

Computational Physics

A distributed-memory tridiagonal solver based on a specialised data structure optimised for CPU and GPU architectures

Semih Akkurt^{a, *}, Sébastien Lemaire^b, Paul Bartholomew^b, Sylvain Laizet^a^a Department of Aeronautics, Imperial College London, SW7 2AZ, United Kingdom^b EPCC, University of Edinburgh, EH8 9BT, United Kingdom

ARTICLE INFO

The review of this paper was arranged by
Prof. David W. Walker

Keywords:

Tridiagonal matrix algorithm
Distributed tridiagonal matrix solver
Compact finite difference schemes

ABSTRACT

Various numerical methods used for solving partial differential equations (PDE) result in tridiagonal systems. Solving tridiagonal systems on distributed-memory environments is not straightforward, and often requires significant amount of communication. In this article, we present a novel distributed-memory tridiagonal solver algorithm, DistD2-TDS, based on a specialised data structure. DistD2-TDS algorithm takes advantage of the diagonal dominance in tridiagonal systems to reduce the communications in distributed-memory environments based on an established strategy. The underlying data structure plays a crucial role for the performance of the algorithm on individual ranks. First, the data structure improves data localities and makes it possible to minimise data movements via cache blocking and kernel fusion strategies. Second, data continuity enables a contiguous data access pattern and results in efficient utilisation of the available memory bandwidth. Finally, the data layout supports vectorisation on CPUs and thread level parallelisation on GPUs for improved performance. In order to demonstrate the robustness of the algorithm, we implemented and benchmarked the algorithm on CPUs and GPUs. We investigated the single rank performance and compared against existing algorithms. Furthermore, we analysed the strong scaling of the implementation up to 384 NVIDIA H100 GPUs and up to 8192 AMD EPYC 7742 CPUs. Finally, we demonstrated a practical use case of the algorithm by using compact finite difference schemes to solve a 3D non-linear PDE. The results demonstrate that DistD2 algorithm can sustain around 66% of the theoretical peak bandwidth at scale on CPU and GPU based supercomputers.

1. Introduction

In numerical linear algebra and applied mathematics, a tridiagonal system is a type of matrix that has non-zero entries only on the main diagonal, the diagonal directly above the main, and the diagonal directly below the main. Certain numerical discretisations that are used for solving partial differential equations (PDE) on structured grids in various fields often result in batches of tridiagonal systems that are independent from each other. Typically, a batch of tridiagonal systems represents a one-directional coupling in a particular spatial direction, and there is at least one batch of tridiagonal systems per spatial direction. For example, discretising the Navier-Stokes equations using high-order compact finite difference schemes [1] result in a batch of tridiagonal systems per term in the equation. In addition to compact finite difference schemes, batches of tridiagonal systems also arise when using alternating direction implicit (ADI) solvers [2].

The performance of tridiagonal matrix solvers have been investigated both for single rank and multiple rank strategies including CPU and GPU architectures [3–6]. A typical sub-optimal performance characteristic found in many of these strategies is the reduced performance in solving the x -directional tridiagonal systems compared to y - and z -directional systems. This is typically due to the Cartesian data structure that is used in these strategies where data points along the x -direction are stored next to each other in memory, however there are jumps in memory between data points that are neighbours in the physical space along the y -direction or the z -direction. Storing the physically neighbouring entries next to each other in memory eliminates the possibility of certain optimisations due to the sequential operations found in most tridiagonal matrix algorithms. In particular, on CPU architectures vectorisation is harder to achieve due to the dependencies between neighbouring elements, and on GPU architectures it becomes harder to utilise all the threads in a thread block efficiently. Laszlo et al. [3] pro-

* Corresponding author.

E-mail address: s.akkurt18@imperial.ac.uk (S. Akkurt).<https://doi.org/10.1016/j.cpc.2025.109747>

Received 20 November 2024; Received in revised form 13 May 2025; Accepted 2 July 2025

Available online 10 July 2025

0010-4655/© 2025 The Author(s). Published by Elsevier B.V. This is an open access article under the CC BY license (<http://creativecommons.org/licenses/by/4.0/>).

posed a methodology to solve this problem based on carrying out local transposes in CPU cache or GPU shared memory. Although this strategy performs significantly better than a naive approach where the discrepancy in performance between x - and y -directional operations can be up to 8x, there is still a meaningful discrepancy of up to 2x especially when the domain size is 512 and above on GPUs.

In this paper, we propose a specialised data structure and develop a novel algorithm for improving performance of tridiagonal matrix solvers both on CPU and GPU architectures. Most importantly, the proposed data structure enables a linear data access pattern, makes it easy to vectorise the operations, and improves data locality which in turn results in a better cache utilisation regardless of the spatial direction of the tridiagonal systems. Moreover, we also demonstrate the suitability of this new data structure for kernel fusion strategies, where the mathematical operations are combined at source code level for reducing data movement requirements and improving performance. This strategy is essential for performance for bandwidth bound algorithms especially on GPU architectures due to their relatively smaller per-core cache memory compared to CPUs.

Furthermore, as the domain where these batches of tridiagonal systems are solved are often quite large, ranging from a billion degree-of-freedom (DOF) to up to a few trillion DOF, employing efficient parallel strategies that can leverage large scale supercomputers is very important. Therefore, we extend our novel strategy to include a distributed-memory tridiagonal matrix algorithm named DistD2-TDS. For simplicity, we focus on diagonally dominant tridiagonal systems for the distributed-memory strategy, taking advantage of the simplified communication pattern of such systems, as discussed in [7]. In order to provide real world use cases of diagonally dominant tridiagonal systems, we use higher-order compact finite difference schemes, also known as spatially implicit schemes, to discretise PDEs of interest and apply the proposed algorithm. It should be noted that the proposed strategy based on the specialised data structure can be extended for solving generic tridiagonal systems as well. Considering generic tridiagonal systems would require a more elaborate approach when dealing with the contribution from other subdomains, see for instance the strategy suggested in [5]. It would require an additional step in our approach between the decoupling phase and substitution phase of the DistD2 algorithm. Further details regarding the limitations of our strategy and potential extensions are discussed in detail in Section 4.

The remainder of the paper is structured as follows. Section 2 focuses on how to discretise PDEs with tridiagonal systems. In section 3 we discuss tridiagonal matrix algorithms including serial and distributed-memory strategies. Section 4 discusses performance considerations focusing on operations typically found in tridiagonal matrix algorithms. Then we propose our novel strategy in Section 5, describing the DistD2-TDS algorithm in detail. Section 6 presents performance results and also provides a practical application of the proposed strategy on a 3D nonlinear PDE. Finally, conclusions are discussed in Section 7.

2. Numerical discretisation of 3D PDEs

In this paper, the primary focus is on solving tridiagonal systems arising from the numerical discretisation of 3D PDEs on structured grids. Tridiagonal systems are a direct result of a one-directional coupling obtained from certain types of numerical discretisations. Typically, they require multiple steps in which contributions from x -, y -, and z -directions and the associated one-directional couplings are considered separately. For example, high-order compact finite difference schemes result in a batch of tridiagonal systems for each term in a PDE, and one-directional coupling in these systems depend on the direction of the partial derivative. Therefore, compact schemes result in a number of batches of tridiagonal systems such that there is at least one batch of tridiagonal systems per spatial direction. In order to examine the challenges associated with dealing with these batches of tridiagonal systems on large-scale supercomputers, we first give a brief description of

high-order compact finite difference schemes in the following subsection and then discuss the potential problems of applying these schemes in distributed-memory environments.

2.1. Compact finite difference schemes

A compact finite difference scheme is a numerical method for approximating derivatives that achieves higher accuracy without increasing the computational stencil width. Unlike standard finite difference schemes, which often use a wider stencil for higher accuracy, compact schemes use implicit formulations that relate the derivatives at multiple grid points. This creates a system of equations that must be solved simultaneously, typically resulting in a tridiagonal system. As an example, a compact scheme for a sixth-order accurate approximation to a first-order derivative can be given as

$$au'_{i-1} + u'_i + au'_{i+1} = c_a \frac{u_{i+1} - u_{i-1}}{2h} + c_b \frac{u_{i+2} - u_{i-2}}{4h}, \quad (1)$$

where u are the field entries, u' is the derivative, h is the grid spacing, α , c_a , and c_b are the coefficients that are obtained from the Taylor series expansions to satisfy sixth-order accuracy ($\alpha = 1/3$, $c_a = 14/9$, $c_b = 1/9$). The formulation on the right hand side of the equation (RHS) is the $\mathbf{d}$ vector in Equation (2). Compact schemes can also be used to obtain higher-order derivatives, interpolations and filters [1]. High-order implicit finite difference schemes offer several advantages in computational fluid dynamics (CFD) for solving partial differential equations (PDEs). They can achieve greater accuracy for the same grid resolution compared to lower-order schemes. This is particularly beneficial for capturing fine details and complex flow structures in fluid dynamics simulations. They generally exhibit low numerical dissipation and dispersion, preserving the physical fidelity of the simulation. Although high-order implicit schemes involve more complex computations per grid point, their ability to achieve high accuracy with fewer grid points can lead to overall computational savings [1]. A good compromise in terms of cost and accuracy are sixth-order implicit schemes which are based on tridiagonal solvers. These types of schemes are used, for instance, in Xcompact3d, a framework of finite-difference solvers, designed to study fluid flow problems on supercomputers [8]. As an example, the solvers in Xcompact3d require approximately 150 tridiagonal systems to be solved at each time step when simulating a fluid flow problem. One important feature of tridiagonal systems resulting from high-order compact finite difference schemes is that they are always diagonally dominant, and the proposed algorithm is taking advantage of the diagonal dominance to minimise the communication requirements between ranks.

