

Profiling and Resource Monitoring

Introduction

This documentation is designed to provide a comprehensive guide on monitoring and managing resources on our cluster. It covers the use of profiling tools, strategies for monitoring memory, and best practices for resource allocation to optimize job scheduling and performance. By following these guidelines, users can enhance their workflow efficiency and overall experience on the cluster.

Profiling Tools

Profiling tools are essential for understanding and optimizing the performance of your applications. Two key tools to mention are Valgrind and the seff script.

Valgrind is a powerful tool for profiling and debugging C/C++ applications, particularly for identifying memory issues. Although Valgrind is not currently installed on the cluster, users have the flexibility to install it under the /scratch directory, allowing them to leverage its capabilities for their specific needs.

The seff script is a custom tool developed in-house, inspired by Slurm's original command. It provides users with insights into their resource usage by comparing the resources allocated to a job with the actual resources consumed during its execution. To use the seff script, simply invoke it with the command `/cluster/slurm/seff <jobid>`, replacing `<jobid>` with the ID of your job. This tool is invaluable for optimizing resource utilization and ensuring that jobs run efficiently.

Monitoring Memory with Slurm

Monitoring memory usage is crucial for managing resources effectively. Slurm provides the MaxRSS metric, which measures the maximum resident set size of a job. However, it's important to understand the limitations of this metric. MaxRSS is recorded at fixed intervals, approximately every three seconds. This means that short, intense memory spikes may not be captured, leading to potential inaccuracies in memory usage reporting. As a result, MaxRSS should be interpreted with caution, especially for applications that experience brief but significant memory usage peaks.

Resource Allocation Best Practices

Requesting the appropriate amount of resources is key to optimizing job scheduling and performance. Over-allocation of resources, such as requesting more memory or processing power than necessary, does not incur penalties but can hinder scheduling efficiency. Slurm prioritizes jobs that fit well within available resources, so it's beneficial to request only what is truly needed.

On the other hand, when using a GPU, specifying a high memory request may be unnecessary. Consider a scenario where you request a fixed amount of RAM on a node equipped with an A100 GPU. In this case, the GPU is likely to be the primary constraint in scheduling your job. Other users typically won't utilize the remaining memory on a GPU node, so it's advantageous to use the full memory

capacity to prevent out-of-memory (OOM) errors. While tight resource requests generally help minimize waiting times, in this case, a stricter memory allocation doesn't necessarily reduce wait times but could lead to job cancellations due to OOM errors. Therefore, it's more practical to avoid limiting the memory, ensuring that your job can utilize the node's full resources without interruptions.

Additionally, leveraging CPU nodes can be an effective strategy for certain tasks. While CPU nodes may not offer the same processing power as GPU nodes for computationally intensive tasks, they are often available immediately and can be used in parallel to expedite job completion. This approach can be particularly beneficial when GPU wait times are long, as it allows you to utilize available resources effectively and reduce overall processing time.

Conclusion

By utilizing profiling tools like Valgrind and the seff script, and adhering to resource allocation best practices, users can significantly enhance their job performance and scheduling efficiency. Understanding the limitations of memory monitoring tools like MaxRSS and strategically requesting resources can lead to better cluster utilization and a more seamless user experience. This documentation serves as a guide to help users make informed decisions and optimize their workflow on the cluster.

From:
<https://hpc.wiki.tu-berlin.de/> - **HPC-Cluster-Dokumentation**

Permanent link:
https://hpc.wiki.tu-berlin.de/doku.php?id=hpc:scheduling:profiling_and_resource_monitoring

Last update: **2025/01/31 17:16**

