



Towards Unified Power and Efficiency Monitoring Across the Worldwide LHC Computing Grid

Natalia Szczepanek,
Domenico Giordano
CERN, Geneva

Sustainable HEP 2026 - 5th edition

Importance of sustainability in WLCG

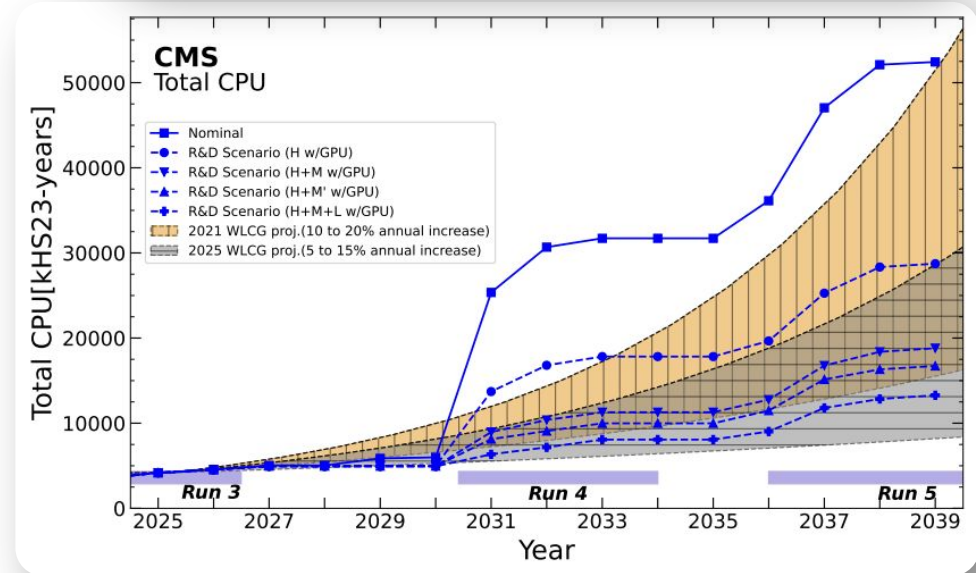
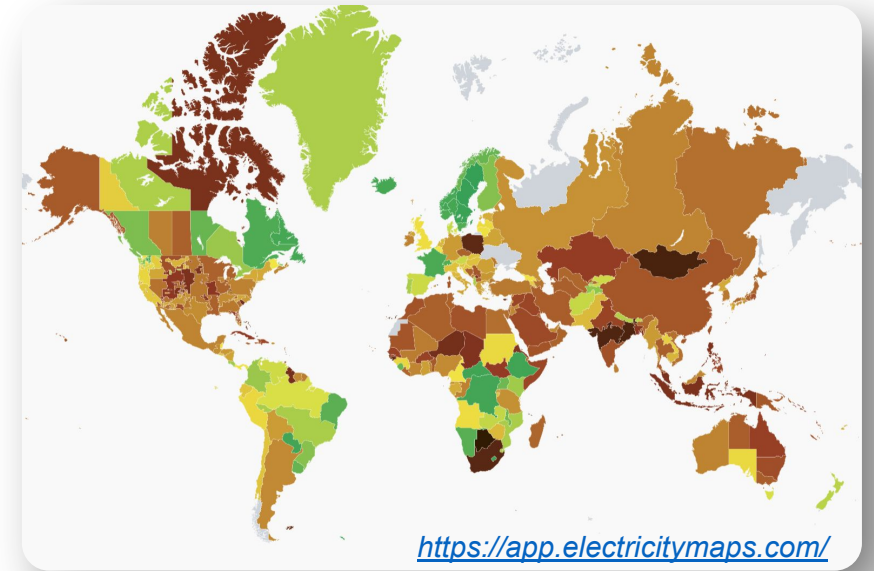
~200 Sites
across 42 countries

80%
of computing power
sits outside CERN

~1.4M CPU cores
used across WLCG

- HL-LHC Computing Needs:
 - Pessimistic scenario: x5 higher than in Run-3
 - Optimistic scenario: 40% higher than in Run-3
- Environmental impact of computing is strictly connected to **financial savings**

A better understanding of power consumption in WLCG will allow to make better, data-driven decisions.



<https://cds.cern.ch/record/2957472/files/LHCC-G-186.pdf>

Where are we? Three pillars of WLCG strategy

WLCG Strategy 2024-2027 – three environmental recommendations guiding the community

ENV-1

Measure & Monitor

Agree on metrics and a framework to collect energy-efficiency data.

Where we are:

- Ensure Project enabling power-consumption data collection
- WLCG starting to collect PUE per site

ENV-2

Efficient hardware

Facilitate more energy-efficient hardware as experiment software matures.

Where we are:

- GPU-capable generators, simulation & reconstruction
- Task Force discussing how to pledge GPU resources

