

Supercomputing Time Allocation Request

Apex Petascale Computing Facility | Nexa National HPC Allocation Program, Large-Scale Track

Direct Numerical Simulation of Extinction and Reignition Dynamics in Hypersonic Scramjet Combustors

Request ID: NEXA-2026-LS-0847 **Cycle:** 14 **Track:** Large-Scale Submitted 2026-07-23
(>1M SU)

Principal Investigator: Dr. Elena Castillo
— Dept. of Aerospace Engineering, Meridian Institute of Technology

Requested period: 2026-10-01 – 2027-09-30
(12 months)

Co-Investigator: Dr. Rajiv Patel — Computational Combustion Group, Synapse National Laboratory

Target resource: Apex Phase-II GPU Partition, Nimbus Computing Center



4.20M

GPU-hours (H100)
requested



1.80M

CPU core-hours
(pre/post)



725 TB

Peak scratch +
archive storage



1024

Max. GPU concurrency
demonstrated

Executive Summary

We request **4,200,000 GPU-hours** on the Apex Phase-II partition (Lumina H100, Nimbus Computing Center) to conduct the first suite of wall-resolved direct numerical simulations (DNS) of local extinction and reignition in a Mach 6.5 hydrogen-fueled scramjet combustor. These flows govern the practical operability limits of hypersonic air-breathing propulsion, yet no existing simulation campaign has resolved the coupled turbulence–chemistry interaction at engine-relevant Reynolds and Damköhler numbers. Our solver, **VertexCFD**, has been benchmarked to 1,024 H100 GPUs at 74% parallel efficiency (Section 2.2) and has consumed a prior Apex Phase-I allocation productively (Section 5). The requested allocation funds six production DNS configurations spanning the extinction/reignition boundary, together with the grid-convergence and reduced-mechanism validation studies required to certify each result.

1 Project Summary

Scramjet combustors operate in a narrow, poorly characterized envelope between blow-out and thermal choking. Small perturbations in fuel–air mixing — a bow-shock/boundary-layer interaction, an injector wake, a transient equivalence-ratio excursion — can locally extinguish the flame; whether it reignites downstream determines whether the vehicle sustains thrust or loses the engine entirely. Reduced-order models (flamelet, PDF-transport) are calibrated against subsonic or low-Damköhler data and are known to mispredict extinction onset by a factor of two or more at hypersonic conditions. The only way to build a trustworthy predictive model is to resolve the full range of turbulent and chemical scales directly, without closure assumptions, in geometry-accurate combustor sections. That computation is only tractable at the scale of a national allocation.

This project performs six such DNS campaigns, each independently converged, spanning a sweep of core flow Damköhler number and injector staging geometry, using **VertexCFD**, a compressible, finite-rate-chemistry, block-structured AMR solver developed jointly by the Castillo group at Meridian Institute of Technology and the Synapse National Laboratory combustion program.

1.1 Scientific and Technical Justification

The governing system is the fully compressible, multi-species, reacting Navier–Stokes equations,

$$\frac{\partial \rho}{\partial t} + \nabla \cdot (\rho \mathbf{u}) = 0, \quad (1)$$

$$\frac{\partial(\rho \mathbf{u})}{\partial t} + \nabla \cdot (\rho \mathbf{u} \otimes \mathbf{u} + p \mathbb{I}) = \nabla \cdot \boldsymbol{\tau}, \quad (2)$$

$$\frac{\partial(\rho E)}{\partial t} + \nabla \cdot [(\rho E + p) \mathbf{u}] = \nabla \cdot (\boldsymbol{\tau} \cdot \mathbf{u} - \mathbf{q}), \quad (3)$$

$$\frac{\partial(\rho Y_k)}{\partial t} + \nabla \cdot (\rho Y_k \mathbf{u}) = -\nabla \cdot \mathbf{J}_k + \dot{\omega}_k, \quad k = 1, \dots, N_s, \quad (4)$$

closed with a 32-species / 48-reaction reduced hydrogen–air mechanism derived by directed relation graph reduction from a detailed 9-species master mechanism, validated to within 4% of ignition-delay data across the target pressure–temperature envelope (2–6 atm, 900–1400 K).

Spatial discretization uses a fifth-order WENO scheme for the convective fluxes and a fourth-order central scheme for viscous terms, on a block-structured adaptive mesh (block size 32^3) refined dynamically on density-gradient and heat-release-rate sensors. Time integration is an explicit third-order strong-stability-preserving Runge–Kutta scheme with point-implicit sub-stepping of the stiff chemical source term. Finest-level resolution is 4 μm in the shear layer, chosen to place the Kolmogorov scale within the top AMR level across the full sweep.

