



Job-specific performance monitoring on HPC clusters: Challenges and Solutions

Job-specific performance monitoring on HPC clusters: Challenges and Solutions

presented by
Christian Iwainsky

- RSE expert
- Researcher
- HPC specialist



TECHNISCHE
UNIVERSITÄT
DARMSTADT



NHR for
Computational
Engineering
Science

christian.iwainsky@tu-darmstadt.de



Hessisches Kompetenzzentrum
für Hochleistungsrechnen

brainware
for science



Hessisches Ministerium
für Wissenschaft und Kunst

Agenda

- Motivation:
 - Why performance monitoring is important ...
 - Compromising between conflicting parties and constraints
- Knowledge gap: admin view vs user view vs performance view
 - Did you know ...
- Tools that help:
 - Cluster Cockpit, Patho Jobs
- State of the art:
 - What else is out there
- Where should we go?
- Summary!
- Q&A

S.th. to look at while I talk ...



My background ...

- Research & Development:
 - Performance-analysis
 - Modelling
 - OpenMP/MPI/C++ tooling
 - Instrumentation Compiler
- Affinity to performance analysis tools:
 - Scalasca, Vampir, LIKWID, HPCToolkit
- Primary research question:
 - What's happening inside this job?
 - How do we make it faster?
 - Still active in research!



ORCID

Where I am now?

- Day-job:
 - HPC expert
 - Last-level HPC support
 - RSE engineer
 - Trainer
 - HPC skills & Softwareengineering-coach
- In-between users and system administrators
- Reality:
 - Screen 10.000 Jobs/day
 - ~ 70% outbound interactions
 - Extreme diversity in codes, users and issues
- **"Did you know that..." is how most user conversations start.**



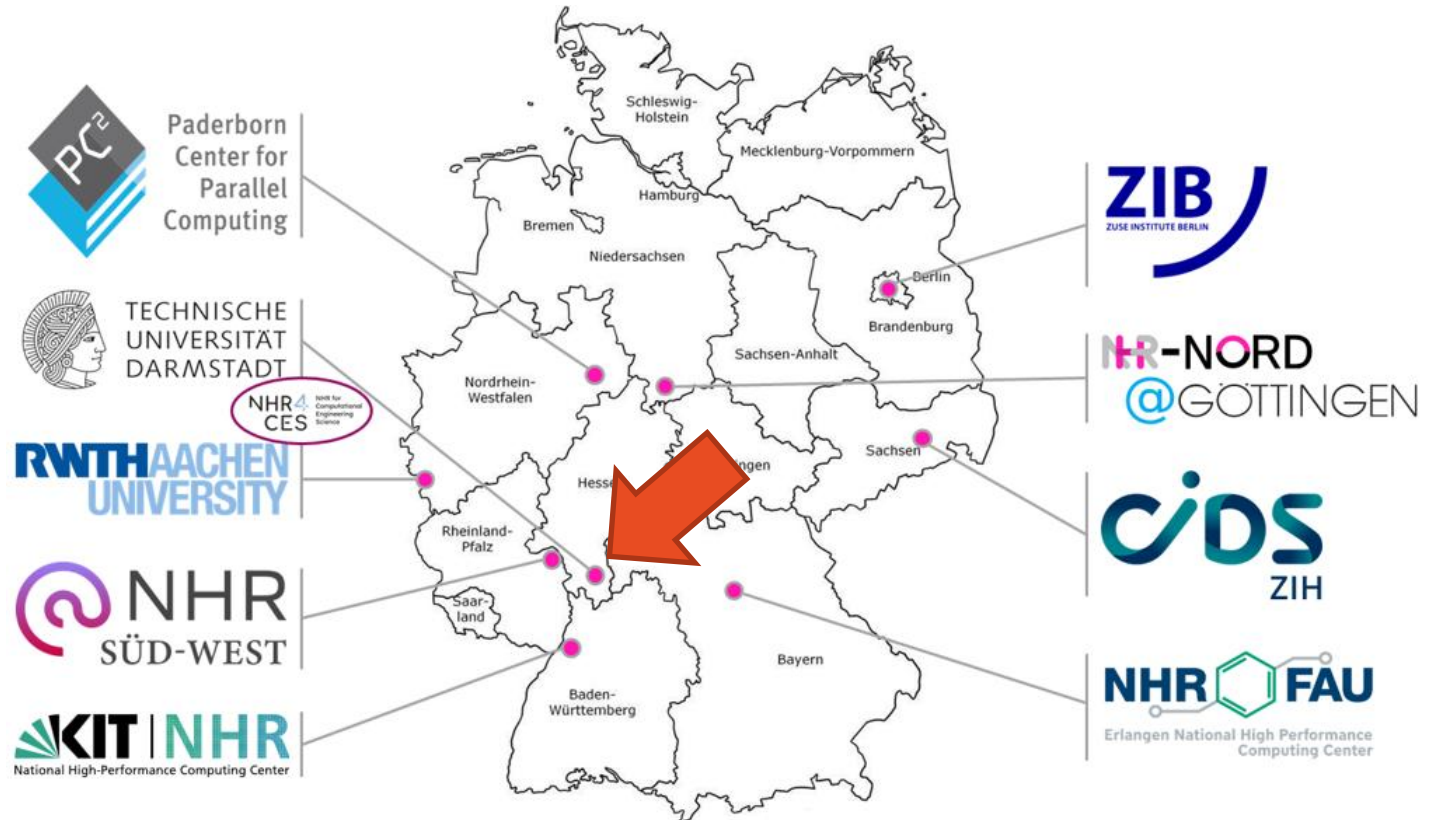
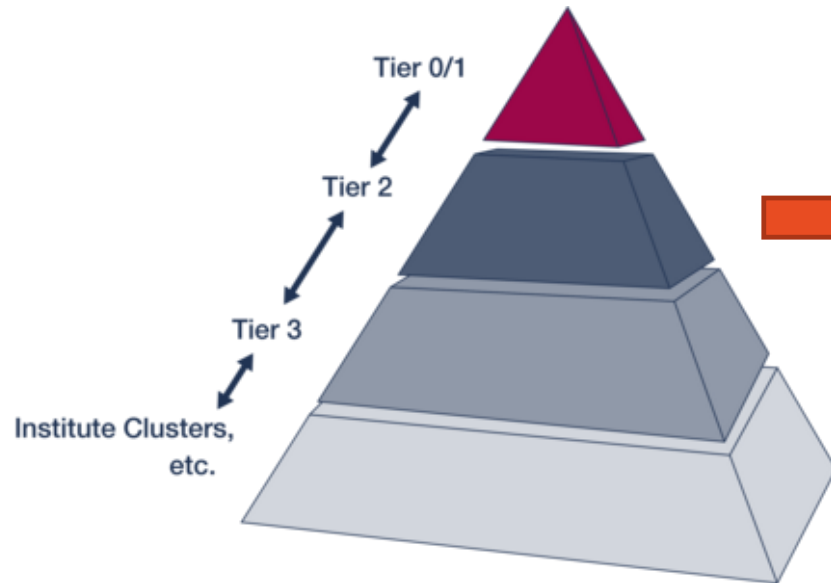
Background of TU Darmstadt

Tier-0 PRACE und EuroHPC

Tier-1 Gauss Center for Supercomputing

Tier-2 NHR

Tier-3 lokale HPC at Universities

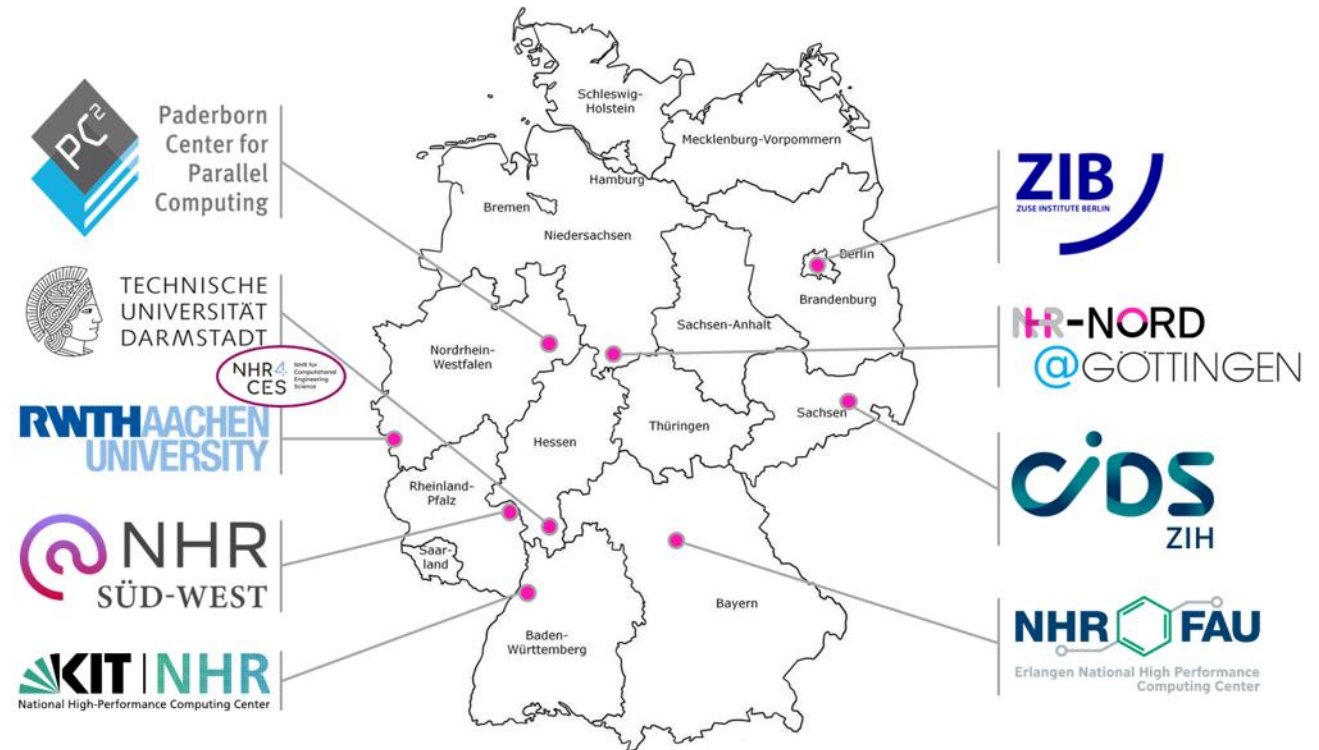


Background of TU Darmstadt

- HPC @ TU Darmstadt

One of 9 NHR centers in Germany

- O(1,000) compute nodes
- O(10,000) total cores @
- at least 4GB RAM/ core
- O(500) active users
 - students,
 - PHD students,
 - post-doc, researchers and
 - profs.
- 2 HPC (user) support staff
- ~ 10 HPC system admins



What I consider DevOps?

DevOps principle		What I think it means
Observability		Jobs and system behavior observable at runtime
Automation		Measurement and observation is automated
Feedback Loops		Information enables to incrementally improve HPC use
Scalability		Mechanisms scale to all HPC levels
Collaboration		Enables teamwork of admins, support staff and end-users
Tooling		Data and formats allow integration and application of more tools

HPC ain't free, it comes at a price

- Lichtenberg at TU Darmstadt
 - $O(10,000)$ cores (on average for a full system)
 - 15 Mio.€ (cluster incl. storage)
 - 5 years of operation (per installment)
 - 10 full-time staff
 - 5 years of power and cooling
 - Building and infrastructure



Can't share specific numbers

- ~ about 75 cents / core*h
- 72€ / node*h

HPC ain't free, it comes at a price

- Mock-up cluster:

- Node:

- CPU: 2x AMD EPYC 9965 (192 cores) ~ 8.000€ / piece
 - Motherboard: Dual Socket Motherboard ~ 1.600€ / piece
 - RAM: 24x 32GB DDR5-4800 (ECC) ~ 270€ / piece
 - Network: InfiniBand HDR100 ~ 850€ / piece
 - ~~GPGPU: NVIDIA H100 ~ 40.000€ / piece~~
 - Misc: Cables, chassis, infrastructure, storage ~ 5000€ / node

- Operational costs (1000 nodes):

- Staff: 7 FTE (full time employees) ~ 80k€ / (FTE * y)
 - Power: ~ 3.4 MW ~ 67,70 € / MWh

- Node/h ~ 0,97€ Node/y ~ 8562€
- Cluster/h ~ 977,44€ Cluster/y ~ 8.562 M€

What is the challenge?

Are users

- using all invested resources?
- fully using all invested resources?
- being stalled by some resource?

~~• Understand performance of a code with a dataset on a machine!~~

- Understand performance
 - of a **specific** code
 - with a **specific** dataset
 - on a **specific subset** of a machine!

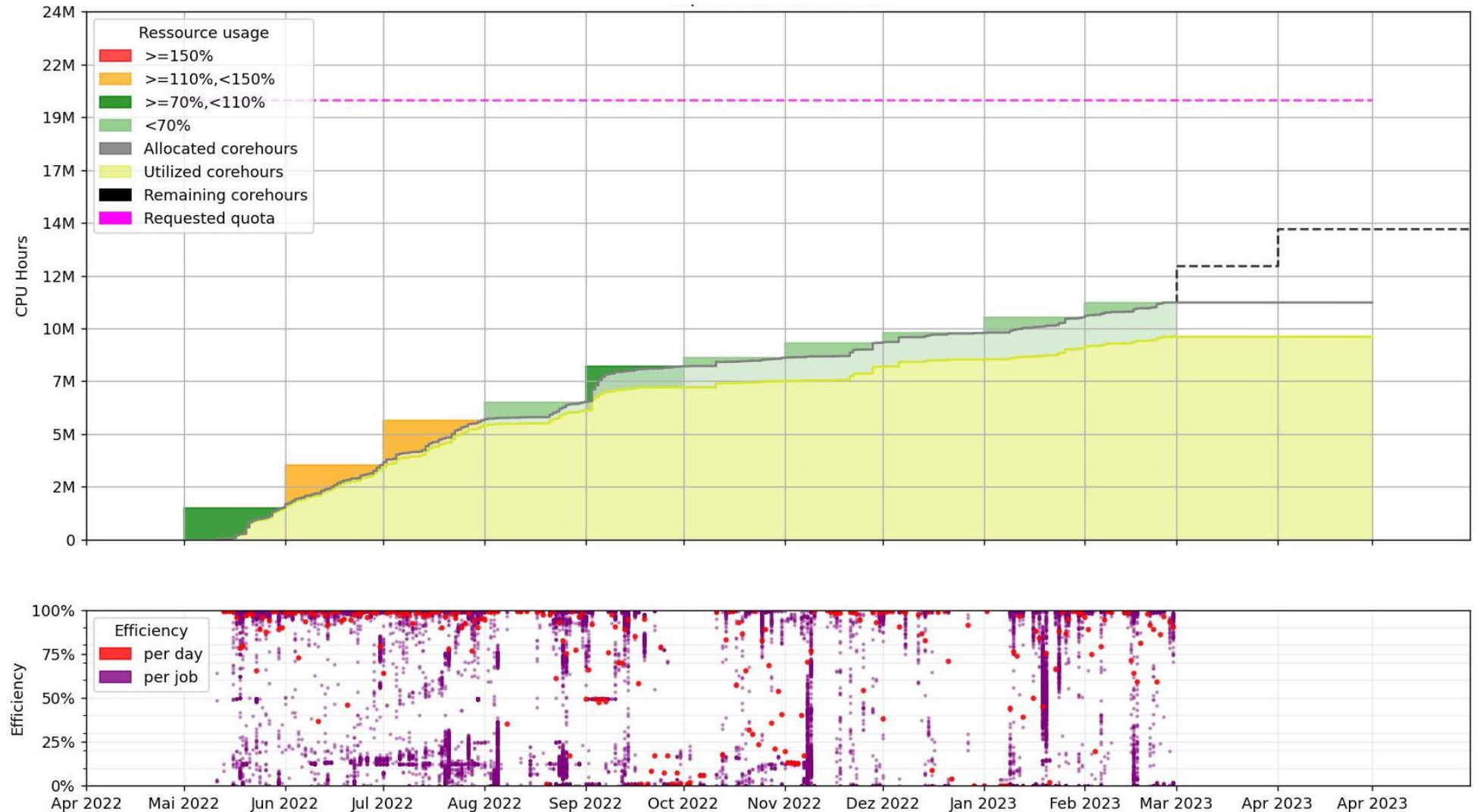
Trust ...

Know your performance ...

- You can only manage what you observe – and in HPC, our users are often flying blind.
- “Know the performance you get”
- $\text{Performance} = \text{Work} / \text{Time}$
- We’re not just measuring performance — we’re integrating it into the workflow, at scale, for everyone

Things to catch ...

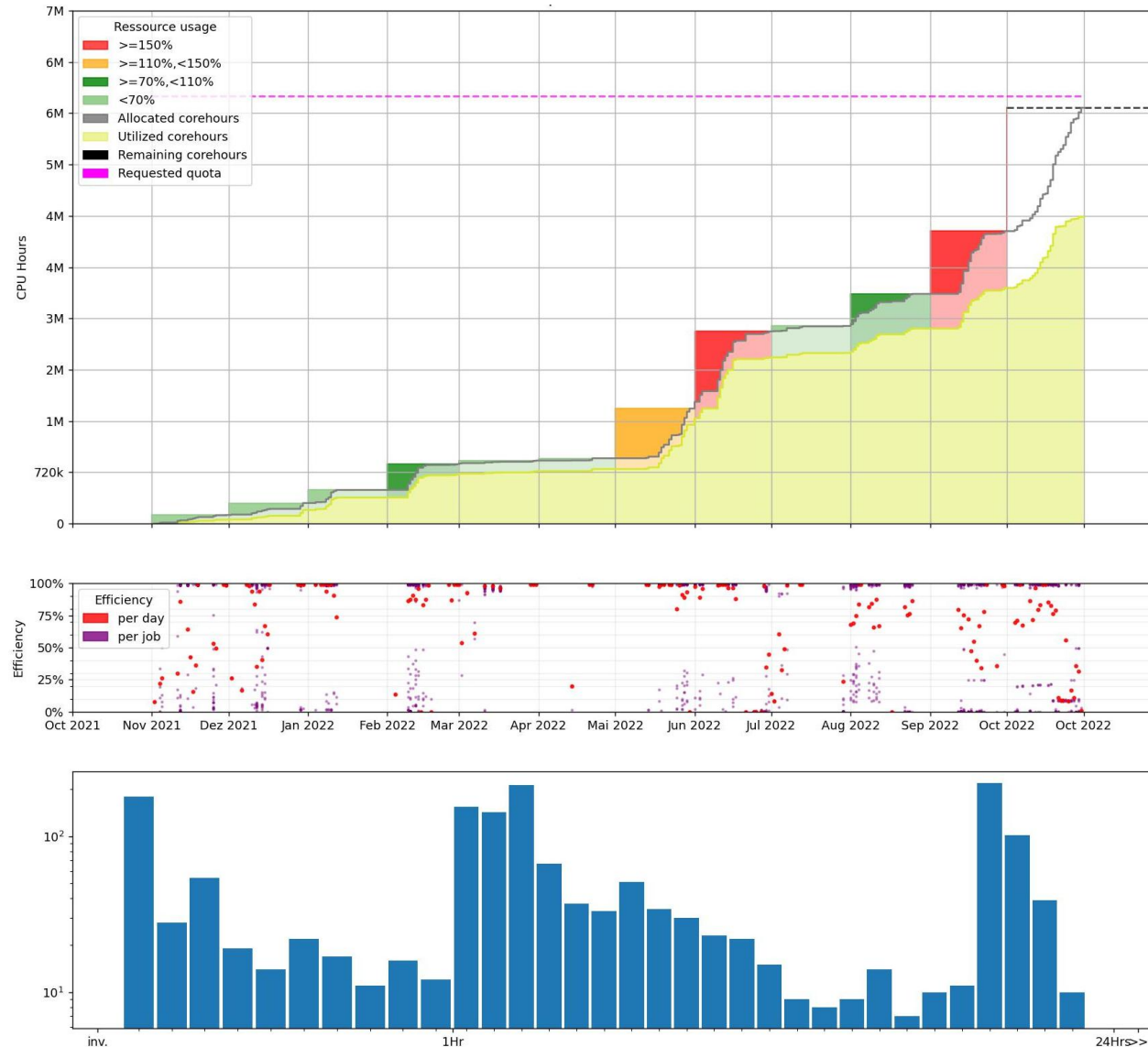
- Utilization
- Usage
- Efficiency



*Deprecated visualization

Things to catch ...

