

Multichannel Analysis of Surface Waves Accelerated (MASWAccelerated): Software for efficient surface wave inversion using MPI and GPUs

Joseph Kump^{*}, Eileen R. Martin¹

Virginia Polytechnic Institute and State University, Department of Mathematics, McBryde Hall, 225 Stanger St., Blacksburg, VA, 24060, USA

ARTICLE INFO

Keywords:

Surface waves
Near-surface imaging
GPU
Parallel computing
MPI

ABSTRACT

Multichannel Analysis of Surface Waves (MASW) is a technique frequently used in geotechnical engineering and engineering geophysics to infer 1D layered models of seismic shear wave velocities in the top tens to hundreds of meters of the subsurface. We aim to accelerate MASW calculations by capitalizing on modern computer hardware available in the workstations of most engineers: multiple cores and graphics processing units (GPUs). We propose new parallel and GPU accelerated algorithms for computing 1D MASW inversion, and provide software implementations in C using Message Passing Interface (MPI) and CUDA. These algorithms take advantage of sparsity that arises in the problem, and the work balance between processes considers typical data trends. We compare our methods to an existing open source Matlab MASW tool. Our serial C implementation achieves a 2x speedup over the Matlab software, and we continue to see improvements by parallelizing the problem with MPI. We see nearly perfect strong and weak scaling for uniform data, and improve strong scaling for realistic data by repartitioning the problem to process mapping. By utilizing GPUs available on most modern workstations, we observe an additional 1.3x speedup over the serial C implementation on the first use of the method. We typically repeatedly evaluate theoretical dispersion curves as part of an optimization procedure, and on the GPU the kernel can be cached for faster reuse on later runs. We observe a 3.2x speedup on the cached GPU runs compared to the serial C runs. This work is the first open-source parallel or GPU-accelerated software tool for MASW imaging, and should enable geotechnical engineers to fully utilize all computer hardware at their disposal.

1. Introduction and background

Multichannel Analysis of Surface Waves (MASW) is a seismic exploration technique used to infer a layered 1D model of the subsurface. In MASW a seismic source is recorded by a linear array of geophones or other vibration sensors, then a dispersion curve is calculated from that data, indicating that Rayleigh waves at each of the wavelengths of interest, W , travel at phase velocities, C_e (Park et al., 1999), (Louie, 2001). To infer whether any one-dimensional test velocity model of the subsurface would explain the observed surface wave dispersion, C_e , a theoretical dispersion curve, C_t is modeled based on how surface waves would be expected to propagate in that test velocity model. This procedure is repeated many times in an optimization loop, minimizing the difference between C_e and C_t , the current test velocity's theoretical dispersion curve. This technique is widely used in geotechnical

engineering for site investigation, particularly because it does not require drilling cores or samples, permitting is simpler, and it can be applied to 15 both active controlled-source and passively recorded seismic signals (Park et al., 2007).

There exists open-source code available to geotechnical engineers, for example, MASWaves which is implemented in Matlab (Olafsdottir et al., 2018). However, to our knowledge there is not currently any open-source MASW software that specifically focuses on optimizing performance of the theoretical dispersion curve calculation present in MASW inversion. There have been algorithms developed to improve the outer optimization problem that determines the order in which velocity models are tested, particularly linearization with *a priori* information, swarm or genetic algorithms, and pattern matching (Cercato, 2009; Song et al., 2008, 2015). Thus we are motivated to implement MASW inversion in C with a focus on fast performance in the theoretical

^{*} Corresponding author.

E-mail address: josek97@vt.edu (J. Kump).

URL: <https://github.com/jlk9/MASWA> (J. Kump).

¹ Authorship Statement: Joseph Kump was the primary software and algorithm developer and contributed to writing this manuscript. Eileen Martin determined the initial problem statement, advised on software development and testing and contributed to writing this manuscript.

dispersion curve analysis, which could then be matched with any improved optimization search technique. Faster analysis of MASW data means that a single theoretical dispersion curve could be evaluated faster, but more importantly it would reduce the cost of repeated evaluations of theoretical dispersion curves. We always perform MASW inversion as part of a larger optimization problem to find the “best” velocity model to explain our data. Further, engineers are interested in quantifying the uncertainty in subsurface velocity models. Both scenarios require repeated evaluations of MASW inversion, so any inefficiencies in existing codes repeatedly accumulate into significant additional electricity cost and wait time for engineers.

As we expand our problem sizes to larger datasets (potentially having finer scale dispersion curves, or extending over wider frequency ranges), we can take advantage of computer clusters which have many cores to drastically reduce the time to actionable solutions. A cluster environment is already the current scenario for rapidly evaluating uncertainty in dispersion images (Dou et al., 2017). Further, the majority of laptops today have multiple cores, so parallelizing MASW will allow larger datasets to be analyzed in the field. In addition to parallelizing algorithms over CPUs, we can take advantage of other hardware: graphics cards. General purpose graphics processing units (GPUs), can be programmed to perform scientific computing simulations with significantly lower electricity cost and time per unit of computation. The majority of engineering workstations in offices today have a graphics card, even many gaming laptops that could be brought into the field, so a typical engineer could benefit from MASW software running on a GPU.

In this paper, we propose new algorithms to accelerate MASW inversion using MPI and GPUs, and introduce MASWAccelerated, a new open-source software package implementing these algorithms. In section 2, we give an overview of the serial algorithm for MASW, the MPI parallel algorithm, and the GPU accelerated algorithm. In section 3, we report on performance tests to show the speedup compared to the existing Matlab code MASWaves (Olafsdottir et al., 2018), to show the scalability of the MPI algorithm, and to show speedups when utilizing a uniform velocity dispersion curve and a more typical variable velocity dispersion curve. Further, we describe performance optimizations utilizing the sparsity structure of stiffness matrices, strategies to tackle load imbalancing which account for typical dispersion curve trends, and improved GPU performance when repeatedly evaluating the kernel as part of an optimization procedure. With these optimizations, the serial C implementation obtains a 2x speedup over MASWaves, the MPI implementation shows near perfect strong and weak scalability with uniform data, a modified partition of the problem for more realistic data shows near-perfect strong scaling up to 8 processes, and the GPU implementation obtains a 3.2x speedup over the serial C implementation. These speedups will enable engineers to perform MASW inversion significantly faster, potentially even in the field.