In summary, the application of compact schemes to solve 3D PDEs results in batches of tridiagonal systems where each batch represents a one-directional coupling in a particular spatial direction. Moreover, each batch contains many tridiagonal systems that correspond to a numerical operation such as first derivatives. Therefore, these independent tridiagonal systems within a batch have common coefficients as described in Equation (1). Thus, the resulting equations can be described in matrix form $\mathbf{A}\mathbf{u}' = \mathbf{d}$ and can be expanded as

$$\begin{bmatrix} b_1 & c_1 & & & & \\ a_2 & b_2 & c_2 & & & \\ & a_3 & b_3 & c_3 & & \\ & & \ddots & \ddots & \ddots & \\ & & & a_{n-1} & b_{n-1} & c_{n-1} \\ & & & & a_n & b_n \end{bmatrix} \begin{bmatrix} u'_1 \\ u'_2 \\ u'_3 \\ \vdots \\ u'_{n-1} \\ u'_n \end{bmatrix} = \begin{bmatrix} d_1 \\ d_2 \\ d_3 \\ \vdots \\ d_{n-1} \\ d_n \end{bmatrix}. \quad (2)$$

The coefficient matrix $\mathbf{A}$ consists of 3 bands where we refer a_i as the lower diagonal, b_i as the diagonal, and c_i as the upper diagonal coefficients, u'_i is the unknown derivative vector, d_i is the right hand side (RHS) vector corresponding to the RHS of Equation (1), and n is the size of the tridiagonal system. The batch dimension corresponds to the individual tridiagonal systems in the domain. The banded structure in

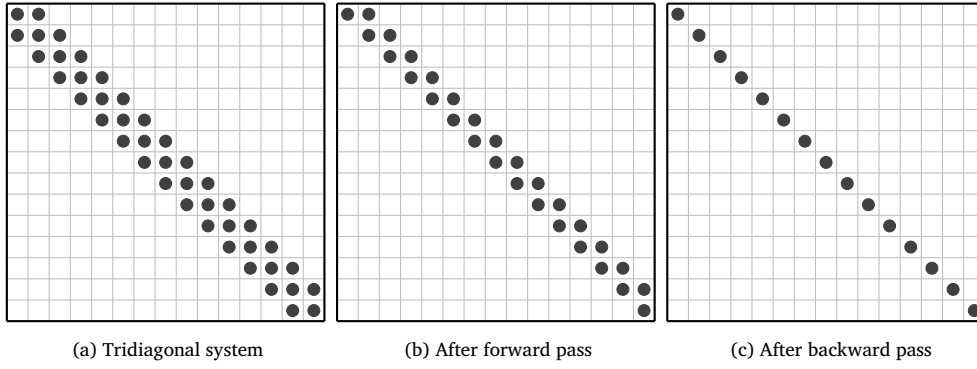


Fig. 1. Thomas algorithm for solving tridiagonal matrices. Non-zero entries in the banded matrix are represented by grey circles. Forward pass eliminates the lower diagonal, and backward pass eliminates the upper diagonal. RHS and the unknown vectors are omitted for simplicity.

tridiagonal matrices allows us to use simplified algorithms to solve these systems. They are examined in detail in Section 3.

3. Tridiagonal matrix algorithms

3.1. Serial algorithm - Thomas algorithm

The Thomas algorithm [9] is effectively a simplified Gauss elimination, consisting of forward and backward passes to eliminate the lower and upper diagonals in the coefficient matrix A as described in Algorithm 1.

Algorithm 1 Thomas algorithm.

```

1:  $u'_1 \leftarrow d_1/b_1$ 
2:  $c_1 \leftarrow c_1/b_1$ 
3: for  $i = 2 : n$  do                                ▷ Forward pass
4:    $w \leftarrow 1/(b_i - a_i c_{i-1})$ 
5:    $u'_i \leftarrow w(d_i - a_i u'_{i-1})$ 
6:    $c_i \leftarrow w c_i$ 
7: end for
8: for  $i = n - 1 : 1$  do                                ▷ Backward pass
9:    $u'_i \leftarrow u'_i - c_i u'_{i+1}$ 
10: end for

```

The changes in the non-zero structure in the coefficient matrix A as a result of forward and backward passes are shown in Fig. 1. Additionally, physical systems with periodic boundary conditions lead to a cyclic tridiagonal systems, however, a simple variation of the Thomas algorithm (periodic Thomas) based on the Sherman-Morrison formula can be used to solve these systems [10]. Periodic Thomas algorithm implementations require an additional forward pass operation in addition to the forward and backward passes in the Thomas algorithm to obtain the solution. All these forward and backward passes have a dependency on the previous iteration, thus the Thomas algorithm is a sequential algorithm. Therefore, it requires all the entries along an individual tridiagonal system to be present in a shared-memory environment such as a single rank. However, there are various studies where the Thomas algorithm is used for solving batches of tridiagonal systems in 3D domains on multiple ranks [8]. Typically, a 1D or 2D domain decomposition as shown in Fig. 2 is used in these strategies to distribute batches of tridiagonal systems across multiple ranks. Additionally, 1D and 2D decomposition strategies alternate between 2 and 3 different states with transposed layouts to have the entirety of the domain in a particular direction in a single rank so that the sequential Thomas algorithm can be applied. These transposition operations between decomposition states are also found in multi-rank fast Fourier Transform algorithms [11]. 1D and 2D decompositions and all the states the decomposition alternates in between are demonstrated in Fig. 2. This strategy necessitates MPI

all-to-all type communications to carry out transposes across ranks between different states and may cause bottlenecks especially on modern GPU clusters.

3.2. Distributed algorithms

The fundamental difference between a serial algorithm such as the Thomas algorithm and various distributed algorithms is that distributed algorithms divide individual tridiagonal systems into multiple subdomains as shown in Fig. 3. This allows distributed-memory algorithms to have up to 3D domain decomposition as shown in Fig. 2, and the decomposition state is constant and does not require frequent alterations between different states as in Thomas algorithm based strategies. Therefore, the most important advantage of distributed-memory algorithms is that all-to-all type MPI communications involving the entire domain are completely eliminated.

In this subsection, we examine two different distributed algorithms. The proposed distributed-memory algorithm in this paper for solving tridiagonal systems combines some of the strategies in the two algorithms described in this section and also introduces novel techniques.

3.2.1. Hybrid algorithms

The hybrid Thomas-PCR algorithm was developed by Lazslo et al. in 2016 [3]. The algorithm consists of 3 phases. The first phase obtains a reduced sized tridiagonal system where each subdomain contributes two unknowns each. The primary operations in the first phase of the algorithm are similar to forward and backward passes in the Thomas algorithm, however, they are carried out in the subdomains rather than the entire domain. Algorithm 2 provides the pseudo-code for the first phase of the hybrid algorithm, modified Thomas algorithm, and Fig. 4 demonstrates the corresponding changes in the non-zero structure of the coefficient matrix A as well as highlighting the contributions from each subdomain to the final reduced system.

After the first phases are carried out in each subdomain, a reduced system is obtained where the first and last entries from each subdomain are coupled, forming a smaller tridiagonal system as shown in Fig. 4. Then, this reduced sized system is solved by using a parallel cyclic reduction (PCR) based strategy [12]. Because the implementation considers a shared memory environment where cores can access data in any subdomain, the PCR based strategy to solve the reduced sized system does not require MPI communications, and the solution can be obtained efficiently. The final phase after solving the reduced sized system is trivial, and only requires substituting the first and last entries in each subdomain as shown in Algorithm 3.

Relying on a shared memory environment to solve the reduced sized system effectively limits the algorithm to a single rank, however, extension of the algorithm for the distributed-memory systems was also studied by Balogh et al. [4]. In the distributed-memory versions, the systems are divided into subdomains and these subdomains are located

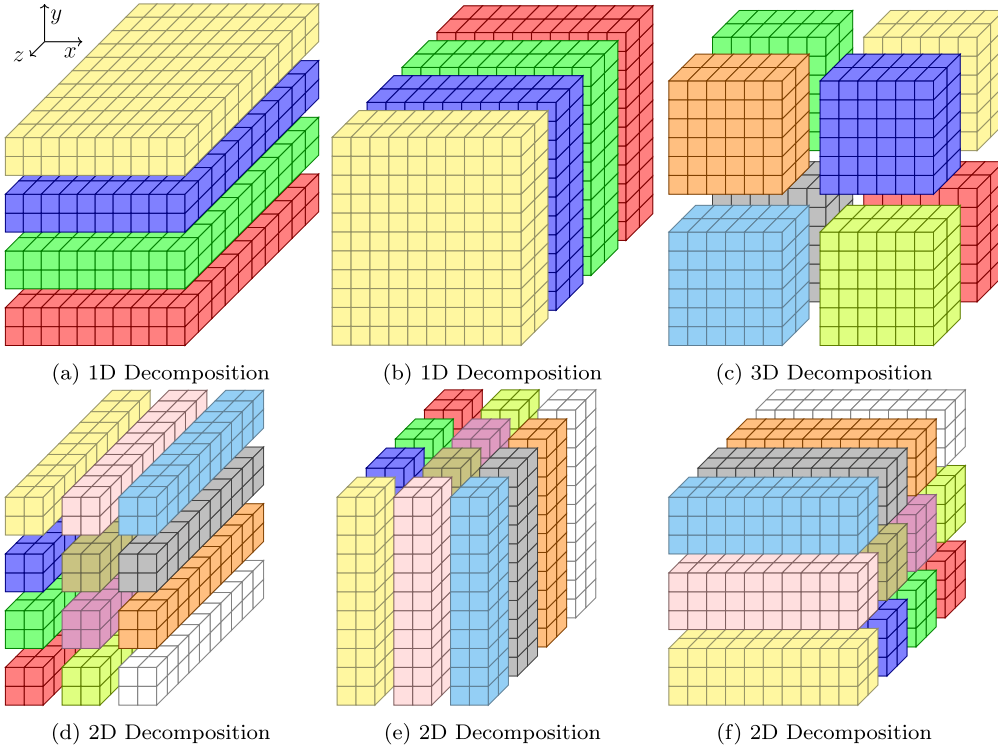


Fig. 2. 1D/2D/3D decomposition strategies with all their possible states.

