

Scaling and performance portability of the particle-in-cell scheme for plasma physics applications through mini-apps targeting exascale architectures

Sriramkrishnan Muralikrishnan^a, Matthias Frey^b, Alessandro Vinciguerra^c,
Michael Ligotino^c, Antoine J. Cerfon^d, Miroslav Stoyanov^e, Rahul Kumar
Gayatri^f, Andreas Adelman^a

^aPaul Scherrer Institut, Forschungsstrasse 111, 5232 Villigen, Switzerland.

^bMathematical Institute, University of St Andrews, KY16 9SS, UK.

^cETH Zurich, Switzerland.

^dCourant Institute of Mathematical Sciences, New York University, New York NY 10012, USA.

^eOak Ridge National Laboratory, Oak Ridge, USA.

^fNational Energy Research Scientific Computing Center, Berkeley, CA, USA.

Abstract

We perform a scaling and performance portability study of the particle-in-cell scheme for plasma physics applications through a set of mini-apps we name “Alpine”, which can make use of exascale computing capabilities. The mini-apps are based on Independent Parallel Particle Layer, a framework that is designed around performance portable and dimension independent particles and fields.

We benchmark the simulations with varying parameters such as grid resolutions (512^3 to 2048^3) and number of simulation particles (10^9 to 10^{11}) with the following mini-apps: weak and strong Landau damping, bump-on-tail and two-stream instabilities, and the dynamics of an electron bunch in a charge-neutral Penning trap. We show strong and weak scaling and analyze the performance of different components on several pre-exascale architectures such as Piz-Daint, Cori, Summit and Perlmutter. While the scaling and portability study helps identify the performance critical components of the particle-in-cell scheme in the current state-of-the-art computing architectures, the mini-apps by themselves can be used to develop new algorithms and optimize their high performance implementations targeting exascale architectures.

Keywords: PIC, Exascale, Performance portability, Plasma physics, Mini-apps, Kokkos

1. Introduction

Heterogeneous computing architectures are unavoidable as scientific computing moves towards the era of exascale computing. This is already evident

Name	Location	CPUs	GPUs
Perlmutter	NERSC, USA	AMD Milan	NVIDIA A100
Summit	ORNL, USA	IBM POWER9	NVIDIA V100
Sunway TaihuLight	NSCC-Wuxi, China	Sunway SW26010	-
Fugaku	RIKEN, Japan	ARM A64FX	-
Frontier	ORNL, USA	AMD EPYC	AMD Instinct
Aurora	ANL, USA	Intel Xeon	Intel Xe

Table 1: An incomplete list of some of the existing (in red) and future (in blue) computing systems showing the diversity of HPC landscape.

from some of the existing and future supercomputers listed in Table 1 which consist of different types of CPUs and GPUs.

The programming paradigms or languages used with these architectures can differ significantly from one another. As a result, naively written codes could require rewrites of considerable portions of the codes simply to achieve compatibility with a given architecture, to say nothing of efficiency. In that context, performance portability is a key criteria for current and future simulation codes, if they are to be able to take advantage of the newest developments in computing architectures with minimal maintenance.

In the plasma physics community, particle-in-cell (PIC) schemes have been widely used for the simulation of kinetic plasmas since their inception [1, 2, 3]. The attractive features of PIC schemes include simplicity, ease of parallelization and robustness for a wide variety of physical scenarios [4]. Because of their flexibility and versatility, PIC schemes are employed in many production level plasma simulation and particle accelerator codes such as TRISTAN-MP [5, 6], ORB5 [7], XGC [8], OSIRIS [9], IMPACT-T [10], OPAL [11] and Warp-X [12], to name a few.

Some of these production codes have already begun their journey towards performance portability as, evidenced in [13, 14]. It is expected that a lot more will also do so sometime soon, in order to benefit from the high performance of existing and future advanced computing architectures. Mini-applications (or mini-apps) which are performance portable can greatly help in this respect [15]. These are light-weight proxy codes which contain performance critical components of the application codes. They can be used for implementing new algorithms, optimizing implementations, and provide reliable performance expectations for the real application of interest on different computing architectures [15]. As such, mini-apps serve as an interface between applied mathematics, high performance computing and production codes.

Our objective in this work is to study the performance portability and scaling of the components of PIC schemes through a set of mini-apps (Alpine¹) with applications in plasma physics and particle accelerator modeling, target-

¹ALPINE: A set of performance portable pLasma physics Particle-in-cell mINi-apps for Exascale

ing exascale architectures. In the present article, we consider physical situations for which the electrostatic assumption is justified, and for which collisions can be neglected. In future work, we plan to extend the study by enriching the collection of mini-apps with electromagnetic examples, which may or may not include collisions and electron pair production. The mini-apps are built using the performance portable library Independent Parallel Particle Layer (IPPL) which is described in detail in Section 2. The contributions of the current work are summarized below:

- Alpine provides a test bed for implementing new algorithms and/or novel implementations of existing algorithms related to PIC schemes in the context of exascale architectures in a portable way.
- The performance study in this work provides important insights on how different components of PIC schemes function on different architectures.
- This work serves as guidance for the plasma PIC community to identify the major causes of performance limitations and better prepare for exascale architectures.
- So far, portable exascale PIC studies have mostly been conducted for PIC schemes designed for electromagnetic plasma models [14, 16]. To the best of our knowledge, this is the first study which considers the performance of a PIC scheme for an electrostatic plasma model in the portable exascale context.

This paper is organized as follows. Section 2 describes the IPPL library. The electrostatic PIC scheme implemented in the mini-apps is described in Section 3. Section 4 describes the mini-apps along with their physical parameters and verification studies. The strong and weak scaling results, as well as the performance analysis of different components on four different pre-exascale architectures, are presented in Section 5. Finally, we summarize our results in Section 6, and propose directions for future work.

2. IPPL

The Independent Parallel Particle Layer (IPPL) C++ library was developed about 20 years ago and was inspired by and partially based on POOMA [17]. The general framework is designed to enable the rapid development of Lagrangian, Eulerian, and hybrid Eulerian-Lagrangian solvers such as solvers based on the PIC method. IPPL uses the Message Passing Interface (MPI) paradigm to distribute fields and particles across multiple processes. It also makes use of expression templates [18] to speed up the computation of mathematical operations on field and particle data.

Besides revising IPPL to the latest C++17 standard, we replaced the core data structures of IPPL with Kokkos [19, 20] data structures to enable performance portability across various hardware architectures. Although Kokkos

allows more flexibility, for simplicity as well as to avoid costly host-device copies IPPL currently allocates all particles and fields on the target device which is either GPU or CPU, but not both.

Figure 1 summarizes how IPPL treats particles. It uses the paradigm of *struct of arrays* where each particle attribute (`ParticleAttrib`) is in principle a Kokkos View enhanced with expression templates. Note that the underlying data type of a particle attribute is not restricted to scalar types; hence vector attributes (e.g. position and velocity) themselves are *arrays of struct*. The basic collection of particles consisting of a position (`R`) and unique identity number (`ID`) attributes is represented with `ParticleBase`. Its parallel distribution among processes is managed by the class `ParticleLayout`. Applications must derive their own particle container class from `ParticleBase`. An exam-

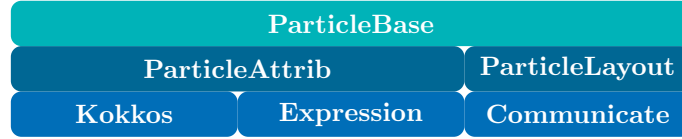


Figure 1: General layout of the particle container class (`ParticleBase`) in IPPL. The underlying container of each particle attribute (`ParticleAttrib`) is a Kokkos View. The memory layout of the particle container is therefore a *struct of arrays*. The communication among processes is managed by its layout (`ParticleLayout`).

ple is shown in Lst 1. Supplementary particle attributes are easily added with `addAttribute`. This makes sure that the new attributes are included in MPI send and receive operations. In order to send particles from one process to another, all particle attributes are serialized into one buffer.

These buffers are allocated with an over-allocation factor which is a parameter in IPPL and can be specified by the users according to their application needs to avoid performance degradation, especially on GPU systems. A high value of this factor is usually favored, as it reduces the number of future allocations after the initial allocation, which are costly on GPUs. However, we may not have enough memory on the GPUs to accommodate such over-allocation, in which case we need to reduce it. It should also be noted that the allocated buffers are never freed until the simulation ends, which may also pose memory constraints during the simulation. With an over-allocation factor of 1.0, the space allocated for each buffer is exactly the requested size.

Once the data is sent, the receiver then unpacks and appends the particles to its attribute containers (see Figure 2). The local particle container of each MPI process is also over-allocated to further reduce the amount of dynamic allocations.

To ensure an even workload across all MPI processes, the particles are redistributed regularly using the orthogonal recursive bisection algorithm [21]. For this purpose, the particle densities are interpolated onto the grid, which is then divided into subregions, each with approximately the same number of particles. The particles are then passed to the MPI process that occupies the subregion

into which they fall.

Listing 1: Example of a particle container class with additional particle attributes (mass and velocity). The design of IPPL allows the user to easily add new attributes. Mathematical expressions compile to a single Kokkos kernel without storing any temporaries, thanks to the use of expression templates.

```

using namespace ippl;
template<class PLayout>
struct Bunch : public ParticleBase<PLayout> {
    Bunch(PLayout& playout) : ParticleBase<PLayout>(playout) {
        // add application attributes
        this->addAttribute(mass);
        this->addAttribute(velocity);
    }
    ParticleAttrib<double> mass;
    ParticleAttrib<Vector<double>> velocity;
};

// compiles to single Kokkos kernel
bunch->R = bunch->R + dt * bunch->velocity;

```

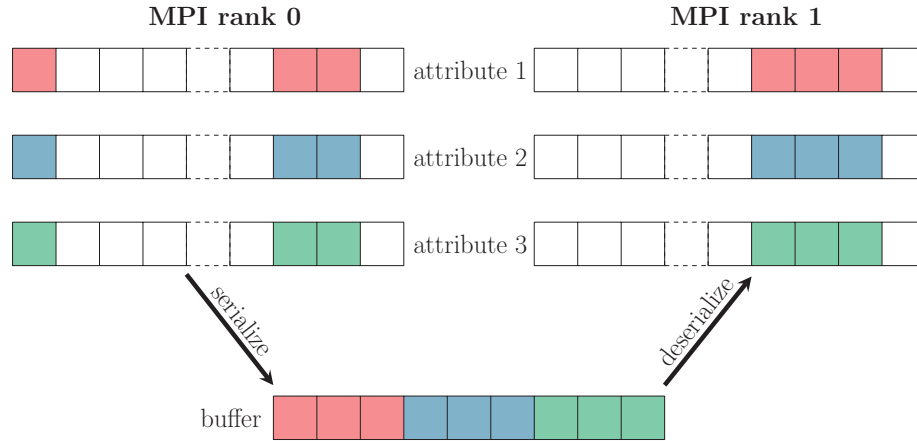


Figure 2: Sending and receiving particles. MPI rank 0 packs all attributes (in this case 3) of all the particles to be sent to rank 1 (again 3 in this case) into a single pre-allocated buffer that is sent to MPI rank 1, which unpacks and appends the received particle data.

The design of fields follows a similar structure to that of the particle container class and is shown in Figure 3. The basic structure is again a Kokkos View. We however distinguish between two types of fields: **BareField** and **Field**. The former represents a field without spatial reference. The communication of field data is managed by the field layout (**FieldLayout**). IPPL uses halo (guard or ghost) cells to exchange field data among MPI processes and it follows the same pack-send-receive-unpack set of operations we described regarding particle communication. The number of halo cells can be set for each field independently.

An interface to heFFTe [22] enables us to compute fast Fourier transforms (FFT) in a portable manner. Currently, we have to perform a copy between IPPL fields and the data structures being passed to heFFTe for computing FFTs. This is because of the presence of halo cells in IPPL fields, which have to

be stripped out before passing the field data to heFFTe. Because of this copy operation, we currently have a memory overhead which we will try to remove in future IPPL versions.



Figure 3: Design layout of fields in IPPL. The underlying data structure is a Kokkos View. IPPL provides an interface to the Fourier transform routines of heFFTe.

3. Particle-in-cell method

3.1. Particle-in-cell for the Vlasov-Poisson system

In this work, we consider the non-relativistic Vlasov-Poisson system with a fixed magnetic field, and introduce the PIC method in that setting. The electrons are immersed in a uniform, immobile, neutralizing background ion population, and the electron dynamics is given by

$$\frac{\partial f}{\partial t} + \mathbf{v} \cdot \nabla_{\mathbf{x}} f + \frac{q_e}{m_e} (\mathbf{E} + \mathbf{v} \times \mathbf{B}_{ext}) \cdot \nabla_{\mathbf{v}} f = 0, \quad (1)$$

where $\mathbf{E} = \mathbf{E}_{sc} + \mathbf{E}_{ext}$, and the self-consistent field due to space charge is given by

$$\mathbf{E}_{sc} = -\nabla \phi, \quad -\Delta \phi = \rho / \varepsilon_0 = (\rho_e - \rho_i) / \varepsilon_0. \quad (2)$$

In equation (1), $f(\mathbf{x}, \mathbf{v}, t)$ is the electron phase-space distribution and q_e , m_e and ε_0 are the electron charge, mass, and permittivity of free space respectively. The total electron charge in the system is given by $Q_e = q_e \int \int f d\mathbf{x} d\mathbf{v}$, the electron charge density by $\rho_e(\mathbf{x}) = q_e \int f d\mathbf{v}$ and the constant ion density by $\rho_i = Q_e / \int d\mathbf{x}$. Throughout this paper we use bold letters for vectors and non-bold letters for scalars. The numerical methods and algorithms we discuss in this article can be easily applied to situations involving non-uniform external magnetic fields with finite curvature. On the other hand, physical systems for which relativistic or electromagnetic effects play a significant role will in general require different approaches, and the corresponding codes will perform differently on large scale high performance computing architectures.

The particle-in-cell method discretizes the phase space distribution $f(\mathbf{x}, \mathbf{v}, t)$ in a Lagrangian way by means of macro-particles (hereafter referred to as “particles” for simplicity). At time $t = 0$, the distribution f is sampled, which leads to the creation of the computational particles. Subsequently, a typical computational cycle in PIC consists of the following steps:

1. Assign a shape function - e.g. cloud-in-cell [2] - to each particle p and deposit the electron charge onto an underlying mesh. This is known as “scatter” in PIC.