2. Overview of algorithms

Multichannel Analysis of Surface Waves (MASW) consists of two separate processes. In MASW dispersion, an array of receivers record the surface wave energy that travels along the ground, called the ground roll, that is generated by an impulsive source. Spectral analysis is performed on the recorded ground roll to generate a dispersion curve, which plots the relationship between its frequencies and velocities. This is possible because the ground roll consists primarily of Rayleigh waves and exhibits dispersion: each of its frequency components has a different propagation velocity. Methods to automate dispersion curve selection are being developed. One approach is to test an unsupervised machine learning technique using the Gaussian mixture model clustering method to identify points of dispersion energy within the image, and then apply spatial clustering to determine the dispersion curve (Wang et al., 2021). Another is to apply narrow filtering to the signal data to make identifying curves within the dispersion image easier (Granados et al., 2019). However, MASW dispersion currently requires manual selection to

identify the dispersion curve within a two-dimensional plot of the dispersion image, and is typically done only once per site and its corresponding data. Thus, it is not in the scope of our software.

In MASW inversion, we test the viability of a proposed ground model for the near surface by generating a dispersion curve of the ground roll using the model parameters. This theoretical dispersion curve is compared to the experimental dispersion curve produced in MASW dispersion; if the numerical error between the two dispersion curves is small, then the proposed model is likely an accurate estimate of the near surface. MASW inversion is a suitable method for parallelization, and is typically run many times for a single site, thus it is the focus of our software.

We review the standard serial algorithm for MASW inversion from (Olafsdottir et al., 2018), and note optimizations we have made. We propose new algorithms that parallelize MASW modeling using MPI, and that perform MASW modeling on GPUs.

2.1. Serial implementation

Before parallelizing the algorithm for MASW inversion, we need to implement a C serial version of the code which can then be modified with MPI and CUDA. We used MASWaves, an existing implementation written in Matlab as our reference to compare against (Olafsdottir et al., 2018). Initially our C serial implementation was a simple port of MASW to C modeled after the MASWaves implementation, but we later made a number of changes to the algorithm to improve its efficiency.

The model evaluation process in MASW inversion consists of two main algorithms. First, the model parameter inputs are used to compute the most likely velocities for each wavelength in the experimentally

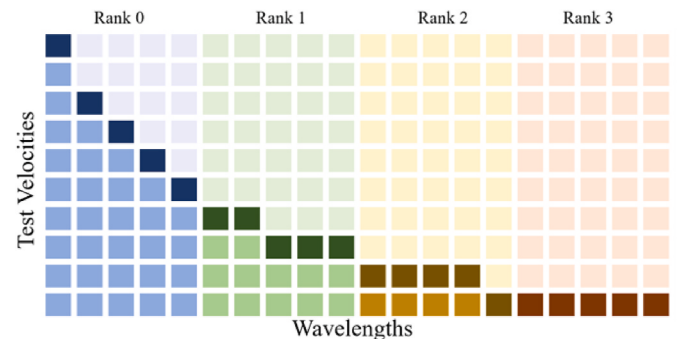


Fig. 1. Visualization of original MPI partition. Darker colors correspond to determinants that were computed, with the darkest highlighting the theoretical curve. Each rank has its own color scheme. Note the number of computed determinants varies heavily for each rank.

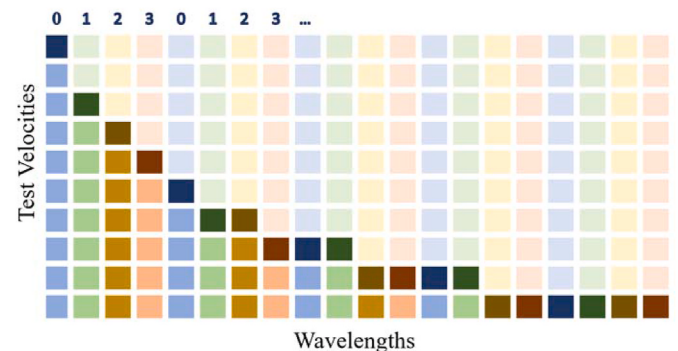


Fig. 2. Visualization of modular MPI partition, with the same color rules in place. In this case the computed determinants are balanced more evenly.

derived dispersion curve. This is done in the function MASW Theoretical Dispersion Curve, illustrated in Algorithm 1. The model parameter inputs (M) are the number of finite thickness layers, the thickness and density of each layer, and the compressional (P) and shear (S) wave velocities through each layer.

Algorithm 1. MASW Theoretical Dispersion Curve

Algorithm 1 MASW Theoretical Dispersion Curve

```

1:  $M$  = Guessed model parameters
2:  $W$  = Wavelength values from experimental dispersion curve
3:  $C_t$  = Velocity corresponding to each wavelength given  $M$ 
4:  $V$  = Range of potential velocity values to test
5:  $d_{old}$  = Determinant value for previous entry in  $V$ 
6:  $d_{new}$  = Determinant value for current entry in  $V$ 
7: for  $w$  in  $W$  do
8:    $d_{old} = \text{stiffness\_matrix}(M, V[0])$ 
9:    $d_{new} = \text{stiffness\_matrix}(M, V[1])$ 
10:   $n = 1$ 
11:  while  $\text{sign}(d_{old}) == \text{sign}(d_{new})$  do
12:     $n = n + 1$ 
13:     $d_{old} = d_{new}$ 
14:     $d_{new} = \text{stiffness\_matrix}(M, V[n])$ 
15:  end while
16:   $C_t[w] = V[n]$ 
17: end for

```

each test velocity (in V) in increasing order, until one has a determinant (d_{new}) with a different sign than its predecessor. This test velocity ($V[n]$) is then stored as the theoretical velocity corresponding to that particular wavelength, i.e. it is stored as $C_t[w]$. The process is repeated for each wavelength in the experimental dispersion curve to generate a theoretical dispersion curve of Rayleigh wave velocities and their wavelengths.