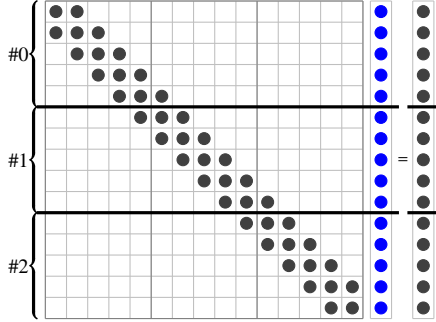


Fig. 3. Subdomains in a distributed-memory tridiagonal matrix algorithm. A tridiagonal system with 15 points are divided into 3 subdomains numbered from 0 to 2 each with 5 points. Blue colour represents the unknowns. (For interpretation of the colour(s), the reader is referred to the web version of this article.)

in different ranks. In this strategy, solution of the reduced systems can be more costly because they require multiple steps of communication depending on the number of ranks involved. A different strategy was suggested in [5], where the reduced systems are distributed across the participating ranks using all-to-all type communications which significantly reduces the amount of data sent and received. In this strategy, every rank solves a batch of reduced systems in its local memory, with no need for a PCR type multiple step communication across ranks. However the solutions of the reduced systems required to be communicated with all-to-all type communications.

The DistD2 algorithm follows a similar strategy with the first phase of modified Thomas algorithm, and obtains a similar non-zero structure. However, our approach can efficiently carry out this procedure on a distributed-memory environment, and the resulting reduced system after the first phase is further simplified down to an independent set of 2x2 systems across each subdomain boundary. Moreover, the specialised data structure enable our approach to utilise the vectorisation on CPUs

Algorithm 2 Modified Thomas algorithm - first phase.

```

1: for  $i = 1 : 2$  do
2:    $u'_i \leftarrow d_i / b_i$ 
3:    $a_i \leftarrow a_i / b_i$ 
4:    $c_i \leftarrow c_i / b_i$ 
5: end for
6: for  $i = 3 : n$  do
7:    $w \leftarrow 1 / (b_i - a_i c_{i-1})$ 
8:    $u'_i \leftarrow w(d_i - a_i u'_{i-1})$ 
9:    $a_i \leftarrow -w a_i a_{i-1}$ 
10:   $c_i \leftarrow w c_i$ 
11: end for
12: for  $i = n - 2 : 2$  do
13:   $u'_i \leftarrow u'_i - c_i u'_{i+1}$ 
14:   $a_i \leftarrow a_i - c_i a_{i+1}$ 
15:   $c_i \leftarrow -c_i c_{i+1}$ 
16: end for
17:  $w \leftarrow 1 / (1 - c_1 a_2)$ 
18:  $u'_1 \leftarrow w(u'_1 - c_1 u'_2)$ 
19:  $a_1 \leftarrow w a_1$ 
20:  $c_1 \leftarrow -w c_1 c_2$ 

```

▷ Forward pass

▷ Backward pass

Algorithm 3 Modified Thomas algorithm - substitution phase.

```

1: for  $i = 2 : n - 1$  do
2:    $u'_i \leftarrow u'_i - a_i u'_{i-1} - c_i u'_{i+1}$ 
3: end for

```

and thread level parallelism on GPUs without requiring on-the-fly local transpose operations as described in [3].

3.2.2. Parallel diagonally dominant algorithm

The parallel diagonally dominant algorithm (PDD) was first developed by Sun in 1995 [7]. As the name implies, this algorithm can only be used if the tridiagonal system is diagonally dominant. When a diagonally dominant tridiagonal system is divided into multiple subdomains, the global coupling across these regions can be broken down, and this depends on the strength of the diagonal dominance in the system as well

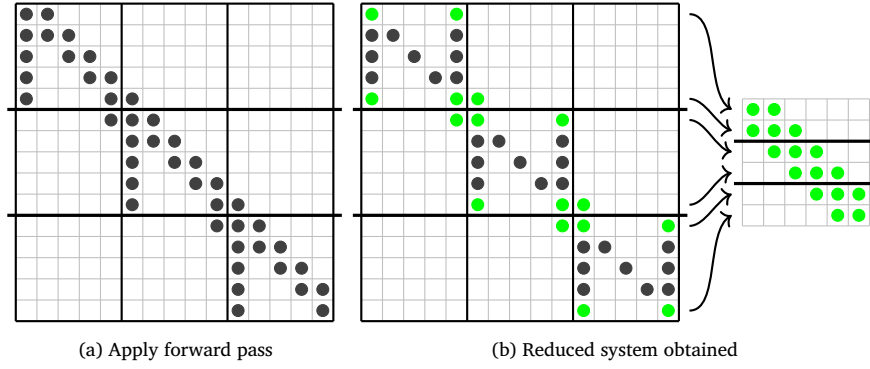


Fig. 4. First phase of the modified Thomas algorithm. Each subdomain carries out forward and backward passes locally resulting in a reduced system shown with ●.

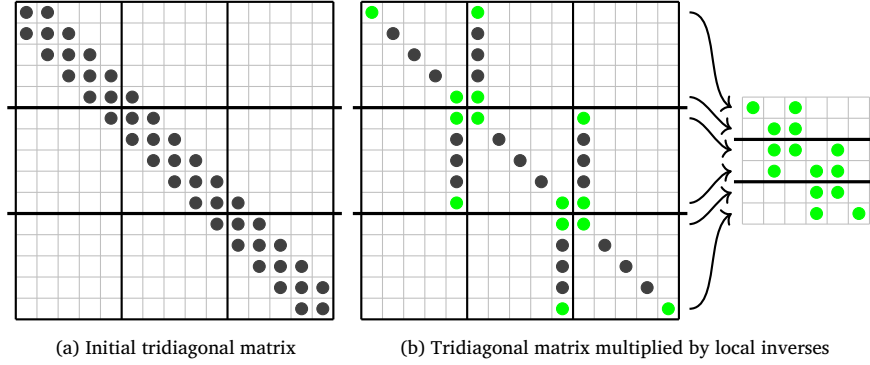


Fig. 5. Parallel diagonally dominant algorithm. Each region is multiplied by its local inverse which results in a reduced penta-diagonal system with first and last data points from each rank.

as the sizes of the subdomains as investigated in [7] in detail. The algorithm requires inverting the local region of the tridiagonal system that belongs to a rank and multiplying the system and the corresponding section of the RHS array by the local inverse matrix. The specific steps are provided in Algorithm 4.

Algorithm 4 Parallel diagonally dominant algorithm.

```

1:  $\mathbf{A}^{inv} \leftarrow \text{invert}(\mathbf{A}[1 : n, 1 : n])$ 
2: for  $i = 1 : n$  do
3:    $a_i \leftarrow \mathbf{A}^{inv}[i, n] \cdot c_n$ 
4:    $c_i \leftarrow \mathbf{A}^{inv}[i, 1] \cdot c_1$ 
5:    $u'_i \leftarrow \mathbf{A}^{inv}[i, 1 : n] \times d[1 : n]$ 
6: end for

```

As a result of a multiplication by the inverse, the boundary points in each subdomain decouple from the interior points as shown in Fig. 5. The points in the boundary from each region then form a penta-diagonal system. However, further analysis in [7] demonstrates that particular entries in the resulting penta-diagonal system are below the zero-machine value, which is $\approx 10^{-16}$ in double precision arithmetic, for diagonally dominant systems and can be disregarded without affecting the accuracy of the solution. Thus, the penta-diagonal systems are transformed into a series of 2x2 systems such that each system is across a subdomain boundary. As a result, the communication requirements to solve the reduced systems are significantly less compared to an alternative strategy to solve the reduced penta-diagonal systems. After the 2x2 systems are solved, the final results are obtained simply by substituting the known values from each neighbouring subdomain, $u'_{n,prev}$ from the previous subdomain and $u'_{1,next}$ from the next subdomain as shown in Algorithm 5.

The proposed approach in this manuscript uses diagonal dominance to decouple the subdomains as in PDD algorithm. As a result, the DistD2 algorithm can reduce communication related latencies and achieves ex-

Algorithm 5 PDD algorithm - substitution phase.

```

1: for  $i = 2 : n - 1$  do
2:    $u'_i \leftarrow u'_i - a_i u'_{n,prev} - c_i u'_{1,next}$ 
3: end for

```

cellent scaling. Furthermore, the DistD2 algorithm decouples the subdomains without requiring matrix inversion and matrix multiplication operations, mainly to avoid high FLOP requirements. Additionally, our approach reduces data movement between the processor and memory while also lowering FLOP requirements, thanks to the efficient use of advanced features in modern hardware, as explained in detail in Section 5.

4. Performance considerations

There are two main characteristics in tridiagonal matrix algorithms that are crucial for computational performance. First, the FLOP requirements of these algorithms are very low compared to data movement requirements, and therefore they are heavily bandwidth bound. Thus, achieving high performance requires efficient utilisation of the available bandwidth on a given system. Second, the algorithms described in Section 3 include forward and backward passes, where the data points in the domain are accessed and operated on in a sequential way. The forward and backward passes result in accessing a given data point twice in a short time interval. Thus, this provides an opportunity to utilise the CPU cache or GPU shared memory to store the intermediate state between forward and backward passes to save unnecessary data movements from and to the main memory.