ENV-3

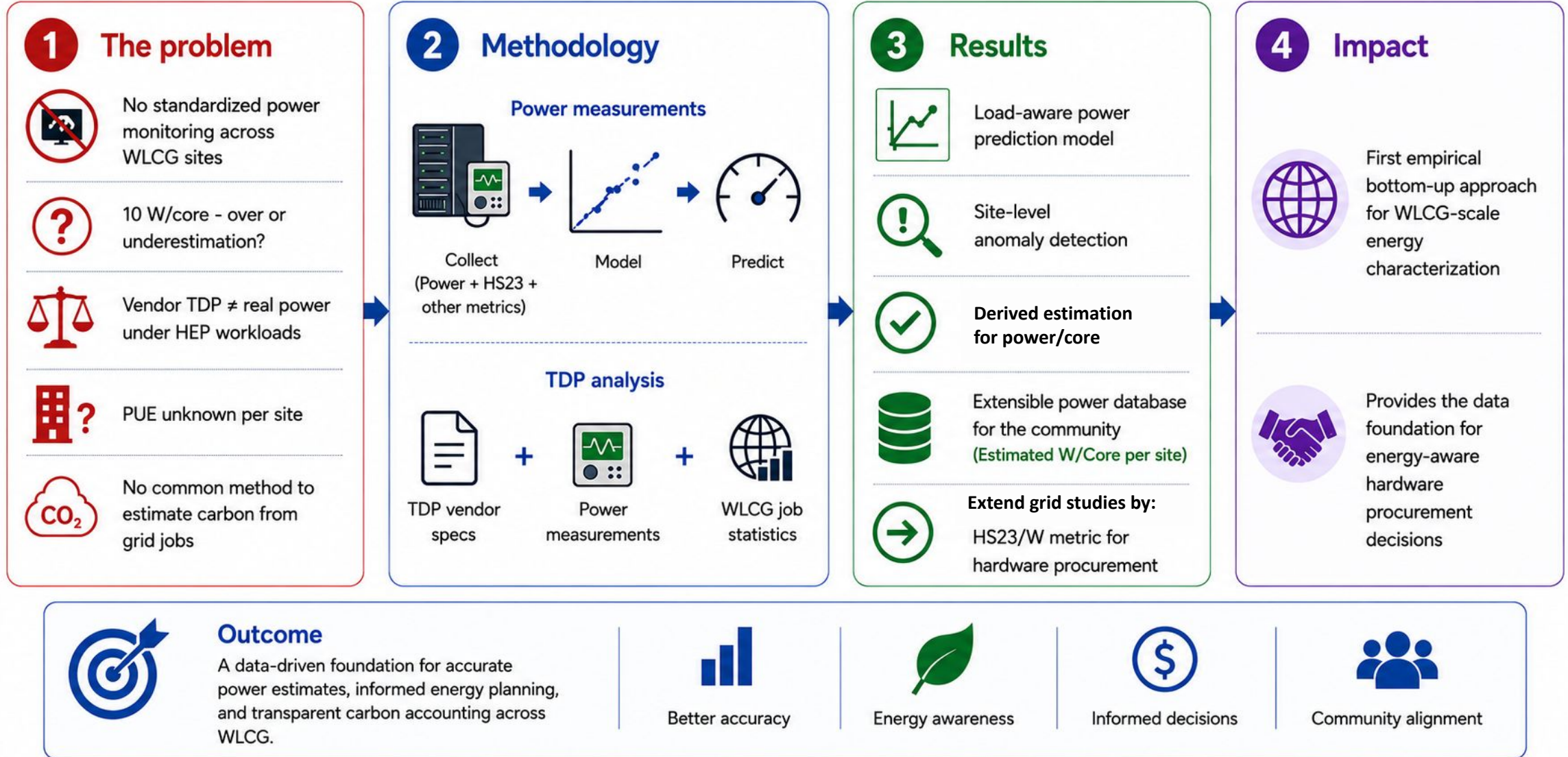
Sustainability plan

Cut carbon footprint across software, computing models, facilities & hardware lifecycle.

Where we are:

- Longer-term – defining scope for a global collaboration
- Early work: heat re-use, ARM/GPU tech, optimal HW lifetime

Guided by the four 'M's – **Measure, Model, Monitor, Moderate**.
What can we measure and how? Model the system to understand optimizations.
Develop monitoring of performance. Moderate usage where possible.



Power Measurements

HEP Benchmark Suite



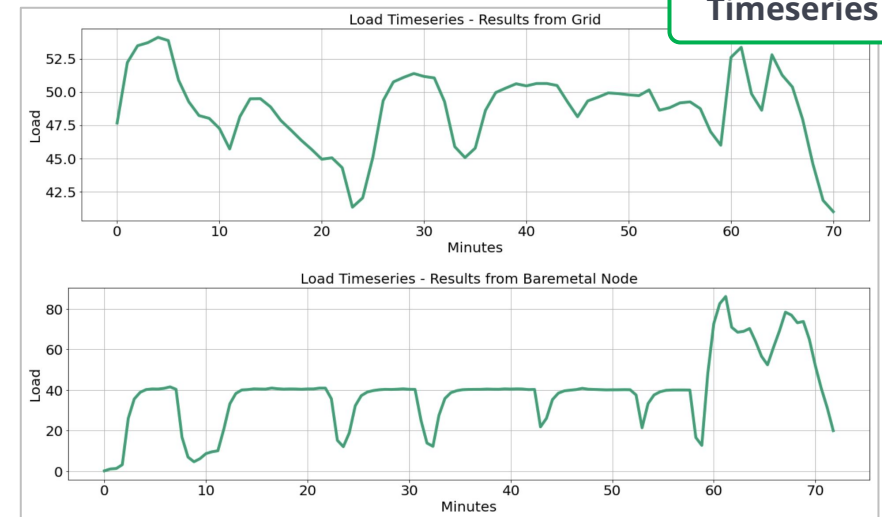
[hep-benchmark-suite](https://hep-benchmark-suite.github.io)

- The Suite is an orchestrator of multiple benchmarks. The suite incorporates metrics such as machine load, memory usage, memory swap, and notably, power consumption*
- Suite Plugins:
 - Run alongside benchmarks
 - Flexible modification and addition of collected metrics
- Only requirement: The command needs to be executed by the suite user

Configuration

```
1 plugins:
2   CommandExecutor:
3     metrics:
4       cpu-frequency:
5         command: >-
6           cpupower -c all frequency-info -f | grep 'current CPU frequency:' |
7           grep -o '[0-9]\{7,\}' | awk '{s+=$1; c++;} END {print (s/c)/1000}'
8         interval_mins: 1
9         regex: (?P<value>\d+\.\d+).*
10        unit: MHz
11      load:
12        command: uptime
13        interval_mins: 1
14        regex: 'load average: (?P<value>\d+\.\d+),'
15        unit: ''
16      power-consumption:
17        command: ipmitool dcmi power reading
18        description: >-
19          Retrieves power consumption of the system. Requires elevated
20          privileges.
21        interval_mins: 1
22        regex: 'Instantaneous power reading:\s*(?P<value>\d+) Watts'
23        unit: W
```

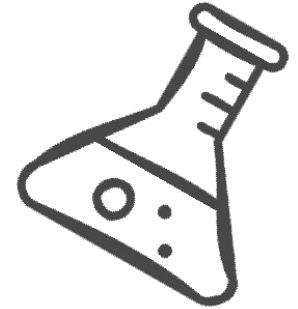
Timeseries data



Measuring Performance Under Diverse Conditions

- **Previous examples mainly covered isolated, controlled conditions:**
the **pledge** use case

- Server executes **exclusively** the benchmark application
- Ensure predictable behaviour and consistent measurement conditions



