



# MPI Application Binary Interface Standardization

Jeff R. Hammond  
NVIDIA Helsinki Oy  
Helsinki, Finland  
jeffpapers@nvidia.com

Lisandro Dalcin  
Extreme Computing Research Center  
KAUST  
Thuwal, Saudi Arabia  
dalcinl@gmail.com

Erik Schnetter  
Perimeter Institute for Theoretical  
Physics  
Waterloo, Ontario, Canada  
eschnetter@perimeterinstitute.ca

Marc Pérache  
CEA DAM  
Arpajon, France  
marc.perache@cea.fr

Jean-Baptiste Besnard  
ParaTools  
Bruyères-le-Châtel, France  
jbbesnard@paratools.fr

Jed Brown  
University of Colorado Boulder  
Boulder, Colorado, USA  
jed@jedbrown.org

Gonzalo Brito Gadeschi  
NVIDIA GmbH  
Munich, Germany  
gonzalob@nvidia.com

Joseph Schuchart  
University of Tennessee, Knoxville  
Knoxville, Tennessee, USA  
schuchart@utk.edu

Simon Byrne  
California Institute of Technology  
Pasadena, California, USA  
simonbyrne@caltech.edu

Hui Zhou  
Argonne National Laboratory  
Lemont, Illinois, USA  
zhouh@anl.gov

## ABSTRACT

MPI is the most widely used interface for high-performance computing (HPC) workloads. Its success lies in its embrace of libraries and ability to evolve while maintaining backward compatibility for older codes, enabling them to run on new architectures for many years. In this paper, we propose a new level of MPI compatibility: a standard Application Binary Interface (ABI). We review the history of MPI implementation ABIs, identify the constraints from the MPI standard and ISO C, and summarize recent efforts to develop a standard ABI for MPI. We provide the current proposal from the MPI Forum's ABI working group, which has been prototyped both within MPICH and as an independent abstraction layer called Mukautuva. We also list several use cases that would benefit from the definition of an ABI while outlining the remaining constraints.

## CCS CONCEPTS

• **Computing methodologies** → **Massively parallel algorithms**;  
• **Software and its engineering** → **Massively parallel systems**;  
*Cooperating communicating processes*; **Interoperability**; *Software libraries and repositories*.

## KEYWORDS

MPI

## ACM Reference Format:

Jeff R. Hammond, Lisandro Dalcin, Erik Schnetter, Marc Pérache, Jean-Baptiste Besnard, Jed Brown, Gonzalo Brito Gadeschi, Joseph Schuchart, Simon Byrne, and Hui Zhou. 2023. MPI Application Binary Interface Standardization. In *Proceedings of EuroMPI2023: the 30th European MPI Users' Group Meeting (EUROMPI '23)*, September 11–13, 2023, Bristol, United Kingdom. ACM, New York, NY, USA, 12 pages. <https://doi.org/10.1145/3615318.3615319>

## 1 INTRODUCTION

MPI [32] has always been an Application Programming Interface (API) standard, which means that it is standardized in terms of the C and Fortran programming languages. Implementations are not constrained in how they define opaque types (for example, MPI\_Comm), which means they compile into different binary representations. This is fine for users who only use one implementation, or are content to recompile their software for each of these. Many users, including those building both traditional C/C++/Fortran libraries and new languages that use MPI via the C ABI, are tired of the duplication of effort required because MPI lacks a standard Application Binary Interface (ABI).

The potential for implementation agnosticism [15, 39] and specifically an ABI [31], has been recognized for many years. However, no serious effort was made to standardize an ABI, for a variety of reasons. Some of the forces acting against ABI standardization were the diversity of HPC systems, the prevalence of static linking, and the lack of adoption of third-party languages. Over the past 20 years, the HPC hardware and software ecosystem has changed dramatically. Distributing software packages through shared libraries is now common. Package managers, including HPC-oriented ones such as Spack [14], distribute binaries that depend on MPI. There is increasing adoption of MPI by applications written in languages other than C and Fortran [4, 5, 9]. The MPICH ABI Initiative [29, 40]

Permission to make digital or hard copies of part or all of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. Copyrights for third-party components of this work must be honored. For all other uses, contact the owner/author(s).

EUROMPI '23, September 11–13, 2023, Bristol, United Kingdom

© 2023 Copyright held by the owner/author(s).

ACM ISBN 979-8-4007-0913-5/23/09.

<https://doi.org/10.1145/3615318.3615319>

was the first serious effort to create mutually interoperable MPI implementations, by reconciling small differences between the ABIs of MPICH and MPICH-based implementations. This allows applications compiled against appropriate versions of MPICH, Intel MPI, Cray MPI, MVAPICH2 and other implementations to run using the shared libraries from any of the other implementations. This is especially useful to leverage the level of platform-specific specialization that goes into some of these libraries.

Since 2014, the appetite for MPI implementation compatibility has grown dramatically for at least two reasons. First, containers are an increasingly popular mechanism for distributing HPC software. Singularity [34, 45] and Shifter [3, 33], among others, now allow complex scientific applications to be shared more easily by packing them as self-sustained software images. However, container portability is hindered [26] by both the lack of a common launch methodology<sup>1</sup> and the absence of an MPI ABI – preventing the advent of portable containers featuring MPI programs. Second, MPI is now used by applications written in languages like Python, Julia, and Rust, which are currently required to build and test against all supported implementations and support the end-user installation of their MPI support against the implementation of the user's system. A standard ABI would eliminate the  $O(N)$  cost of packaging and simplify testing. The  $O(N)$  costs due to MPI implementation ABIs are not unique to these languages.

In the rest of this paper, we describe the constraints associated with an MPI ABI, the potential benefits for the HPC ecosystem, and the proposed ABI implementation as defined by the ABI working group. Performance experiments demonstrate that a high-quality implementation of the standard ABI in MPICH has negligible overhead, while the third-party implementation in Mukautuva has a tolerable overhead. We also discuss important considerations for compatibility besides the C ABI, including library naming, launchers, and Fortran.

## 2 BACKGROUND AND RELATED WORK

The HPC user community has been actively working to address the issue of ABI compatibility in MPI implementations. For a long time, the requirements associated with ABI compatibility in MPI have led to complexities in terms of software deployment, particularly in large computing centers.

Wi4MPI is a wrapper interface that implements ABI interoperability for MPI, supporting both Fortran and C languages [30]. It can be used in two ways. First, users can compile their applications against the generic MPI interface from Wi4MPI and then redirect them to their implementation of choice. Alternatively, they can redirect one implementation to another. For Wi4MPI to work, its wrapper interface needs to be compiled for each source and target MPI. The performance overhead has been shown to be minimal, making it an effective tool for running containers in a portable manner. Wi4MPI is leveraged in the Extreme-scale Scientific Software Stack [24] as a support tool in the e4s-c1 container launcher tool, which implements on-the-fly MPI detection and library translation at container launch time [37].

A similar effort to Wi4MPI was undertaken at the Perimeter Institute for Theoretical Physics, leading to MPItrampoline [35],

which defines its own ABI enabling applications' portability on several MPI runtimes.

It is known that a patent [41] exists for a specific method of interoperating different MPI ABIs, preventing its use by the open-source community.

In general, the availability of a standard ABI will simplify the tasks of these converters. Instead of having to implement conversions between the two APIs, these adaptation layers will primarily focus on compilation-related tasks, such as fixing dependency detection and enabling the replacement of one MPI with another.

After the MPI ABI working group was formed, two efforts were started to prototype the proposed designs, to understand their feasibility. The first of these was Mukautuva [22], which is a standalone ABI abstraction layer that maps from its own ABI (i.e. an approximation to the one under discussion in the working group) to MPICH and Open MPI by redirecting MPI symbols through a translation layer to the underlying MPI implementation, with renamed symbols (via `dlsym`) to avoid conflicts. The final design of Mukautuva is unintentionally quite similar to MPItrampoline; this convergence may be an indication of the suitability of their design. Meanwhile, a prototype was developed in MPICH [46]. Working together, these efforts revealed the relative ease of implementing the ABI proposal both internally and externally to an existing implementation. They also exposed non-portable assumptions in various MPI test suites.

## 3 CURRENT ABI DESIGNS

There are multiple aspects to an MPI ABI. Here are a few:

- (1) The integral types of `MPI_Aint`, `MPI_Offset`, and `MPI_Count`.
- (2) The `MPI_Status` object. This is a C struct with three standard members as well as hidden fields used by the implementation.
- (3) Opaque handles such as `MPI_Comm`. Implementations can define these to be anything that satisfies the required properties.
- (4) Callback functions, e.g., `MPI_User_function`. These callback functions usually do not allow registering any data with the function pointers, which is a challenge to intercepting and forwarding registered functions.
- (5) Values for both integer and handle constants, as well as predefined callbacks. Some of these are arbitrary, while others must be chosen carefully.

