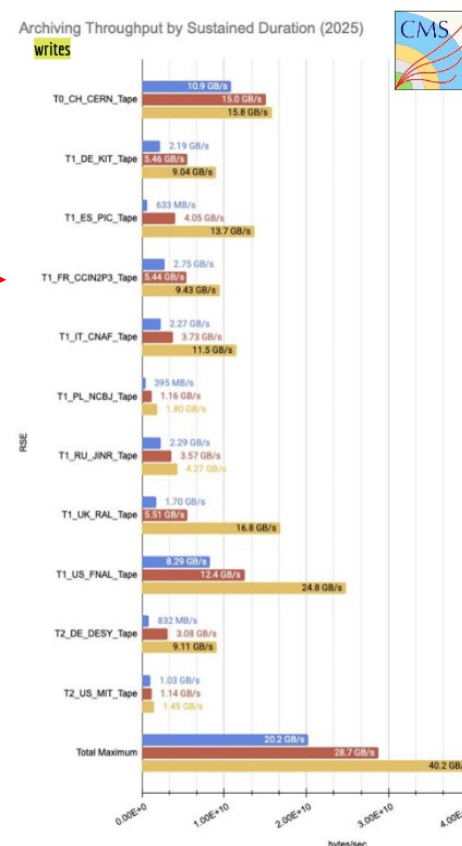


The following charts detail the sustained **read and write throughputs** from **2025 production data**, establishing a **performance baseline** to guide future optimizations



CC-IN2P3

9.2 GB/s over a full week
12.2 GB/s over a day
18.0 GB/s over an hour



⇒ Voir la track « storage »
(« tape jamboree »)

20.2 GB/s over a full week
28.7 GB/s over a day
40.2 GB/s over an hour

- **Perspectives (issues des discussions)**
 - **Nouveau sondage sur les efforts opérationnels des sites**
 - Finaliser en mai
 - Réponse attendue en septembre (au plus tard)
 - Présentation & discussion lors du HEPiX d'automne (octobre 2026)
 - **Quelques pistes de discussions à poursuivre**
 - Tape performances : Capacity Vs Throughput
 - To cope with the HL-LHC expected performances
 - Site evolution
 - Storage (concentration du stockage sur certains sites, ie Tier-1s ?)
 - Coût important d'après un précédent sondage...
 - Orchestration (steep learning curve, « choose the right tool »)
 - Ne remet pas en cause les conteneurs, juste la couche d'abstraction en plus



Technology Watch

Second topical session

Hardware technology and cost in 2026 (A. Sciabà)

- « The objective of this talk is (as usual) to report on the status of hardware technology but with a strong focus on cost. »

- L'IA a l'origine de la crise du « hardware »

- Les investissements explosent

- ~2500 milliards \$ estimés en 2026 !

- Prix de RAM et disques explosent

- DRAM x 4~5

- SSD x 470 % (see VDURA storage cost, slide 30)

- Impact on WLCG (« existential threat for WLCG flat budget »)

- Améliorations des perfs ne vont pas compenser la hausse des coûts

- CPU : should be ok as most sites bought CPU last year

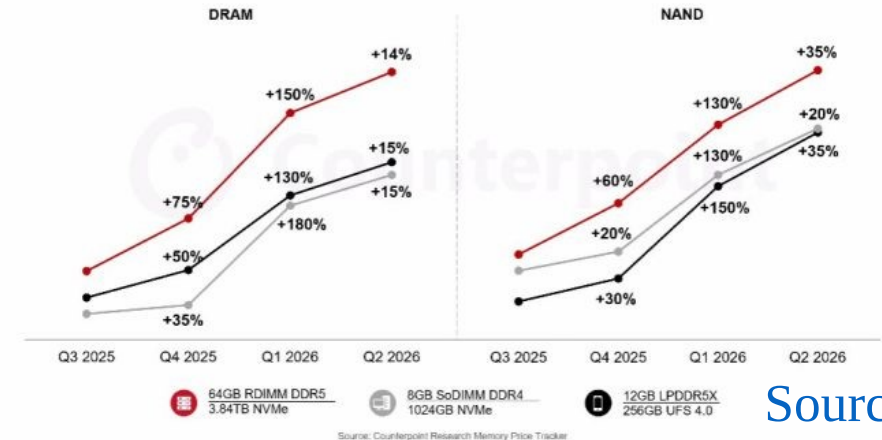
- Storage : critical

	Previous (Circa 2021)	New (June 2025)	CMS Actual (Run1-3)
CPU	(15+/-5)%	(10+/-5)%	13%
Disk	(15+/-5)%	(5+/-5)%	15%
Tape	(15+/-5)%	(10+/-5)%	20%

