

GPU-accelerated Finite-Element Method for the Three-dimensional Unstructured Mesh Atmospheric Dynamic Framework

Leisheng Li¹, Ximeng Fu^{1,2}, Xiyu Zheng^{1,2}, Huiyuan Li¹, and Jinxi Li³

¹Institute of Software, Chinese Academy of Sciences, Beijing 100190, China

²University of Chinese Academy of Sciences, Beijing, 100190, China

³State Key Laboratory of Atmospheric Environment and Extreme Meteorology, Institute of Atmospheric Physics, Chinese Academy of Sciences, Beijing 100029, China

Correspondence: Jinxi Li (ljx2311@mail.iap.ac.cn)

Abstract. The three-dimensional unstructured-mesh finite-element atmospheric dynamical framework is gaining significance owing to its flexibility in representing complex topography and capability for multi-scale simulations in high resolutions. However, this framework has substantial bottlenecks. Unlike structured-grid models, the unstructured finite element method (FEM) must frequently access irregular mesh connectivity among nodes, edges, and elements, causing indirect memory addressing, inadequate data locality, and substantial memory bandwidth bottlenecks on conventional CPU architectures. Consequently, element-wise computations and global assembly are among the primary contributors alongside the sparse linear solver to the runtime in high-resolution simulations.

This study develops a GPU-parallel implementation of the Fluidity-Atmosphere dynamical core to address these challenges. The GPU-oriented data structures and optimized kernels are designed to efficiently leverage the computing power of GPUs. These kernels enable parallelized element integration and are efficient solvers for specific size matrices; a parallel assembly strategy enhances memory throughput during global sparse matrix construction. On the NVIDIA A100 GPU, the optimized kernels achieve speeds over 100× compared to the single CPU core baseline for element-wise computations and up to 389.02 times for global matrix assembly, resulting in an overall acceleration of 8.57 times with four messages passing interface (MPI) processes. The proposed framework demonstrates that tailored GPU parallelization is effective in overcoming the computational bottleneck of unstructured FEM-based atmospheric models, facilitating high-resolution simulations on heterogeneous architectures.

1 Introduction

Atmospheric dynamic frameworks are the fundamental tools for weather forecasting and climate simulation. In recent years, growing demands for the accurate prediction and management of extreme pollution and weather events have driven atmospheric models toward higher spatial resolutions and more sophisticated physical parameterizations. These advances create formidable computational challenges: doubling the spatial resolution can lead to exponential growth in computational cost, substantially increasing the demand for computing power and storage resources (Michalakes, 2020). For the discretization methods, the choice of scheme is primarily dependent on the topological structure of the underlying mesh. For high-resolution simulations,

models are now able to resolve increasingly complex underlying surface features, such as terrains and urban structures. Owing
25 to their geometric constraints, traditional terrain-following coordinate grids typically produce significant mesh distortions near
steep terrain, introducing significant computational errors. Therefore, conventional atmospheric models based on structured
"horizontal + vertical" grid configurations encounter inherent limitations in representing complex terrain boundaries, offering
inadequate flexibility for localized adaptation.

To address this challenge, fully three-dimensional (3D) unstructured grids have emerged as an effective solution. Replacing
30 the conventional quasi-3D approach that separates horizontal and vertical discretization, unstructured grids provide superior
geometric adaptability and more flexible mesh generation capabilities. The finite element method (FEM) proves particularly
well-suited to such grid architectures, offering inherent advantages in handling complex geometrical configurations. This syn-
ergistic combination has established FEM as an increasingly prominent methodology in atmospheric modeling (Farrell et al.,
2009; Li et al., 2018). FEM not only can be formulated to satisfy discrete conservation properties and numerical stability on ir-
35 regular geometries but also enables local mesh refinement in critical regions, thereby improving simulation accuracy (Takle and
Russell, 1988; Li et al., 2021). To leverage these advantages, the Institute of Atmospheric Physics, Chinese Academy of Sci-
ences, and AMCG group at Imperial College London jointly developed Fluidity-Atmosphere, a 3D adaptive atmospheric model
(Davies et al., 2010). Based on the unstructured FEM, Fluidity-Atmosphere tightly couples the Navier-Stokes equations, com-
plete advection-diffusion dynamics, and anisotropic adaptive mesh algorithms, facilitating dynamic mesh optimization during
40 simulations to capture multi-scale flow phenomena.

However, while the FEM on unstructured meshes provides superior geometric flexibility, it also introduces inherent compu-
tational challenges that are absent in structured-grid models. In unstructured FEM frameworks such as Fluidity-Atmosphere,
each element must frequently access irregularly connected nodes, edges, and faces during numerical integration and global
sparse matrix assembly. This results in non-contiguous memory access, indirect addressing, and load imbalance across ele-
45 ments. These challenges are typically absent in structured grids, where data are stored in regular arrays with predictable access
patterns. Moreover, the adaptive mesh refinement (AMR) feature of Fluidity-Atmosphere dynamically modifies the mesh topol-
ogy during simulations, further complicating the memory layout. Consequently, the two most time-consuming components,
namely, element-wise computations and global sparse matrix assembly, become dominated by irregular data movement rather
than floating-point arithmetic, causing substantial memory bandwidth bottlenecks on CPU architectures (Sulyok et al., 2019).

50 Meanwhile, the slowdown of Moore's law and energy-efficiency bottlenecks of CPUs render CPU-only optimization inad-
equate for high-resolution simulations. In this context, GPUs have emerged as the widely used heterogeneous accelerators in
high-performance computing. With massive parallelism and high-bandwidth device memory, GPUs have delivered orders-of-
magnitude speedups in many numerical applications (Czarnul et al., 2020). For atmospheric models, particularly those based
on FEM, computational intensity primarily arises from solving large-scale linear systems, parameterizing physical processes,
55 and performing massive element-level operations, each amenable to GPU acceleration. Reengineering these components to run
efficiently on GPUs has therefore become a key research direction.

Recent studies have reported significant advances in GPU acceleration of atmospheric models, including Navier-Stokes
solvers (Goddeke et al., 2009; Thibault and Senocak, 2012), advection-diffusion equations (Simek et al., 2009), and iterative

solvers for implicit schemes (Müller et al., 2015), typically with speedups of several orders of magnitude. Some Earth system models such as ocean model (Petersen et al., 2025), tsunami model (Conde et al., 2025), cloud-resolving atmosphere model (Bogenschütz et al., 2025) and sea-ice model (Jendersie et al., 2025) also were accelerated using GPUs. FEM modules such as numerical integration (Macioł et al., 2010; Sanfui and Sharma, 2020), matrix assembly (Kiran et al., 2020; Sanfui and Sharma, 2021), and linear solvers (Ratnakar et al., 2021; Kiran et al., 2020) have been explored, achieving speedups from tens to hundreds of times. Typically, the solution of large sparse linear systems is regarded as the dominant cost in FEM-based simulations and has been extensively studied, with numerous GPU-parallel implementations already achieving mature performance. However, the computation of elements and subsequent matrix assembly may pose significant challenges, primarily due to indirect memory addressing and low arithmetic intensity, particularly in large-scale unstructured atmospheric models. During each iteration, local element-wise matrices and right-hand sides must be computed and assembled into global sparse matrices. As the simulation domain expands, the number of mesh elements grows exponentially, leading to a proportional increase in computational cost. Without targeted optimization, the efficiency of the entire model can be substantially reduced.

In practice, element-wise computations and global matrix assembly involve frequent access to unstructured mesh data, where the topological relationships among points, edges, faces, and cells are maintained through multiple interlinked index arrays. These indirect data access operations disrupt spatial locality and cause irregular memory traffic, resulting in low cache utilization and high latency. Therefore, despite substantial progress in GPU acceleration of structured or semi-structured atmospheric models, unstructured FEM-based frameworks—characterized by irregular data dependencies and adaptive mesh refinement—still lack systematic and efficient GPU solutions. To address this research gap, this study focuses on the GPU parallelization and optimization of the two most time-consuming components of Fluidity-Atmosphere: **element-wise computations** (e.g., evaluating local mass and stiffness matrices via numerical integration) and **global matrix assembly**. The primary contributions are as follows.

1. **GPU-based element-wise computations:** Multiple optimization techniques are applied to accelerate compute-intensive kernels via GPU offloading, achieving more than $100\times$ speedups compared with CPU versions.
2. **GPU-based global matrix assembly:** Parallel strategies are designed and implemented for element-wise assembly.
3. **Integrated GPU implementation:** The GPU kernels are integrated into Fluidity-Atmosphere with pinned memory to optimize data transfer. The results show that GPU kernels facilitate speedups exceeding two orders of magnitude, with 4 MPI processes and one GPU, achieving an overall acceleration of 8.57 times compared to a single CPU process.

The remainder of this paper is organized as follows. Section 2 introduces the program structure of Fluidity-Atmosphere, with emphasis on element-wise computations and global matrix assembly. Section 3 presents the GPU implementation and optimization of element-wise computations. Section 4 describes GPU-based global matrix assembly strategies. Section 5 evaluates the integrated GPU implementation and its performance. Section 6 concludes the paper, highlighting the future scope.