- **Real-world conditions: *the **batch job** use case***

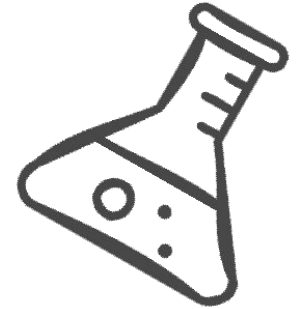
- Thanks to the **free license**, the benchmark can be executed over the Grid
- Running with **limited cores via Grid job submission**, while the rest of the server runs diverse, unpredictable **production** applications
- Allows sampling of virtually all grid nodes under **practical workload scenarios**



Measuring Performance Under Diverse Conditions

- **Previous examples mainly covered isolated, controlled conditions:**
the **pledge** use case

- Server executes **exclusively** the benchmark application
- Ensure predictable behaviour and consistent measurement conditions

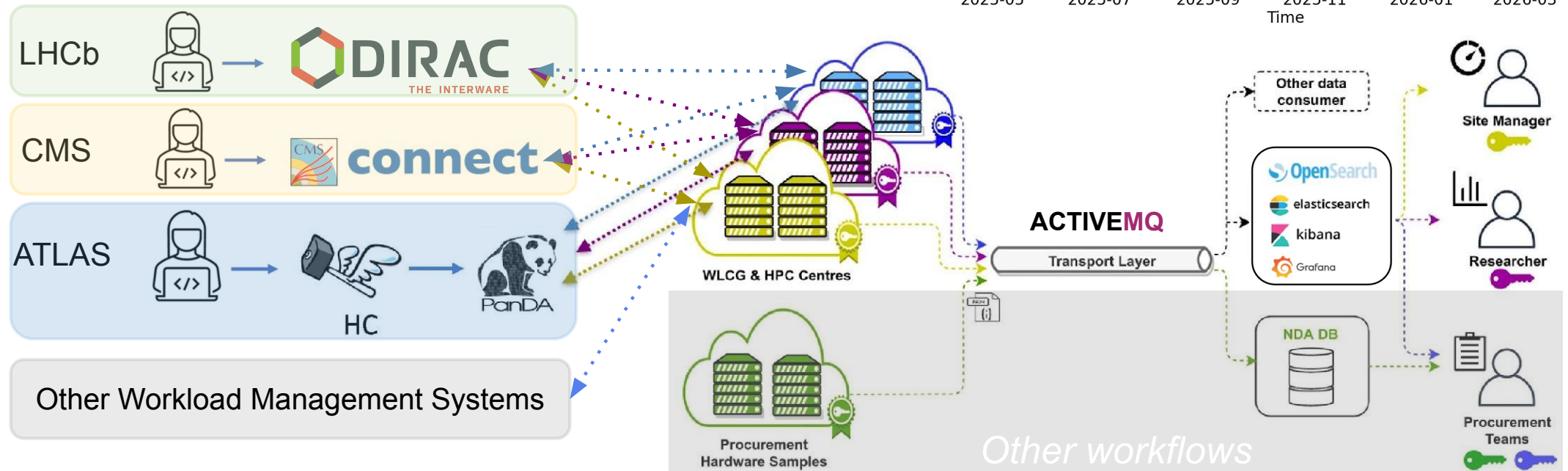
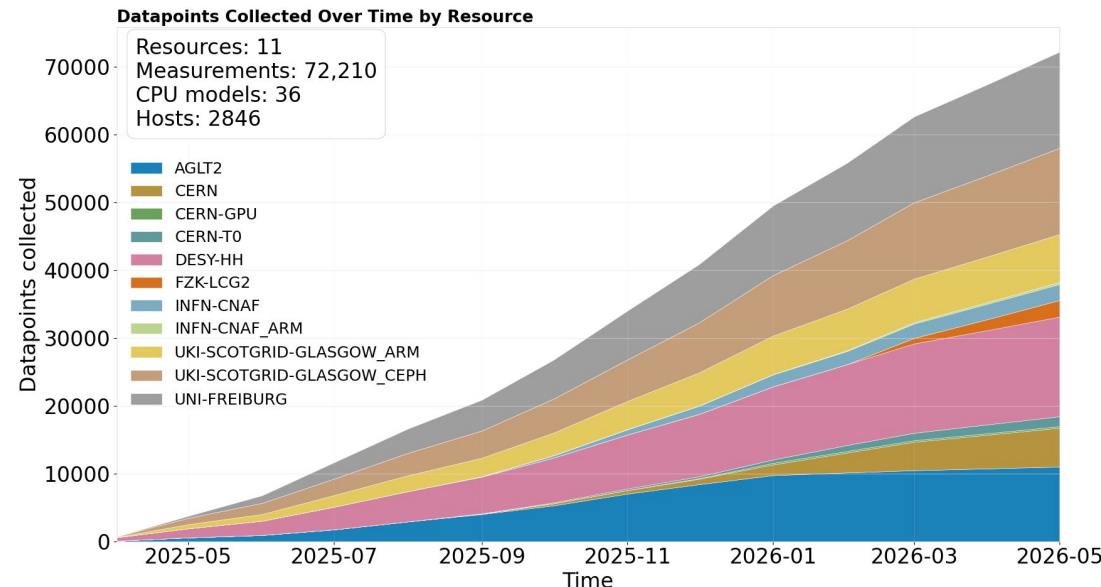


- **Real-world conditions: *the **batch job** use case***

- Thanks to the **free license**, the benchmark can be executed over the Grid
- Running with **limited cores via Grid job submission**, while the rest of the server runs diverse, unpredictable **production** applications
- Allows sampling of virtually all grid nodes under **practical workload scenarios**