2. Use a grid-based Poisson solver to compute ϕ by solving $-\Delta\phi = \rho/\varepsilon_0$ and differentiate ϕ to get the electric field $\mathbf{E} = -\nabla\phi$ on the mesh with appropriate boundary conditions.
3. Interpolate $\mathbf{E}$ from the grid points to particle locations $\mathbf{x}_p$ using the same interpolation function as in the scatter operation. This is typically known as “gather” in PIC.
4. By means of a time integrator, advance the particle positions and velocities using

$$\frac{d\mathbf{v}_p}{dt} = \frac{q_e}{m_e} (\mathbf{E} + \mathbf{v} \times \mathbf{B}_{ext})|_{\mathbf{x}=\mathbf{x}_p}, \quad (3)$$

$$\frac{d\mathbf{x}_p}{dt} = \mathbf{v}_p. \quad (4)$$

3.2. Numerical implementation

In our implementation we use the non-dimensional form of the Vlasov-Poisson system and we follow the same non-dimensionalization as in [23] (see Table I). In this case, the equations for the PIC scheme look very similar to the dimensional form except that in equation (3) the factor q_e/m_e becomes -1 and ε_0 in equation (2) is taken as 1. To avoid introducing new notations, in the rest of the article, we denote the non-dimensional quantities with the same notations that we introduced before for their dimensional counterparts. Now, we succinctly describe the numerical algorithms that are used for the four PIC steps enumerated in the previous section in all our mini-apps.

First, we do a random sampling of the particles based on the distribution function f using inverse transform sampling [24]. It is important to create the particles locally in a load balanced manner² by each MPI rank or GPU. Otherwise, the very high communication required after the particle creation can have a cost which is prohibitive, even if it is a one time cost. It can also lead to a crash in the simulation due to memory requirements during communication or in the creation stage itself, especially on GPUs. Moreover, for simulations on GPUs, we create the particles directly on the device using the Kokkos random number generator rather than first creating them on the host (CPU) and then copying them to the device. This host-to-device copy would also have very significant cost for the large numbers of particles that we are targeting in this work, and be prohibitively costly even if it is a one time cost. However, it should be noted that the random number generation on GPUs is not deterministic even with a fixed seed and this can present some difficulties with respect to reproducibility. Nevertheless, the results are still reproducible in a statistical sense.

We use a cell-centered grid and cloud-in-cell shape function for interpolation from particles to grid and vice versa. We use periodic boundary conditions in all

²In this paper, when we mention load balanced or imbalanced configurations it is only with respect to particles unless we explicitly mention the fields.

directions and since our grid spacings in the x , y and z directions are constants (although potentially different in each direction), we use an FFT-based spectral solver for the computation of the potential and electric field from the charge density. The FFTs in the field solver are taken using the heFFTe library. For the time integration, we use the synchronized form of the Boris scheme as given in equations 4(a) - 4(c) of [25], which gives both the velocity and position of the particles at integer time steps. Having explained the numerical algorithms used in our mini-apps we move on to describe the physical parameters for each of the examples considered.

4. The Mini-Apps

In this section, we briefly describe the electrostatic plasma physics problems that we consider for the scaling and performance study.

4.1. Landau damping

The first problem we consider is Landau damping in the weak and strong damping regimes. It is a classical problem which has been studied extensively in the literature [4, 26, 27, 28, 29]. The availability of analytical results makes it an excellent candidate for our verification and performance study. Similar to [4], we consider the following initial distribution

$$f(t=0) = \frac{1}{(2\pi)^{3/2}} e^{-|\mathbf{v}|^2/2} (1 + \alpha \cos(kx)) (1 + \alpha \cos(ky)) (1 + \alpha \cos(kz)), \quad (5)$$

in the domain $[0, L]^3$, where $L = 2\pi/k$ is the length in each dimension. We choose the following parameters for our weak Landau damping tests: $k = 0.5$, $\alpha = 0.05$. The total electron charge based on our initial distribution is $Q_e = -L^3$. For the case of strong Landau damping the difference is that we use a stronger perturbation parameter $\alpha = 0.5$. The other parameters are the same as those for the weak damping case. The presence of a stronger perturbation parameter in the strong damping case necessitates particle load balancing as will be shown in Figures 15, 16. As a verification, we show in Figure 4 the damping of the electric field energy in the x -direction for the weak and strong damping cases. Our results agree well with the analytical rates as well as the results in [26].

4.2. Bump-on-tail/Two-stream instability

The second mini-app we consider is the two-stream or bump-on-tail instability problem. Similar to Landau damping, this is another classical benchmark problem studied in the literature [26, 27, 28, 29], which also has analytical estimates for the growth rates, which can be calculated from the dispersion relation derived from the linearized equations.

We consider the following initial distribution of electrons

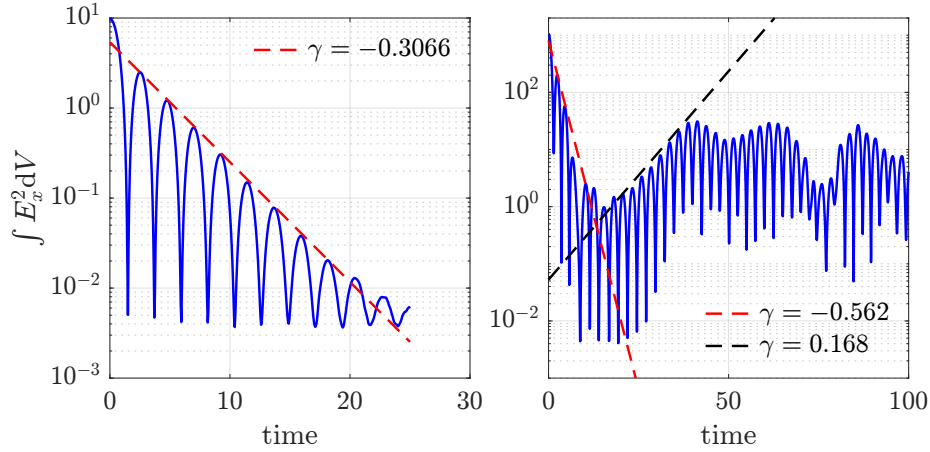


Figure 4: Landau damping of the electric field energy in the x -direction ($\int_V E_x^2 dV$, where V is the total simulation volume) as a function of time for weak (left) and strong (right) damping cases. The rates match well with the analytical values shown by the dashed lines as well as the results in [26]. The number of mesh points, number of particles and the time step in these simulations are 32^3 ; 83,886,080 and 0.05 respectively.

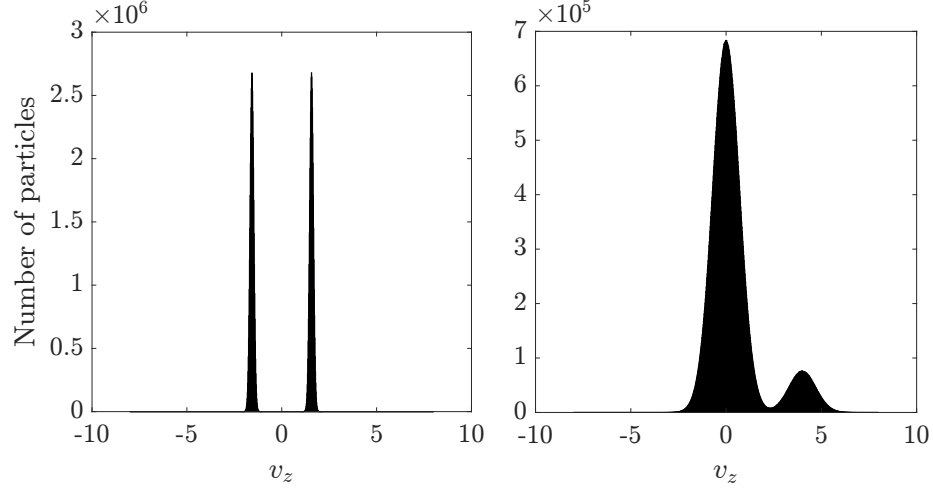
$$f(t=0) = \frac{1}{\sigma^3 (2\pi)^{3/2}} \left\{ (1-\epsilon) e^{-\frac{|\mathbf{v}-\mathbf{v}_{b1}|^2}{2\sigma^2}} + \epsilon e^{-\frac{|\mathbf{v}-\mathbf{v}_{b2}|^2}{2\sigma^2}} \right\} (1 + \alpha \cos(kz)), \quad (6)$$

in the domain $[0, L]^3$, where $L = 2\pi/k$ is the length in each dimension. Depending on the choice of parameters, we get two flavors of this example. With $\epsilon = 0.5$, $\sigma = 0.1$, $k = 0.5$, $\alpha = 0.01$, $\mathbf{v}_{b1} = \{0, 0, -\pi/2\}$ and $\mathbf{v}_{b2} = \{0, 0, \pi/2\}$ we get the two-stream instability problem studied in [26]. On the other hand, choosing $\epsilon = 0.1$, $\sigma = 1/\sqrt{2}$, $k = 0.21$, $\alpha = 0.01$, $\mathbf{v}_{b1} = \{0, 0, 0\}$ and $\mathbf{v}_{b2} = \{0, 0, 4\}$ we get the bump-on-tail instability problem similar to [30]. The total charge Q_e is chosen in the same way as in the Landau damping example.

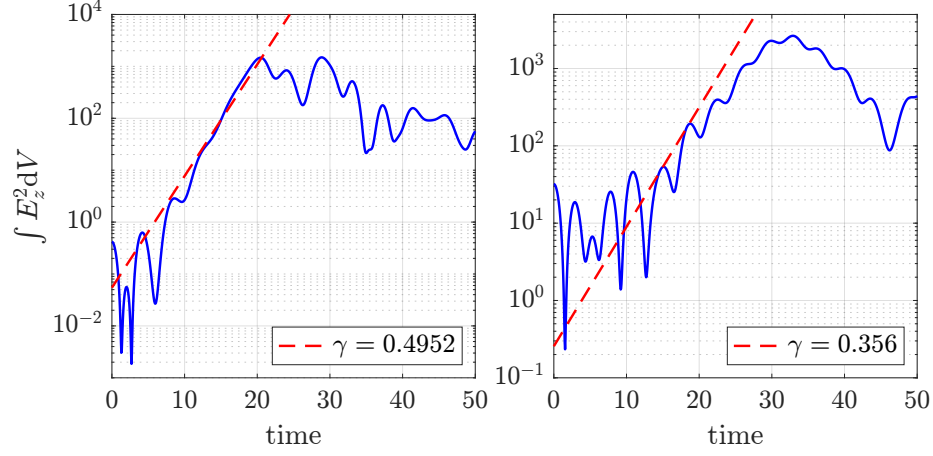
In the top row of Figure 5, we show the velocity distribution in the z -direction for these two sets of parameters. In both cases, the electric field energy grows as a function of time as shown in the bottom row of Figure 5. Similar to Landau damping, we see very good agreement with the analytical rates as well as the results in [26]. We would like to mention that even though we simulate the bump-on-tail or two-stream instability in 3D-3V in our mini-app, the essential physics for the initial distribution selected occurs predominantly in the z -direction, at least for early times.

4.3. Electron dynamics in a charge neutral Penning trap

Our next mini-app corresponds to the dynamics of electrons in a Penning trap with a neutralizing static ion background, as we already considered in [31]. This problem involves bunching of electrons in configuration space, and therefore presents challenges in terms of field and particle load balancing. The



(a) Initial velocity distribution in z -direction for two-stream (left) and bump-on-tail (right) instabilities.



(b) Electric field energy in z -direction for two-stream (left) and bump-on-tail (right) instabilities.

Figure 5: Initial velocity distribution in the z -direction (top row) and electric field energy ($\int_V E_z^2 dV$, where V is the total simulation volume) in the z -direction as a function of time (bottom row) for the two-stream and bump-on-tail instability test cases. The growth rates agree well with the analytical values shown by the dashed lines as well as the results in [26]. The number of mesh points, number of particles and the time step in these simulations are 32^3 ; 83,886,080 and 0.05 respectively.

initial conditions for this example, as well as the electron dynamics they lead to, are very similar to that of cyclotrons [32, 33, 31]. Since the particle accelerator library OPAL [11] will include the portable version of IPPL in the near future, this example is of interest from the point of view of cyclotron simulations.

Regarding the parameters for this problem, we follow the same setup as in [31]. The domain is $[0, L]^3$, where $L = 20$. The external magnetic field is given by $\mathbf{B}_{ext} = \{0, 0, 5\}$ and the quadrupole external electric field by

$$\mathbf{E}_{ext} = \left\{ -\frac{15}{L} \left(x - \frac{L}{2} \right), -\frac{15}{L} \left(y - \frac{L}{2} \right), \frac{30}{L} \left(z - \frac{L}{2} \right) \right\}.$$

For the initial conditions, we sample the phase space using a Gaussian distribution in all the variables. The mean and standard deviation for all the velocity components are 0 and 1, respectively. While the mean for all the configuration space variables is $L/2$, the standard deviations are $0.15L$, $0.05L$ and $0.2L$ for x , y and z , respectively. The total electron charge is $Q_e = -1562.5$.

4.4. Uniform plasma test

The uniform plasma test is an idealized, non-physical test case. It randomly moves all the particles in the range $[-h, h]$ during each time step and carries out the other steps as in a typical PIC cycle. This mini-app can be used as a gold standard for the performance of different components with which the other mini-apps can be gauged and understood.

5. Scaling results

5.1. Setup

In this section, we present representative strong and weak scaling results as well as the performance of different components for the mini-apps described in Section 4. We do not present the scaling results for the two-stream and bump-on-tail instabilities or the uniform plasma test, as they are very similar to those of the weak Landau damping study; the timings differ by at most ± 10 percent for the most part.

For the strong scaling, we consider the mesh and particle setups shown in Table 2. For the scaling studies up to Section 5.4, the simulations are done

Case	Grid (N_c)	Total number of Particles (N_p)	Particles per cell (P_c)
A	512^3	1,073,741,824	8
B	1024^3	8,589,934,592	8
C	256^3	1,073,741,824	64
D	512^3	8,589,934,592	64

Table 2: Cases considered for the strong scaling study with different number of grid points and particles. The total number of particles for the two sets of cases (A,B) and (C,D) are the same, however, the latter has 8 times more particles per cell than the former due to reduced grid resolutions.

on the CPU and GPU partitions of the Piz Daint supercomputer at the Swiss National Supercomputing Centre (CSCS), Switzerland. The architectures of Piz Daint’s GPU and CPU partitions are as follows:

- GPU nodes: Each Cray XC50 compute node has 1 NVIDIA Tesla P100 GPU with 16 GB memory and 12 Intel Xeon E5-2690 v3 cores at 2.6 GHz with 64 GB RAM.
- CPU nodes: Each Cray XC40 compute node has 36 Intel Xeon E5-2695 v4 cores at 2.1 GHz with either 64 or 128 GB RAM.
- The interconnect configuration is Aries and the network topology is Dragonfly.

For the GPU simulations, we use 1 MPI rank per GPU and no OpenMP thread-based parallelism, whereas for the CPU simulations we use 1 MPI rank and 32 OpenMP threads per node. Basically, we use MPI for communication between GPUs both within a node and between nodes, whereas for CPU-based simulations, we use MPI only for communication between nodes while OpenMP-based shared memory parallelism is used within a node. This setup helps to minimize the particle and field communication costs.