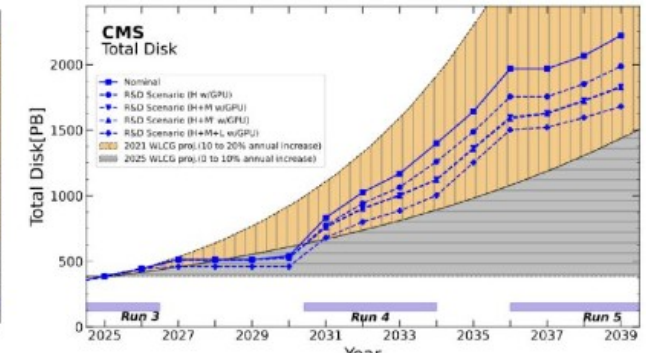
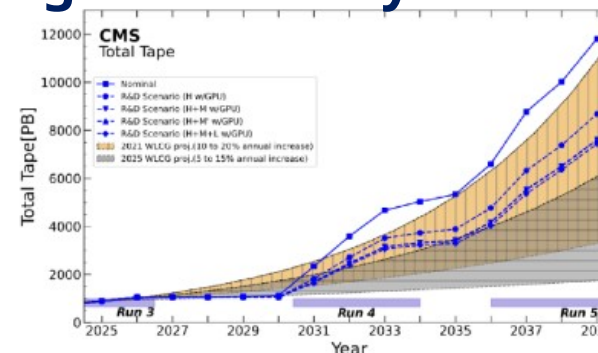
Prices More Than Double in Q1 2026 vs. Q4 2025

Counterpoint

Price Increase by Quarter (Memory Price Tracker)



Source



Mitigation measures (personal thoughts)

- **Delay / advance procurement**
 - Hoping for AI bubble to burst to pop soon seems an increasingly bad bet
- **Reduce hardware specifications**
 - Sites which were generous with memory cannot afford it anymore
 - Maybe even moving memory modules from old to new servers?
- **Consider using DDR4**
 - Cheaper than DDR5, it might be an option in some specific cases
- **Extend hardware lifetime**
 - Do not retire old hardware unless necessary
- **Find more opportunistic resources**
 - Exploit as many HPCs as possible
- **Resource requests reduction by experiments**
 - Push for as much R&D as possible
 - Make GPUs a cost-efficient alternative to CPUs
- **Convince funding agencies to allocate more money to hardware**
 - As difficult as it sounds, it might make the difference between being able to process Run4 data or not



Computing and batch system

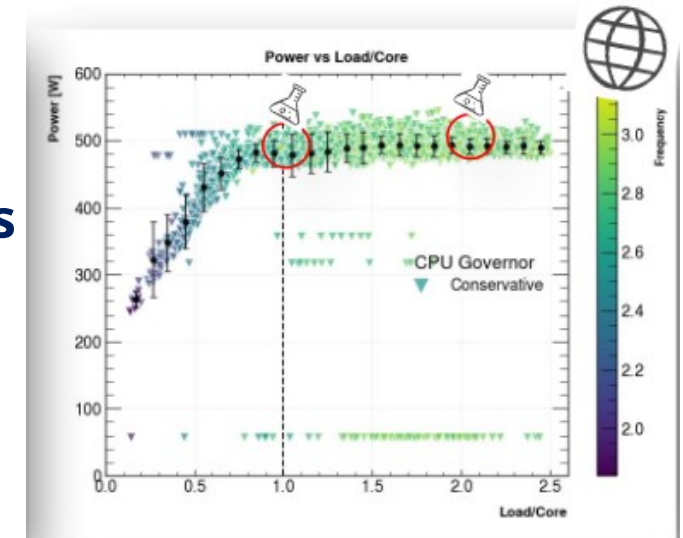
Benchmarking WG status report (D. Giordano)

- Le benchmarking GPU devient concret
 - 8 workloads (réels) identifiés
 - Validation sur infra réelles sur sites avec GPU
 - P100, V100, H100, A100, L40S, L4, ...
 - Call for more α -testers (see [Wiki instructions](#))
- Corrélation Performances ET énergie dépensée
 - Base de données : ~275 configurations issues de 33 sites
 - $\Rightarrow$ Performance Vs Efficiency (HS23/W)