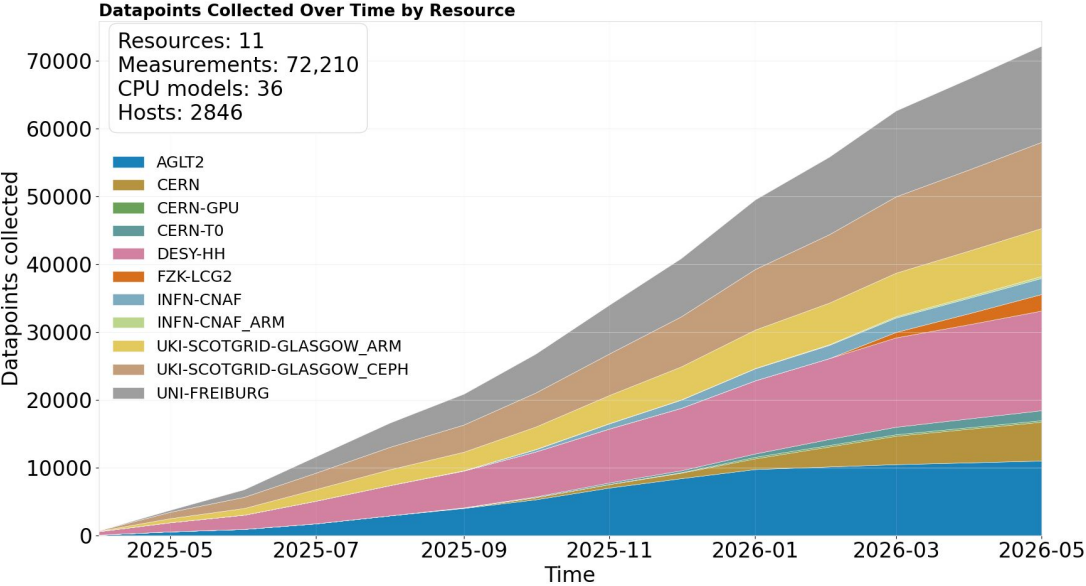
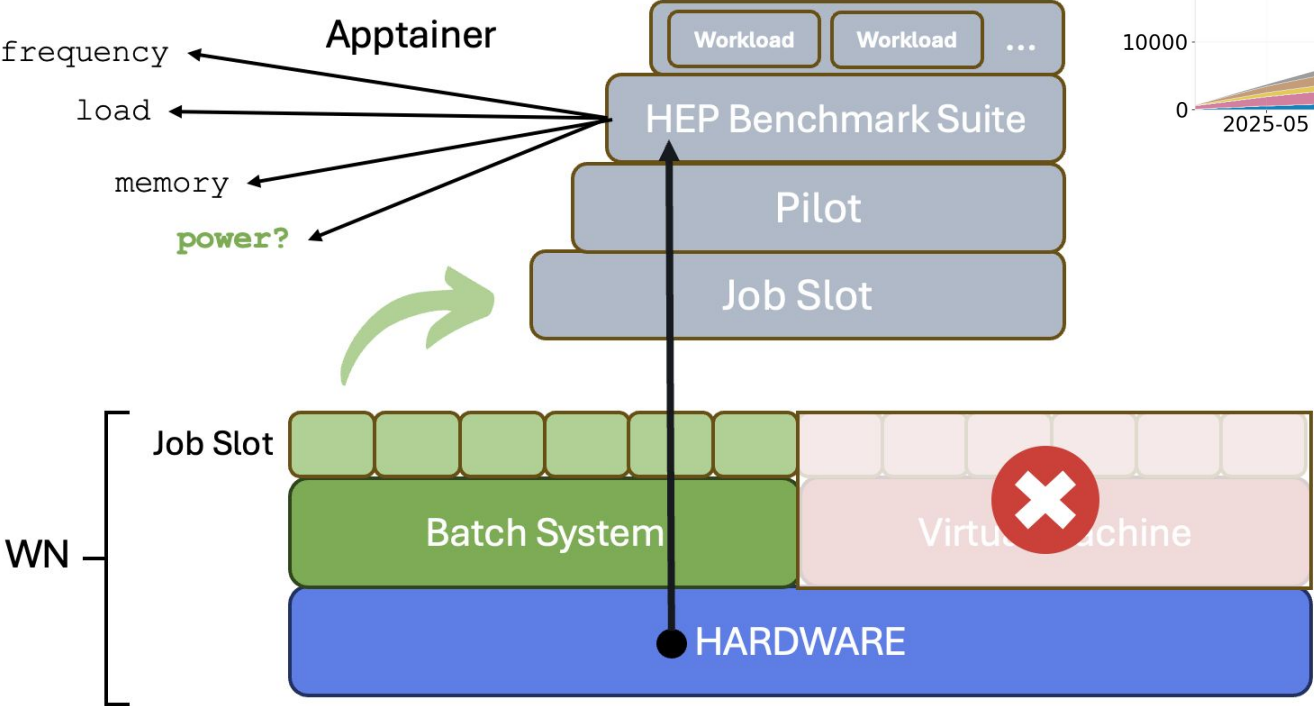


Leverage existing submission infrastructure



Source: <https://link.springer.com/article/10.1007/s41781-021-00074-y/figures/6>

Leverage existing submission infrastructure



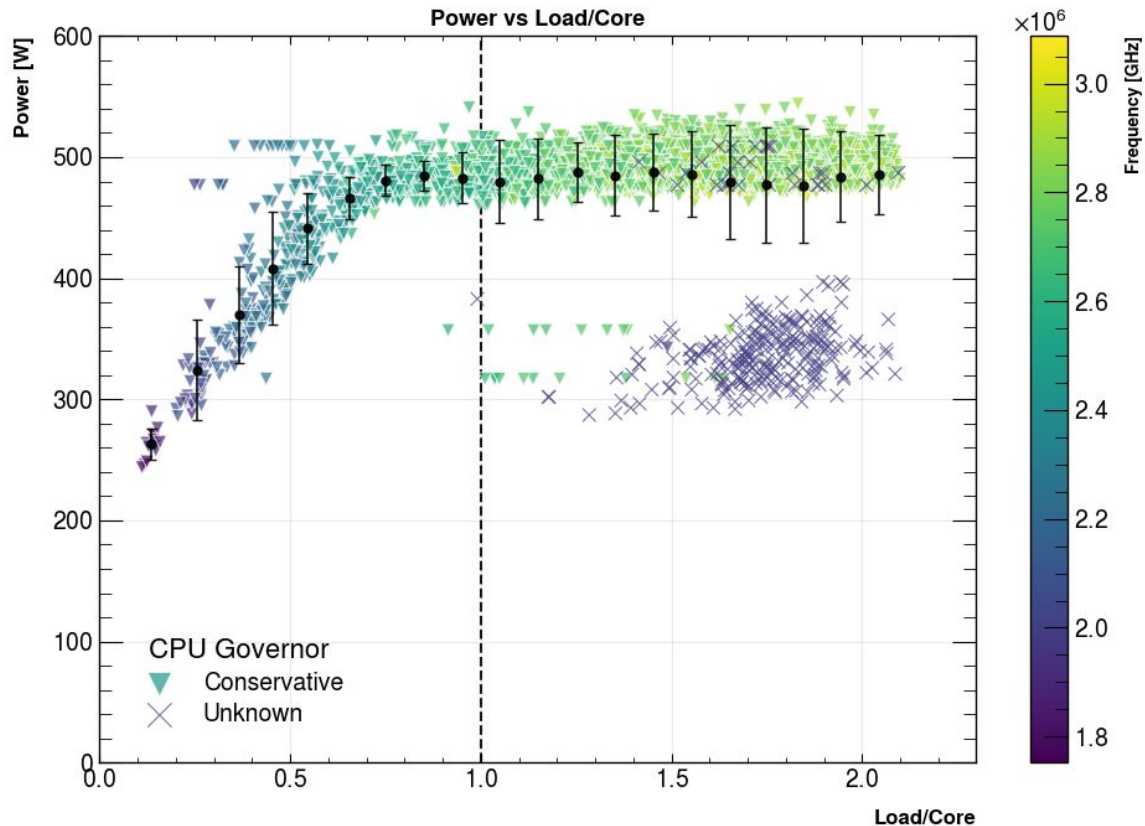
**Plugins run
alongside the
benchmark and
collect information
about the whole
server**