We conduct the strong scaling study for cases A and C with the following sets of GPU and CPU nodes: number of GPU nodes $\in \{16, 32, 64, 128, 256\}$, number of CPU nodes $\in \{4, 8, 16, 32, 64\}$. The minimum number of GPU or CPU nodes is based on the memory requirement for the given number of particles and grid points. For the maximum number of nodes, we stop after increasing the number of MPI ranks or GPUs five times, each time by a factor of 2, as is done in [14]. The reason is that the runtimes are mostly communication dominated at this stage and the problems stop scaling beyond that. For cases B and D, the number of grid points and particles are both 8 times larger and hence the number of CPU and GPU resources are increased accordingly, i.e. number of GPU nodes $\in \{128, 256, 512, 1024, 2048\}$, number of CPU nodes $\in \{32, 64, 128, 256, 512\}$.

In terms of the time step, for the Landau damping mini-app we choose it to be proportional to the mesh size h as $\Delta t = 0.5h_{min}$. We also make sure that it is below the stable time step requirement of $\Delta t \leq 2\omega_{pe}^{-1}$, where ω_{pe} is the electron plasma frequency. Since the Penning trap simulations involve a lot more particle communication than the Landau damping problem, we choose the time step based on the finest mesh used in our scaling studies, i.e. $\Delta t = 0.5(L/2048) \approx 0.005$, just to have the same dynamics, and therefore similar particle communication, for different grid resolutions in a weak scaling study.

For the mini-apps we choose an over-allocation factor of 2.0 for the weak and strong Landau damping tests, whereas for the Penning trap simulations we choose a value of 1.0 due to its high memory requirement per rank. These values are chosen based on numerical experiments.

For the ease of performance analysis, in Table 3 we split the significant kernels in our code into computation kernels, from which we can expect parallel efficiency, and communication kernels, which are required because of domain

decomposition. However, this separation is not perfect since the FFTs required for the field solve are computed using heFFTe and this includes communication in addition to the purely local operations. This is because we use heFFTe as a black box for Fourier transforms and hence do not use IPPL timers inside its source code. The total time which is included in the computation kernels column is the time taken for the entire simulation to finish including the communication kernels. However, we exclude the time spent writing output data to files and an initial warm up call to the field solve. The first call to the field solve takes significantly more time than the subsequent ones owing to the initializations performed by heFFTe. In the communication kernels, the “Fill halo cells” and “Accumulate halo cells” operations are required during the gather and scatter stages of the PIC cycle, and the “Particle update” sends the particles to the appropriate ranks once they leave the local subdomain of the current rank. The components not included in Table 3 account for less than 10 percent of the total time in most of our simulations. We therefore do not consider them in the performance study.

Computation kernels	Communication kernels
Gather (gather)	Particle update (updateParticle)
Scatter (scatter)	Fill halo cells (fillHalo)
Push position (pushPosition)	Accumulate halo cells (accumulateHalo)
Push velocity (pushVelocity)	Particle load balance (loadBalance)
Particle BCs (particleBC)	
FFT-based field solve (solve)	
Total time (total)	

Table 3: List of computation and communication kernels used for the performance study. Inside the parentheses are the labels which are used to represent the components in the figures.

We run the simulations for 20 time steps and report in figures the wall time per simulation time step for each of the kernels in Table 3. As such, our performance figures are not indicative of production runs, as one typically needs to run for thousands of time steps in real plasma simulations. Furthermore, depending on the long time dynamics of the problem under consideration the performance can be significantly different. Our objective here is to assess the performance of different components in the mini-apps across different architectures without spending too many node hours or having to wait for long periods in the job submission queues.

The versions of the compilers, MPI, Kokkos, heFFTe and IPPL used for the simulations on each of the computing architectures considered in this work is mentioned in Appendix A. For the parameters in heFFTe we chose pencil decomposition, pipelined point-to-point communication with no-reordering. We refer the readers to [22] for descriptions of these parameters. These options were chosen based on our heFFTe benchmarking study with all the possible parameter combinations, selecting the best one in terms of scalability and time to solution. In terms of domain decomposition for our mini-apps, we use parallelization in

all three directions for the fields, which gives each processor a brick of field data along with one layer of halo cells.

5.2. Weak Landau damping

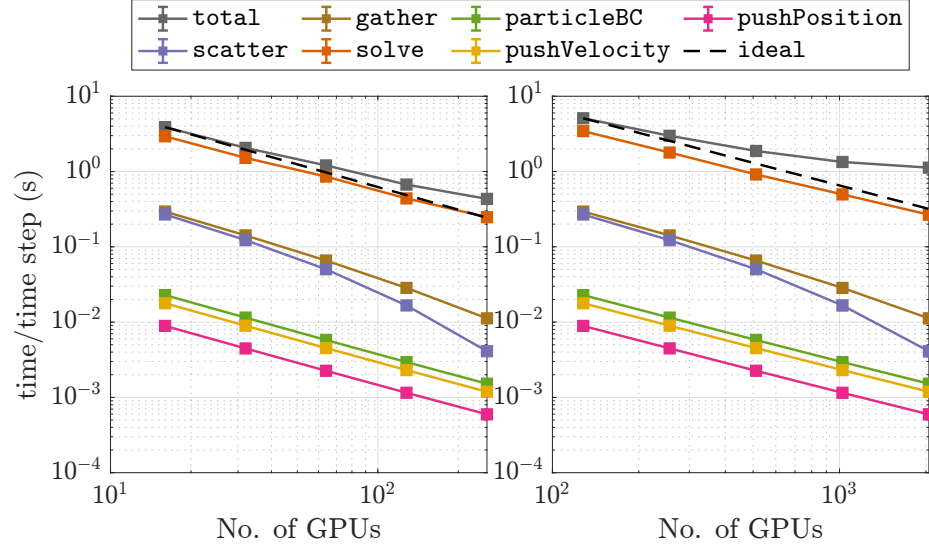
5.2.1. Strong scaling: 8 particles per cell

In this section we present the scaling results for the weak Landau damping problem. We first consider the strong scaling study. In Figure 6, we show the timings of computation and communication kernels listed in Table 3 for cases A and B on GPUs. In all the figures we show the average, minimum and maximum wall times across the MPI ranks or GPUs by means of square markers and errorbars. It is clear from the figure that the computation kernels scale almost ideally, except for the `solve`. Since the FFT-based field solver also includes communication between GPUs, the scaling deteriorates when we increase the number of GPUs. We also observe that the field solve is the dominant computation kernel. It takes an order of magnitude more time than the next dominant kernel which is the `gather`. The runtimes of the field solve are also two orders of magnitude higher than those of the pure particle kernels.

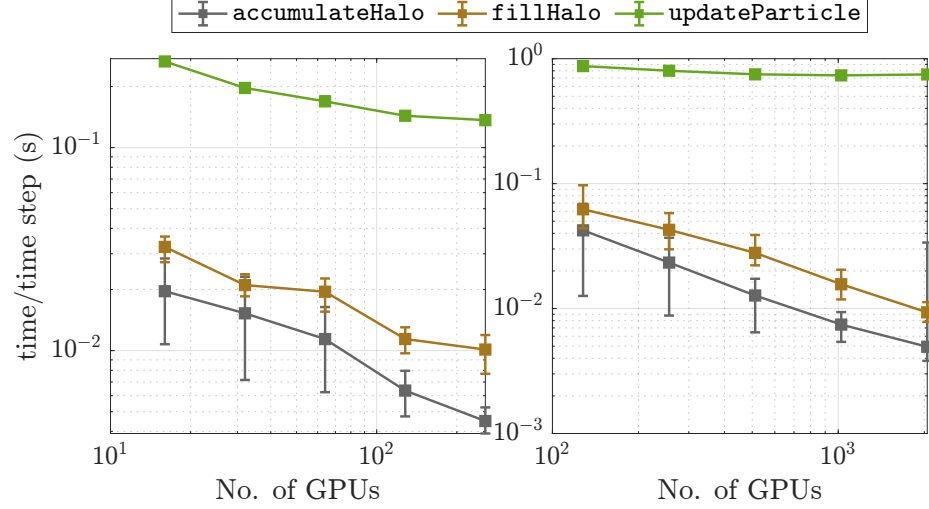
One way to improve the field solver performance is by using a 2D or 1D domain decomposition which gives pencils or slabs for the fields rather than bricks as we use in this study. This will reduce the number of transposes required in the FFT computation, which in turn leads to a reduction in runtime of the field solve. However, it can have adverse effects on the field and particle communication costs in the rest of the code. A thorough study of the different domain decompositions, in order to determine the decomposition leading to the best overall performance is beyond the scope of this article. We intend to investigate this question in our future work.

Regarding the communication kernels, the `updateParticle` takes the most time, as seen in Figure 6(b). Let us therefore examine its components in more detail:

1. The first one is the search for all the local particles across all the MPI ranks to figure out the mapping between particles and ranks after the push. Since each MPI rank knows the global information about the domains owned by the other ranks, this search happens locally without any need for communication. Then, each processor communicates to all the other processors how many particles they receive from it. These two operations do not scale in our current implementation, as they involve two loops: one over the local number of particles and another over the total MPI ranks. The product of these quantities is a constant in a strong scaling setup, and increases linearly in a weak scaling setup.
2. The next component involves packing and sending the particles which are leaving the local field domain of the current processor, and receiving and unpacking the particles which are coming from the other processors. This



(a) Computation kernels for case A (left) and case B (right)



(b) Communication kernels for case A (left) and case B (right)

Figure 6: Weak Landau damping. Strong scaling on GPUs: scaling of computation and communication kernels for cases A and B.

part scales in both the strong and weak scaling sense³.

When the second component dominates, we see a decrease in `updateParticle` runtime as illustrated by the first four data points in the left column of Figure 6(b) for case A. This is because the weak Landau damping simulations have a very small particle load imbalance. In our experiments, we noticed that it does not require load balancing. However, the first component will eventually start dominating, and we then see a flat regime in the particle update cost. For case B this happens almost from the beginning, as we can observe from the right column of Figure 6(b). This is due to both the higher particle count and the higher GPU count compared to case A.

From Figure 6(b), the field communication kernels `fillHalo` and `accumulateHalo` take much less time compared to the `updateParticle`. They also decrease in a strong scaling setup with increasing number of GPUs, as the number of halo cells to be communicated with the field neighbors decreases.

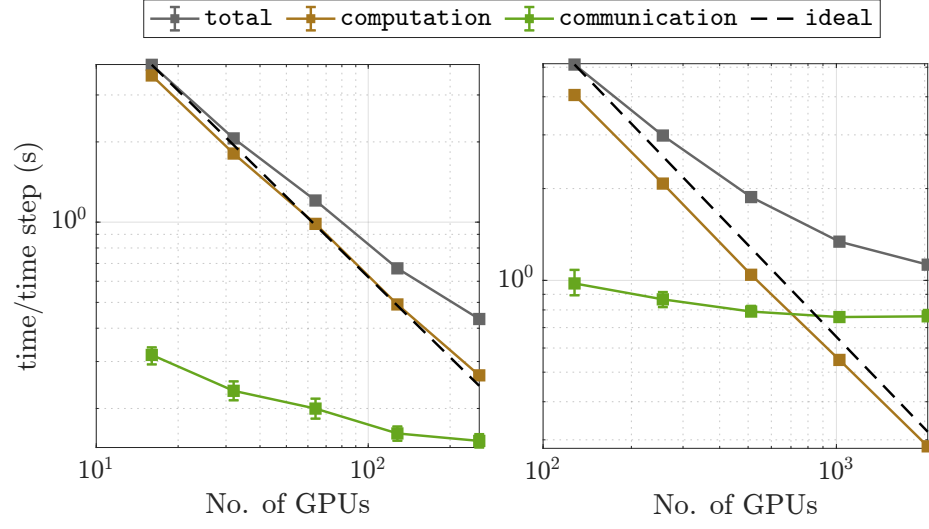
In Figure 7(a) the cumulative time of computation kernels (except total time) and the communication kernels are compared with the total time, for increasing number of GPUs. While in case A the cross over point between communication and computation costs is not seen for our set of parameters, in case B the communication kernels dominate over the computation kernels for the last two data points. The efficiency is shown in Figure 7(b) for the computation kernels. We see that the total efficiency is comparable to the field solver efficiency until the communication kernels become comparable or start dominating, at which point we observe a significant drop in the total efficiency. All other computation kernels except for the field solver exhibit near ideal efficiency, as visually observed in Figure 6.

In Figures 8 and 9 we show the results for the CPU-based strong scaling study. Here, we only highlight the differences with the GPU simulations. We first observe from Figure 8(a) that even though field solve is still the dominant computation kernel, the difference between `solve` and the next dominant kernel, `gather`, as well as the pure particle kernels, is much less than on the GPUs. Next, we note that the field solve efficiency and the total efficiency are higher for the CPU cases than the GPU ones. This is partly due to the fact that we can use fewer MPI ranks in the CPU simulations than the GPU simulations, because the CPU nodes have more memory. This in turn leads to lower communication costs and increased efficiency. Apart from these differences, the other observations we made for the GPU simulations hold for the CPU cases too.

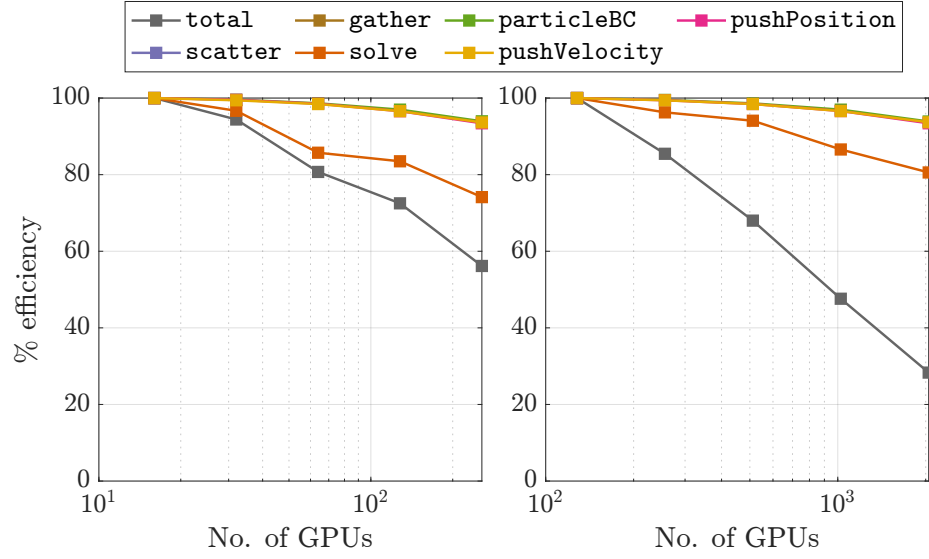
5.2.2. Strong scaling: 64 particles per cell

So far, we have considered the strong scaling study on GPUs and CPUs with 8 particles per cell for two different mesh sizes. However, many of the plasma physics codes typically use a much higher number of particles per cell to reduce the statistical noise in the simulations. Hence, to obtain representative results

³It should be noted that the statements about the scaling of the components are valid only when there is not much imbalance in the particle loads across ranks.