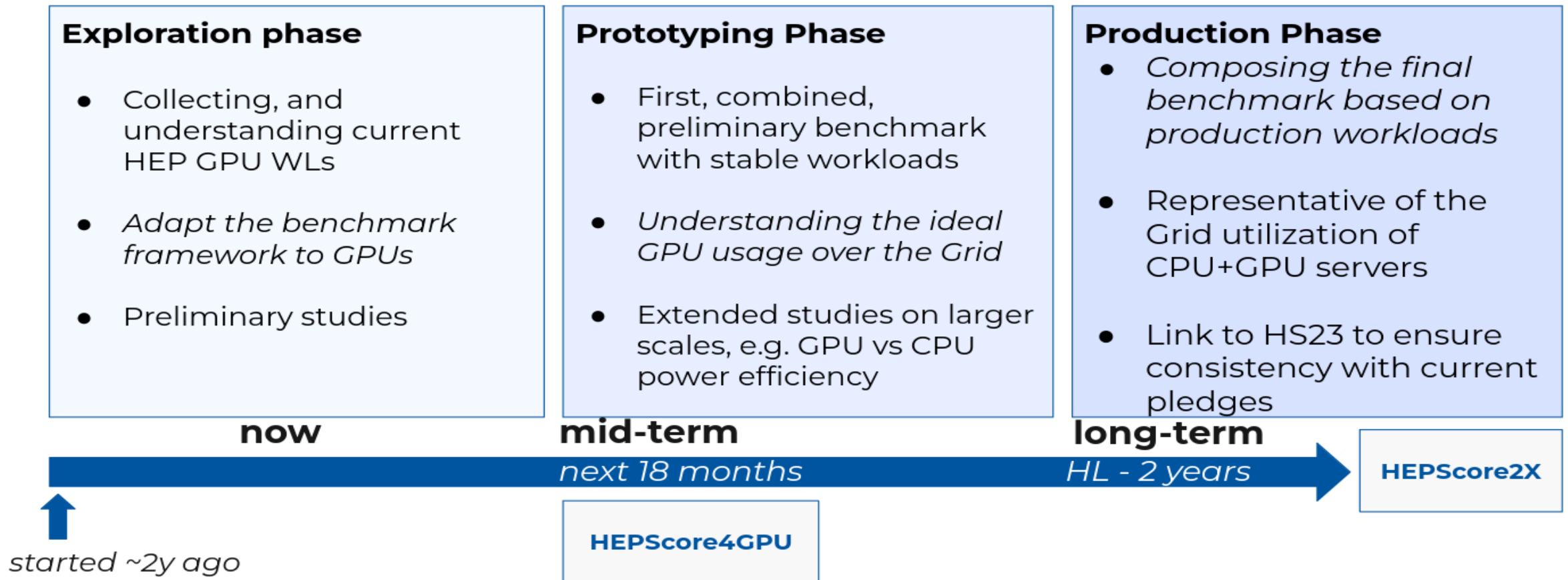
Workload	Type	Current Version
GEN-Pepper-GPU	Event Gen	✓ v0.4
GEN-MadGraph-GPU	Event Gen	✗ ci-v0.2
ATLAS-FullSim-GPU	Sim	✓ v0.1
CMS-FlowSim-GPU	Inference	✓ v0.3
CMS-HLT-GPU	HLT	✗ ci-v0.2
Alice-O2-GPU	HLT	🔍
CMS-MLPF-GPU*	Training/Inference	🔍
LHCb-Fullsim-GPU*	SIM	🔍

NB: CMS FlowSim replicates the inference step of the CMS FlashSim simulation framework

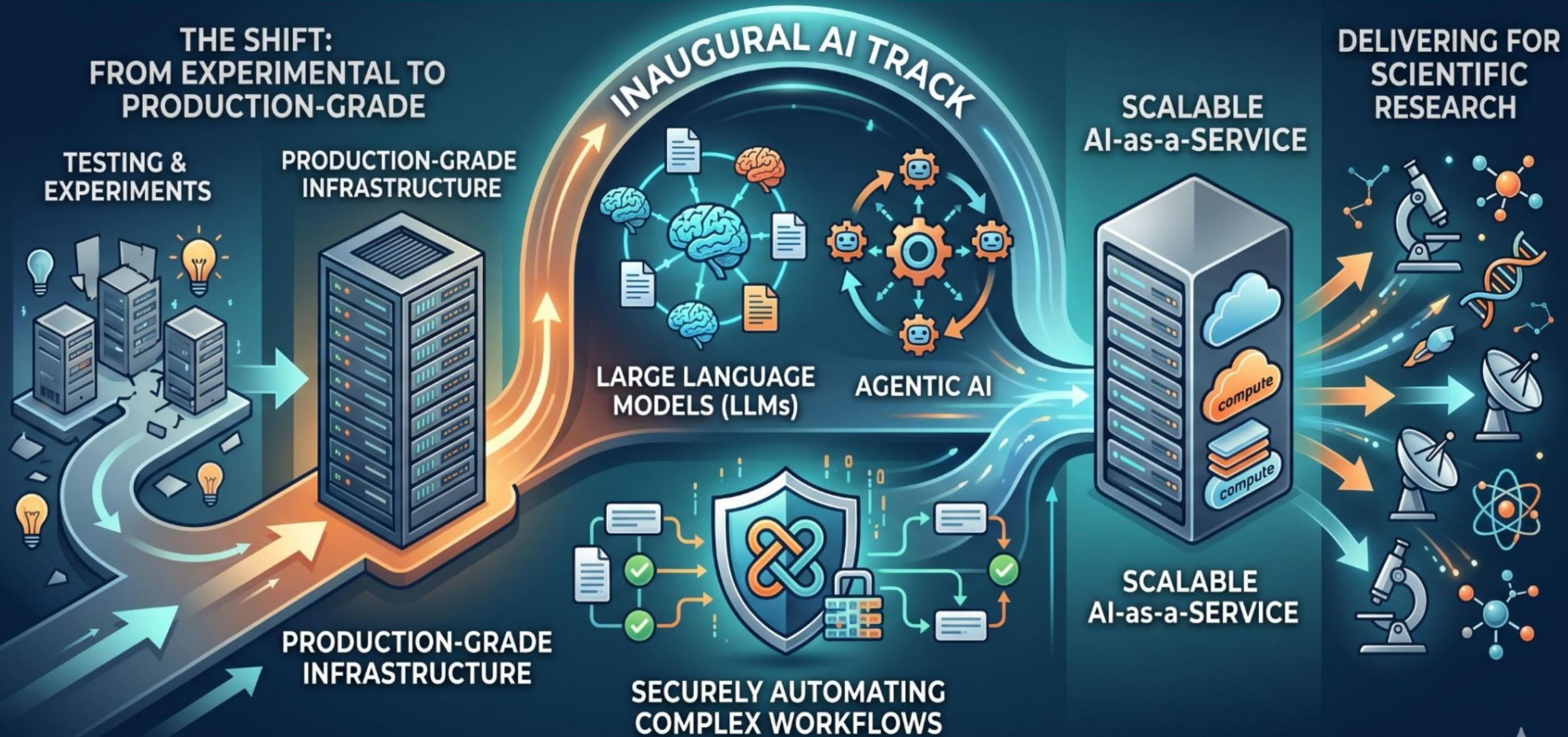
- ✓ – stable (pre-)release version
- ✗ – development version
- 🔍 – under investigation (no release available yet)