- Utilization
- Usage
- Efficiency

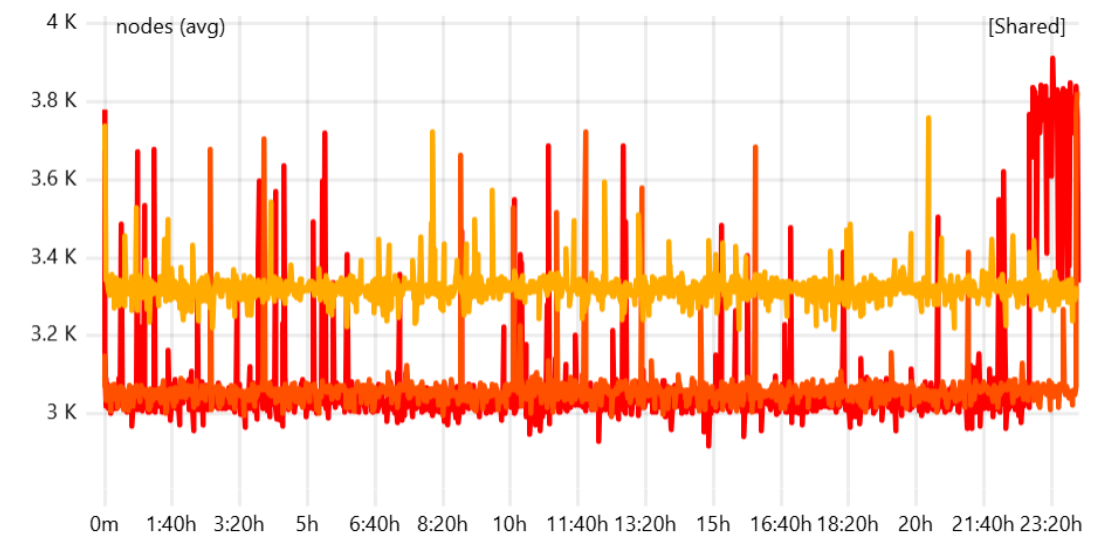
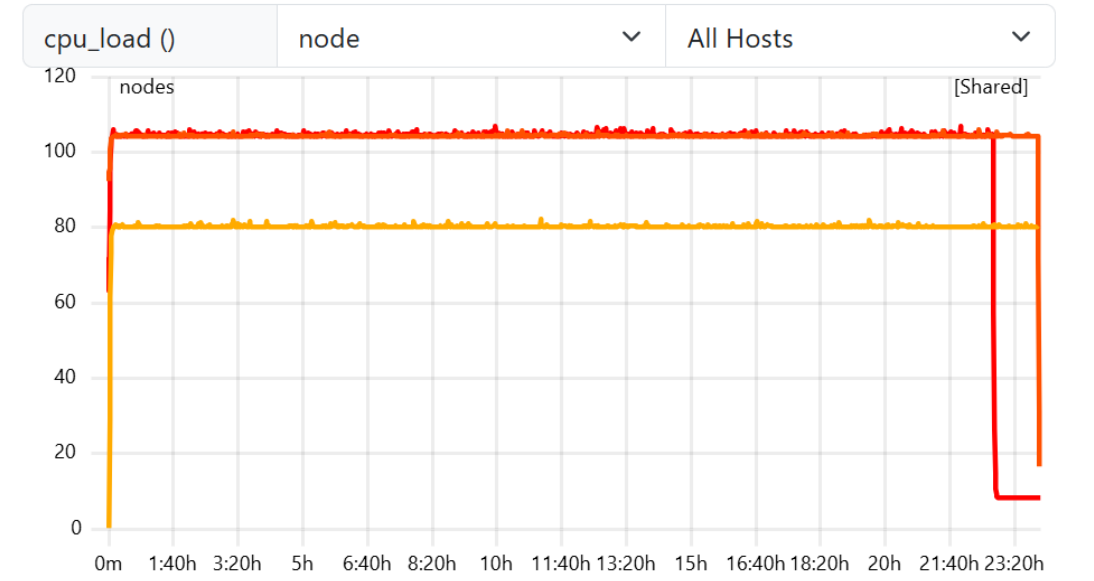


Example: unexpected runtime

System errors & User mistakes:

Job: Runtime unexpectedly slow

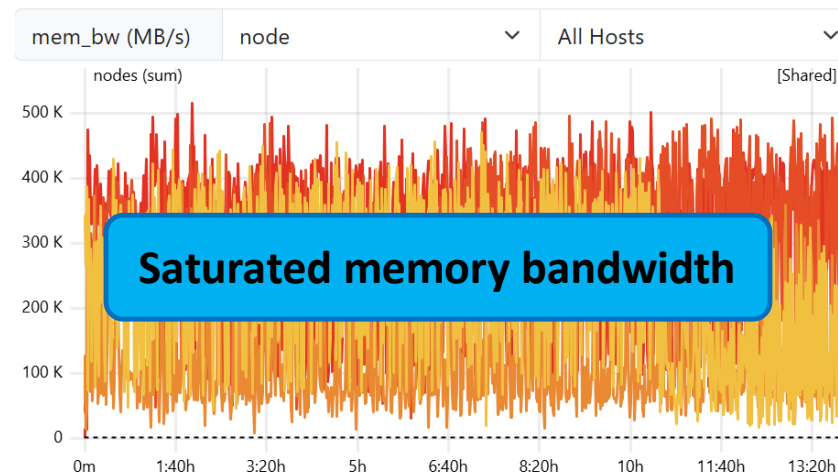
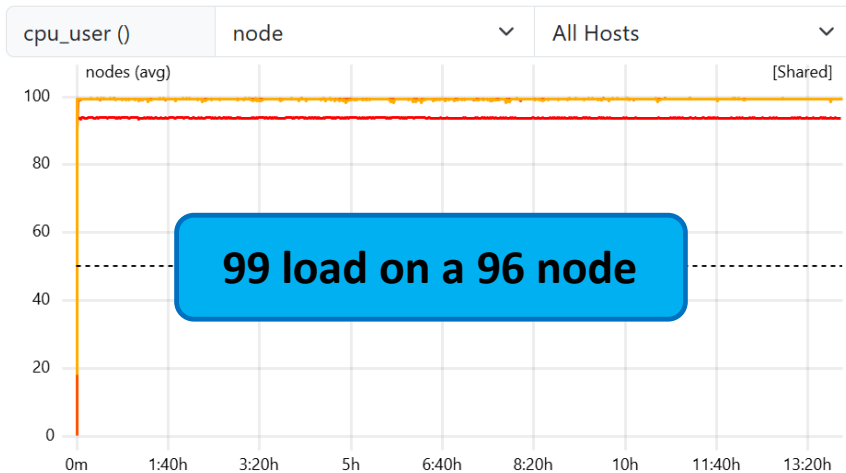
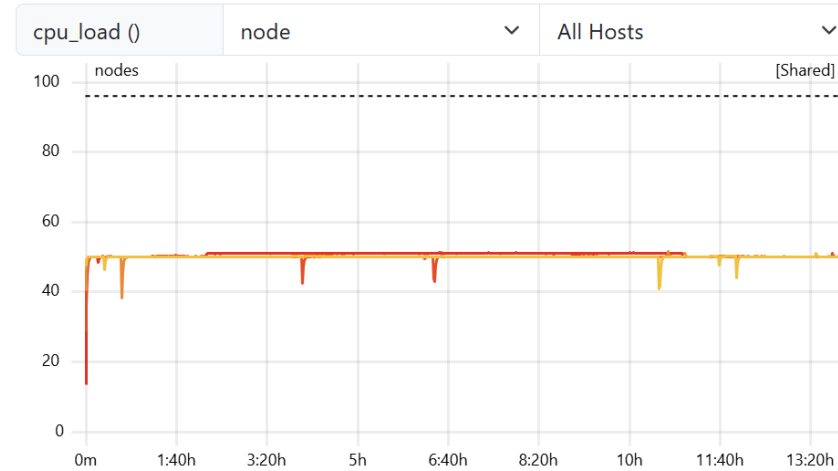
- 96 Processes
- 1 Thread / Core
- Mem-per-cpu=3000 MB
- Runtime use:
 - 3 nodes
 - 48 on node A
 - 24 on node B
 - 24 on node C



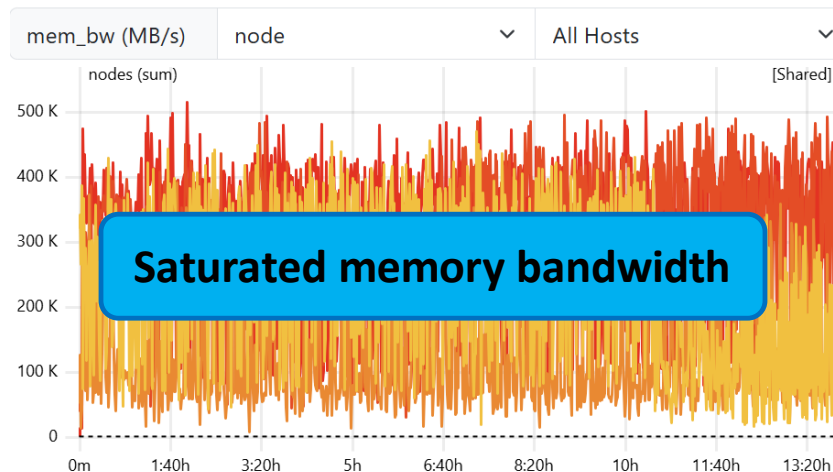
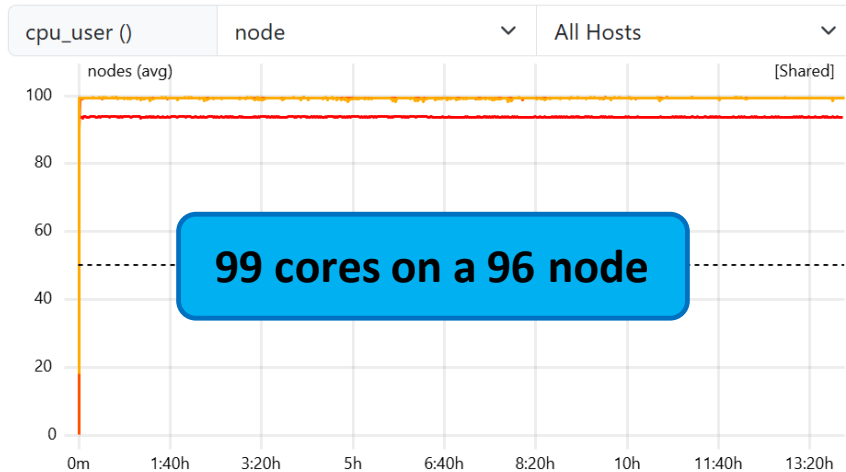
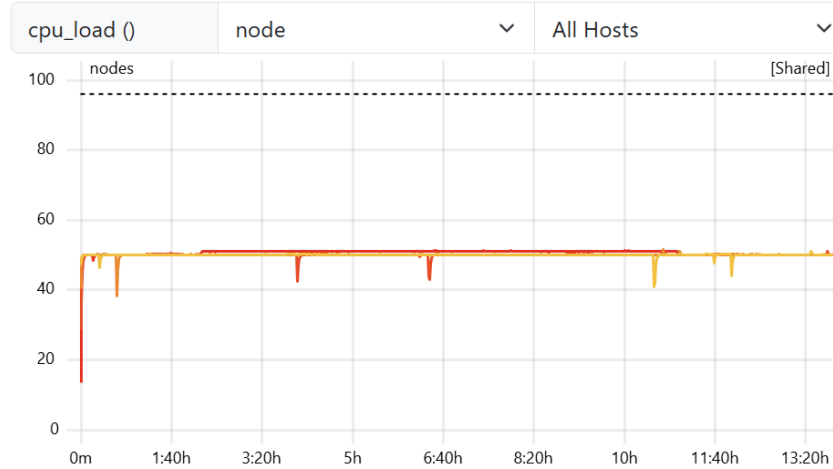
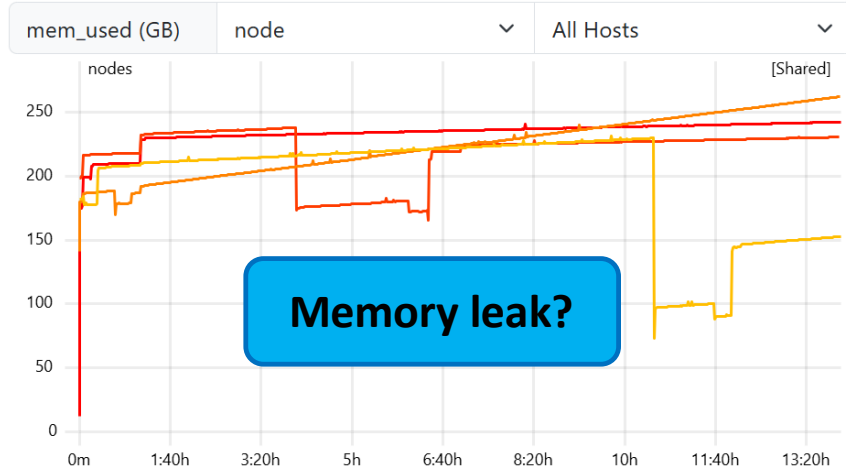
Example: low load

- 96 processes
 - 4 nodes
 - 7GB / rank
- ➔ 24 processes / node

How can job be improved?



Example: low load



Trust ...

But monitor!

What would a monitoring tool need!

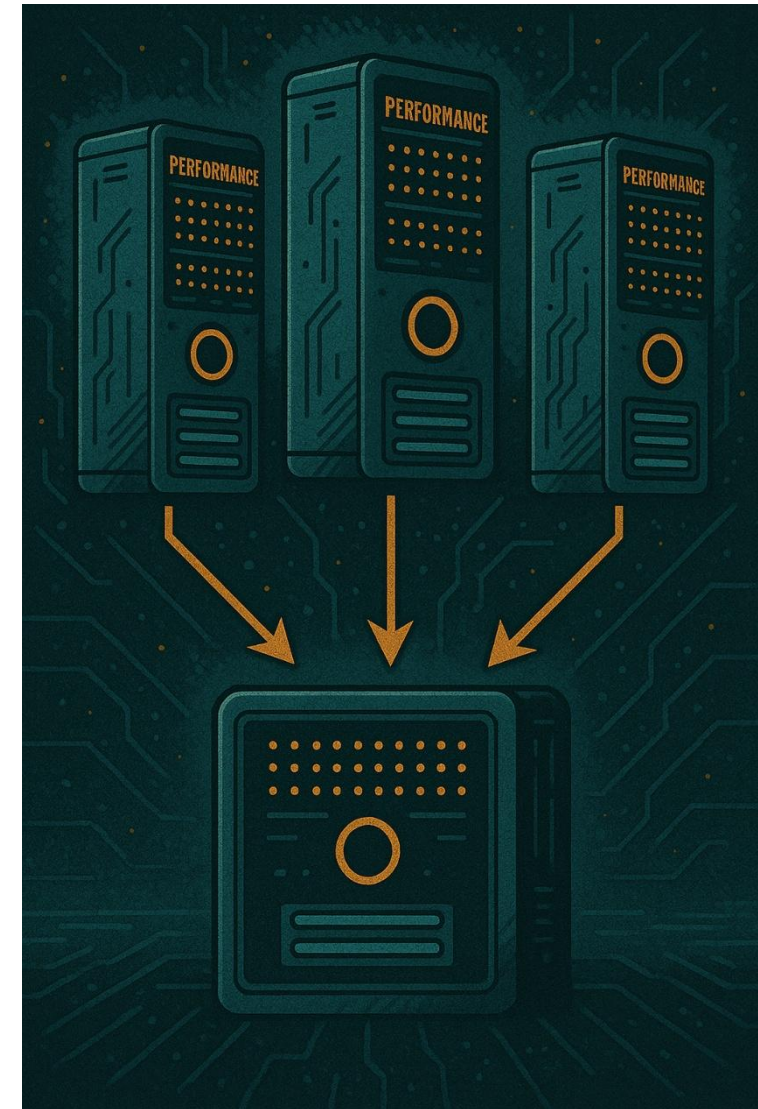
- Scheduling information:
 - Which processes ran where?
 - Which resources were used by a specific job?
 - Resources per parallel entity?
 - What aspects of the hardware was used?
- Hardware information:
 - Which cores were saturated?
 - What are the functional units doing?
 - Which network connection was full?
- Context information:
 - Which mathematical equations were being solved?
 - Which phase of the code was operational?
- Temporal information:
 - What was the previous state?

Background: System Level Monitoring

- Goal: track health-status of hardware
 - Zabbix
 - Prometheus
 - Grafana
 - Vendor specific tools
- Operation:
 - Node-level daemon gathers data locally, and sends
 - System-level daemon gathers data remotely
- Criteria:
 - Local overhead
- Data:
 - Operational information:
 - Load, hardware-failurs, energy consumption

How does it work?

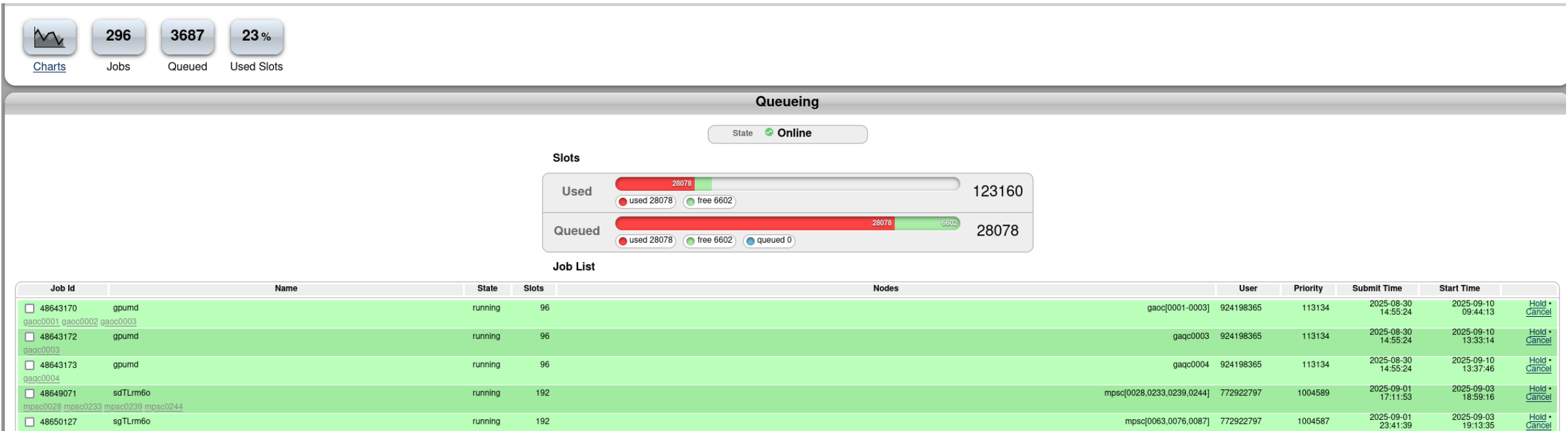
- Node agent gather all data
 - push: active send it to an aggregation agent
 - pull: data gathered from aggregation agent
- Aggregation agent / service
 - Processes data and vizualize
- Challenges:
 - Perturbation by the node agent
 - 1 ms / core @ 394 cores -> 0.4 s / node / sample
 - Network-traffik
 - $1000 \text{ nodes} * 394 \text{ cores} * 1\text{kb} = 394\text{MB} / \text{sample}$
 - Aggregation work (DDOS)
 - synchronous delivery: $O(1 \text{ million})$ packets per query



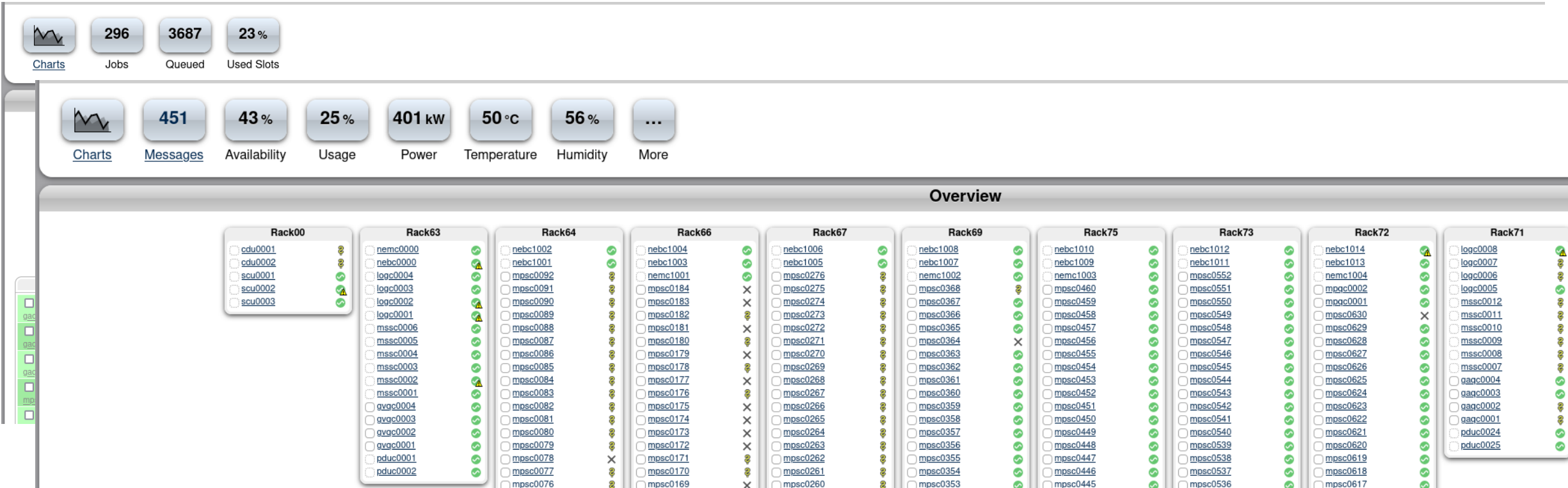
System Level Monitoring:



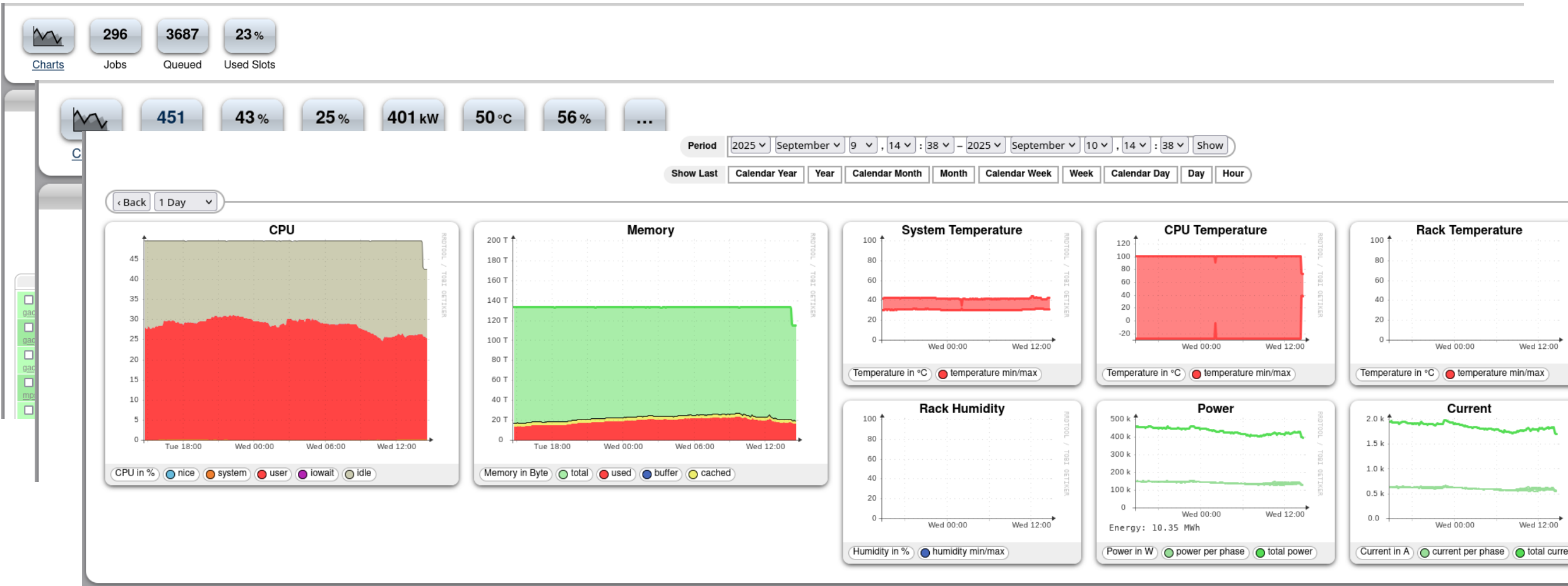
System Level Monitoring:



System Level Monitoring:



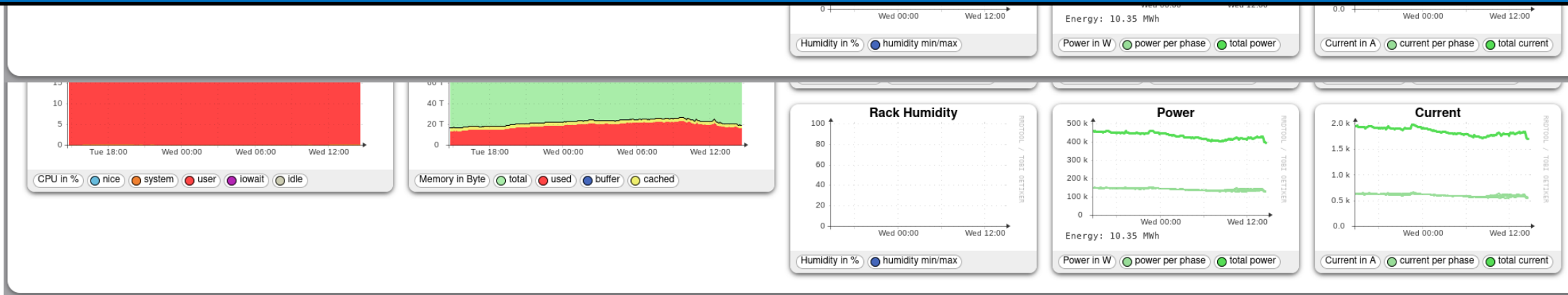
System Level Monitoring:



Manual correlation of data ...

sacct (SLURM accounting):

```
4      1      96 02:01:34 COMPLETED 0:0    336K      0      0      0
594.24K 130.23M  mpsc0150      0    130.23M  7.57M      mpsc0150
0      7.57M      cpu=96,me+ cpu=02:00:34,+ cpu=02:00:34,+ cpu=mpsc0150,ener+
cpu=0,fs/disk=0,m+ cpu=02:00:34,+ cpu=mpsc0150,ener+ cpu=0,fs/disk=0,m+
cpu=02:00:34,+ energy=81,fs/d+ energy=mpsc0150,fs+      fs/disk=0 energy=70,fs/d+
energy=70,fs/d+
```

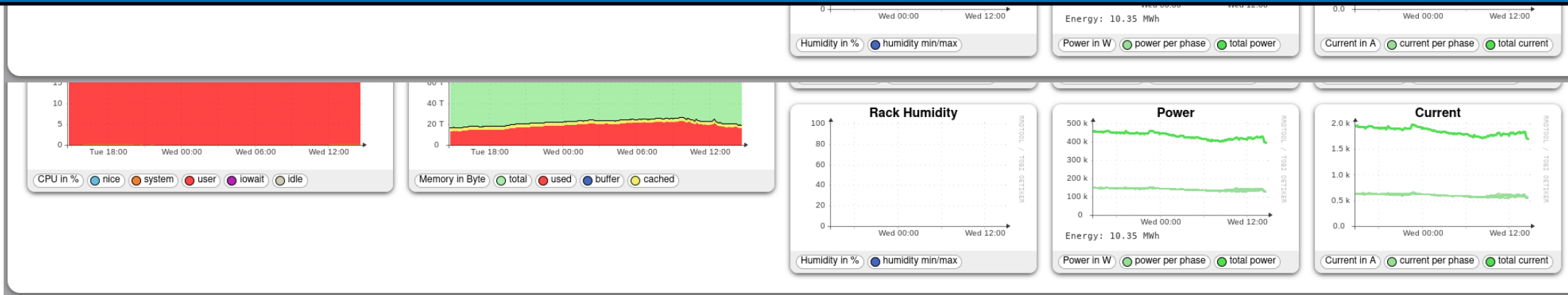


Manual correlation of data ...

sacct (SLURM accounting):

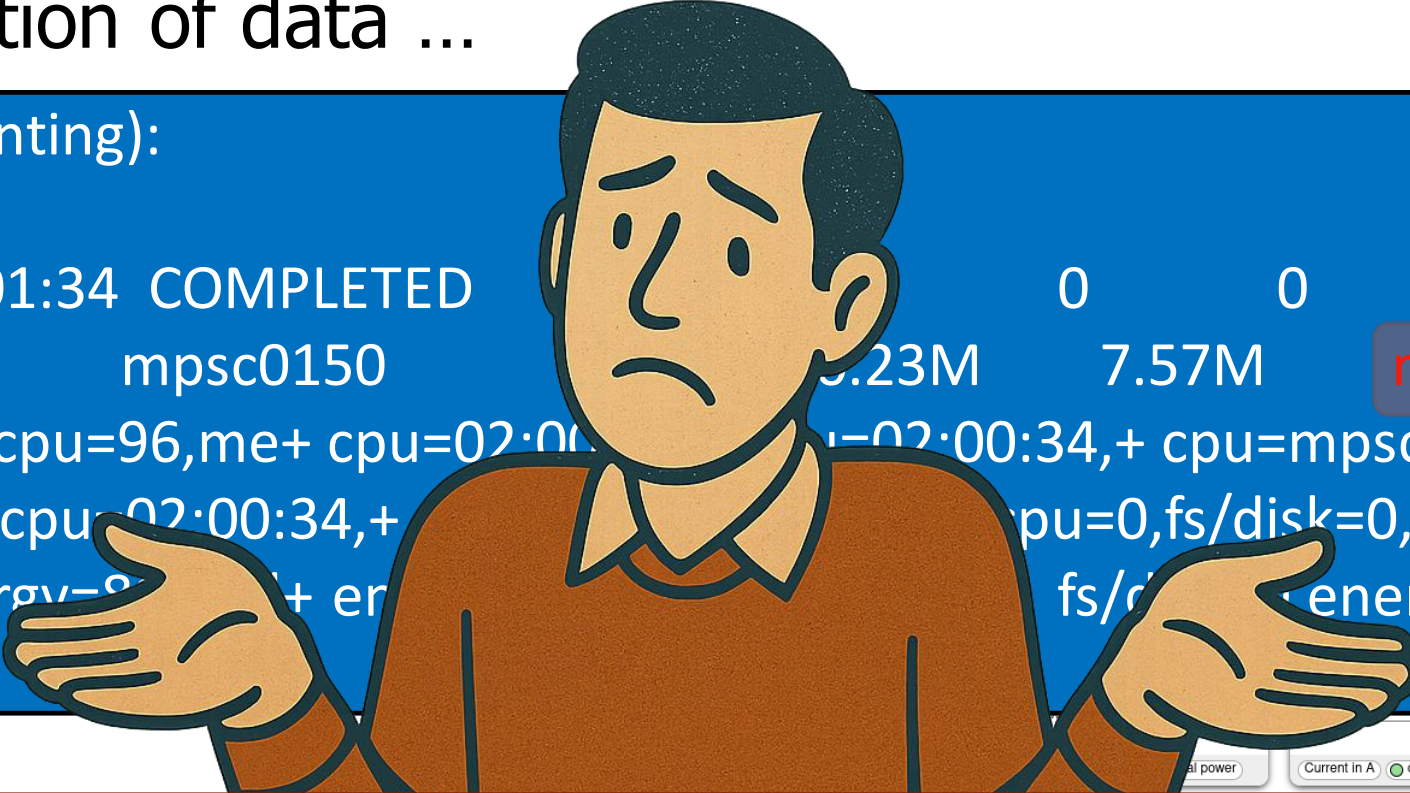
```

4      1      96 02:01:34 COMPLETED 0:0 336K 0 0 0
594.24K 130.23M mpsc0150 0 130.23M 7.57M mpsc0150
0      7.57M      cpu=96,me+ cpu=02:00:34,+ cpu=02:00:34,+ cpu=mpsc0150,ener+
cpu=0,fs/disk=0,m+ cpu=02:00:34,+ cpu=mpsc0150,ener+ cpu=0,fs/disk=0,m+
cpu=02:00:34,+ energy=81,fs/d+ energy=mpsc0150,fs+ fs/disk=0 energy=70,fs/d+
energy=70,fs/d+
  
```



Manual correlation of data ...

sacct (SLURM accounting):



```
4      1      96 02:01:34 COMPLETED      0      0      0
594.24K 130.23M mpsc0150 5.23M 7.57M mpsc0150
0      7.57M      cpu=96,me+ cpu=02:00:34,+ cpu=02:00:34,+ cpu=mpsc0150,ener+
cpu=0,fs/disk=0,m+ cpu=02:00:34,+ cpu=0,fs/disk=0,m+
cpu=02:00:34,+ energy=90,fs/disk=0,m+ energy=70,fs/d+
energy=70,fs/d+
```



Doing this without help is ... difficult!

What would we like to see?

- Domain of job:
 - Which processes ran where?
 - Which resources were used by a specific job?
 - Resources per parallel entity?
 - What aspects of the hardware was used?
- What was the hardware doing:
 - Which cores were saturated?
 - Which network connection was full?
- What was the software trying to achieve
 - Which mathematical equations were being solved?
 - Which phase of the code was operational?

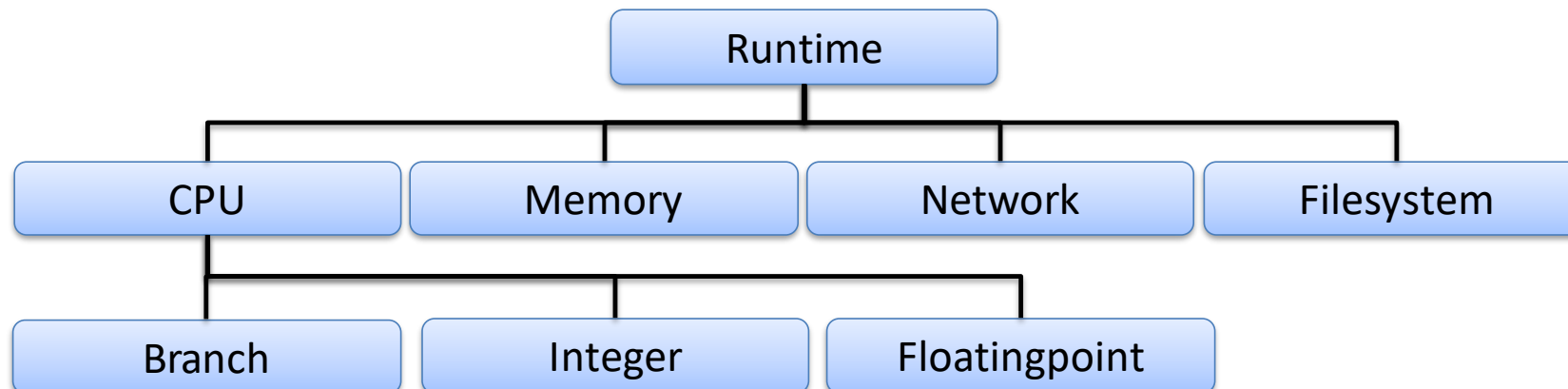


What metrics?

Simple model

$$Runtime = \frac{\text{Number of Instructions}}{\frac{\text{Instructions}}{\text{Second}}} + \frac{\frac{\text{Memory}}{\text{Instruction}}}{\frac{\text{Memory}}{\text{Second}}} + \frac{\text{Memory to Communicate}}{\frac{\text{Memory}}{\text{Second}}} + \dots$$

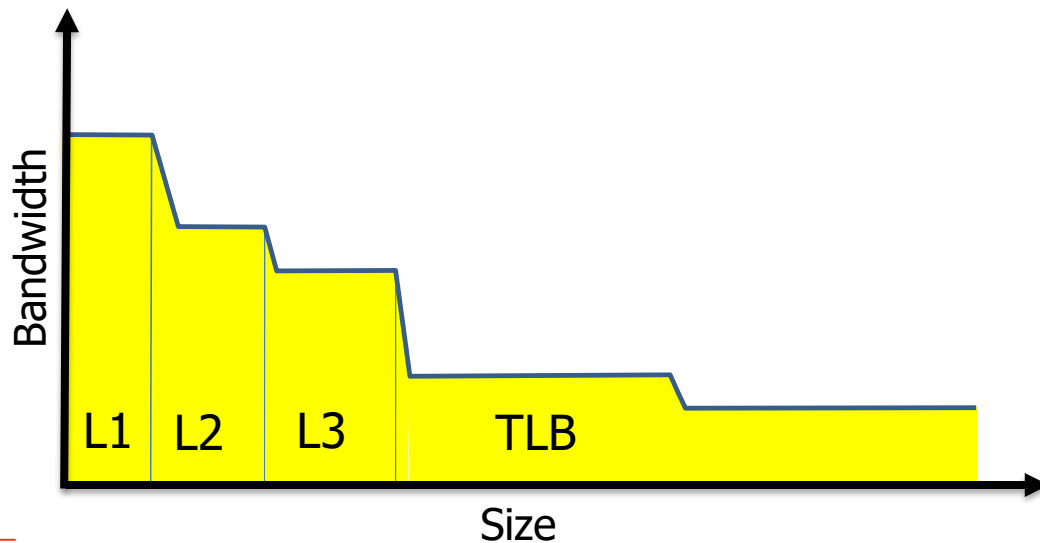
- Metrics aim to capture and overt aspects of this expression
- Each of the terms above can (and should) be split into more precise sub-terms



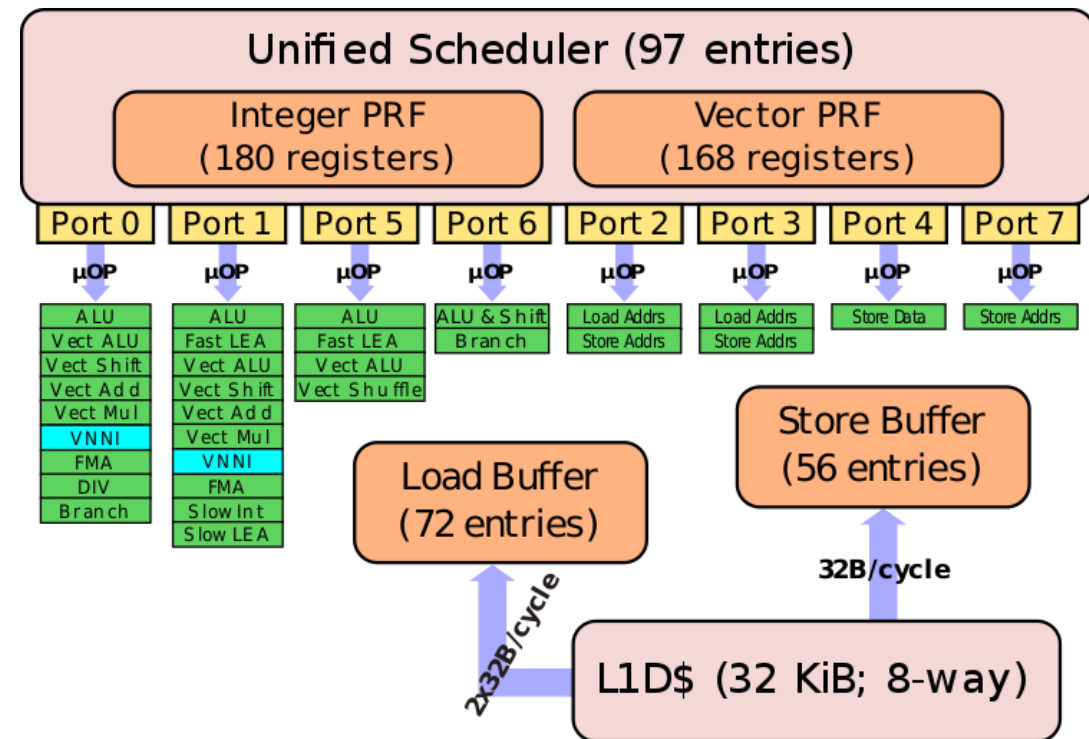
Architecture knowledge

- Understanding of the systems architectural capabilities: everything is complex

Memory



CPU design



Src.: https://en.wikichip.org/wiki/intel/microarchitectures/cascade_lake

Hardware Performance Counters

- More than 2000 events, limited counter registers

Counter-group	Details
CYCLE_STALLS	
FLOPS_DP	Double Precision MFLOP/s
L2	L2 cache bandwidth in MBytes/s
DATA	Load to store ratio
MEM	Main memory bandwidth in MBytes/s
L3CACHE	L3 cache miss rate/ratio
Power	Power and Energy consumption
BRANCH	Branch prediction miss rate/ratio
...	
CACHES	CACHES



MEM:

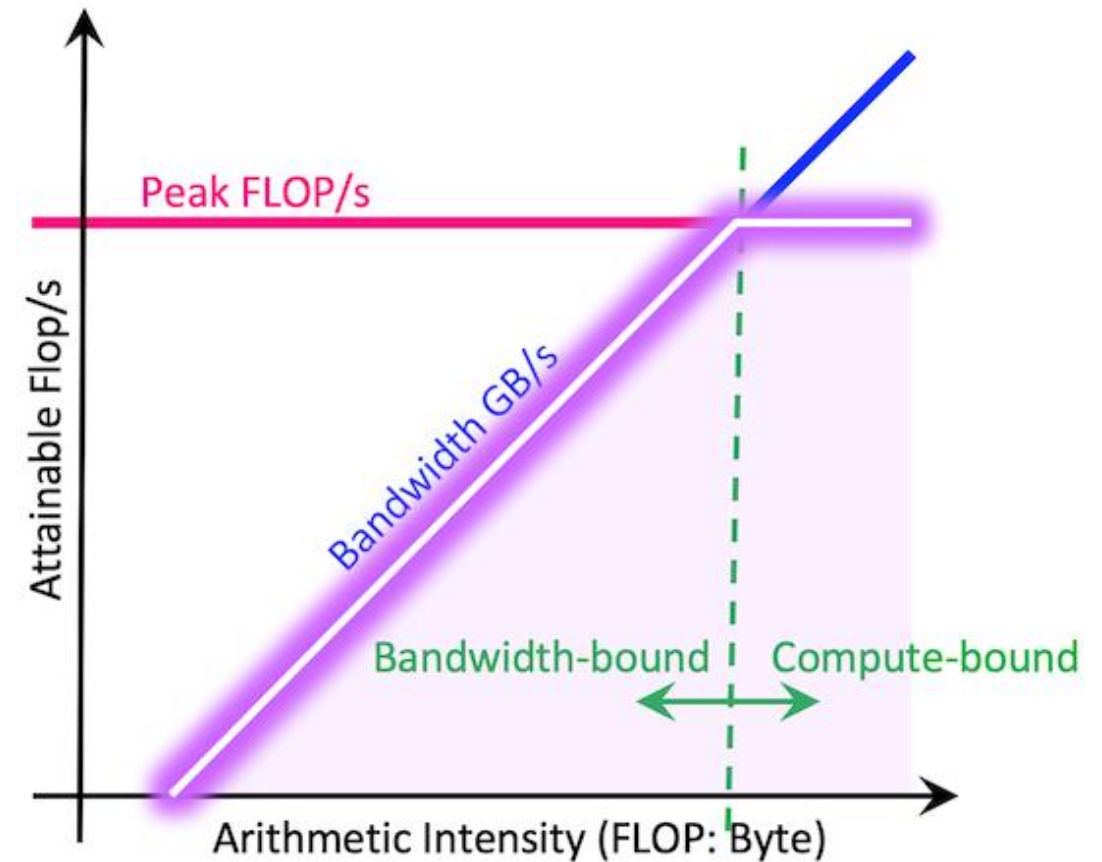
- Memory read bandwidth [MBytes/s]
- Memory read data volume [GBytes]
- Memory write bandwidth [MBytes/s]
- Memory write data volume [GBytes]
- Memory bandwidth [MBytes/s]
- Memory data volume [GBytes]



Roof-line model

- Lay-mans summary:
Performance limited either by
 - compute or
 - memory
- Arithmetic intensity $\frac{Work}{Byte}$
- Goal: Get to the ridge-line
 - Move to the right: do more work with the same data
 - Move to the left: do less work with the same data

Applies to a steady-state

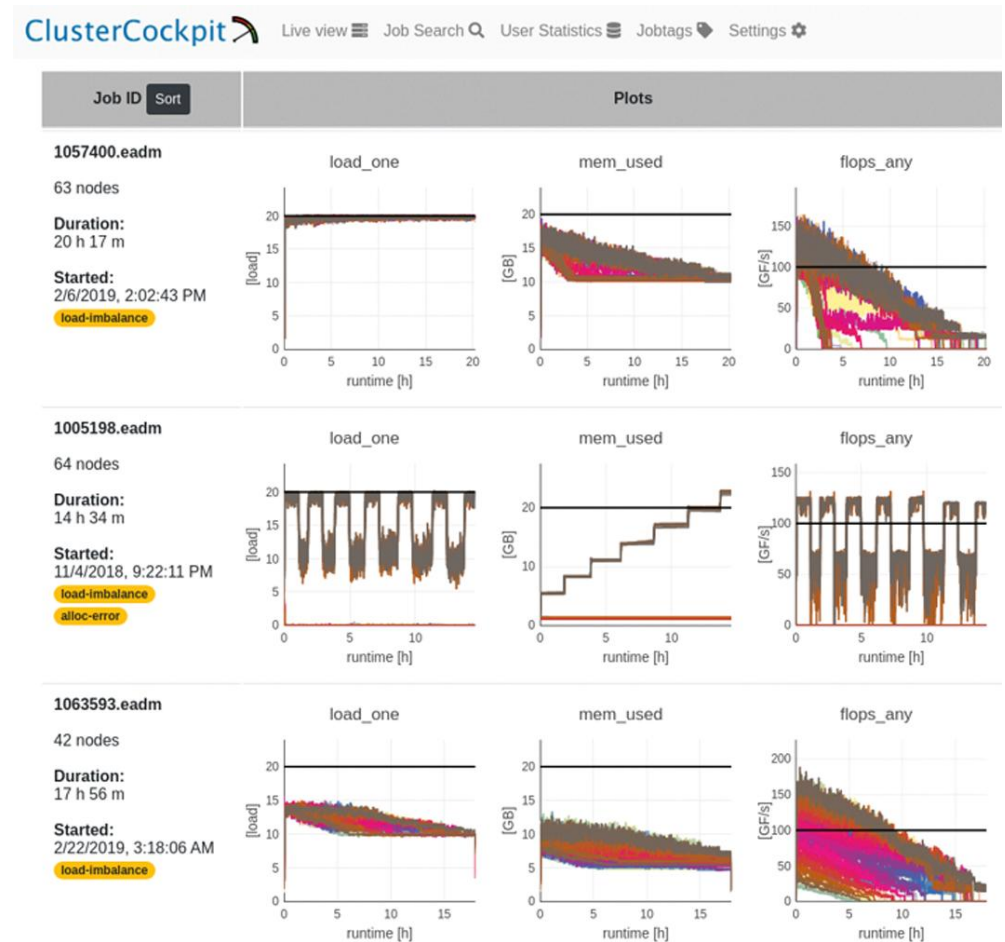


Src. <https://docs.nersc.gov/tools/performance/roofline/>



Understanding and Context

- Metrics gathered for a „whole run“ have little expressiveness
- Context during which a metric was gathered
- Correlate change in metric to specific code regions
- Requirements:
 - How it „should“ run
 - What the hardware can do
 - Observe actual execution



Src: <https://www.rze.fau.de/2020/01/webanwendung-zum-jobspezifischen-performance-monitoring-eine-erfolgsgeschichte-keine-ressourcen-verschwendung-gezielte-optimierung/>



Measurement Techniques

- Process isolation prevents direct observation of a programs state
 - Modify the target program before execution
 - Use a loader to gain control of the application
 - Keyword: LD_PRELOAD mechanism
- Terminology:
 - **Hook**: Facility to potentially measure
 - **Probe**: A measurement device, connecting to a hook
 - **Measurement**: Gathering of data and context
- Resolving the temporal state of the program
 - 1) Sampling
 - 2) Instrumentation

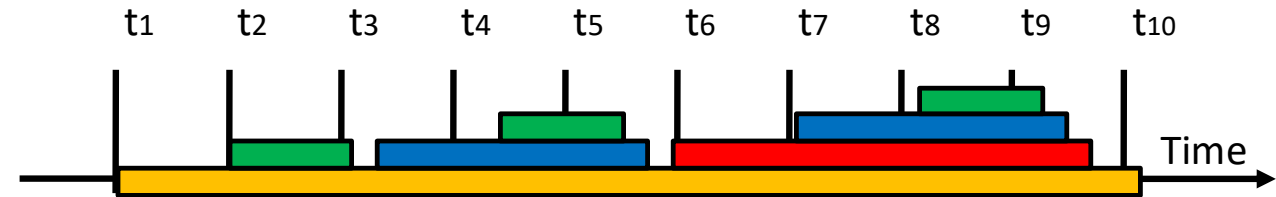


Sampling without unwinding

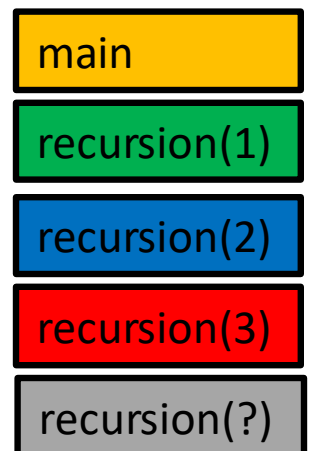
```
void recursion(int depth){  
    if (depth>1)  
        return recursion(depth-1);  
    return;  
}  
  
int main(int argc, char ** argv){  
    for (int depth=1;depth<4;depth++)  
    {  
        recursion(depth);  
    }  
}
```

- Unwinding exposes stack
- „Thunking/Trunking“ accurately tracks function exits

Call-stack view:



Measurement view:



Basic Instrumentation Concept

```
void recursion(int depth){  
    if (depth>1)  
        return recursion(depth-1);  
    return;  
}  
  
int main(int argc, char ** argv){  
    for (int depth=1;depth<4;depth++)  
    {  
        recursion(depth);  
    }  
}
```

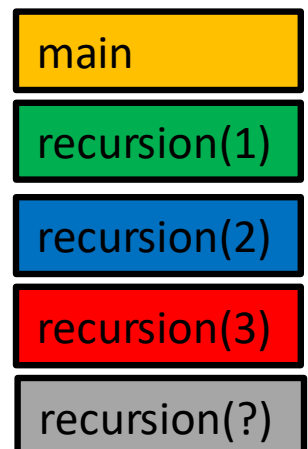
Call-stack view:



Measurement view:

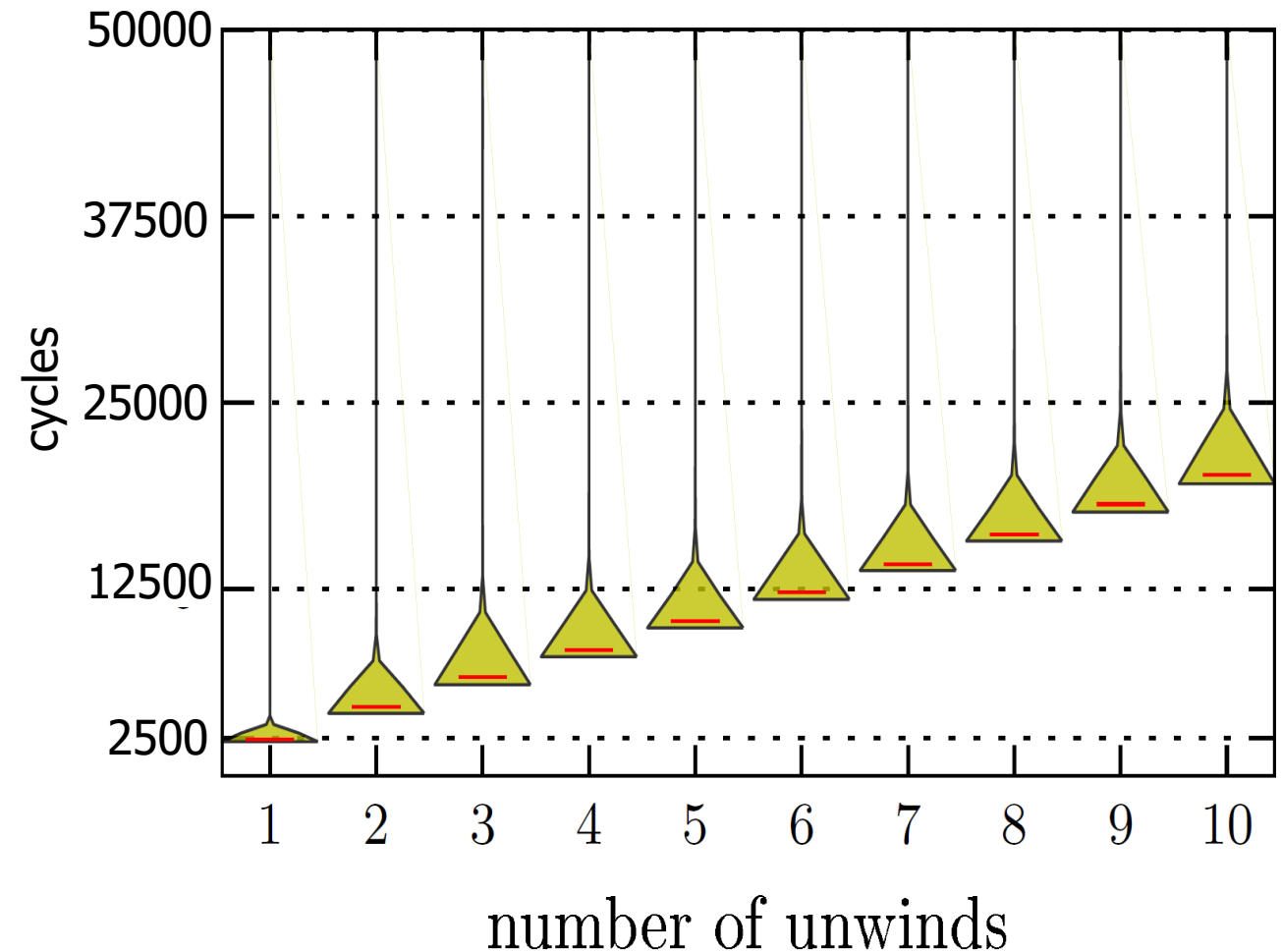


- Instrumentation before or after optimization



Overhead of sampling and unwinding

- $O_{\text{Sampling}} = \sum_{\text{locations}} O_{\text{Unwind}}(\text{loc.}) * n(\text{loc.})$
 - » $O_{\text{Unwinding}} \approx 4 * 10^{-6} \text{ s}$,
 - » $\sum_{\text{locations}} n(\text{loc.}) / \text{second} \approx 200 \text{ Hz}$
 - Unwinding is not for free
 - Example:
 - Runtime: 10000s
 - Samplerate: 200Hz
 - Average unwinding-depth: 10
- Overhead:
 $100,000 \text{ s} * 200 \text{ Hz} * 7.5 \mu\text{s} = 15 \text{ s}$



Instrumentation overhead considerations

Benchmark	Runtime	Number of Functions	Number of Function Calls
444.namd	392	100	1,62E+06
453.povray	155	809	1,41E+06
464.h264ref	62	319	2,55E+08
403.gcc	288	2734	3,40E+08
447.dealII	302	5491	1,17E+08
458.sjeng	488	73	1,11E+10
473.astar	338	129	7,66E+09
433.milc	437	112	1,73E+07
450.soplex	205	769	4,72E+09
482.sphinx3	532	203	1,25E+09
lulesh	102	254	6,50E+10
miniFE	46	552	2,80E+10
DROPS	63	7257	1,70E+11

- $O_{Instrumentation} = O_{Probe} * n_{invocations}$
- Cost of a single probe:
 - $O_{Probe} \approx 4 * 10^{-9} \text{ s}$
- Number of probe-calls:
 - $1,41 \text{ E}+06 < n_{invocations} < 170 \text{ E}+09$



Instrumentation overhead considerations

Benchmark	Runtime	Number of Functions	Number of Function Calls	Overhead	Overhead in %
444.namd	392	100	1,62E+06	1,30E-02	0%
453.povray	155	809	1,41E+06	1,13E-02	0%
464.h264ref	62	319	2,55E+08	2,04E+00	3%
403.gcc	288	2734	3,40E+08	2,72E+00	1%
447.dealll	302	5491	1,17E+08	9,36E-01	0%
458.sjeng	488	73	1,11E+10	8,88E+01	18%
473.astar	338	129	7,66E+09	6,13E+01	18%
433.milc	437	112	1,73E+07	1,38E-01	0%
450.soplex	205	769	4,72E+09	3,78E+01	18%
482.sphinx3	532	203	1,25E+09	1,00E+01	2%
lulesh	102	254	6,50E+10	5,20E+02	510%
miniFE	46	552	2,80E+10	2,24E+02	487%
DROPS	63	7257	1,70E+11	1,36E+03	2159%



Take-Away:

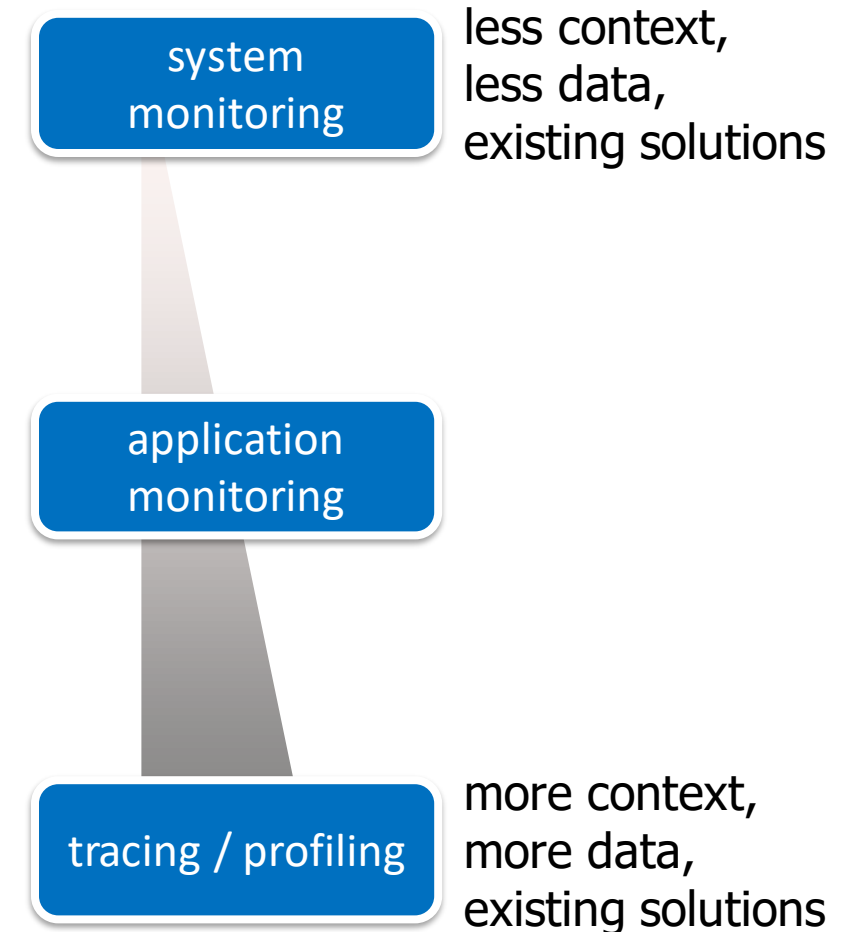
- Performance analysis techniques helpful
- Too expensive for job level monitoring
- User work to tailor measurement

 Gathering program state without help not likely



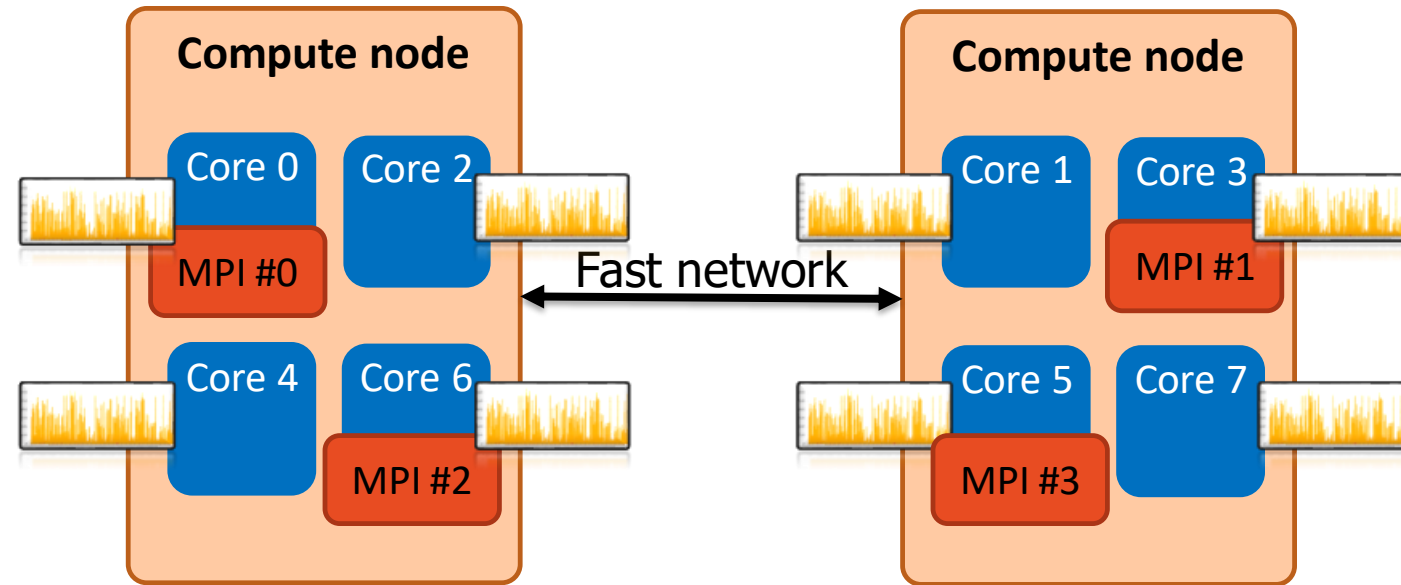
Is there a middle ground?

- System operation:
 - system healthy?
 - all hardware operational?
 - providing peak performance?
- Usage monitoring:
 - what resources are used ?
 - which resources are the limiting factor?
(compute units, memory bandwidth, memory capacity ...)?
 - are resources utilized, or just occupied?
- Performance monitoring / analysis:
 - what is the software doing?
 - why is the software behaving like that?