The frequent accesses to the same memory locations in the algorithm make it possible to use cache memory and reduce the data movements from the main memory. However, the sequential nature of the arithmetic operations within the forward and backward passes can disable the vectorisation on CPUs and thread level parallelism on GPUs if not

Table 1

Specifications of state of the art CPUs in 2024 by AMD, Intel, and Nvidia.

	AMD EPYC 9754	Intel 8490H	Nvidia Grace CPU
Microarchitecture	Zen 4c	Sapphire Rapids	Neoverse V2
# of cores	128	60	72
L1d Cache	32 KiB	48 KiB	64 KiB
L2 Cache (private)	1 MiB	2 MiB	1 MiB
Peak Memory BW	460.8 GB/s	307.2 GB/s	512 GB/s

Table 2

Specifications of state of the art GPUs in 2024 by AMD, Intel, and Nvidia.

	AMD MI300X	Intel GPU Max	Nvidia GH200
Architecture	CDNA3	GPU Max	Hopper
# of SM/CU/XeC	304	128	132
# of DP cores per SM	32	32	64
Total Memory	192 GB	128 GB	96 GB
Total L1 Cache	80 KiB	512 KiB	264 KiB
Shared L1 (Scratch)	64 KiB	512 KiB	228 KiB
L2 Cache (shared)	8 × 4 MiB	288 MiB	50 MiB
Peak Memory BW	5.3TB/s	3.2TB/s	4TB/s

handled correctly. Although the characteristics at the algorithmic level are promising for exploiting the memory hierarchy and obtaining performance gains, there are a few challenges for achieving these improvements in practice. First of all, our primary focus is solving tridiagonal systems in 3D domains. Therefore, both the sizes of the domains and the data structure are very important considerations for enabling locality optimisations. For example, a typical data structure for storing a field in a 3D domain is based on a 3D array where x , y , and z indices enable accessing a location in the domain easily. However, this results in quite significant jumps in memory between some of the entries that are next to each other in physical space. Therefore, when we access these entries that are far from each other in memory to perform a numerical operation, it results in a separate read operation per entry, potentially limiting the performance.

Exploiting the memory hierarchy via locality optimisations requires the processor to have enough cache memory to store the intermediate results. Modern CPUs have up to 2 MiB private L2 cache per core, and it is large enough to be used when solving a tridiagonal system in a 3D domain. Typically, domain sizes along a direction range between 512 and 4096 entries for computational fluid dynamics applications. An individual tridiagonal system along a given line with 4096 data points in double precision is only 32 KiB in size, and most CPUs can store such systems entirely in cache. Three new generation CPUs and their specifications are provided in Table 1 to illustrate current trends. However, the current generation of GPUs has smaller cache sizes compared to CPUs as shown in Table 2. Depending on the size of cache requirement of the operations which need to be carried out, GPU caches may not be large enough due to their smaller per-core caches, leaving kernel fusion as the only option for reducing data movements on GPUs.

Finally, vectorisation on CPUs and thread level parallelisation on GPUs are also important factors to make efficient use of a given hardware. For example, sequential operations in the Thomas and Modified Thomas algorithms can prevent vectorisation and thread level parallelisation based on the data layout. If the data points that need to be operated on sequentially are stored next to each other in memory, CPUs cannot issue vector instructions such as AVX512. Moreover, threads in a GPU thread block are effectively parallelised and work efficiently only when they can operate concurrently, which does not align well with sequential operations.

To summarise, the data access pattern and exploiting the memory hierarchy are critical for achieving a high performance implementation for tridiagonal matrix solvers. First, a better data structure can lead to a more linearised data access pattern, enabling a higher percentage of peak bandwidth. Also, a well designed data structure can enable vectori-

sation on CPUs and thread level parallelisation on GPUs for improved performance.

5. Proposed strategy

In this section, we provide a detailed description of a specialised data structure along with a new distributed-memory tridiagonal solver algorithm, DistD2-TDS, to achieve high performance on CPU and GPU hardware. Furthermore, we also discuss the advantages of the data structure in terms of performance improvements both for a Thomas algorithm implementation and for our custom distributed-memory tridiagonal solver implementation based on DistD2 algorithm.

5.1. Data structure

As mentioned in Section 1 and further discussed in Section 3, tridiagonal solver algorithms often require sequential operations. However, we are interested in solving batches of tridiagonal systems in a 3D domain such that the individual tridiagonal systems are independent from each other. Although this provides a decent level of parallelism that can be taken advantage of, a naive implementation might still lead to underutilisation of the hardware resulting in low performance due to the sequential operations required by these algorithms. In particular, x -directional tridiagonal systems are the most vulnerable because of the typical Cartesian data structure where data points along the x -direction are stored next to each other in memory, while in y - and z -directions there is a jump in memory between data points that are adjacent to each other in physical domain. In case of sequential data dependencies between data points that are stored next to each other in memory, CPUs cannot use vector instructions and GPUs cannot easily take advantage of thread level parallelism. Thus, x -directional tridiagonal systems typically suffer from low performance. On the other hand, y - and z -directional tridiagonal systems can use vector instructions on CPUs and also fully utilise the available thread parallelism on GPUs. However, because the physically adjacent data points are separated by significant jumps in memory, the data access pattern in y - and z -directional tridiagonal systems is non-contiguous and this may cause a reduced performance.

In order to enable a linear memory access pattern regardless of the spatial direction of the tridiagonal systems, we subdivide the computational domain into groups of individual tridiagonal systems and tightly pack these individual systems so that their data is contiguous in memory. It is important to note that the specifics of packing depend on the spatial direction of the tridiagonal systems. However, once the specified data layout is obtained, the memory access pattern when solving the tridiagonal systems is independent of the spatial direction. A schematic is provided in Fig. 6 where the domain size is $32 \times 8 \times 4$, here a single group consists of 4 individual tridiagonal systems and thus the domain can be mapped in 8 groups as numbered in the figure. Thus, a Cartesian mesh with n_x , n_y , and n_z entries is mapped as $(SZ, n_x, n_y \cdot n_z / SZ)$ with the proposed data structure in the x -directional data layout. The optimum number of tridiagonal systems in a group, SZ , depends on the specifics of the hardware, and is 4 in Fig. 6 for the simplicity of demonstration. Typically, for a double precision simulation on CPUs, $SZ = 8$ as vector registers are typically 512 bit long and they can be used to carry out 8 double precision FLOPs per cycle, and for a double precision simulation on GPUs, $SZ = 32$ as a single streaming multiprocessor (SM) typically has 32 double precision cores in total, where each core is effectively assigned an individual tridiagonal system.

In order to enable vectorisation and thread level parallelism the proposed data structure packs the n^{th} entries of all the tridiagonal systems in a single group next to each other in memory. Therefore, the sequential operations in the algorithm as we apply the forward and backward passes can be concurrently executed for SZ systems at once per core on a CPU and per SM on a GPU. Furthermore, the substitution phase of the hybrid algorithm also benefits from the data structure in a similar

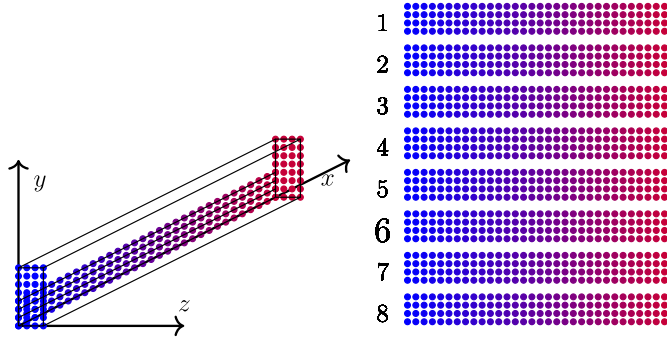


Fig. 6. The proposed data structure for an x-directional tridiagonal system. Data continuity in memory is in column-major order.

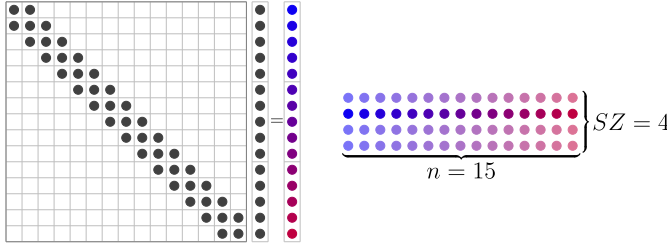


Fig. 7. A single group with $SZ = 4$ and $n = 15$. The tridiagonal system is demonstrated with a single RHS vector, and the location of the RHS vector in the specialised data structure is highlighted in colour to indicate the exact distribution of a single RHS in memory. The data layout is in column-major order, thus each subsequent data point in a single tridiagonal system is separated from each other by SZ many data points in memory.

way. Fig. 7 demonstrates how a tridiagonal system is stored in memory with the proposed data structure. One important note here is that the data layout is specific to a direction in the domain, and it requires reordering based on the direction of the tridiagonal solver. However, reordering operations account for only a small percentage of the total runtime as demonstrated later in Section 6.3.