GPU Benchmarking Strategy



INAUGURAL AI TRACK: ACCELERATING SCIENTIFIC DISCOVERY



EMPOWERING THE HEPiX COMMUNITY WITH VALUABLE REAL EXPERIENCES

Integrating Spark, Mlflow, and Agentic AI in Maxwell Cluster at DESY : Advancing reproducible and intelligent scientific workflows



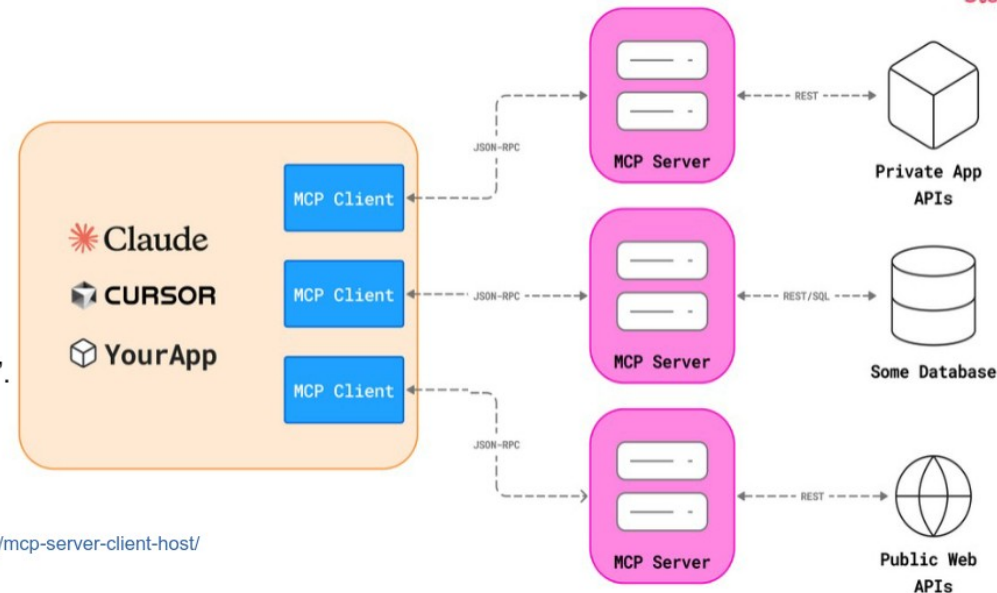
- **IA-workflows reproducible**

- **Spack pour la reproductibilité**
- **MLflow pour la sauvegarde des (meta-)données d'apprentissage (IA)**
- **Automatisation avec agents IA**

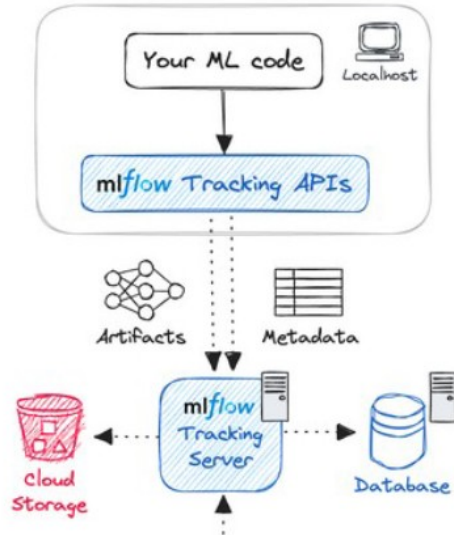
- **Executes** an action, decision (how to refine the query) is done by agent.
- **Iterates** through **reasoning** and **acting** (ReAct) until **goal** is achieved.
- Follows **Model Context Protocol (MCP)**, an open-source standard for connecting AI applications to external systems.
- <https://modelcontextprotocol.io/examples>

Example:

“Find the most recently modified Python file in \$HOME, create a Slurm batch script for it, submit it with sbatch, and **report back the job ID and output** once it completes”.