Correlation game:

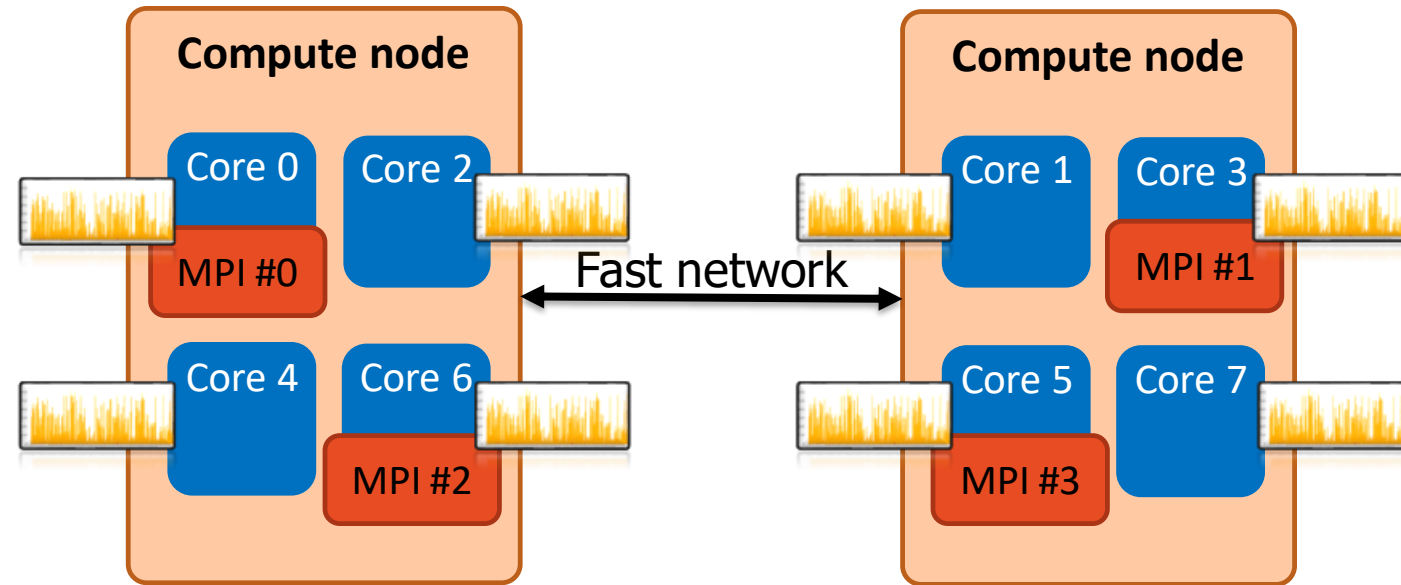
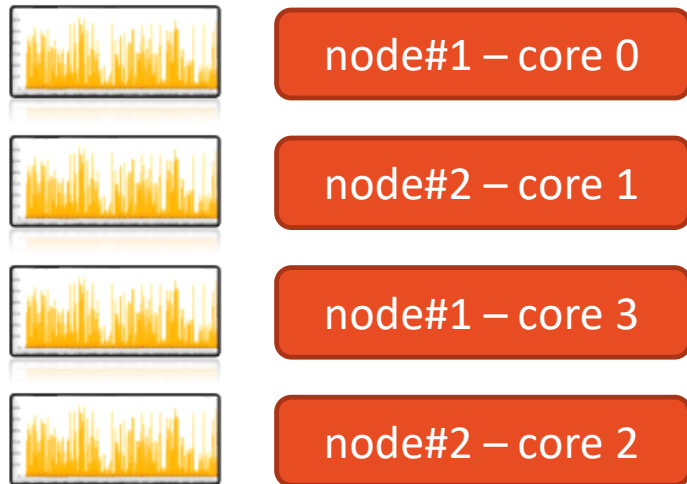
- What can we know:
 - Which software (job) ran when & where
 - Batch systems knows
 - E.g.: Which MPI rank ran when and where?
 - Time-series of metrics at granularity
 - Metrics per scheduled entity



Correlation game:

- What can we know:
 - Which software (job) ran when & where
 - Batch systems knows
 - E.g.: Which MPI rank ran when and where?
 - Time-series of metrics at granularity
 - Metrics per scheduled entity

Tool to match and correlate data!

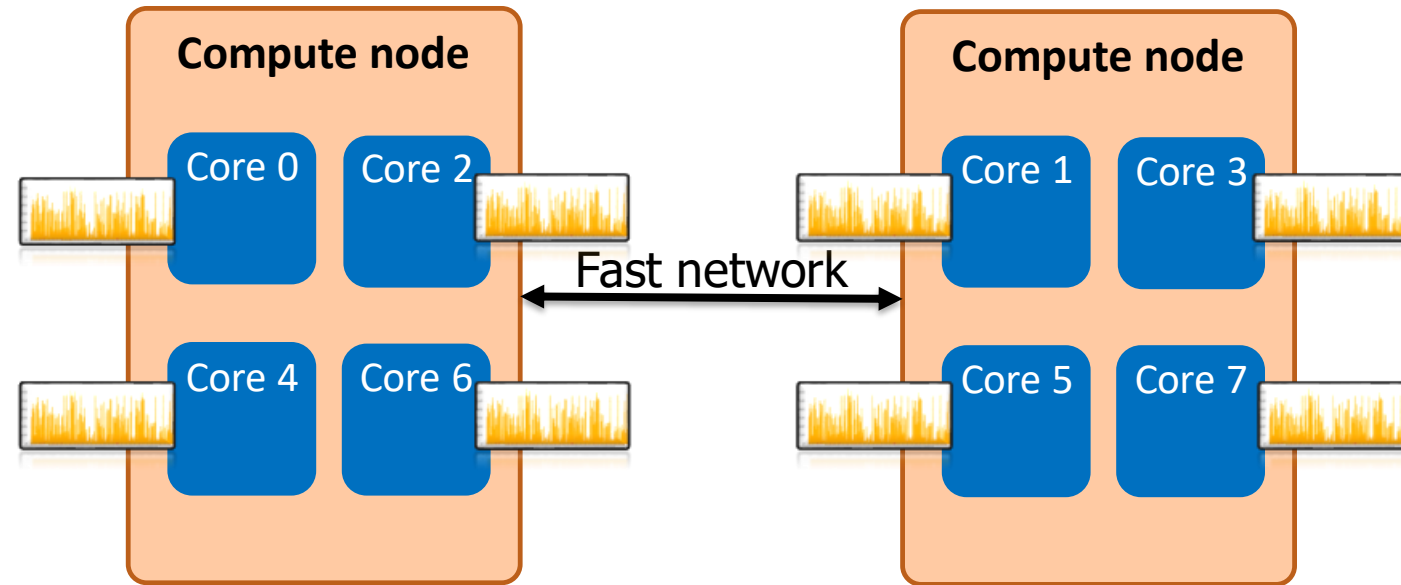


Correlation game:

- What can we know:
 - Which software (job) ran when & where
 - Batch systems knows
 - E.g.: Which MPI rank ran when and where?
 - Time-series of metrics at granularity
 - Metrics per scheduled entity

Tool to match and correlate data!

MPI #0		node#1 – core 0
MPI #1		node#2 – core 1
MPI #2		node#1 – core 3
MPI #3		node#2 – core 2



Metrics to measure ... ?

- Mock-up cluster:

- Node:

- **CPU:** ¹ 2x AMD EPYC 9965 (192 cores) ~ 8.000€ / piece
- Motherboard: Dual Socket Motherboard ~ 1.600€ / piece
- **RAM:** ² 24x 32GB DDR5-4800 (ECC) ~ 270€ / piece
- **Network:** ³ InfiniBand HDR100 ~ 850€ / piece
- ~~GPGPU: ~~NVIDIA H100~~ ~ ~~40.000€ / piece~~~~
- Misc: Cables, chassis, infrastructure, **storage** ⁴ ~ 5000€ / node

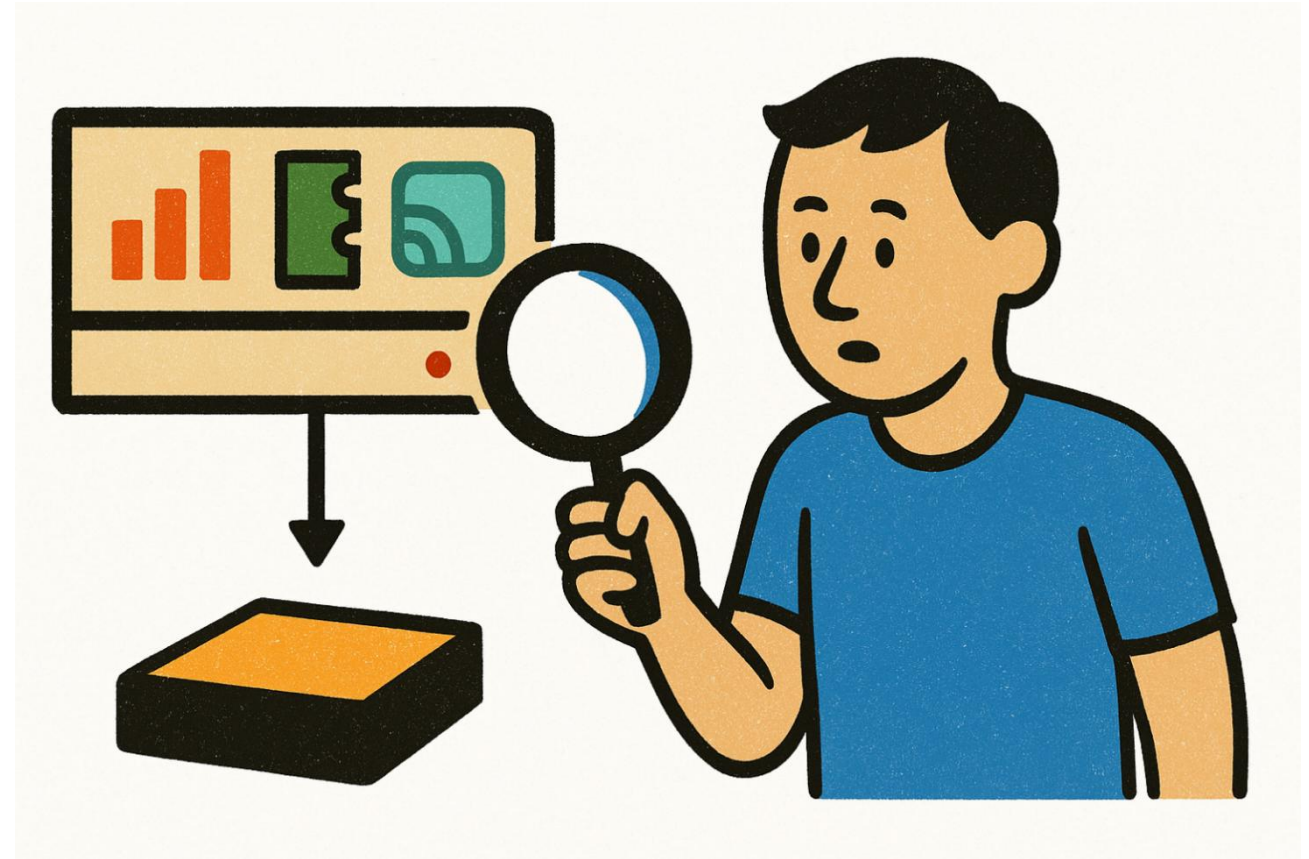
- Operational costs (1000 nodes):

- Staff: 7 FTE (full time employees) ~ 80k€ / FTE / y
- Power: ~ 3.4 MW ~ 67,70 € / MWh

- Node/h ~ 0,97€ Node/y ~ 8562€
- Cluster/h ~ 977,44€ Cluster/y ~ 8.562 M€

Metrics to measure?

- Any metric, that is related to individual hardware and usual bottlenecks:
 - **Utilization**
 - Allocation vs actual use
 - **CPU**
 - Compute unit utilization
 - Cache use
 - CPI / IPC
 - **Memory**
 - Memory use
 - Memory bandwidth
 - **Network**
 - Saturation
 - Packet size
 - **Storage**
 - IO operations



Cluster Cockpit

ClusterCockpit: Concept and Purposes

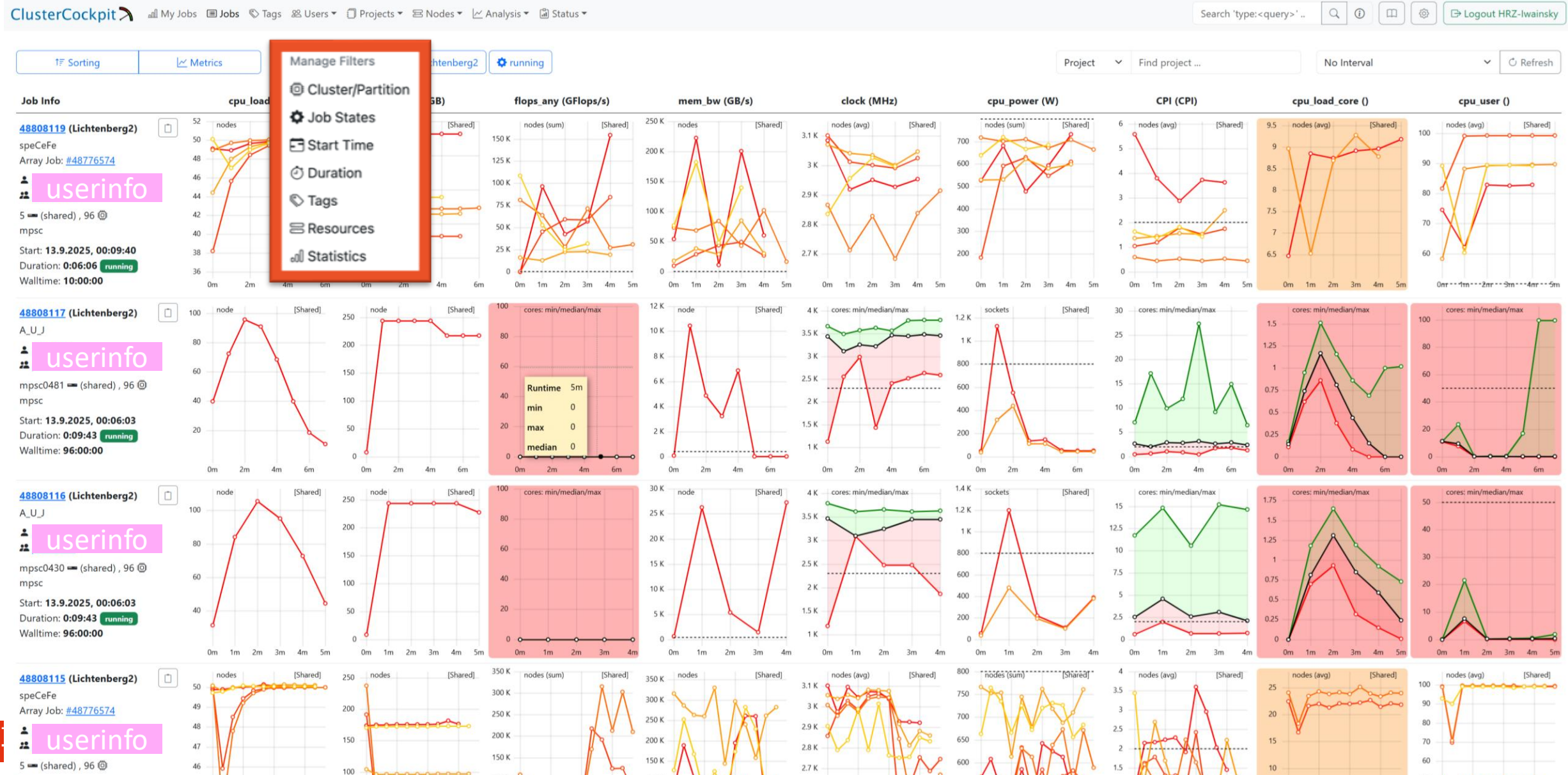
- Cluster-wide continuous monitoring
 - Hardware Performance Monitoring (HPM) metrics: CPI, GFlops/s, Memory bandwidth; load, memory usage, File IO, Network utilization, RAPL energy metrics, GPUs ...
 - Fixed frequency measurements for native granularities (core, memory domain, socket, node)
- UI for
 - job monitoring,
 - early detection of performance or runtime issues
 - access to aggregate statistics for users and jobs
 - Web-interface (different views for roles: admin, support, manager, user)
 - Search, filter, tag, and sort jobs, users, and projects
 - Job-specific views with metric plots, job statistics, job meta information
 - Cluster-specific views with overall utilization, job and user statistics



<https://github.com/ClusterCockpit/>



Cluster Cockpit: Job-list



Cluster Cockpit: Job-list

ClusterCockpit  My Jobs  Jobs  Tags  Users  Projects  Nodes  Analysis  Status 

Search 'type:<query>' ..



Logout HRZ-Iwinsky

Sort rows

- Start Time
- Duration
- Number of Nodes
- Max. Memory Used
- Avg. FLOPs
- Avg. Memory Bandwidth
- Avg. Network Bandwidth

Close

mpsc

Start: 13.9.2025, 00:06:03

Duration: 0:09:43 running

Walltime: 96:00:00

48808116 (Lichtenberg2)

A_U_J

  **userinfo**

mpsc0430 == (shared) , 96 @ mpsc

Start: 13.9.2025, 00:06:03

Duration: 0:09:43 running

Walltime: 96:00:00

48808115 (Lichtenberg2)

speCeFe

Array Job: #48776574

  **userinfo**

5 == (shared) , 96 @

Lichtenberg2

running

Project

Find project ...

No Interval

Refresh

used (GB)

flops_any (GFlops/s)

mem_bw (GB/s)

clock (MHz)

cpu_power (W)

CPI (CPI)

cpu_load_core ()

cpu_user ()

[Shared]

[Shared]

[Shared]

[Shared]

[Shared]

[Shared]

[Shared]

[Shared]

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (avg)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (avg)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (avg)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (avg)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (avg)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (avg)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (avg)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (avg)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (avg)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (avg)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (avg)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (avg)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (avg)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (avg)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (avg)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (avg)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (avg)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (avg)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (avg)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (avg)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (avg)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (avg)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (avg)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (avg)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (avg)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (avg)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (avg)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (avg)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (avg)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (avg)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (avg)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (avg)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (avg)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum)

nodes (avg)

nodes (avg)

nodes (avg)

nodes (sum)

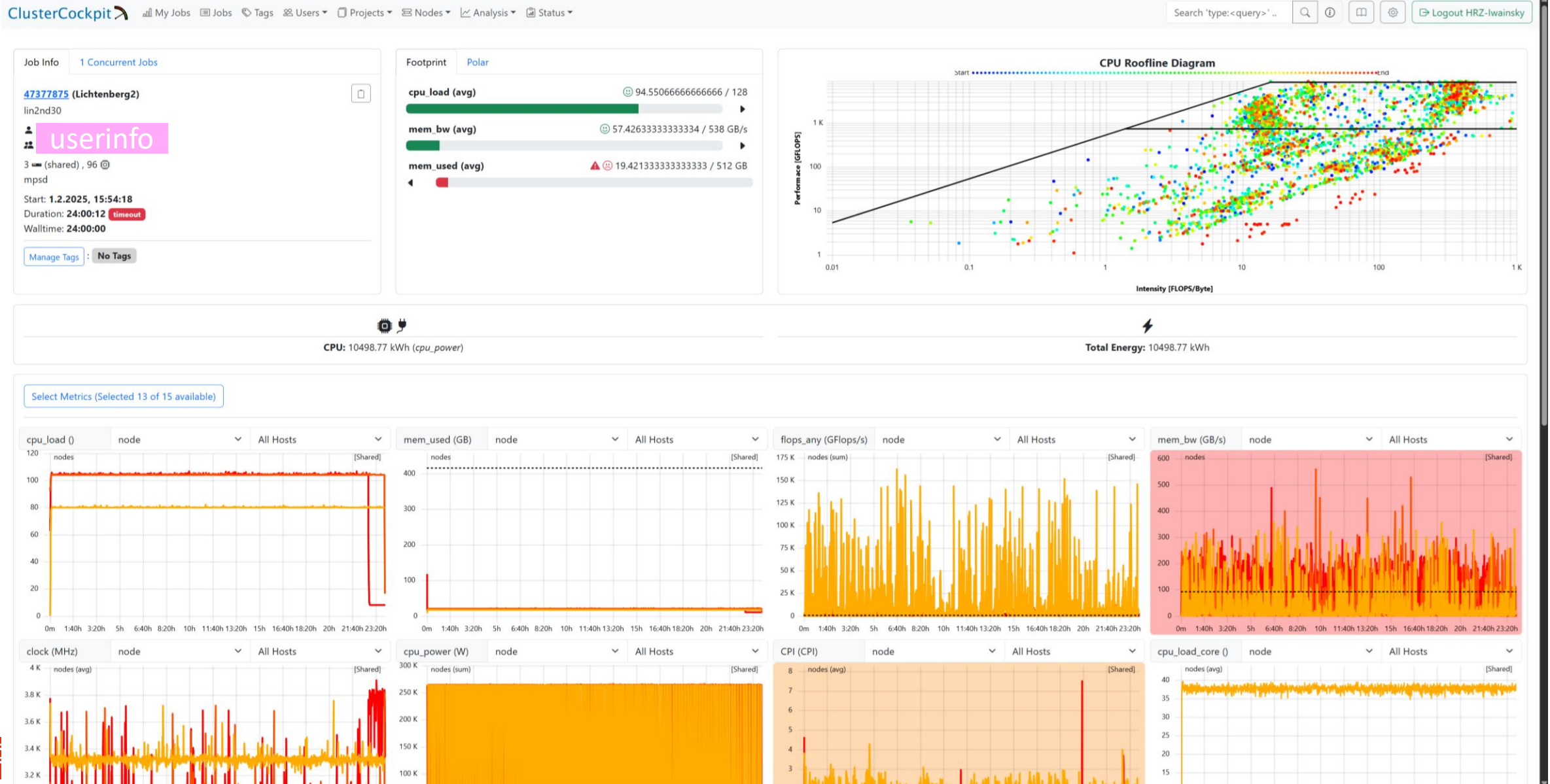
nodes (avg)

nodes (sum)

nodes (avg)

nodes (sum

Cluster-Cockpit: Job-view



Cluster-Cockpit: Job-view



```
#!/bin/bash
#SBATCH -A userinfo
#SBATCH -t 24:00:00
#SBATCH --mem-per-cpu=3600
#SBATCH -n 96
#SBATCH -c 1
#SBATCH -d singleton
```