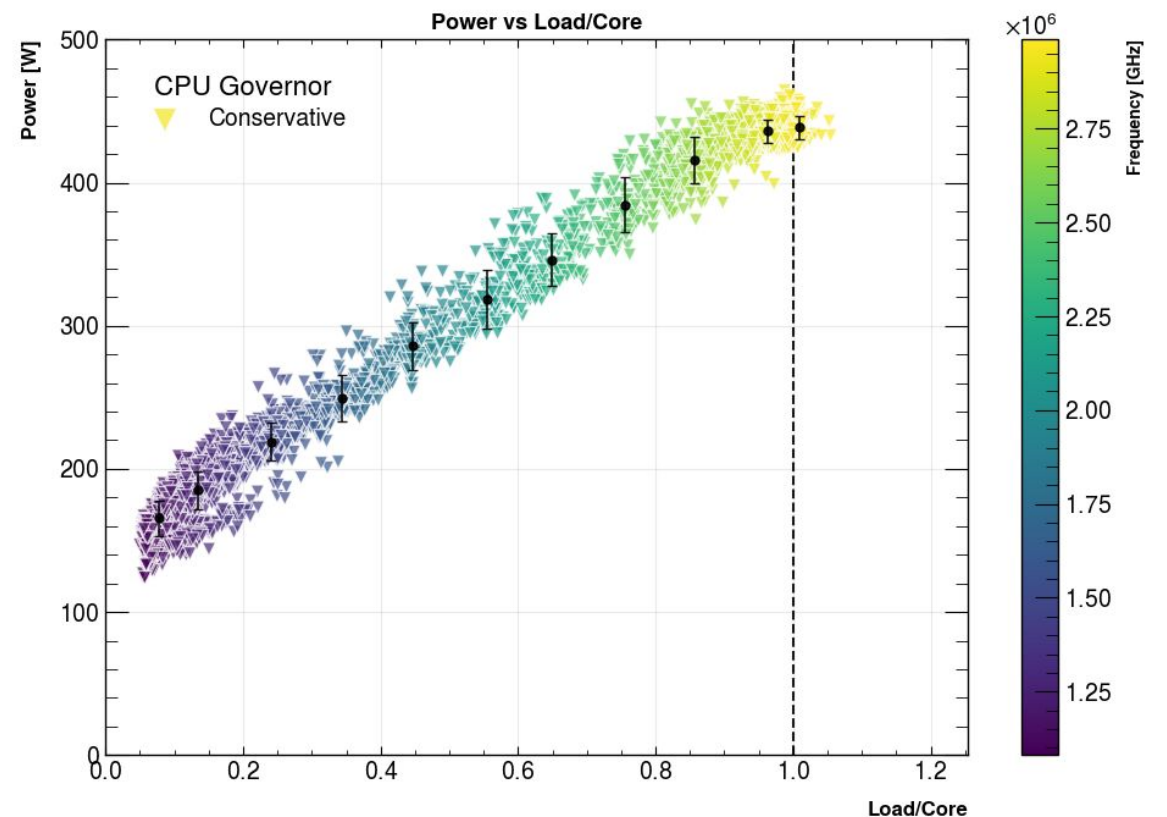
<https://gitlab.cern.ch/hep-benchmarks/hep-power-collector>

Collected measurements from Grid nodes

AMD EPYC 7513 32-Core



Neoverse-N1



Power shows a saturating trend with load ~ 1 (x86 - HT enabled) in load range (0-2)
We need to differentiate noise and outliers from deliberate server configurations

Looking ahead: Dynamic Power Prediction

- We intend to provide the community with a dynamic power prediction tool
- **Proof of concept**, to be fully developed
- The goal is to collect as much reliable fitted models as possible to cover WLCG needs
- Given a CPU model and load, the tool can predict idle, current and full-load power without direct measurements
 - An accurate validation on a test dataset is needed



TDP and Power Analysis

* TDP - Thermal Design Power

Strategy

1 Collect CPU TDP data

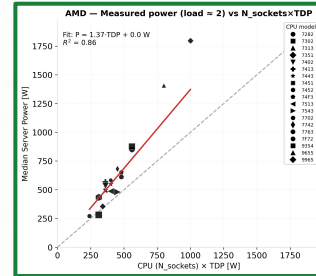
Collect CPU TDP data from official vendor specifications for the broadest possible coverage of deployed CPU models.



Model	Architecture	TDP (W)
Xeon 8480+	Sapphire Rapids	350
EPYC 9654	Genoa	360
EPYC 9554	Genoa	360
...	...	...

2 Derive scaling factors

Use real power measurements from previous step to derive architecture-specific scaling factors that translate TDP into estimated server power.



3 Determine usage weights

Use 2025 production job statistics, based on walltime x cores, to determine the relative usage weight of each CPU model at every site.

Relative usage by CPU model (at a site)



CPU model	Weight
Xeon 8480+	45%
EPYC 9654	30%
EPYC 9554	15%
Others	10%

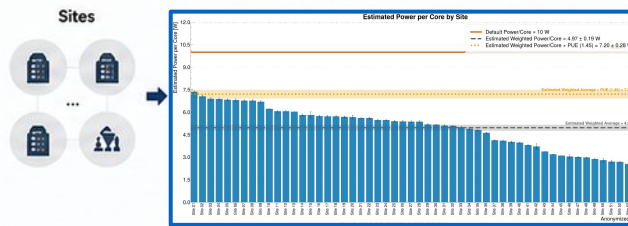
4 Compute weighted power/core per site

Compute the weighted estimated power per core for each site using the derived scaling factors and site-specific CPU fleet composition.