MASWaves, like many implementations of MASW, utilizes the stiffness matrix method (Kausel, 1981) to compute the theoretical dispersion curve. Given a wavelength, sparse stiffness matrices are computed for

A stiffness matrix may be generated for each wavelength and each test velocity. Realistically a dispersion curve may have up to 100 wavelengths and 1000 test velocities, requiring a total of 100,000 stiffness matrices to be generated (in practice this number is often lower, since the ideal test velocity is usually found before all velocities are checked for a given wavelength). These stiffness matrices also have a sparse banded structure, which is not utilized by MASWaves but is in our algorithms.

The algorithm also requires the determinants for each of these matrices to be computed. The stiffness matrices are of size $2(N+1) \times 2(N+1)$, where N is the number of finite-thickness layers in M . Since they have a symmetric banded structure, computing the entries of a stiffness matrix and its determinant requires $O(N)$ operations.

Once the theoretical dispersion curve C_t is computed, the second part of the algorithm compares its velocities to the velocities experimentally derived from the data, labelled C_e . The average relative error is labelled the misfit, and indicates how accurate M serves as a model of the ground structure. Essentially the misfit behaves as a loss function for MASW.

The model parameters M , along with the given test velocity and wavelength, are used to compute each stiffness matrix. If the theoretical velocity for each wavelength is close to the experimental velocity (i.e. the misfit is small), then the model parameters are more likely to be a good approximation of the ground truth.

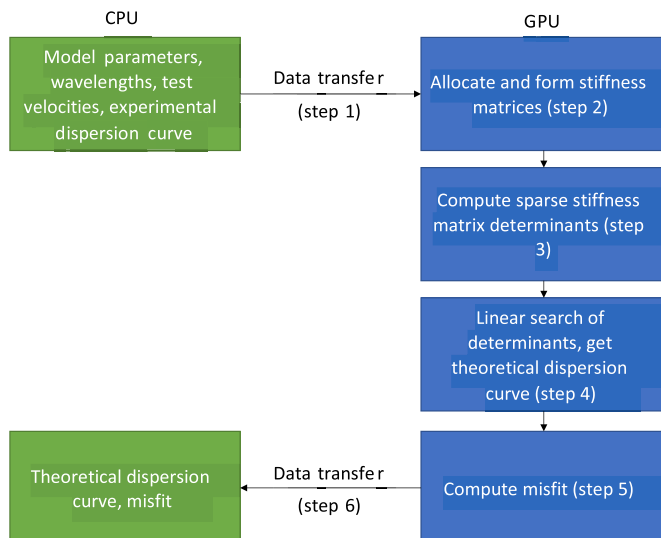


Fig. 3. General outline of GPU implementation.

Algorithm2. MASW Misfit

need to be computed if the first sign change (and thus correct test velocity) has already been found. Thus this approach may lead to several

Algorithm 2 MASW Misfit

C_t = Velocity corresponding to each wavelength given M

2: C_e = Experimentally derived velocities

e = Relative errors

4: l = Length of W , C_t , and C_e

m = the average misfit

6: **for** $i = 1$ to l **do**

$e = e + \frac{|C_t[i] - C_e[i]|}{C_e[i]}$

8: **end for**

$m = \frac{e}{l}$

A downside to MASW is that it is nonlinear and does not have any form of backpropagation to accompany its loss function. Therefore the only way to minimize the misfit is to compute theoretical dispersion curves for a large quantity of plausible model parameters, each of which require forming and finding the determinants of up to 100,000 small matrices. While methods have been used to more efficiently select which models should be tested, such as linearized algorithms based on inequality constraints, pattern search algorithms, and bee swarm algorithms, these still require a large number of model theoretical dispersion curves to be computed (Cercato, 2009; Song et al., 2008, 2015). Since each of these curves and their misfits can be computed independently, this algorithm can benefit from parallelization, both between different test models and within individual models.

2.2. MPI parallelism

It is possible to visualize MASW inversion as a mesh computation, although it is not a literal ground mesh. Stiffness matrices must be computed with several different values of wavelengths and test velocities and the same values of model parameters. If the wavelength and velocity values are thought of as the x and y axis, then MASW Theoretical Dispersion Curve can be viewed as computing data points along a two dimensional “grid”. It is reasonable to partition this grid of computations into multiple processes in an MPI implementation. This design philosophy is analogous to mesh computations in some parallel finite-difference method implementations, where each process is assigned one subsection of the mesh covering a specific range of spatial (x and y) coordinates (Londhe et al., 2021).

There are two obvious ways to partition this grid. The stiffness matrix method heavily utilizes the test velocities in computing the individual matrices. Therefore, one could partition the grid along the velocity axis: given S processes and V test velocities, compute stiffness matrix determinants for all wavelengths and the first $\frac{V}{S}$ velocities in process 0, then the next $\frac{V}{S}$ velocities in process 1, and so on. This approach would allow components of the stiffness matrix dependent on the test velocity to be pre-computed, reducing the number of repetitive computations.

Partitioning along the velocity axis would have significant drawbacks, however. Once the determinants are computed, a linear search must be performed along each wavelength to find the first sign change. Since the determinants for each wavelength are split along multiple processes, a significant amount of communication would be required between each process to find the first sign change.

In addition, determinants for higher test velocities might not even

unnecessary computations.

Because of these difficulties, the method used was partitioning along the wavelength axis, shown in Fig. 1. The stiffness matrix computations do not feature the wavelength input as much as the test velocity, so there is little potential to precompute components of the matrices for each wavelength. But finding the ideal test velocity for a given wavelength has no dependence on other wavelength values, so no communication between processes is required for the dispersion curve. In addition, once the ideal test velocity is found for a particular wavelength, the process can begin computing matrix determinants for the next wavelength. This preserves the serial implementation’s advantage of avoiding unnecessary matrix and determinant computations. Partitioning along the wavelengths makes computing the misfit parallel as well, with only one reduction operation required to combine the errors from each process.