MPI 4.0 requires that most constants be usable in C for initialization and assignments, but not case statements, which means they need not be compile-time constants. Fortran requires they be compile-time constants, which constrains the C ABI when constants are the same in both languages. Buffer address constants cannot be used for initialization/assignment, while string length constants must be suitable as sizes in array declarations.

MPICH [10] has elected to provide compile-time constants, which is necessary on some operating systems that do not support link-time constants, and works in both C and Fortran. Open-MPI [13] does not have compile-time constant predefined handles in C, and has an indirection table from Fortran integer handles to the C ones.

<sup>1</sup>Note the PMIx standard [6] has made important progress on addressing this issue.

### 3.1 MPI integer types

The types `MPI_Aint` and `MPI_Offset` are used to store addresses and file offsets, respectively. `MPI_Count` was added in MPI-3 for the large-count effort, and this type is required to hold values of `MPI_Aint` and `MPI_Offset`, so it is at least as large as these. `MPI_Aint` is somewhat challenging since it must hold both absolute addresses and relative displacements of pointers, so it is similar to `(u)intptr_t` and `ptrdiff_t` from C. However, because it must also work in Fortran as `INTEGER(KIND=MPI_ADDRESS_KIND)`, it must be treated as if it is signed (because Fortran does not support unsigned integers). Another complication is that pointers, addresses, and differences of pointers may not always be the same size. In the past, segmented addressing meant that addresses could be larger than pointers, whereas there are now platforms where the reverse is true, and `MPI_Aint` must be able to hold a pointer [21] to support struct datatypes, for example.

### 3.2 The status object

This section describes multiple implementations of the `MPI_Status` object and their history.

**3.2.1 New MPICH (MPICH ABI Initiative).** Below is the status object in MPICH, which was made consistent with Intel MPI, in order to establish the MPICH ABI initiative. This meant that applications and libraries compiled against Intel MPI could be run using many implementations.

```
typedef struct MPI_Status {
    int count_lo;
    int count_hi_and_cancelled;
    int MPI_SOURCE;
    int MPI_TAG;
    int MPI_ERROR;
} MPI_Status;
```

**3.2.2 Old MPICH.** Prior to being consistent with Intel MPI, MPICH had the following status object. This definition included unused fields as a hedge against future needs, but also allowed for platform-specific fields, which meant that MPICH builds on different platforms could be ABI-incompatible.

```
...
typedef struct MPI_Status {
    int MPI_SOURCE;
    int MPI_TAG;
    int MPI_ERROR;
    MPI_Count count;
    int cancelled;
    int abi_slush_fund[2];
    @EXTRA_STATUS_DECL@
} MPI_Status;
```

**3.2.3 Open MPI.** The status object from recent versions of Open MPI is shown below. The status used by Wi4MPI has the same layout.

```
typedef struct ompi_status_public_t MPI_Status;
struct ompi_status_public_t {
    int MPI_SOURCE;
    int MPI_TAG;
    int MPI_ERROR;
    int _cancelled;
    size_t _ucount;
};
```

**3.2.4 MPItrampoline.** MPItrampoline defines a status object that holds the three public fields as well as a union of structs equivalent to the status objects of MPICH and Open MPI.

This definition is not space efficient but convenient for converting between the trampoline definition and the underlying implementation one, although it stores the public fields redundantly.

We see here that all variants have the required fields, `MPI_SOURCE`, `MPI_TAG` and `MPI_ERROR`, and the old MPICH ABI matched the Open MPI ABI in having both at least one bit for the canceled state plus a count field that supports at least 63 bit values. The question for ABI standardization is what sort of hidden fields may need to exist in the future, since there is little to no slack space to add new fields in the current implementations.

### 3.3 MPI handle types

MPI datatypes are opaque objects although the constraints on them limit the implementation choices. The MPI standard requires that opaque objects can be compared for equality and inequality. For the C language, this means that they need to have a built-in type, which reasonably only allows integer and pointer types, and excludes union and struct types.

The other important constraint on handles is related to attributes: *“Attributes in C are of type void\* [...] Attributes are scalar values, equal in size to, or larger than a C-language pointer. Attributes can always hold an MPI handle.”* Because MPI handles must be able to be held in a type `void*`, they cannot be larger than a pointer.

Since Fortran only supports signed integers, and older versions of C provide a limited set of integer types, one can expect implementations to use a 32-bit integer, a 64-bit integer, or a pointer for handles, although an 8- or 16-bit integer would be permitted. We see that MPICH uses a C `int` (32-bits on all supported platforms) and Open MPI uses incomplete struct pointers. The utility of incomplete struct pointers is that they allow for compiler type-checking. That is, `MPI_Comm` and `MPI_Group`, for example, are recognizable as different types and the compiler can issue warnings about invalid handle arguments. On the other hand, the MPICH design allows for zero-overhead conversion between C and Fortran, as well as the encoding of information in the handle values themselves. Open MPI does not utilize this capability since handles to C objects are not compile-time constants.

Below are some of the MPICH datatype handles, which reveal how information is encoded within them:

```
typedef int MPI_Datatype;
#define MPI_CHAR ((MPI_Datatype)0x4c000101)
#define MPI_SHORT ((MPI_Datatype)0x4c000203)
#define MPI_INT ((MPI_Datatype)0x4c000405)
#define MPI_LONG ((MPI_Datatype)0x4c000807)
#define MPI_FLOAT ((MPI_Datatype)0x4c00040a)
#define MPI_DOUBLE ((MPI_Datatype)0x4c00080b)
```

These handles encode the size of built-in datatypes that can be queried trivially with this macro:

```
#define MPIR_Datatype_get_basic_size(a) (((a)&0x0000ff00)>>8)
```

There are other macros that take advantage of the hidden structure of the `MPI_Datatype` handle that the reader can study in `mpir_datatype.h`.

Open MPI’s `mpi.h` defines the datatype handle to be a pointer to an incomplete struct, which is resolved externally at link-time. The definition of the structure is only visible when building the MPI library itself; otherwise, the compiler only knows its name. This means that the data pointed to by a handle need not be the

same at runtime, because the MPI application or library does not depend on it.

```
#define OMPI_PREDEFINED_GLOBAL(type, global) ((type) ((void *) &(global)))
...
typedef struct ompi_datatype_t *MPI_Datatype;
...
#define MPI_CHAR OMPI_PREDEFINED_GLOBAL(MPI_Datatype, ompi_mpi_char)
#define MPI_DOUBLE OMPI_PREDEFINED_GLOBAL(MPI_Datatype, ompi_mpi_double)
...
extern struct ompi_predefined_datatype_t ompi_mpi_char;
extern struct ompi_predefined_datatype_t ompi_mpi_double;
```

The runtime cost of querying handles is different in Open MPI relative to MPICH. Open MPI has to look up the size of the datatype inside of a 352-byte struct, which is not a concerning overhead since the type of MPI code that will notice such an overhead is going to pass the same datatype over and over, in which case the CPU is going to cache and correctly branch-predict the lookup and associated use every time.

```
static inline int32_t
opal_datatype_type_size(const opal_datatype_t *pData, size_t *size) {
    *size = pData->size;
    return 0;
}
```

Wi4MPI defines all the opaque handles to be `size_t`. This ensures they are at least as large as MPICH's `int` handles and Open MPI's pointer handles on most platforms (technically, `intptr_t` must be used for this to be strictly true but the exceptions are obscure [43]).

Wi4MPI defines the built-in datatypes to be sequential integers, which means they are not attempting to encode useful information the way MPICH do, although they are compile-time constants, unlike Open MPI.

```
/* C datatypes */
#define MPI_DATATYPE_NULL 0
#define MPI_BYTE 1
#define MPI_PACKED 2
#define MPI_CHAR 3
#define MPI_SHORT 4
#define MPI_INT 5
#define MPI_LONG 6
#define MPI_FLOAT 7
#define MPI_DOUBLE 8
```

MPITrampoline uses `uintptr_t` internally in its ABI, and incomplete struct pointers in its public API for type safety:

```
typedef struct MPITrampoline_Comm *MPI_Comm;
typedef struct MPITrampoline_Datatype *MPI_Datatype;
```

*Analysis.* There are advantages to both approaches. MPICH optimizes for the common case of built-in types, and does a lookup for others, while Open MPI always performs a pointer lookup, but then has what it needs in both cases.

The other advantage of the MPICH approach is with Fortran. In Fortran, handles are `INTEGER` or a type with a single member that is an `INTEGER`. MPICH conversions between C and Fortran are trivial. Open MPI has to maintain a lookup table to map Fortran handles to C objects.