(a) Computation vs communication kernels for case A (left) and case B (right)



(b) Efficiency of computation kernels for case A (left) and case B (right)

Figure 7: Weak Landau damping. Strong scaling on GPUs: cross over point between computation and communication kernels and parallel efficiencies as a function of the number of GPUs for cases A and B.

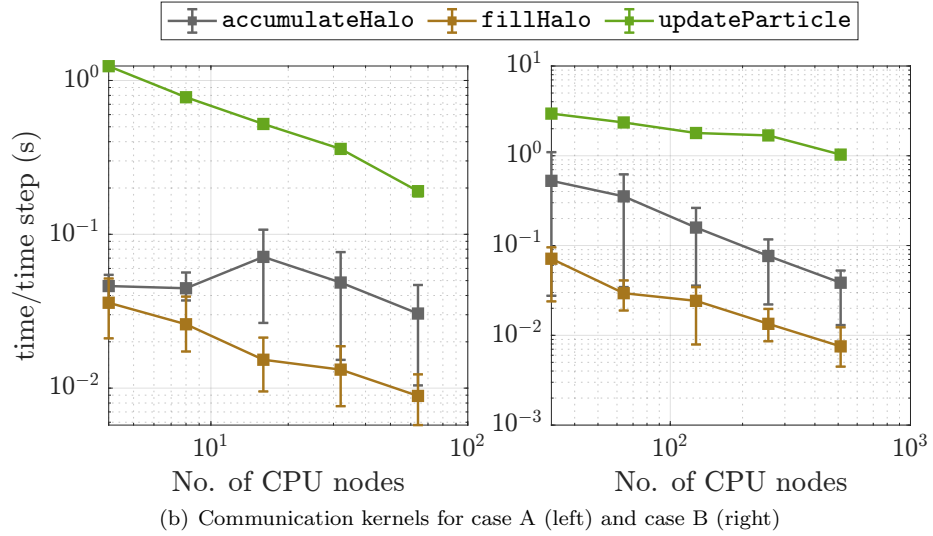
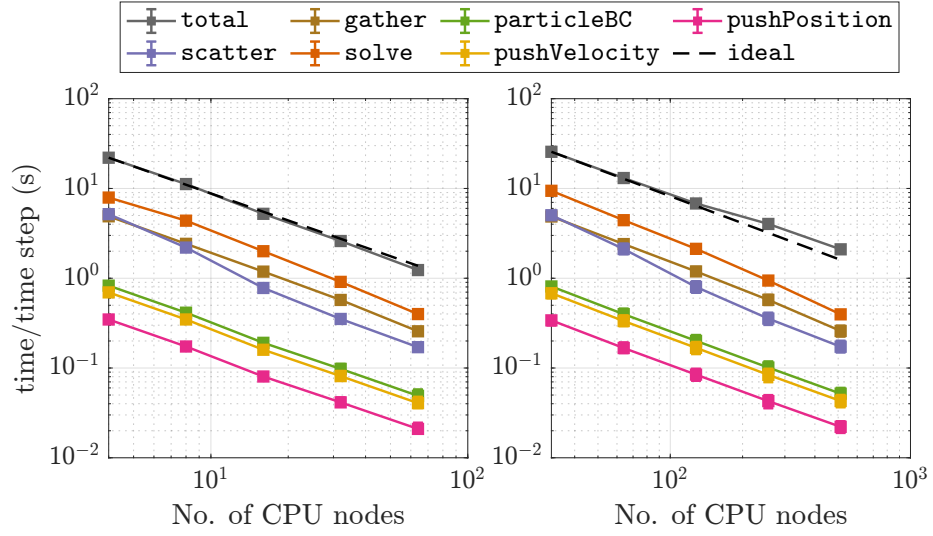


Figure 8: Weak Landau damping. Strong scaling on CPUs: scaling of computation and communication kernels for cases A and B.

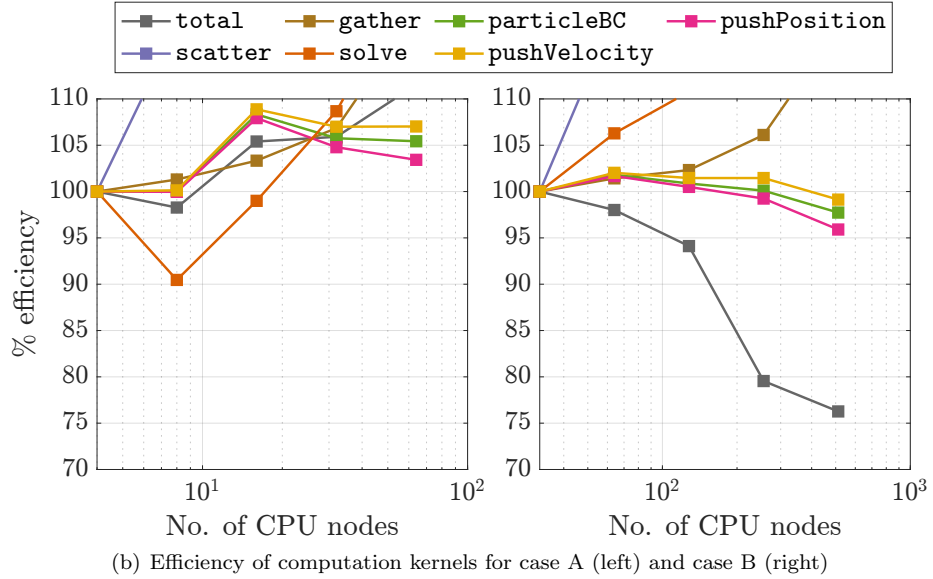
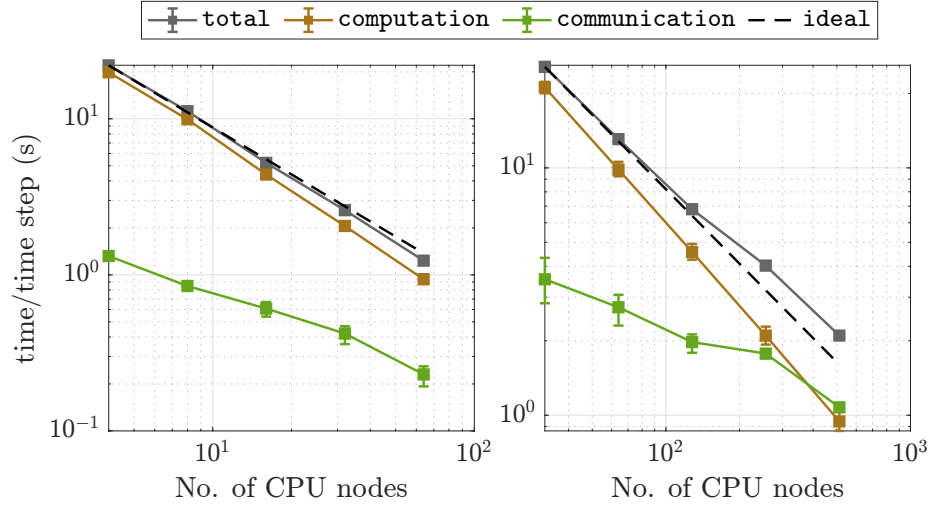
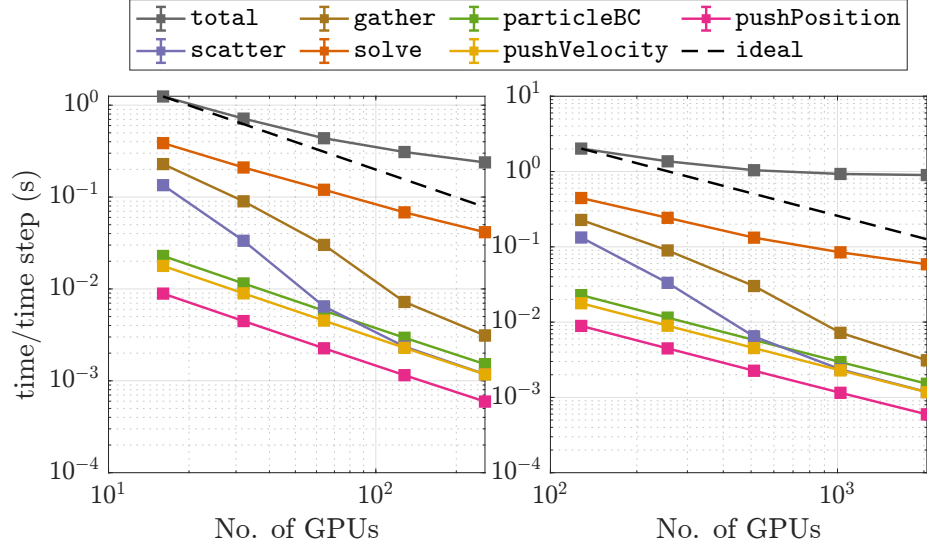
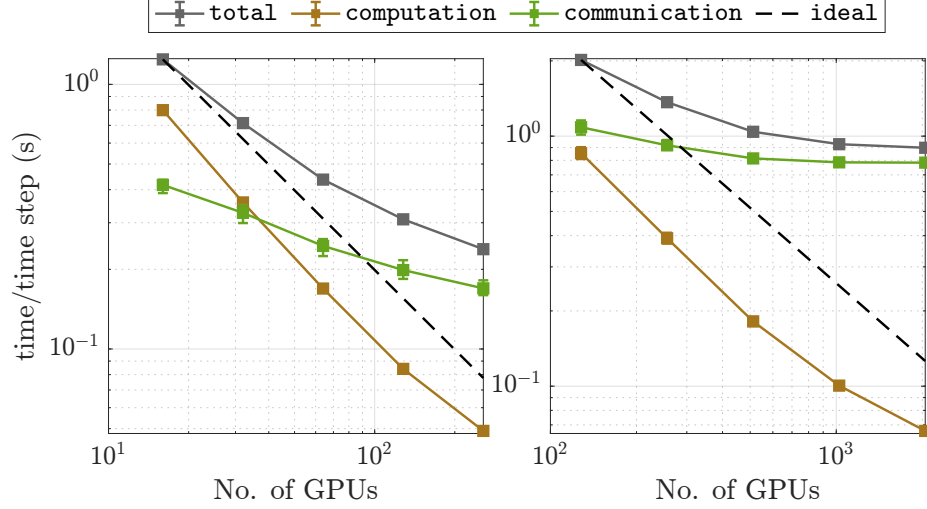


Figure 9: Weak Landau damping. Strong scaling on CPUs: cross over point between computation and communication kernels and parallel efficiencies as a function of the number of CPU nodes for cases A and B.



(a) Computation kernels for case C (left) and case D (right)



(b) Computation vs communication kernels for case C (left) and case D (right)

Figure 10: Weak Landau damping. Strong scaling study on GPUs with higher number of particles per cell ($P_c = 64$).

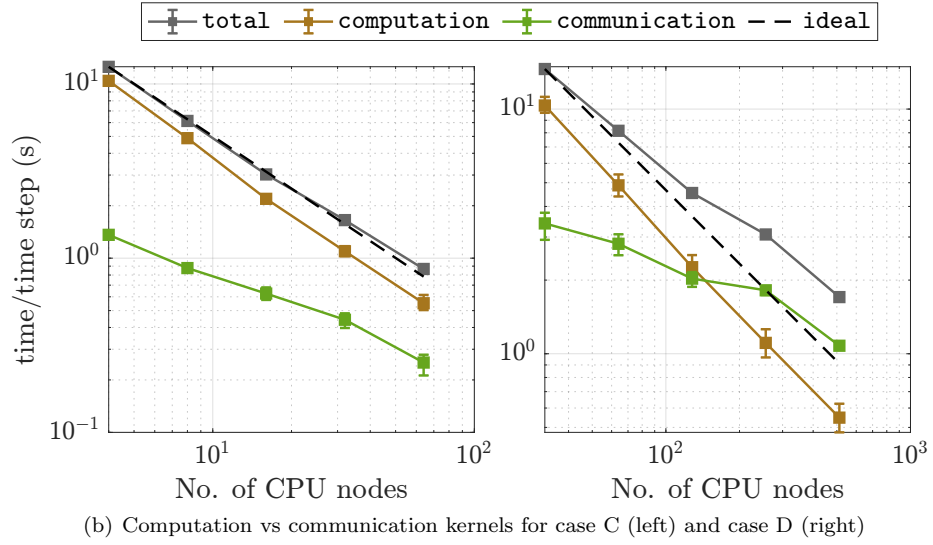
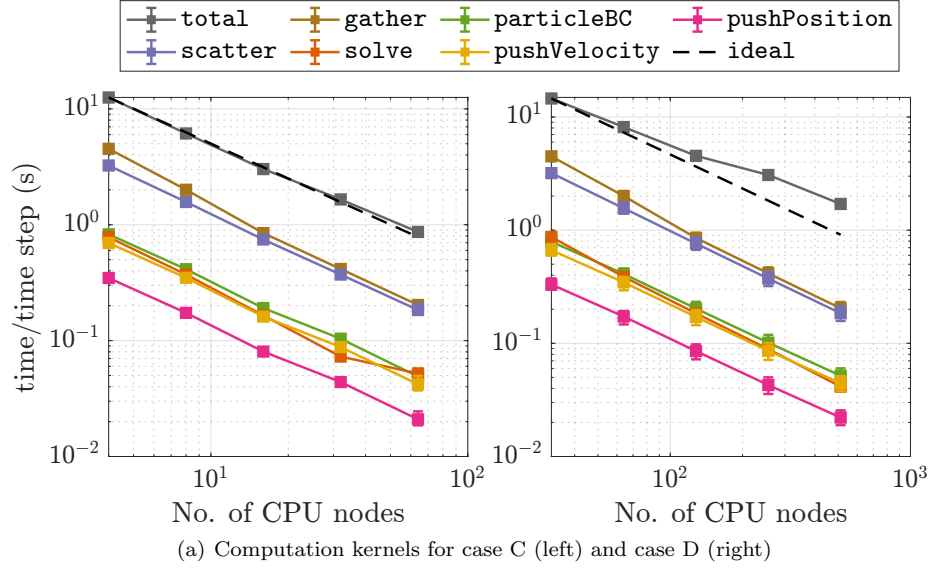


Figure 11: Weak Landau damping. Strong scaling study on CPUs with higher number of particles per cell ($P_c = 64$).

for this scenario we consider the strong scaling study in the cases C and D with 64 particles per cell and reduced grid resolutions of 256^3 and 512^3 , respectively. This is because we have an upper bound on the total number of particles that we can simulate with the considered number of GPUs and CPU nodes due to memory requirements. Hence, to have a setup with a higher number of particles per cell we keep the same total number of particles for cases C and D as in cases A and B but with reduced grid resolutions.