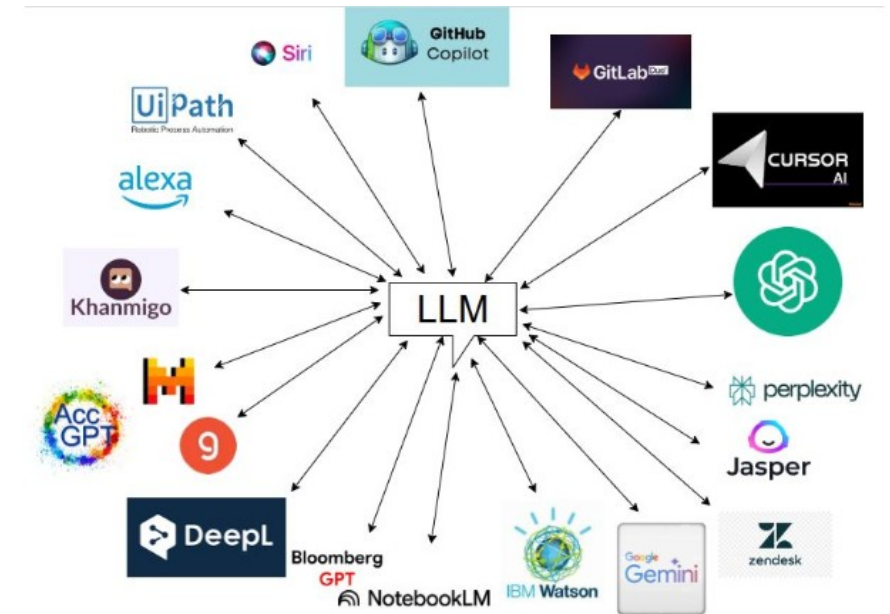


3. Remote Tracking w/ Tracking Server



A centralized LLM and Agentic AI infrastructure for CERN

- From PoC to production
 - +80 use cases identifiés
 - +600 utilisateurs utilisent AccGPT (CERN chatbot, 2023)
 - ⇒ central IT service
- Shared platform components
 - Portail vers grands modèles de langues (LLM)
 - Création d'agents
 - MCP (multi context protocol)
 - RAG



LLM / Agentic AI Service: Roadmap

➤ Phase 1 — Q2 2026: First PoC and Use Cases

- Inference gateway, monitoring, accounting
- Cataloging, versioning and discovery of LLM models
- Onboarding selected use cases from initial requests

➤ Phase 2 — Q4 2026: Agentic AI

- Infrastructure for AI agents and MCP tools
- Management of RAG pipelines (scraping, vectors, embeddings)
- Onboarding of selected Agentic AI use cases

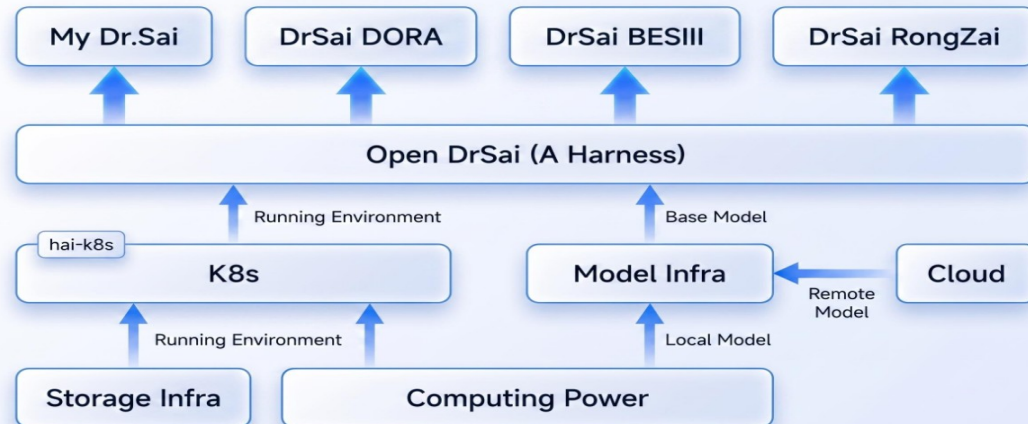
➤ Phase 3 — Q2 2027: Production & General Availability

- Long-term sustainability, day-2 operations
- Integration with resources portal for lifecycle & costing
- Extensive documentation and training courses