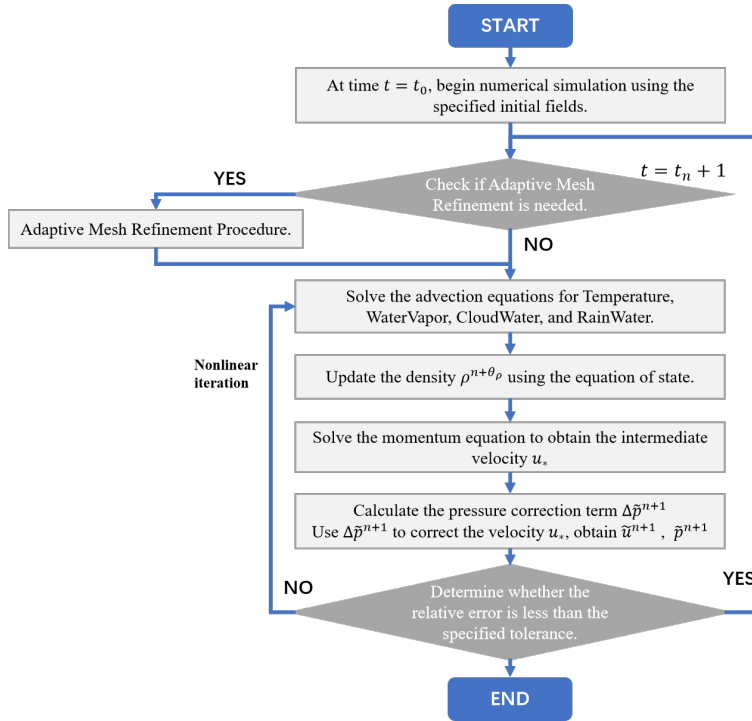


Figure 1. Time loop of Fluidity-Atmosphere.

90 2 Target Scientific Application Definition: Fluidity-Atmosphere

Fluidity-Atmosphere employs a mixed finite element/finite volume framework that incorporates the Navier-Stokes momentum equations, advection-diffusion equations for potential temperature and water vapor, and the compressible continuity equation within an anisotropic adaptive mesh algorithm. This design facilitates dynamic co-optimization of the mesh and physical fields. Fluidity-Atmosphere solves a coupled system of nonlinear equations with (typically) time-varying solutions. The time-
 95 marching algorithm employed uses a nonlinear iteration scheme known as Picard iteration in which each equation is solved using the currently optimal solution for the other variables. The dynamical framework supports both the continuous Galerkin (CG) and discontinuous Galerkin (DG) finite element formulations. This study is focused exclusively on the CG scheme. Specifically, in the performance evaluations and benchmarks presented in this study, the dynamic framework predominantly utilizes linear Lagrange elements (P_1) for the continuous Galerkin (CG) discretization. This choice inherently yields 4 degrees
 100 of freedom per tetrahedral element, resulting in the 4×4 local element matrices evaluated during the simulation.

Figure 1 illustrates the general iteration loop of Fluidity-Atmosphere (Li et al., 2021). Fluidity-Atmosphere could invoke the adaptive algorithm at regular timesteps to ensure that the dynamics do not extend beyond the zone of adapted resolution. The adaptive algorithm has been parallelized using MPI processes, which run efficiently on CPUs. In our performance profiling, the nonlinear iterations were predominant. In each timestep iteration, the model solves governing equations such as the advection-

105 diffusion and momentum equations. The processes of solving these equations in Fluidity-Atmosphere typically are divided into two major computational stages:

- (1) **Construct matrices:** For each element, numerical integration is performed using Gaussian quadrature to evaluate local stiffness, mass matrices, and right-hand side (RHS) vectors. These local contributions are subsequently inserted into the global sparse matrix system.
- 110 (2) **Solution of linear systems:** The assembled sparse matrices are solved using parallel numerical libraries such as PETSc, employing iterative solvers and preconditioners well-suited for large-scale sparse problems. The PETSc library has good performance and scalability. In the simulated case studies, its contribution to the overall runtime was a significant portion (37%) in the CPU-only baseline, and as detailed later, it becomes the dominant bottleneck once the matrix assembly is accelerated on the GPU. In addition, PETSc already supports GPUs (Mills et al., 2021). Hence, future studies will be
115 focused on the integration of Fluidity-Atmosphere and PETSc GPU computing.

The computation times of constructing matrices contributes significantly to the overall runtime. For instance, the Mountain-wave3D simulation, in which the number of mesh nodes and mesh elements are 103,635 and 554,394, respectively. Table 1 shows the computational times of constructing matrices in a timestep. In particular, the pressure diffusion matrix is used to calculate the stabilization term, which is only calculated at the first timestep or when the grid changes because of the adaptive
120 process. In the Mountainwave3D simulation, the pressure diffusion matrix is computed every ten time steps. Without calculating the stabilization term, the total time of the timestep will decrease, but the proportion of time required to construct the matrices will slightly increase. As quantified in Table 1, the processes of constructing matrices account for approximately 60% of the total timestep duration in both scenarios (with and without the stabilization term). This high proportion confirms that accelerating **element-wise computations** and **global assembly** is the most effective strategy for improving the overall
125 performance of the Fluidity-Atmosphere model.

The processes of constructing matrices share highly similar loop patterns. Hence, their general workflow can be summarized as follows:

- Step 1 **Preparing Data:** Initialize element mass matrices, right-hand sides, and local physical variables (such as velocity, density, and viscosity).
- 130 Step 2 **Coordinate transformation:** The `transform_to_physical` function maps shape function gradients into physical space and computes Gaussian weights (`detwei`).
- Step 3 **Setup test function.**
- Step 4 **Contribution evaluation:** Call physics modules to accumulate contributions into element matrices and right-hand sides. Several functions, each corresponding to a distinct physical process, such as mass, advection, diffusion, source,
135 and viscosity terms, are called. In addition to the calculation of physical variables, these functions extensively call the integration calculation function on the elements such as `dshape_tensor_dshape`, `shape_dshape`, and

shape_shape. In the process of assembling the pressure diffusion matrix, a special function named `get_edge_lengths` (in which the small matrix solver is called) is used to calculate the element length scale in the physical space. All these functions are well-suited for GPU parallel computing.

140 Step 5 **Assembly**: Local contributions are inserted into the global matrix and RHS. The subsequent irregular insertion operations into the global matrix makes it highly data-intensive and memory-bound. Owing to the massive throughput of GPU device memory, these functions can be accelerated on GPUs.

Therefore, the following sections describe the implementation and performance optimization methods of the element-wise computations and global matrix assembly on the GPU.

145 Although the current implementation is specifically optimized for the NVIDIA A100 GPU utilizing the CUDA programming model (NVIDIA, 2025a), the proposed optimization strategies are largely hardware-agnostic at the algorithmic level. Specifically, the element-wise parallelization strategy (one-thread-per-element), the data layout design, and the sparse matrix assembly workflow are general and can be applied to other GPU architectures.

However, certain implementation details rely strictly on CUDA-specific hardware features, such as double-precision atomic
150 operations in the L2 cache, memory hierarchy optimization, and specific kernel launch configurations. Porting the implementation to other platforms, such as AMD GPUs, would primarily involve adapting CUDA-specific APIs to alternative frameworks like HIP (AMD, 2025). Since HIP provides similar abstractions for thread hierarchy, memory management, and atomic operations, the overall syntax porting effort is expected to be moderate.

Nevertheless, achieving true performance portability would require additional tuning to account for architectural differences
155 in memory bandwidth, cache structure, and execution models. For instance, while NVIDIA V100 and A100 GPUs handle massive atomic updates highly efficiently, evaluating similar unstructured-grid applications on AMD CDNA architectures (e.g., MI100) has shown that standard atomic updates can incur higher penalties, sometimes necessitating alternative data restructuring or register-based aggregation to achieve optimal efficiency (Stone et al., 2021). Overall, the proposed algorithmic approach is expected to maintain its fundamental effectiveness across heterogeneous architectures, although achieving peak
160 optimal performance on different platforms will inevitably require platform-specific tuning.

3 GPU Parallelization of Element-wise Computations

In Fluidity-Atmosphere, element-level operations such as numerical integration and coordinate transformation constitute the core of unstructured finite element-wise computations. Profiling results indicate that these routines are invoked millions of times per timestep. Each nonlinear iteration involves re-evaluating element Jacobians, transforming basis functions, and accumulating
165 physical contributions for all elements, which cumulatively dominate the computational workload. These observations identify numerical integration and coordinate transformation as the primary performance bottlenecks and thus the key focus of GPU parallelization in this study. The computation for each mesh element can be performed independently of others, making this step an ideal candidate for parallelization (Georgescu et al., 2013). GPUs feature a massively parallel processor architecture,

Table 1. The computation times of constructing matrices

Functional module	With stabilization term		Without stabilization term	
	CPU time(ms)	Percentage	CPU time(ms)	Percentage
The advection-diffusion matrix (Temperature)	2477.545	4.84%	2466.594	5.95%
The advection-diffusion matrix (WaterVapor)	2477.194	4.82%	2471.302	5.96%
The advection-diffusion matrix (CloudWater)	2466.853	4.82%	2494.424	6.02%
The advection-diffusion matrix (RainWater)	1785.827	3.49%	1799.104	4.34%
The momentum matrix	13029.408	25.47%	12093.145	29.19%
The divergence matrix	3662.082	7.16%	3634.111	8.77%
The pressure diffusion matrix	3505.105	6.85%		
The projection matrix	1126.747	2.20%	1125.822	2.72%
Total of constructing matrices	30530.760	59.68%	26084.502	62.96%
Others	20625.871	40.32%	15348.450	37.04%
Total timestep duration	51156.631	100.00%	41432.952	100.00%

which enables the simultaneous launch of a large number of parallel threads. These threads can be associated with different mesh elements to execute element-wise computations. These results are stored for subsequent assembly into the global matrix. Building on the workflow analysis in Section 2, this section presents the GPU implementation and optimization strategies for key subroutines, followed by a summary of general acceleration techniques.

3.1 Data Structures and Parallelization Strategy

In unstructured meshes, element connectivity must be explicitly stored. Each tetrahedral element in Fluidity-Atmosphere records four vertex indices. Physical variables are stored in a unified Fields data structure organized as (number of vertices $\times$ dimension) arrays: scalars (1D), vectors (3D), and tensors (9D). The same layout is adopted in GPU memory to ensure contiguous access and consistent indexing (Figure 2). Common fields such as density and temperature are stored as scalars, positions and velocities as vectors, and viscosity terms as tensors.

