

Table of Contents

Heterogeneous Jobs using inter-module MPI communication	2
Application-dependent Tuning	2
Modularity-aware Collectives	2
Feature Description	2
Feature Usage	2
CUDA Support by ParaStation MPI	3
What is CUDA awareness for MPI?	3
Some external Resources	3
Current status on the DEEP system	3
Usage on the DEEP system	3
NAM Integration for ParaStation MPI	4
Documentation	4
API Prototype Implementation	4
Extensions to MPI	4
Code Example for a ?Hello World? workflow	4
Usage Example on the DEEP-ER SDV	5
Cleaning up of persistent memory regions	7

...a modularity-enabled MPI library.

Heterogeneous Jobs using inter-module MPI communication

ParaStation MPI provides support for inter-module communication in federated high-speed networks. Therefore, so-called gateway (GW) daemons bridge the MPI traffic between the modules. This mechanism is transparent to the MPI application, i.e., the MPI ranks see a common MPI_COMM_WORLD across all modules within the job. However, the user has to account for these additional GW resources during the job submission. An example SLURM Batch script illustrating the submission of heterogeneous pack jobs including the allocation of GW resources can be found [here](#).

Application-dependent Tuning

The GW protocol supports the fragmentation of larger messages into smaller chunks of a given length, i.e., the Maximum Transfer Unit (MTU). This way, the GW daemon may benefit from pipelining effect resulting in an overlapping of the message transfer from the source to the GW daemon and from the GW daemon to the destination. The chunk size may be influenced by setting the following environment variable:

```
PSP_GW_MTU=<chunk size in byte>
```

The optimal chunk size is highly dependent on the communication pattern and therefore has to be chosen for each application individually.

Modularity-aware Collectives

Feature Description

In the context of DEEP-EST and MSA, ParaStation MPI has been extended by modularity awareness also for collective MPI operations. In doing so, an MSA-aware collective operation is conducted in a hierarchical manner where the intra- and inter- module phases are strictly separated:

- i. First do all module-internal gathering and/or reduction operations if required.
- ii. Then perform the inter-module operation with only one process per module being involved.
- iii. Finally, distribute the data within each module in a strictly module-local manner.

This approach is here exemplarily shown in the following figure for a Broadcast operation with nine processes and three modules:

Besides Broadcast, the following collective operations are currently provided with this awareness:

- MPI_Bcast / MPI_Ibcast
- MPI_Reduce / MPI_Ireduce
- MPI_Allreduce / MPI_Iallreduce
- MPI_Scan / MPI_Iscan
- MPI_Barrier

Feature Usage

For using this feature, the following environment variables must be set:

```
- PSP_MSA_AWARENESS=1      # Generally activate the consideration of modular topologies (NOT enabled by default)
- PSP_MSA_MODULE_ID=xyz    # Pass the respective Module ID (Integer) explicitly to the processes (NOT set automatically)
```

Attention: Please note that the environment variable for the respective Module ID (PSP_MSA_MODULE_ID) is currently *not* set automatically! This means that the user has to set and pass this variable explicitly, for example, via a bash script:

```
#!/bin/bash
# Script (script0.sh) for Module 0: (e.g. Cluster)
APP="./IMB-MPI1 Bcast"
export PSP_MSA_AWARENESS=1
export PSP_MSA_MODULE_ID=0 # <- set an arbitrary ID for this module!
./${APP}
```

```
#!/bin/bash
# Script (script1.sh) for Module 1: (e.g. ESB)
```

```
APP="./IMB-MPI1 Bcast"
export PSP_MSA_AWARENESS=1
export PSP_MSA_MODULE_ID=1 # <- set a different ID for this module!
./${APP}
```

```
> srun ./script0 : ./script1
```

In addition, the following variable can be used to influence the behavior of the MSA awareness for the collective operations:

```
- PSP_MSA_AWARE_COLLLOPS=0|1|2
  0: Disable MSA awareness for collective MPI operations
  1: Enable MSA awareness for collective MPI operations (default)
  2: Apply MSA awareness recursively in multi-level topologies
```

In the recursive application of MSA awareness (`PSP_MSA_AWARE_COLLLOPS=2`), a distinction is first made between inter- and intra-module communication and then, in a second step, likewise between inter- and intra-node communication within the modules.

CUDA Support by ParaStation MPI

What is CUDA awareness for MPI?

In brief, *CUDA awareness* in an MPI library means that a mixed CUDA + MPI application is allowed to pass pointers to CUDA buffers (these are memory regions located on the GPU, the so-called *Device* memory) directly to MPI functions such as `MPI_Send()` or `MPI_Recv()`. A non CUDA-aware MPI library would fail in such a case because the CUDA-memory cannot be accessed directly, e.g., via `load/store` or `memcpy()` but has to be transferred in advance to the host memory via special routines such as `cudaMemcpy()`. As opposed to this, a CUDA-aware MPI library recognizes that a pointer is associated with a buffer within the device memory and can then copy this buffer prior to the communication into a temporarily host buffer — what is called *staging* of this buffer. Additionally, a CUDA-aware MPI library may also apply some kind of optimizations, e.g., by means of exploiting so-called *GPUDirect* capabilities that allow for direct RDMA transfers from and to the device memory.

