

NEWS RELEASE

For Immediate Release



E4S RELEASE 25.06 NOW AVAILABLE WITH SIGNIFICANTLY EXPANDED AI AND HPC PORTFOLIO

EUGENE, OR., June 6, 2025— The E4S Project, today announced the immediate availability of Release 25.06. This exciting release includes a wide range of new features and improvements, including a significantly expanded AI Portfolio and Support for the NVIDIA Blackwell GPU Architecture with CUDA 12.8.

E4S, an [HPSE](#) project, is an HPC-AI software ecosystem for science. It's a curated, [Spack](#) based collection of scientific libraries and tools that form the foundation of some of the world's most advanced scientific applications. E4S curates the largest single collection of open-source GPU-enabled libraries and tools for scientific applications, including support for equation-based modeling and simulation, and AI for science. Since the project launch in 2018, E4S has experienced widespread adoption as a globally recognized ecosystem of open-source software packages to accelerate the development, deployment and use of HPC-AI applications.

"We are excited to announce that Adaptive Computing's [Heidi AI](#) Supercomputers and Heidi Workshops use E4S as the core software stack for HPC/AI applications for On-Premises Datacenters, AWS, Azure, Google Cloud, and OCI, for all industry verticals, including the K-12 and Higher Education verticals," said Art Allen, President and CEO, Adaptive Computing Enterprises, Inc.

E4S 25.06 includes over 950 binaries and over 130 HPC and AI packages. All E4s public-domain software is thoroughly tested for interoperability and portability to multiple computing architectures and will continue to be enhanced and expanded to address architectural changes and emerging new architectures. While E4S supports many products and distributions, users can confidently select any subset of functionality. We build and test the whole so you can select what you need.

The key features of E4S 25.06 support a timely expansion of the ecosystem's AI portfolio and include:

- Support for the NVIDIA Blackwell GPU architecture with CUDA v12.8.
- [NVIDIA NeMo™ Framework](#) v2.3.0rc0, a comprehensive framework for building, customizing, deploying, and maintaining generative AI models. It includes support for large language models (LLMs), video models, vision language models (VLMs), and speech AI.
- [NVIDIA BioNeMo Framework](#) v2.6, an open source, AI platform to accelerate drug discovery and biopharmaceutical research providing a comprehensive suite of frameworks, pretrained models, generative AI tools, and microservices for computational biology and chemistry. It includes

support for modeling DNA, RNA, and proteins and offers optimized AI models and curated training recipes for biopharma supporting the entire drug discovery process.

- Vllm v0.8.3, an inference engine for efficiently serving LLMs at scale and supporting GPUs.
- Hugging Face Hub v0.32.3 and CLI tools, a leading open-source platform for AI/ML often described as the “GitHub of ML”. These tools provide access to a model hub with over a million pre-trained models for text classification, language generation, translation, summarization, and image generation. It features a Transformers Python library simplifying the use of deep learning models for NLP tasks.
- E4S also includes TensorFlow, PyTorch, JAX, Scikit-Learn, Pandas, Lbann, OpenCV, and Torchbraid optimized for GPUs.

"E4S delivers a robust, well-tested, and integrated ecosystem for developing and deploying HPC-AI applications across heterogeneous architectures, including CPUs and GPUs from NVIDIA, AMD, and Intel," said Prof. Sameer Shende, University of Oregon and ParaTools, Inc. "E4S scales seamlessly from laptops to supercomputers and supports major cloud platforms such as AWS, Azure, Google Cloud, and Oracle Cloud, with optimized multi-node GPU acceleration. The E4S 25.06 release emphasizes generative AI capabilities, featuring tools for deploying and customizing large language models (LLMs), direct integration with Hugging Face's model repository, frameworks to build and deploy AI chatbots, and support for NVIDIA's BioNeMo framework, enabling biopharma research and drug discovery workflows. This release strengthens E4S's position as a unified platform for scientific computing and AI innovation, bridging traditional HPC with cutting-edge generative AI applications."

Additional New Features in Release 25.06

Other new features of the E4S release include support for PyTorch and Vllm for Intel GPUs, Intel 25.1 compilers, ROCm 6.3.3 integration for AMD GPUs; the high-order API library, [libCEED](#); and [Julia](#) with support for AMD and NVIDIA GPUs with MPI.

[Chapel](#) has been upgraded with support for NVIDIA and AMD GPUs.

Other applications of note in E4S 25.06 include CP2K, DealII, FFTX, GROMACS, LAMMPS, Nek5000, Nekbone, NWChem, OpenFOAM, WarpX, WRF, Quantum Espresso, and Xyce, with support for GPUs where available.

VSCodium provides the integrated development environment, Jupyter notebook provides a web-based Python notebook interface.

3D graphics tools include ParaView, VisIt, and TAU.