The element-wise computations in the FEM are intrinsically independent, making them particularly suitable for massively parallel execution on GPU architectures. Unlike spectral methods that involve global communication or high-order finite-difference schemes relying on extended stencils, finite-element computations are confined strictly within each element boundary. It is important to clarify that while Fluidity-Atmosphere utilizes a mixed CG/DG discretization, the GPU acceleration in

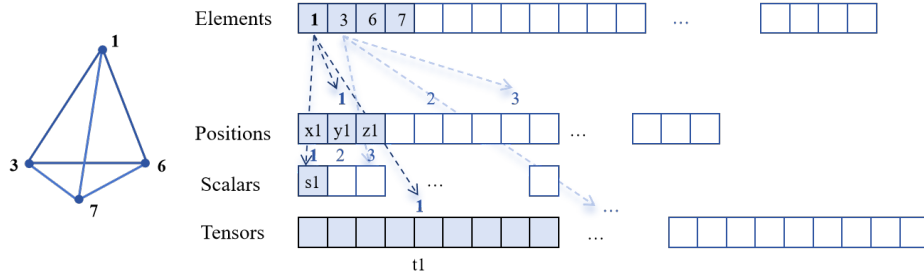


Figure 2. Data layout of elements and variables.

this study focuses exclusively on the Continuous Galerkin (CG) operators (e.g., for the momentum and advection-diffusion equations). In the CG formulation, the numerical integration to evaluate local matrices remains mathematically entirely local. Unlike DG methods, which require calculating inter-element fluxes and jumps over facets, the CG element-wise computations do not introduce cross-element mathematical dependencies during the integration phase. The global coupling arising from shared mesh nodes is manifested exclusively during the global matrix assembly phase through indirect memory access.

This strong locality offers two principal advantages: (1) all element-wise computations can proceed independently, eliminating inter-thread communication during the compute phase; and (2) the workload is naturally balanced across elements, minimizing synchronization overhead. Consequently, a “one-thread-per-element” parallelization strategy is adopted. Each GPU thread handles one element. In our implementation, thread block sizes of 128 and 256 were both utilized, with the specific parameter for each kernel determined empirically via NVIDIA Nsight Compute. For the target mesh comprising 554,394 elements, a block size of 128 yields 4,332 blocks, while 256 yields 2,166 blocks. Given that the NVIDIA A100 GPU features 108 Streaming Multiprocessors (SMs), each with a theoretical maximum of 2,048 resident threads, both configurations generate sufficient numbers of thread blocks to provide approximately 2.5 execution waves across the GPU, which helps maintain high occupancy and effectively hide memory latency.

3.2 GPU Parallel Coordinate Transformation

With the data layout on GPUs established, the next step is to implement the core finite element operation: coordinate transformation. Using a reference element (in this study, a tetrahedron), shape functions and Gauss points are consistently defined. In Fluidity-Atmosphere, each physical element employs 11 Gauss points: four vertices, six edge midpoints, and the centroid (Figure 3). For each Gauss point, the transformation routine evaluates the integrand and multiplies it by the predefined weight.

Even though element-independent and theoretically parallelizable, this kernel is computationally intensive and one of the most frequently executed routines in the model. For every iteration and every element, it performs numerous small matrix operations: matrix multiplications, adjoint and determinant evaluations, and transformations of derivative matrices.

Algorithm 1 outlines the `transform_to_physical` procedure, where input and output arrays are stored contiguously to maximize memory efficiency. Crucially, the algorithm incorporates a conditional optimization for linear (P_1) elements:

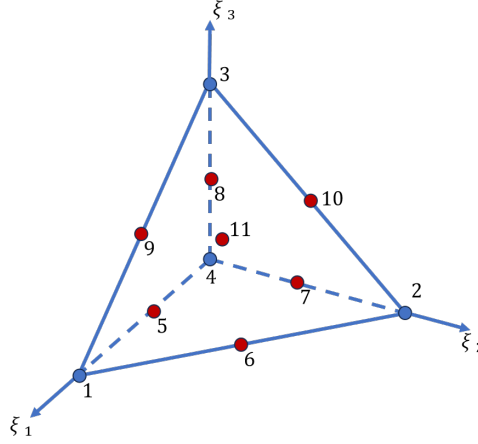


Figure 3. The reference element and Gauss points.

because standard affine mappings yield a constant Jacobian matrix, expensive operations (Jacobian construction, determinant, and inverse) are evaluated only at the first Gauss point ($g_i = 1$) and cached.

Algorithm 1 Transform_to_physical

Input: Element nodal coordinates X , shape functions N , reference derivatives $\nabla_\xi N$, Gauss points g_i

Output: Physical derivatives $\nabla_x N$, weights detwei , J , J^{-1} , $\det J$

```

for  $g_i = 1$  to  $n_{gi}$  do
  compute  $\nabla_\xi N(g_i)$ 
  if mapping is nonlinear or  $g_i == 1$  then
    construct Jacobian matrix  $J \leftarrow X \cdot (\nabla_\xi N)^T$ 
    compute  $\text{adj}(J)$  and  $\det J$ 
    normalize inverse matrix  $J^{-1} \leftarrow \text{adj}(J) / \det J$ 
  end
  transform to physical derivatives  $\nabla_x N \leftarrow J^{-1} \cdot \nabla_\xi N$ 
  compute quadrature weights  $\text{detwei}(g_i) \leftarrow |\det J| \cdot w_{gi}$ 
end

```

The inner loop over Gauss points is deliberately retained to preserve framework generality for higher-order elements. Furthermore, rather than precomputing and storing these geometric quantities, which would exacerbate memory traffic in this heavily memory-bound framework, we employ a "compute-on-the-fly" strategy, trading abundant GPU arithmetic cycles to conserve precious memory bandwidth (Kronbichler and Kormann, 2012; Böhm et al., 2025).

3.3 GPU Parallelization of Element-wise Integration

215 Element-wise integration routines constitute the computational core of unstructured finite-element models, as they are responsible for evaluating the contributions of each element to the global system. Profiling of Fluidity-Atmosphere indicates that these kernels are among the most frequently executed routines. During each nonlinear iteration, they are called once per element per physical equation, resulting in millions of invocations per timestep. As each call involves multi-dimensional tensor operations and numerical integration over multiple Gauss points, the accumulated cost dominates both the element-computation phase and the overall runtime of the model. Therefore, optimizing element integration is critical to achieving substantial end-to-end
220 performance gains. The principal integration functions are as follows:

- **dshape_tensor_dshape** couples shape function gradients with tensor fields, supporting stabilization terms and physical field coupling.
- **shape_shape** integrates shape functions to construct the mass matrix, required in momentum and energy conservation equations.
- 225 – **shape_dshape** performs weighted integration of shape functions and their gradients at nodal points.

Among these functions, the `dshape_tensor_dshape` function is the most computationally intensive, as it involves both tensor-vector and vector-vector multiplications for every integration point. On the GPU, this procedure is parallelized at the thread level, with each thread processing one element. Beyond kernel-level parallelization, a series of GPU-specific performance tuning techniques are applied to completely exploit the hardware potential. These optimizations are guided by profiling
230 by employing NVIDIA Nsight Compute, focusing on reducing memory latency, register pressure, and control overhead:

- **Memory optimization:** The `__restrict__` qualifier is applied to pointer variables to eliminate aliasing and enable more aggressive memory-access optimizations. Small `__device__` functions are annotated with `__forceinline__` to reduce function-call overhead.
- **Loop optimization:** Loop-invariant expressions (particularly global memory accesses) are hoisted outside the loop body; manual loop unrolling or `#pragma unroll` directives are applied to innermost loops to reduce control overhead and
235 increase instruction-level parallelism.
- **Template metaprogramming:** As the dimensions of the different matrices and vectors are already known at compile time, the `__device__` functions such as matrix-matrix product and dot product are implemented as templates, facilitating the compiler to apply deeper optimizations for specific instantiations.

240 These measures are reflected in Nsight Compute reports as improved register utilization and reduced memory bottlenecks. Loop unrolling and templating yield substantial enhancements in small-scale computations. In combination with GPU-parallel element integration and small-matrix solvers, these optimizations form a cohesive strategy that maximizes performance while maintaining full consistency with the original CPU implementation.

It should be noted that the CPU baseline used for comparison represents the official, unmodified production code of Fluidity-
 245 Atmosphere, compiled with aggressive optimization flags `-O3 -ffast-math`. Under these flags, modern CPU compilers
 automatically apply substantial loop unrolling and vectorization. However, due to the warp-based execution model and strict
 register allocation of GPUs, the CUDA compiler frequently requires explicit manual directives (e.g., `#pragma unroll`) and
 compile-time size resolution (via templates) to achieve optimal instruction scheduling and register reuse. Thus, these GPU-
 specific manual optimizations are implemented to fully unlock the hardware’s architectural potential rather than to introduce
 250 an algorithmic disparity. The optimized code is listed in **Appendix A**.

3.4 GPU Parallelization of Small-Scale Linear Algebra Solvers

In addition to numerical integration, the construction of stabilization terms involves numerous small-scale linear algebraic
 operations, such as solving 6×6 systems or computing eigen-decompositions of 3×3 matrices. Even though these operations
 are apparently lightweight, they are called repeatedly for every element and Gauss point, resulting in a substantial cumulative
 255 cost. On CPUs, relying on standard library routines such as LAPACK’s `DSPEV` or Fortran’s intrinsic solvers for these operations
 is highly inefficient. While a small 6×6 matrix easily fits into the L1 data cache, gathering these matrices from an unstructured
 mesh via indirect addressing disrupts spatial locality. Furthermore, these independent tiny matrices offer minimal opportunity
 for data reuse, limiting cache efficiency. From a software perspective, pre-compiled external library routines typically cannot
 be automatically inlined by the compiler within the tight element loop. Invoking them millions of times introduces profound
 260 function-call overhead, and their internal argument validation and branching dwarf the actual floating-point arithmetic.

To address this challenge, this study implements dedicated GPU-based small-matrix solvers and eigen-decomposition ker-
 nels as inline `__device__` operators.