Since the algorithm does not require communication between different wavelengths of the dispersion curve, there is freedom to choose exactly how to partition the problem along the wavelength axis. The original approach was to assign wavelengths contiguously: given s processes and wavelengths, assign wave-lengths 0, 1, ..., $\frac{W}{s}$ to process 0, wavelengths $\frac{W}{s} + 1$, ..., $\frac{2W}{s}$ to process 1, and so on. We determined a more efficient approach was to partition wavelengths in a modular pattern: assign every wavelength equivalent to 0 mod s to rank 0, every wavelength equivalent to 1 mod s to rank 1, and so on. This partition is visualized in Fig. 2.

2.3. GPU acceleration

Many workstations used by geotechnical engineers contain CPUs with multiple cores as well as graphics cards. Thus, in addition to writing MASW inversion for MPI, we have also implemented the MASW inversion algorithm for graphics processing units (GPUs). MASW requires calculating many determinants of sparse matrices, a problem which has not previously been adapted to GPUs.

There are five main steps to the GPU implementation that are distinct from the serial and MPI versions of MASW.

First, the initial input values to the algorithm - wavelengths, test velocities, experimental dispersion curve, and the ground model parameters - are transferred from the CPU “host” to the GPU “device” as shown in step 1 of Fig. 3.

Copying memory between the host and device is costly, but the volume of data copied here is relatively small (the test velocity array may have ≈ 1000 entries, and the other inputs are scalars or much shorter arrays), so this transfer is not problematic.

Once the inputs are on the device, the stiffness matrices are allocated

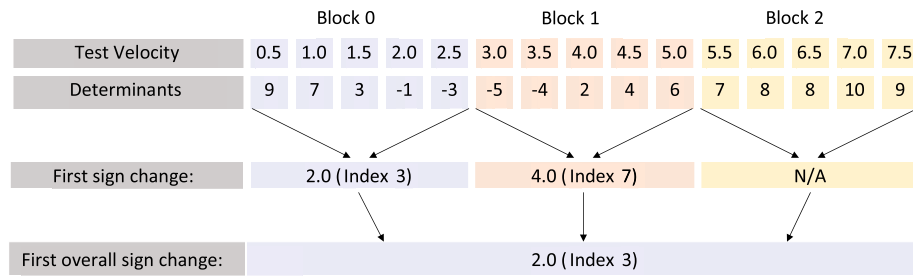


Fig. 4. Outline of GPU determinant search along one wavelength. The first sign change within each block is found, then a search over the blocks finds the first overall sign change.

as global device memory and their entries are filled in based on the model parameters M , shown in step 2 of Fig. 3. Forming the stiffness matrices is one of the most costly steps in MASW inversion, so effective parallelization is critical. Since a GPU contains thousands of cores (as opposed to the dozens that may be available on one or more CPUs), it is feasible to partition the problem along both the wavelength and velocity axis dimensions of the grid. In fact, it is reasonable to compute all the stiffness matrices concurrently since they are mutually independent. Using CUDA, each thread is assigned to fill in the entries for one stiffness matrix.

Computing values for a stiffness matrix is $O(N)$ as described in the serial implementation, where N (the number of finite thickness layers in M) is typically small ($N \leq 10$). Entries in the stiffness matrices are incrementally increased multiple times, so trying to compute a stiffness matrix across multiple threads can lead to race conditions. Thus assigning a single thread to each stiffness matrix is reasonable.

Once the stiffness matrices are formed, Gaussian elimination is performed so their determinants can be easily computed, shown in step 3 of Fig. 3.

This is roughly equal to forming the stiffness matrices in terms of time cost.

Initially, we used the function `cublasZgetrfBatched()` to perform LU factorizations on the stiffness matrices. This function took over 50% of the runtime of the GPU implementation, most likely because it did not take advantage of the stiffness matrices' banded structure and was therefore $O(N^3)$ for each factorization. Because of this, we replaced it with a kernel that assigned one thread to each stiffness matrix to perform a banded Gaussian elimination. This function was roughly ten times faster than `cublasZgetrfBatched()`, and approximately equal in runtime to the kernel that formed the stiffness matrices.

The next step is to find the first test velocity whose corresponding

determinant is the opposite sign of its predecessor for each wavelength, which is step 4 in Fig. 3. Normally this would require a linear search along all the test velocities for each wavelength - a serial $O(N)$ process that does not lend well to GPU architecture. However, there is still some potential to partition the problem along the device cores. The search kernel breaks up the stiffness matrices along its blocks by wavelength, and along threads within each block by test velocity.

This is illustrated in Fig. 4:

Each wavelength will be paired with multiple blocks, since the block size is assumed to be 256 (for compatibility with older GPUs) and often MASW is run with ≥ 500 test velocities. Within each block, thread i computes the determinant of its respective stiffness matrix by multiplying the diagonal entries, then compares it to the determinant of thread $i+1$. It then stores the result of the sign comparison into shared memory.

Once this is complete for all threads, the first thread of each block then performs a linear search for the first sign change in its shared memory block, and stores this as the index of the first sign change within that range of test velocities. These results are placed in an $l \times b$ matrix, where l is the number of wavelengths and b is the number of blocks assigned to each wavelength. The next kernel then iterates over S to find the first recorded sign change for each wavelength, which is a small linear search over approximately 4 entries. The test velocity whose stiffness matrix determinant produced the first sign change is then labelled as the velocity corresponding to that wavelength in the theoretical dispersion curve.

The final steps are to average the errors between the theoretical dispersion curve and experimental curve to get the overall misfit, and to send the theoretical dispersion curve and misfit to the host, which are steps 5 and 6 of Fig. 3. The former is effectively equivalent to a vector dot product (in terms of number of FLOPs), which is straightforward to implement in CUDA. The latter is simply a CUDA memory copy involving a vector of length ≤ 100 and a scalar. Steps 1 and 6 require

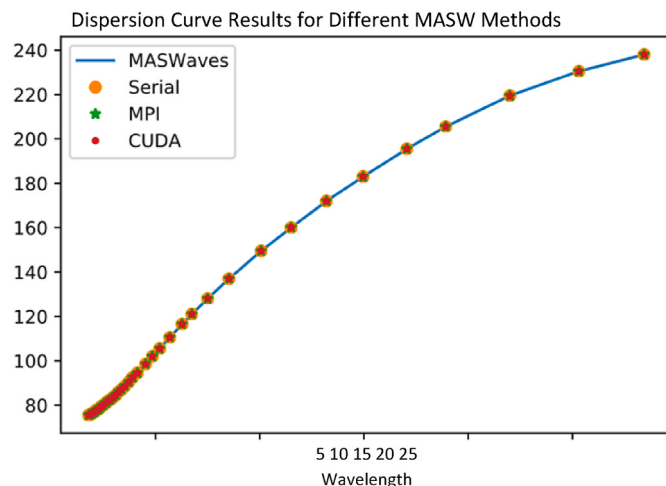


Fig. 5. Results from different implementations of MASW show that all implementations accurately calculate the same dispersion curve.