From Figure 10(a) we observe that the field solve is still the dominant computation kernel on GPUs, although the runtime difference between the `solve` and the other particle-based kernels is reduced compared to cases A and B due to the reduced number of grid points. We also notice that the reduced grid sizes adversely affect the scaling of the field solver, due to the smaller amount of local work compared to the communication costs in the FFTs. On the other hand, for the CPU simulations we observe from Figure 11(a) that the field solve is in fact one of the least dominant computation kernels. It is at a level comparable to the particle pushes. The scaling is also better than for the GPUs since fewer MPI ranks are involved. This is possible because CPU nodes have more memory compared to GPUs. For the GPU simulations, we indeed utilize only the device memory without transferring field or particles to the host, because this operation is very costly. This forces us to involve more GPUs.

We do not show the figures for communication costs here, because the particle communication costs are very similar to cases A and B, given that the total number of particles is the same in both sets of cases. The field communication costs are smaller than for cases A and B due to the reduced grid sizes and lower number of halo cells that need to be communicated.

Finally, from Figure 10(b) we see that the cross over point between computation and communication happens much sooner on GPUs compared to cases A and B, especially for case D which is completely communication dominated. This is because the field solver now takes less time due to the smaller numbers of grid points. As a result, the efficiency is significantly reduced compared to the other cases, especially for case D. For CPUs this effect is somewhat less pronounced, with case C looking very similar to case A. For case D the cross over point happens a bit earlier than for case B as seen from Figure 11(b). This can be understood as follows. The field solver is not the dominant kernel anymore, but the gather operation, which takes only slightly less time than the field solver in Figure 8(a), now dominates the computation kernels and has similar runtimes for both sets of cases.

We can see from Figure 10 that cases C and D show poor scaling. This is due to the following reason. Because of the small memory capacity of Piz Daint GPUs (16 GB/GPU), we have to use a lot of GPUs to have enough memory for the particles. However, we do not have enough work to utilize the computing power of these GPUs, due to the reduced field sizes. Hence, it does not make much sense to run the higher particles per cell cases like D on large numbers of GPUs (like 1024 and 2048) as we get very little efficiency out of it. That is why for the other examples in the remainder of this article, we focus only on cases A and B, as we observed very similar results for cases C and D in

the other mini-apps. In future work, we will reduce the communication costs corresponding to the particle search, by first searching over the field neighbors, and then only searching for the leftover particles among the remaining ranks. Since in explicit PIC the time step is selected such that most of the particles do not travel more than one mesh cell during a time step, this can help to reduce the particle communication cost.

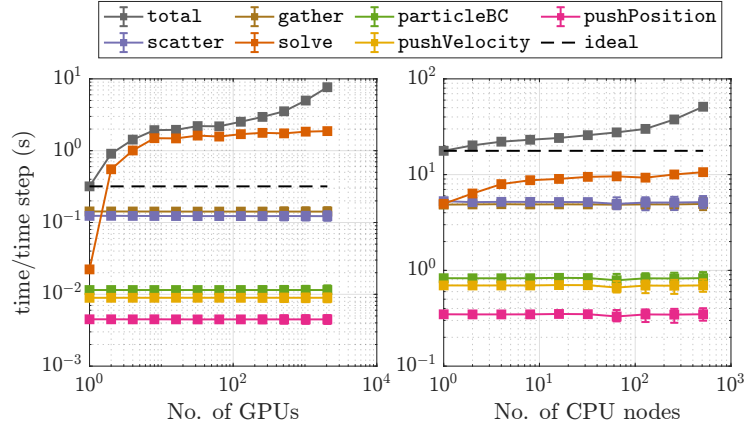
5.2.3. Weak scaling

Having studied the strong scaling so far, we will now turn our attention to the weak scaling. We first describe the weak scaling setup for the weak Landau damping test problem before discussing the results. For GPU simulations, we take the base case for 1 GPU as a 256×128^2 grid whereas for CPUs we take a 512×256^2 grid, because of the greater memory availability. For both CPUs and GPUs our simulations have 8 particles per cell. The maximum grid size and number of particles for the GPU simulations are $N_c = 2048^3$ and $N_p = 68,719,476,736$ on 2048 GPUs, whereas for CPU simulations they are $N_c = 4096 \times 2048^2$ and $N_p = 137,438,953,472$ on 512 nodes = 16,384 cores.

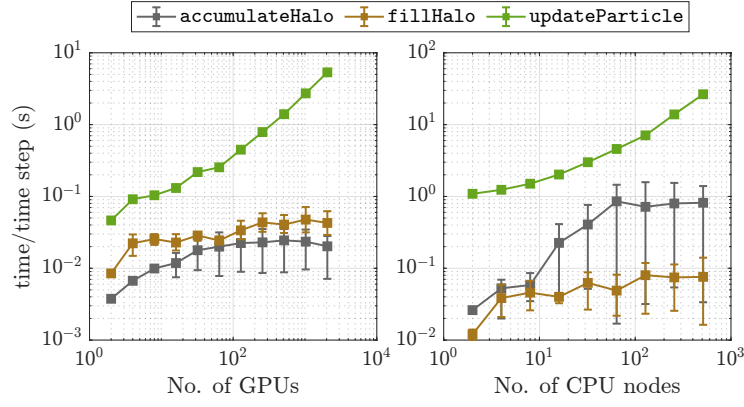
In Figure 12, all the computation kernels scale ideally, except for the field solve and total simulation, as we had already seen for the strong scaling study. The scaling of the field solver is also close to ideal starting from 8 GPUs or CPU nodes. This can be explained as follows. Starting from 8 MPI ranks, each processor has a brick of field data and the number of transposes performed during the FFT remains the same. On the other hand, for 1, 2 and 4 ranks, we have either no communication or fewer transposes due to slab or pencil decompositions. Hence they take less time than the 8 ranks case. The total time is mostly dictated by the field solve, except for the last few data points, where communication kernels dominate. We also see a linear increase asymptotically in the particle communication cost in Figure 12(b), due to the particle search and communication as explained in the strong scaling study. In Figure 13, we show the number of particles pushed per second on GPUs and CPUs as measured in this weak scaling study. This metric is obtained by dividing the total number of particles by the wall time per simulation time step. We get a maximum of approximately 10^{10} particles per second on GPUs and 3×10^9 particles per second on CPUs.

5.3. Strong Landau damping

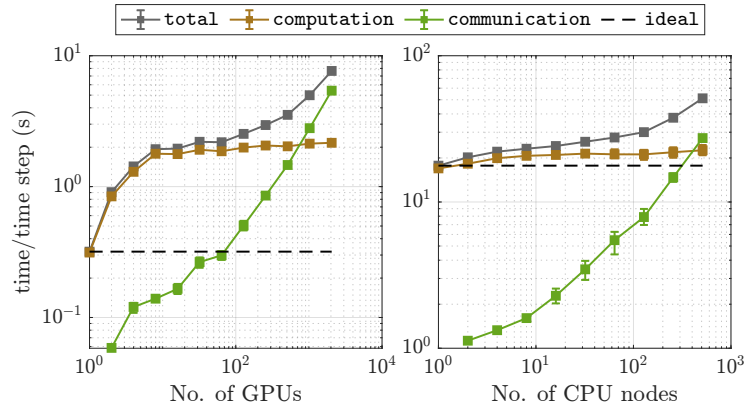
In this section, we consider the case of strong Landau damping, which corresponds to a larger perturbation parameter α in equation (5) as compared to the weak damping case. This leads to a higher particle load imbalance among the MPI ranks with equal distribution of field domains. This in turn leads to an increased total time for the simulation compared to the load balanced case, as the MPI rank or GPU which takes the maximum time for the computation and communication kernels determines the overall simulation time. An even more important problem is the memory requirement, as the GPU or CPU node which has the greatest number of particles may not have enough memory to store them



(a) Computation kernels on GPUs (left) and CPUs (right)



(b) Communication kernels on GPUs (left) and CPUs (right)



(c) Computation vs communication kernels on GPUs (left) and CPUs (right)

Figure 12: Weak Landau damping. Weak scaling of computation and communication kernels and cross over point between computation and communication on GPUs and CPUs.

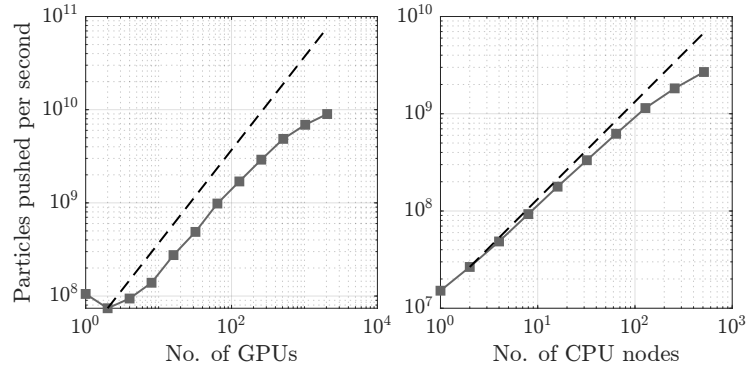


Figure 13: Weak Landau damping. Weak scaling: Number of particles pushed per second on GPUs (left) and CPUs (right).

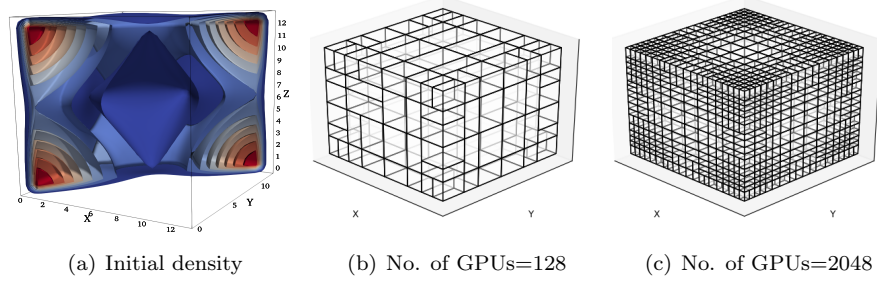


Figure 14: Strong Landau damping: Cross sectional view of the initial particle density profile showing high (red) and low (blue) densities, and particle load balanced domain decompositions for 128 and 2048 GPUs.

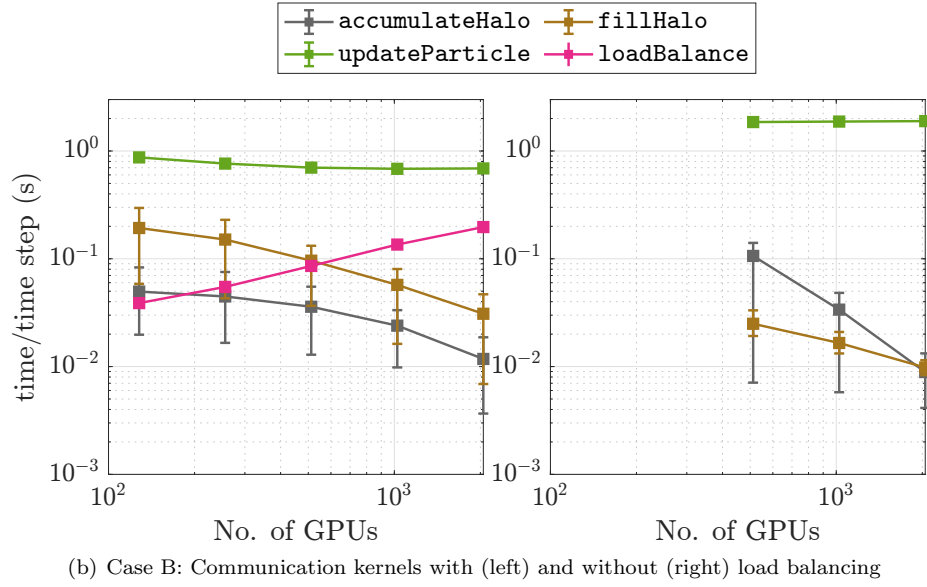
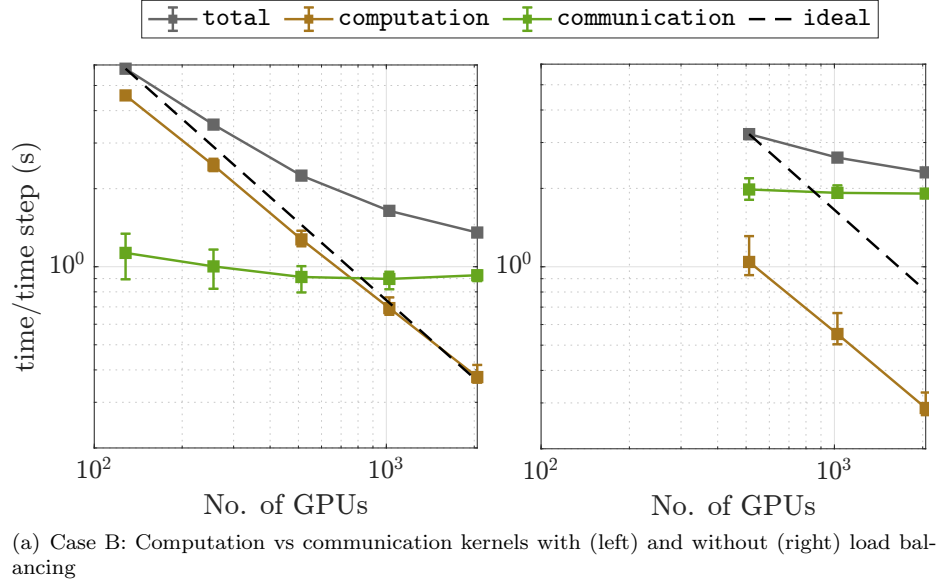
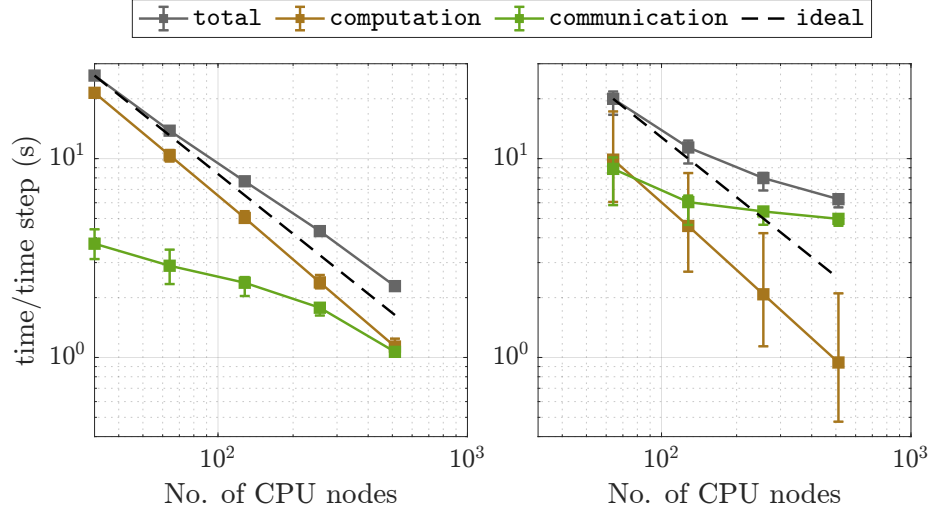
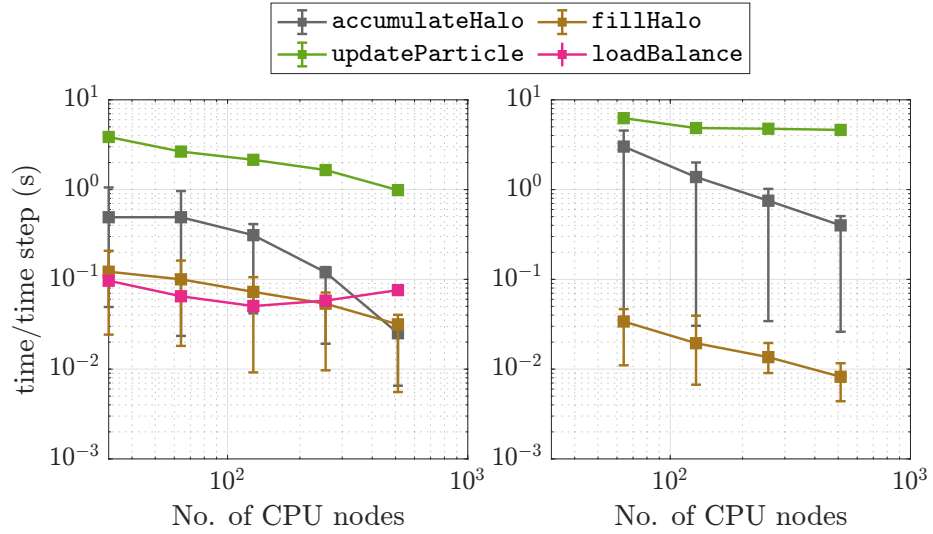


