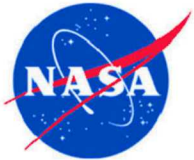


Performance Portable Magnetohydrodynamics for the Exascale Era

Forrest W. Glines^{1,2,3},
Brian W. O'Shea^{1,2}, Philipp Grete¹



¹Department of Physics and Astronomy, Michigan State University, East Lansing, MI

²Department of Computational Mathematics, Science and Engineering, Michigan State University, East Lansing, MI

³Sandia National Laboratory, Albuquerque, NM



Sandia National Laboratories is a multimission laboratory managed and operated by National Technology & Engineering Solutions of Sandia, LLC, a wholly owned subsidiary of Honeywell International Inc., for the U.S. Department of Energy's National Nuclear Security Administration under contract DE-NA0003525.



The Renaissance Simulations (O'Shea et al. 2015)

Movie c/o Donna Cox, Bob Patterson, NCSA Advanced Visualization Laboratory

CPU



Xeon Phi



Stampede 2



Cori (12)
LBNL



NVIDIA

BLUE WATERS
SUSTAINED PETASCALE COMPUTING



Summit (1)



Perlmutter

GPU



Aurora
ANL



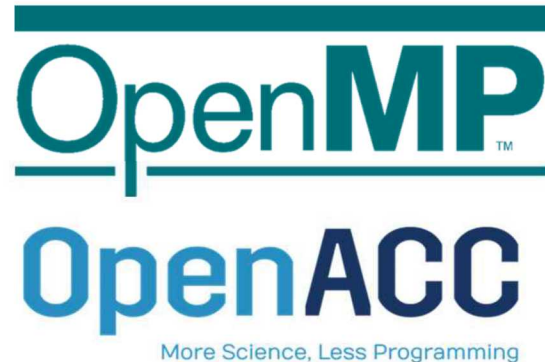
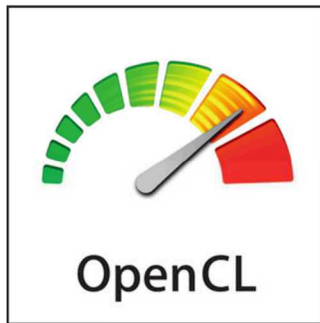
AMD



Performance Portability APIs

Typical Solutions

- OpenCL
- OpenACC
- OpenMP 4.5, 5.0



Abstraction Layers

- OCCA
- RAJA
- Kokkos

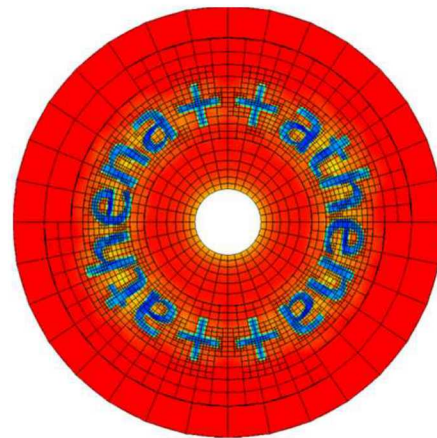


K-Athena: Performance Portable MHD

- Conversion of Athena++ with Kokkos
- Minimal code changes
- Performance Portable
 - Ready for CPUs, GPUs
- Freely available on GitLab

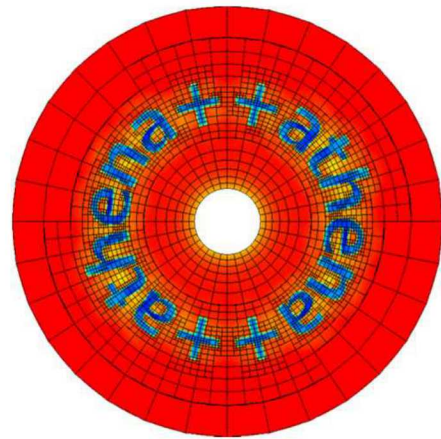
<https://gitlab.com/pgrete/kathena>

<https://arxiv.org/abs/1905.04341>



Pathway of Conversion

1. Conversion of data to Kokkos Views
 - Unified GPU memory helps
2. Converted kernels one by one to Kokkos
3. Ensured data remained on device
 - Turning off Unified GPU memory helps



Kokkos Conversion

Athena++

```
for( int k = ks; k < ke; k++){  
  for( int j = js; j < je; j++){  
    #pragma omp simd  
    for( int i = is; i < ie; i++){  
      /* Loop Body */  
      u(k,j,i) = ...  
    }  
  }  
}
```

K-Athena

```
using namespace Kokkos;  
parallel_for( MDRangePolicy<Rank<3>>  
  ({ks,js,is},{ke,je,ie}),  
  KOKKOS_LAMBDA(int k, int j, int i){  
    /* Loop Body */  
    u(k,j,i) = ...  
  });
```