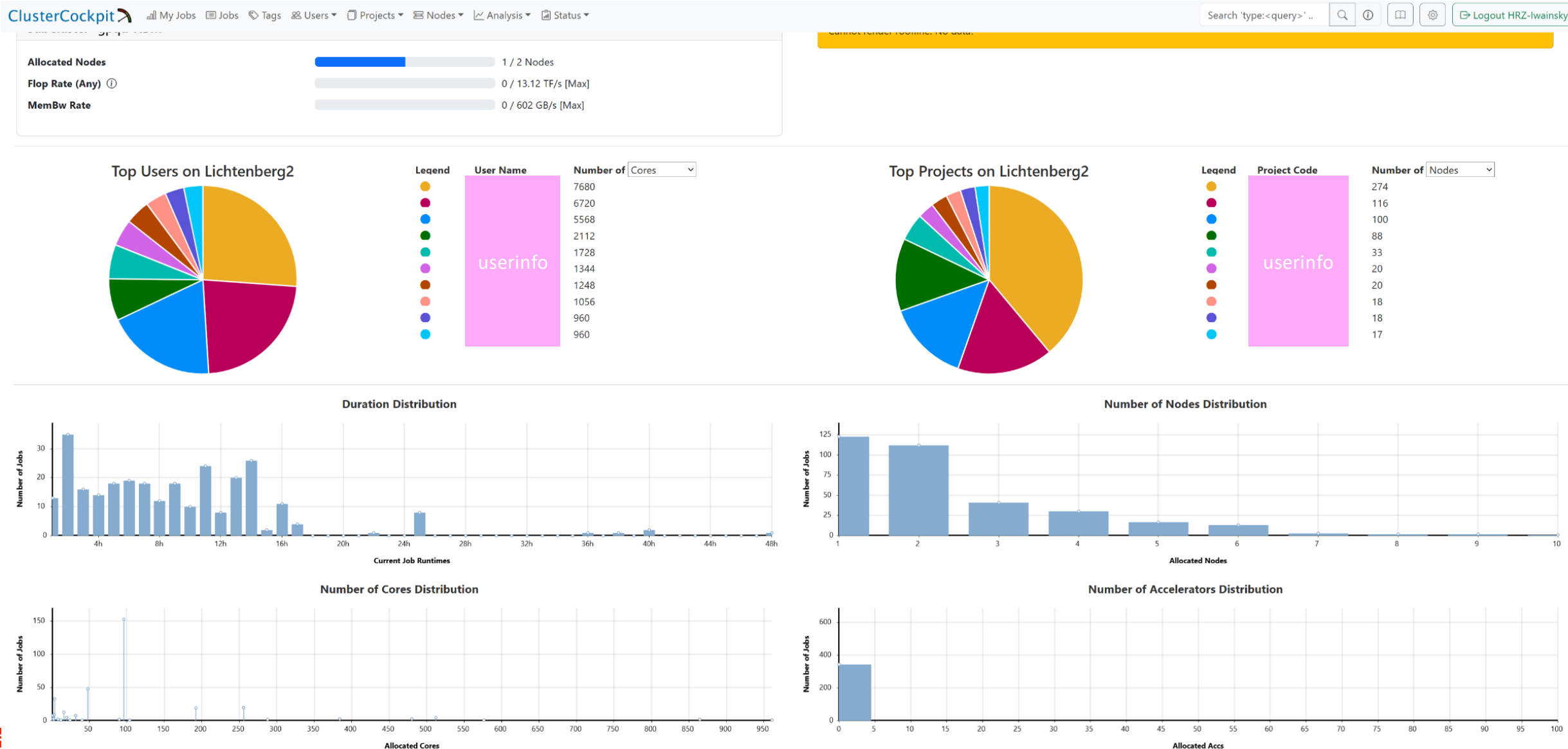
Cluster-Cockpit: Job-view



CusterCockpit: User-list (a project list is also available)

Username	Name	Total Jobs	Total Waltime	Total Core Hours	Total Accelerator Hours
userinfo	userinfo	16664	28376	28376	0
userinfo	userinfo	4159	25052	25052	0
userinfo	userinfo	1838	23682	378909	23681
userinfo	userinfo	119916	19073	32914	0
userinfo	userinfo	11302	18481	295709	18481
userinfo	userinfo	716	13705	986793	0
userinfo	userinfo	4234	13432	967142	0
userinfo	userinfo	1446	12146	194313	12146
userinfo	userinfo	3562	10032	10032	0
userinfo	userinfo	2826	9880	59281	0
userinfo	userinfo	533	9864	631300	0
userinfo	userinfo	763	9373	299900	0
userinfo	userinfo	3573	8197	3541186	0
userinfo	userinfo	1023	8186	130986	8186
userinfo	userinfo	3479	7811	141933	8870
userinfo	userinfo	5967	7431	118900	7431
userinfo	userinfo	761	7318	29274	0
userinfo	userinfo	314	6996	223874	13992

CusterCockpit: Cluster status view (aka management view)



ClusterCockpit: Cluster-view



ClusterCockpit: User-view



ClusterCockpit: Components

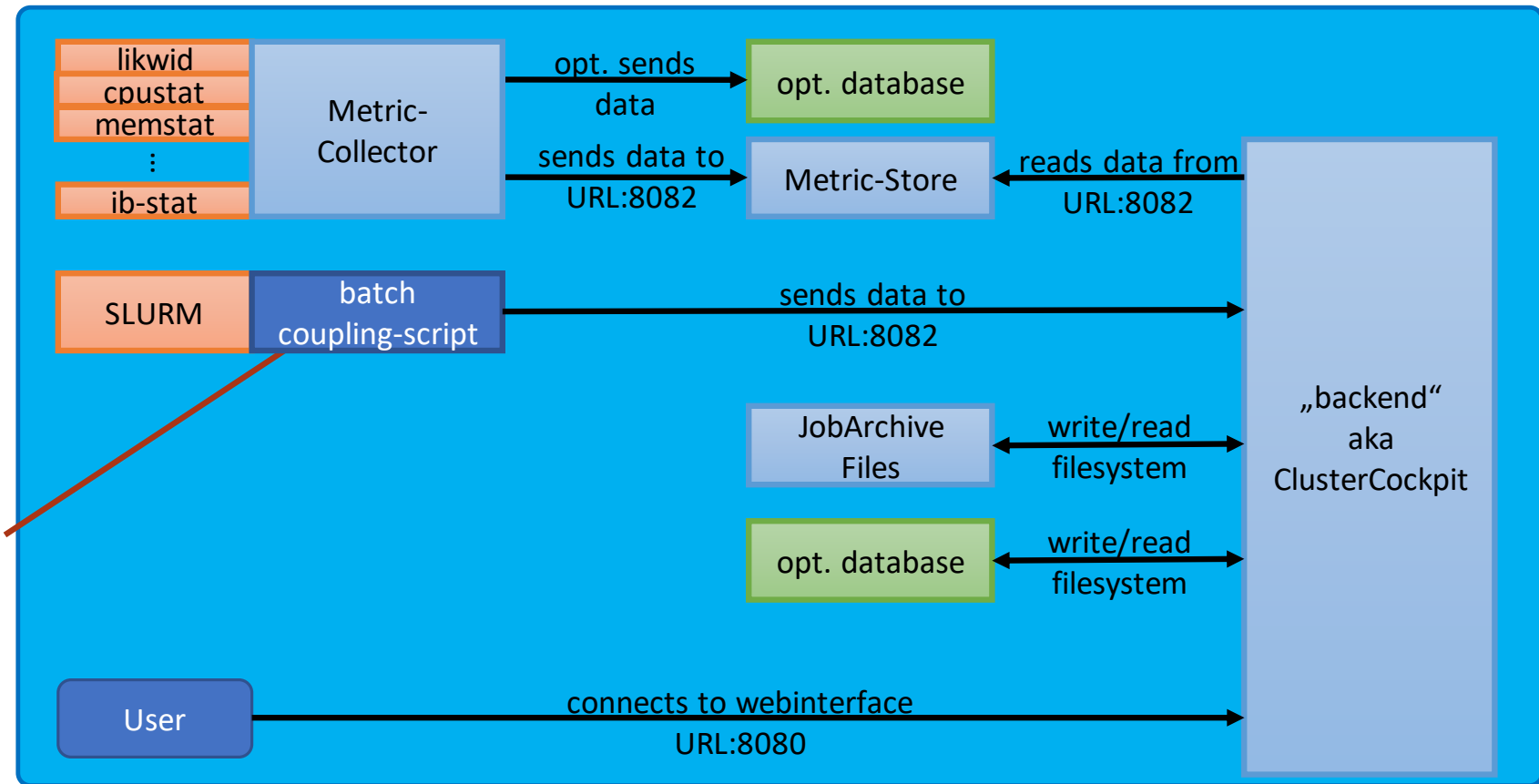
■ Components

- cc-backend
- cc-metric-collector
- cc-metric-store

■ Open-source MIT license

You must provide adapter, e.g.:

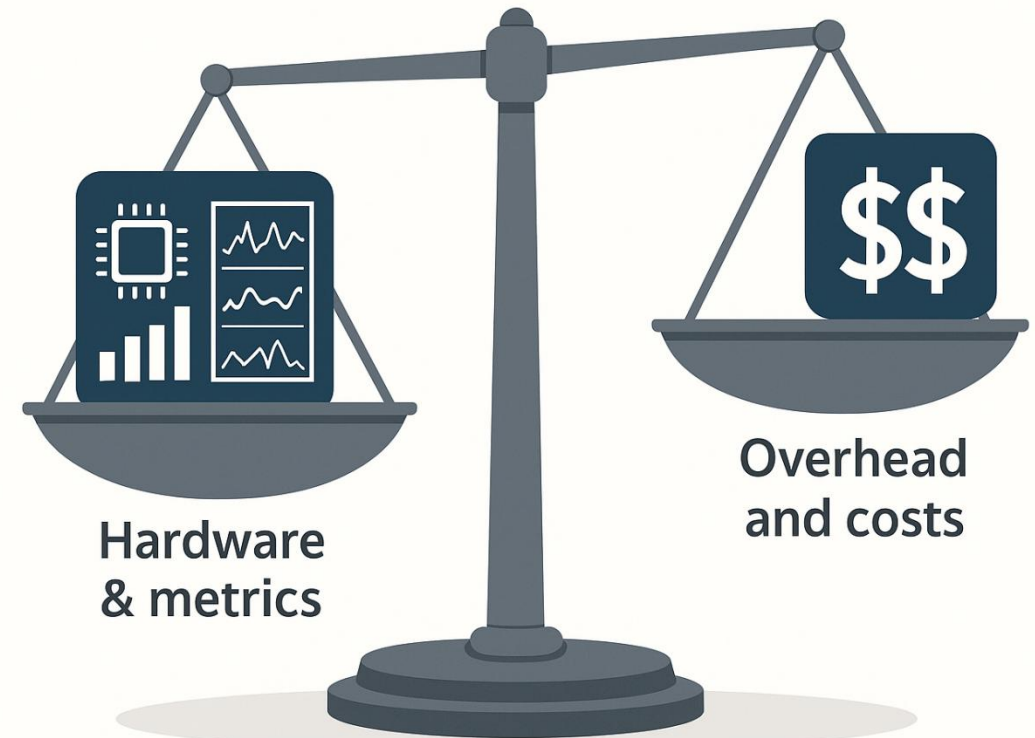
- Slurm commands
- Slurm REST API
- PrEp Plugin
- Existing SLURM coupling



<https://github.com/ClusterCockpit/>

... nice tool, but at what price?

- Metrics collectors at the node level
 - Collector Hardware Performance Counters / Core
 - 96 / 106 / 392 times the base overhead
 - System level metrics
 - Once per system
 - Data-transfer / node
- Metric Store
 - Data-sink for all nodes
 - O(2kB) data per node
 - Data retention for max job length
- Backend / Visualization

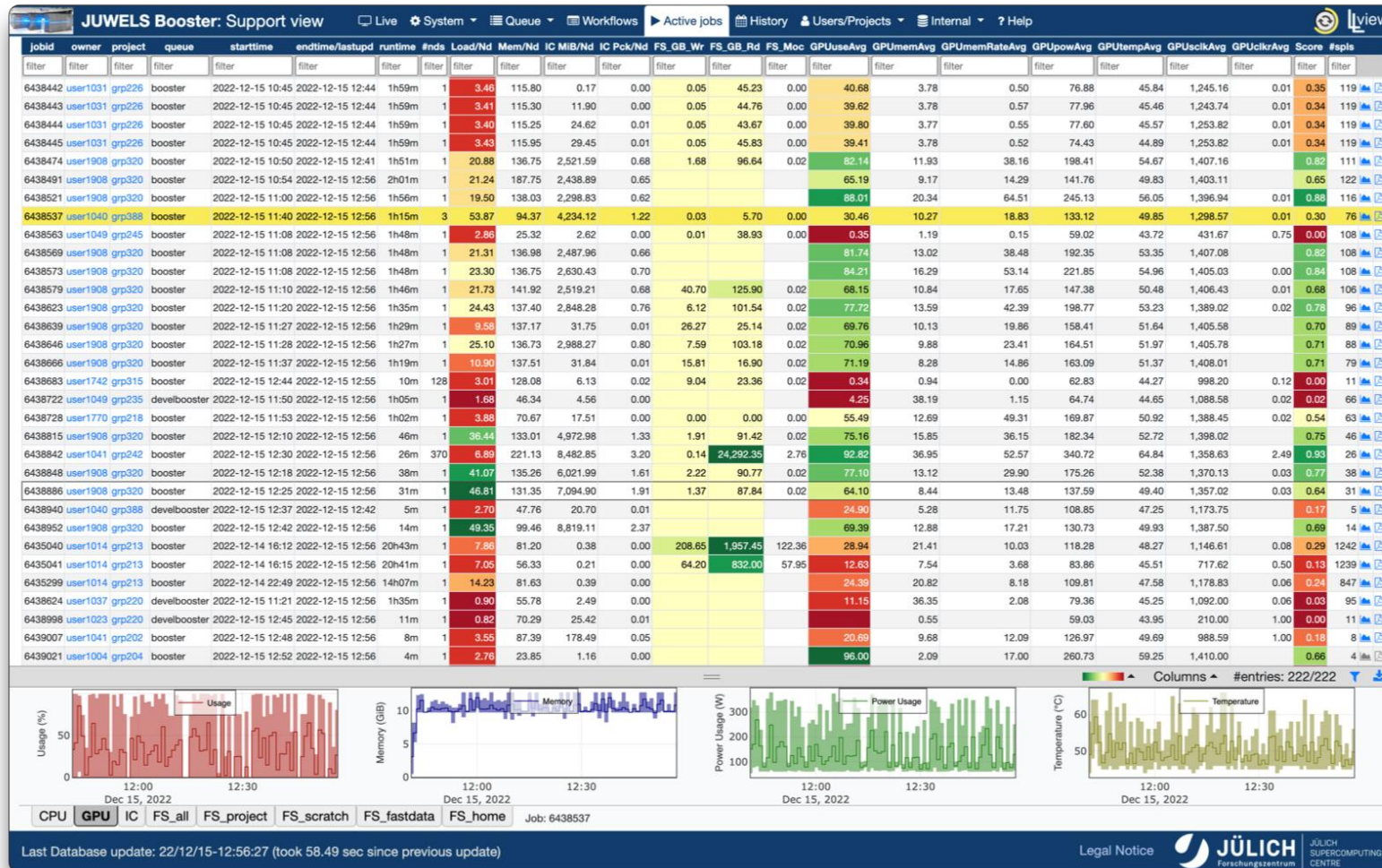


Alternatives?

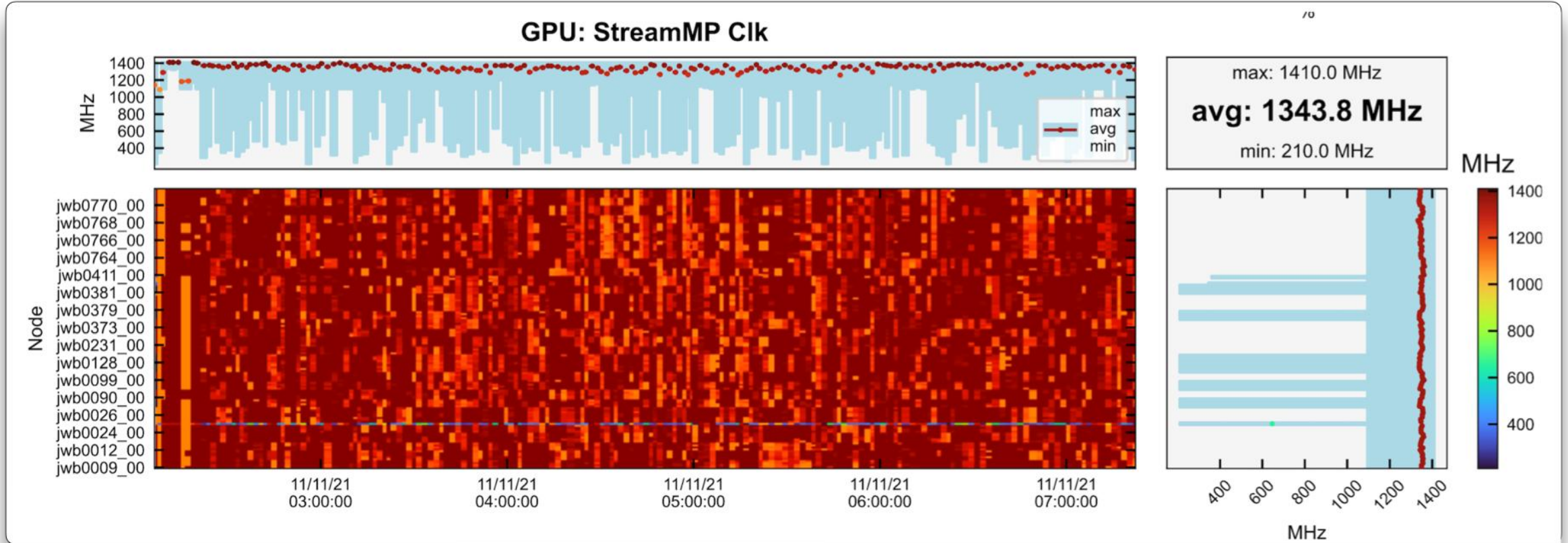
Other tools?