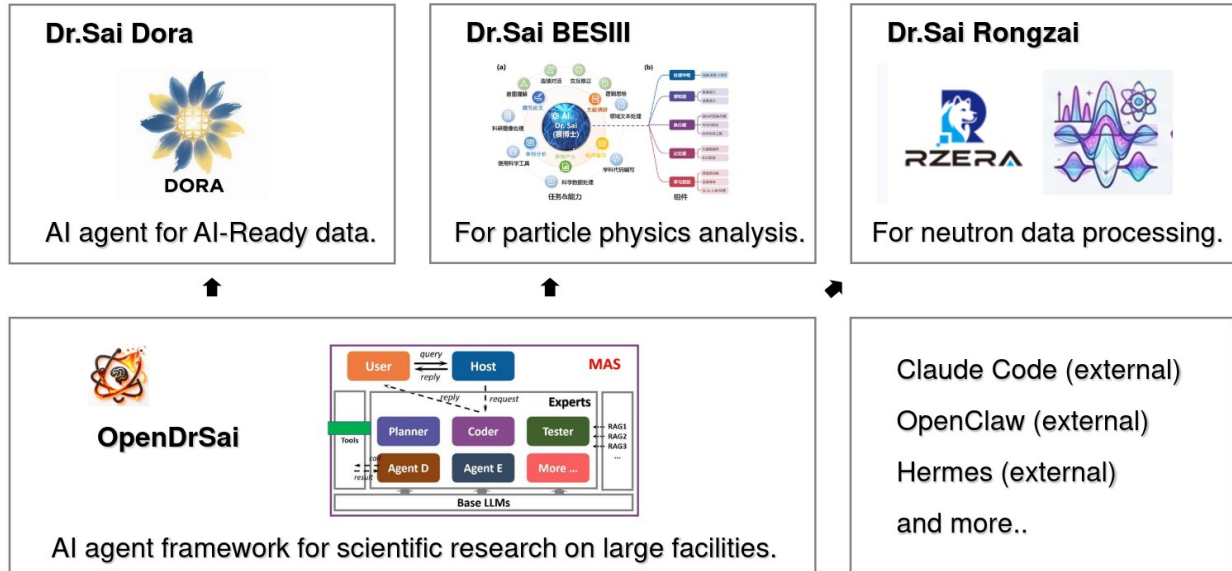
New AI infrastructure for Dr.Sai Agents at IHEP

Architecture



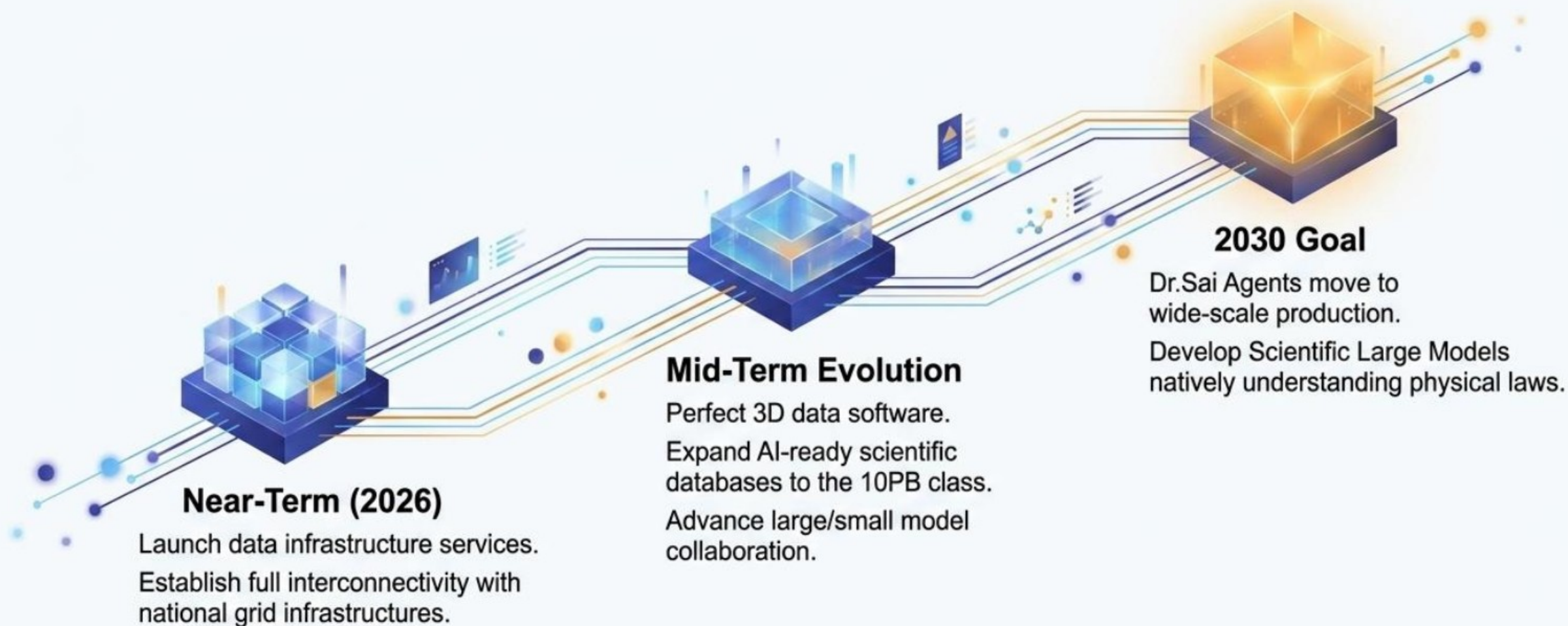
- From AI models to AI agents
 - Construction d'un écosystème complet basé sur l'IA
 - Data infra DOMAS
 - Models serving Qionwu
 - Elastic AI computing services
 - Agents thématiques déjà disponibles