Loop Heading Changes



Data Changes to Kokkos::View



Loop Body Unchanged



Parallelism Flexibility

Flat Parallelism

```
using namespace Kokkos;  
parallel_for( MDRangePolicy<Rank<3>>  
    ({ks, js, is}, {ke, je, ie}),  
    KOKKOS_LAMBDA(int k, int j, int i){  
        /* Loop Body */  
        u(k, j, i) = ...  
    });
```

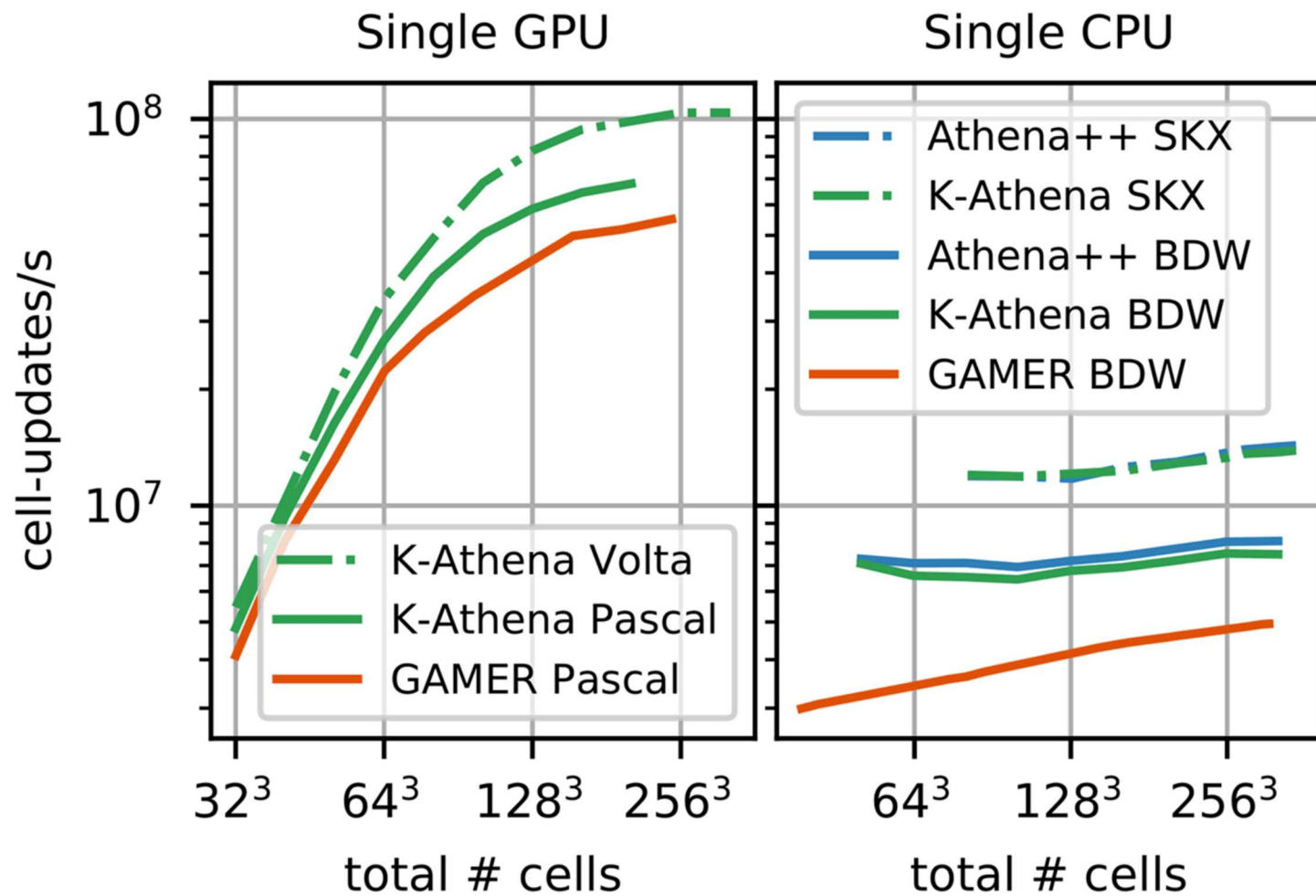
Teams: CPU Cores on a Socket, SMs on GPU