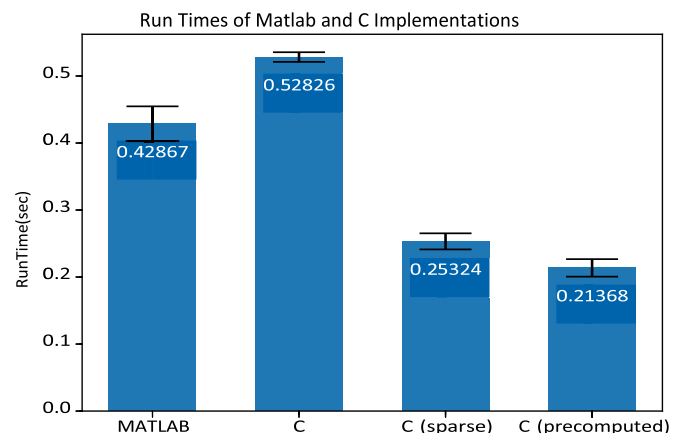


Fig. 6. Comparison of Matlab and C on a variable dataset.

only a small amount of data to be transferred between the host and device, while the stiffness matrices, which take up much more memory, are allocated and freed exclusively on the device.

3. Test results and performance optimizations

We developed a variety of tests for correctness of the code (unit tests and end-to-end tests), as well as performance and scalability tests to understand code efficiency. All of our new implementations of MASW inversion - serial C, MPI, and CUDA - appear to compute theoretical dispersion curves correctly. A plot of each of their results for one set of model parameters is shown in Fig. 5. Since the stiffness matrix method only uses the signs of matrix determinants to determine the entries of the theoretical dispersion curve and is not dependent on the determinants' actual values, it is robust to potential sources of error. Each implementation selected the same test velocities for every wavelength of the theoretical dispersion curve as MASWaves, which we treat as our baseline MASW implementation.

We used two main test cases to evaluate MASWAccelerated's performance. A synthetic dispersion curve of identical wavelength values was used to provide a "uniform" dataset. This is useful because of the design of the serial and MPI algorithms - since every test velocity is evaluated in increasing order, varying wavelengths with different theoretical velocities will have different runtimes. Thus a uniform test allows us to observe other factors that may affect how MASWAccelerated scales with more data. The second test case was a more realistic "variable" dataset, with decreasing wavelengths, matching a typical dispersion curve. We present these test cases in the code as testScaling and testProcess, respectively. Both of these functions also have wrappers (testScaling_full and testProcess_full) that enable test cases to be run multiple times in a loop.

3.1. Serial tests

The primary purpose of the serial C implementation is to enable usage of MPI and CUDA to parallelize MASW inversion. However, it is still useful to compare its speed to the original Matlab version, MASWaves. We timed the speed of MASWaves on our realistic dataset, and compared it to MASWAccelerated's speed on the variable dataset when run in serial. Both of these tests were run ten times on a laptop with a 3.1 GHz Intel Core i7 CPU, and the mean results are shown in Fig. 6, along with errors denoting one sample standard deviation.

Initially the serial C implementation, while mathematically correct, was slower than MASWaves, which is not ideal. This is likely because MASWaves made use of Matlab's vectorized operations to compute stiffness matrix determinants, while the initial Gaussian elimination

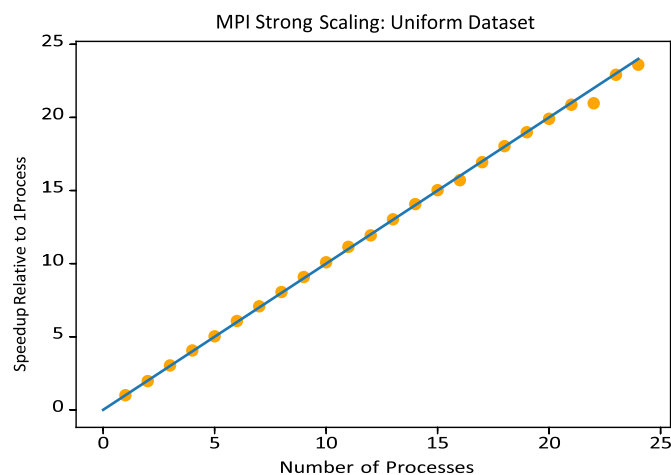


Fig. 7. MPI strong scaling on uniform data.

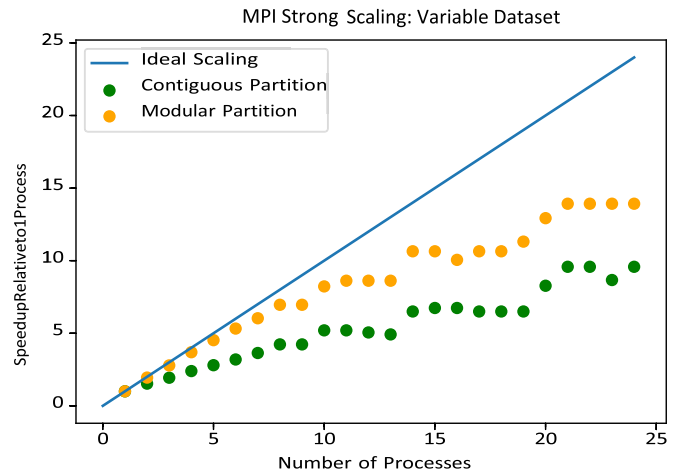


Fig. 8. MPI strong scaling on variable data, using different partitions. The variable data highlights the two features that can make MASW problematic for MPI: large variations in wavelength values for the dispersion curve, and short dispersion curve length.

algorithm written for MASWAccelerated was not vectorized.

