

WorkVisualizer

**A User Support Tool to Facilitate General Computing
Using Heterogeneous Architectures at Scale**

Pierre Pébay - Barcelona Supercomputing Center (BSC-CNS) - NexGen Analytics (NGA)

Bio

- MSc ISIMA, France
- Currently: PhD at Barcelona Supercomputing Center (BSC-CNS) in computational fluid dynamics
- Worked three years at NexGen Analytics
- Latest publications: ML-assisted decentralized load-balancing algorithms



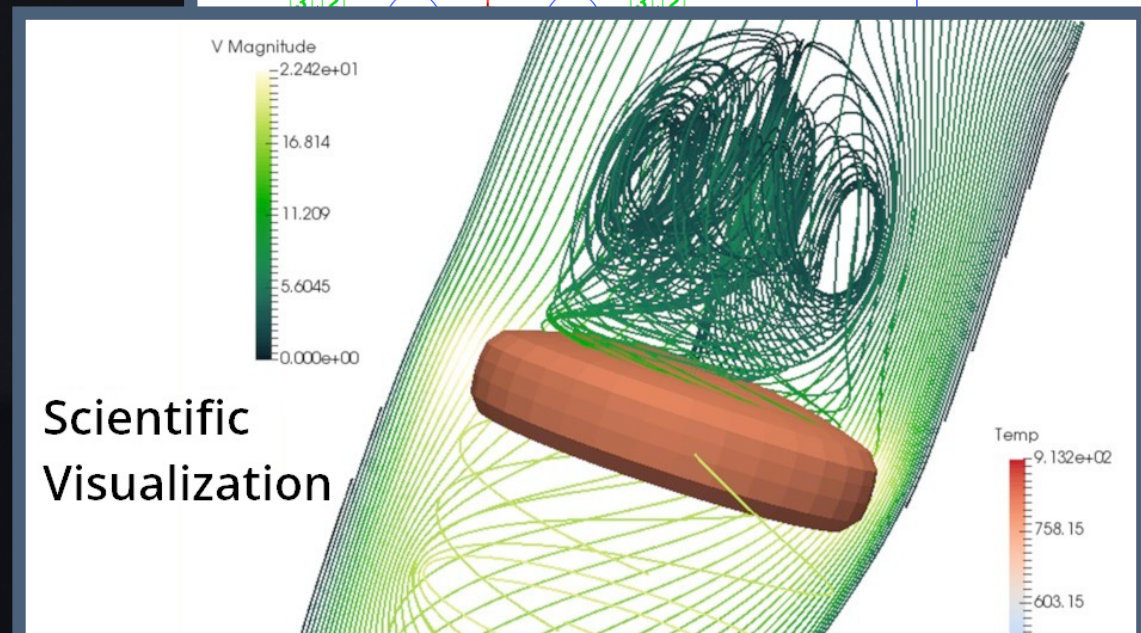
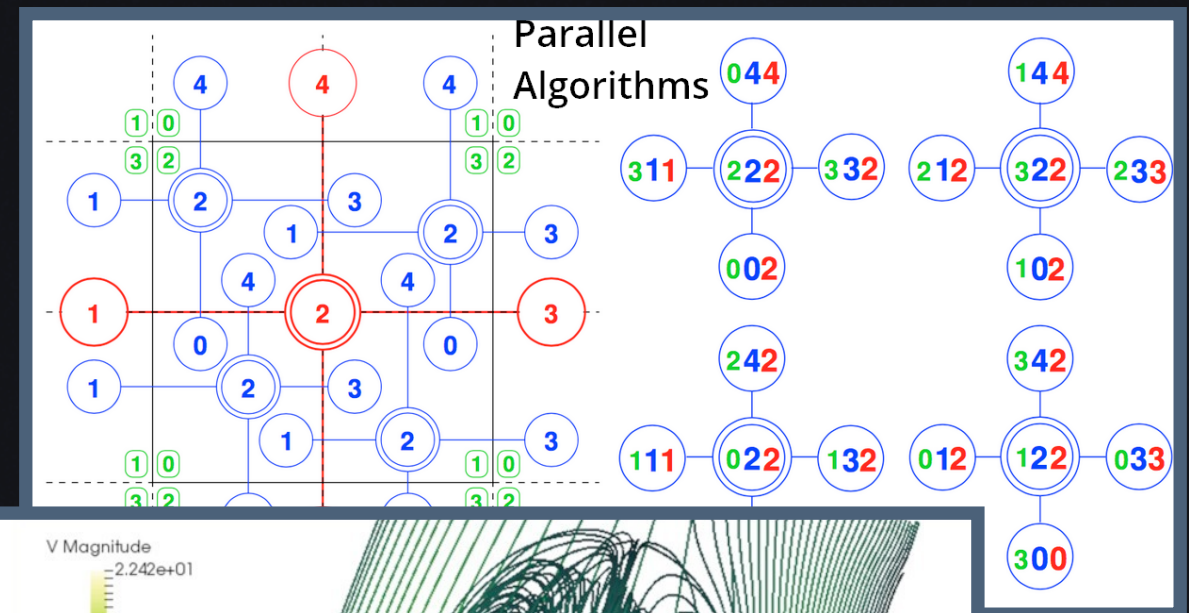


NEXGEN ANALYTICS

- Office in North Carolina, USA
- Mission: support R&D projects
- Private and public clients
- Partnership with Sandia National Laboratories
- OVIS/LDMS

Our HPC simulation background

- Multiple physics simulation codes: electromagnetics, solid mechanics, fluid dynamics
- Parallel programming paradigms, large-scale HPC
- Scalable load-balancing