An advantage of the Open MPI approach of using pointer types to represent opaque types is increased type safety. This enables the compiler to flag type mismatches, e.g. an `MPI_Comm` and an `MPI_Datatype` argument have accidentally been swapped.

### 3.4 Functions

Function prototypes in MPI follow naturally from the definitions of their arguments, which are either opaque handles, MPI integer

types, or intrinsic language types. What is essential for ABI purposes is that the calling convention be fixed. This can be done by specifying the aforementioned types and defining the calling convention to be “as if” compiled by the platform C compiler. In most cases, all of the C compilers on a given platform share a calling convention but there are at least historical cases where this was not true. As long as the MPI library uses the platform C compiler calling convention, it will be compatible with libraries and applications built with it, or another compatible compiler.

## 4 ECOSYSTEM IMPACT

One of the main motivations for an ABI is the ability to simplify the end user's life, thus improving the usability of the various MPI implementation through standardization. In this section, we detail particular points of interest for the community which would directly benefit from the availability of an ABI.

### 4.1 Python

The Python language provides MPI bindings through the `mpi4py` package [9]. `mpi4py` uses Cython [2], a super-set of the Python language with C extensions. The Cython compiler generates C code calling into the Python C-API and the MPI C-API. The wrapper C code has to be compiled and linked against a specific MPI implementation to generate a Python extension module.

The lack of a standardized MPI ABI presents several drawbacks. The `mpi4py` testing infrastructure built on publicly available services like GitHub Actions and Azure Pipelines requires adding both MPICH and Open MPI to the build matrix, effectively duplicating the required resources for running continuous integration. The `mpi4py` maintainers cannot distribute pre-built binary Python wheels via the Python Package Index, effectively forcing end users to set up a working C and Python development environment and build `mpi4py` from a sources distribution. The `conda-forge` [8] project somewhat alleviates these issues by featuring the `conda` package manager and its ability to install different variants of pre-built binaries in user-defined non-system locations. Nonetheless, the lack of a standardized MPI ABI prevents `conda-forge` binaries from using MPI implementations that are not ABI-compatible with either MPICH or Open MPI. In addition, `conda-forge` also suffers from the doubling of required resources to generate binaries for every downstream application or library using MPI.

A standardized MPI ABI would allow `mpi4py` to explore alternative implementations based on the runtime loading of dynamic/shared libraries and C foreign function interface (FFI) mechanisms. Such an approach would circumvent the generation of platform-specific binaries, allowing any pure Python code to access the MPI library and its features in a platform-agnostic way.

### 4.2 Julia

The Julia language provides MPI bindings through the `MPI.jl` package [5]. Unlike Python, Julia does not make use of a C compiler to call into external libraries. Instead, the user provides the corresponding types to the function signature to the `ccall` command. To support this, the developers of `MPI.jl` had to (a) define the constants and type definitions for each MPI ABI, (b) develop heuristics to detect which ABI a particular MPI library is using, and (c) provide

a mechanism to switch between the ABI definitions, invalidating Julia’s cache of pre-compiled code. This code has been a significant source of issues, hampering its usability and requiring significant engineering effort on the part of its volunteer maintainers.

A key benefit of a standardized ABI will be making it easier to provide downstream binaries. The Julia package manager provides prebuilt binaries of many MPI-enabled libraries, such as ADIOS2, HYPRE, P4est, and PETSc, but the support is rather cumbersome, especially when users wish to use non-bundled MPI implementations. The ABI would boost usability, especially for the long-tail of users on lower-end systems.

### 4.3 Rust

The Rust programming language provides MPI bindings through the libraries in the `rsmpi` project [38]. `rsmpi` combines a thin static library that re-exports underspecified identifiers (providing symbols where MPI implementations are allowed to use macros) and uses `bindgen` to create the raw Rust interface. `bindgen` relies on `libclang` to parse the header (of the thin library and `mpi.h`) to generate Rust bindings that conform to the C ABI of the MPI implementation at hand. This approach works for any compliant MPI and does not require tedious definition or maintenance of stubs. The disadvantages are many and include long initial compilation times due to having to fetch and build the dozens of dependencies of `bindgen`, a need for pre-installed `libclang`, long testing times by building against multiple different MPI implementations, and need for users to rebuild to pick a different MPI implementation.

Rust is known for reliable package management and tooling for binary distribution, including cross-compilation (across OS and ISA). A standardized MPI ABI would allow `rsmpi` to provide a thin and stable Rust binding that can be built without any dependencies and simply links against a dynamic library on the target platform. For example, a single CI/CD job for an application could publish binaries for MacOS (x86-64 and ARM), Android, Windows, and Linux (x86-64, POWER, and ARM) without needing to think about MPI idiosyncrasies. ABI stability would make it easier to support more MPI features and would also enable low-level idiomatic Rust features that improve safety and static analysis for C FFI bindings, but that are hard to incorporate into the current `bindgen` approach.

### 4.4 Fortran

Currently, all implementations of MPI Fortran wrappers are integrated with MPI implementations. Vapaa [20] is the first attempt to write the MPI Fortran 2008 interface as a standalone project, based on calling the C interface, without any use of the internal state. Because Fortran constants must be compile-time constants, not just link-time constants, when Open MPI is used, Fortran interfaces must define their own set of constants and translate them to the C ones at runtime. Furthermore, to handle status objects, the status ABI must be known. Thus, Vapaa ends up implementing its own ABI and translating all constants and status objects. A standard ABI would simplify the translation process and eliminate the need for implementation-specific status handling. If the Fortran 2008 interface had the type of `MPI_VAL` equivalent to handle types in C, then no translation would be necessary. However, this would

be both an ABI and an API change for `mpi_f08.mod`; it also offers nothing to users of older MPI Fortran interfaces.

In addition, due to differences in name mangling among Fortran compilers, running a Fortran program that calls MPI functions inside a container can result in linker dependencies that are not fixed and depend on the specific Fortran module and compiler convention used. This prevents an MPI Fortran program from having its MPI implementation replaced through interposition (i.e., `LD_PRELOAD`). Having an external Fortran implementation that relies on the ABI would enable the static linking of the Fortran adaptation layer in the target binary, abstracting away from these language-dependent variations and restoring ABI compatibility.

### 4.5 Packaging

The availability of an ABI is highly important for packaging MPI applications. MPI is a fundamental package for most scientific software. However, there are several libraries that provide MPI support, including vendor-specific ones. Therefore, building a binary for MPI can become cumbersome when managing a long chain of dependencies between packages, resulting in repetitive building. While the ABI alone is insufficient to solve packaging dependency issues, it is a significant step in the right direction. When running an executable, the loader is responsible for locating various dynamic shared objects (DSOs) to fulfill execution dependencies, based on the system’s configuration and environment (e.g., `LD_LIBRARY_PATH` and search paths). Thus, the ABI alone has no impact on the loading of libraries when running a program. Nonetheless, it is still a crucial step towards achieving drop-in replacement for MPI, which involves changing the MPI of a given binary. This goal can be attained by defining a common library naming scheme or developing specific stub libraries in charge of bridging implementations – the binary being linked to the stubs.

Linux package managers such as APT and RPM ship binaries for two dependency chains, with packages like `hdf5-openmpi` and `hdf5-mpich`. APT manages these through the `/etc/alternatives` mechanism while RPM (Fedora) deliberately rejects that usage due to their rule [12]: “If a non-root user would gain value by switching between the variants then alternatives MUST NOT be used.” As such, Debian/Ubuntu users developing code that depends on an MPI-enabled HDF5 can have a default implementation (that might change unexpectedly at the whims of their sysadmin) while Fedora/Red Hat users must use verbose paths unlikely to be found by configure scripts. In this example, Arch Linux provides `hdf5-openmpi` in the default repository (binary distribution) while `hdf5-mpich` (and `mpich` itself) are in the user repository that must be installed from source. Homebrew provides only `hdf5-mpi`, which uses Open MPI. This complexity requires maintenance and communication to users, increases testing time and frequency of bugs, and harms reproducibility. With an ABI, there could be one `hdf5-mpi` and let the `mpiexec` or `ldconfig/LD_LIBRARY_PATH` (since these distributions eschew `RPATH`) determine which MPI to execute with.

A possibility to retarget binaries would be changing embedded `RPATHs` inside the executable. `Spack` [14], which relies extensively on `RPATHs` has implemented such rewriting techniques [44] when deploying binary packages to systems, relocating the `RPATHs`

through carefully planned compilation and clever rewriting techniques. This approach was required as the end-user may have deployed his spack tree at a different location from where it was initially compiled. By doing so, Spack then manages to restructure complex dependency chains, it is a process analogous to what would be required to change the runtime of a given MPI-enabled binary. The Anaconda [27] software distribution and its conda package manager also rely on RPATH rewriting to allow binary relocation.