The stiffness matrices formed by MASW are always sparse, and moreover have a banded heptadiagonal structure (nonzero entries only on the main diagonal and the three above and below it). We used this fact to improve the Gaussian elimination algorithm to be only $O(N)$ instead of $O(N^3)$, which increased the algorithm's speed by 2.0 times, as shown in the C (sparse) column. We also noticed components of the entries in the stiffness matrices could be pre-computed before iterating over the matrix entries in a loop, reducing the number of arithmetic operations required and increasing the algorithm speed a further 18% as shown in the time for C (precomputed). These improvements enable the serial C implementation to be over twice as fast as MASWaves without any parallelization or vectorization calls. The precomputed serial C version is used as the basis for the MPI algorithm, and for performance comparison with the GPU algorithm.

It is worth noting the runtime for the first instance of MASWaves was significantly slower than subsequent runs, causing the increased error. This may be due to some type of caching effect in Matlab, which is not present in C.

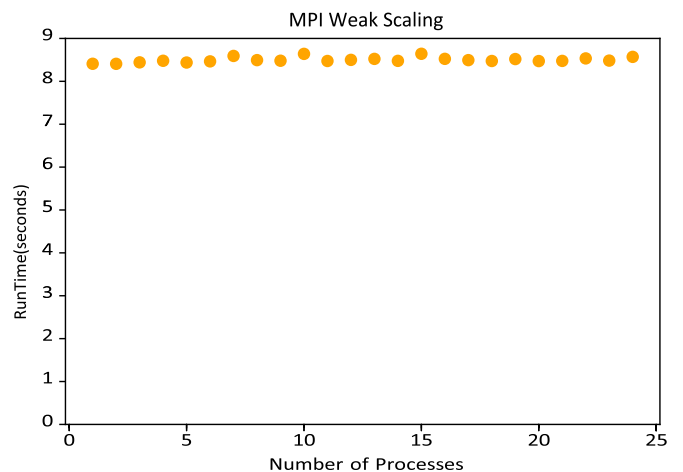


Fig. 9. MPI weak scaling on uniform data, with a dispersion curve of length $1000 \times$ number of processes.

3.2. MPI tests

As the MPI implementation with one process is virtually identical to the serial implementation, the main purpose of testing is to evaluate how it scales with more processes and larger datasets. For this purpose we ran a few strong and weak scaling studies.

For the strong scaling study, we first tested the algorithm on the uniform dataset with 1000 wavelengths, allowing for distinctions in runtime to be more noticeable. The results are shown in Fig. 7, which shows the average runtime for each process count on ten runs.

This test was run on the NewRiver computing cluster at Virginia Tech, using two Haswell E5-2680v3 2.5 GHz processors with 12 cores each. The algorithm scaled almost exactly linearly with more cores (except at 22 cores, which may be due to some error), which is expected given the minimal amount of communication required between processes and the highly parallel method used to partition MASW.

In practice, dispersion curves usually have a couple features that make the MPI algorithm scale less than linearly, as seen in the strong scaling test for the uniform dataset.

One problem is that dispersion curves do not have identical wavelength values throughout, but rather varying wavelength values which correspond to different velocities. Since MASW must evaluate test velocities in increasing order to identify the determinant sign change, entries with larger velocities will require more stiffness matrix determinants to be computed, resulting in a longer run time. Because of this, the MPI algorithm is prone to load imbalancing on realistic data, even though the entries of the dispersion curve are partitioned evenly.

For near-surface imaging, most dispersion curves are decreasing in both wavelength and velocity, and usually resemble a continuous curve (this is called a normally dispersive profile, and is guaranteed in the common condition where deeper layers have increasing shear and compressional wave velocities (Foti et al., 2018)). A naive contiguous partition of the dispersion curve will place entries with similar wavelengths and velocities on the same process, thus resulting in a few processes receiving all of the high-velocity entries, exacerbating the load imbalance. The modular partition of the dispersion curve mitigates this problem, and improves the MPI algorithm's strong scaling as a result. The comparison of speedup for these partitions on the variable dataset can be seen in Fig. 8. A line has been added to compare both to ideal linear scaling.

This test was run on the same hardware as the uniform strong scaling, and again the average of ten runs per size was used. Although still not linear, the modular partition scales significantly better than the original contiguous partition: with 3 processes it has a speedup of nearly

2.8 compared to 1.9 for the naive approach, and with 8 processes it has a speedup of nearly 7.0 compared to only 4.2. Overall, the MPI algorithm with a modular partition will have near-linear scaling for most datasets with relatively few processes.

The reduction in speedup with more processes is likely due to another problem - most dispersion curves have a relatively short length (the variable dispersion curve has 40 entries, which is fairly typical). When the size of the partition is small, approximately 10 or less, each additional process reduces the number of dispersion curve entries computed for all of the other processes. For example, at size 3 each process is computing velocities for 13 or 14 wavelengths, while at size 4 each process is computing velocities for only 10 wavelengths. This is a significant reduction in workload and results in major speedup as seen in Fig. 8. But at larger sizes there is not a reduction in workload for every process. For example, at size 20 each process is computing 2 entries, while at size 24 16 processes are still computing 2 entries and the last 8 are computing 1. Since many processes have no reduction in workload, the overall runtime of the algorithm is not reduced. Therefore, the MPI partition scales near linearly with relatively few processes (depending on the size of the dispersion curve), but experiences diminishing returns with more processes.

A weak scaling study was also performed on the NewRiver cluster using the uniform dataset, again with the same hardware and taking the average of ten runs per size. The size of the dispersion curve was $1000 \times$ the number of processes. As seen in Fig. 9, the MPI implementation scaled efficiently by this measure, having no significant increase in time with more data and processes. There are a few sizes with marginally higher runtimes - 7, 10, and 15 - but these are minor and likely due to external factors. Like the strong scaling study with uniform data, this test highlights the minimal communication requirements for the MPI implementation.