Threads: HT on CPU, Warps on GPU

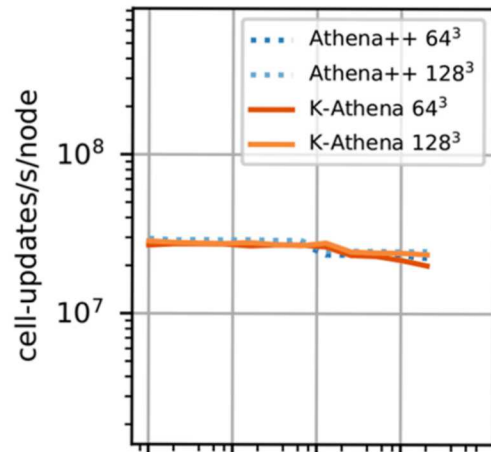
Vectors: SIMD on CPU, Threads on GPU

Hierarchical Parallelism

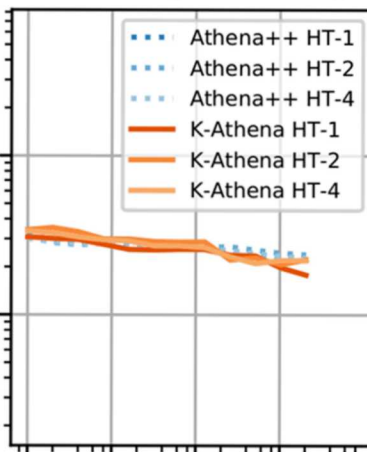
```
using namespace Kokkos;  
parallel_for(team_policy(nk*nj, AUTO),  
    KOKKOS_LAMBDA(member_type team_mem){  
        int lr = team_mem.league_rank();  
        int k = lr / nj + ks;  
        int j = lr % nj + js;  
        parallel_for(  
            TeamThreadRange<>(team_mem, is, ie),  
            [&](int i) {  
                /* Loop Body */  
                u(k, j, i) = ...  
            });  
    });
```



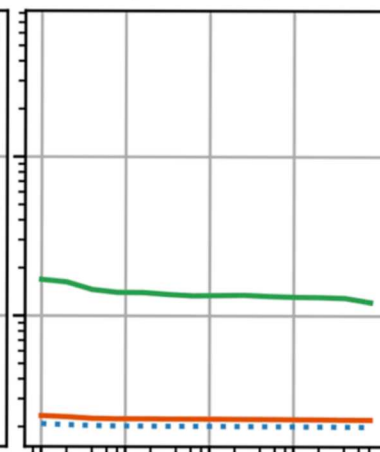
Electra Skylake CPU



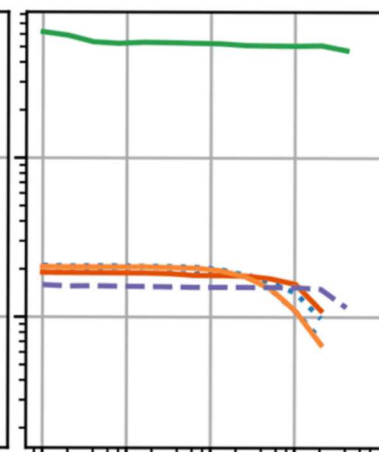
Theta Knights Landing



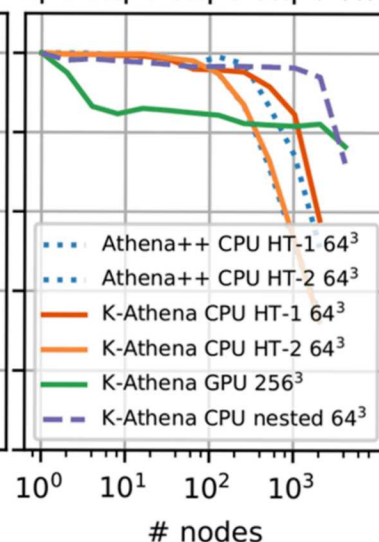
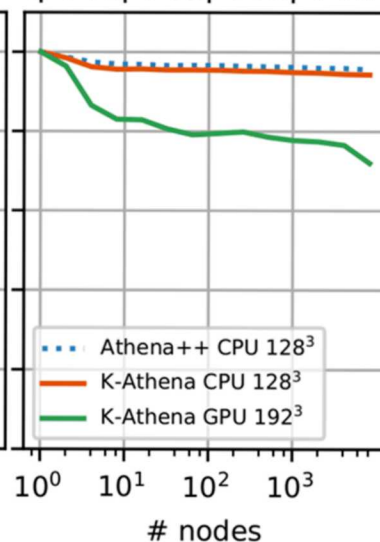
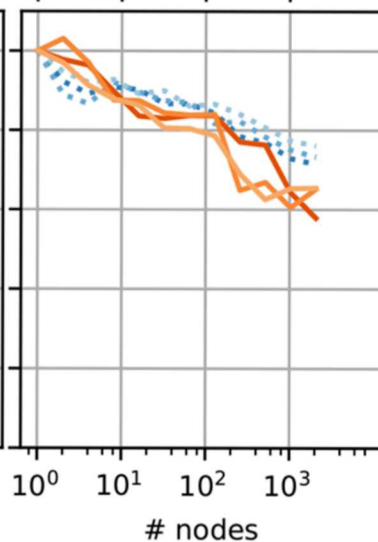
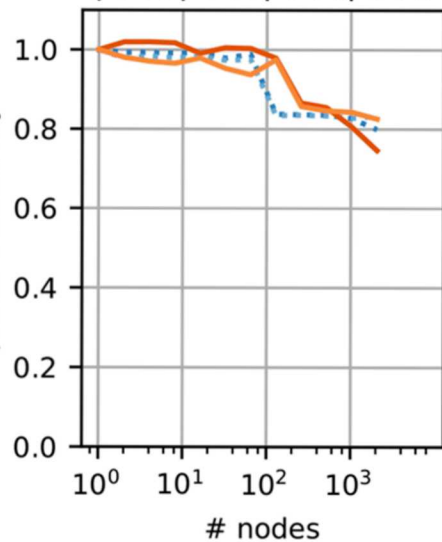
Titan Opteron/Kepler GPU