Site CPU mix (weights)	Scaling factor (Server Power / TDP)	Estimated power/core (W/core)
Xeon 8480+ 45%	Xeon (SPR) 1.45	X
EPYC 9654 30%	EPYC 9654 1.40	
EPYC 9554 15%	EPYC 9554 1.38	
Others 10%	Others ...	

5 Derive WLCG-wide power/core and compare

Aggregate weighted site-level values to obtain a WLCG-wide power/core estimate, then apply an assumed average PUE to compare against the current default assumption of 10 W/core.



6 Develop an automated calculator

Develop an automated calculator to estimate site power consumption from fleet composition.

Site Power Calculator		
Fleet composition (weights)		Results
Xeon 8480+	45 %	Power / core (IT)
EPYC 9654	30 %	6.7 W/core
EPYC 9554	15 %	Power / core (facility)
Others	10 %	10.6 W/core
PUE (assumed)	1.58	Total site power
		2.3 MW



Outcome

Data-driven, site-specific and WLCG-wide power estimates for informed energy planning and realistic comparisons.



Better accuracy



Energy awareness



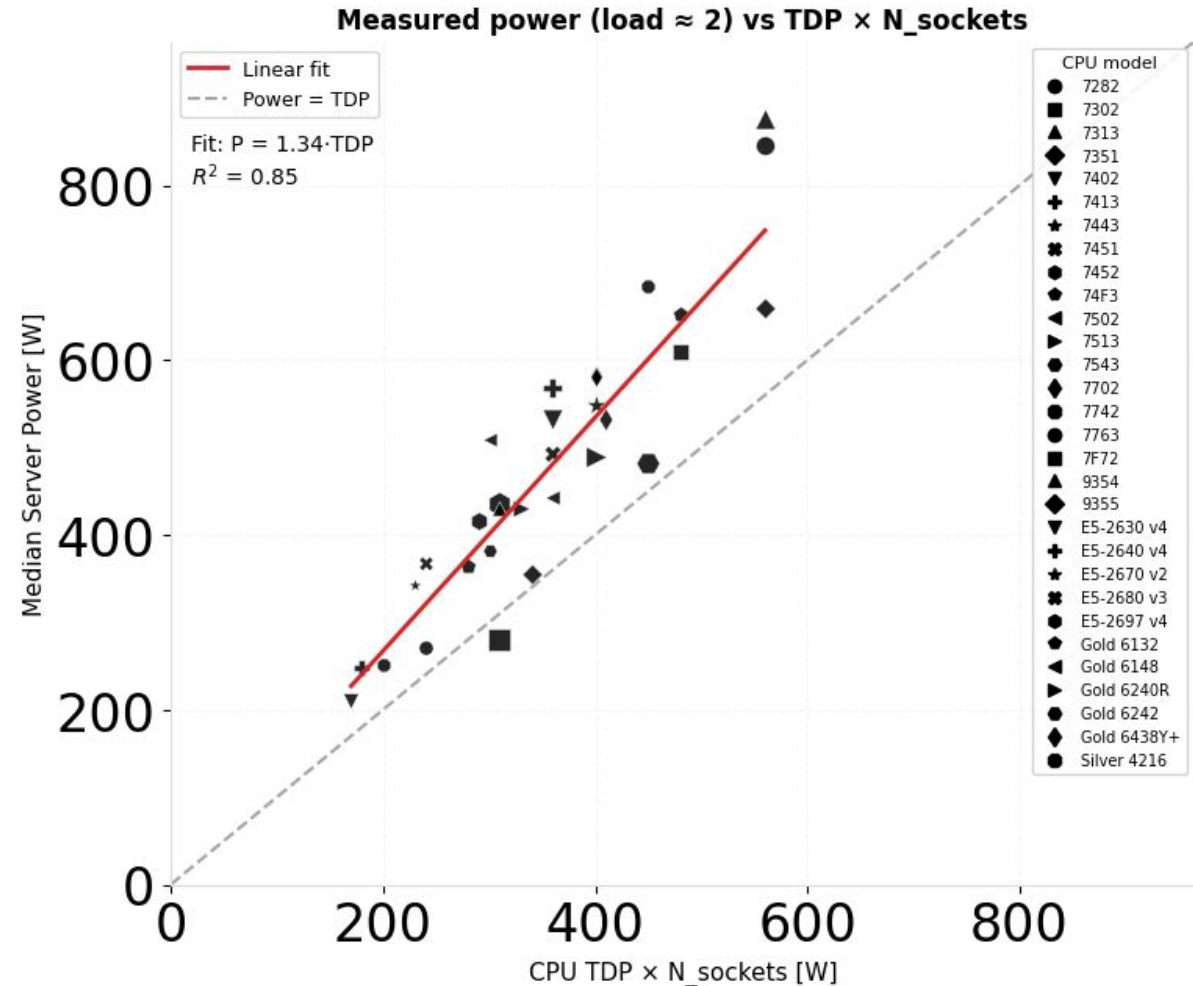
Informed decisions



Community alignment

CPU TDP vs server IPMI

- TDP stands for Thermal Design Power, in watts, and refers to the power consumption under the maximum theoretical load.
- TDP applies only to the chip itself (CPU or GPU), not the whole system.
- Measured power at fully loaded machines vs TDP shows linear relation that could be used if power measurements are not possible



CPU TDP dataset

- Collected 212 CPU model specs from official data sources
- Tables with CPU and GPU info, visualizations, API; **very useful for analysis**
- PoC, plan to be extended in the future



<https://computespecsdb.com/>

Compute Specs DB

HPC and Datacenter Hardware Specifications

Search by model, family, codename...