E4S tools updated in release 25.06

E4S tools updated in release 25.06 include e4s-chain-spack.sh to chain two Spack instances to allow users to install new packages while leveraging pre-installed E4S packages and dependencies, e4s-alc to customize containers from base container images, and e4s-cl to support binary distribution of MPI

applications by substituting the MPI library embedded in the containerized application with the system/vendor MPI for near native performance.

E4S 25.06 has separate CPU (x86_64, aarch64, ppc64le) and GPU containers for NVIDIA, Intel, and AMD GPUs optimizing the size of the containers for each GPU architecture.

Rounding out the expanded capabilities of E4S 25.06, E4S is the availability of [ParaTools Pro for E4S™](#) on commercial cloud platforms and their marketplaces on Azure, AWS, Google Cloud, and Oracle Cloud Infrastructure with support for NVIDIA GPUs.

Additional details on Adaptive Computing and E4S

Art Allen, President and CEO, Adaptive Computing Enterprises, Inc. also commented that, *“Being browser based, Heidi AI provides a uniform environment across all user devices and Cloud platforms for deploying HPC/AI workloads that are CPU/GPU enabled. With a highly performant VNC-based Linux desktop and compute nodes, Heidi AI also uses the fastest network interconnects available, often reaching near theoretical speeds. Heidi AI uses containerized E4S deployments as well, both on-premises and in the Cloud. Heidi AI has four (4) US patents with three (3) more pending.”*

For more information:

Sameer Shende, sameer@cs.uoregon.edu

Michael Heroux, mheroux@paratools.com

=====

Editors Note: Supplemental Visual Information attached



Figure 1: [Zero-shot prediction of BRCA1](#) variant effects with NVIDIA BioNeMo Evo 2

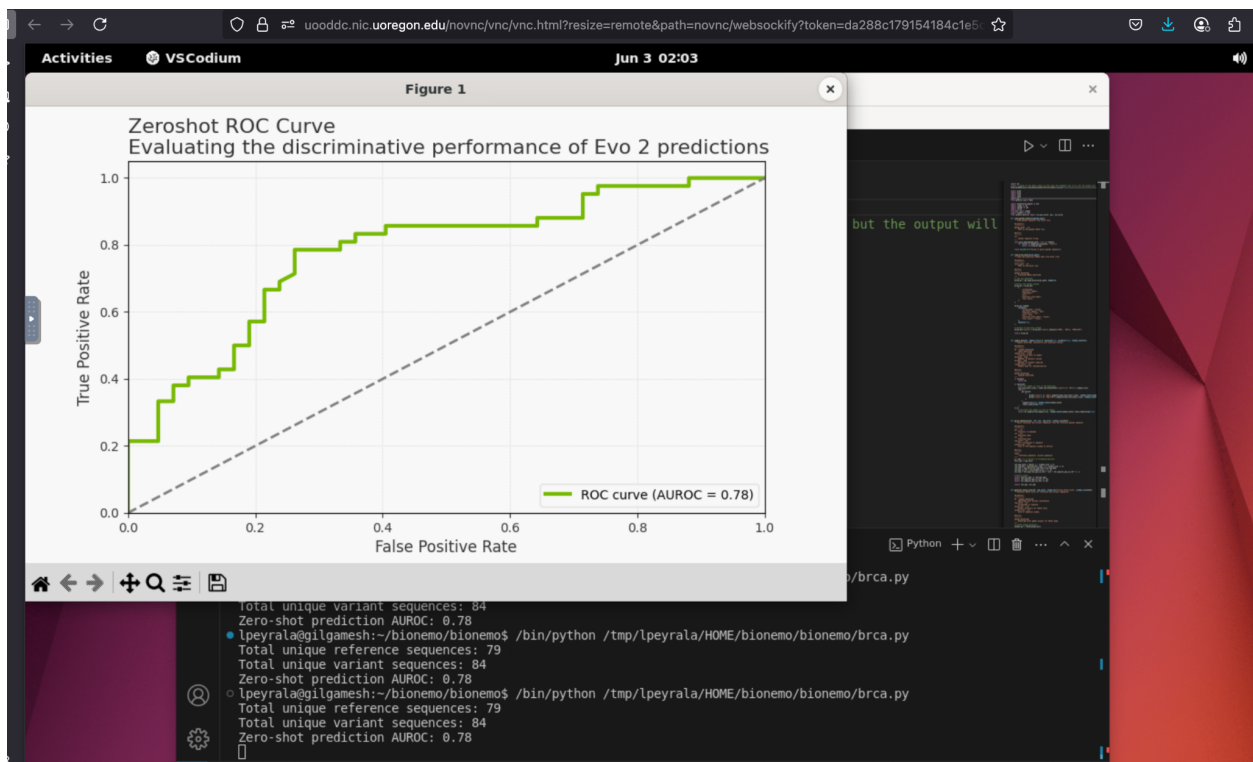


Figure 2: NVIDIA BioNeMo Framework's tutorial available [here](#).