- **Small linear systems (e.g., 6×6):** These are solved on GPUs using Gaussian elimination. The input is an augmented
 matrix A (in column-major layout, including the right-hand side) with dimensions $m \times n$. Pivoting and row exchanges are
 265 applied to control numerical errors, whereas column-major storage and in-place operations reduce register pressure and
 facilitate the alignment with GPU memory characteristics.
- **Eigen-decomposition of 3×3 real symmetric matrices:** Both iterative and direct GPU solvers are implemented, based
 on the methods employed in prior research. To improve efficiency, the implementation employs template-based array
 classes and includes sorting functions for eigenvalues and eigenvectors. By exploiting symmetry, only six entries (the
 270 diagonal and upper triangular part) are stored, minimizing register usage and avoiding spills to local memory.

These GPU-implemented solvers are lightweight and reusable across kernels. The modular `__device__` design also pro-
 motes code reuse across physical modules and supports further fusion with integration kernels.

4 GPU Parallelization of Matrix Assembly

In the FEM, matrix assembly serves as the critical bridge between local element-wise computations and the solution of global linear systems. As this stage involves extensive operations on sparse data structures and frequent memory access, it typically hinders the performance of the entire program. Therefore, migrating the matrix assembly procedure to GPUs and applying targeted optimizations is crucial for enhancing the computational efficiency of atmospheric models. This section introduces the global matrix storage formats, the element-wise parallel assembly strategy, and the GPU implementation of RHS vector assembly.

Before detailing our assembly strategy, it is worth noting that a prominent trend in high-performance finite-element simulations is to discard global matrix assembly in favor of matrix-free (or partially assembled) operator evaluations (Kronbichler and Kormann, 2012, 2019; Rudi et al., 2015). Matrix-free algorithms were fundamentally developed to address the performance bottlenecks inherent in high-order finite elements. In high-order discretizations, full matrix assembly leads to an exponential explosion in memory footprint and floating-point operations per degree of freedom (DoF). Matrix-free approaches, typically coupled with sum factorization, suppress this complexity, thereby drastically increasing computational intensity and parallel efficiency (Fischer et al., 2020).

However, the dynamic framework evaluated in this study predominantly employs low-order linear elements (P_1). For low-order FEM, the memory access per DoF is intrinsically small and remains comparable whether using fully assembled operators or on-the-fly matrix-free evaluations. Since low-order evaluations are heavily memory-bound, transitioning to matrix-free methods does not yield the profound bandwidth savings observed in high-order regimes. Furthermore, explicitly assembling the global sparse matrix provides a critical advantage: it enables the use of highly optimized, generic Sparse Matrix-Vector multiplication (SpMV) kernels and sophisticated algebraic preconditioners, such as Algebraic Multigrid (AMG) and Incomplete LU (ILU) provided by PETSc, which are indispensable for the stability and efficiency of the implicit time-stepping solvers in Fluidity-Atmosphere. Therefore, accelerating the global sparse matrix assembly remains the most rational and impactful optimization pathway for this class of low-order dynamical cores.

4.1 Global Matrix Assembly

Following local integration, the algorithm accumulates (A^e, F^e) into the global matrix A and vector F according to the connectivity information of the mesh. This procedure, widely recognized as the global assembly, transforms element-level contributions into a consistent global representation based on the mapping between local and global degrees of freedom, as shown in **Algorithm 2**. Here, ε is the set of elements and e is an element. The *elem* function performs element-wise computations to get element matrices. The *L* function retrieves the node numbering.

Algorithm 2 The Global Matrix Assembly

Output: Global matrices A , F Initialize A , F to zero

```
forall  $e \in \varepsilon$  do
  Compute  $(A_e, F_e) = \text{elem}(e)$ 
  forall local degrees of freedom  $d_1$  of  $e$  do
     $F(L(e, d_1)) += F_e(d_1)$ 
    forall local degrees of freedom  $d_2$  of  $e$  do
       $A(L(e, d_1), L(e, d_2)) += A_e(d_1, d_2)$ 
    end
  end
end
```

In contrast to structured-grid models, where mesh nodes are arranged in a regular (i, j, k) indexing order and memory access can be performed sequentially and predictably, unstructured meshes lack geometric regularity. The connectivity of each element varies depending on its neighbours, and the mapping between local and global indices must be retrieved indirectly through lookup tables. Consequently, the assembly phase in unstructured FEM involves frequent non-contiguous global memory reads and writes, with irregular strides between consecutive memory locations. This destroys spatial locality and drastically increases cache miss rates, thereby amplifying the ratio of memory access time to floating-point computation.

On GPU architectures, these irregular access patterns exacerbate performance degradation. Concurrent threads typically need to update overlapping entries of the global sparse matrix, requiring synchronization or the use of atomic operations to ensure correctness. The combined effects of irregular global memory traffic, low cache reuse, and write conflicts make matrix assembly a highly memory-bound process.

Consequently, global matrix assembly presents two major challenges: (1) the imbalance between memory traffic and arithmetic intensity, caused by the dominance of irregular global data movement over computation; (2) the difficulty of ensuring data consistency during concurrent updates. Cumulatively, these factors make matrix assembly another critical performance bottleneck, and thus a central focus of GPU parallel optimization in this work.

4.2 Sparse Matrix Storage Formats

Fluidity-Atmosphere employs different sparse storage formats for different types of physical fields. For scalar field matrices, the **compressed sparse row (CSR)** format is used. For vector field matrices such as three-component velocity, the **block-CSR (BCSR)** format is adopted.

For a scalar field, the local stiffness matrix can be expressed as Eq. (1):

$$K_{ij} = \int_{\Omega_e} (\nabla N_i)^T C (\nabla N_j) d\Omega \quad (1)$$

where N_i and N_j are basis functions, Ω_e is the integration domain of the element, and C is the associated tensor. For scalar fields, N_i is a scalar and K_{ij} is a scalar integral, resulting in a 4×4 local element matrix. By obtaining the global node indices n_i and n_j corresponding to the local nodes i and j , the result K_{ij} is updated in the global matrix entry (n_i, n_j) .

325 The global matrix is typically stored in the CSR format, which consists of three arrays: **values** (non-zero entries), **col_index** (column indices), and **row_pointers** (row offsets) (NVIDIA, 2025b). Within a row range, the target column can be located using binary search.

For vector fields, each node contains three degrees of freedom (DOFs). Thus, while assembling the local stiffness matrix, all interactions among the nodes and their DOFs must be considered. This results in a set of 4×4 blocks, each of size 3×3 :

$$330 \quad K = \begin{bmatrix} K_{11} & K_{12} & K_{13} & K_{14} \\ K_{21} & K_{22} & K_{23} & K_{24} \\ K_{31} & K_{32} & K_{33} & K_{34} \\ K_{41} & K_{42} & K_{43} & K_{44} \end{bmatrix} \quad (2)$$

Each element of the element matrix K_{ij} is actually a 3×3 matrix.

$$K_{ij} = \begin{bmatrix} a & b & c \\ d & e & f \\ g & h & i \end{bmatrix} \quad (3)$$

Here, Fluidity adopts the BCSR format, which is structurally similar to CSR, but stores continuous block data in the value array. This block-based representation better supports sparse operations for vector fields.

335 4.3 GPU Parallel Element-Wise Assembly

In our GPU parallelization, the contributions of each mesh element are computed and stored in parallel, with each GPU thread responsible for one element. A dedicated kernel is launched to assemble element matrices.

In the implementation, each thread iterates over the entries K_{ij} of the local element matrix and looks up the corresponding global non-zero indices. Given the global node indices n_i and n_j , the **row_pointers** array is first used to locate the row range [340 row_pointers(n_i), row_pointers($n_i + 1$)). Within this range, the column index n_j is searched using binary search. To enhance the efficiency, the implemented GPU assembly function consistently adopts binary search.

A significant challenge arises during parallel write-back, where race conditions may occur if multiple threads attempt to update the same non-zero entry simultaneously. An atomic operation is indivisible and guarantees consistency by preventing interference from other threads. Modern GPUs provide hardware-level support for atomic addition, enabling conflict-free 345 concurrent updates without explicit locks. In this work, the atomic add approach is adopted for its efficiency and simplicity. The CUDA parallel assembly code is listed in **Appendix B**.

To summarize the complete execution pipeline, the global matrix assembly occurs entirely within the GPU device memory. Following the element-wise computations, the local element matrices are accumulated into the global sparse matrix arrays (i.e.,

the `values` array in CSR/BCSR format), which are pre-allocated on the GPU. This is achieved using the aforementioned parallel assembly kernel equipped with hardware-level `atomicAdd` operations. Upon the completion of the assembly kernel, the fully constructed global matrix arrays are transferred back to the CPU host memory via PCIe. The assembled matrix is then handed over to the CPU-bound PETSc library to perform the subsequent large-scale linear system solve.

4.4 GPU Parallel Assembly of Right-Hand-Side (RHS) Vectors

Similar to the global matrix, the RHS vector is assembled using an element-wise parallel strategy. Each GPU thread computes the nodal contributions of one element and accumulates them into the corresponding global vector entries.

The implementation first determines the indices of the nodes belonging to each element, then loads the nodal values, and finally applies `atomicAdd` operations to ensure correctness during concurrent accumulation. For cases where multiple variables are associated with a node, appropriate offsets are added to the corresponding positions. The **Appendix C** illustrates the procedure of RHS assembly, which achieves efficient large-scale parallel accumulation while preserving correctness.

5 Results and Evaluation

5.1 Experimental Environment

This section presents a systematic analysis of the performance of the proposed GPU parallelization strategies. Experiments were conducted to first validate correctness on both CPU and GPU platforms, and then to evaluate performance in four aspects: element-wise computation, global matrix assembly, data transfer with computation overlap, and overall model performance. The GPU and CPU experimental environments are shown in Table 2. All CPU baseline evaluations were executed using the specified number of CPU cores (MPI processes) on a single physical AMD EPYC 7713 processor.

GPU acceleration was applied to the momentum equation, advection equation, and stabilization term construction of Fluidity-Atmosphere. Then, the GPU-accelerated version was validated for correctness and evaluated for performance to ensure both accuracy and efficiency.