The Dr.Sai Agent Matrix



IHEP construit une couche transverse basée sur l'IA : AaaS (AI-As-A-Service)

Roadmap to 2030: The High Energy AI Ecosystem



Connect & Build With Us:

- HepAI Platform: ai.ihep.ac.cn
- Dr.Sai Agent: opendrsai.ihep.ac.cn
- Open Source: github.com/hepai-lab/drsai
- Contact: zdzhang@ihep.ac.cn

Storage

https://indico.cern.ch/event/1598655/contributions/7028375/attachments/3260206/5820713/HEPIX_Spring_2026_OTF8_Report_JFlix.pdf



CTA - Increasingly adopted across WLCG sites (CERN, RAL, FNAL, PIC)

- Support for archive metadata to improve file colocation
- Improved scheduling and drive orchestration
- Separating repacking workflows from production traffic
- Tokens!



dCache

- Improved integration with CTA
- QoS-based access policies (which remain to be seen how they will be used)
- Bulk operations on large file collections
- Feature request in progress to prevent dCache from dropping archive metadata (currently, it accepts but drops anything not in the standard header)



Storm

- Ongoing development to support future requirements (e.g. StoRM now HTTP-based)
- Planning to look into support for archive metadata hints
- Possibly further automation of scheduling and drive allocation

CMS:



- Very **high tape reliability!**
- Many **migrations** (CTA, HTTP) and **new sites** commissioned during Run-3
- Challenges include **slow recalls** and **inconsistent performance** across sites
- We do need **improved prioritization** and **scheduling**
 - Increase average file size and hopefully, improve tape efficiency through archive metadata
 - **Expected HL-LHC T0 rates: 70GB/s**

LHCb:



- Typical **pattern**: write during runs, recall during reprocessing periods
- But concern over some sites when trying to do both simultaneously
- Tape staging seen as a “black box” operationally. Limited monitoring
 - **monitoring is difficult**: the chain of communication is extremely long
- Future work includes improving monitoring, using archival metadata and looking into possible prioritization mechanisms



ALICE

ALICE:

- Tape mainly used as **long-term archive** (~90% of stored data is RAW)
- **LS3** will involve **large-scale reprocessing campaigns**
- ALICE writes files in the specific order they should be read, **repacking can randomize this carefully planned order** - problems for future recalls

ATLAS:



- Major usage of **Data Carousel** to use tape actively (integrated into PanDA)
- Slow **tails cause issues**, similar to CMS
- Mostly likely will remain **below 10GB average file size in Run-4**
- Expected tape usage: **2.5 EB by 2031**

- **Target 80 %+ with read efficiency (= specific time-based metric that compares active data transfer time against the total time the tape drive is occupied) :**
 - Use co-location and archive metadata to close the gap between current write efficiency (>90%) and recall efficiency (40-50 %)
- **Define clear SLAs : establish concrete read/write rate commitments for all tape sites.**



Software

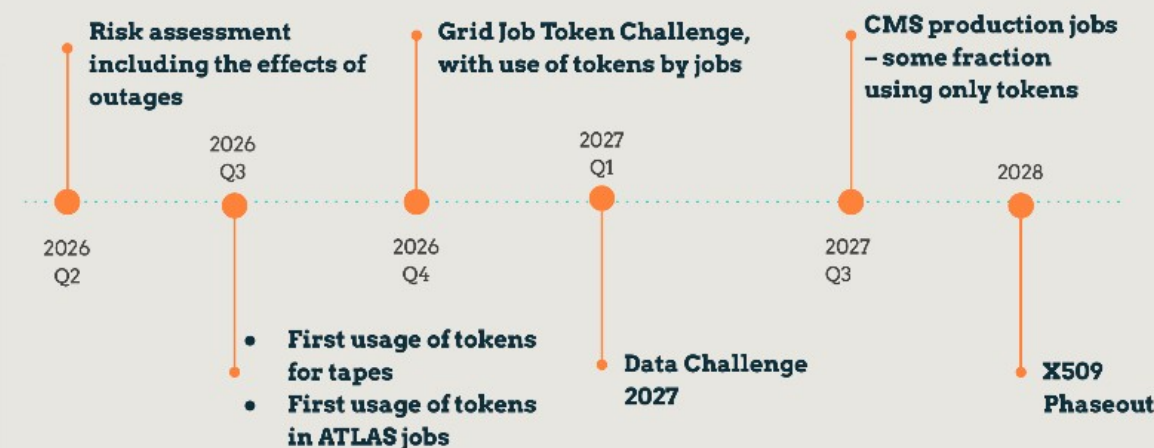
[https://indico.cern.ch/event/1598655/contributions/6988057/attachments/3260536/5821909/WLCG%20Operations%20Coordination%20towards%20HL-LHC_%20Current%20Focus%20Areas%20and%20Near-Term%20Plans%20-%20HEPiX%20Spring%202026%20\(1\).pdf](https://indico.cern.ch/event/1598655/contributions/6988057/attachments/3260536/5821909/WLCG%20Operations%20Coordination%20towards%20HL-LHC_%20Current%20Focus%20Areas%20and%20Near-Term%20Plans%20-%20HEPiX%20Spring%202026%20(1).pdf)