#### 4.6 Testing

The advantage of the ABI support for testing is less direct. Indeed, when running a program with a given MPI, the MPI is also part of the equation to be validated. Gains could be envisioned in the building of the test cases but it is not obvious that MPI implementation will be directly interoperable even in the event of a unified ABI. Indeed, as for packaging, the build chain for MPI involves compiler wrappers and implementations are free to choose the name of the MPI Dynamic Shared Object (DSO) – preventing drop-in replacement of MPI. Overall, for testing, the ability to retarget MPI programs to another implementation requires more than an ABI but either a clearly defined object layout for MPI or a dedicated redirection layer similar to what is provided in “trampoline” interfaces [30, 35]. Analogous dependencies on the naming of the DSO are also present for containers as discussed in the next section.

#### 4.7 Containers

Containers are an abstraction built on Linux namespaces. Containers are the systematic use of such namespaces to run “software images” in a custom environment. The main advantage of containers is the ability to move images around systems to avoid recreating complex software environments. In HPC networks, it is common to use OS bypass techniques to optimize network performance. This involves creating a separate networking layer that operates independently of the operating system, allowing for faster and more efficient communication between nodes. As a result, the networking namespace (IP bound) is not commonly used in HPC networks which prefer faster fabrics than TCP/IP. Similarly, to prevent security issues such as privilege escalation, the user namespace, which allows mimicking root behavior inside the container, is not used for containers in non-virtualized environments[7]. Overall, the namespace leveraged by HPC container runtimes such as Singularity and Shifter is the mount namespace.

Using the mount namespace, it is possible to change the mount point seen by the running program. In HPC, the user’s home is often bind-mounted inside the container, stacking the container’s view of the file system atop of the preexisting shared one. With a containerized program compiled against MPI, the corresponding MPI is likely present in the container image. This MPI has to be compatible with a wide range of interconnects, unlike the host MPI from the system, since it cannot anticipate its target environment. While this can be mitigated with communication libraries such as libfabric[16], which manage to unify high-speed network interfaces, there may be features that are not possible outside of the native MPI environment, perhaps because they are proprietary and not generally available, or because they require system awareness

(e.g., network topology information) that cannot be included in the widely distributed implementations of MPI.

On this aspect, using the host MPI (as opposed to a container MPI) would allow the guest binary to take advantage of all the custom features of the system. It also obviates the need for application containers to redistribute MPI at all. For this purpose, approaches such as e4s-cl [37] recursively locate all dependencies of MPI and inject them into the container. The target binary may depend indirectly on libraries such as hwloc that are required by MPI. There are ways to mitigate this issue, like embedding all dependencies or using symbol versioning for standard HPC libraries to reduce possibilities for symbol conflicts. Another method is to ensure that MPI-shared libraries do not cause transitive dependencies so that the binary only requires MPI, and the MPI implementation takes care of its dependencies directly. A second challenge faced by the MPI container is launching the application. Indeed, MPI is in many cases relying on the Process Management Interface (PMI) to wire up its processes, which also has to be mapped into the container. There have been studies on this point as part of a complete rework of this interface in the PMIx standard [6] enabling PMI portability.

To summarize, containers need support from MPI to allow binary retargeting, i.e., the ability to change the MPI implementation on a binary compiled against another MPI implementation. Note that changing the guest MPI to the host MPI also allows PMI disambiguation – removing complexity on the launch side. Having an ABI is compulsory as retargeting does not allow recompilation of the application. This last point is the main blocker for accepting and distributing MPI containers.

#### 4.8 Performance and Debugging Tools

MPI tools often use the profiling interface (PMPI) to intercept function calls and extract the current state of MPI and to time operations, for performance and debugging purposes [25, 28]. Since this interception operates on the compiled library code, all MPI tools must be compiled against the relevant implementation ABIs. A standard ABI makes it possible for PMPI interposition tools to be compiled only once and reused with different MPI implementations.

The Tools working group in the MPI Forum is working on the QMPI interface [11]. This interface is designed to support multi-instrumentation, mimicking what has been previously pioneered with  $P^n$ MPI [36]. As with PMPI, the absence of a standard ABI requires each QMPI tool to be compiled for every ABI, and potentially more, if any of these tools modify ABI-related properties of the interface. One of the advantages of the proposed status object for the standard ABI is that it has additional space that allows tools to hide state in the reserved fields. Managing this in a layered context is complicated and is left as an exercise for the implementers of such tools.

### 5 PROPOSAL

This section outlines the current proposal for the MPI standard ABI, based upon detailed analysis of requirements from MPI as well as the behavior of platforms MPI should support. Following Section 3, the ABI proposal defines MPI integer types, the status object, opaque handles, and constant values. The calling convention must be equivalent to the platform C compiler.

## 5.1 MPI integer types

The purpose of `MPI_Aint` is to hold addresses or pointers, whichever is larger, because its usage requires both. It should also be signed because Fortran does not support unsigned integers. The only standard C type that meets this requirement is `intptr_t`. It is necessary to use this integer type on platforms with so-called “wide pointers” [43], although this situation is rare. There is no C integer type associated with filesystem offsets, but all modern systems should use at least 64-bit integers. There are some platforms where the underlying filesystem offset may be 128-bits, but there is no need for MPI to define `MPI_Offset` this way since MPI files greater than 8 EiB are unlikely.<sup>2</sup> Additionally, 128-bit integers are not implemented natively on most systems and thus may perform poorly, so it is undesirable to force the use of 128-bit integers for offset and count to support impossibly large files. On the other hand, most systems with 32-bit addressing have 64-bit filesystems, so there are at least some scenarios where the MPI ABI should be flexible enough to support different sizes of address and offset types.

To ensure all relevant target platforms can be supported, the MPI ABI should be described in terms of the size of `MPI_Aint` and `MPI_Offset`, while `MPI_Count` matches the larger of these two (which will be `MPI_Offset` on most systems). The integer sizes of the MPI ABI can be denoted  $An_{Address}On_{Offset}$  to denote the number of bits in the `MPI_Aint` and `MPI_Offset` types, respectively. This is similar to how platform ABIs are described using  $In_iLn_lLLn_{ll}Pn_p$  notation, to denote the size of C `int`, `long`, `long long`, and `void*`, respectively. For example, modern Linux platforms are described as LP64, meaning that `long` and `void*` are 64-bit. Today, essentially all MPI ABIs are A32O64 or A64O64 ABIs, because we have 32- or 64-bit addresses, but most filesystems are 64-bit. An A64O128 ABI is possible, although, for the aforementioned reasons, it is neither necessary nor desirable.

The potential for more than one MPI ABI on a given platform is undesirable. Current trends in filesystem technology suggest that a `MPI_Offset` larger than 64 bits will not be necessary for at least 20 years. For these reasons, we propose to prescribe the MPI ABI for platforms with 32- and 64-bit pointers as follows:

```
typedef intptr_t MPI_Aint;
typedef int64_t MPI_Offset;
typedef int64_t MPI_Count;
```

This ABI definition covers essentially all relevant platforms since the introduction of LFS [1] until the availability of filesystems far in excess of 8 exabytes. These types are part of C99 and C++11 but implementations can use older equivalents for compiler portability.

One may observe that `intptr_t` is optional in C and, in theory, a system may lack an integer type capable of satisfying its requirements. This is uncommon, and exists to accommodate systems with 128-bit pointers but where supporting `intptr_t` would force a change in `intmax_t`, which would be a breaking change in the platform ABI [17]. We note that MPICH requires `intptr_t` and platforms that do not provide it are not supported, so a reasonable portion of the MPI ecosystem is unconcerned with this situation.

<sup>2</sup>At current prices of \$10/TB, one such file would require more than \$90M in filesystem hardware.

While we have considered the case of 128-bit pointers, the current proposal will only include A32O64 and A64O64. It is appropriate for the MPI community to gain more experience with such platforms before attempting to standardize for them. For example, while CHERI [42] has 128-bit pointers but doesn’t necessarily require 128-bit file offsets, but if `MPI_Aint` and therefore `MPI_Count` have to be 128b, it might be prudent to make offsets the same width, so that there is only one MPI ABI for all 128-bit platforms.

The one MPI integer that cannot be prescribed like the others is `MPI_Fint`, since this corresponds to Fortran `INTEGER`, which is not fixed, but varies as a function of Fortran compiler flags. It seems appropriate to have a runtime query to allow C code to know the size of a Fortran integer and work with it appropriately. This requires code changes compared to the current situation where `MPI_Fint` is known at compile-time, but the C code that relies on this is rare. Alternatively, the standard ABI could force `MPI_Fint` to be a C `int`, and disallow MPI Fortran interfaces from supporting larger integer sizes. This would please Fortran purists who loathe the compiler feature that allows changing the Fortran default integer size, but displease users of existing implementations that support it.