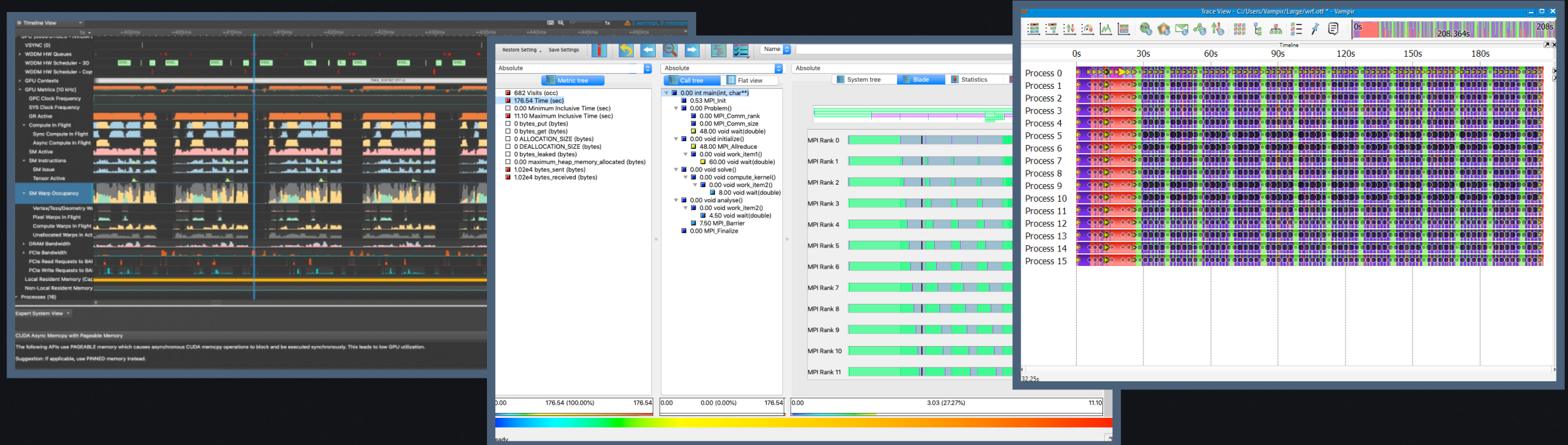


$$\begin{aligned}
 \tilde{M}_p &= \sum_{i=1}^{n_{\mathcal{A}}} (x_i - \tilde{x}_W)^p + \sum_{i=n_{\mathcal{A}}+1}^n (x_i - \tilde{x}_W)^p && \text{Computational Statistics} \\
 &= \sum_{i=1}^{n_{\mathcal{A}}} \left(x_i - \tilde{x}_{W,\mathcal{A}} - \frac{W_{\mathcal{B}}}{W} \delta_{\mathcal{B},\mathcal{A}} \right)^p + \sum_{i=n_{\mathcal{A}}+1}^n \left(x_i - \tilde{x}_{W,\mathcal{B}} + \frac{W_{\mathcal{A}}}{W} \delta_{\mathcal{B},\mathcal{A}} \right)^p \\
 &= \sum_{k=0}^p \binom{p}{k} \delta_{\mathcal{B},\mathcal{A}}^k \left[\tilde{M}_{p-k}^{\mathcal{A}} \left(\frac{-W_{\mathcal{B}}}{W} \right)^k + \tilde{M}_{p-k}^{\mathcal{B}} \left(\frac{W_{\mathcal{A}}}{W} \right)^k \right]
 \end{aligned}$$

New to fusion – Opportunity

- Fusion is new to us - we are learning
- Here to learn
- Show a potentially interesting tool
- Share technical approach to handling large-scale data and creating visually pleasing tools

Market gap: profiling/tooling in HPC



- Why existing tools may fall short: complexity, steep learning curve, fragmented workflows
- No analysis
- Gap in the profiler/tooling market: easy to use, modern, quick, low additional work, helps improve HPC code performance

SBIR Funding Context

Department Of Energy Small Business Innovation Research

- Awarded **DOE SBIR Phase I** last year to build an MVP and demonstrate feasibility
- Successfully completed Phase I: validated approach, proved technical viability
- Recently awarded **Phase II** (just last month) → major expansion of scope



U.S. DEPARTMENT
of **ENERGY**

Office of
Science

Context: What Makes HPC Simulations Hard to Run Efficiently?

- Clusters: hundreds–thousands of nodes working together
- MPI + parallel programming: Nodes exchange data constantly → communication patterns matter
- Imbalance is common: If one rank slows down, *everyone* waits
- Heterogeneous hardware: Mix of CPUs / GPUs → different speeds, different bottlenecks
- Data volume is massive: Traces generate millions of events across thousands of ranks

Phase I objectives (MVP, feasibility)

- **Collection Pipeline:** low-friction data capture with zero source-code changes
- **Analysis Pipeline:** turn large datasets into actionable insights
- **Visualization Layer:** fast, interactive, high-level insights

WorkVisualizer Phase I In Practice

Call Tree

Select visualization

▼

?

Summary

Analysis

Settings

Program Information

▼

Comm Size

250

Program Runtime

62.4077 s

Call Distribution

🖱️

Call Type	Average Counts per Rank	Unique Counts
Program	42007	50
MPI Collective	9038	7
MPI Point-To-Point	43581	26
Kokkos	20072	32

Largest Calls

<

Select a visualization from the dropdown menu.