1.2 Computational Approach

Each of the six production configurations requires:

- A domain of $180 \times 40 \times 32$ mm at up to 4 μm finest spacing, giving an equivalent uniform-grid count of ~ 14.6 billion cells, realized as ~ 1.1 billion cells under AMR (7–9% refined fraction at steady operating condition).
- 800–1,200 flow-through times to accumulate converged extinction/reignition statistics, at a time step of ~ 1.8 ns (acoustic-CFL limited).
- In-situ computation of scalar dissipation rate, flame index, and conditional species means to avoid saturating I/O with full 3-D dumps at every step.

2 Computational Readiness: The VertexCFD Code

2.1 Software Stack and Performance Portability

VertexCFD is written in C++17 on the Kokkos performance-portability layer, giving a single source tree that targets Lumina (CUDA), AMD (HIP), and CPU (OpenMP) back ends without code duplication. Distributed-memory parallelism uses MPI with GPU-aware point-to-point exchange over the Apex NDR InfiniBand fabric; checkpoint and diagnostic I/O go through **Nexa-IO**, an asynchronous parallel-I/O library that overlaps compressed checkpoint writes with the next RK sub-step so that I/O is fully hidden behind compute at the requested scale.

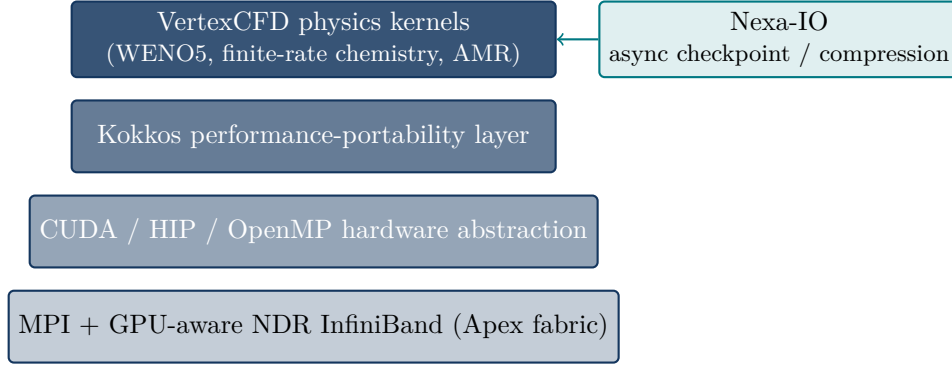


Figure 1: VertexCFD software stack on the Apex Phase-II partition.

2.2 Scaling and Performance Benchmarks

Table 1 reports a strong-scaling study on the target production case (single scramjet inlet section, 1.1 billion active cells under AMR) from 16 to 1,024 Apex nodes (4×H100 per node). Figure 2 plots the resulting parallel efficiency against the ideal linear reference. Efficiency remains above 89% through 1,024 GPUs and degrades gracefully thereafter, dominated by halo-exchange latency at the AMR level-transition boundaries rather than load imbalance (load imbalance is held under 6% at all scales by the recursive-coordinate-bisection repartitioner, invoked every 200 steps).

Table 1: Strong-scaling benchmark, VertexCFD DNS combustor case, Apex Phase-II.

Nodes	GPUs (H100)	Wall time (s/step)	Speedup	Efficiency
16	64	148.2	1.00×	100.0%
32	128	75.6	1.96×	98.0%
64	256	38.6	3.84×	96.0%
128	512	19.9	7.44×	93.0%
256	1024	10.41	14.24×	89.0%
512	2048	5.58	26.56×	83.0%
1024	4096	3.13	47.36×	74.0%