## 5.2 The status object

The proposed standard status object is:

```
typedef struct MPI_Status {
    int MPI_SOURCE;
    int MPI_TAG;
    int MPI_ERROR;
    int mpi_reserved[5];
} MPI_Status;
```

This object is 32 bytes in size, which leads to good alignment when arrays of statuses are used, and includes at least two extra fields more than current implementations.

## 5.3 Handle types

In order to have type-safety in handles, incomplete struct pointers are proposed; Open MPI has used this design and its properties are well understood. The incomplete struct name will become part of the ABI, so that compiler warning messages are clear:

```
typedef struct MPI_ABI_Comm * MPI_Comm;
typedef struct MPI_ABI_Request * MPI_Request;
```

## 5.4 Constants

Constants in MPI come in different forms. They include:

- Error codes, which start with `MPI_SUCCESS=0`.
- Buffer address constants, e.g., `MPI_BOTTOM`, which must have special values distinguishable from user buffers.
- Handle constants.
- Integer constants that must have special values to avoid conflicts; for example, `MPI_ANY_SOURCE` can never be a valid rank, and thus should be a negative number.
- Integer constants that must be powers of two, to support combination using XOR.
- Integer constants that correspond to string lengths.
- Integer constants that can have any value.
- Predefined attribute callback functions.

Some of the desirable properties brought forth by users and implementers include a desire for unique integer constants, so that

errors can be identified precisely. For example, if a user passes `MPI_ANY_TAG` as a rank, this can be identified precisely if the constant value is unique with respect to all other constants, especially `MPI_ANY_SOURCE`. Another desirable property is the ability to encode information in handle constants, as MPICH does. For maximum portability, integer constants cannot be larger than 32767, because that is the largest value of type `int` guaranteed by the C standard. This constraint is strictly academic for the relevant systems but there was no reason to violate it either.

For handle constants, the working group discussed designs with and without unique values as well as the use of one or more lookup tables versus a Huffman code. The current proposal uses a Huffman code but is sufficiently compact so as to require a relatively small lookup table, for implementations that choose to use one. The Huffman code uses 10 bits and therefore fits into the zero page of common operating systems; as a result, implementations that allocate user handles from the heap need not verify that they do not conflict with predefined constants.

As datatypes make up the majority of MPI's predefined handles, half of the Huffman code bits are reserved for datatypes, although less than 100 values are used. The language, numerical properties, and sizes of all *fixed-size* types are encoded in the handles. For example, `MPI_CHAR` can be determined to be a 1-byte type immediately. Unfortunately, `MPI_INT` is not a fixed-size type, so its size is not encoded, as that would mean that the constant value was a function of the platform ABI. While it would be possible for some use cases to handle this, it is undesirable to force higher-level languages like Julia to determine the platform ABI in order to use MPI.

Other handles can be decoded quickly using the bit pattern alone. The value zero is always an invalid handle, which allows uninitialized handles to be detected as errors instead of being confused as legal null handles. Legal null handles use the non-zero bits of the handle kind followed by zeros. The current Huffman code has a sufficient amount of free space to allow for many new handle types and new handle constants for existing types to be added, without requiring breaking changes.

The values of integer constants for string lengths, e.g., `MPI_MAX_LIBRARY_VERSION_STRING`, and constants that can be combined with XOR, e.g., `MPI_MODE_NOCHECK`, are not particularly interesting. For the former, the largest known values used in existing implementations were chosen. There was some concern that stack allocation of 8192 bytes could be a problem, but (1) nothing prevents users from allocating such strings on the heap and (2) no issues with this value (used by MPICH) have ever been reported.

The other integer constants are unique negative numbers, which means that implementation can tell the user by name what constant they passed, when the user passes an incorrect constant.

For simplicity, predefined attribute callbacks were set to `0x0` for `MPI_XXX_NULL_COPY_FN` and `MPI_XXX_NULL_DELETE_FN`, and `0xD` for `MPI_XXX_DUP_FN`. Since compilers can detect incompatible function pointer arguments there is little need to detect errors at runtime.

The encoding of operation handles is provided in Appendix A.1. The gaps in the ranges for the different operation types are intentional since they provide room for future extensions. Moreover, the modified Huffman encoding enables fast error checking by implementations, simply by applying a bitmask.

Handles for opaque objects are encoded in a similar way, as shown in Appendix A.2. The encoding leaves room for future extensions for each handle type, making it possible to add new handles without requiring special case handling.

Examples of datatype handles are provided in Appendix A.3. Types with variable size (e.g., C `int`, `float`) are encoded with the prefix `0b1000XXXXXX`. Fixed-size types are encoded with the prefix `0b1001XXXXXX`, with the size encoded in the lower bits at position 4–6. For example, types with size 1 are encoded with prefix `0b1001000XXX` (e.g., `MPI_BYTE` with `0b1001000111`; size  $2^{000b}$ ) while types with size 4 are encoded with prefix `0b1001010XXX` (e.g., `MPI_INT32_T` with `0b1001010000` and size  $2^{010b} = 2^2$ ).

The full definition of the Huffman code for handle constants can be found in [18], while the other constants are listed in [19].

## 6 EXPERIMENTS

In this section, we present three experiments regarding the implementation of the standard ABI. First, we measure the performance impact of different ABIs for querying the size of a type. Second, we measure the message rate for MPICH-based implementations, with and without standard ABI support. Third, we describe Mukautuva, which demonstrates the feasibility of implementing the standard ABI outside of any existing implementation. Finally, we mention the effort required to implement the standard ABI in MPICH. For both implementations – the one outside of an MPI implementation and the one within MPICH – we see that the cost of ABI translation is small.

### 6.1 Performance

Historically, there has been a performance argument in favor of MPICH's integer handles for datatypes because information like type size is encoded directly in handles, whereas with Open MPI, it must be fetched from the internal state. We measured the throughput of `MPI_Type_size` to be  $\approx 11.5$  nanoseconds with both implementations on an AMD EPYC 7413 CPU. Not only is the difference between the two implementations negligible, both are negligible relative to the network cost of sending a single message, which is at least 500 nanoseconds.

Table 1 shows the message rate determined by the OSU MPI Benchmarks 7.0.1 for three different builds of MPICH: the latest Intel MPI and MPICH development versions built with UCX using the MPICH ABI<sup>3</sup> and the standard ABI prototype<sup>4</sup>, with and without Mukautuva. We see that adding the indirection from Mukautuva has a noticeable impact, but it is likely acceptable as a worst-case implementation of the standard ABI.

### 6.2 Mukautuva

Mukautuva [22] ("Adaptable" in Finnish) was created both as an ABI compatibility layer and as a way to prototype the ABI proposal being developed for the MPI Forum. It represents a worst-case scenario implementation for the standard ABI, if implementers refuse to support it.

Mukautuva (MUK) consists of two shared libraries. The first library provides the MPI interface symbols. The second library is

<sup>3</sup><https://github.com/pmodels/mpich/commit/0f72280c70214411e00e5c03e1f4111972fab4b7>

<sup>4</sup>[https://github.com/hzhou/mpich/tree/2302\\_abi\\_4ce7cd3b6c6a92e99090957ae1c4a9166efadedf](https://github.com/hzhou/mpich/tree/2302_abi_4ce7cd3b6c6a92e99090957ae1c4a9166efadedf)

**Table 1: Message rate (8-byte messages) determined by osu\_mbw\_mr on an Intel i7-1165G7 CPU running Linux 5.19.0 (Ubuntu 22.04). Build options unrelated to ABI – the shared-memory performance of UCX versus OFI – have a significant impact on message rate. The MPICH dev UCX results show no difference between the MPICH ABI and the proposed standard ABI.**

| MPI                   | Messages/second |
|-----------------------|-----------------|
| Intel MPI 2021.9.0    | 4658939.64      |
| + Mukautuva           | 4606473.95      |
| MPICH dev UCX [1]     | 13643117.42     |
| + Mukautuva           | 12278837.03     |
| MPICH dev UCX ABI [2] | 13643378.98     |