Summit Power9/Volta GPU



parallel efficiency



Performance Portability Metric

$$P(a, p, H) = \begin{cases} \frac{H}{\sum_{i \in H} \frac{1}{e_i(a, p)}} & \text{if } i \text{ is supported } \forall i \in H \\ 0 & \text{otherwise} \end{cases}$$

a is an application

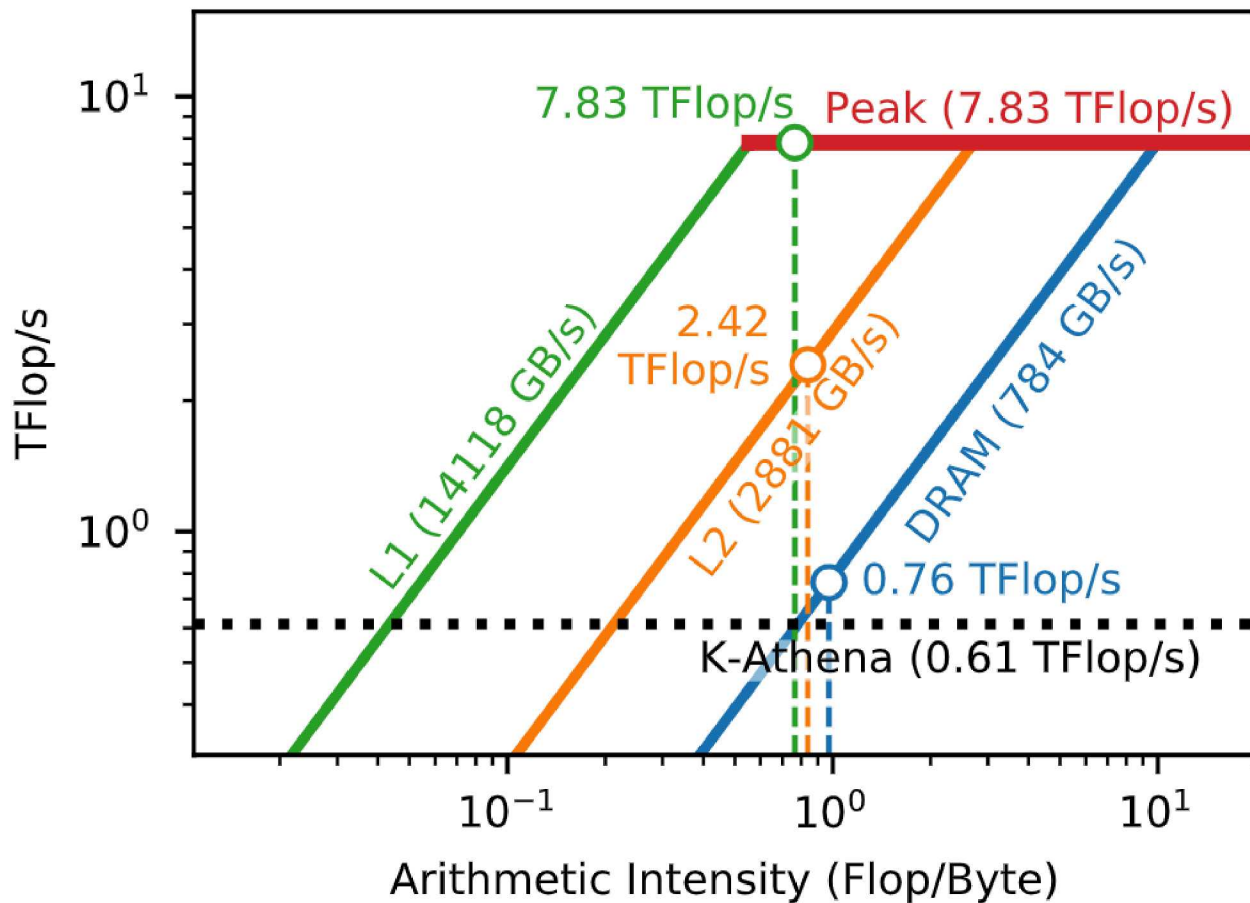
H is a set of computational platforms

p is a problem solvable by the application

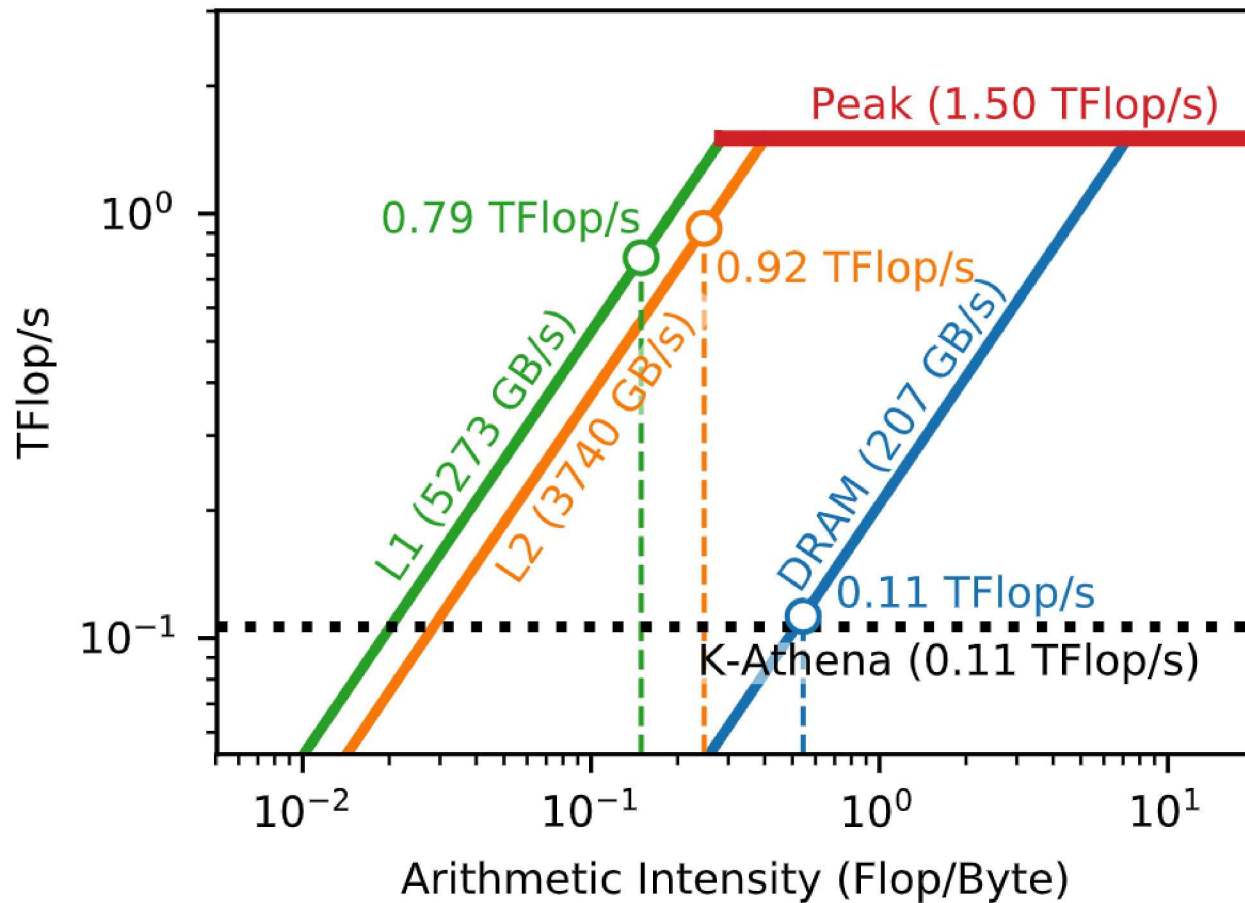
e_i is the efficiency of a to solve p on platform i