Figure 15: Strong Landau damping. Strong scaling on GPUs: comparison of scaling and performance with and without particle load balancing for case B. The right column, corresponding to the case without load balancing, is missing the first two data points because these simulations fail due to high memory imbalance leading to lack of memory.



(a) Case B: Computation vs communication kernels with (left) and without (right) load balancing



(b) Case B: Communication kernels with (left) and without (right) load balancing

Figure 16: Strong Landau damping. Strong scaling on CPUs: comparison of scaling and performance with and without particle load balancing for case B. The right column, corresponding to the case without load balancing is missing the first data point because this simulation fails due to high memory imbalance leading to lack of memory.

all, while the memory in the other GPUs or CPU nodes is under-utilized. Hence, particle load balancing is of critical importance in this example. We create the particles in a balanced way by means of orthogonal recursive bisection using the initial analytical density profile as shown in Figure 14(a). Figures 14(b) and 14(c) show the representative initial particle load balanced domain decompositions we obtain for 128 and 2048 GPUs. After the initial load balancing, we perform the balancing if the imbalance percentage given by

$$\text{imbalance\%} = \left(\frac{|n_{loc} - n_{ideal}|}{N_p} \right) \times 100$$

in any rank exceeds a given threshold. Here, n_{loc} is the local number of particles in each rank or GPU and $n_{ideal} = N_p / N_{ranks}$ is the ideal local number of particles. We set the threshold to 1% for this problem, as well as the Penning trap example in the next section, based on numerical experiments.

In Figure 15, we compare the results for case B with and without load balancing on GPUs. We do not show results for case A, as the scaling and performance of cases with and without load balancing are similar, and they are also comparable to the weak Landau damping results shown in Figures 6 and 7, due to the smaller total number of particles. We can see from Figure 15(b) that the particle communication cost is higher in the case without load balancing. This leads to an earlier cross over point between computation and communication, and therefore loss of scaling as seen from Figure 15(a). This is because without load balancing some ranks contain a lot more particles than the others, and also have to communicate more. The particle communication has a synchronization step, and since it has to wait for the last arriving rank, this leads to an increase in the `updateParticle` cost. Moreover, due to particle imbalance, the ranks which contain a lot of particles take more time in the computation kernels than the others, and this additional time also gets reflected in the synchronization steps, which in our case correspond to the `updateParticle`. We also notice from the right column of Figure 15 that the case without load balancing is missing data points for 128 and 256 GPUs as these runs fail due to lack of memory. As explained before, this is due to the need for higher memory in some GPUs, which have to store a lot of particles. We draw similar conclusions from the CPU simulations for case B in Figure 16.

Thus we can see from this example that our particle load balancing strategy is effective in cases with non-uniform particle distributions, without which the simulations either fail due to lack of memory, or have poor scaling. We should however note that the load balancing strategy comes with a price which we describe now. First, the additional overhead has to be small relative to the costs of the other kernels for it to be beneficial. This is the case in our tests shown in Figures 15 and 16 as, apart from the initial load balancing, the routine is never invoked for the set imbalance threshold of 1%. Second, our strategy of particle load balancing creates a load imbalance in the fields, as every rank now owns a brick of varied size in contrast to the uniform field distribution in the case without particle load balancing. This is visible from Figures 14(b) and 14(c). Although this does not affect the field solver scaling much in the

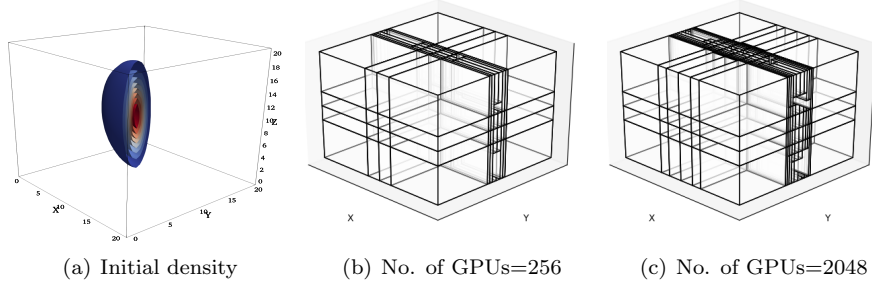


Figure 17: Penning trap: Cross sectional view of the initial particle density profile showing high (red) and low (blue) densities, and particle load balanced domain decompositions for 256 and 2048 GPUs.

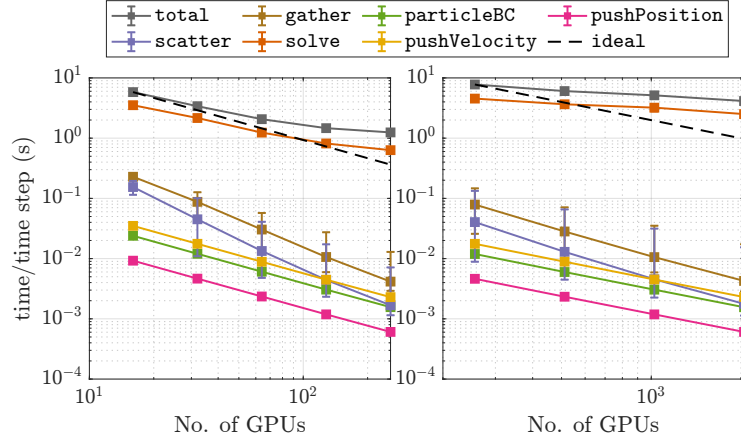
case of strong Landau damping, for problems with highly non-uniform particle distributions, such as our Penning trap simulations, it can have a significant impact, as will be explained in Section 5.4.

The weak scaling results for the strong Landau damping simulations with load balancing are very similar to those of the weak Landau damping simulations shown in Figure 12. We therefore do not show them, to avoid repetition.

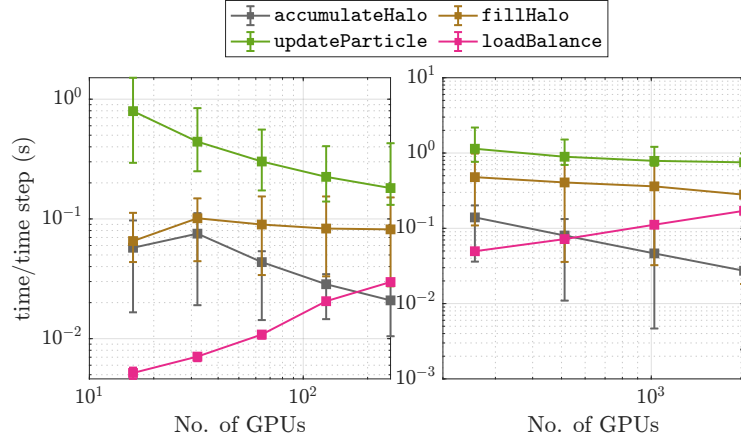
5.4. Electron dynamics in a charge neutral Penning trap

In this section, we present the strong and weak scaling results for the Penning trap mini-app. In Figure 17(a), the initial electron density is shown, which is a Gaussian bunch. In contrast to the Landau damping examples, this highly non-uniform particle distribution leads to a highly non-uniform field distribution after particle load balancing, as shown in Figures 17(b) and 17(c). In that respect the Penning trap mini-app is a harder test case in terms of field and particle load balancing compared to the Landau damping examples.

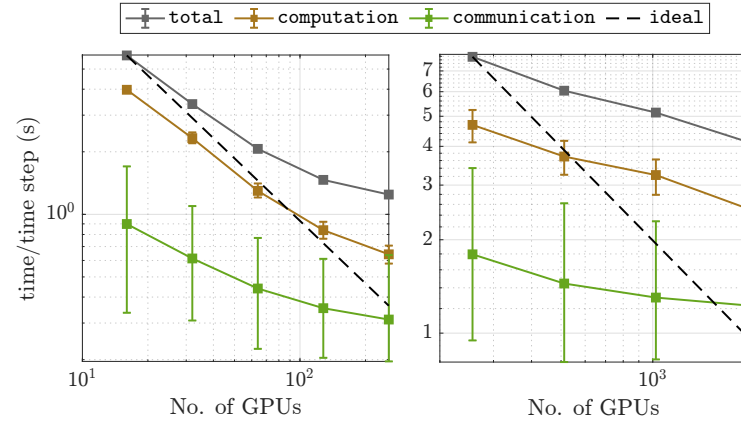
In Figure 18, we show the strong scaling results for cases A and B with load balancing on GPUs. For this example, we are unable to run most of the simulations without load balancing, due to memory requirements. Even in the few cases for which we are able to run the simulations, we did not observe scaling. We therefore only present results with load balancing. We observe from Figure 18(a) that the field solver takes more time and that the scaling is much worse than for the Landau damping cases. This is more significant for case B due to the large level of load imbalance in the fields, as shown in Figures 17(b) and 17(c). We also observe that even with particle load balancing, we are unable to run our simulation for case B with 128 GPUs due to insufficient memory for particle communication operations during time stepping. These observations clearly show that for highly non-uniform particle distributions, as in this Penning trap example, our current load balancing strategy has to be improved in order to effectively carry out these simulations with large numbers of particles on large numbers of nodes. We will investigate this in future work.



(a) Computation kernels for case A (left) and case B (right)



(b) Communication kernels for case A (left) and case B (right)



(c) Computation vs communication kernels for case A (left) and case B (right)

Figure 18: Penning trap. Strong scaling on GPUs: scaling of computation kernels, communication kernels, and cross over point between them for cases A and B with load balancing.

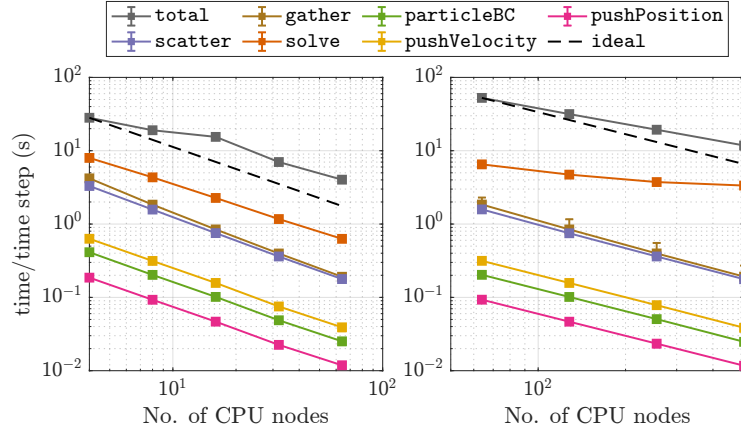
We notice from Figure 18(b) that even with particle load balancing, we have significant variation in timings across the GPUs for the `updateParticle`, which we did not observe in the strong Landau damping simulations. This again is a symptom of the highly non-uniform electron distribution in the Penning trap simulations.

In Figure 19, we show the strong scaling study on CPUs. Because of the higher memory per node, we are able to run case A without load balancing, whereas for case B we still require load balancing. We observed loss of scaling and increased communication costs for case A without load balancing. Since the results are very similar to the results we obtained for the strong Landau damping simulations, we do not show them in this article.

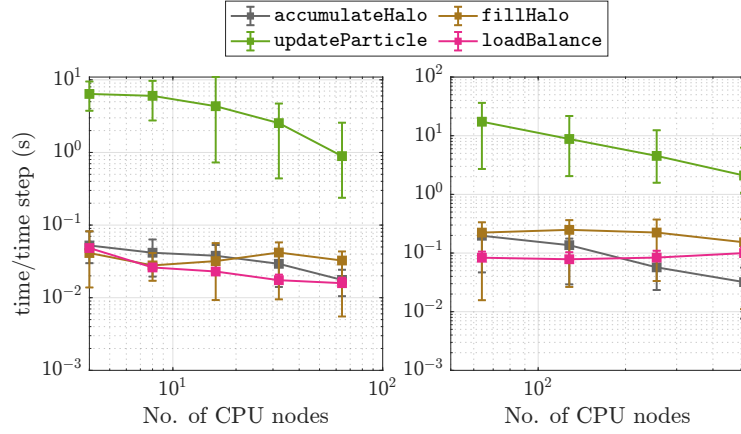
For case B, from the right column of Figure 19(a) we see that the field solver efficiency is lower than in previous examples, as it is for the GPU simulations. Furthermore, even with load balancing, we are unable to run it on 32 nodes, due to the high memory requirements. Comparing the right column of Figure 19(b) and the left column of Figure 16(b), we see that the `updateParticle` cost is higher in the Penning trap simulations than in the strong Landau damping simulations. The decrease of the time per time step as a function of the number of CPU nodes suggests that the particle sending and receiving operations dominate the `updateParticle` in this case, as explained in Section 5.2.