3.3. GPU tests

First we compared the CPU and GPU implementations on the variable dataset. This was done eleven times in the wrapper loop testProcessFull using the MPI implementation (with one process), then eleven times both in a loop and separately with the CUDA implementation. The resultant run times are shown in Fig. 10. Note the light blue bars are the first run for each method, while the dark blue bars are the means of subsequent runs with their sample standard deviation posted. The CPU algorithm is unchanged from the serial and MPI tests, but its runtime is different since it was run on a different machine.

All three of these implementations were run on the same desktop, using an Intel Xeon CPU E3-1271 v3 @ 3.60 GHz and an Nvidia Quadro K620 GPU. The first run is the run time of the first instance of

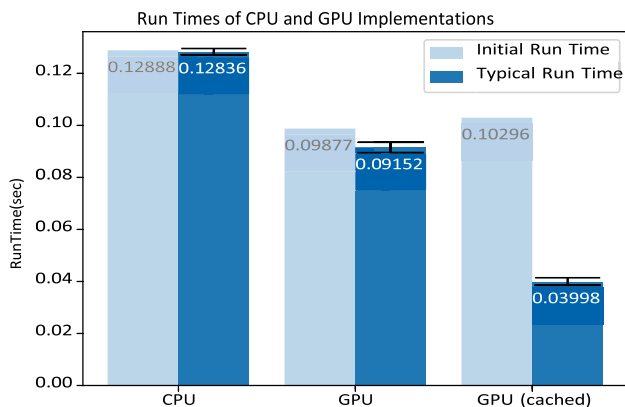


Fig. 10. Comparison of CPU and GPU run times on the variable dataset. The light blue bars show the time for the first run of the MASW implementation, while the dark blue bars show the typical run times excluding the first run. Running the GPU implementation multiple times in a for-loop produces significant speedup for all runs after the first.

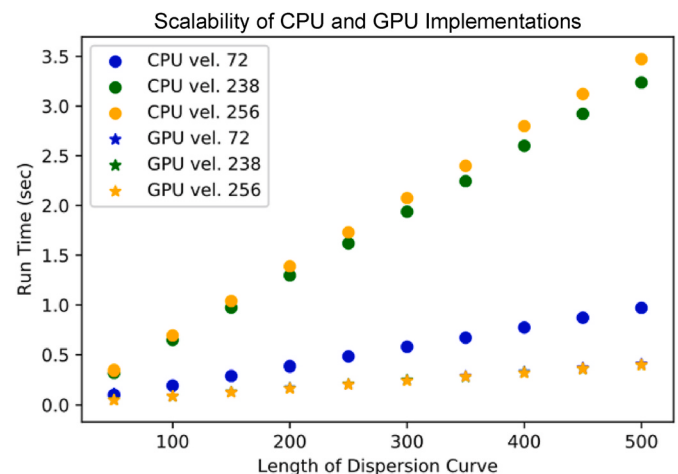


Fig. 11. Comparison of CPU and GPU on increasing uniform datasets.

MASW_inversion, while the subsequent runs denote the mean runtime of all other instances. The error bars denote one standard deviation for the subsequent runs. The GPU implementation of MASW_inversion is about 25% faster than the CPU implementation when run once, but over 3.2 times faster when it is run multiple times in a for loop. This is likely because of just-in-time compilation, used by Nvidia to allow CUDA kernels to benefit from new device architectures. When a kernel is run multiple times within a function call, it only needs to be compiled for the first kernel run while it is “cached” for subsequent runs. Since MASW inversion is often run multiple times with different test velocity models, it is reasonable to design MASWAccelerated to run inversions on multiple models to take advantage of this caching effect. It is worth noting that, while not as dramatic, there is still roughly 8% speedup when running the inversion multiple times individually on the GPU, but no significant speedup for the CPU implementation.

We next use the uniform dispersion curve to evaluate how the GPU implementation performs on progressively larger datasets. This test ran the MPI (one rank) and CUDA implementations of MASW_inversion on dispersion curves of lengths 50–500. The wavelength values of these curves were designated to match test velocities of 72, 238, and 256 (each implementation and each dispersion length was run three times with three different wavelength values). The results are shown in Fig. 11, run on the same desktop used for the previous test.

As expected, the CPU implementation scaled linearly with the length of the dispersion curve. The exact run time is highly dependent on the theoretical velocity values, as shown in the MPI scaling studies. The GPU implementation was significantly faster (since the caching effect was recognized in testing the variable data, we made use of it here), and had no dependence on the theoretical velocities since all stiffness matrices are computed regardless.

It is worth noting that too large a dispersion curve or too many test velocities can overload the GPU memory on the CUDA implementation, since all the stiffness matrices are formed concurrently in global memory and the dispersion curve and test velocities determine the number of stiffness matrices. The exact upper limit depends on the memory space of the GPU and the number of finite thickness layers in the model M (which determines the size of the stiffness matrices), but typical problem sizes for MASW will not take up too much memory, even for older GPUs. We found dispersion curves larger than 500 typically caused memory problems for the Quadro K620 GPU, which has 2 GB of global memory. This is because each stiffness matrix has 196 entries, each of which is a CuDoubleComplex datatype that takes 16 bytes, so $500 \text{ wavelengths} \times 1000 \text{ test velocities} \times 13136 \text{ bytes} \approx 1.6 \text{ GB}$, close to the memory limit of the GPU.

4. Discussion and conclusions

There are two major priorities to improve the usability of MASWAccelerated.

First, an I/O system would make it much easier to input model data and collect results from the inversion process. It would also eliminate the need to recompile the code each time new data is being inputted. The other priority is enabling command line arguments, such as the option to evaluate multiple input files (with separate data models) at once, or to run the suite of test functions. These two enhancements are necessary to enable the next step of MASWAccelerated being easily applied to a wide variety of use cases.

Another useful enhancement would be integrating the MPI and Cuda implementations of MASWAccelerated into one program. Theoretically, this is feasible, as each process in MPI is essentially MASW inversion carried out only on a portion of the dispersion curve, which can still be done by a GPU. Solving this problem would enable further optimization of the algorithm to various hardware configurations.

In this paper, we have proposed two new algorithms for MASW inversion: one using MPI parallelism, and one accelerating the code with Cuda for GPUs.