1. `-enable-error-checking=no -enable-fast=0s -enable-g=none -with-device=ch4:ucx`
2. Same as 1 plus `-enable-mpi-abi`

compiled against MPICH or Open MPI and provides the underlying implementation (IMPL). Applications compiled against MUK are relying on its ABI, which is a proxy for a future MPI standard ABI. At runtime, the first shared library determines which implementation will be used, and activates it via `dlopen` and `dlsym`. MPI symbols call a wrapper layer with the MUK namespace. MUK symbols are function pointers to the WRAP namespace in the implementation-specific shared library. WRAP functions call the implementation, with the appropriate conversion of handles and constants. An excerpt for the case of `MPI_Comm_size` follows.

```
libmuk.so:

typedef union {
    void * p; // Open-MPI
    int i; // MPICH
    intptr_t ip;
} MUK_Handle;
typedef MUK_Handle MPI_Comm;

// during initialization
...
MUK_Comm_size = MUK_DLSYM(wrap_so_handle, "WRAP_Comm_size"); // wraps dlsym()
...

int MPI_Comm_size(MPI_Comm comm, int * size) {
    return MUK_Comm_size(comm, size);
}

impl-wrap.so:

#include <mpi.h> // implementation details

static inline MPI_Comm CONVERT_MPI_Comm(WRAP_Comm comm) {
    if (comm.ip == (intptr_t)MUK_COMM_WORLD) { return MPI_COMM_WORLD; } else
    if (comm.ip == (intptr_t)MUK_COMM_SELF) { return MPI_COMM_SELF; } else
    if (comm.ip == (intptr_t)MUK_COMM_NULL) { return MPI_COMM_NULL; } else
    {
#ifdef MPICH
        return comm.i;
#elif OPEN_MPI
        return comm.p;
#else
#error NO ABI
#endif
    }
}

// success is the common case, so static inline it.
int ERROR_CODE_IMPL_TO_MUK(int error_c);
static inline int RETURN_CODE_IMPL_TO_MUK(int error_c) {
    if (error_c == 0) return 0;
    return ERROR_CODE_IMPL_TO_MUK(error_c);
}
```

```
int WRAP_Comm_size(WRAP_Comm comm, int *size) {
    MPI_Comm impl_comm = CONVERT_MPI_Comm(comm);
    int rc = IMPL_Comm_size(impl_comm, size);
    return RETURN_CODE_IMPL_TO_MUK(rc);
}
```

The vast majority of MPI features can be translated from one ABI to another with trivial overhead. The exceptions to this come in two forms: first, when callbacks are involved, and second, when vectors of handles are required. For callbacks, MUK must translate to IMPL handles to call IMPL functions, but then translate IMPL handles back to MUK handles, because the callback functions compiled as user code utilize the MUK ABI. The callback interfaces do not always make this easy, but it can be done in all cases, using methods described in the `README.md`. The situation with vector arguments is similar to [23], where vectors of datatype handles must be converted from one ABI to another, and freed upon completion, which is tricky in the case of nonblocking `alltoallw` operations. For these cases, like with callbacks, we use a map, currently implemented with `std::map` from the C++ standard library, to associate a temporary state with a handle. Callback function trampolines or request completion operations lookup the temporary state associated with handles when needed. The worst-case overhead will arise when the user has initiated a nonblocking `alltoallw` operation, followed by a large number of nonblocking point-to-point operations to be completed via `MPI_Testall`, for example. In this case, every call to `MPI_Testall` will look up every request in the map associated with nonblocking `alltoallw` operations. This is not currently optimized, due to the low probability of such a scenario in real applications.

During the development of MUK, we identified flaws in the early ABI proposals as well as in MPI test suites. The MPICH test suite, for example, assumed the MPICH ABI in many places<sup>5</sup>, which meant that it could not be used to test other implementations, or ABI translation layers such as `Wi4MPI`, `MPITrampoline`, and `MUK`. Most if not all of these issues have been resolved in the meantime.

MUK now passes all of the MPICH test suite tests except for a handful that uses dynamic process management, which appears to be related to environmental problems, yet to be investigated. MUK also passes all tests associated with `ARMCI-MPI`, the `Intel MPI Benchmarks (IMB)`, and the `OSU MPI Benchmarks (OMB)`. It complies with MPI-4 except for sessions, which are expected to suffer from the same issues observed in dynamic process management functionality. Calling functions before initialization or after finalization is not fully supported, and will be fixed in the future.

### 6.3 MPICH

While it has been demonstrated that the standard ABI can be implemented without any change to existing implementations, doing the translation inside of an MPI implementation has lower overheads. Hui Zhou has implemented support for the standard ABI in MPICH [46]. The changes consist primarily of abstracting away prior assumptions about the types of handles and callback signatures and inserting the appropriate conversions, where necessary. Most of the changes are in the interface code generator or guarded by a preprocessor token, hence having no impact on execution time.

<sup>5</sup>e.g. <https://github.com/pmodels/mpich/issues/6398>

The most expensive conversions are for datatype and reduce ops, with a worst-case that requires  $O(N_{\text{predefined}})$  comparisons.

## 7 OTHER CONSIDERATIONS

A standard ABI is necessary but insufficient to provide seamless compatibility of MPI software across implementations. For example, MPI applications often require a parallel launcher, e.g., `mpiexec`, which is not part of the ABI, but interacts with the MPI program in non-standard ways, such as environment variables.

There are at least two solutions for portable launching. The first is that the launcher determines the MPI shared library to be used, in which case the launcher and the library will be compatible. Another is the use of a launcher that is supported by multiple MPI implementations, such as the ones provided by popular schedulers like SLURM and PBS.

Applications also need to know what shared library to use. As `libmpi.so` is used by a number of implementations already, the name `libmpi_abi.so` is proposed for implementations of the standard ABI. Standardizing a new, descriptive name is especially important since it is expected that implementations will continue to support their existing ABIs, using the existing library name(s). It is expected that `libmpi_abi.so` will follow the platform-specific conventions for versioning to allow for future – hopefully backwards-compatible – changes.

Obviously, much of the MPI ABI is contained in the header file, `mpi.h`. The same filename will be used for the standard ABI, to ensure source compatibility, but applications must use exactly one ABI, and therefore every component of an application will need to be compiled against the same header. We expect that the standard ABI will be implemented in a header file provided by the MPI Forum that can be used with any implementation that supports the standard ABI, to ensure consistency in its definition. Implementations can provide this header in a different path from their own header, and perhaps help users with appropriate `pkg-config` definitions or compiler wrapper scripts, e.g., `mpi_abi.pc` or `mpicc_abi`, but these aspects of MPI are not standardized, nor are they part of the ABI.

### 7.1 Fortran

This paper focused on a standard C ABI for MPI, but many codes use MPI from Fortran. Fortran presents its own ABI challenges, not the least of which is that `INTEGER`, used for MPI handles in `mpif.h` and `mpi.mod` (and the `MPI_VAL` in typed handles defined by `mpi_f08.mod`) varies in size depending on compiler flags. Furthermore, each Fortran compiler has its own ABI and each has their own runtime library, in contrast to C, where it is common for C compilers to reuse the system C runtime, and thus be ABI compatible (e.g., Intel and GCC on Linux).

The current ABI proposal for Fortran follows the C one; many constants are required to be the same in both languages anyways. While Fortran handles may be too small to hold the C handle values in general, implementations can optimize for the case of predefined handles because the C constants will be representable in Fortran integers and do not require a translation table.

The overhead of translation for user-defined handles could be achieved with a new implementation of `mpi_f08.mod`, where `MPI_VAL`

is `INTEGER(kind=c_intptr_t)`, although this is a breaking change and would require a new module, which could be called `mpi_f08_abi`. One could also imagine a module `mpi_abi` that requires handles be `INTEGER(kind=c_intptr_t)`. No new MPI Fortran interfaces or modules are currently proposed.

## 8 CONCLUSIONS

We have reviewed the current practices for MPI ABIs in the popular implementations, MPICH and Open MPI, as well as ABI abstraction layers like Wi4MPI and MPItrampoline. The motivations for standardizing an MPI ABI come from multiple sources, including the packaging and distribution of MPI applications and libraries in binary form, the use of MPI from languages other than C (or C++), and the development of implementation-agnostic MPI performance and debugging tools. The MPI ABI working group has developed a proposal for a standard MPI ABI, which satisfies all of the requirements and relies only on ISO C language features. The standard ABI has been prototyped in both MPICH and Mukautuva, and is determined to be both practical and performant. We identified issues with compatibility and portability not related to the ABI that are expected to be solved by the MPI ecosystem.

The next steps for the proposed MPI ABI are (1) it must be standardized by the MPI Forum, (2) it must be implemented either by the major implementations and/or ABI abstraction layers. (3) users of MPI must recompile against the standard ABI. Work towards 1 is underway, and this paper has provided sufficient evidence that 2 is either already done or straightforward.

We cannot predict the behavior of all MPI users, and certainly, some may be reluctant, either because they expect the MPI standard ABI to be less reliable than existing ABIs or that it will break in the near future. There is obviously a large one-time cost of recompiling everything against the MPI ABI, but it is no worse than the cost of compiling everything against a new major release of Open MPI, for example. Fortunately, there is no immediate need for users to adopt the MPI ABI. It is expected that both MPICH and Open MPI will support their existing ABIs for as long as users require them, and will consider a translation to using the standard ABI natively only after there is sufficient understanding of its use across a wide range of platforms.