Apart from enabling the vectorisation on CPUs and thread parallelism on GPUs, there are two more main benefits of the proposed data structure. First, the data structure packs only a small number of lines together, and the tridiagonal solver implementation makes sure that all the lines packed together are being solved concurrently as the data is accessed from main memory. Thus, the data access is linear and predictable by the hardware prefetching mechanism, with no intermittent memory accesses that would be present in y- or z-directional tridiagonal solvers in a Cartesian data structure. Second, the tight packing of lines that are concurrently being solved also makes it possible to enable a more efficient cache utilisation based on temporal locality between the forward and backward passes of the tridiagonal solver algorithms. A similar tight data packing strategy is used in [13] to take advantage of temporal localities and enable cache blocking.

5.2. DistD2-TDS - customised tridiagonal matrix solver algorithm

In this subsection, we describe our customised distributed-memory tridiagonal solver algorithm in detail, and discuss its advantages over existing methodologies. First of all, we focus on solving tridiagonal systems resulting from high-order compact finite difference schemes [1] as described in Section 2. Before describing the details of our customised strategy, we first investigate the possibility of preprocessing the coefficient arrays to avoid recalculating them on the fly. Applying a compact scheme formulation in a 3D domain along a given direction requires solving a tridiagonal system with many RHS arrays. Therefore, we start by preprocessing the tridiagonal system that defines the operation into a set of coefficient arrays such that these preprocessed arrays can be used

to avoid recalculating them on the fly for each RHS array. Algorithm 6 describes the procedure, which has to be carried out only once for each distinct operation that correspond to a tridiagonal matrix.

Algorithm 6 Preprocessing for the DistD2 algorithm.

```

1: for  $j = 1 : 2$  do
2:    $s_j^a \leftarrow s_j^a / b_j$ 
3:    $s_j^c \leftarrow s_j^c / b_j$ 
4:    $w_j \leftarrow s_j^c$ 
5:    $f_j \leftarrow 1 / b_j$ 
6: end for
7: for  $j = 3 : n$  do
8:    $w_j \leftarrow 1 / (b_j - s_j^a s_{j-1}^a)$ 
9:    $r_j \leftarrow s_j^a$ 
10:   $s_j^a \leftarrow -w_j s_j^a s_{j-1}^a$ 
11:   $s_j^c \leftarrow w_j s_j^c$ 
12: end for
13: for  $j = n - 2 : 2 : -1$  do
14:   $s_j^a \leftarrow s_j^a - s_j^c s_{j+1}^a$ 
15:   $w_j \leftarrow s_j^c$ 
16:   $s_j^c \leftarrow -s_j^c s_{j+1}^c$ 
17: end for
18:  $w_1 \leftarrow 1 / (1 - s_1^a s_2^a)$ 
19:  $s_1^a \leftarrow w_1 s_1^a$ 
20:  $s_1^c \leftarrow -w_1 s_1^c s_2^c$ 

```

The proposed strategy starts by dividing a batch of tridiagonal systems into multiple subdomains as in Fig. 3, where the subdomains are located across multiple ranks in a distributed memory environment. There are two novel aspects in DistD2 algorithm. First, it uses a specialised data structure to enable vectorisation on CPUs and thread level parallelism on GPUs. Second, we fuse the RHS construction based on Equation (1) with the forward pass in the decoupling phase of the algorithm such that the data movement requirements are minimised. Using the preprocessed coefficient arrays described in Algorithm 6, decoupling phase of DistD2 algorithm based on the specialised data structure can be described as shown in Algorithm 7, where $c_{[-2:2],j}^r$ are the RHS coefficients in Equation (1). The corresponding changes in the non-zero structure of the tridiagonal matrix are shown in Fig. 8. Additionally, Fig. 9 illustrates the fused RHS construction and forward pass, and also emphasises the concurrent execution of SZ many systems at once per core on a CPU or per SM on a GPU.

Algorithm 7 DistD2 - decoupling phase.

```

1: for  $j = 1 : 2$  do
2:   for  $i = 1 : SZ$  do
3:      $u'_{i,j} \leftarrow c_{-2,j}^r u_{i,j-2} + c_{-1,j}^r u_{i,j-1} + c_{0,j}^r u_{i,j} + c_{1,j}^r u_{i,j+1} + c_{2,j}^r u_{i,j+2}$ 
4:      $u'_{i,j} \leftarrow u'_{i,j} r_j$ 
5:   end for
6: end for
7: for  $j = 3 : n$  do
8:   for  $i = 1 : SZ$  do
9:      $u'_{i,j} \leftarrow c_{-2,j}^r u_{i,j-2} + c_{-1,j}^r u_{i,j-1} + c_{0,j}^r u_{i,j} + c_{1,j}^r u_{i,j+1} + c_{2,j}^r u_{i,j+2}$ 
10:     $u'_{i,j} \leftarrow f_j(u'_{i,j} - r_i u'_{i,j-1})$ 
11:  end for
12: end for
13: for  $j = n - 2 : 2 : -1$  do
14:   for  $i = 1 : SZ$  do
15:      $u'_{i,j} \leftarrow u'_{i,j} - w_j u'_{i,j+1}$ 
16:   end for
17: end for
18: for  $i = 1 : SZ$  do
19:    $u'_{i,1} \leftarrow f_1(u'_{i,1} - w_1 u'_{i,2})$ 
20: end for

```

After the decoupling phase of the algorithm is complete, the next step is solving the reduced systems. As we are focusing on diagonally

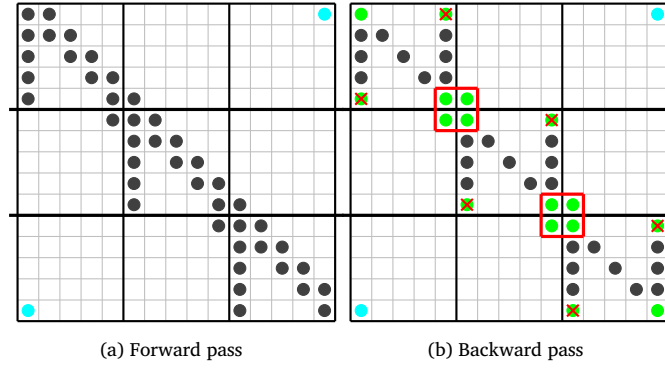


Fig. 8. Distributed-memory tridiagonal matrix algorithm. Forward and backward passes as in modified Thomas algorithm result in reduced system shown with ●. ● indicate the non-zero entries in a cyclic tridiagonal system, and if they are present they remain unchanged in forward and backward passes. A diagonally dominant system results in below zero-machine values in certain locations, marked with × to indicate that they are disregarded. Red squares show the decoupled 2×2 systems across the domain. In case of a cyclic system, the non-zero entries at the corners form an additional 2×2 system between the first and last ranks.

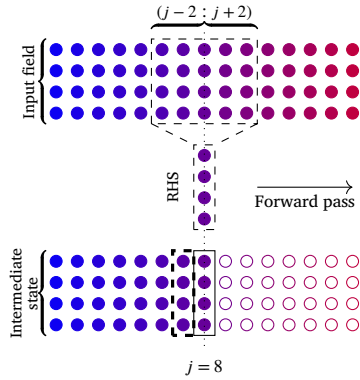


Fig. 9. Visualisation of the fused forward pass and RHS construction operation. Input field is shown on top, and the intermediate state is at the bottom. The RHS at each j is constructed using the input field values at $j - 2$ to $j + 2$. The forward pass is then applied using the intermediate state values at $j - 1$. SZ lines are processed concurrently in a CPU core or in a GPU SM.

dominant systems, certain entries in the reduced tridiagonal system are significantly below the zero-machine value, and can be eliminated without any loss in numerical accuracy. The entries that can be eliminated are indicated with red crosses in Fig. 8. Eliminating these entries simplifies the reduced systems further down to independent 2×2 systems as shown in red boxes in Fig. 8, and each of these systems are coupled across an MPI boundary. Although these systems are located in different ranks, solving them is trivial, and requires only local communication with first level neighbours. The communication phase of the DistD2 algorithm therefore carries out an MPI communication with the two neighbouring ranks, and solves the 2 × 2 systems. Single step MPI communication that involves only the previous and next neighbours is a significant advantage over existing strategies such as [3], [4], and [5] where solving the reduced tridiagonal systems require considerably more communications.

Finally, the substitution phase of the DistD2 algorithm requires a simple algebraic substitution. The first and last entries in each subdomain, $u'_{i,1}$ and $u'_{i,n}$, are solved in the communication phase. Therefore, the remaining interior entries can be solved via Algorithm 8.

Algorithm 8 DistD2 - substitution phase.

```

1: for  $j = 2 : n - 1$  do
2:   for  $i = 1 : SZ$  do
3:      $u'_{i,j} \leftarrow u'_{i,j} - s_j^a u'_{i,j-1} - s_j^c u'_{i,j+1}$ 
4:   end for
5: end for

```

▷ Vectorised Loop

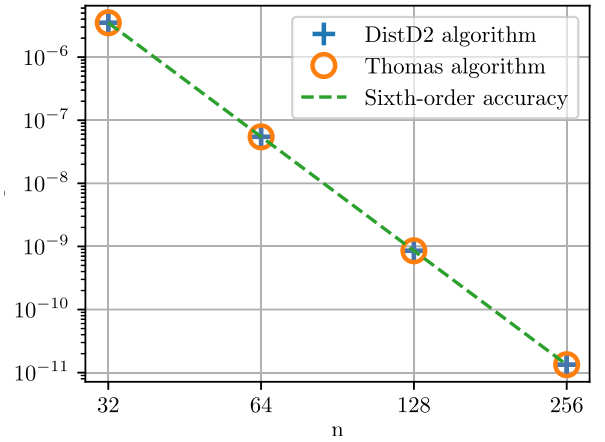


Fig. 10. Order of accuracy of DistD2 algorithm compared against Thomas algorithm.