WorkVisualizer Phase I In Practice

Call Tree

Select visualization

?

Summary

Analysis

Settings

Program Information

Comm Size250

Program Runtime62.4077 s

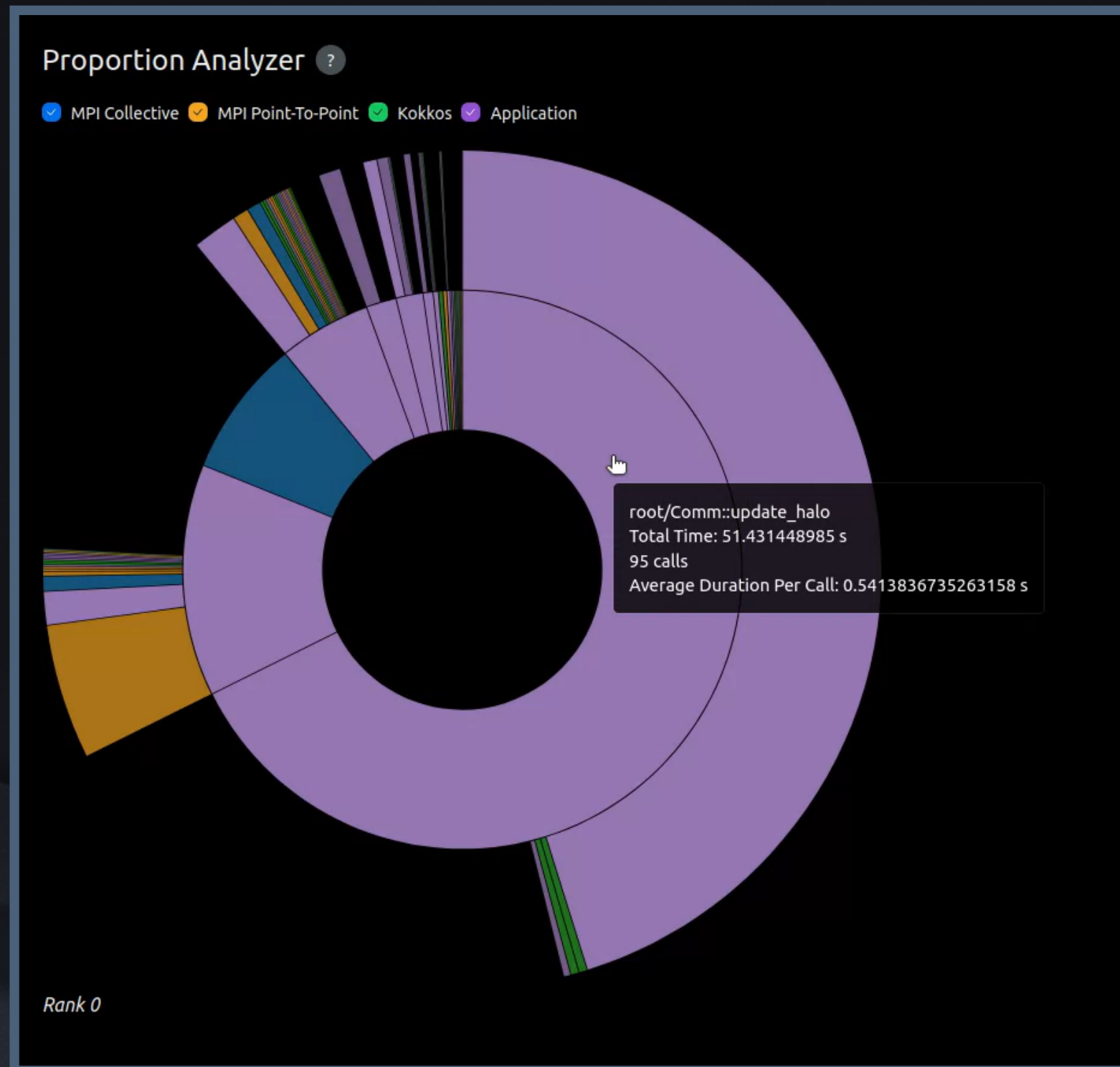
Call Distribution

Largest Calls

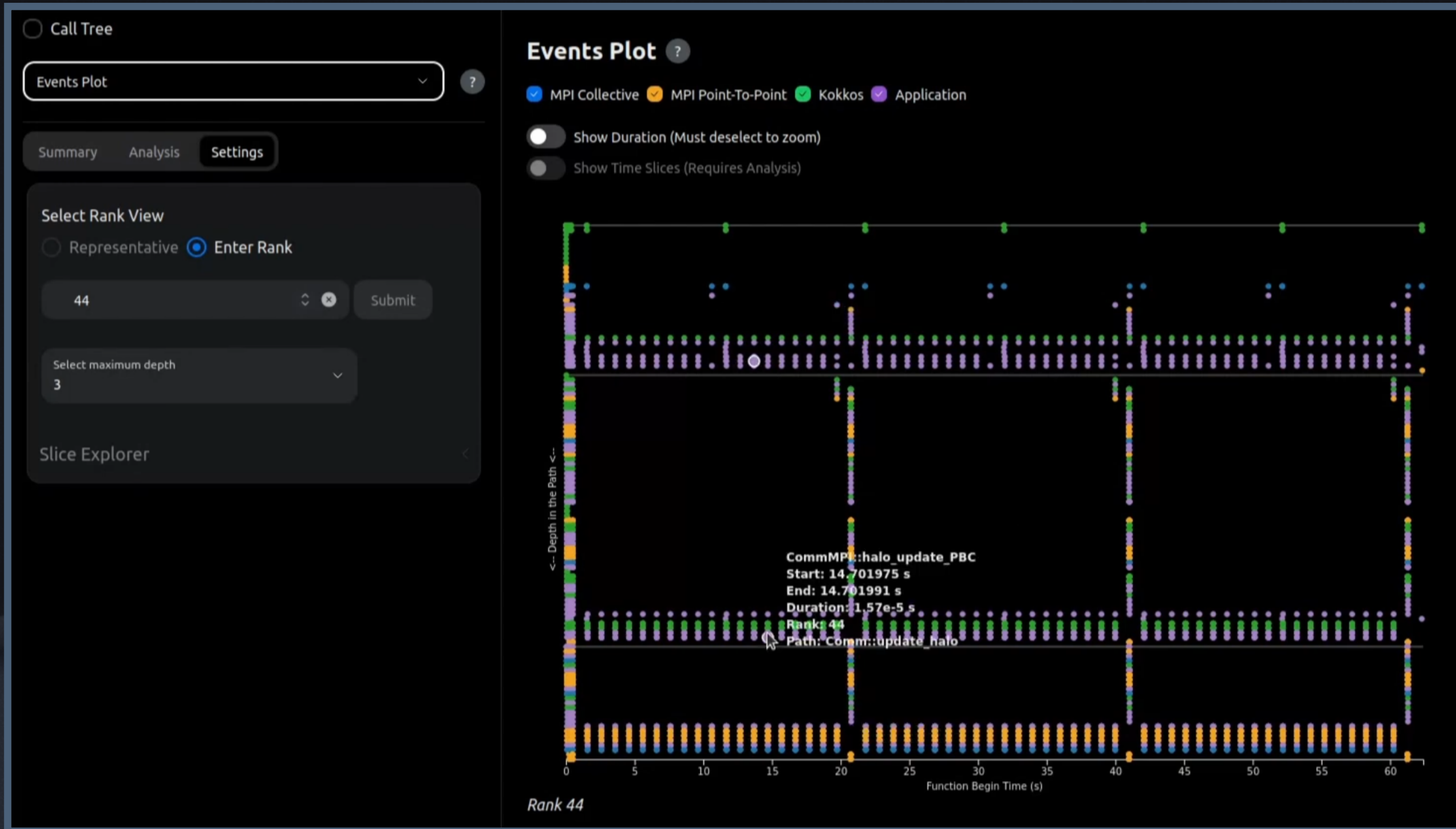
Function	Average Time
Comm::update_halo	51.380871 s
Cabana::gather	51.286209 s
MPI_Barrier	50.430176 s
MPI_Allreduce	6.995763 s
Comm::exchange	3.086395 s
MPI_Comm_dup	3.002727 s
Comm::exchange_halo	0.147836 s
MPI_Send	0.318127 s
Cabana::neighbor_parallel_for	0.082310 s
ForceLJCabanaNeigh::compute_full	0.079758 s