Regardless of how long it takes to realize the full potential of a standard ABI, we expect that it will significantly reduce the pain of using MPI in a variety of contexts, and encourage greater use of MPI in new domains.

## ACKNOWLEDGMENTS

The authors are indebted to numerous people who contributed to the MPI ABI standardization effort, including, but not limited to, Brian Barrett, Jim Dinan, Dan Holmes, Julien Jaeger, Quincey Koziol, Ken Raffanetti, Jeff Squyres, and Bill Williams. Simon Cooksey was a valuable resource regarding the ChERI architecture, which exposes some of the sharper edges of the C standard regarding pointers and integer constants. Jeff thanks Tim Costa for encouraging him to work on ABI standardization in the first place. Jed's work was partially supported by the U.S. Department of Energy, Office of Science, Office of Advanced Scientific Computing Research, applied mathematics program. Research at Perimeter Institute is supported

in part by the Government of Canada through the Department of Innovation, Science and Economic Development and by the Province of Ontario through the Ministry of Colleges and Universities. This work was supported partly by the Exascale Computing Project (17-SC-20-SC), a collaborative effort of the U.S. Department of Energy Office of Science and the National Nuclear Security Administration.

## REFERENCES

- [1] [n. d.]. Large-file support. [https://en.wikipedia.org/wiki/Large-file\\_support](https://en.wikipedia.org/wiki/Large-file_support) Accessed: April 28, 2023.
- [2] Stefan Behnel, Robert Bradshaw, Craig Citro, Lisandro Dalcin, Dag Sverre Seljebotn, and Kurt Smith. 2011. Cython: The Best of Both Worlds. *Computing in Science & Engineering* 13, 02 (2011), 31–39. <https://doi.org/10.1109/MCSE.2010.118>
- [3] Lucas Benedictic, Felipe A Cruz, Alberto Madonna, and Kean Mariotti. 2017. Portable, high-performance containers for HPC. *CoRR* abs/1704.03383 (2017), 11 pages. arXiv:1704.03383 <http://arxiv.org/abs/1704.03383>
- [4] Andrey Bychkov and Vsevolod Nikolskiy. 2021. Rust Language for Supercomputing Applications. In *Supercomputing*, Vladimir Voevodin and Sergey Sobolev (Eds.). Springer International Publishing, Cham, 391–403.
- [5] Simon Byrne, Lucas C Wilcox, and Valentin Churavy. 2021. MPIjl: Julia bindings for the Message Passing Interface. In *Proceedings of the JuliaCon Conferences* 1, 1, 68. <https://doi.org/10.21105/jcon.00068>
- [6] Ralph H Castain, David Solt, Joshua Hursey, and Aurelien Bouteiller. 2017. PMIx: process management for exascale environments. In *Proceedings of the 24th European MPI Users' Group Meeting*, 1–10.
- [7] CEA-HPC. [n. d.]. Private Cloud On a Compute Cluster (PCOCC). <https://github.com/cea-hpc/pcocc>. Accessed 08/2023.
- [8] conda-forge community. 2015. The conda-forge Project: Community-based Software Distribution Built on the conda Package Format and Ecosystem. <https://doi.org/10.5281/zenodo.4774216>
- [9] Lisandro Dalcin and Yao-Lung L Fang. 2021. mpi4py: Status update after 12 years of development. *Computing in Science & Engineering* 23, 4 (2021), 47–54.
- [10] MPICH developers. [n. d.]. MPICH: High-Performance Portable MPI. <https://www.mpi4py.org>
- [11] Bengisu Elis, Dai Yang, and Martin Schulz. 2019. QMPI: A next generation MPI profiling interface for modern HPC platforms. In *Proceedings of the 26th European MPI Users' Group Meeting*, 1–10.
- [12] Fedora. 2023. Alternatives: Usage within Fedora. <https://docs.fedoraproject.org/en-US/packaging-guidelines/Alternatives/>
- [13] Edgar Gabriel, Graham E. Fagg, George Bosilca, Thara Angskun, Jack J. Dongarra, Jeffrey M. Squyres, Vishal Sahay, Prabhakaran Kambadur, Brian Barrett, Andrew Lumsdaine, Ralph H. Castain, David J. Daniel, Richard L. Graham, and Timothy S. Woodall. 2004. Open MPI: Goals, Concept, and Design of a Next Generation MPI Implementation. In *Proceedings, 11th European PVM/MPI Users' Group Meeting*. Budapest, Hungary, 97–104.
- [14] Todd Gamblin, Matthew LeGendre, Michael R Collette, Gregory L Lee, Adam Moody, Bronis R De Supinski, and Scott Futral. 2015. The Spack package manager: bringing order to HPC software chaos. In *Proceedings of the International Conference for High Performance Computing, Networking, Storage and Analysis*, 1–12.
- [15] William Gropp. 2002. Building Library Components That Can Use Any MPI Implementation. In *Proceedings of the 9th European PVM/MPI Users' Group Meeting on Recent Advances in Parallel Virtual Machine and Message Passing Interface*. Springer-Verlag, Berlin, Heidelberg, 280–287.
- [16] OFI Working Group. [n. d.]. Libfabric. <https://ofiwg.github.io/libfabric/> Accessed 08/2023.
- [17] Jens Gustedt. 2021. Pointers and integer types. <https://www.open-std.org/jtc1/sc22/wg14/www/docs/n2873.htm>
- [18] Jeff Hammond. 2023. <https://github.com/mpiwg-abi/specification-text-draft/blob/main/print-handle-constants.py>
- [19] Jeff Hammond. 2023. <https://github.com/mpiwg-abi/specification-text-draft/blob/main/IntegerConstants.md>
- [20] Jeff Hammond. 2023. A standalone implementation of the MPI Fortran 2018 module. <https://github.com/jeffhammond/vapaa>
- [21] Jeff Hammond. 2023. MPI\_Aint needs to be wide enough to hold a pointer, not just an address. <https://github.com/mpi-forum/mpi-issues/issues/709>
- [22] Jeff Hammond. 2023. Mukautuva: An MPI ABI compatibility layer. <https://github.com/jeffhammond/mukautuva>
- [23] Jeff R. Hammond, Andreas Schäfer, and Rob Latham. 2014. To INT\_MAX... and Beyond! Exploring Large-Count Support in MPI. In *2014 Workshop on Exascale MPI at Supercomputing Conference (ExaMPI)*. IEEE Computer Society, Los Alamitos, CA, USA, 1–8. <https://doi.org/10.1109/ExaMPI.2014.5>
- [24] M Heroux, J Willenbring, S Shende, C Coti, W Spear, J Peyralans, J Skutnik, and E Keever. 2020. E4S: Extreme-scale Scientific Software Stack. In *2020 Collegeville Workshop on Scientific Software Whitepapers*.
- [25] Tobias Hilbrich, Martin Schulz, Bronis R de Supinski, and Matthias S Müller. 2010. MUST: A scalable approach to runtime error detection in MPI programs. In *Tools for High Performance Computing 2009: Proceedings of the 3rd International Workshop on Parallel Tools for High Performance Computing, September 2009, ZIH, Dresden*. Springer, 53–66.
- [26] Joshua Hursey. 2020. Design Considerations for Building and Running Containerized MPI Applications. In *2020 2nd International Workshop on Containers and New Orchestration Paradigms for Isolated Environments in HPC (CANOPIE-HPC)*, 35–44. <https://doi.org/10.1109/CANOPIEHPC51917.2020.00010>
- [27] Anaconda Inc. 2023. Anaconda Software Distribution. <https://anaconda.com>
- [28] Andreas Knüpfer, Christian Rössel, Dieter an Mey, Scott Biersdorff, Kai Diethelm, Dominic Eschweiler, Markus Geimer, Michael Gerndt, Daniel Lorenz, Allen Malony, et al. 2012. Score-P: A joint performance measurement run-time infrastructure for Periscope, Scalasca, TAU, and Vampir. In *Tools for High Performance Computing 2011: Proceedings of the 5th International Workshop on Parallel Tools for High Performance Computing, September 2011, ZIH, Dresden*. Springer, 79–91.
- [29] Argonne National Laboratory. 2013. MPICH ABI Compatibility Initiative. <https://www.mpi4py.org/abi/>
- [30] Edgar A. León, Marc Joos, Nathan Hanford, Adrien Cotte, Tony Delforge, François Diakhaté, Vincent Ducrot, Ian Karlin, and Marc Pérache. 2021. On-the-Fly, Robust Translation of MPI Libraries. In *2021 IEEE International Conference on Cluster Computing (CLUSTER)*. IEEE, Portland, OR, USA, 504–515. <https://doi.org/10.1109/Cluster48925.2021.00026>
- [31] Greg Lindahl. 2005. The Case for an MPI ABI. In *The 6th International Conference on Linux Clusters/Pittsburgh Supercomputing Center*. [https://www.pbm.com/~lindahl/Case\\_for\\_an\\_MPI\\_ABI.pdf](https://www.pbm.com/~lindahl/Case_for_an_MPI_ABI.pdf)
- [32] Message Passing Interface Forum. 2021. MPI: A Message-Passing Interface Standard. Version 4.0. <https://www.mpi-forum.org/docs/mpi-4.0/mpi40-report.pdf>
- [33] NERSC. [n. d.]. Using Shifter at NERSC. <https://docs.nersc.gov/development/shifter/how-to-use/> Accessed: April 25, 2023.
- [34] Victor Sande Veiga, Manuel Simon, Abdulrahman Azab, Carlos Fernandez, Giuseppa Muscianisi, Giuseppe Fiameni, and Simone Marocchi. 2019. Evaluation and Benchmarking of Singularity MPI containers on EU Research e-Infrastructure. In *2019 IEEE/ACM International Workshop on Containers and New Orchestration Paradigms for Isolated Environments in HPC (CANOPIE-HPC)*, 1–10. <https://doi.org/10.1109/CANOPIE-HPC49598.2019.00006>
- [35] Erik Schnetter. 2022. MPItrampoline: A forwarding MPI implementation. <https://doi.org/10.5281/zenodo.6174409>
- [36] Martin Schulz and Bronis R de Supinski. 2008. *PnMPI*. Technical Report. Lawrence Livermore National Lab.(LLNL), Livermore, CA (United States).
- [37] Jean-Baptiste Skutnik and Frederick Deny. 2021. E4S Container Launcher's 1.0.3 documentation. <https://e4s-cl.readthedocs.io/en/latest/how.html>. Accessed: April 25, 2023.
- [38] Benedikt Steinbusch, Andrew Gaspar, and Jed Brown. 2022. MPI bindings for Rust. <https://github.com/rsmipi/rsmipi>
- [39] Shinji Sumimoto, Kota Nakajima, Akira Naruse, Koichi Hisakado, Takashi Yasui, Yoshikazu Kamoshida, Hiroya Matsuba, Atsushi Hori, Yutaka Ishikawa, et al. 2009. Design and performance evaluation of MPI-Adapter for seamless MPI environment. *Research Report High Performance Computing (HPC)* 2009, 12 (2009), 1–8. <http://id.nii.ac.jp/1001/00062773>
- [40] Alexander Supalov and Artem Yalozo. 2014. 20 Years of the MPI Standard: Now with a Common Application Binary Interface. *The Parallel Universe Magazine* (2014), 28. Issue 18. [https://www.qtsoftware.de/intel/TheParallelUniverseIssue18\\_final1.pdf#page=28](https://www.qtsoftware.de/intel/TheParallelUniverseIssue18_final1.pdf#page=28)
- [41] Alexander V Supalov. 2011. Using message passing interface (MPI) profiling interface for emulating different MPI implementations. US Patent 7,966,624.
- [42] Robert N. M. Watson, Simon W. Moore, Peter Sewell, and Peter G. Neumann. 2019. *An Introduction to ChERI*. Technical Report UCAM-CL-TR-941, University of Cambridge, Computer Laboratory. <https://doi.org/10.48456/tr-941>
- [43] Robert N. M. Watson, Alexander Richardson, Brooks Davis, John Baldwin, David Chisnall, Jessica Clarke, Nathaniel Filardo, Simon W. Moore, Edward Napierala, Peter Sewell, and Peter G. Neumann. 2020. *ChERI C/C++ Programming Guide*. Technical Report UCAM-CL-TR-947, University of Cambridge, Computer Laboratory. <https://doi.org/10.48456/tr-947>
- [44] Farid Zakaria, Thomas RW Scogland, Todd Gamblin, and Carlos Maltzahn. 2022. Mapping out the HPC dependency chaos. In *SC22: International Conference for High Performance Computing, Networking, Storage and Analysis*. IEEE, 1–12.
- [45] Jie Zhang, Xiaoyi Lu, and Dhabaleswar K. Panda. 2017. Is Singularity-Based Container Technology Ready for Running MPI Applications on HPC Clouds?. In *Proceedings of The10th International Conference on Utility and Cloud Computing (Austin, Texas, USA) (UCC '17)*. Association for Computing Machinery, New York, NY, USA, 151–160. <https://doi.org/10.1145/3147213.3147231>
- [46] Hui Zhou. 2023. ABI: add option to build libmpi\_abi.so. <https://github.com/pmodels/mpich/pull/6390> Accessed: May 11, 2023.