- Many tools
- Highlight:
LLView by Jülich Research Center
<https://github.com/FZJ-JSC/LLview>
 - System level views
 - Per Job statistics
 - No additional hardware counters
 - Not suited for shared nodes

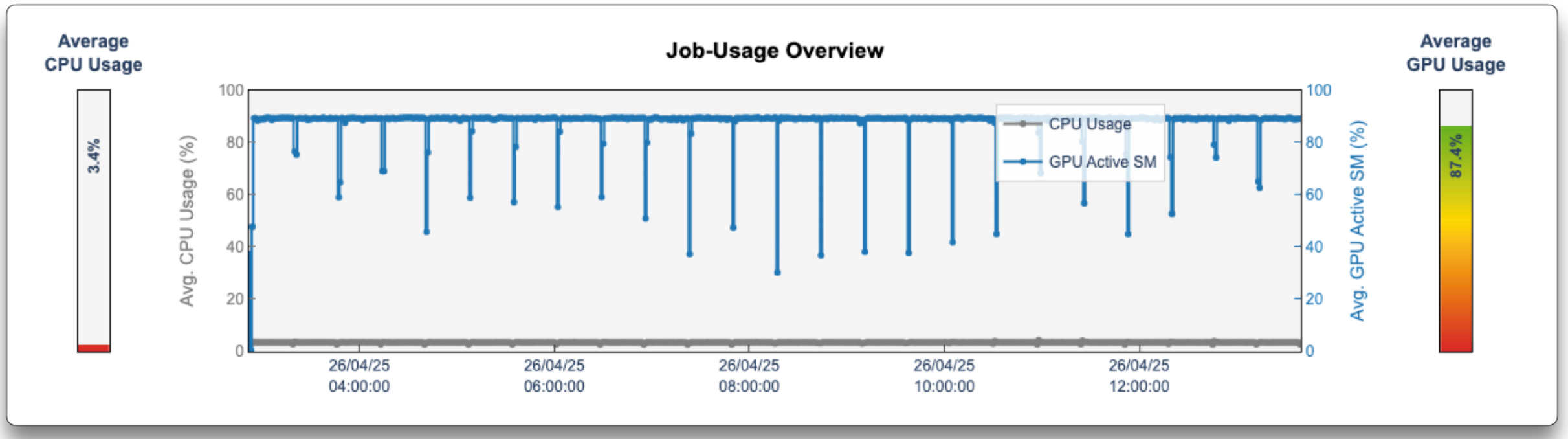
LLView



LLView: GPU stats



LLView: Job Usage



LLView: Node list

Node List					
1 GPU0: 74.9% 12.8GiB 1333.0MHz 50.0C GPU1: 76.2% 11.0GiB 1337.0MHz 49.9C GPU2: 73.7% 11.0GiB 1329.0MHz 52.8C GPU3: 76.1% 11.0GiB 1336.0MHz 52.5C Interconnect group: 6	2 GPU0: 86.4% 11.0GiB 1369.0MHz 49.5C GPU1: 86.5% 11.1GiB 1370.0MHz 50.6C GPU2: 86.6% 11.0GiB 1370.0MHz 49.3C GPU3: 86.6% 11.0GiB 1370.0MHz 49.3C Interconnect group: 6	3 GPU0: 84.2% 11.0GiB 1361.0MHz 49.4C GPU1: 86.0% 11.1GiB 1366.0MHz 49.7C GPU2: 85.6% 11.0GiB 1364.0MHz 50.4C GPU3: 84.6% 11.0GiB 1361.0MHz 49.7C Interconnect group: 6	4 GPU0: 84.6% 11.0GiB 1361.0MHz 48.9C GPU1: 87.0% 11.1GiB 1368.0MHz 49.4C GPU2: 85.6% 11.0GiB 1364.0MHz 50.3C GPU3: 86.1% 11.0GiB 1366.0MHz 49.5C Interconnect group: 6	5 GPU0: 87.5% 11.0GiB 1370.0MHz 50.1C GPU1: 87.9% 11.1GiB 1371.0MHz 52.6C GPU2: 87.5% 11.0GiB 1370.0MHz 50.0C GPU3: 87.8% 11.0GiB 1370.0MHz 49.4C Interconnect group: 6	6 GPU0: 82.5% 11.0GiB 1354.0MHz 49.0C GPU1: 84.5% 11.1GiB 1360.0MHz 49.9C GPU2: 82.6% 11.0GiB 1354.0MHz 49.8C GPU3: 83.7% 11.0GiB 1358.0MHz 48.7C Interconnect group: 6
7 GPU0: 86.8% 11.0GiB 1368.0MHz 50.8C GPU1: 86.5% 11.1GiB 1367.0MHz 50.7C GPU2: 86.2% 11.0GiB 1366.0MHz 49.9C GPU3: 86.8% 11.0GiB 1368.0MHz 49.2C Interconnect group: 6	8 GPU0: 88.9% 11.0GiB 1374.0MHz 52.5C GPU1: 89.3% 11.1GiB 1375.0MHz 52.5C GPU2: 88.8% 11.0GiB 1375.0MHz 53.2C GPU3: 88.3% 11.0GiB 1374.0MHz 51.5C Interconnect group: 6	9 GPU0: 55.6% 11.0GiB 1279.0MHz 48.8C GPU1: 64.6% 11.1GiB 1304.0MHz 49.5C GPU2: 61.6% 11.0GiB 1296.0MHz 49.0C GPU3: 62.4% 11.0GiB 1299.0MHz 48.4C Interconnect group: 7	10 GPU0: 89.2% 11.0GiB 1375.0MHz 49.3C GPU1: 90.0% 11.1GiB 1378.0MHz 51.1C GPU2: 88.8% 11.0GiB 1374.0MHz 50.5C GPU3: 90.0% 11.0GiB 1378.0MHz 50.2C Interconnect group: 7	11 GPU0: 69.1% 11.1GiB 1318.0MHz 50.2C GPU1: 73.1% 11.1GiB 1329.0MHz 49.9C GPU2: 70.9% 11.1GiB 1322.0MHz 50.2C GPU3: 72.8% 11.0GiB 1327.0MHz 50.0C Interconnect group: 8	12 GPU0: 82.6% 11.1GiB 1358.0MHz 50.9C GPU1: 82.0% 11.1GiB 1357.0MHz 50.2C GPU2: 83.6% 11.1GiB 1361.0MHz 51.0C GPU3: 83.0% 11.0GiB 1360.0MHz 50.7C Interconnect group: 8
13 GPU0: 87.6% 11.0GiB 1373.0MHz 49.9C GPU1: 88.2% 11.1GiB 1374.0MHz 49.7C GPU2: 88.2% 11.0GiB 1373.0MHz 50.5C GPU3: 88.5% 11.0GiB 1375.0MHz 49.3C Interconnect group: 9	14 GPU0: 87.8% 11.0GiB 1371.0MHz 50.7C GPU1: 87.2% 11.1GiB 1372.0MHz 49.8C GPU2: 87.2% 11.0GiB 1371.0MHz 50.4C GPU3: 87.8% 11.0GiB 1374.0MHz 49.9C Interconnect group: 9	15 GPU0: 85.5% 11.0GiB 1364.0MHz 49.4C GPU1: 86.9% 11.1GiB 1369.0MHz 49.8C GPU2: 86.9% 11.0GiB 1369.0MHz 49.8C GPU3: 88.5% 11.0GiB 1374.0MHz 50.1C Interconnect group: 9	16 GPU0: 84.6% 11.0GiB 1361.0MHz 49.3C GPU1: 86.5% 11.1GiB 1367.0MHz 50.0C GPU2: 86.2% 11.0GiB 1367.0MHz 49.9C GPU3: 86.6% 11.0GiB 1368.0MHz 48.8C Interconnect group: 9	17 GPU0: 48.0% 11.0GiB 1250.0MHz 48.6C GPU1: 52.6% 11.1GiB 1264.0MHz 48.8C GPU2: 57.1% 11.0GiB 1277.0MHz 49.4C GPU3: 54.5% 11.0GiB 1270.0MHz 48.0C Interconnect group: 9	18 GPU0: 90.9% 11.0GiB 1382.0MHz 50.3C GPU1: 89.8% 11.1GiB 1377.0MHz 49.0C GPU2: 89.6% 11.0GiB 1377.0MHz 51.9C GPU3: 90.3% 11.0GiB 1379.0MHz 49.7C Interconnect group: 9
19 GPU0: 88.3% 11.0GiB 1371.0MHz 49.8C GPU1: 87.1% 11.1GiB 1369.0MHz 49.4C GPU2: 87.8% 11.0GiB 1371.0MHz 48.6C GPU3: 87.9% 11.0GiB 1372.0MHz 49.8C Interconnect group: 9	20 GPU0: 82.4% 11.0GiB 1355.0MHz 50.1C GPU1: 82.5% 11.1GiB 1363.0MHz 50.2C GPU2: 82.6% 11.0GiB 1357.0MHz 49.7C GPU3: 83.2% 11.0GiB 1357.0MHz 49.4C Interconnect group: 9	21 GPU0: 84.8% 11.0GiB 1360.0MHz 49.6C GPU1: 84.6% 11.1GiB 1363.0MHz 50.2C GPU2: 84.7% 11.0GiB 1363.0MHz 50.1C GPU3: 85.9% 11.0GiB 1367.0MHz 49.0C Interconnect group: 9	22 GPU0: 89.5% 11.0GiB 1377.0MHz 49.5C GPU1: 90.0% 11.1GiB 1381.0MHz 49.8C GPU2: 89.1% 11.0GiB 1378.0MHz 50.7C GPU3: 89.1% 11.0GiB 1378.0MHz 50.4C Interconnect group: 9	23 GPU0: 85.6% 11.0GiB 1366.0MHz 49.9C GPU1: 85.8% 11.1GiB 1367.0MHz 49.6C GPU2: 85.4% 11.0GiB 1366.0MHz 49.5C GPU3: 85.9% 11.0GiB 1367.0MHz 49.7C Interconnect group: 9	24 GPU0: 89.6% 11.0GiB 1378.0MHz 50.6C GPU1: 89.4% 11.1GiB 1377.0MHz 50.0C GPU2: 89.6% 11.0GiB 1377.0MHz 50.3C GPU3: 89.5% 11.0GiB 1377.0MHz 49.7C Interconnect group: 10
25 GPU0: 69.9% 11.0GiB 1318.0MHz 49.5C GPU1: 75.7% 11.1GiB 1335.0MHz 49.9C GPU2: 74.3% 11.0GiB 1331.0MHz 49.0C GPU3: 71.7% 11.0GiB 1323.0MHz 49.6C Interconnect group: 10	26 GPU0: 87.3% 11.0GiB 1370.0MHz 50.5C GPU1: 87.6% 11.1GiB 1371.0MHz 50.1C GPU2: 88.1% 11.0GiB 1372.0MHz 52.1C GPU3: 88.4% 11.0GiB 1373.0MHz 49.6C Interconnect group: 10	27 GPU0: 87.7% 11.0GiB 1371.0MHz 49.8C GPU1: 88.4% 11.1GiB 1373.0MHz 49.7C GPU2: 88.8% 11.0GiB 1374.0MHz 49.3C GPU3: 88.2% 11.0GiB 1373.0MHz 49.9C Interconnect group: 10	28 GPU0: 87.2% 11.0GiB 1370.0MHz 50.6C GPU1: 88.7% 11.1GiB 1377.0MHz 49.5C GPU2: 88.8% 11.0GiB 1376.0MHz 51.1C GPU3: 88.3% 11.0GiB 1374.0MHz 49.9C Interconnect group: 10	29 GPU0: 84.7% 11.0GiB 1364.0MHz 49.4C GPU1: 85.6% 11.1GiB 1366.0MHz 49.6C GPU2: 84.9% 11.0GiB 1365.0MHz 48.7C GPU3: 85.9% 11.0GiB 1367.0MHz 49.5C Interconnect group: 10	30 GPU0: 93.2% 11.0GiB 1389.0MHz 50.1C GPU1: 93.2% 11.1GiB 1389.0MHz 50.3C GPU2: 93.1% 11.0GiB 1390.0MHz 50.8C GPU3: 94.1% 11.0GiB 1394.0MHz 51.4C Interconnect group: 12
31 GPU0: 92.7% 11.0GiB 1386.0MHz 48.9C GPU1: 92.8% 11.1GiB 1386.0MHz 50.0C GPU2: 93.0% 11.0GiB 1387.0MHz 49.7C GPU3: 93.2% 11.0GiB 1387.0MHz 49.6C Interconnect group: 12	32 GPU0: 91.8% 11.0GiB 1385.0MHz 49.6C GPU1: 92.6% 11.1GiB 1388.0MHz 48.9C GPU2: 92.5% 11.0GiB 1388.0MHz 49.7C GPU3: 92.9% 11.0GiB 1389.0MHz 49.1C Interconnect group: 12	33 GPU0: 92.9% 11.0GiB 1387.0MHz 50.2C GPU1: 93.4% 11.1GiB 1389.0MHz 50.1C GPU2: 93.6% 11.0GiB 1391.0MHz 50.0C GPU3: 93.4% 11.0GiB 1390.0MHz 48.9C Interconnect group: 12	34 GPU0: 93.2% 11.0GiB 1389.0MHz 49.8C GPU1: 93.6% 11.1GiB 1390.0MHz 50.1C GPU2: 93.8% 11.0GiB 1391.0MHz 49.8C GPU3: 94.0% 11.0GiB 1393.0MHz 52.0C Interconnect group: 15	35 GPU0: 92.4% 11.0GiB 1385.0MHz 49.7C GPU1: 93.0% 11.1GiB 1387.0MHz 49.8C GPU2: 93.2% 11.0GiB 1385.0MHz 49.6C GPU3: 92.4% 11.0GiB 1385.0MHz 49.0C Interconnect group: 15	36 GPU0: 92.7% 11.0GiB 1386.0MHz 49.9C GPU1: 93.1% 11.1GiB 1387.0MHz 49.9C GPU2: 93.4% 11.0GiB 1388.0MHz 50.0C GPU3: 94.0% 11.0GiB 1390.0MHz 49.4C Interconnect group: 15
37 GPU0: 92.7% 11.0GiB 1387.0MHz 50.1C GPU1: 92.6% 11.1GiB 1388.0MHz 48.8C GPU2: 93.7% 11.0GiB 1391.0MHz 49.7C GPU3: 93.4% 11.0GiB 1390.0MHz 50.2C Interconnect group: 15	38 GPU0: 93.0% 11.0GiB 1388.0MHz 49.1C GPU1: 93.0% 11.1GiB 1389.0MHz 49.6C GPU2: 92.5% 11.0GiB 1390.0MHz 49.2C GPU3: 93.2% 11.0GiB 1390.0MHz 49.4C Interconnect group: 17	39 GPU0: 92.3% 11.0GiB 1385.0MHz 51.4C GPU1: 92.1% 11.1GiB 1386.0MHz 50.1C GPU2: 92.5% 11.0GiB 1387.0MHz 51.0C GPU3: 93.5% 11.0GiB 1391.0MHz 50.1C Interconnect group: 17	40 GPU0: 93.8% 11.0GiB 1389.0MHz 48.7C GPU1: 93.4% 11.0GiB 1388.0MHz 48.6C GPU2: 93.6% 11.0GiB 1389.0MHz 50.7C GPU3: 94.0% 11.0GiB 1389.0MHz 49.8C Interconnect group: 17	41 GPU0: 92.4% 11.0GiB 1384.0MHz 50.2C GPU1: 93.6% 11.1GiB 1391.0MHz 50.8C GPU2: 93.2% 11.0GiB 1390.0MHz 51.2C GPU3: 92.5% 11.0GiB 1387.0MHz 50.4C Interconnect group: 17	42 GPU0: 92.9% 11.0GiB 1387.0MHz 49.0C GPU1: 93.1% 11.1GiB 1387.0MHz 50.1C GPU2: 93.3% 11.0GiB 1388.0MHz 50.0C GPU3: 93.3% 11.0GiB 1388.0MHz 49.2C Interconnect group: 17
43 GPU0: 93.4% 11.0GiB 1387.0MHz 49.4C GPU1: 92.7% 11.1GiB 1388.0MHz 49.7C GPU2: 93.6% 11.0GiB 1391.0MHz 50.4C GPU3: 93.1% 11.0GiB 1389.0MHz 49.5C Interconnect group: 17	44 GPU0: 93.2% 11.0GiB 1389.0MHz 50.7C GPU1: 92.2% 11.1GiB 1387.0MHz 48.7C GPU2: 92.2% 11.1GiB 1388.0MHz 50.1C GPU3: 92.9% 11.0GiB 1391.0MHz 48.8C Interconnect group: 17	45 GPU0: 93.3% 11.0GiB 1389.0MHz 50.1C GPU1: 94.5% 11.1GiB 1393.0MHz 51.3C GPU2: 94.5% 11.0GiB 1392.0MHz 50.1C GPU3: 94.1% 11.0GiB 1392.0MHz 50.3C Interconnect group: 18	46 GPU0: 92.7% 11.0GiB 1386.0MHz 50.4C GPU1: 93.0% 11.1GiB 1387.0MHz 50.8C GPU2: 93.2% 11.0GiB 1387.0MHz 50.0C GPU3: 93.1% 11.0GiB 1387.0MHz 49.6C Interconnect group: 18	47 GPU0: 93.5% 11.0GiB 1388.0MHz 49.8C GPU1: 93.7% 11.1GiB 1389.0MHz 50.0C GPU2: 93.0% 11.0GiB 1387.0MHz 49.4C GPU3: 93.2% 11.0GiB 1388.0MHz 49.3C Interconnect group: 18	48 GPU0: 92.4% 11.0GiB 1385.0MHz 50.8C GPU1: 92.9% 11.0GiB 1386.0MHz 50.9C GPU2: 92.8% 11.0GiB 1386.0MHz 50.1C GPU3: 92.8% 11.0GiB 1386.0MHz 52.4C Interconnect group: 18

Other tools?

- LLView by Jülich Research Center
<https://github.com/FZJ-JSC/LLview>
 - Sophisticated system level breakdown
 - Not for shared nodes
 - No hardware performance counters
- Different choices available, pick your solution

PathoJobs

Patho-Jobs

- Goal:
 - Automated rule-based detection system
 - Detection of pathological HPC jobs
 - jobs with previously encountered inefficiencies
 - Target audience: cluster operators and users

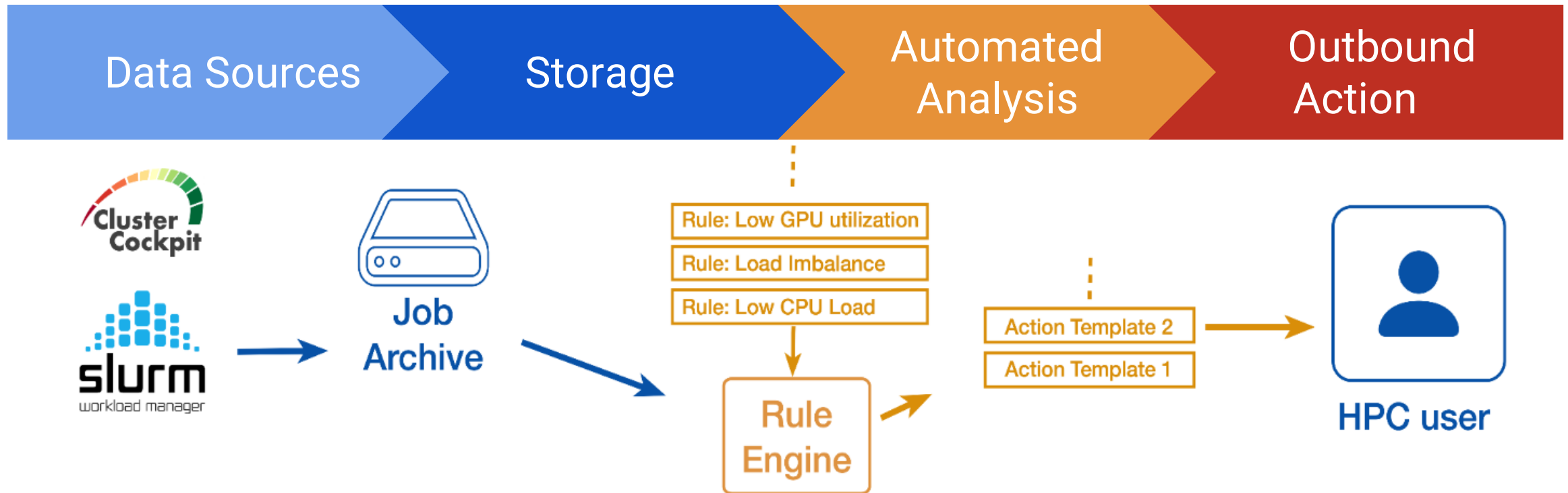
Collaboration with:



Paderborn Center for
Parallel Computing



<https://git-ce.rwth-aachen.de/pathojobs/>



- Cluster Cockpit generates a job-archive for each job;
 - alternative sources experimented with, but not ready to use
- Rule-Engine applies anti-pattern and triggers appropriate response for each job
- Detected anti-patterns trigger action templates (Email or job-archive tags)

PathoJobs: Prototype antipatterns

- Low CPU load
- Load-Imbalance
- Job Payload Overhead
- Excessive CPU load
- Problematic memory usage
- Uncoordinated multi-process GPU-usage
- Low GPU utilization
- Unnecessary job distribution
- Unnecessary PFS read

Patho-Jobs: **Low-Load** Pattern

```

{
  "name": "Low CPU load",
  "tag": "lowload",
  "parameters": ["threshold_factor"],
  "metrics": ["cpu_load, job"],
  "terms": [
    { "load_mean": "cpu_load.mean('all')"},
    { "load_threshold": "job.numHwthreads * threshold_factor"},
    { "lowload_nodes": "load_mean < load_threshold"},
    { "lowload": "lowload_nodes.any('all')"},
    { "load_perc": "1.0 - (load_mean / load_threshold)"}],
  "output": "lowload",
  "output_scalar": "load_perc",
  "template": "Job ({{job.jobId}}) was detected as the mean cpu
load {{load_mean}} falls below {{load_threshold}}."
}