MPI_Waitall	5.187821 s

Select a visualization from the dropdown menu.

WorkVisualizer Phase I In Practice



WorkVisualizer Phase I In Practice



WorkVisualizer Phase I In Practice

☐ Call Tree

Analysis Viewer

Summary

Analysis

Settings

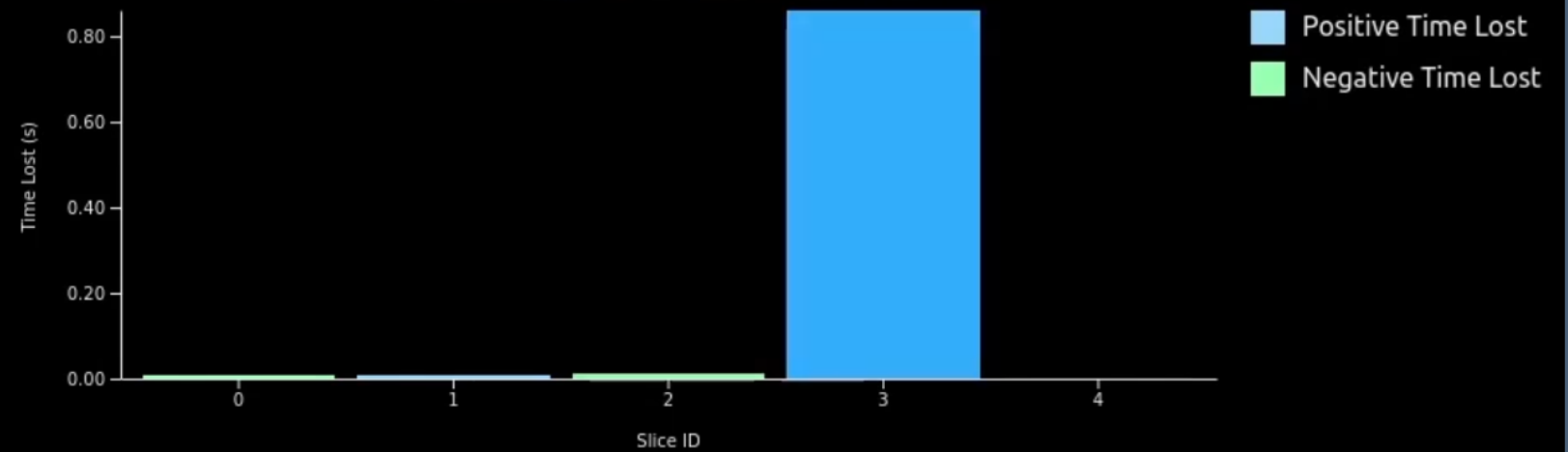
Analysis Results

# of Ranks with Significant Time Lost	4
Most Time-Losing Slice	3
Most Time-Losing Rank	Rank 4 on Slice 3
Longest Slice	3
% of Total Runtime Lost on Rank 4	93.23%
% of Total Runtime Lost Across All Ranks	1.34%

Tip: Visualize these results with the Analysis Viewer in the visualizations dropdown.