5.2 Validation Methodology

For validation of correctness, the 3D idealized mountain wave test case (Li et al., 2021) was used. The mountain wave test serves as a crucial validation procedure for assessing the dynamic framework model of the Fluidity-Atmosphere model in simulating orography-induced airflow. By comparing numerical results against theoretical solutions of idealized orographically forced flows, this test effectively evaluates the capability of the model to capture mountain wave generation and propagation phenomena under complex topographic conditions. The computational domain spans 60 km in both horizontal dimensions with a vertical extent of 16 km. Adaptive mesh resolution dynamically ranges from 125 m to 10 km throughout the simulation domain, with mesh refinement actively guided by variations in velocity magnitude and potential temperature. The mesh consists of 554,394 elements and 103,635 nodes. A 3D bell-shaped mountain profile is mathematically described as follows:

Table 2. GPU & CPU and compiler information.

GPU NVIDIA	A100
Peak perf. (FP64)	9.7 TFLOPS
Max Clock Freq	1410 MHz
L1 Cache	192 KiB
L2 Cache	40960 KiB
SM per GPU	108
Register file size per SM	256 KiB
Device Memory Bandwidth	1555 GB/s
PCIe Bandwidth	32 GB/s
nvcc	13.0
CPU	AMD EPYC 7713 64-Core Processor
ISA	X86_64
Max Clock Freq	3720 Mhz
L1 Cache	8 MiB
L2 Cache	64 MiB
Cores	64
Memory Bandwidth	204.8 GB/s
gcc	13.3
CMake	3.22.1
gfortran	13.3
mpicc	MPICH 3.4.2
Optimization	-O3 -ffast-math

$$h(x, y) = \frac{h_0}{\left(1 + \frac{x^2 + y^2}{a^2}\right)^{\frac{3}{2}}} \quad (4)$$

where $h_0 = 400m$ represents the peak elevation and $a = 1000m$ denotes the characteristic half-width parameter. The stratified atmospheric background is characterized by $N = 0.01s^{-1}$, while the surface potential temperature initializes at $\theta_0 = 293.15K$. The incoming flow maintains a constant velocity of $u = (10, 0, 0)^T m/s$. For numerical stability, an absorbing layer is implemented in the upper atmospheric region (10-16 km-altitude), with additional dissipative layers extending for 10 km inward from all lateral boundaries (comprehensive specifications available in Li et al., 2021).

In scientific applications demanding high-precision floating-point operations, microscopic rounding differences between CPU and GPU architectures are inevitable. In a fully dynamic simulation, these bit-wise differences can cause the Adaptive

Mesh Refinement (AMR) algorithm to split elements differently over time. Comparing results across divergent meshes necessitates spatial interpolation, which introduces algorithmic noise and obscures the underlying arithmetic consistency. To eliminate structural variations and rigorously validate the GPU implementation, we conduct evaluations under two complementary fixed-
 390 mesh scenarios.

First, to validate the core operators during the non-stationary initial stage, we extract the refined unstructured mesh generated at $t = 3000$ s from an adaptive run and employ it as a fixed, static grid right from the beginning of the simulation ($t = 0$). Both the CPU and GPU versions simulate the transient phase up to $t = 3000$ s with the AMR module deactivated. While this static mesh is not dynamically optimized for the early stages, it ensures an identical degree-of-freedom layout from $t = 0$ and
 395 completely avoids spatial interpolation errors. Second, to verify long-term stability in the steady-state regime, the simulation is run up to $t = 3000$ s under the standard CPU-AMR configuration until the wave profile stabilizes. The mesh is then fixed, and both versions advance the physical time independently from $t = 3000$ s to $t = 5000$ s. The computational results for both scenarios are presented in Figure 4. For the transient fixed-mesh run (Figs. 4b and 4c), the absolute differences between the CPU and GPU fields at $t = 3000$ s peak at 1.44×10^{-14} for the vertical velocity (w) and 4.73×10^{-15} for the potential temperature
 400 perturbation (θ'). For the steady-state continuation run (Figs. 4d and 4e), the differences remain bounded within a similar magnitude at $t = 5000$ s. These minuscule variations strictly correspond to the double-precision machine epsilon, confirming that the GPU-accelerated operators are fundamentally and mathematically consistent with the original CPU baseline across all simulation phases.

5.3 Results and Discussion of Performance Optimization

405 5.3.1 GPU Element-wise Computation Performance

In unstructured meshes, element node values are typically stored non-contiguously in memory. Thus, the performance of element-wise computation is primarily limited by memory access efficiency. While CPUs incur high memory overhead for non-contiguous access, GPUs leverage high bandwidth and massive thread parallelism to alleviate this bottleneck. Table 3 compares GPU and CPU performance in scalar, vector, and tensor data access. Specifically, functions such as `ele_val_scalar` serve
 410 as data-gathering routines that collect scattered node-based field values from global arrays into contiguous element-local arrays for each element prior to numerical integration. The results demonstrate the GPU’s efficiency in accelerating these heavily memory-bound ‘gather’ operations. This is primarily because arrays storing scalar fields have a simpler, lower-dimensional layout, facilitating GPUs to completely leverage high-bandwidth access. All the performance profiling data were collected with the mesh in the Mountainwave3D simulation, in which the number of mesh nodes and elements are 103,635 and 554,394,
 415 respectively.

Furthermore, Table 4 presents the GPU performance of numerous core functions in Fluidity-Atmosphere (coordinate transformation, integration, and auxiliary routines). Results indicate that parallelizing dense element-wise computations on GPUs significantly accelerates the performance. All listed functions achieved more than **103** $\times$ speedups compared with CPU versions. These functions are computationally intensive, involving repeated local linear system and eigenvalue/eigenvector solves

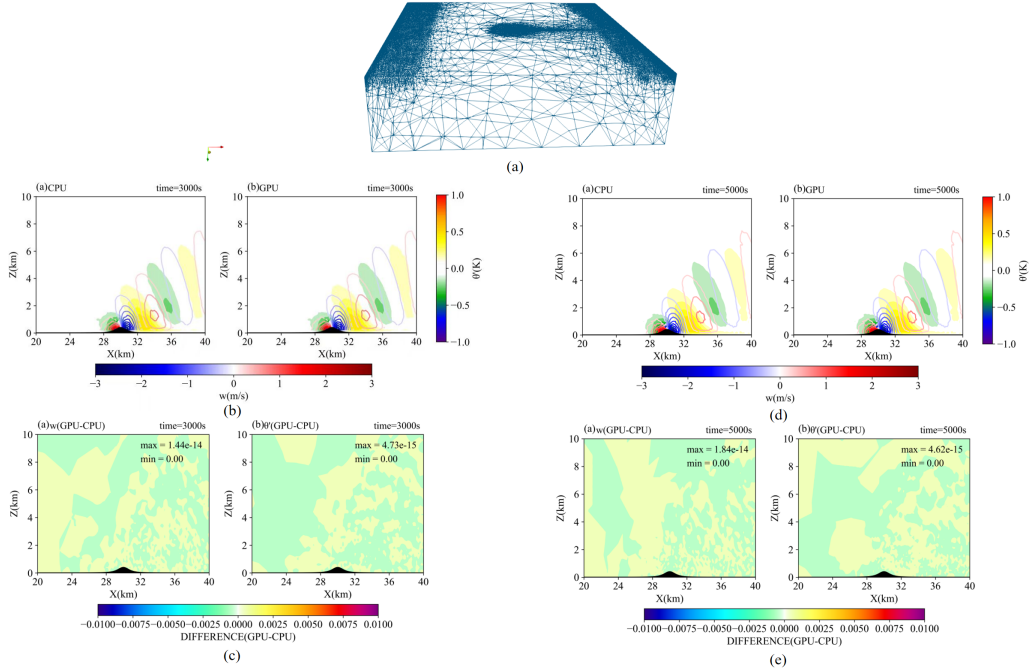


Figure 4. Comparison between the CPU baseline and the GPU-accelerated computation results under fixed-mesh configurations. (a) The unstructured mesh layout extracted at $t = 3000$ s; (b) Simulated vertical velocity (w) and potential temperature perturbation (θ') fields at $t = 3000$ s from the transient validation run (simulated continuously from $t = 0$ using the fixed-mesh); (c) Absolute differences between CPU and GPU results at $t = 3000$ s for the transient fixed-mesh run; (d) Simulated fields at $t = 5000$ s from the steady-state stability validation run (advanced from $t = 3000$ s on the fixed-mesh); (e) Absolute differences between CPU and GPU results at $t = 5000$ s for the steady-state stability run.

Table 3. GPU Performance of GPU parallel element data access (Mesh: 103,635 nodes, 554,394 elements; timings are averaged over 1,000 invocations).

Function	CPU time (ms)	GPU time (ms)	Speedup
ele_val_scalar	6.628	0.0276	240.14
ele_val_vector	13.633	0.0768	177.51
ele_val_tensor	32.345	0.1976	163.69

420 for each element in single-core CPU execution. In contrast, the GPU can execute such highly independent tasks concurrently across thousands of threads, comprehensively leveraging parallelism.

To analyse the effects of GPU-oriented optimizations, we take the `dshape_tensor_dshape` function as an example. This function is a hotspot in finite-element computation, involving frequent access to gradients of shape functions at Gauss

Table 4. GPU kernel performance in element-wise computation (Mesh: 103,635 nodes, 554,394 elements).

Function	CPU time (ms)	GPU time (ms)	Speedup
transform_to_physical	525.575	3.964	132.59
dshape_tensor_dshape	1865.290	1.969	947.33
shape_shape	104.221	1.012	103.01
shape_dshape	246.047	1.371	179.41
get_edge_lengths	1369.171	7.476	183.14

Table 5. Performance of dshape_tensor_dshape with optimizations (Mesh: 554,394 elements).

