

Srinath Vadlamani, Ph.D.

GPU & HPC Performance Scientist — making exascale science run fast

Email: srinath.vadlamani@pm.me · Location: Lafayette, CO Web: svac-llc.com · LinkedIn: [/in/srinathvadlamani](https://in.linkedin.com/in/srinathvadlamani) · GitHub: [srinathv](https://github.com/srinathv) U.S. Citizen · Eligible for DOE/NNSA security clearance

- **El Capitan** — owned GPU application performance on the **#1 system on the TOP500** (~44,500 AMD MI300A APUs, 2.79 exaflops), the NNSA's exascale flagship.
- **+40% at exascale** — recovered over 40% performance on critical NNSA benchmarks by diagnosing and mitigating system noise / OS jitter (SC'25).
- **2.7% → 89%** — lifted a production GPU kernel from 2.7% to 89% of peak DRAM bandwidth through roofline-driven optimization, portably across NVIDIA and AMD.
- **20 years · Ph.D. Applied Math · 11 refereed publications · cross-vendor (NVIDIA + AMD) at every scale.**

PROFILE

I make mission-critical scientific codes run fast on the world's largest machines. Over twenty years I've owned performance on systems from climate supercomputers to **El Capitan**, working the full stack — from **silicon and microarchitecture** to **numerical algorithms** and **application science**. I pair deep, **cross-vendor GPU performance methodology** (NVIDIA *and* AMD) with a Ph.D. applied-mathematician's grasp of the science itself — and I prove every optimization is *real*, not just plausible.

CORE EXPERTISE

- **Performance at scale** — roofline analysis, hardware-counter profiling, occupancy & register-pressure tuning, memory-bandwidth optimization, collective/communication tail-latency, OS-noise/jitter mitigation at exascale.
- **Heterogeneous architectures** — AMD Instinct (MI250, MI300A), NVIDIA (V100/A100/H100), x86 & aarch64; ISA/microarchitecture-level modeling.
- **Cross-architecture modernization** — performance-portable C++/Fortran; CPU↔GPU porting in CUDA, HIP, OpenMP-target, OpenACC.
- **Scientific computing** — numerical analysis, PDE/finite-volume methods, particle-in-cell & continuum plasma methods, coupled multiphysics.
- **AI for science** — physics-informed / geometry-aware ML, symmetry-based verification, AI-assisted performance workflows.
- **Reproducible science** — Spack, ReFrame, CI/CD, containers; open-source; standards leadership.

Toolchain: C · C++ · Fortran · Python · Bash | MPI (+OpenSHMEM) · OpenMP · HIP · CUDA · OpenCL | rocprof / rocprofiler-compute · Nsight Compute/Systems · perf · TAU · VTune · CrayPat | Spack · ReFrame · Slurm · Flux · Lustre/GPFS · InfiniBand/Slingshot · Apptainer/Docker

EXPERIENCE

Founder & Principal Performance Consultant — SV Advanced Computing LLC · 2025–Present

Independent GPU/HPC performance-engineering and applied-research practice (svac-llc.com). - **Published a portable-GPU performance methodology** that took a hydrodynamics proxy (PENNANT) from **2.7% to 89% of peak DRAM bandwidth**, ported across **NVIDIA T4/H100 and AMD MI300X** (CUDA / HIP / OpenMP-target) with a single roofline-driven approach. - **Advancing AI-for-science**: physics-informed and **geometry-aware (Lie-symmetry) neural methods** for fusion/plasma simulation, using symmetry- and conservation-consistency as a *label-free verification oracle*. - Authored an **adaptive-parallelism** study that reconfigures data/tensor/pipeline parallelism over the interconnect topology to sustain large-model training throughput as workloads grow.

Principal Performance Engineer — HPE Center of Excellence, LLNL “El Capitan” · 2022–2025

Owned application performance for the NNSA’s exascale flagship (AMD MI300A APU) — **El Capitan, #1 on the TOP500**. - **Recovered >40% performance** on critical NNSA benchmarks at scale by diagnosing and mitigating **system noise / OS jitter** across the ~44,500-APU machine — the basis of a **2025 Supercomputing (SC’25)** paper. - **Directed the MI300A power-management campaign** — per-chip thermal/DVFS characterization and full-machine power-cap policy (rocm-smi) — extracting maximum sustained performance within a fixed power envelope. - Set **application-readiness criteria** for mission-critical DOE/NNSA codes and built reproducible Spack/ReFrame benchmarking pipelines; advised HPE architects and LLNL scientists on heterogeneous-APU tuning strategy.

Senior Staff Field Engineer — Arm Inc. · 2017–2022

- **Drove aarch64 adoption** across DOE/NNSA scientific workflows; benchmarked A64FX / Graviton / Ampere with production codes and delivered **ISA/microarchitecture-level** performance guidance.
- Prototyped HIP kernels and ran comparative GPU profiling on AMD MI250/MI300A. **Arm’s representative to the INCITS Fortran Standards Committee.**

Computational Scientist — ParaTools Inc. · 2014–2017

- Built performance-portable C++/MPI applications (GraviT); extended the **TAU** performance-tools ecosystem for CUDA/GPU profiling.

Software Engineer III — NCAR · 2011–2014

- Modernized **CESM** climate-model kernels for vectorization and many-core/GPU (OpenACC); mentored interns and led community sessions.

Research Scientist — Tech-X Corp. · 2006–2011

- Co-developed gyrokinetic **fusion** modeling tools (GYRO, GEM); led componentization of the **FACETS** coupled core-edge multiphysics framework.

Postdoctoral Fellow — University of Washington · 2005–2006

- Authored **MH4D**, a finite-volume MHD code for fusion-plasma modeling.

EDUCATION

Ph.D., Applied Mathematics — University of Colorado Boulder, 2005 · *Dissertation: Unification of Particle-In-Cell and Continuum Methods for ETG Instability* **M.S., Applied Mathematics** — CU Boulder, 2003 · **M.S., Mathematics** — UNC Wilmington, 2001 · **B.S., Physics** — UNC Chapel Hill, 1996

RECOGNITION & LEADERSHIP

- **Author** — “Breaking the System Noise Barrier at Exascale,” Proc. **SC’25**, 2025.

- **Reviewer** — IEEE Workshop on Performance, Portability & Productivity in HPC (**P3HPC**), Supercomputing series, 2021–Present.
- **INCITS Fortran Standards Committee** representative · **Maintainer/contributor** — Spack, ReFrame.
- **Presenter** — Arm HPC User Group @ ISC'21 & SC'21; SC'21 Hackathon.

SELECTED PUBLICATIONS (11 REFEREED; FULL LIST ON REQUEST)

- S. Vadlamani et al., “Breaking the System Noise Barrier at Exascale,” Proc. **SC'25**, 2025.
- S. Vadlamani, Y. Kim, J. Dennis, “DG-kernel: A Climate Benchmark on Accelerated and Conventional Architectures,” IEEE XSW'13.
- A. Hakim, ..., S. Vadlamani et al., “Coupled core-edge simulations of H-mode buildup using the FACETS code,” *Phys. Plasmas* **19**, 032505 (2012).
- A. Pletzer, D. McCune, S. Muszala, S. Vadlamani, S. Kruger, “Exposing Fortran Derived Types to C and Other Languages,” *Computing in Science & Engineering*, 2008.
- S. Vadlamani, S. E. Parker, Y. Chen, C. Kim, “The particle-continuum method: an algorithmic unification of particle-in-cell and continuum methods,” *Comp. Phys. Comm.* **164**, 209–213 (2004).