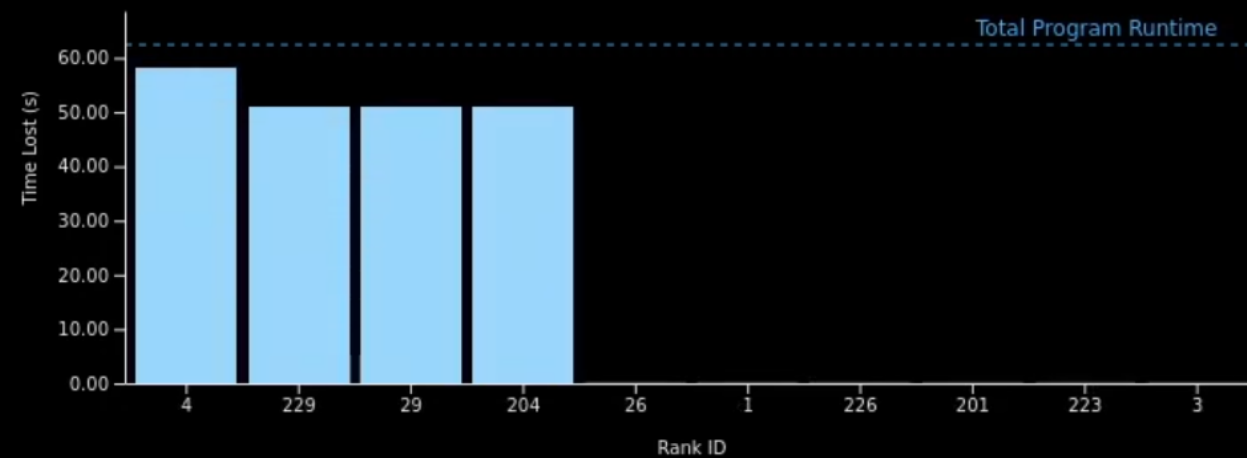
Analysis Viewer

Time Lost Per Slice Averaged Over All Ranks



☒ Show Time Lost Per Slice

Time Lost Per Rank Per Slice



Slice 3 Duration: 61.92 s (0.42 - 62.34)