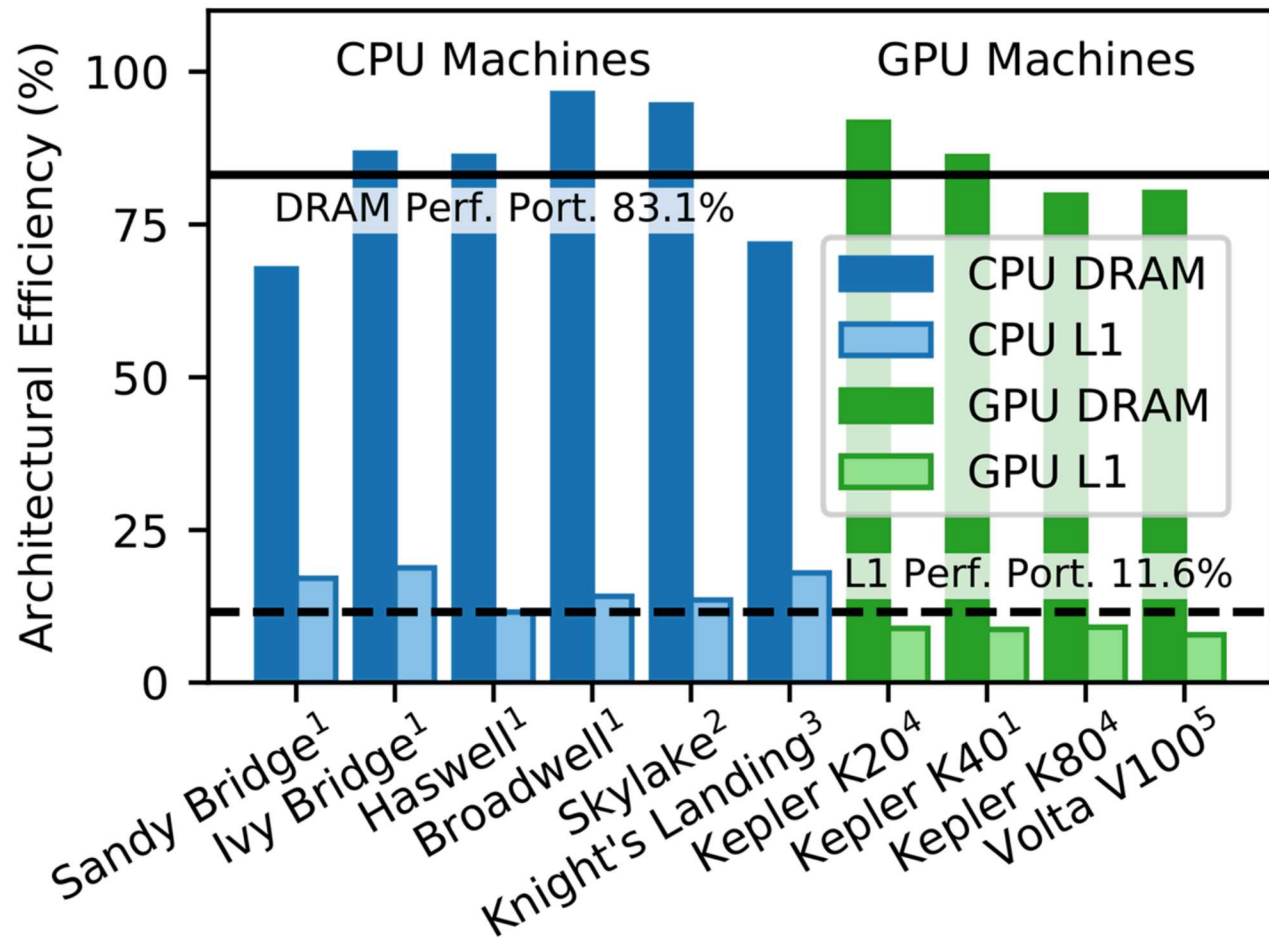
(Pennycook 2019)