Finally, we consider the weak scaling study for the Penning trap test case. For this, we take the same base case of a 256×128^2 grid per GPU as we had for the Landau damping tests, whereas for CPUs we take a 256^3 grid because of memory requirements imposed by field imbalance as well as field and particle communications. Similar to Landau damping for both CPUs and GPUs we take 8 particles per cell. The runs on 1024 and 2048 GPUs fail due to memory requirements. Hence, the maximum grid size and number of particles for the GPU simulations are $N_c = 2048 \times 1024^2$ and $N_p \approx 1.6 \times 10^{10}$ on 512 GPUs, whereas for CPU simulations they are $N_c = 2048^3$ and $N_p \approx 6.9 \times 10^{10}$ on 512 nodes = 16,384 cores.

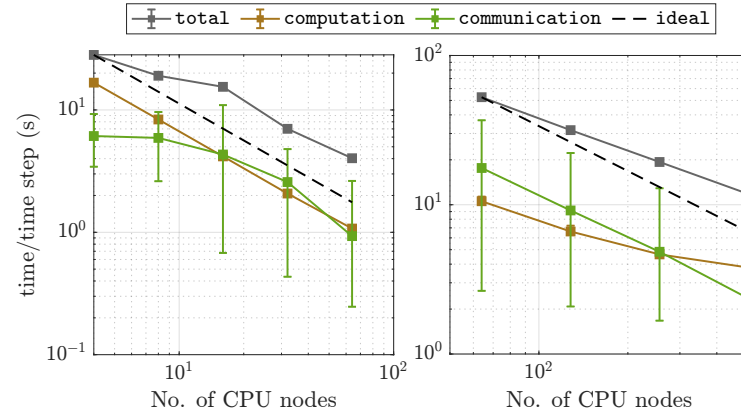
In Figure 20, we show the weak scaling results on GPUs and CPUs for the Penning trap case. As noticed in the strong scaling study, the field solver scaling is affected due to the imbalance in the fields on both GPUs and CPUs. This leads to higher total time than the Landau damping problem, for which the results were shown in Figure 12. It should also be noted that the increase in time on CPUs is despite having half as many particles and grid points than the Landau damping case. We also see from the right column of Figure 20(c) that the communication costs become comparable to the computation costs much earlier than we observed for the Landau damping simulations in Figure 12(c). Thus in cases in which the particle density has strong non-uniformity, such as this Penning trap test case, the field and particle load balancing have to be carefully investigated in order to determine the setup that would lead to the least overall time. This requires improvements in the current load balancing strategy and also partly depends on the number of particles, grid points and MPI ranks or GPUs used in the simulation.



(a) Computation kernels for case A (left) and case B (right)

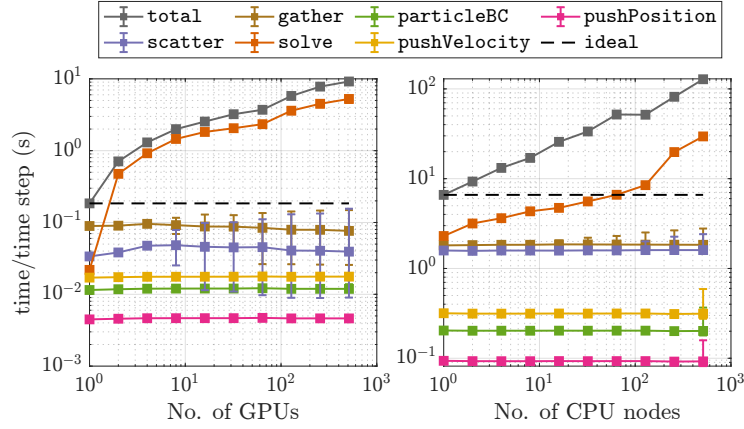


(b) Communication kernels for case A (left) and case B (right)

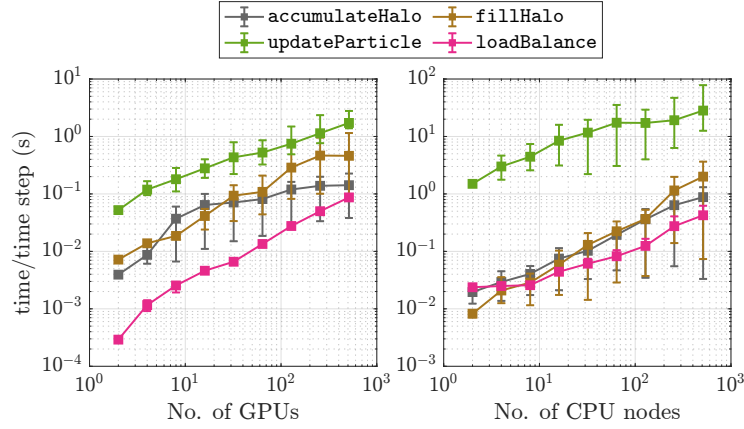


(c) Computation vs communication kernels for case A (left) and case B (right)

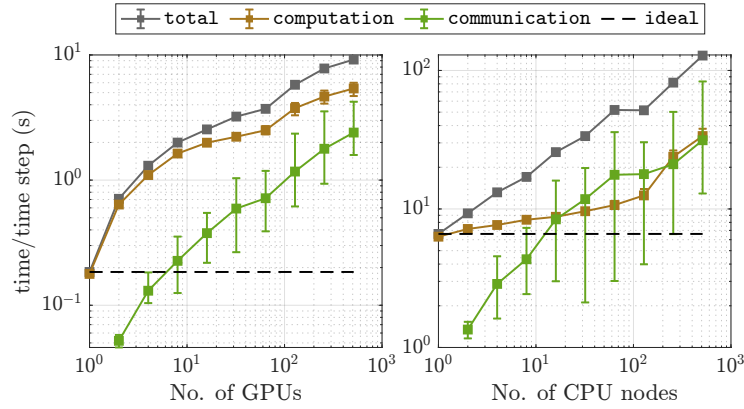
Figure 19: Penning trap. Strong scaling on CPUs: scaling of computation kernels, communication kernels, and cross over point between them for case A and case B with load balancing.



(a) Computation kernels on GPUs (left) and CPUs (right)



(b) Communication kernels on GPUs (left) and CPUs (right)



(c) Computation vs communication kernels on GPUs (left) and CPUs (right)

Figure 20: Penning trap. Weak scaling of computation kernels, communication kernels, and cross over point between them for GPUs and CPUs.

5.5. Performance comparison across different architectures

Until now we have considered the performance of different mini-apps on Piz Daint CPU and GPU partitions. In this section, we consider the weak Landau damping mini-app and evaluate its performance across different architectures. For this purpose we consider the CPU and GPU architectures in Table 4.

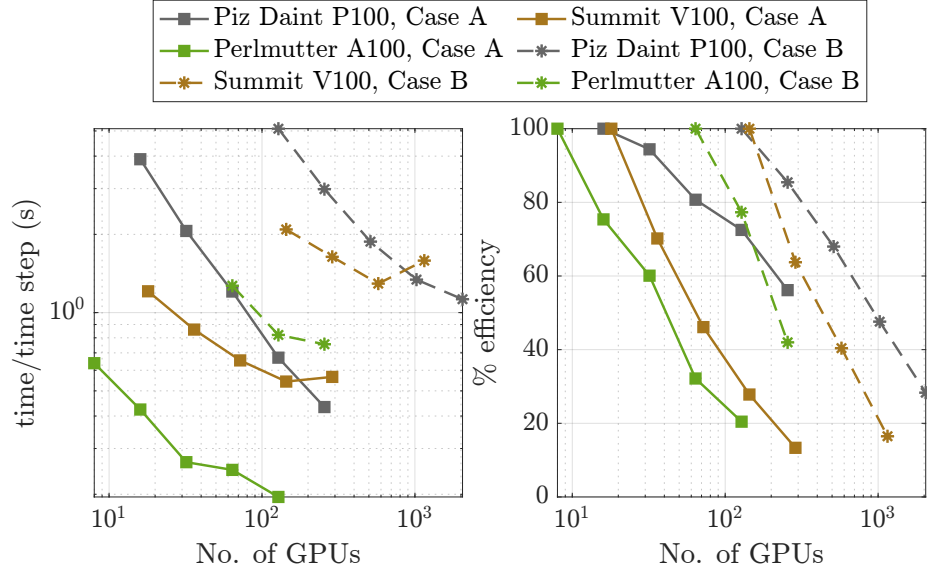
GPU architectures	CPU architectures
Piz Daint P100	Piz Daint (Intel Xeon E5-2695 v4 @ 2.1GHz)
Summit V100	Cori KNL (Intel Xeon Phi 7250 @ 1.4 GHz)
Perlmutter A100	

Table 4: List of CPU and GPU partitions used for the performance comparison across different architectures.

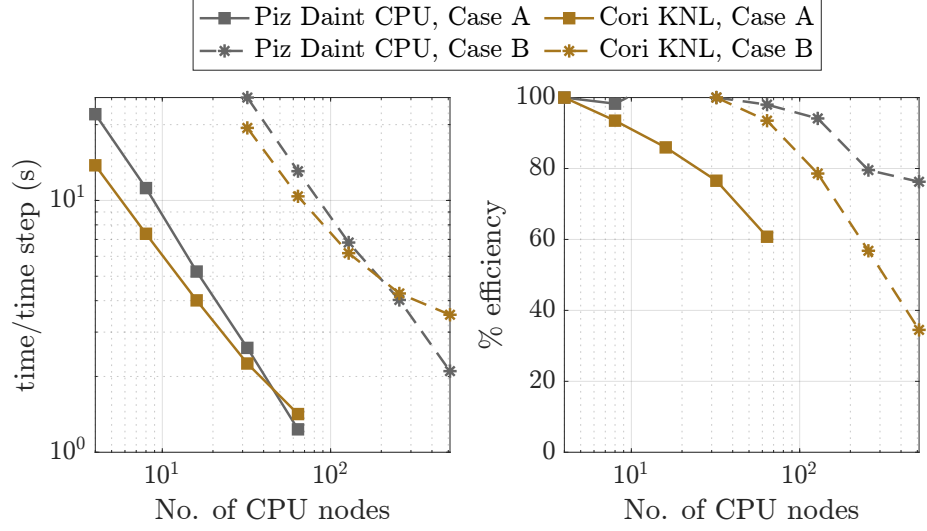
Each node of the Piz Daint GPU partition consists of 1 Tesla P100 GPU, whereas the Summit and Perlmutter nodes consist of 6 Tesla V100 GPUs and 4 Ampere A100 GPUs, respectively. The memory per GPU is 16 GB for both P100 and V100 GPUs, whereas for A100 it is 40 GB. In the CPU partitions, each node of the Piz Daint CPU partition consists of 36 cores, whereas the Cori KNL nodes have 68 cores each. The interconnect configuration is Aries and the network topology is Dragonfly for Piz Daint CPU, Piz Daint GPU and Cori KNL architectures. For Perlmutter, the interconnect configuration is HPE Cray Slingshot with a three-hop Dragonfly topology whereas for Summit, it is Mellanox EDR 100G InfiniBand with a Non-blocking Fat Tree topology.

We utilize all the GPUs per node in the GPU architectures, whereas for the CPU ones we use 32 and 64 cores per node on the Piz Daint and Cori systems, respectively. The multithreading is turned off for the CPU runs. Similar to the previous study on Piz Daint, we use 1 MPI rank per GPU for the GPU architectures, whereas for the CPU ones we use 1 MPI rank per node and use 32 and 64 OpenMP threads on Piz Daint and Cori respectively.

In Figure 21(a) the total time and efficiencies are compared across the three GPU architectures for the strong scaling study corresponding to cases A and B. In the case of Perlmutter, we can start the scaling study with half as many GPUs as the Piz Daint and Summit partitions, thanks to the higher memory configuration of A100 GPUs (40 GB) compared to those of P100 and V100 GPUs (16 GB). We make the following observations. In terms of the absolute wall time per simulation time step, Perlmutter is the fastest, followed by Summit and then Piz Daint. In terms of scaling efficiencies, Piz Daint has the highest efficiency followed by Perlmutter and then Summit. We can understand the total time and efficiencies better by looking at the cumulative computation and communication kernels across the three architectures for cases A and B in Figure 22(a). The computation kernels of Summit have a speedup of 3-4 compared to Piz Daint until the scaling stops, whereas Perlmutter computation kernels have a speedup of roughly 10 compared to Piz Daint. This is because of the architecture and CUDA compute capability of the GPUs: the latest A100 GPUs are more powerful than the V100, which in turn are more powerful than



(a) Total time (left) and efficiency (right) on different GPU architectures



(b) Total time (left) and efficiency (right) on different CPU architectures

Figure 21: Comparison of different GPU architectures (top row) and CPU architectures (bottom row) with respect to strong scaling for the weak Landau damping mini-app.

the P100 GPUs. In terms of communication costs, Perlmutter has the least communication cost for both cases A and B. Comparing Summit and Piz Daint, for case A Summit has a higher communication cost than Piz Daint, whereas for case B both are comparable. The higher communication cost of Summit is also reflected in the scaling of the computation kernels, where the field solver loses scaling earlier than Piz Daint and Perlmutter.

We now consider the CPU architectures, and compare the strong scaling performance on Piz Daint and Cori KNL nodes. In Figure 21(b) we show the total time and efficiency for these two systems. For small number of nodes, the wall time per simulation time step of Cori nodes are slightly less than for Piz Daint nodes. However, the better scaling and efficiency of Piz Daint eventually leads to a lower runtime than Cori. From the computation and communication kernels in Figure 22(b) we notice that the computation kernels perform slightly better on Cori than on Piz Daint, but the communication time of Cori is higher, which eventually leads to a lower efficiency than Piz Daint. However, overall the performance per number of nodes of the two systems are very similar.

6. Conclusion

In this work, we performed a scaling and performance portability study of the particle-in-cell scheme for plasma physics applications by means of a set of mini-apps, namely “Alpine”, targeting exascale architectures. The mini-apps include weak and strong Landau damping, the dynamics of an electron bunch in a quasi-neutral Penning trap, and the two-stream and bump-on-tail instabilities, which are commonly used as benchmarks for electrostatic PIC studies. Our scaling and performance analysis shows that the weak Landau damping simulations perform the best among the mini-apps in terms of scalability and time to solution. This is because the particle distribution is relatively uniform in this case, and the particle communication timings therefore remain small compared to the timings of computation kernels. We obtained a maximum particle push rate of approximately 10^{10} particles per second on GPUs and 3×10^9 particles per second on CPUs in the weak scaling study. The scaling and performance of strong Landau damping simulations with particle load balancing are very similar to those of the weak Landau damping simulations, whereas without particle load balancing the particle communication costs as well as memory imbalance are much higher, which leads to poor scaling. The Penning trap mini-app corresponds to the toughest test case, due to the highly non-uniform particle distribution. In this case, particle load balancing leads to a significant imbalance in the fields, which then affects the scaling of the field solve. Our current load balancing strategy needs to be improved to handle such cases.