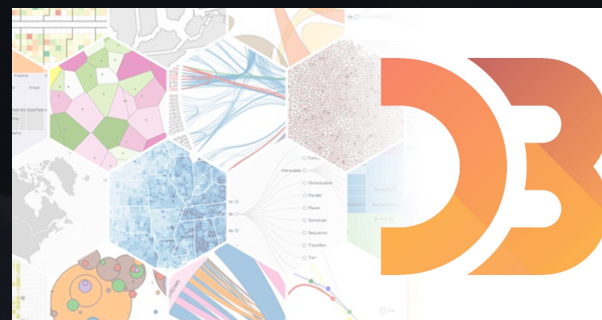
Tech stack

- Backend: Python + FastAPI



- Frontend: Next.js + D3

NEXT.js



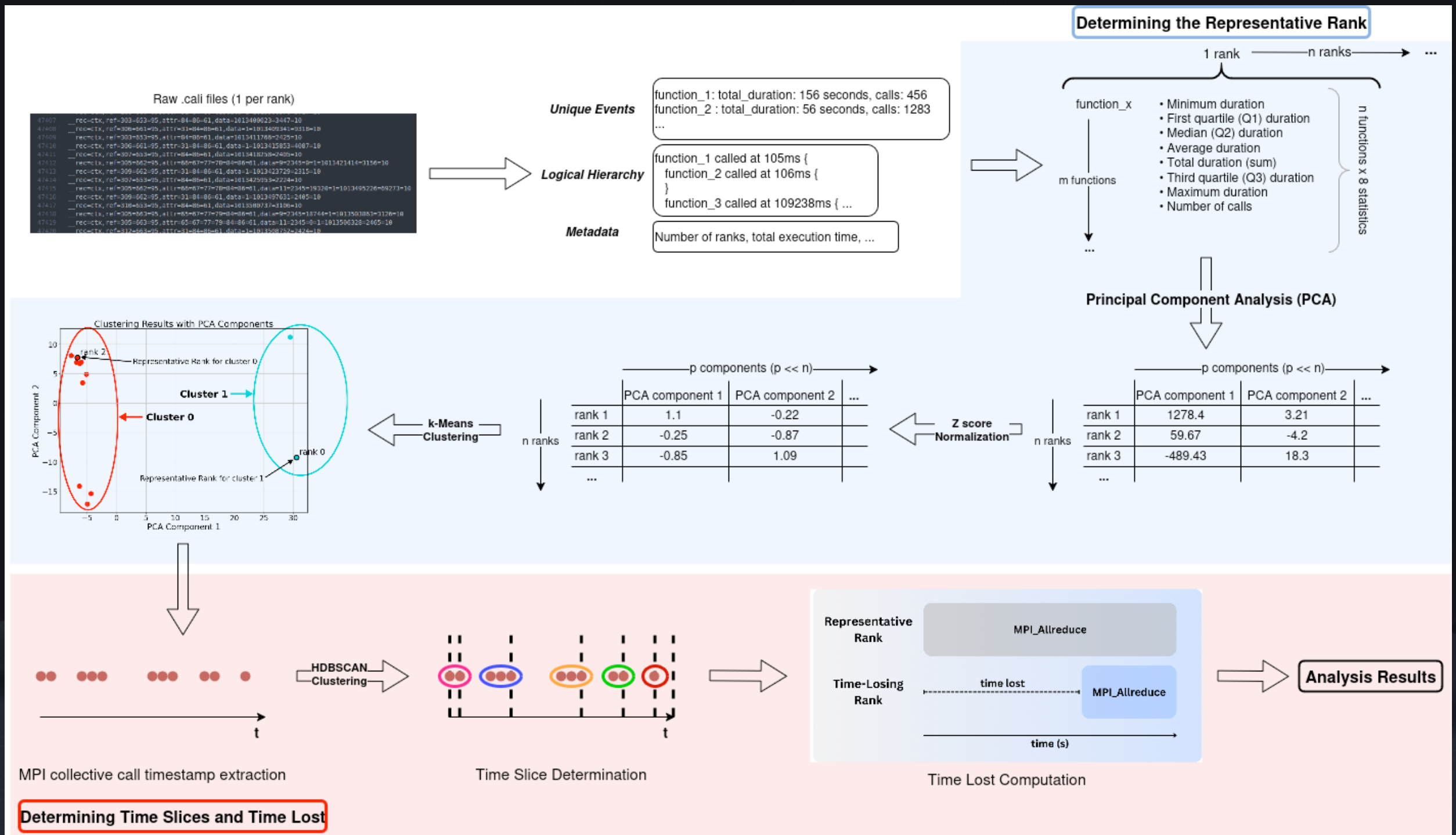
Data Collection Pipeline

- Built on **Caliper**, a DOE-funded performance analysis toolkit
 - Seamless integration with **Kokkos** for performance-portable MPI applications
- Zero source-code changes required: Users just load Caliper + our custom `caliper.config`
- Automatic capture of:
 - Kernel execution (start/end times, durations)
 - MPI communications (source, destination, size)
- Each run produces `data.<rank>.cali` files: one per MPI rank

Challenges: high-dimensional data & visualization

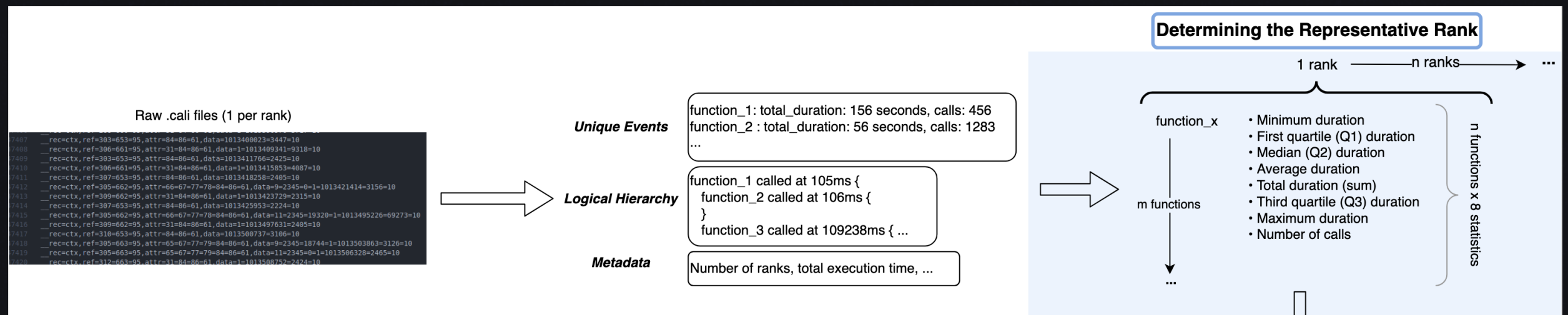
- MPI ranks, functions, call tree, time, communication
- Multiple well-chosen visualizations can capture a lot of information
- But: we need analysis