NVIDIA Volta 100 GPU Roofline Model

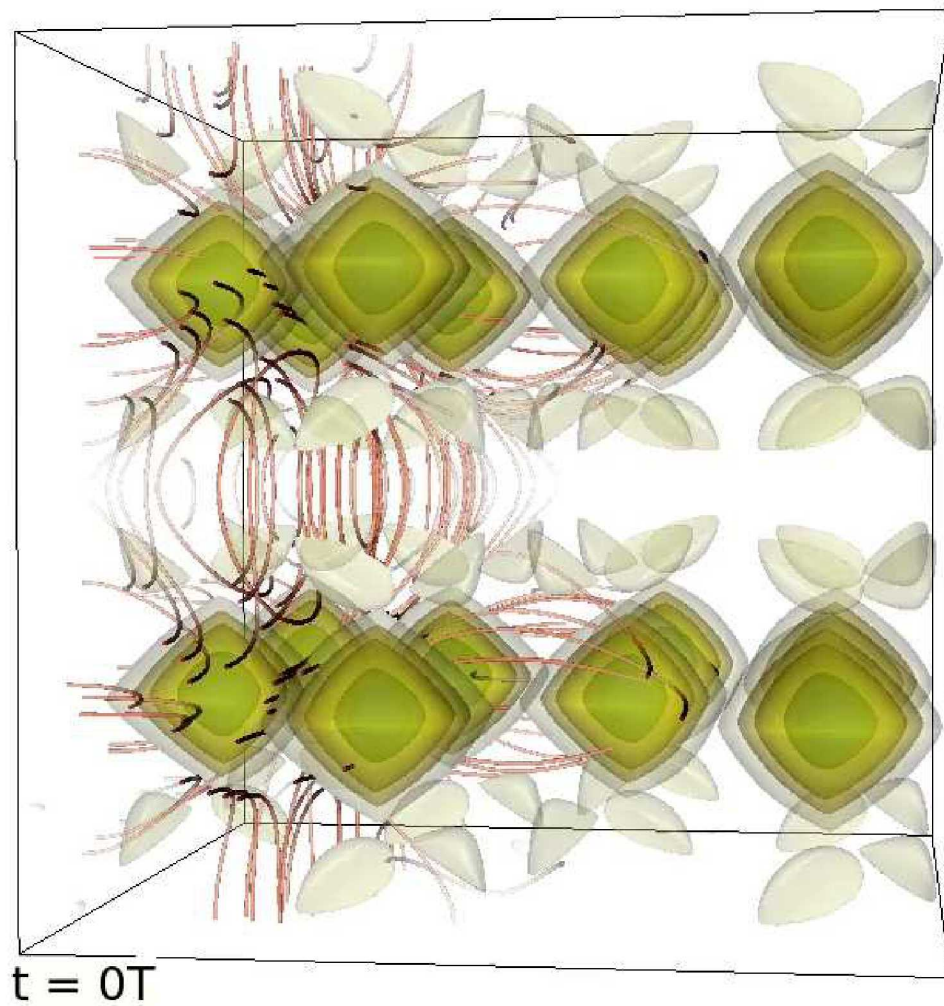


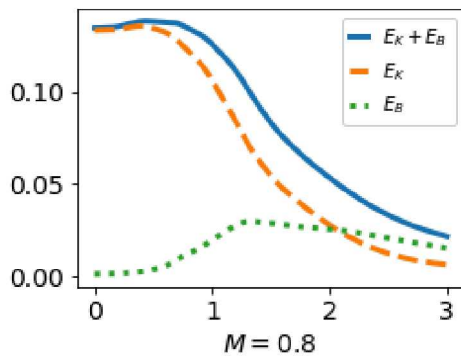
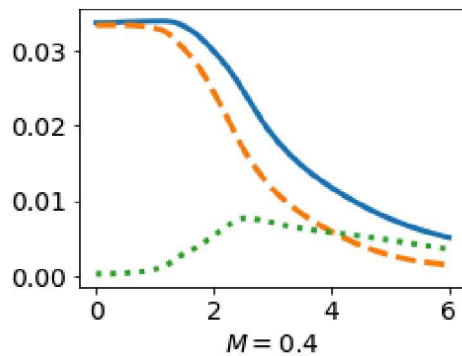
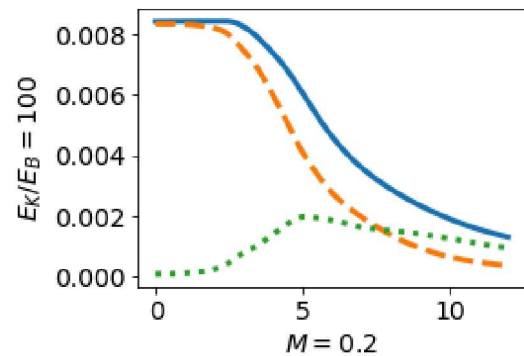
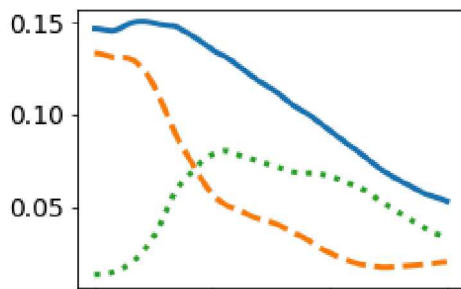
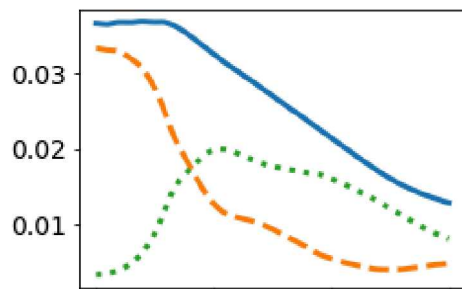
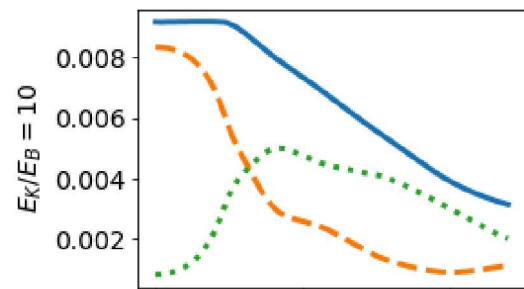
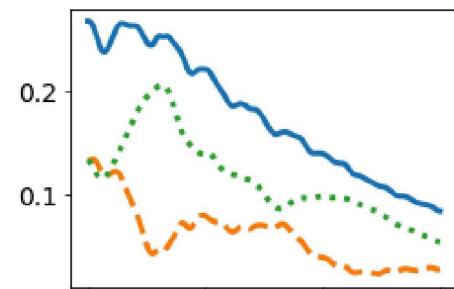
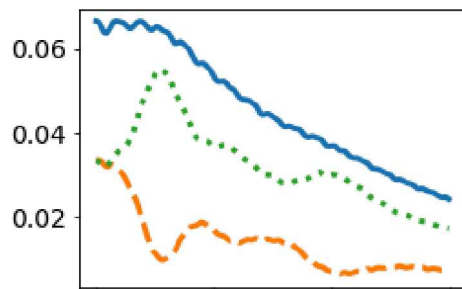
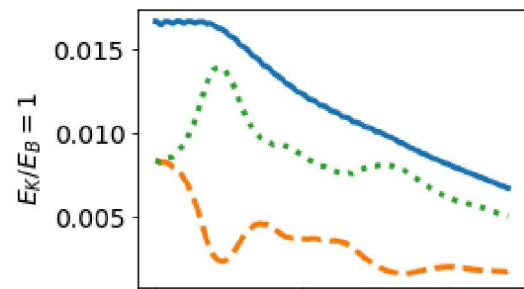
Intel Xeon Gold 6148 "Skylake" CPU Roofline Model