## A HANDLE ENCODING

This appendix presents in more detail some implementation specifics of the ABI with respect to operations, handles and datatypes. The integer constants below are all in binary.

### A.1 Operations

The encoding for handle types is provided below:

```
0b0000000000 invalid (uninitialized)
0b0000000000 reserved handle
0b0000100000 MPI_OP_NULL
// arithmetic ops
0b0000100001 MPI_OP_SUM
0b0000100010 MPI_OP_MIN
0b0000100011 MPI_OP_MAX
0b0000100100 MPI_OP_PROD
0b0000100101 reserved arithmetic op
// binary ops
0b0000101000 MPI_OP_BAND
0b0000101001 MPI_OP_BOR
0b0000101010 MPI_OP_BXOR
0b0000101011 reserved bit op
// logical ops
0b0000110000 MPI_OP_LAND
0b0000110001 MPI_OP_LOR
0b0000110010 MPI_OP_LXOR
0b0000110011 reserved logical op
// other ops
0b0000111000 MPI_OP_MINLOC
0b0000111001 MPI_OP_MAXLOC
0b0000111010 reserved other op
0b0000111011 MPI_OP_REPLACE
0b0000111101 MPI_NO_OP
0b0000111110 reserved other op
0b0000111111 reserved handles
```

### A.2 Other Handles

The encoding of opaque handles is shown below:

```
// communicator
0b0100000000 MPI_COMM_NULL
0b0100000001 MPI_COMM_WORLD
0b0100000010 MPI_COMM_SELF
0b0100000011 reserved comm
// group
0b0100000100 MPI_GROUP_NULL
0b0100000101 MPI_GROUP_EMPTY
0b0100000110 reserved group
// windows
0b0100001000 MPI_WIN_NULL
0b0100001001 reserved win
// files
0b0100001100 MPI_FILE_NULL
0b0100001101 reserved file
// sessions
0b0100010000 MPI_SESSION_NULL
0b0100010001 reserved session
// messages
0b0100010100 MPI_MESSAGE_NULL
0b0100010101 MPI_MESSAGE_NO_PROC
0b0100010110 reserved message
// error handler
0b0100011000 MPI_ERRHANDLER_NULL
0b0100011001 MPI_ERRORS_ARE_FATAL
0b0100011010 MPI_ERRORS_RETURN
0b0100011011 MPI_ERRORS_ABORT
0b0100011100 reserved handle
// requests
0b0100100000 MPI_REQUEST_NULL
0b0100100001 reserved request
0b0100100010 reserved handle
```

### A.3 Datatypes

Examples for datatype handles are represented below:

```
0b1000000000 MPI_DATATYPE_NULL
// variable-size types
0b1000000001 MPI_AINT
0b1000000010 MPI_COUNT
0b1000000011 MPI_OFFSET
0b1000000100 reserved datatype
```

```
0b1000000111 MPI_PACKED
0b1000001000 MPI_SHORT
0b1000001001 MPI_INT
0b1000001010 MPI_LONG
0b1000001011 MPI_LONG_LONG
0b1000001100 MPI_UNSIGNED_SHORT
0b1000001101 MPI_UNSIGNED_INT
0b1000001110 MPI_UNSIGNED_LONG
0b1000001111 MPI_UNSIGNED_LONG_LONG
0b1000010000 MPI_FLOAT
...
// fixed-size types
0b1001000000 MPI_INT8_T
0b1001000001 MPI_UINT8_T
0b1001000010 <float 8b>
0b1001000011 MPI_CHAR
0b1001000100 MPI_SIGNED_CHAR
0b1001000101 MPI_UNSIGNED_CHAR
0b1001000110 reserved datatype
0b1001000111 MPI_BYTE
0b1001001000 MPI_INT16_T
0b1001001001 MPI_UINT16_T
0b1001001010 <float 16b>
0b1001001011 <C complex 2x8b>
0b1001001100 reserved datatype
0b1001001101 <C++ complex 2x8b>
0b1001001110 MPI_INT32_T
0b1001001111 MPI_UINT32_T
0b1001010010 <C float 32b>
0b1001010011 <C complex 2x16b>
...
0b1001011000 MPI_INT64_T
0b1001011001 MPI_UINT64_T
0b1001011010 <C float64>
0b1001011011 <C complex 2x32b>
...
```