Analysis Pipeline



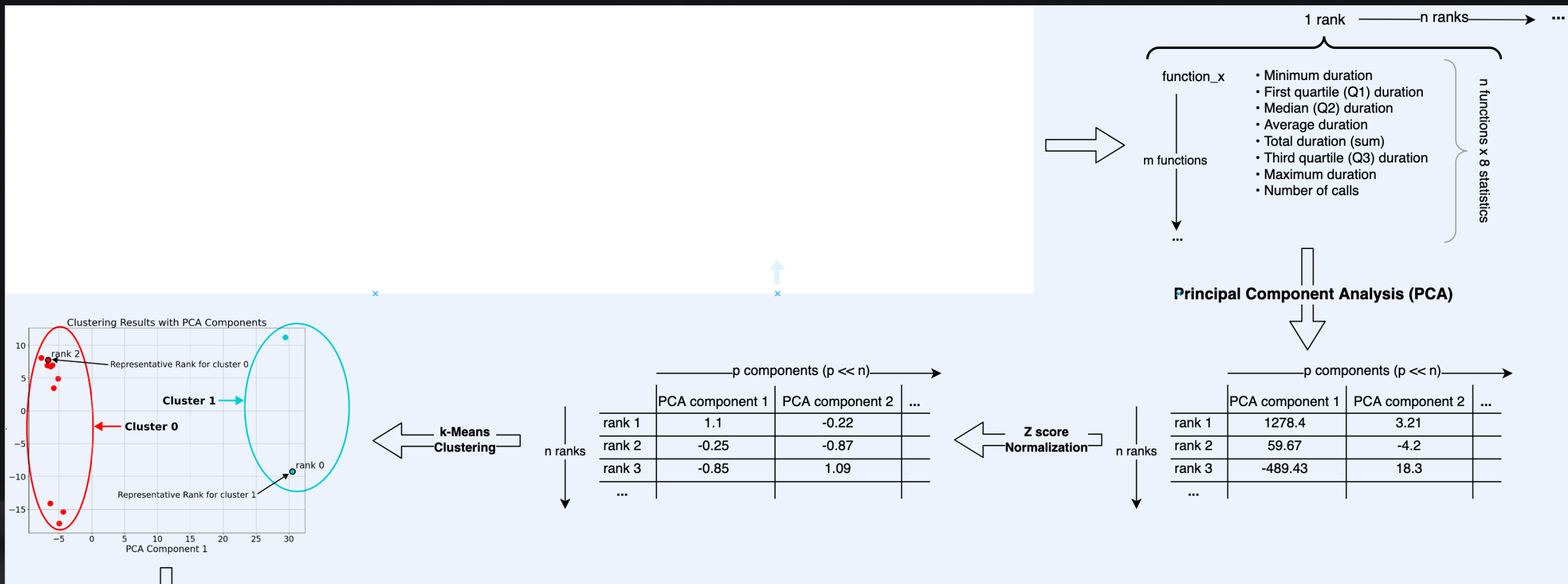
Analysis Pipeline

Preprocessing



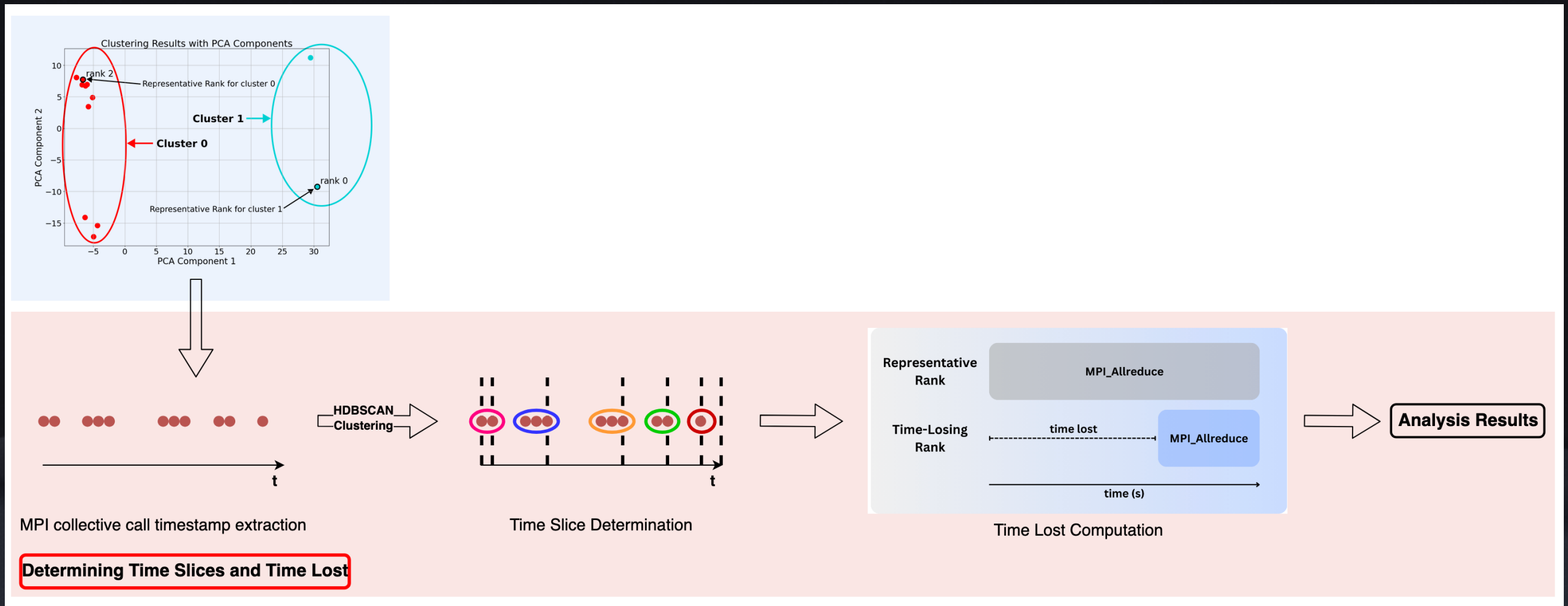
Analysis Pipeline

Determining the Representative Rank



Analysis Pipeline

Determining the Representative Rank



Phase I Takeaways

- Delivered full end-to-end MVP (open-source: <https://github.com/NexGenAnalytics/WorkVisualizer>)
- Proven low-overhead data collection with Caliper + Kokkos
- Novel analysis pipeline: rank clustering + time slicing + “time lost” metric
- Visualizations are fast, interactive, and actionable
- Validated on real HPC apps (e.g., CabanaMD) → successfully detected injected performance issues