The important point is that Algorithms 7 and 8 consider a group of RHS arrays together as shown in Fig. 6, and the inner i loop over SZ many RHS arrays indicates a vectorised execution on CPUs and thread level parallel execution on GPUs. On CPUs the vectorisation can be enabled via OpenMP *simd* directive [14], and the kernel that solves the tridiagonal systems for a group of lines is called within a parallel for loop over all the groups in the domain. On GPUs the thread level parallelism is achieved via executing the kernel with a thread dimension $SZ = 32$, and a block dimension equal to the number of groups in the domain [15]. Therefore, each thread is assigned a single RHS and the sequential operations across the j loop do not block the concurrent execution of individual systems.

Finally, we have carried out an order of accuracy analysis and compared DistD2 algorithm against Thomas algorithm to demonstrate that disregarding the zero-machine terms in the matrix do not reduce the accuracy of the algorithm. The tests are carried out from 32 points up to 256 points for the first derivative using compact schemes in double precision with periodic boundary conditions such that boundary conditions do not impact the accuracy. Periodic Thomas algorithm is executed on a single rank and DistD2 algorithm on 2 ranks. The results are shown in Fig. 10 and demonstrates that DistD2 algorithm obtains identical results compared to the Thomas algorithm.

6. Results and performance

In this section, we first focus on demonstrating the performance of the proposed algorithm on a single rank and at scale, and then assess the performance of the overall strategy on a multi-dimensional PDE.

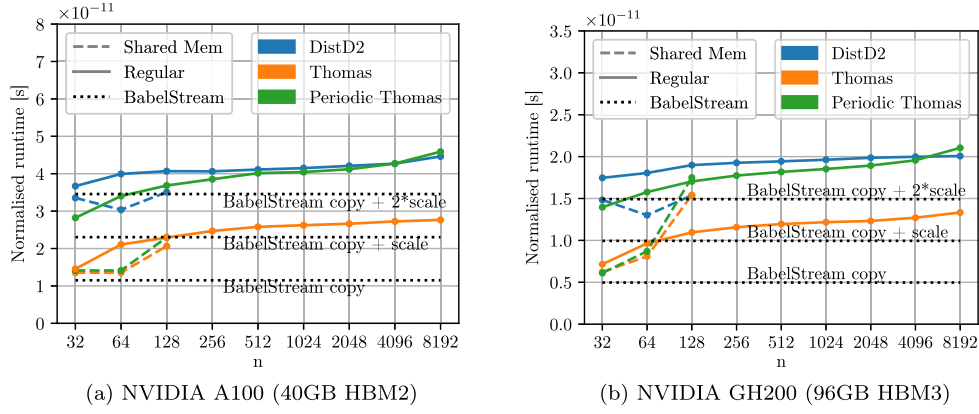


Fig. 11. Single rank performance of the DistD2 and Thomas algorithms on NVIDIA A100 (40 GB HBM2)(a) and GH200 (96 GB HBM3) (b) GPUs.

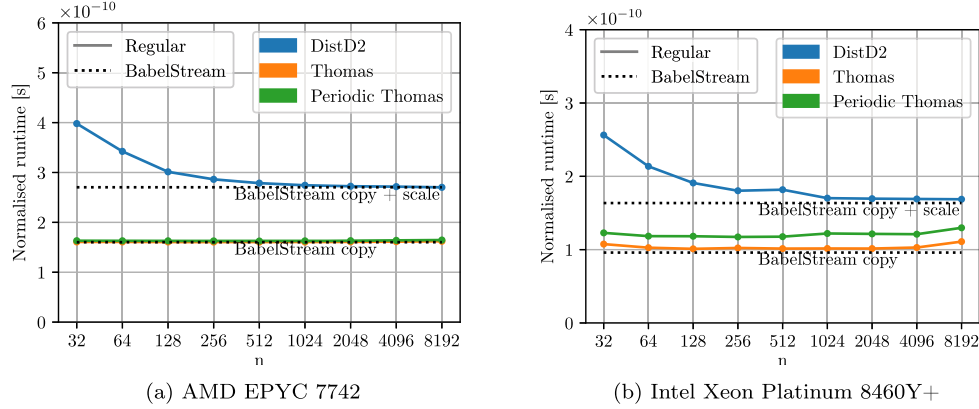


Fig. 12. Single rank performance of the DistD2 and Thomas algorithms on AMD EPYC 7742 (a) and Intel Xeon Planitum 8460Y+ (b) CPUs.

The tests carried out in Sections 6.1 and 6.2 correspond to performing a numerical calculation that represents a single partial differential term in a multi-dimensional PDE discretised using compact finite difference schemes. Section 6.3 on the other hand investigates solving a complete 3-dimensional PDE with all terms discretised using compact schemes and including the reorder operations that are required due to the specialised data structure when switching between spatial dimensions. Therefore, this section presents the performance of the proposed strategy in a complete range starting from the underlying building blocks and up to a complete PDE.

6.1. Single-rank performance

First, we investigate the single-rank performance both for the DistD2 algorithm and for the Thomas algorithm implemented using the proposed data structure on CPU and GPU architectures. In order to provide a baseline for comparison, we worked out the data movements these algorithms require in terms of domain sized read, write, and read&write operations when solving a batch of tridiagonal systems in a given domain. Data movement requirements based on the algorithm are given in Table 3 both for cached and standard implementations where cached indicates the intermediate states as the algorithm progresses are stored in cache, and standard indicates the intermediate states are not stored in cache, thus requiring additional data movements from the main memory.

The FLOP requirements of these algorithms are very low as shown in Table 3, thus, we include copy and scale benchmarks from STREAM [16,17] that closely match the data movement requirements of the tridiagonal matrix solver algorithms. The GPU and CPU benchmarks for all three tridiagonal matrix solver algorithms implemented based on the proposed data structure are shown in Fig. 11 and Fig. 12 respectively.

Table 3

Data movement requirements of all three algorithms examined in the manuscript including cached and standard implementations. Data movement requirements for the STREAM Copy and Scale benchmarks are also provided for comparison. All the numbers are normalised with respect to number of grid points.

	FLOP	Standard			Cached		
		R	W	R&W	R	W	R&W
DistD2 decoupling phase	~14	1	1	1	1	1	0
DistD2 substitution phase	~4	0	0	1	0	0	1
Thomas	~14	1	1	1	1	1	0
Periodic Thomas	~16	1	1	2	1	1	0
STREAM [16] copy	-	1	1	0	-	-	-
STREAM [16] scale (in-place)	~1	0	0	1	-	-	-

The figures also include copy and scale benchmarks executed with BabelStream [18] with a minor change in the source code such that the multiply (scale) benchmark carries out the operation in-place to reflect the data movement characteristics of DistD2 substitution phase as shown in Table 3. The benchmarks are carried out in a range of sizes of tridiagonal systems starting from 32 grid points per tridiagonal system up to 8192 grid points, and then the runtime is normalised with respect to the number of tridiagonal systems that are being solved in a 3D domain and the size of the individual tridiagonal systems. The number of systems in a 3D domain is varied such that the total number of grid points are always constant and the problem require around 75% of the available memory. All the benchmarks are carried out in double precision and we observed less than 1% variability in repeated tests. Because the FLOP and data movement requirements scale linearly with the size of the tridiagonal systems and number of systems in a 3D domain, we anticipate a

Table 4

Bandwidth utilisation on various hardware. The theoretical maximum available bandwidths are provided for each processor ('BW'). The achieved bandwidth is given as percentage of maximum available bandwidth ('achv') alongside the incurred data movement requirements per grid point ('req') based on the algorithm and hardware.

	BW	DistD2		Thomas		Per. Thomas	
		req	achv	req	achv	req	achv
AMD EPYC 7742	205 GB/s	5	70.1%	3	73.1%	3	72.1%
Intel Xeon 8460Y+	307 GB/s	5	71.6%	3	77.0%	3	66.4%
Nvidia A100 GPU	1550 GB/s	6	73.3%	4	77.4%	6	77.9%
Nvidia V100 GPU	900 GB/s	6	68.7%	4	75.6%	6	73.6%
Nvidia GH200 GPU	4023 GB/s	6	61.4%	4	66.6%	6	65.7%

relatively constant normalised runtime across the range of tridiagonal system sizes. The results shown in Fig. 11 are consistent with our expectations. First of all, due to the limited size of GPU shared memory, the shared memory implementation cannot go beyond a limited number of points per tridiagonal system. However, in the range where the cache size is large enough to store the intermediate state, the performance of the kernels using shared memory are on par with the BabelStream copy benchmark for the periodic and non-periodic Thomas algorithms as indicated in Table 3. Furthermore, the regular implementation of the DistD2 algorithm achieves a performance that matches the combination of one copy and two scale benchmarks, and the shared memory implementation of the DistD2 algorithm is on par with the combination of one copy and one scale benchmark as indicated in Table 3. The regular implementations do not explicitly store the intermediate states in shared memory, instead rely on data to stay in L1 cache whenever possible. Thus, we observe a better performance from the shared memory implementations. In summary, as the trends in all three algorithms demonstrate up to $n=8192$, increasing the tridiagonal system size reduces the relative portion of the system that can be stored in cache, and a small reduction in performance is observed. However, the implementations achieve a very high utilisation of the available bandwidth considering the required data movements. Furthermore, the runtimes are on par with BabelStream benchmarks that require equivalent levels of data movements.