A performance comparison across different GPU and CPU architectures for the weak Landau damping mini-app shows that Perlmutter with the latest NVIDIA A100 GPUs performs the best in terms of wall time per simulation time step, with almost an order of magnitude speedup compared to the Piz Daint P100 GPUs, and approximately a three times speedup as compared to Summit with V100 GPUs. In the comparison of CPU architectures, the wall

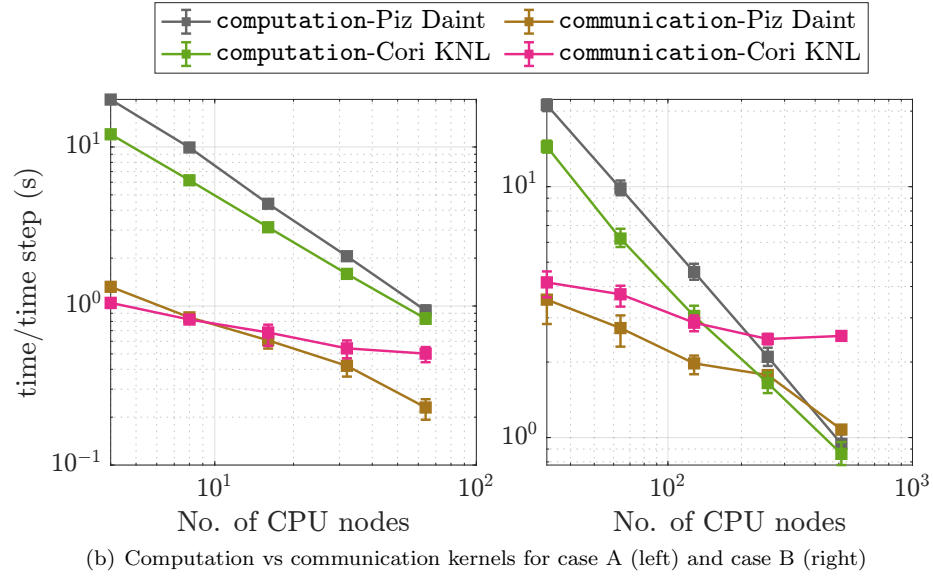
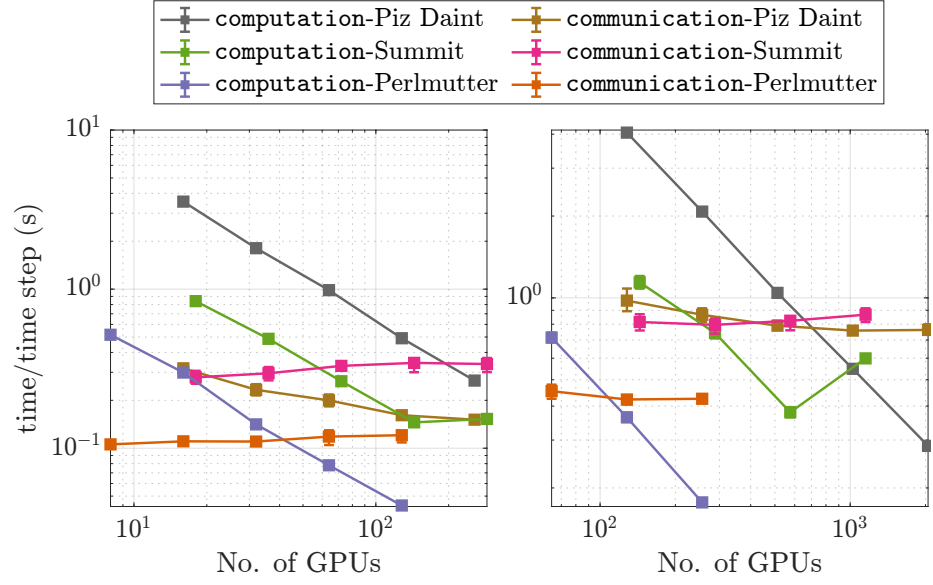


Figure 22: Timings of computation and communication kernels for different GPU architectures (top row) and CPU architectures (bottom row) with respect to strong scaling for the weak Landau damping mini-app.

time per simulation time step of Piz Daint and Cori KNL nodes are very similar, with Piz Daint showing better scaling than Cori.

In terms of future work, we will optimize and improve the current load balancing strategy and particle communication. We will continue our benchmarking studies with Perlmutter and other upcoming architectures, to test for higher numbers of particles, grid points and GPUs and CPU cores. Finally, we would also like to extend our current study by adding more mini-apps to the Alpine collection, with test problems requiring the inclusion of collisions and the implementation of an electromagnetic PIC scheme.

Availability

Alpine and IPPL are open source projects. The sources can be downloaded from <https://gitlab.psi.ch/OPAL/Libraries/ippl>.

Acknowledgments

The authors would like to thank the Kokkos team for helping us with all the queries during the development of IPPL and Alpine. We would also like to thank Marc Caubet Serrabou from PSI for his help with all the installations during the development of IPPL. This project has received funding from the European Union’s Horizon 2020 research and innovation program under the Marie Skłodowska-Curie grant agreement No. 701647 and from the United States National Science Foundation under Grant No. PHY-1820852. This research used resources of the National Energy Research Scientific Computing Center (NERSC), a U.S. Department of Energy Office of Science User Facility located at Lawrence Berkeley National Laboratory, operated under Contract No. DE-AC02-05CH11231 using NERSC award ASCR-ERCAPM888. We acknowledge access to Piz Daint at the Swiss National Supercomputing Centre, Switzerland under the PSI’s share with the project IDs psi07 and psigpu. Finally, this research also used resources of the Oak Ridge Leadership Computing Facility at the Oak Ridge National Laboratory, which is supported by the Office of Science of the U.S. Department of Energy under Contract No. DE-AC05-00OR22725.

Appendix A. Compilers and libraries used for the benchmarks on different computing architectures

We used Kokkos version 3.5.0 and heFFTe version 2.2.0 for all our simulations. With respect to IPPL, we used the tag `Scaling_study_for_Alpine_paper` for all the scaling studies. It can be obtained from the IPPL gitlab repository <https://gitlab.psi.ch/OPAL/Libraries/ippl>.

For each of the architectures, the compiler type, version and the MPI used are specified in Table 5 below.

For the Piz Daint CPUs, we also conducted the benchmarking study with the following Intel compiler.

Architecture	Compiler	MPI
Piz Daint GPU	gcc/9.3.0	OpenMPI/4.1.2 with CUDA/11.2
Piz Daint CPU	gcc/11.2.0	cray-mpich/7.7.18
Cori KNL	intel/19.1.1.217	cray-mpich/7.7.18
Summit GPU	gcc/9.1.0	spectrum-mpi/10.4.0 with CUDA/11.0
Perlmutter GPU	gcc/11.2.0	OpenMPI/4.1.2 with CUDA/11.4

Table 5: Compiler type, version and MPI used for each of the architectures used for the scaling study in Section 5.

- intel/2021.3.0
- cray-mpich/7.7.18

The results were very comparable to those obtained with gcc. However, we sometimes observed an inconsistent “Bus error”, the reason of which is yet unknown and under investigation.

References

- [1] R. W. Hockney, J. W. Eastwood, Computer simulation using particles, CRC Press, 1988.
- [2] C. K. Birdsall, A. B. Langdon, Plasma physics via computer simulation, CRC press, 2004.
- [3] J. M. Dawson, Particle simulation of plasmas, Reviews of modern physics 55 (2) (1983) 403.
- [4] L. F. Ricketson, A. J. Cerfon, Sparse grid techniques for particle-in-cell schemes, Plasma Physics and Controlled Fusion 59 (2) (2016) 024002.
- [5] A. Spitkovsky, Simulations of relativistic collisionless shocks: shock structure and particle acceleration, in: AIP Conference Proceedings, Vol. 801, American Institute of Physics, 2005, pp. 345–350.
- [6] O. Buneman, Computer space plasma physics, simulation techniques and softwares, ed, H. Matsumoto and Y. Omura (Terra Scientific, Tokyo, 1993) p 67.
- [7] S. Jolliet, A. Bottino, P. Angelino, R. Hatzky, T.-M. Tran, B. Mcmillan, O. Sauter, K. Appert, Y. Idomura, L. Villard, A global collisionless PIC code in magnetic coordinates, Computer Physics Communications 177 (5) (2007) 409–425.
- [8] C.-S. Chang, S. Ku, Spontaneous rotation sources in a quiescent tokamak edge plasma, Physics of Plasmas 15 (6) (2008) 062510.

- [9] R. A. Fonseca, L. O. Silva, F. S. Tsung, V. K. Decyk, W. Lu, C. Ren, W. B. Mori, S. Deng, S. Lee, T. Katsouleas, et al., Osiris: A three-dimensional, fully relativistic particle in cell code for modeling plasma based accelerators, in: International Conference on Computational Science, Springer, 2002, pp. 342–351.
- [10] J. Qiang, S. Lidia, R. D. Ryne, C. Limborg-Deprey, Three-dimensional quasistatic model for high brightness beam dynamics simulation, *Physical Review Special Topics-Accelerators and Beams* 9 (4) (2006) 044204.
- [11] A. Adelmann, P. Calvo, M. Frey, A. Gsell, U. Locans, C. Metzger-Kraus, N. Neveu, C. Rogers, S. Russell, S. Sheehy, J. Snuvernik, D. Winklehner, OPAL a versatile tool for charged particle accelerator simulations, arXiv preprint arXiv:1905.06654.
- [12] J.-L. Vay, A. Almgren, J. Bell, L. Ge, D. Grote, M. Hogan, O. Kononenko, R. Lehe, A. Myers, C. Ng, et al., Warp-x: A new exascale computing platform for beam-plasma simulations, *Nuclear Instruments and Methods in Physics Research Section A: Accelerators, Spectrometers, Detectors and Associated Equipment* 909 (2018) 476–479.
- [13] S. M. Mniszewski, J. Belak, J.-L. Fattebert, C. F. Negre, S. R. Slattery, A. A. Adedoyin, R. F. Bird, C. Chang, G. Chen, S. Ethier, et al., Enabling particle applications for exascale computing platforms, *The International Journal of High Performance Computing Applications* 35 (6) (2021) 572–597.
- [14] A. Myers, A. Almgren, L. Amorim, J. Bell, L. Fedeli, L. Ge, K. Gott, D. P. Grote, M. Hogan, A. Huebl, et al., Porting warpx to gpu-accelerated platforms, *Parallel Computing* 108 (2021) 102833.
- [15] M. A. Heroux, D. W. Doerfler, P. S. Crozier, J. M. Willenbring, H. C. Edwards, A. Williams, M. Rajan, E. R. Keiter, H. K. Thornquist, R. W. Numrich, Improving performance via mini-applications, Sandia National Laboratories, Tech. Rep. SAND2009-5574 3.
- [16] R. Bird, N. Tan, S. V. Luedtke, S. L. Harrell, M. Taufer, B. Albright, Vpic 2.0: Next generation particle-in-cell simulations, *IEEE Transactions on Parallel and Distributed Systems* 33 (4) (2021) 952–963.
- [17] J. V. W. Reynders, J. Cummings, P. F. Dubois, The POOMA Framework, *Computers in Physics* 12 (5) (1998) 453–459.
- [18] T. Veldhuizen, Expression templates, *C++ Report* 7 (5) (1995) 26–31.
- [19] H. C. Edwards, C. R. Trott, D. Sunderland, Kokkos: Enabling many-core performance portability through polymorphic memory access patterns, *Journal of Parallel and Distributed Computing* 74 (12) (2014) 3202 – 3216, domain-Specific Languages and High-Level Frameworks for High-Performance Computing.

- [20] C. R. Trott, D. Lebrun-Grandie, D. Arndt, J. Ciesko, V. Dang, N. Ellingwood, R. Gayatri, E. Harvey, D. S. Hollman, D. Ibanez, et al., Kokkos 3: Programming model extensions for the exascale era, *IEEE Transactions on Parallel and Distributed Systems* 33 (4) (2021) 805–817.
- [21] D. J. Quinlan, M. Berndt, MLB: Multilevel load balancing for structured grid applications, Tech. rep., Los Alamos National Lab.(LANL), Los Alamos, NM (United States) (1997).
- [22] A. Ayala, S. Tomov, A. Haidar, J. Dongarra, hefft: Highly efficient fft for exascale, in: *International Conference on Computational Science*, Springer, 2020, pp. 262–275.
- [23] D. Rodríguez-Patiño, S. Ramírez, J. Salcedo-Gallo, J. Hoyos, E. Restrepo-Parra, Implementation of the two-dimensional electrostatic particle-in-cell method, *American Journal of Physics* 88 (2) (2020) 159–167.
- [24] L. Devroye, Nonuniform random variate generation, *Handbooks in operations research and management science* 13 (2006) 83–121.
- [25] K. Tretiak, D. Ruprecht, An arbitrary order time-stepping algorithm for tracking particles in inhomogeneous magnetic fields, *Journal of Computational Physics: X* 4 (2019) 100036.
- [26] A. Ho, I. A. M. Datta, U. Shumlak, Physics-based-adaptive plasma model for high-fidelity numerical simulations, *Frontiers in Physics* 6 (2018) 105.
- [27] A. Myers, P. Colella, B. V. Straalen, A 4th-order particle-in-cell method with phase-space remapping for the Vlasov–Poisson equation, *SIAM Journal on Scientific Computing* 39 (3) (2017) B467–B485.
- [28] K. Kormann, A semi-lagrangian vlasov solver in tensor train format, *SIAM Journal on Scientific Computing* 37 (4) (2015) B613–B632.
- [29] G. Chen, L. Chacón, D. C. Barnes, An energy-and charge-conserving, implicit, electrostatic particle-in-cell algorithm, *Journal of Computational Physics* 230 (18) (2011) 7018–7036.
- [30] S. Sarkar, S. Paul, R. Denra, Bump-on-tail instability in space plasmas, *Physics of Plasmas* 22 (10) (2015) 102109.
- [31] S. Muralikrishnan, A. J. Cerfon, M. Frey, L. F. Ricketson, A. Adelman, Sparse grid-based adaptive noise reduction strategy for particle-in-cell schemes, *Journal of Computational Physics: X* (2021) 100094.
- [32] S. Adam, Space charge effects in cyclotrons-from simulations to insights, in: *Proc. of the 14th Int. Conf. on Cyclotrons and their Applications*, (World Scientific, Singapore, 1996), Vol. 446, 1995.

- [33] J. Yang, A. Adelman, M. Humbel, M. Seidel, T. Zhang, et al., Beam dynamics in high intensity cyclotrons including neighboring bunch effects: Model, implementation, and application, *Physical Review Special Topics-Accelerators and Beams* 13 (6) (2010) 064201.