Phase II: What's Next

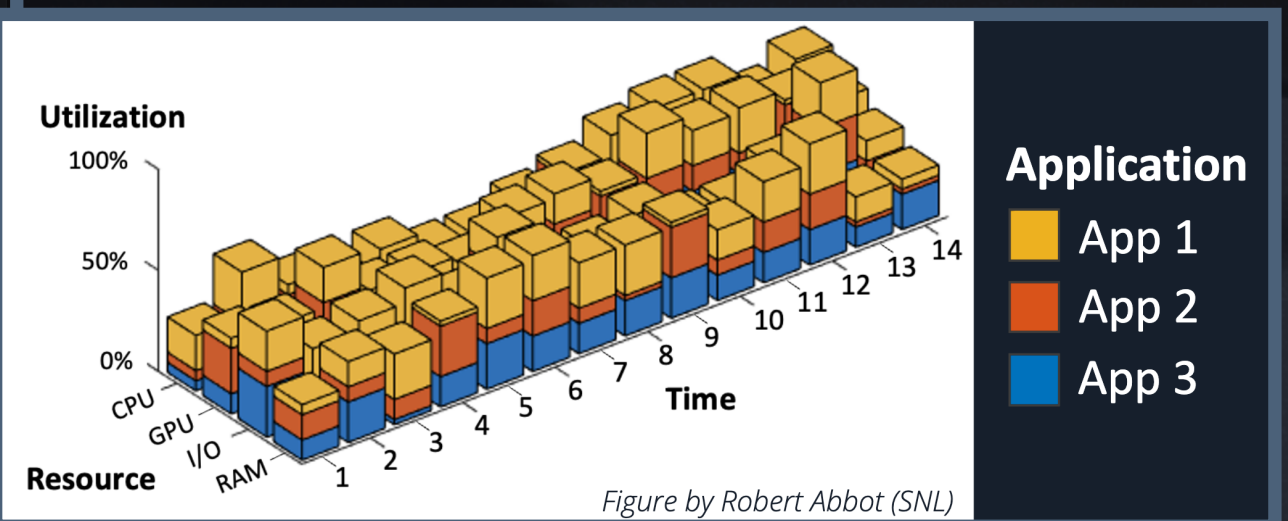
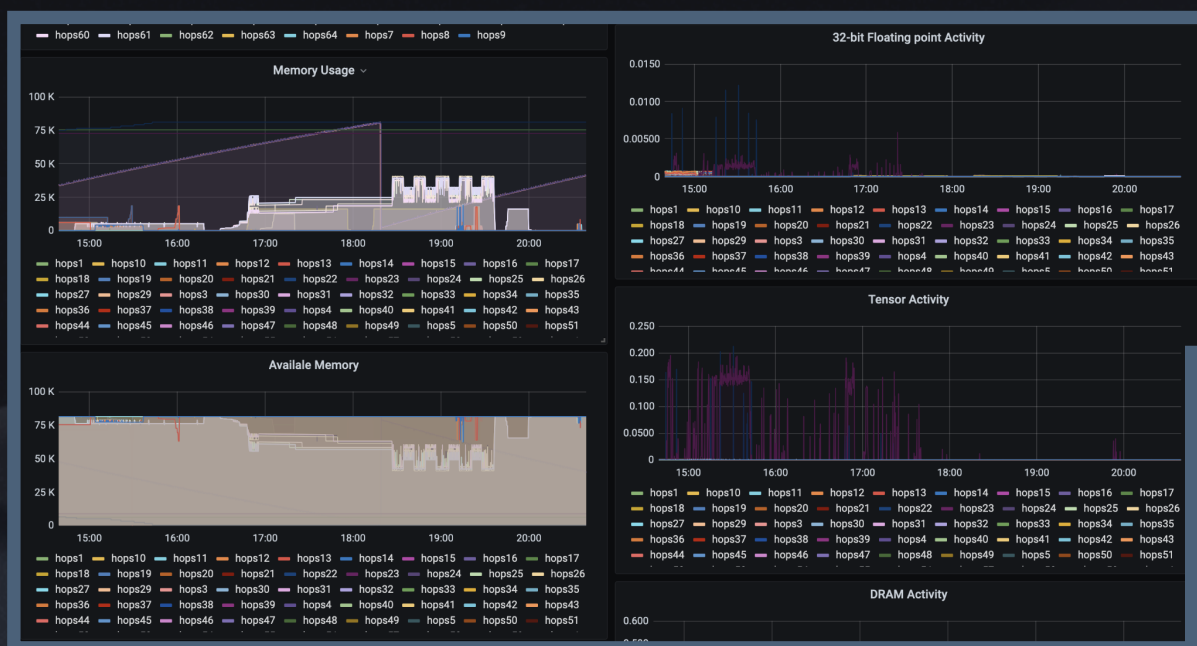
- Focus on scaling, speed, and usability
 - **Accelerate time-to-results** (faster ingestion, real-time interaction)
- Expand analysis capabilities & GPU support
 - Smarter detection of bottlenecks (communication, imbalance, transfers)
- Unified views: select a function → highlighted everywhere
- Move toward deployment + in-situ workflows
- Some work will remain open-source; advanced features become commercial

Phase II: Research Directions

- AI/ML methods for pattern & anomaly detection
- Generalized, improved time-slicing
- Real-time / *in-situ* WorkVisualizer
- OVIS/LDMS partnership

OVIS/LDMS Collaboration

- Framework for collecting, aggregating, and transporting system metrics in HPC environments
- Partnership: Predicting active memory usage, automating scheduling, visualizing the results



Summary & Q&A

- Phase I proved feasibility with a full open-source prototype
- Phase II focuses on speed, scalability, GPUs, and deeper analysis
- Aim: make HPC performance insights accessible & actionable
- Looking forward to engaging with fusion applications!