Next, results in Fig. 12 are also consistent with our data movement based performance model. The per core private caches in modern CPUs are large enough for the sizes of tridiagonal systems that we are interested in. For example, a group of tridiagonal systems with 4096 points and $SZ=4$ would only require around 256 KiB of storage for the input and output arrays, and an additional 64-128 KiB for the coefficient arrays depending on the algorithm in double precision. The benchmarks in Fig. 12 are carried out on AMD Zen2 and Intel Sapphire Rapids CPUs with 0.5 MiB and 2 MiB L2 cache per core, respectively. Therefore, the intermediate states of the periodic and regular Thomas algorithms can be stored in cache, which makes the data movement requirements and performance on par with the BabelStream copy benchmark as shown in Fig. 12. The DistD2 algorithm on the other hand consists of decoupling and substitution phases with MPI communications in between. The decoupling phase writes the intermediate state back to main memory, and then the substitution phase reads the intermediate state from main memory. This results in additional data movements and is reflected in the performance results. The data movements of the decoupling phase of the algorithm are similar to the Thomas algorithm, however the substitution phase operates on the intermediate state in place. Therefore, the performance of the DistD2 algorithm is roughly equivalent to combination of copy and scale benchmarks as shown in Fig. 12. The higher normalised runtime on smaller tridiagonal systems on CPUs can be explained by halo data processing requiring relatively more work.

It is important to note that the performance of the DistD2 algorithm is slightly worse compared to the Thomas algorithm implementation for non-periodic case on GPUs, and for both periodic and non-periodic cases

on CPUs. However, the DistD2 algorithm can be run on distributed settings, with an excellent scaling as demonstrated in the following section. Ideally, both implementations can be available in a domain application and the decision to use one or the other can be made automatically at runtime.

Finally, all three algorithms achieve a decent utilisation of the available bandwidth across a range of tridiagonal system sizes on a variety of architectures. The bandwidth utilisations provided in Table 4 are based on the data movement requirements on each system for all three algorithms at 512 grid points per tridiagonal system.

It is important to note that AMD EPYC 7742 and Intel Xeon 8460Y+ use a write-allocate cache policy like many other modern x86 chips, which incurs an extra read operation to load the entries in cache before writing back to main memory [16]. The other systems in Table 4 do not exhibit such behaviour, and this is reflected in the incurred data movements. In summary, the results demonstrate that the proposed strategy is highly efficient and can achieve a very high utilisation of the available bandwidth near the practical maximum which is around 75-85% as shown in [18] with various STREAM benchmarks.

6.2. Multi-rank performance - strong and weak scalability

One of the key features of the proposed algorithm is that there is a significant reduction in the amount of communication between ranks compared to alternative approaches. Solving a tridiagonal system in a distributed-memory environment with DistD2 algorithm requires only 2 sets of MPI communications. The first one is to communicate the halo data which are data points located in a neighbouring rank that contribute to the RHS for the data points near the boundaries. Depending on the stencil of the RHS, the depth of halo data is typically between 2 to 4. The second set of MPI communication is required in between the decoupling and substitution phases of the DistD2 algorithm. The diagonally dominant tridiagonal systems are reduced down to 2×2 systems across each MPI boundary with the decoupling phase of the DistD2 algorithm, and then the first and last grid points of each tridiagonal system are sent and received between the neighbouring ranks. The DistD2 algorithm deviates from existing strategies at this step. Most tridiagonal solvers use a PCR-type strategy to obtain the solution of the reduced system, or solve the reduced systems via a Thomas algorithm after gathering the reduced systems on a single rank. Because we use compact schemes and always solve a diagonally dominant system, the reduced systems in our case are not coupled across the entire domain. Instead, they are coupled across 2 individual ranks across an MPI boundary. This is allowing us to carry out a single set of MPI communication and obtain the final result by carrying out the second phase of the DistD2 algorithm without any further MPI communications as in PCR strategies.

Fig. 13 demonstrate the strong scaling of our strategy on CPU and GPU based supercomputers. In all benchmarks we used a sixth-order compact scheme to obtain the first derivative along a spatial dimension in double precision. We start with a domain size that is big enough to use around 75% of the available memory, which corresponds to 12.88

Table 5

Data movement requirements for the transport equation on NVIDIA GPUs and x86 CPUs. A kernel fusion strategy is used on NVIDIA GPUs and a cache blocking strategy is used on x86 CPUs to minimise data movement requirements. $a + b$ format reflects the requirements in decoupling and substitution phases of DistD2 algorithm. ‘#’ indicates number of times a kernel is called, and ‘T’ indicates the incurred total data movement requirement based on the specific implementation and hardware.

	#	NVIDIA GPU				x86 CPU			
		R	W	R&W	T	R	W	R&W	T
Abstract kernel ($i \neq j$)	6	2+4	3+1	3+0	96	2+3	3+0	0+1	78
Abstract kernel ($i = j$)	3	1+4	3+1	3+0	45	1+3	3+0	0+1	36
Reordering	6	1	1	0	12	1	1	0	18
Accumulation	6	1	0	1	18	1	0	1	18
Total					171				150

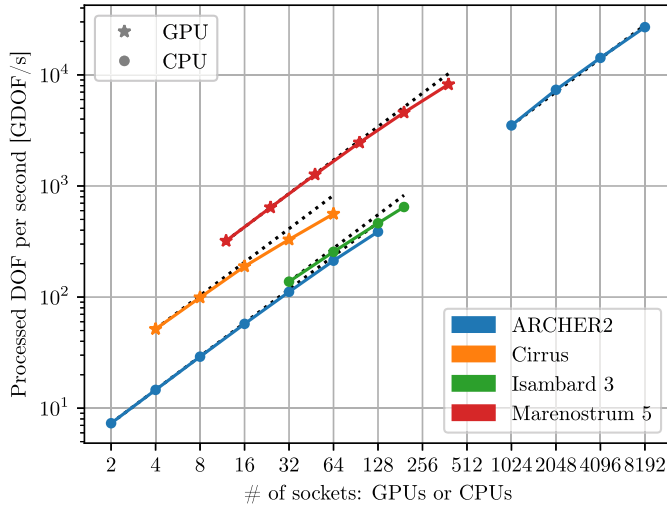


Fig. 13. Strong scaling of DistD2 algorithm on ARCHER2 [19], Cirrus [20], Isambard 3 [21], and MareNostrum 5 [22] supercomputers.

billion grid points on a single node on ARCHER2 (2x AMD EPYC 7742 CPU, 256 GiB DDR4 in total), 3.22 billion grid points on a single node on Cirrus (4x NVIDIA V100 GPU, each GPU with 16 GiB HBM), 11.95 billion grid points on a single node on Isambard 3 (2x NVIDIA Grace CPU, 240 GiB LPDDR5X in total), and 12.88 billion grid points on a single node on MareNostrum 5 (4x NVIDIA H100 GPU, each GPU with 64 GiB HBM2). The domain is decomposed with a 1D decomposition strategy such that all the tridiagonal systems in the domain are divided into as many subdomains as there are sockets and distributed such that each socket operates on one subdomain from all of the tridiagonal systems in the domain. The results demonstrate that the efficiency of strong scalability is 82.7% on ARCHER2 from 2 to 128 CPUs and 95.8% from 1024 to 8192 CPUs, 68.0% on Cirrus from 4 to 64 GPUs, 78.3% on Isambard from 32 to 192 CPUs, and 79.9% on MareNostrum 5 from 12 to 384 GPUs.

6.3. Application - 3D non-linear PDE

In this section we examine the performance of the proposed algorithm on a practical PDE in a 3D domain. We apply compact scheme discretisation to the momentum transport equation, an important step when solving incompressible Navier-Stokes equations with fractional time stepping strategy [8]. The momentum transport equation is often formulated in a skew-symmetric form using the continuity equation and this results in one additional term per equation. The skew-symmetric formulation is useful for the stability when using collocated velocity fields

for the incompressible Navier-Stokes equations. The formulation of the momentum transport equation can be given in skew-symmetric form as,

$$\frac{\partial \mathbf{u}^*}{\partial t} = -\frac{1}{2} [\nabla(\mathbf{u} \otimes \mathbf{u}) + (\mathbf{u} \cdot \nabla)\mathbf{u}] + \nu \nabla^2 \mathbf{u} \quad (3)$$

where $\mathbf{u}$ is the velocity, $\mathbf{u}^*$ denotes the intermediate field in fractional time stepping strategies, and ν is the kinematic viscosity. Equation (3) can be simplified and given in index notation as

$$\frac{\partial u_i^*}{\partial t} = -\frac{1}{2} \left(u_j \frac{\partial u_i}{\partial x_j} + \frac{\partial u_j u_i}{\partial x_j} \right) + \nu \frac{\partial^2 u_i}{\partial x_j^2}, \quad (4)$$

where the terms operated on the same spatial direction are grouped together.

We use compact finite difference discretisation to evaluate the spatial derivative terms on the right hand side. The numerical characteristics of compact schemes have been studied in depth elsewhere [1,8], therefore, in this section we evaluate the performance of the proposed strategy by focusing on the spatial discretisation and its implementation.

As discussed in Section 5, the proposed data structure layout depends on the spatial direction of the derivative operations. Thus, in a multi-dimensional PDE there is a requirement to change the data layout for each direction. In this benchmark we demonstrate that the data layout changes account for only a small percentage of the overall runtime, thus the overall strategy is efficient for solving 3D PDEs.