```

Tag for use in ClusterCockpit

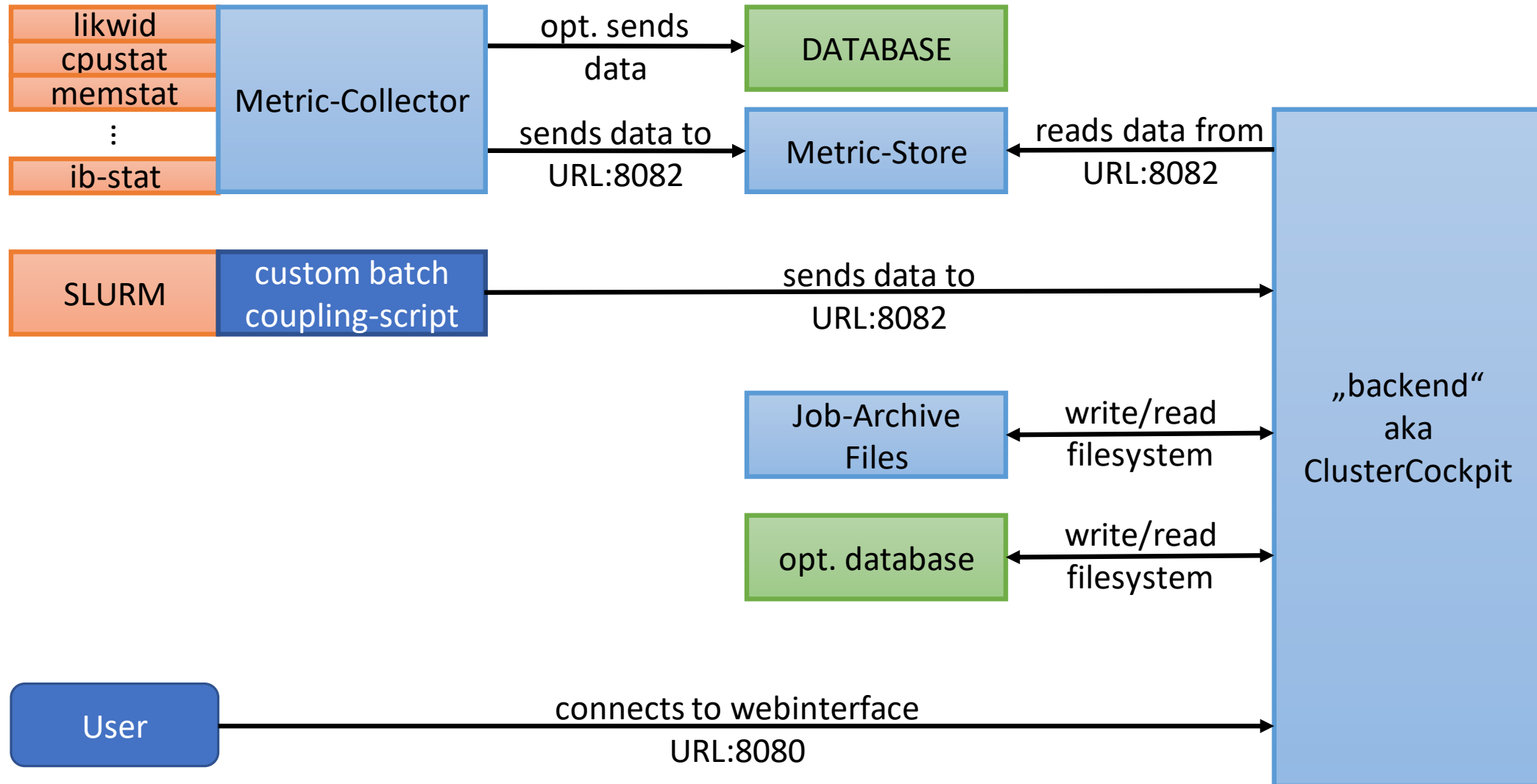
External configurable threshold

Which datapoints to pull from job-archive

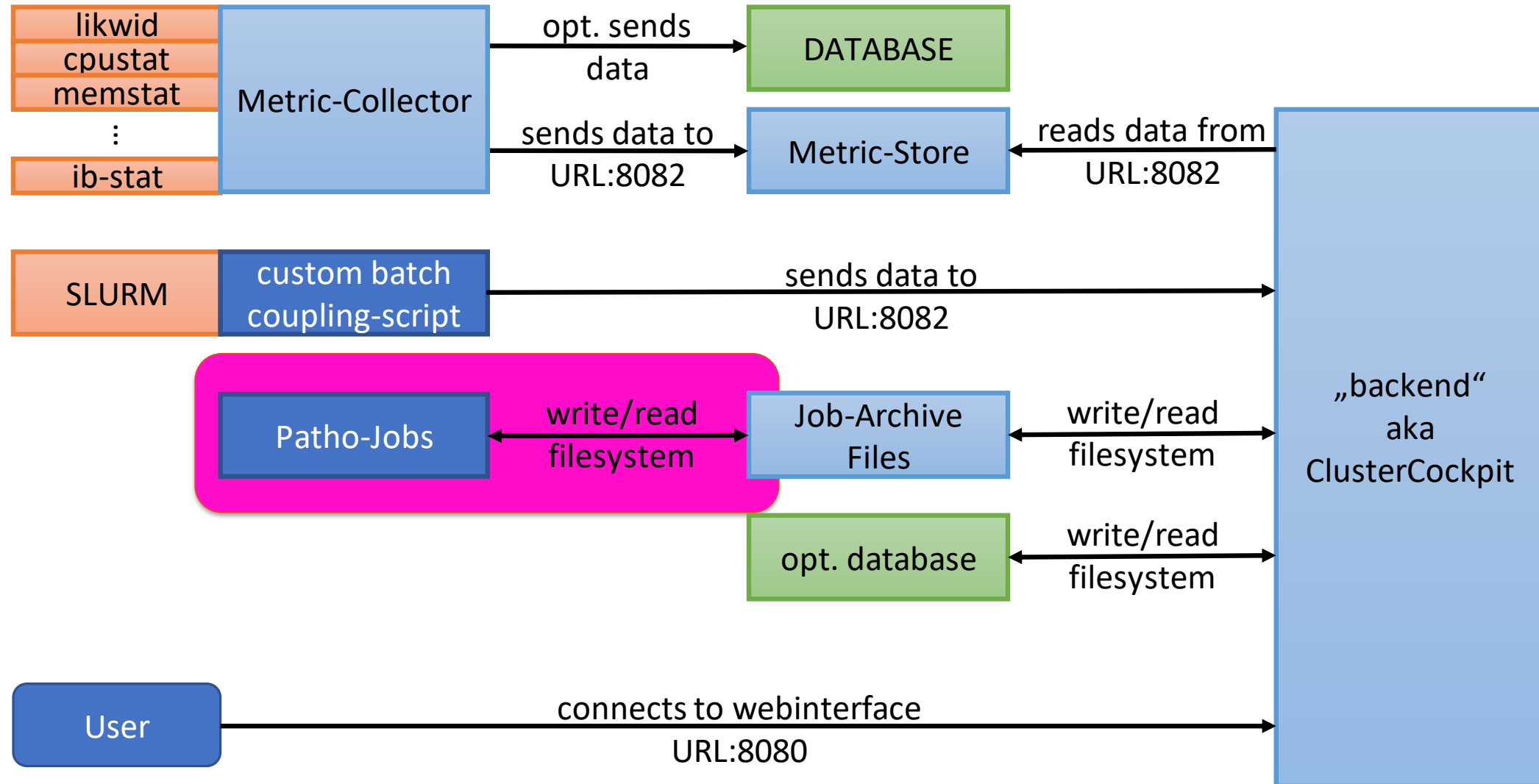
Variables initialized with in sequence evaluation of terms

Pattern output, variables can be used

Classic CusterCockpit:



Patho-Jobs extension:



- **Deployment:**
 - The “PathoJobs” system is deployed by partners on production HPC systems.
 - Each processes job-archives from hundred-thousands to millions of jobs per month.
 - Current performance anti-patterns are evaluated for all HPC jobs.
- **First Experience & Insights:**
 - Rule evaluation does not exhibit noticeable extra-load on the Cluster Cockpit hosts
 - Number of jobs with detected inefficiencies exceeds acceptable rates
 - System is able to detect explicitly malformed test-jobs

Location	Jobs / month	Jobs with issues	Rules with highest hit-rate	Analysis cost
TU Darmstadt	~ 250k	598 / 1000	Low CPU load Load-imbalance	230 ms / job
Paderborn	~ 190k	180 / 1000	Low CPU load CPU oversubscription	560 ms / job
FAU	~ 255k	326 / 1000	Low CPU load Low resource utilization	n.a.

Experience ...

- Software can be directly deployed from binaries:
 - Did not attempt to compile on HPC system
 - 3 types of services:
 - metric-collector per compute node,
 - at least one metric store
 - backend/UI service
- Basic setup:
 - **cluster configuration**: specification of the cluster for the backend; **tools help**
 - **metric collector**: select which metrics are of interest, **challenging**
 - managing CPU registers, trade-off available metric sources, manage access privileges
 - **metric store**: holds metrics; must match with metric collector and backend; ascertaining peak capacity is not trivial

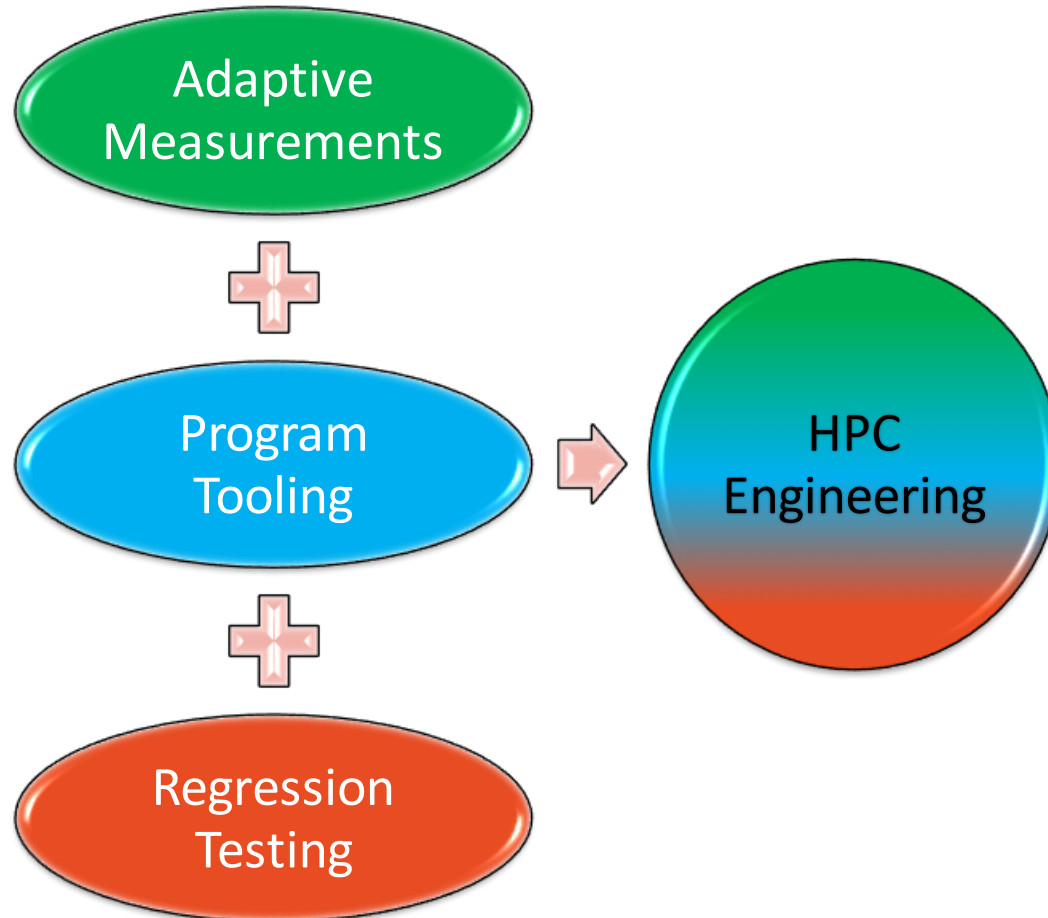
Why is that DevOps?

DevOps principle		How It applies in this talk
Observability		Job-level monitoring with ClusterCockpit and hardware counters
Automation		PathoJobs automatically detects inefficient jobs using rule engines
Feedback Loops		Tagged jobs, performance summaries, and user-facing metrics
Scalability		Designed for 10,000+ jobs/day with minimal overhead
Collaboration		Shared infrastructure for admins, users, and support staff
Open Tooling		Open-source, API-driven, SLURM-integrated - DevOps-ready

Key Takeaways – Boosting HPC with Job-Specific Monitoring

- **HPC systems require more than just system monitoring**
→ Insightful resource usage must happen at the job level, not just the node.
- **Instrumentation is powerful, but doesn't scale**
→ We need lightweight, non-invasive observability for real-world workloads.
- **ClusterCockpit enables scalable job-aware monitoring**
→ Tied to users, jobs, and projects — not just hosts and hardware.
- **PathoJobs adds automation and rule-based diagnosis**
→ Proactive detection of inefficiencies from 10,000+ jobs/day.
- **HPC DevOps is not just CI/CD — it's feedback, insight, and collaboration**
→ Observability and automation must become part of everyday HPC operations.

Where are things going ...



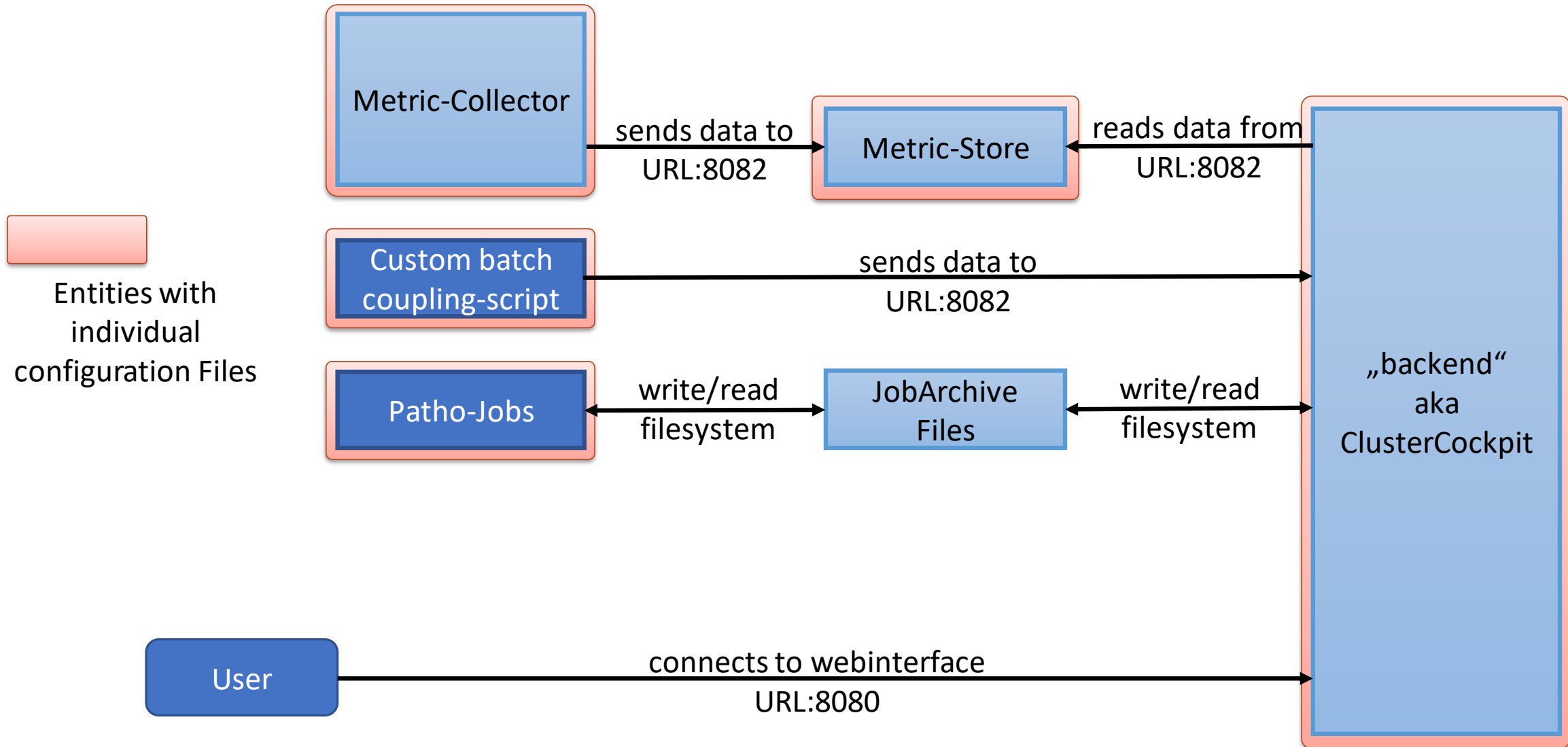
- Adaptive measurements
 - InstRO, CAPI, XRAY
 - All applications come pre-tooled for performance analysis, but inactive
 - When more data is required, it detailed measurements are seamless activated
- Program tooling
 - CALIPER style integration
 - Application reports what it is doing
- Performance regressions tests
 - Integrate „self reflection“ into programs and communicate results to external viewers
 - Keyword: Performance assertions/active program documentation
 - Extend PathoJobs approach

Key Takeaways – Boosting HPC with Job-Specific Monitoring

- **HPC systems require more than just system monitoring**
- **Instrumentation is powerful, but doesn't scale**
- **ClusterCockpit enables scalable job-aware monitoring**
- **PathoJobs adds automation and rule-based diagnosis**
- **HPC DevOps is not just CI/CD — it's feedback, insight, and collaboration**

Questions?

Individual configuration files:



Metric Collector

- Collector config:

```
{  
  "sinks": "sinks-http.json",  
  "collectors" : "collectors-CLX.json",  
  "receivers" : "receivers.json",  
  "router" : "router.json",  
  "interval": "60s",  
  "duration": "1s"  
}
```

Example

- Collectors:

```
View Go Run Terminal Help  
1 {  
2   "memstat": { ...  
3 },  
4  
5 "cpustat": { ...  
6 },  
7  
8 "iostat": { ...  
9 },  
10  
11 "loadavg": { ...  
12 },  
13  
14 "netstat": { ...  
15 },  
16  
17 "ibstat": { ...  
18 },  
19  
20 "cpufreq": { ...  
21 },  
22  
23 "numastats": {},  
24 "schedstat": {},  
25  
26 "tempstat": { ...  
27 },  
28  
29 "likwid": {  
30   "force_overwrite": true,  
31   "invalid_to_zero": true,  
32   "access_mode": "accessdaemon",  
33   "accessdaemon_path": "/usr/sbin/",  
34   "liblikwid_path": "/usr/lib64/liblikwid.so",  
35   "old_liblikwid_path": "/usr/lib64/liblikwid.so.5.4",  
36   "o_lockfile_path": "/tmp/likwid.lock",  
37   "lockfile_path": "/run/cc-metric-collector/likwid.lock",  
38   "eventsets": {  
39     "events": {  
40       "FIXC0": "INSTR_RETIRED_ANY",  
41       "FIXC1": "CPU_CLK_UNHALTED_CORE",  
42       "FIXC2": "CPU_CLK_UNHALTED_REF",  
43       "HBOXBC0": "CAS_COUNT_RD",  
44       "HBOXBC1": "CAS_COUNT_WR",  
45       "HBOXBC2": "CAS_COUNT_RD",  
46     }  
47   }  
48 }  
49 }
```

Example

Metric-Store

■ Metric-Store config:

```
{
  "metrics": {
    "clock": {"frequency": 60, "aggregation": "avg"},
    "cpu_user": {"frequency": 60, "aggregation": "avg"},
    "acc_utilization": {"frequency": 60, "aggregation": "avg"},
    "acc_mem_used": {"frequency": 60, "aggregation": "sum"},
    "mem_used": {"frequency": 60, "aggregation": null }
  },
  "checkpoints": {
    "interval": "12h", "directory": "/work/home/hrzcc/var/metric-store/checkpoints",
    "restore": "168h"
  },
  "archive": {
    "interval": "168h",
    "directory": "/work/home/hrzcc/var/metric-store/archive"
  },
  "http-api": {
    "address": "0.0.0.0:8082",
    "https-cert-file": null,
    "https-key-file": null
  },
  "retention-in-memory": "168h",
  "nats": null,
  "jwt-public-key": "use-your-own-key"
}
```

Example

Backend

2+ Config-files:

- Cluster-configurations:
 - Helper:
generate-subcluster.pl
- Backend-configuration:

Cluster-config

```

1  {
2      "name": "Lichtenberg2",
3  >  "metricConfig": [ ...
96      ],
97      "subClusters": [
98 >  { ...
145      },
146      { ...
147      },
148  ]
149 }
```

Example

Backend-config

```

1  {
2      "addr": "0.0.0.0:8080",
3  >  "archive": { ...
6      },
7      "db": "/work/home/hrzcc/var/backend/job.db",
8  >  "jwts": { ...
10     },
11     "https-cert-file": "/home/hrzcc/etc/cc-ui+ms-HARICA.pem",
12     "https-key-file": "/home/hrzcc/etc/privkey",
13 >  "ldap": { ...
19     },
20     "clusters": [
21         {
22             "name": "Lichtenberg2",
23             "metricDataRepository": {
24                 "kind": "cc-metric-store",
25                 "url": "http://cc-ms:8082",
26                 "token": "use-your-own-jwt-token"
27             },
28             "filterRanges": { ...
30             }
31         },
32         {
33             "name": "Lichtenberg2-Login",
34             "metricDataRepository": { ...
36             },
37             "filterRanges": { ...
39             }
40         }
41     ],
42     "filterRanges": { ...
44     }
45 }
46 }
```

Example