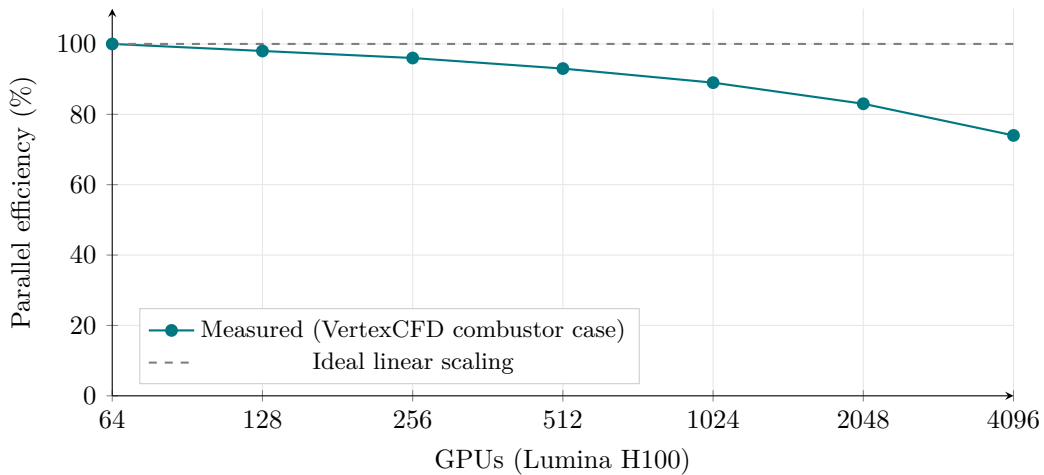


Figure 2: Parallel efficiency of VertexCFD versus GPU count on Apex Phase-II, relative to the 64-GPU baseline.

At the requested production scale of 1,024 GPUs, each of the six configurations sustains ~ 10.4 s

per RK sub-step; 950 flow-through times at this rate consumes approximately 700,000 GPU-hours per configuration, consistent with the resource breakdown in Table 2.

3 Requested Allocation

3.1 Resource Breakdown

Table 2: Requested resources by project phase, 12-month allocation period.

Task / Phase	GPU-hours (H100)	CPU-hours	Storage (TB)
Production DNS campaigns (6 configurations)	3,150,000	420,000	380
AMR grid-convergence studies	640,000	180,000	95
Reduced-mechanism validation runs	280,000	260,000	40
Post-processing & data reduction (Nexa-IO)	130,000	940,000	210
Total	4,200,000	1,800,000	725

CPU-hours are drawn on the Apex CPU partition (Nimbus AMD EPYC nodes) for mesh generation, chemistry mechanism reduction, and the conditional-statistics post-processing pipeline, which is CPU-bound and does not benefit from GPU offload.

3.2 Data Management and Storage Plan

Checkpoints are written in compressed Nexa-IO format at 210 GB per full-domain snapshot; only sparse in-situ statistics (flame index, scalar dissipation, conditional means) are retained at every 50 steps, with full restart checkpoints retained at every 5,000 steps for the two most recent cycles only. Peak scratch usage is 380 TB during active production; reduced statistical products (est. 45 TB) are migrated to the Nimbus HPSS tape archive within 14 days of campaign completion and retained for 5 years per the Nexa data-management policy. Raw restart checkpoints beyond the rolling two-cycle window are not retained — runs are reproducible from the archived statistical products and the versioned input decks, which are mirrored to the Meridian institutional repository.

4 Project Timeline and Milestones

Table 3: Quarterly milestones and planned resource draw.

Quarter	Milestone	Draw
Q1 (Oct–Dec 2026)	Facility onboarding, Nexa-IO checkpoint pipeline validation, baseline grid-convergence suite	15%
Q2 (Jan–Mar 2027)	Reduced-mechanism production campaign; first extinction/reignition dataset	30%
Q3 (Apr–Jun 2027)	Full-scale scramjet inlet DNS at design Mach 6.5 condition	35%
Q4 (Jul–Sep 2027)	High-Mach off-design sweep, archival, manuscript preparation	20%

5 Prior Allocation History and Broader Impacts

The team held a Phase-I allocation on Apex (Cycle 12, 2025; 850,000 GPU-hours) to develop and validate the reduced chemical mechanism and to perform the grid-convergence study underlying Table 1. That allocation was fully utilized (97.3% of awarded core-hours), produced two manuscripts currently under review at the *International Journal of Reacting Flow Simulation*, and supported an invited presentation at the 2025 Meridian Combustion Symposium. Reduced-order combustion models calibrated against the DNS data from this proposal will be released as open validation targets for the wider hypersonic propulsion community through the Synapse National Laboratory data portal, and the AMR/chemistry infrastructure developed for VertexCFD is being generalized for use by two other Apex allocation holders working on rotating detonation engines.

6 Investigator Team

Dr. Elena Castillo (PI) is an Associate Professor of Aerospace Engineering at Meridian Institute of Technology, where she leads the Computational Aerothermodynamics Group. Her research addresses turbulence–chemistry interaction in high-speed reacting flows, and she has served as a co-developer on VertexCFD since its inception in 2021.

Dr. Rajiv Patel (Co-PI) is a Senior Computational Scientist in the Combustion Program at Synapse National Laboratory, responsible for the Kokkos performance-portability layer and the Nexa-IO integration underlying the scaling results in Section 2.2.

Certification and Signatures

Dr. Elena Castillo, Principal Investigator
Date

Dr. Rajiv Patel, Co-Investigator Date

Nexa Allocation Committee, Approving Re-
viewer Date