Search

Access JSON API:

All CPUs

All GPUs

Visualizations

API Info

API Docs

212

Total CPUs

11

Families

28

Codenames

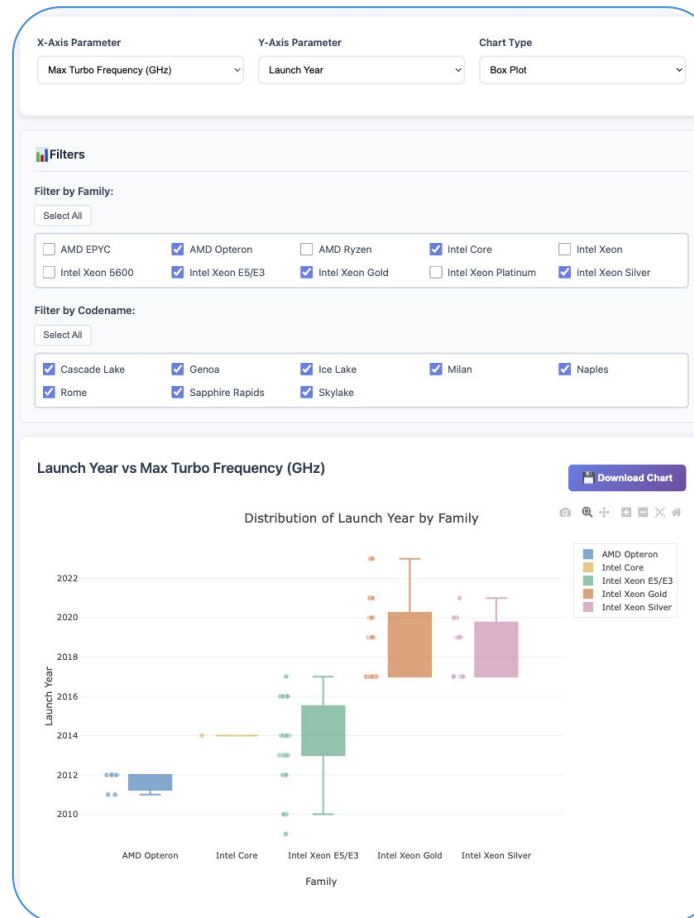
192

Max Cores

2009-2024

Year Range

CPUs											GPUs											
ID	%	CPU MODEL NAME	%	CODENAME	%	FAMILY	%	CPU MODEL	%	CORES	%	THREADS	%	TURBO GHZ	%	L3 MB	%	TDP (W)	%	YEAR	%	MEM
1		AMD EPYC 4564P 16-Core Processor		Genoa		AMD EPYC		EPYC 4564P		16		32		5.7		64		170		2024		1
2		AMD EPYC 7262 8-Core Processor		Rome		AMD EPYC		EPYC 7262		8		16		3.4		128		155		2019		4
3		AMD EPYC 7282 16-Core Processor		Rome		AMD EPYC		EPYC 7282		16		32		3.2		64		120		2019		4
4		AMD EPYC 7301 16-Core Processor		Naples		AMD EPYC		EPYC 7301		16		32		2.7		64		170		2017		4
5		AMD EPYC 7302 16-Core Processor		Rome		AMD EPYC		EPYC 7302		16		32		3.3		128		155		2019		4
6		AMD EPYC 7313 16-Core Processor		Milan		AMD EPYC		EPYC 7313		16		32		3.7		128		155		2021		4
7		AMD EPYC 7351 16-Core Processor		Naples		AMD EPYC		EPYC 7351		16		32		2.9		64		170		2017		2
8		AMD EPYC 7351P 16-Core Processor		Naples		AMD EPYC		EPYC 7351P		16		32		2.9		64		170		2017		2
9		AMD EPYC 7352 24-Core Processor		Rome		AMD EPYC		EPYC 7352		24		48		3.2		128		155		2019		4
10		AMD EPYC 7402 24-Core Processor		Rome		AMD EPYC		EPYC 7402		24		48		3.35		128		180		2019		4
11		AMD EPYC 7413 24-Core Processor		Milan		AMD EPYC		EPYC 7413		24		48		3.6		128		180		2021		4



Launch Year vs Max Turbo Frequency (GHz)

Download Chart

Distribution of Launch Year by Family

AMD Opteron

Intel Core

Intel Xeon E5/E3

Intel Xeon Gold

Intel Xeon Silver

Parameters

Cancel

Name

Description

q

* required

string

(query)

Search query (searches in model name, family, CPU model, and codename)

intel

Execute

Clear

Responses

Curl

curl -X "GET" \

'https://hpc-cpu-specification-api.onrender.com/api/cpus/search?q=intel' \

-H "accept: application/json"

Request URL

https://hpc-cpu-specification-api.onrender.com/api/cpus/search?q=intel

Server response

Code

Details

200

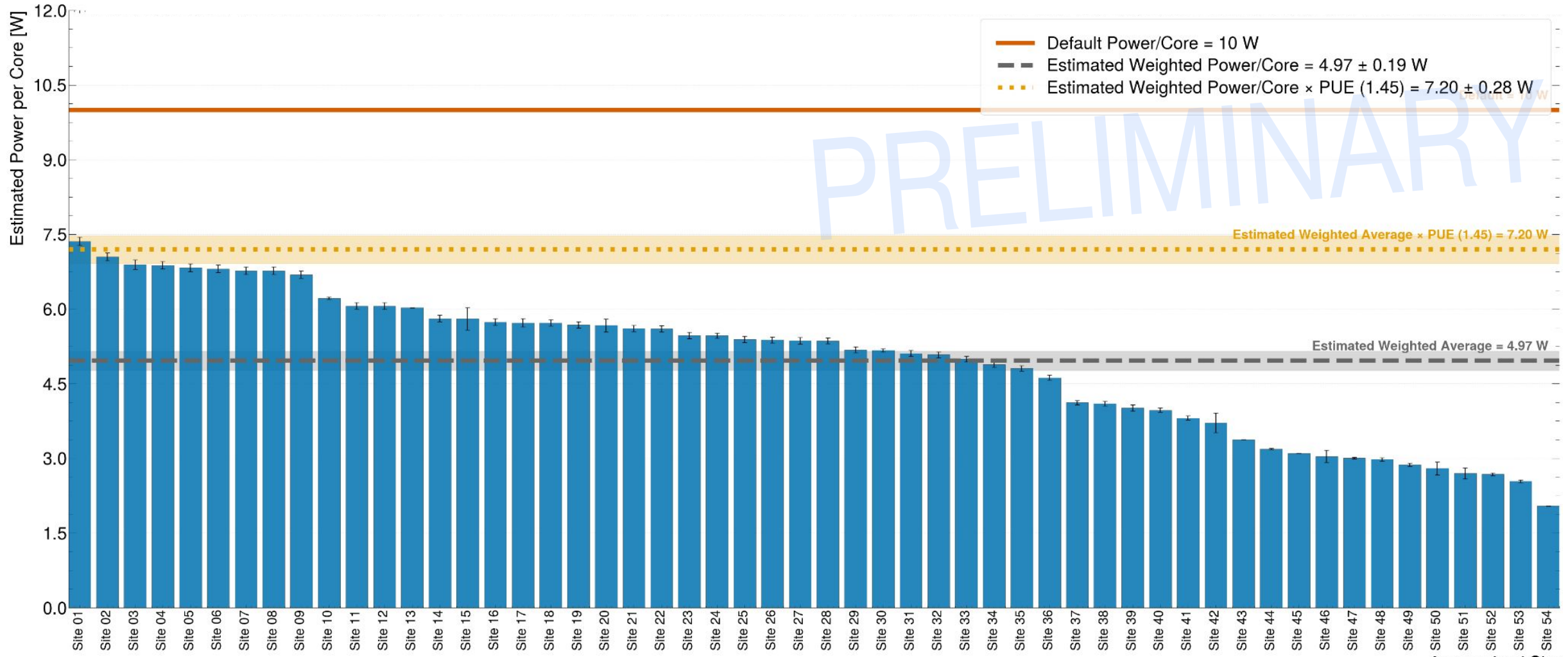
Response body

```
{
  "id": 53,
  "cpu_model_name": "Intel(R) Core(TM) i7-5960X CPU @ 3.00",
  "family": "Intel Core",
  "cpu_model": "i7-5960X",
  "codename": null,
  "cores": 8,
  "threads": 16,
  "max_turbo_frequency_ghz": 3.5,
  "l3_cache_mb": 20,
  "tdp_watts": 140,
  "launch_year": 2014,
  "max_memory_tb": 0.064
}
```