Furthermore, this equation system provides opportunities for fusing multiple operators into one to reduce the data movement requirements even further. For example, x -directional derivatives in Equation (4) all require the u field, thus, it is possible to reduce the overall data movements by implementing a dedicated kernel that reads the u field once and solves for all three x -directional operators concurrently. The abstract implementation based on DistD2 algorithm obtains the combination of three terms in the form of

$$RHS_j^{u_i} = -\frac{1}{2} \left(u_j \frac{\partial u_i}{\partial x_j} + \frac{\partial u_j u_i}{\partial x_j} \right) + \nu \frac{\partial^2 u_i}{\partial x_j^2}, \quad (5)$$

which requires 1 input field when $i = j$ and 2 input fields when $i \neq j$, and is executed three times per equation. The operators in y - and z -directions require all three input fields to be reordered based on the specialised data structure explained in Section 5, and then three solution fields in y - and z -directions need to be accumulated back into x -directional result fields. Time integration can be carried out after this stage, however, as it is trivial and involves only vector additions we do not investigate its performance.

In order to assess the performance of the proposed strategy, we first identify the total amount of data movements our implementation requires including kernels carrying out data layout reorderings and result accumulations. The data movement requirements are given in Table 5 normalised by the size of a scalar field on a given grid, for each kernel, specifying the number of times a kernel is called and the total data move-

Table 6

Runtime per step (t) and the sustained bandwidth utilisation as a ratio of available bandwidth (BW). MareNostrum 5 has 4x Nvidia H100 GPUs per node, and ARCHER2 has 2x AMD EPYC 7742 CPU per node.

	32 nodes		64 nodes	
	t [s]	BW	t [s]	BW
MareNostrum 5	0.662	68.0%	0.337	66.9%
ARCHER2	9.057	66.2%	4.568	65.7%

ment it contributes per full evaluation for two different platforms. The reordering operations that are introduced due to the specialised data structure account for only 7.0% and 12.0% of the total data movement requirement on GPUs and CPUs respectively. Therefore, the impact of the additional operations required due to the specialised data structure are minimal when solving a multi-dimensional PDE based on the proposed strategy.

We have carried out the benchmarks on two different platforms at scale to demonstrate the robustness of the algorithm. On both platforms we used a grid with 4096^3 points resulting in 68.7 billion grid points in total and the domain is a square box with periodic boundary conditions. The GPU benchmarks are carried out on MareNostrum 5 and the CPU benchmarks are carried out on ARCHER2. The available bandwidth on MareNostrum 5 is 4×1600 GB/s per node due to the HBM2 memory used in custom H100 GPUs, and 409 GB/s per node on ARCHER2 due to 8 memory channels with 3200 MHz DDR4 setup. We compare the actual runtimes from benchmarks and calculate the bandwidth utilisation by dividing the data movement requirements to runtimes. The results are given in Table 6. Both systems have 256 GB main memory per node, and the benchmarks use 87.5% and 43.8% of the available memory on 32 and 64 nodes respectively. A 3D decomposition strategy is used due to the large size of the domain, and the decomposition on 32 nodes is $4 \times 4 \times 8$ on MareNostrum 5 such that there is 1 rank per GPU and $4 \times 8 \times 8$ on ARCHER2 such that there are 4 ranks per CPU to account for 4 NUMA zones. The ratio of the runtime per step on MareNostrum 5 and ARCHER2 is around 13.6 and it is reasonable considering that a single MareNostrum 5 node has 15.65x more memory bandwidth but requires 14% more data movements due to minor differences in implementation which brings expected speedup down to 13.73x. The sustained bandwidth ranges from 66% to 68% of the theoretical peak, and the strong scalability from 32 nodes to 64 nodes is 98.3% on MareNostrum 5 and 99.1% on ARCHER2.

7. Conclusion

Solving a batch of tridiagonal systems in distributed-memory environments is a challenging problem. First of all, most existing strategies require more communication as the number of ranks increase, preventing a linear strong scaling. This is particularly a problem with large scale supercomputers with thousands of CPUs or GPUs. The proposed algorithm relies on an established methodology to minimise the communication requirements based on the diagonal dominance in certain tridiagonal systems. Furthermore, our approach is based on a specialised data structure to improve the performance by using data localities and data continuity. As a result, data movements from and to memory are reduced significantly, achieving a near peak bandwidth utilisation, and vectorisation on CPUs and thread level parallelism on GPUs. We also demonstrated the strong scaling on up to 8192 AMD EPYC 7742 CPUs on ARCHER2, and up to 384 NVIDIA H100 GPUs on MareNostrum 5. Finally, we provided a practical application using compact finite difference schemes and solved a 3D PDE to demonstrate the robustness of the algorithm. The performance analysis indicates that DistD2 algorithm sustains around 66% of theoretical peak bandwidth on ARCHER2, and 67% of theoretical peak bandwidth on MareNostrum 5 at scale when solving a practical 3D PDE.

CRediT authorship contribution statement

Semih Akkurt: Software, Conceptualization, Writing – review & editing, Methodology, Writing – original draft, Investigation. **Sébastien Lemaire:** Conceptualization, Writing – review & editing, Software. **Paul Bartholomew:** Conceptualization, Writing – review & editing, Software. **Sylvain Laizet:** Writing – review & editing, Conceptualization, Supervision, Funding acquisition.

Declaration of competing interest

The authors declare that they have no known competing financial interests or personal relationships that could have appeared to influence the work reported in this paper.

Acknowledgements

The authors would like to thank the Engineering and Physical Research Council (EPSRC) for their support via grant number EP/W026686/1, the Edinburgh Parallel Computing Centre (EPCC) for their support via grants ARCHER2 CPUeCSE10-1 and GPUeCSE0110, EuroHPC for their support via grant number EHPC-DEV-2024D06-021 providing development access to MareNostrum5, the Bristol Supercomputing Centre (BrICS) for their support via early access calls to Isambard 3 and Isambard-AI, and the UK Turbulence Consortium for providing access to ARCHER2 and Cirrus via EPSRC grant number EP/X035484/1. For the purpose of open access, the author has applied a Creative Commons Attribution (CC BY) licence to any Author Accepted Manuscript version arising from this submission.

Data availability

Data will be made available on request.

References

- [1] S.K. Lele, Compact finite difference schemes with spectral-like resolution, *J. Comput. Phys.* 103 (1) (1992) 16–42.
- [2] D.W. Peaceman, H.H. Rachford, The numerical solution of parabolic and elliptic differential equations, *J. Soc. Ind. Appl. Math.* 3 (1) (1955) 28–41.
- [3] E. László, M. Giles, J. Appleby, Manycore algorithms for batch scalar and block tridiagonal solvers, *ACM Trans. Math. Softw.* 42 (4) (2016) 1–36.
- [4] G.D. Balogh, T.S. Flynn, S. Laizet, G.R. Mudalige, I.Z. Reguly, Scalable many-core algorithms for tridiagonal solvers, *Comput. Sci. Eng.* 24 (1) (2022) 26–35.
- [5] K.-H. Kim, J.-H. Kang, X. Pan, J.-I. Choi, PaScaLTDMA: a library of parallel and scalable solvers for massive tridiagonal systems, *Comput. Phys. Commun.* 260 (2021) 107722.
- [6] H. Song, K.V. Matsuno, J.R. West, A. Subramaniam, A.S. Ghate, S.K. Lele, Scalable parallel linear solver for compact banded systems on heterogeneous architectures, *J. Comput. Phys.* 468 (2022) 111443.
- [7] X.-H. Sun, Application and accuracy of the parallel diagonal dominant algorithm, *Parallel Comput.* 21 (8) (1995) 1241–1267.
- [8] S. Laizet, N. Li, Incompact3d: a powerful tool to tackle turbulence problems with up to $O(10^5)$ computational cores, *Int. J. Numer. Methods Fluids* 67 (11) (2011) 1735–1757.
- [9] L.H. Thomas, Elliptic problems in linear difference equations over a network, *Watson Sc. Comp. Lab. Rep.*, Columbia University, New York, 1949.
- [10] J. Sherman, W.J. Morrison, Adjustment of an inverse matrix corresponding to a change in one element of a given matrix, *Ann. Math. Stat.* 21 (1) (1950) 124–127.
- [11] S. Rolfo, C. Flageul, P. Bartholomew, F. Spiga, S. Laizet, The 2DECOMP&FFT library: an update with new CPU/GPU capabilities, *J. Open Source Softw.* 8 (91) (2023) 5813.
- [12] W. Gander, G.H. Golub, Cyclic reduction - history and applications, in: *Proceedings of the Workshop on Scientific Computing*, 1997.
- [13] S. Akkurt, F. Witherden, P. Vincent, Cache blocking strategies applied to flux reconstruction, *Comput. Phys. Commun.* 271 (2022) 108193.
- [14] OpenMP Reference Guide, <https://www.openmp.org/resources/refguides/>, 2021.
- [15] CUDA Programming Model, <https://docs.nvidia.com/cuda/cuda-c-programming-guide/>, 2024.
- [16] J.D. McCalpin, Memory bandwidth and machine balance in current high performance computers, in: *IEEE Computer Society Technical Committee on Computer Architecture (TCCA) Newsletter*, 1995, pp. 19–25.
- [17] J.D. McCalpin, Stream2, <https://www.cs.virginia.edu/stream/stream2/>, 1999.

- [18] T. Deakin, J. Price, M. Martineau, S. McIntosh-Smith, Evaluating attainable memory bandwidth of parallel programming models via babelstream, *Int. J. Comput. Sci. Eng.* 17 (3) (2018) 247–262.
- [19] ARCHER2 Hardware Specifications, <https://www.archer2.ac.uk/about/hardware.html>, 2021.
- [20] Cirrus Hardware Specifications, <https://www.cirrus.ac.uk/about/hardware.html>.
- [21] Isambard 3 Hardware Specifications, <https://docs.isambard.ac.uk/specs/#system-specifications-isambard-3-grace>, 2024.
- [22] Marenostrum 5 Hardware Specifications, <https://www.bsc.es/ca/marenostrum/marenostrum-5>, 2023.