Notable token developments

Two WLCG Working Groups:
Authorization WG
Token Trust Traceability WG

- WLCG Common JWT Profiles v1.2 was published on Dec 15
 - Has changes in support of using tokens for tape operations
- WLCG standard system JWKS cache specification v1.0 was published on Jan 27
 - Will be used e.g. by the SciTokens libraries used by HTCondor and XRootD
- The token **usage** risk assessment is being finalized in the Token Trust & Traceability WG – report expected this month!
 - To be complemented by a risk assessment of IAM **downtimes**
- IAM v1.14 will permit access tokens no longer to be stored in the DB – ETA May
 - Will finally allow ATLAS FTS to use tokens for **all** disk-to-disk transfers
- IAM v2.0 will be based on the modern Spring Authorization Server and a modern React GUI – ETA autumn

Updated tentative milestones



Xrootd monitoring status and plans

- A new architecture and plan was presented in the last FTS/XRootD workshop and an update was given during the CMS O&C Spring Week Open Session.
- The development of all components is largely complete
 - The new Go-based collector is **deployed at CERN to evaluate** and improve!
 - xroot over dCache monitoring can already be achieved using dCache bookkeeping
 - Thanks to **PIC** and **FNAL** that are already looking into sending us traces
 - More **volunteers** are very welcome!
- Our primary focus now shifts to **data validation** and **evaluation** of the components, with particular attention to operational aspects.
 - Deployment campaign comes right after!
- The goal is **DC27 with XRootD monitoring** deployed at the majority of WLCG sites

19

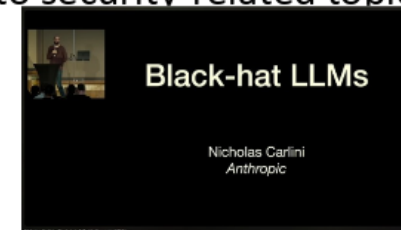
XRootD Monitoring update

WLCG Monitoring Task Force

- While FTS traffic is well monitored and trusted, we are still missing complete and reliable monitoring of xrootd traffic
 - XRootD traffic represents a substantial fraction of WLCG network usage in particular for remote data reading
- Currently only ALICE has a reliable implementation of XRootD monitoring
- There are two use cases :
 - Native monitoring flow of the xrootd servers
 - dcache with xrootd door where monitoring data is based on dcache bookkeeping information

Security and networking

- Likely already heard about Claude Mythos: “During our testing, we found that Mythos Preview is capable of identifying and then exploiting zero-day vulnerabilities in every major operating system and every major web browser when directed by a user to do so.” (<https://red.anthropic.com/2026/mythos-preview>)
 - Hype or real issue?
 - Can’t ignore the advantages AI gives to malicious actors (mass scanning of open source projects for bugs, black-box testing, AI-assisted phishing, malware development) to lower the cost and effort of cyber offense.
- **Good news: we can use it the same way they can!**
 - If you don’t do it now, someone else eventually will
- **Tips:**
 - Do it regularly. AI models improve over time
 - Try different models. Every model is better at something else, some are really sensitive to security-related topics.
 - Right now, most models are not thorough. Point it at specific subsystems and/or files
 - Watch this great Unprompted talk: <https://www.youtube.com/watch?v=1sd26pWhfmg>



- **Scanning and Understanding your Attack Surface:**

- Do you know all the services on your network? Create a system to keep an up-to-date list of targets BEFORE a critical vulnerability comes around

- **Scrapers and mitigation**

- Block using IP/CIDR/ASN reputation lists, apply rate limits, 2nd service instance without FW opening, Anubis

- **Supply Chain Attacks – Mitigation**

- Configure a cooldown (minimum age of a package), understand your package manager (e.g. ``npm ci`` strictly enforce lock files, ``npm install`` doesn't)

- **IoT Security – Vulnerabilities and Configuration**

- IoT devices are a huge source of vulnerabilities
- Yet, can't always be blamed on the vendor => deployment requirements are often specified by vendors and not followed by installers/contractors

- **Phishing – still a problem!**

- A user was tricked into handing over their password and 2FA token
- Shows 2FA is not a silver bullet

- **Suite de l'OTF « mid-long term evolution of facilities » à l'automne**
 - Sondage à venir, et conclusions données et discutées lors du prochain workshop
- **Benchmark $\Rightarrow$ HEPscore2x (HL - 2 ans)**
 - Préparation d'un bench GPU, et prise en compte de la performance Vs puissance
- **Software**
 - Notable developments concerning tokens, accounting, XRootD monitoring
- **Prochains workshops**
 - Automne 26 : Nebraska, 19-23 octobre 2026
 - Printemps 27 : KIT (dates non fixées)

Merci !

 HEPiX

[More Information](#)

See you at the next meeting in

University of Nebraska

Lincoln, USA

19-23 October 2026