Version	dshape1 and dshape2 are the same array.	dshape1 and dshape2 are different arrays.
	Time (ms)	Time (ms)
Baseline	35.799	36.314
Unroll	20.847	26.890
Templated matrix product and dot product	1.969	3.689

points and dense matrix multiplications/dot products. Table 5 shows performance enhancement due to loop unrolling and optimization via the use of template functions. Notably, Fluidity-Atmosphere typically calls this routine with identical dshape1 and dshape2 arrays. However, for generality we also tested cases with two different dshape1 and dshape2 arrays. While the same array case performs better, both scenarios show consistent performance gains.

Nsight Compute profiling further indicates that templated matrix product and dot product increased memory throughput from 70 % to 82.35 %, slightly improved compute throughput, and enhanced L1 cache hit rates. SASS code inspection revealed that baseline __device__ functions were only locally reordered, thereby limiting the performance. After templating, the compiler applied deeper optimizations, interleaving SASS instructions between __device__ and __global__ functions, yielding significant performance enhancements.

These results confirm that GPU parallelization achieves order-of-magnitude acceleration in element-wise computations, while targeted optimizations further unlock the hardware potential.

5.3.2 Performance of Global Matrix Assembly Optimization

Table 6 shows assembly performance for the advection-diffusion matrix (Temperature) and the right-hand sides. Compared with CPU execution, GPU parallel assembly achieved speedups of up to $389.02\times$ and $415.64\times$, respectively, highlighting the advantages of massive threading and ultra-high device memory bandwidth. Compared with existing GPU studies on sparse ma-

Table 6. Global matrix assembly performance (Mesh: 103,635 nodes, 554,394 elements).

Task	CPU assembly time (ms)	GPU assembly time (ms)	Speedup
Matrix	327.538	0.842	389.02
RHS	48.254	0.116	415.64

trix assembly, a substantially higher acceleration is achieved in the proposed method, demonstrating the superior performance of our kernel designs and data structures.

Regarding the parallel assembly, it is worth noting that multiple elements inevitably share nodes in unstructured meshes, which can theoretically introduce memory contention when using atomic operations. To address this, we evaluated alternative algorithms that avoid atomic operations. While this contention-free approach further reduced the sparse matrix update time (e.g., from approximately 5 ms to 3 ms in specific test stages), it required an explicit mesh preprocessing step. This preprocessing consumed several seconds, introducing an overhead that far outweighed the modest kernel-level savings. Therefore, the direct atomic operation algorithm was selected for its simplicity and overall efficiency. This design choice is further supported by recent evaluations of unstructured-grid CFD applications on GPUs, which demonstrate that NVIDIA V100 and A100 architectures inherently deliver exceptional performance on kernels dominated by double-precision atomic updates, often outperforming complex register-based aggregation or data restructuring methods (Stone et al., 2021).

5.3.3 Hardware Utilization and Bottleneck Analysis

To further comprehend the GPU execution efficiency, we analyzed the hardware utilization of our core kernels using NVIDIA Nsight Compute. The theoretical FP64 peak performance of the A100 GPU is 9.7 TFLOP/s. However, in unstructured FEM frameworks, performance is predominantly bounded by memory bandwidth rather than compute capabilities.

For the element-wise integration kernel `dshape_tensor_dshape`, the Compute (SM) Throughput reached approximately 8.50% of the theoretical peak. While this kernel involves intensive small-matrix multiplications, its arithmetic intensity is inherently constrained by the necessity to fetch scattered nodal variables and coordinates from the global memory for each element. This is strongly evidenced by its Memory Throughput, which achieved an outstanding 90.44% of the device’s peak memory bandwidth, indicating that the kernel is optimally saturating the hardware’s memory subsystem.

Similarly, the global matrix assembly `assemble_csr_matrix` and RHS construction `rhs_addto_kernel` kernels are archetypal memory-bound operations. Their Compute Throughputs are limited to 7.93% and 18.29%, respectively, primarily because they execute sparse floating-point operations (such as atomic additions) amidst heavy indirect memory accesses. Instead, their execution efficiency is profoundly reflected in their Memory Throughputs, which achieved 79.48% and 82.98%, respectively. Given the highly irregular access patterns dictated by unstructured mesh connectivity (e.g., indirect addressing via `row_pointers` and `ndglno`), maintaining such high memory bandwidth utilization demonstrates that our data layout and atomic-based assembly strategies effectively saturate the hardware limits for these memory-intensive workloads.

Table 7. Pinned vs non-pinned memory transfer performance

Data	Size(Byte)	Non-pinned time(ms)	Non-pinned Bandwidth(GB/s)	Pinned time(ms)	Pinned BW(GB/s)
ndgln array	8870304	0.686	12.94	0.348	25.47
scalar field array	829080	0.086	9.67	0.047	17.60
vector field array	2487240	0.204	12.18	0.105	23.60
tensor field array	7461720	0.586	12.73	0.332	22.44

5.3.4 CPU-GPU Data Transfer and Overlap with Computation

In hotspot acceleration, CPU-GPU data transfer is another critical performance factor. In our system, CPU and GPU are connected via PCIe 4.0 $\times 16$, with a bidirectional bandwidth of 32 GB/s. When the CPU has pinned memory, transfer bandwidth utilization is significantly higher than that with pageable memory. During Fluidity-Atmosphere computation, the type of data transferred frequently between the CPU and GPU includes scalar, vector, and tensor field arrays, the element-node connectivity array (**ndgln**), as well as the fully assembled global sparse matrix arrays returning from the GPU device memory to the host for PETSc solvers. Table 7 compares pinned and non-pinned transfer performance, displaying an evidently superior bandwidth with pinned memory.

In addition, CUDA provides stream concurrency and asynchronous APIs such as `cudaMemcpyAsync`, facilitating simultaneous data transfer and kernel execution. The `transform_to_physical` kernels can run concurrently with transfers of subsequent physical field data. The transfer-computation overlap effectively mitigates communication bottlenecks and further leverages HPC system performance.

5.3.5 Performance of functional modules

After embedding GPU-accelerated element-wise computation and global assembly strategies into Fluidity-Atmosphere, we evaluated the overall performance at three levels: module, single timestep, and multi-process parallelism. In the restarting Mountainwave3D simulation, the pressure diffusion matrix was calculated only once, and the others were the average times of simulation timesteps which ranges from 3000 to 5000 s. Table 8 summarizes performance across major functional modules.

As the global sparse linear solver was not GPU-parallelized, the overall acceleration for the single-process execution remained limited. However, Fluidity-Atmosphere already supports MPI, enabling MPI+GPU hybrid execution for maximal performance. Table 9 compares the one timestep execution performance using different numbers of CPU cores (MPI processes) and GPU-enabled configurations.

The results show that GPU acceleration achieved a **2.44 $\times$ speedup** for single-process execution. With 4 MPI processes (each handling approximately 25,000 nodes and 140,000 elements), the hybrid MPI+GPU version achieved **8.57 $\times$ speedup**

Table 8. GPU acceleration performance of major modules.

Functional module	CPU time (ms)	GPU time (ms)	Speedup
the advection-diffusion matrix (Temperature)	2563.158	11.566	221.60
the advection-diffusion matrix (WaterVapor)	2542.717	11.251	226.00
the advection-diffusion matrix (CloudWater)	2566.037	11.296	227.17
the advection-diffusion matrix (RainWater)	1877.409	9.481	198.02
the momentum matrix	12641.197	46.363	272.66
the divergence matrix	3657.007	21.497	170.12
the pressure diffusion matrix	3505.105	28.507	122.96
the projection matrix	1126.354	10.944	102.92

Table 9. The performance of GPU accelerated Fluidity-Atmosphere.

Version	Time (s)	Speedup (vs 1CPU)
1 CPU core (Baseline)	41.838	1.00
1 CPU core + 1 GPU	17.147	2.44
4 CPU cores	13.376	3.13
4 CPU cores +1 GPU	4.880	8.57

compared with single CPU core execution. Table 9 presents the average times of simulation timesteps which ranges from 3000
490 to 5000 s.

Notably, the 4 CPU cores + 1 GPU configuration achieves an $8.57\times$ speedup, which is more than triple the speedup of the 1
CPU core + 1 GPU configuration ($2.44\times$). This synergistic effect is attributed to two factors: first, the domain decomposition
in multi-process execution allows the remaining CPU-bound linear solvers (PETSc) to benefit from improved cache locality on
smaller sub-domains; second, multiple MPI processes can concurrently issue kernels to the GPU via CUDA streams, leading
495 to higher hardware occupancy on the NVIDIA A100. As detailed in Table 10, this configuration effectively addresses the
"bottleneck shift" dictated by Amdahl's Law. In the 1 CPU core + 1 GPU setup, GPU matrix assembly is compressed to less
than 1% of the total time, leaving the PETSc solver and other unaccelerated modules as the dominant bottlenecks. However,
by leveraging 4 MPI processes, the absolute execution time of the PETSc solver is dramatically reduced from 5.60 s to 2.19 s,
and the time for other unaccelerated modules drops from 11.43 s to 2.62 s. This multi-core mitigation of the remaining CPU
500 bottlenecks, combined with sustained GPU efficiency, strictly validates the $8.57\times$ overall performance gain.

Table 10. Comparison of computational performance and runtime distribution across different configurations (Unit: Time in seconds, Percentage of total timestep).