We provide an open-source implementation called MASWAccelerated, along with test cases to verify code correctness and efficiency. Our comparisons show significant speedups over the existing MASWaves code, and we show several optimizations that take advantage of the problem structure (both typical trends in data and sparsity structure of stiffness matrices) to further improve code efficiency. These modifications will help engineers to perform rapid data analysis by taking advantage of all computer hardware available to them, ideally even performing analysis in the field as data are acquired.

Computer code availability

The MASWAccelerated software, along with examples to produce the results in this paper, are publicly available at <https://github.com/jlk9/MASWA> under an MIT license. The code was primarily developed by Joseph Kump, who can be reached by email at josek97@vt.edu. The MASWAccelerated code was first made publicly available in 2020 upon submission of this paper. Hardware and software requirements, as well as other code features, are available in the repository's README file.

Disclaimer

This report was prepared as an account of work sponsored by an agency of the United States Government. Neither the United States Government nor any agency thereof, nor any of their employees, makes any warranty, express or implied, or assumes any legal liability or responsibility for the accuracy, completeness, or usefulness of any information, apparatus, product, or process disclosed, or represents that its use would not infringe privately owned rights. Reference herein to any specific commercial product, process, or service by trade name, trademark, manufacturer, or otherwise does not necessarily constitute or imply its endorsement, recommendation, or favoring by the United States Government or any agency thereof. The views and opinions of authors expressed herein do not necessarily state or reflect those of the United States Government or any agency thereof.

Declaration of competing interest

The authors declare the following financial interests/personal relationships which may be considered as potential competing interests: Coauthor Eileen Martin is an associate editor for Computers & Geosciences. Therefore, the paper should be handled and reviewed by a different editor.

Acknowledgements

We would like to thank the DOE Office of Geothermal Technologies for financial support through the STTR Phase I Award No. DE-SC0019630, which is a collaboration with Luna Innovations led by P.I. Derek Rountree. This work was supported in part by the National Science Foundation Award No. 1937984, through the Engineering for Civil Infrastructure Program. We thank Virginia Tech Advanced Research Computing for computing resources. This material is based upon work supported by the U.S. Department of Energy, Office of Science, SC-1 under Award Number DE-SC0019630.

References

- Cercato, M., 2009. Addressing non-uniqueness in linearized multichannel surface wave inversion. *Geophys. Prospect.* 57 (1), 27–47. <https://doi.org/10.1111/j.1365-2478.2007.00719.x> arXiv:<https://onlinelibrary.wiley.com/doi/pdf/10.1111/j.1365-2478.2007.00719.x>. <https://onlinelibrary.wiley.com/doi/abs/10.1111/j.1365-2478.2007.00719.x>.
- Dou, S., Lindsey, N., Wagner, A., Daley, T., Freifeld, B., Robertson, M., Peterson, J., Ulrich, C., Martin, E., Ajo-Franklin, J., 2017. Distributed acoustic sensing for seismic monitoring of the near surface: a traffic-noise interferometry case study. *Sci. Rep.* 7 article no. 11620.

- Foti, S., Hollender, F., Garofalo, F., Albarello, D., Asten, M., Bard, P.-Y., Comina, C., Cornou, C., Cox, B., Di Giulio, G., Forbriger, T., Hayashi, K., Lunedei, E., Martin, A., Mercierat, D., Ohrnberger, M., Poggi, V., Renalier, F., Sicilia, D., Socco, V., 2018. Guidelines for the good practice of surface wave analysis: a product of the inter-pacific project. *Bull. Earthq. Eng.* 16, 2367–2420. <https://doi.org/10.1007/s10518-017-0206-7>.
- Granados, I., Cal, M., Ramos, V., 2019. Noisy dispersion curve picking (ndcp): a matlab package for group velocity dispersion picking of seismic surface waves. *Comput. Geosci.* 133, 104315. <https://doi.org/10.1016/j.cageo.2019.104315>. <https://www.sciencedirect.com/science/article/pii/S0098300417312669>.
- Londhe, A., Rastogi, R., Srivastava, A., Khonde, K., Sirasala, K.M., Kharche, K., 2021. Adaptively accelerating fwm2da seismic modelling program on multi-core cpu and gpu architectures. *Comput. Geosci.* 146, 104637. <https://doi.org/10.1016/j.cageo.2020.104637>. <https://www.sciencedirect.com/science/article/pii/S0098300420306142>.
- Louie, J., 2001. Faster, better: shear-wave velocity to 100 meters depth from refraction microtremor arrays. *Bull. Seismol. Soc. Am.* 91 (2), 347–364. <https://doi.org/10.1785/0120000098>.
- Olafsdottir, E., Erlingsson, S., Bessason, B., 2018. Tool for analysis of multichannel analysis of surface waves (masw) field data and evaluation of shear wave velocity profiles of soils. *Can. Geotech. J.* 55 (2), 217–233. <https://doi.org/10.1139/cgj-2016-0302>.
- Park, C., Miller, R., Xia, J., 1999. Multichannel analysis of surface waves. *Geo505 Phys.* 64 (3), 800–808. <https://doi.org/10.1190/1.1444590>.
- Park, C., Miller, R., Xia, J., Ivanov, J., 2007. Multichannel analysis of surface waves (masw)- active and passive methods. *Lead. Edge* 26 (1), 60–64.
- Song, X., Gu, H., Zhang, X., Liu, J., 2008. Pattern search algorithms for nonlinear inversion of high-frequency Rayleigh-wave dispersion curves. *Comput. Geosci.* 34 (6), 611–624.
- Song, X., Gu, H., Tang, L., Zhao, S., Zhang, X., Li, L., Huang, J., 2015. Application of artificial bee colony algorithm on surface wave data. *Comput. Geosci.* 83, 219–230.
- Wang, Z., Sun, C., Wu, D., 2021. Automatic picking of multi-mode surface-wave dispersion curves based on machine learning clustering methods. *Comput. Geosci.* 153, 104809. <https://doi.org/10.1016/j.cageo.2021.104809>. <https://www.sciencedirect.com/science/article/pii/S0098300421001084>.