Response headers

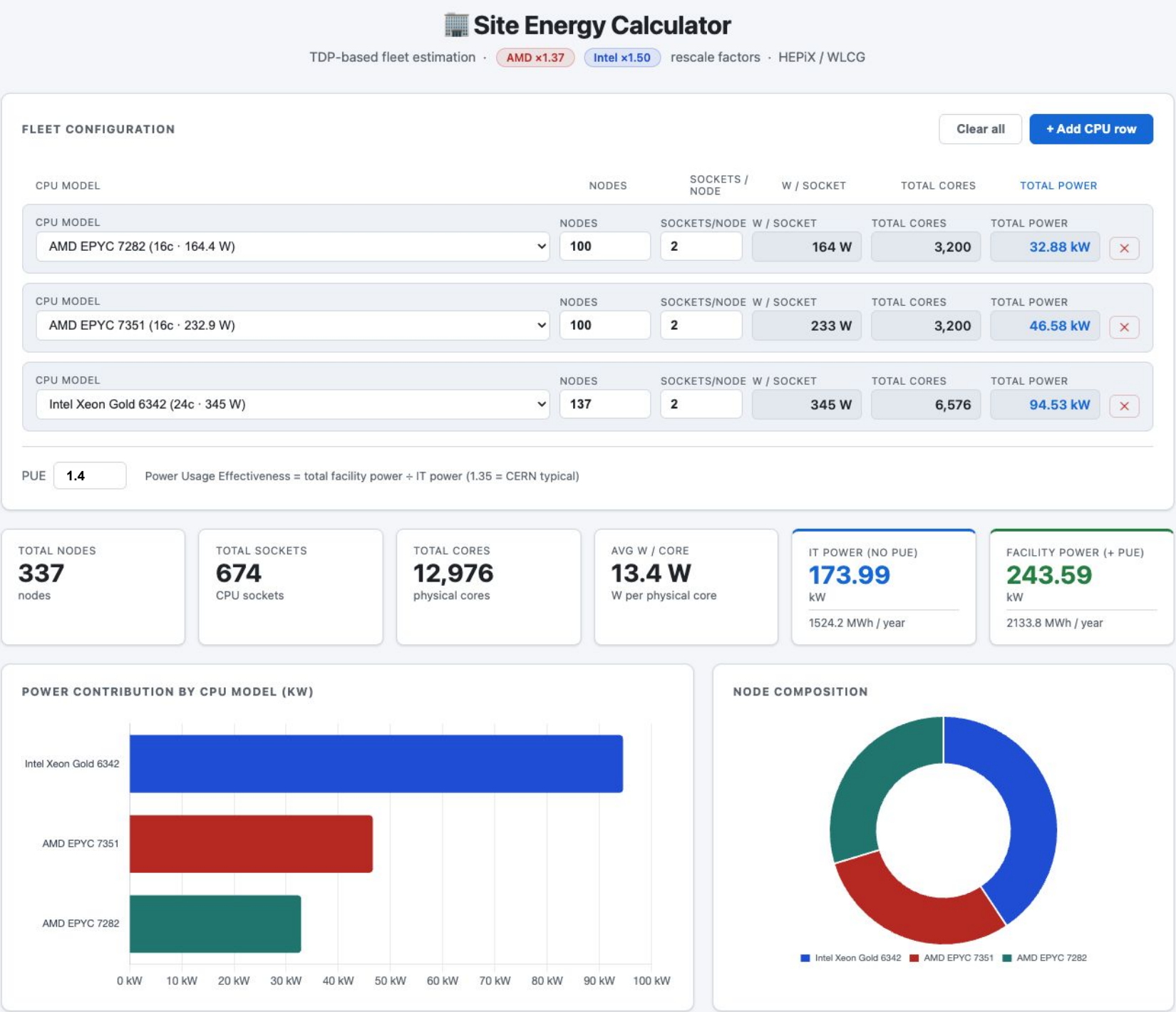
```
alt-svc: h3="443"; ma=86400
cf-cache-status: DYNAMIC
cf-ray: 9c844dc9e1e4-MRS
content-encoding: br
content-length: 2552
content-type: application/json
date: Tue, 03 Feb 2026 21:06:44 GMT
priority: u=1,i
rmdr-id: 20c3825-62b7-483f
server: cloudflare
```

What is the estimated power per core of some WLCG sites?



Empirical bottom-up approach to estimate power/core per site and WLCG Average:
This is a subset of WLCG Sites, using available data from ATLAS → will be extended to entire WLCG
Estimated Weighted Average Power/Core scaled by PUE (1.45) is 7.2W/core

Goal:
Develop a
"Site Energy
Calculator"



Conclusions