Functional module	Baseline (1 CPU core)		1 CPU core + 1 GPU		4 CPU cores + 1 GPU	
	Time (s)	Percentage	Time (s)	Percentage	Time (s)	Percentage
Temperature Matrix	2.467	5.95%	0.0226	0.13%	0.0120	0.25%
WaterVapor Matrix	2.471	5.96%	0.0137	0.08%	0.0041	0.08%
CloudWater Matrix	2.494	6.02%	0.0139	0.08%	0.0038	0.08%
RainWater Matrix	1.799	4.34%	0.0104	0.06%	0.0029	0.06%
Momentum Matrix	12.093	29.19%	0.0313	0.18%	0.0384	0.79%
Pressure Correction Matrix	4.760	11.49%	0.0165	0.10%	0.0092	0.19%
Total PETSc (Setup + Solver)	5.212	12.58%	5.6045	32.68%	2.1921	44.92%
Other unaccelerated modules	10.136	24.46%	11.4341	66.68%	2.6175	53.63%
Total timestep duration	41.433	100.00%	17.147	100.00%	4.880	100.00%

6 Conclusions

This study focuses on the unstructured-mesh finite-element atmospheric model named Fluidity-Atmosphere, thereby addressing the computational bottlenecks in element-wise computations and matrix assembly. Leveraging the architectural features of the NVIDIA A100 GPU and the CUDA programming model, we designed and implemented GPU-oriented parallel optimization strategies. The proposed high-performance GPU kernels substantially enhanced the computational efficiency while solving the momentum equations, advection equations, and stabilization term construction, thereby accelerating critical processes in atmospheric numerical simulations. For element-wise computations, template-based kernel design and deep optimizations achieved speedups ranging from tens to several hundred times. With asynchronous data transfer and MPI parallelization, an overall acceleration that is 8.57 times greater than that of the CPU version was achieved using four processes while ensuring correctness and stability in multi-time step simulations. The proposed parallelization methods substantially enhance the performance of atmospheric models on heterogeneous systems, remarkably extending the support for high-performance implementations of unstructured-mesh models. Future work will focus on further overcoming performance bottlenecks in the overall simulation, particularly by migrating the sparse linear solver (which is still CPU-dependent) to GPUs, thereby facilitating end-to-end acceleration.

Code and data availability. The original Fluidity model is distributed free of charge under the GNU Lesser General Public License (LGPL). The source code is publicly available from its official repository at <https://fluidityproject.github.io/get-fluidity.html> (AMCG, 2014). The version used in this study corresponds to Fluidity release 2025.12.

The GPU-accelerated extension developed in this work is permanently archived at Zenodo. The version used to generate all results presented in this paper corresponds to release v2 and is available at: <https://doi.org/10.5281/zenodo.18799735> (Fu, 2026). All simulation data produced in this study are publicly available at <https://doi.org/10.5281/zenodo.17824051> (Li et al., 2025).

The Zenodo archive includes: (1) the complete GPU-modified source files, (2) CUDA kernels and GPU interface implementation, (3) compilation scripts, (4) a detailed README file providing step-by-step instructions for reproducing the numerical experiments and figures reported in this paper.

Appendix A: Code

```
525 template<int T_M>
    __device__ double dot_product_t(double*__restrict__ a,
    double* __restrict__ b){
        double r=0.0;
        #pragma unroll
    530     for(int i=0;i<T_M;i++){
            r += a[i]*b[i];
        }
        return r;
    }

535 // T_M, T_N, T_K: Matrix dimensions for generic small matrix multiplication
template<int T_M, int T_N, int T_K>
__device__ void __forceinline__ matmul_t(const double* a, const double* b, double* C){
    for (int mi = 0; mi < T_M; ++mi) {
    540         for (int ki = 0; ki < T_K; ++ki) {
            #pragma unroll
            for (int ni = 0; ni < T_N; ++ni) {
                // Column-major layout stride; efficient due to thread-local
                // register/L1 caching
    545                 C[ni*T_M + mi] += a[ki * T_M + mi] * b[T_K * ni + ki];
            }
        }
    }
}

550 // T_dim: Spatial dimension (e.g., 3 for 3D)
```



```

// T_loc1, T_loc2: Number of local nodes per element
// T_ngi: Number of Gauss integration points
template<int T_dim, int T_loc1, int T_loc2, int T_ngi>
555 __device__ void dshape_tensor_dshape_unroll_t(const double *dshape1,
const double *tensor, const double *dshape2, const double *detwei, double *r){
    double tmp[T_dim]={0.0}, mat_mul[T_dim]={0.0}, dotproduct=0.0;
    for(int gi=0; gi< T_ngi; ++gi){
        double gidetwei=detwei[gi];
560     for(int iloc=0; iloc<T_loc1; iloc++){
        #pragma unroll
        for(int n=0; n<T_dim; ++n){
            tmp[n]=dshape1[n*T_ngi*T_loc1+gi*T_loc1+iloc];
            mat_mul[n] = 0.0;
565     }
        // <1,3,3> literals explicitly used to unroll 3D spatial
        // vector-tensor interaction
        matmul_t<1,3,3>(tmp, &tensor[gi*T_dim*T_dim], mat_mul);
        for(int jloc=0; jloc<T_loc2; jloc++){
570     #pragma unroll
        for(int n=0; n<T_dim; ++n){
            tmp[n]=dshape2[n*T_ngi*T_loc2+gi*T_loc2+jloc];
        }
        dotproduct = dot_product_t<T_dim>(mat_mul, tmp);
575     r[jloc*T_loc1+iloc]=r[jloc*T_loc1+iloc]+dotproduct*gidetwei;
    }
    }
}

```

580 Appendix B: Code

```

template<int T_dim, int T_loc, int T_ngi>
__global__ void assemble_csr_matrix(double *result, double* values,
int *row_pointers, int *col_index, int *ndgln, const int ele_num){
    int element_index = blockIdx.x*blockDim.x + threadIdx.x;

```

```

585     if(element_index>=ele_num) return;
        int loc1=T_loc;
        int loc2=T_loc;
        int mpos=0,base=0,upper_j=0,upper_pos=0,lower_j=0,lower_pos=0,
            this_pos=0,this_j=0;
590     int inode[T_loc]={0};
        int j=0;
        double *element_result=result + element_index *T_loc * T_loc;
        int *row=nullptr;
        int iloc=0, jloc=0, n=0;
595     //ndgln stores the connectivity for neighbouring elements.
        for(int i = 0; i < T_loc; ++i){
            inode[i]= ndgln[element_index * loc1 + i];
        }
        for(iloc=0;iloc<loc1;iloc++){
600             //the index in Fluidity-Atmosphere Fortran code is from 1.
            base = row_pointers[inode[iloc]-1]-1;
            n = row_pointers[inode[iloc]] -1 - base;
            for(jloc=0;jloc<loc2;jloc++){
                if(element_result[jloc*loc1+iloc]==0)continue;
605                 row = col_index+base;
                upper_pos=n-1;
                upper_j=row[n-1]-1;
                lower_pos=0;
                lower_j=row[0]-1;
610                 mpos=-2;
                j = inode[jloc]-1;
                if (upper_j<j){
                    mpos=-1;
                }else if (upper_j==j){
615                     mpos=upper_pos+base;
                }else if (lower_j>j){
                    mpos=-1;
                }else if (lower_j==j){
                    mpos=lower_pos+base;

```

```

620         }

        while(((upper_pos-lower_pos)>1)&&(mpos==-2)){
            this_pos=(upper_pos+lower_pos)/2;
            this_j=row[this_pos]-1;
625         if(this_j == (inode[jloc]-1)){
            mpos=this_pos+base;
        }
        else if(this_j > (inode[jloc]-1)){
            upper_j=this_j;
630         upper_pos=this_pos;
        }
        else{
            lower_j=this_j;
            lower_pos=this_pos;
635        }
    }
    if(mpos<0){
    }else{
        atomicAdd(values+mpos,element_result[jloc*loc1+iloc]);
640    }
}
}
}

```

Appendix C: Code

```

645 template<int T_loc>
__global__ void rhs_addto_kernel(GpuScalarField *field, const int ele_num,
double *gpu_ele_val){
    int element_index = blockDim.x * blockIdx.x + threadIdx.x;
    if(element_index >=ele_num) return;
650 //element_index will be used to find the node indexes starting from 1.
    element_index = element_index + 1;
    int nodes[T_loc];

```

```

//get the node indexes of current element.
gpu_ele_nodes_scalar(*field, element_index,nodes);
655 double *ptr = gpu_ele_val + (element_index -1)*T_loc;
    #pragma unroll
    for(int i = 0; i < T_loc;++i){
        atomicAdd(&(*field).val[(nodes[i] -1)],*(ptr+i));
    }
660 }

```

Author contributions. **Conceptualization & Methodology:** Li L., Li J. and Li H. provided technical guidance and supervision.

Software & Investigation: Li L., Zheng X. and Fu X. developed the software and performed the testing.

Writing - Original Draft: Fu X. prepared the manuscript.

Writing - Review & Editing: Li L. and Li J. revised the manuscript. All co-authors reviewed and approved the final manuscript.

665 *Competing interests.* The authors declare there are no conflicts of interest for this manuscript.

Acknowledgements. This work is jointly supported by the National Key R&D Program of China (Grant No. 2023YFC3705701), the State Key Laboratory of Atmospheric Environment and Extreme Meteorology (2024ZD04). Dr. J. Li thanks for the technical support of the National Large Scientific and Technological Infrastructure "Earth System Numerical Simulation Facility" (Grant 2025-EL-PT-000886, <https://cstr.cn/31134.02.EL>).