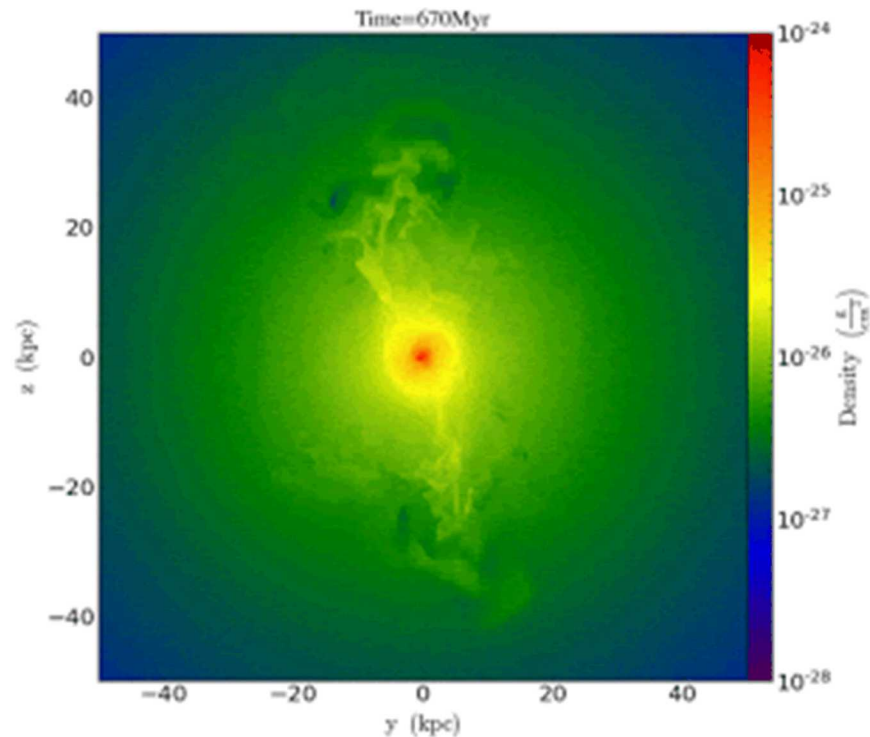
¹Pleiades, ²Electra, ³Stampede 2, ⁴MSU HPCC, ⁵Summit





Future: Isolated Galaxy Clusters with K-Athena

- Physics rich cluster simulation
 - AMR MHD
 - Magnetic AGN feedback
 - Tabulated Cooling
 - Cosmic Ray model
 - Viscosity and Conduction
- How does the interstellar medium behave?
- How does active galactic nuclei feedback thermalize?



K-Athena

- Performance Portable Ideal MHD
 - Available for CPU, GPU, future architectures (Aurora, Frontier, ARM)
 - Adaptive Mesh Refinement soon
 - Special Relativistic and General Relativistic MHD can be made available
- MHD Turbulence Simulations on Blue Waters, XSEDE systems
- Magnetized Galaxy Cluster simulations
- Kokkos for future projects (Enzo-E)

<https://gitlab.com/pgrete/kathena>

<https://arxiv.org/abs/1905.04341>

References

- Carter Edwards, H., Trott, C. R., & Sunderland, D. 2014, Journal of Parallel and Distributed Computing, 74, 3202
- Grete, P., Glines, F. W., & O'Shea, B. W. 2019, arXiv:190504341 [astro-ph, physics:physics], <http://arxiv.org/abs/1905.04341>
- O'Shea, B. W., Wise, J. H., Xu, H., Norman, M. L. 2015, The Astrophysical Journal Letters, 807, 22
- Pennycook, S. J., Sewall, J. D., & Lee, V. W. 2019, Future Generation Computer Systems, 92, 947
- White, C. J., Stone, J. M., & Gammie, C. F. 2016, The Astrophysical Journal Supplement Series, 225, 22

We acknowledge funding by NASA Astrophysics Theory Program grant #NNX15AP39G. Code development, testing, and benchmarking was made possible through various computing grants including allocations on NASA Pleiades (SMD-16-7720), OLCF Titan (AST133), OLCF Summit (AST146), ALCF Theta (Athena performance), XSEDE Comet (TG-AST090040), and Michigan State University's High Performance Computing Center.

Sandia National Laboratories is a multitechnology laboratory managed and operated by National Technology & Engineering Solutions of Sandia, LLC, a wholly owned subsidiary of Honeywell International Inc., for the U.S. Department of Energy's National Nuclear Security Administration under contract DE-NA0003525.