- AMCG: Fluidity manual, <https://fluidityproject.github.io/get-fluidity.html>, 2014.
- AMD: CUDA to HIP API Function Comparison, https://rocm.docs.amd.com/projects/HIP/en/latest/reference/api_syntax.html, 2025.
- Bogenschutz, P. A., Clevenger, T. C., Bradley, A. M., Caldwell, P. M., Beydoun, H., Mahfouz, N., Keen, N. D., Guba, O., Bertagna, L., Foucar, J., Zhang, J., and Donahue, A. S.: High Performance, High Fidelity: A GPU-Accelerated Doubly-Periodic Configuration of the Simple Cloud-Resolving E3SM Atmosphere Model Version 1 (DP-SCREAMv1), *Journal of Advances in Modeling Earth Systems*, 17, e2025MS005127, <https://doi.org/10.1029/2025MS005127>, 2025.
- Böhm, F., Bauer, D., Kohl, N., Alappat, C. L., Thönnies, D., Mohr, M., Köstler, H., and Rüde, U.: Code Generation and Performance Engineering for Matrix-Free Finite Element Methods on Hybrid Tetrahedral Grids, *SIAM Journal on Scientific Computing*, 47, B131–B159, <https://doi.org/10.1137/24M1653756>, 2025.
- Conde, D. A. S., Ferreira, R. M. L., Canelas, R., Ricardo, A. M., and Mendes, L.: A Distributed-Heterogeneous Design for Explicit Hyperbolic Solvers. Application to Tsunami Urban Run-Up Modelling, *Journal of Advances in Modeling Earth Systems*, 17, e2024MS004602, <https://doi.org/10.1029/2024MS004602>, 2025.
- Czarnul, P., Proficz, J., and Drypczewski, K.: Survey of Methodologies, Approaches, and Challenges in Parallel Programming Using High-Performance Computing Systems, *Scientific Programming*, 2020, 4176794, <https://doi.org/10.1155/2020/4176794>, 2020.
- Davies, D. R., Wilson, C. R., Kramer, S. C., Piggott, M. D., Le Voci, G., and Collins, G. S.: Fluidity: A New Adaptive, Unstructured Mesh Geodynamics Model, in: *EGU General Assembly Conference Abstracts*, EGU General Assembly Conference Abstracts, p. 5364, 2010.
- Farrell, P. E., Piggott, M. D., Pain, C. C., Gorman, G. J., and Wilson, C. R.: Conservative interpolation between unstructured meshes via supermesh construction, *Computer Methods in Applied Mechanics and Engineering*, 198, 2632–2642, <https://doi.org/10.1016/j.cma.2009.03.004>, 2009.
- Fischer, P., Min, M., Rathnayake, T., Dutta, S., Kolev, T., Dobrev, V., Camier, J.-S., Kronbichler, M., Warburton, T., Swirydowicz, K., and Brown, J.: Scalability of High-Performance PDE Solvers, *International Journal of High Performance Computing Applications*, 34, 562–586, <https://doi.org/10.1177/1094342020915762>, 2020.
- Fu, X.: GPU-Accelerated Implementation of the Fluidity-Atmosphere Dynamical Core, <https://doi.org/10.5281/zenodo.18799735>, 2026.
- Georgescu, S., Chow, P., and Okuda, H.: GPU Acceleration for FEM-Based Structural Analysis, *Archives of Computational Methods in Engineering*, 20, 111–121, <https://doi.org/10.1007/s11831-013-9082-8>, 2013.
- Goddeke, D., Buijssen, S., Wobker, H., and Turek, S.: GPU acceleration of an unmodified parallel finite element Navier-Stokes solver, in: *Proceedings of the 2009 International Conference on High Performance Computing and Simulation, HPCS 2009*, pp. 12 – 21, <https://doi.org/10.1109/HPCSIM.2009.5191718>, 2009.
- Jendersie, R., Lessig, C., and Richter, T.: A GPU parallelization of the neXtSIM-DG dynamical core (v0.3.1), *Geosci. Model Dev.*, 18, 3017–3040, <https://doi.org/10.5194/gmd-18-3017-2025>, 2025.
- Kiran, U., Gautam, S. S., and Sharma, D.: GPU-based matrix-free finite element solver exploiting symmetry of elemental matrices, *Computing*, 102, 1941–1965, <https://doi.org/10.1007/s00607-020-00827-4>, 2020.
- Kronbichler, M. and Kormann, K.: A generic interface for parallel cell-based finite element operator application, *Computers Fluids*, 63, 135–147, <https://doi.org/10.1016/j.compfluid.2012.04.012>, 2012.
- Kronbichler, M. and Kormann, K.: Fast Matrix-Free Evaluation of Discontinuous Galerkin Finite Element Operators, *ACM Transactions on Mathematical Software*, 45, 1–40, <https://doi.org/10.1145/3325864>, 2019.

- Li, J., Zheng, J., Zhu, J., Fang, F., Pain, C. C., Steppeler, J., Navon, I. M., and Xiao, H.: Performance of Adaptive Unstructured Mesh Modelling in Idealized Advection Cases over Steep Terrains, *Atmosphere*, 9, <https://doi.org/10.3390/atmos9110444>, 2018.
- Li, J., Fang, F., Steppeler, J., Zhu, J., Cheng, Y., and Wu, X.: Demonstration of a three-dimensional dynamically adaptive atmospheric dynamic framework for the simulation of mountain waves, *Meteorology and Atmospheric Physics*, 133, 1627–1645, <https://doi.org/10.1007/s00703-021-00828-8>, 2021.
- Li, L., Fu, X., Zheng, X., Li, H., and Li, J.: Atmospheric Mountain Wave Simulation Dataset (Fluidity-Atmosphere), <https://doi.org/10.5281/zenodo.17824052>, 2025.
- Macioł, P., Płaszewski, P., and Banaś, K.: 3D finite element numerical integration on GPUs, *Procedia Computer Science*, 1, 1093–1100, <https://doi.org/10.1016/j.procs.2010.04.121>, 2010.
- Michalakes, J.: HPC for Weather Forecasting, in: *Parallel Algorithms in Computational Science and Engineering*, edited by Grama, A. and Sameh, A. H., pp. 297–323, Springer International Publishing, Cham, ISBN 978-3-030-43736-7, https://doi.org/10.1007/978-3-030-43736-7_10, 2020.
- Mills, R. T., Adams, M. F., Balay, S., Brown, J., Dener, A., Knepley, M., Kruger, S. E., Morgan, H., Munson, T., Rupp, K., Smith, B. F., Zampini, S., Zhang, H., and Zhang, J.: Toward performance-portable PETSc for GPU-based exascale systems, *Parallel Computing*, 108, 102 831, <https://doi.org/10.1016/j.parco.2021.102831>, 2021.
- Müller, E. H., Scheichl, R., and Vainikko, E.: Petascale solvers for anisotropic PDEs in atmospheric modelling on GPU clusters, *Parallel Computing*, 50, 53–69, <https://doi.org/10.1016/j.parco.2015.10.007>, 2015.
- NVIDIA: CUDA C++ Programming Guide, <https://docs.nvidia.com/cuda/cuda-c-programming-guide/index.html>, 2025a.
- NVIDIA: Sparse Matrix Formats, https://docs.nvidia.com/nvpl/latest/sparse/storage_format/sparse_matrix.html, 2025b.
- Petersen, M. R., Asay-Davis, X. S., Barthel, A. M., Begeman, C. B., Bishnu, S., Brus, S. R., Jones, P. W., Kang, H.-G., Kim, Y., Mametjanov, A., O'Neill, B., Ringel, K. K., Smith, K. M., Sreepathi, S., Van Roekel, L. P., and Waruszewski, M.: The Ocean Model for E3SM Global Applications: Omega Version 0.1.0. A New High-Performance Computing Code for Exascale Architectures, *EGUsphere*, 2025, 1–37, <https://doi.org/10.5194/egusphere-2025-3500>, 2025.
- Ratnakar, S. K., Sanfui, S., and Sharma, D.: Graphics Processing Unit-Based Element-by-Element Strategies for Accelerating Topology Optimization of Three-Dimensional Continuum Structures Using Unstructured All-Hexahedral Mesh, *Journal of Computing and Information Science in Engineering*, 22, 021 013, <https://doi.org/10.1115/1.4052892>, 2021.
- Rudi, J., Malossi, A. C. I., Isaac, T., Stadler, G., Gurnis, M., Staar, P. W. J., Ineichen, Y., Bekas, C., Curioni, A., and Ghattas, O.: An extreme-scale implicit solver for complex PDEs: highly heterogeneous flow in earth’s mantle, in: *SC ’15: Proceedings of the International Conference for High Performance Computing, Networking, Storage and Analysis*, pp. 1–12, <https://doi.org/10.1145/2807591.2807675>, 2015.
- Sanfui, S. and Sharma, D.: A three-stage graphics processing unit-based finite element analyses matrix generation strategy for unstructured meshes, *International Journal for Numerical Methods in Engineering*, 121, 3824–3848, <https://doi.org/10.1002/nme.6383>, 2020.
- Sanfui, S. and Sharma, D.: Symbolic and Numeric Kernel Division for Graphics Processing Unit-Based Finite Element Analysis Assembly of Regular Meshes With Modified Sparse Storage Formats, *Journal of Computing and Information Science in Engineering*, 22, <https://doi.org/10.1115/1.4051123>, 2021.
- Simek, V., Dvorak, R., Zboril, F., and Kunovsky, J.: Towards Accelerated Computation of Atmospheric Equations Using CUDA, in: *2009 11th International Conference on Computer Modelling and Simulation*, pp. 449–454, <https://doi.org/10.1109/UKSIM.2009.25>, 2009.

- Stone, C. P., Walden, A., Zubair, M., and Nielsen, E. J.: Accelerating unstructured-grid CFD algorithms on NVIDIA and
745 AMD GPUs, in: 2021 IEEE/ACM 11th Workshop on Irregular Applications: Architectures and Algorithms (IA3), pp. 19–26,
<https://doi.org/10.1109/IA354616.2021.00010>, 2021.
- Sulyok, A. A., Balogh, G. D., Regul, I. Z., and Mudalige, G. R.: Locality optimized unstructured mesh algorithms on GPUs, *Journal of
Parallel and Distributed Computing*, 134, 50–64, <https://doi.org/10.1016/j.jpdc.2019.07.011>, 2019.
- Takle, E. and Russell, R.: Applications of the finite element method to modeling the atmospheric boundary layer, *Computers & Mathematics
750 with Applications*, 16, 57–68, [https://doi.org/10.1016/0898-1221\(88\)90024-7](https://doi.org/10.1016/0898-1221(88)90024-7), 1988.
- Thibault, J. and Senocak, I.: CUDA Implementation of a Navier-Stokes Solver on Multi-GPU Desktop Platforms for Incom-
pressible Flows, in: 47th AIAA Aerospace Sciences Meeting including The New Horizons Forum and Aerospace Exposition,
<https://doi.org/10.2514/6.2009-758>, 2012.