Some external Resources

- [Getting started with CUDA](#) (by NVIDIA)
- [NVIDIA GPUDirect Overview](#) (by NVIDIA)
- [Introduction to CUDA-Aware MPI](#) (by NVIDIA)

Current status on the DEEP system

Currently (effective October 2019), ParaStation MPI supports CUDA-awareness for Extoll just from the semantic-related point of view: The usage of Device pointers as arguments for send and receive buffers when calling MPI functions is supported but by an explicit *Staging* when Extoll is used. This is because the Extoll runtime up to now does not support GPUDirect, but EXTOLL is currently working on this in the context of DEEP-EST. As soon as GPUDirect will be supported by Extoll, this will also be integrated and enabled in ParaStation MPI. (BTW: For InfiniBand communication, ParaStation MPI is already GPUDirect enabled.)

Usage on the DEEP system

Warning: *This manual section is currently under development. Therefore, the following usage guidelines may be not flawless and are likely to change in some respects in the near future!*

On the DEEP system, the CUDA awareness can be enabled by loading a module that links to a dedicated ParaStation MPI library providing CUDA support:

```
module load GCC
module load ParaStationMPI/5.4.0-1-CUDA
```

Please note that CUDA-awareness might impact the MPI performance on systems parts where CUDA is not used. Therefore, it might be useful (and the other way around necessary) to disable/enable the CUDA-awareness by setting this environment variable:

```
PSP_CUDA=0|1
```

NAM Integration for ParaStation MPI

Documentation

- [Proposal for accessing the NAM via MPI](#)
- [API Prototype Implementation](#)
- [Usage Example on the DEEP-EST SDV](#)

API Prototype Implementation

For evaluating the proposed semantics and API extensions, we have already developed a shared-memory-based prototype implementation where the persistent NAM is (more or less) ?emulated? by persistent shared-memory (with *deep_mem_kind=deep_mem_persistent*).

Advice to users

Please note that this prototype is not intended to actually emulate the NAM but shall rather offer a possibility for the later users and programmers to evaluate the proposed semantics from the MPI application?s point of view. Therefore, the focus here is not put on the question of how remote memory is managed at its location (currently by MPI processes running local to the memory later by the NAM manager or the NAM itself), but on the question of how process-foreign memory regions can be exposed locally. That means that (at least currently) for accessing a persistent RMA window, it has to be made sure that there is at least one MPI process running locally to each of the window?s memory regions.

Extensions to MPI

The API proposal strives to stick to the current MPI standard as close as possible and to avoid the addition of new API functions and other symbols. However, in order to make the usage of the prototype a little bit more convenient for the user, we have added at least a small set of new symbols (denoted with MPIX) that may be used by the applications.

```
extern int MPIX_WIN_DISP_UNITS;
#define MPIX_WIN_FLAVOR_INTERCOMM      (MPI_WIN_FLAVOR_CREATE +      \
                                         MPI_WIN_FLAVOR_ALLOCATE +    \
                                         MPI_WIN_FLAVOR_DYNAMIC +      \
                                         MPI_WIN_FLAVOR_SHARED + 0)
#define MPIX_WIN_FLAVOR_INTERCOMM_SHARED (MPI_WIN_FLAVOR_CREATE +      \
                                         MPI_WIN_FLAVOR_ALLOCATE +    \
                                         MPI_WIN_FLAVOR_DYNAMIC +      \
                                         MPI_WIN_FLAVOR_SHARED + 1)
```

Code Example for a ?Hello World? workflow

The following two C codes should demonstrate how it shall become possible to pass intermediate data between two subsequent steps of a workflow (Step 1: hello / Step 2: world) via the persistent memory of the NAM (currently emulated by persistent shared-memory):

```
/** hello.c **/

/* Create persistent MPI RMA window: */
MPI_Info_create(&win_info);
MPI_Info_set(win_info, "deep_mem_kind", "deep_mem_persistent");
MPI_Win_allocate(sizeof(char) * HELLO_STR_LEN, sizeof(char), win_info, MPI_COMM_WORLD,
                 &win_base, &win);

/* Put some content into the local region of the window: */
if(argc > 1) {
    snprintf(win_base, HELLO_STR_LEN, "Hello World from rank %d! %s", world_rank, argv[1]);
} else {
    snprintf(win_base, HELLO_STR_LEN, "Hello World from rank %d!", world_rank);
}
MPI_Win_fence(0, win);

/* Retrieve port name of window: */
MPI_Info_free(&win_info);
MPI_Win_get_info(win, &win_info);
```

```

MPI_Info_get(win_info, "deep_win_port_name", INFO_VALUE_LEN, info_value, &flag);

if(flag) {
    strcpy(port_name, info_value);
    if(world_rank == root) printf("(%d) The Window's port name is: %s\n", world_rank, port_name);
} else {
    if(world_rank == root) printf("(%d) No port name found!\n", world_rank);
}

```

```

/** world.c */

/* Check for port name: (to be passed as a command line argument) */
if(argc == 1) {
    if(world_rank == root) printf("[%d] No port name found!\n", world_rank);
    goto finalize;
} else {
    strcpy(port_name, argv[1]);
    if(world_rank == root) printf("[%d] The Window's port name is: %s\n", world_rank, port_name);
}

/* Try to connect to the persistent window: */
MPI_Info_create(&win_info);
MPI_Info_set(win_info, "deep_win_connect", "true");
MPI_Comm_set_errhandler(MPI_COMM_WORLD, MPI_ERRORS_RETURN);
errcode = MPI_Comm_connect(port_name, win_info, root, MPI_COMM_WORLD, &inter_comm);
printf("[%d] Connection to persistent memory region established!\n", world_rank);

/* Retrieve the number of remote regions: (= former number of ranks) */
MPI_Comm_remote_size(inter_comm, &remote_size);
if(world_rank == root) printf("[%d] Number of remote regions: %d\n", world_rank, remote_size);

/* Create window object for accessing the remote regions: */
MPI_Win_create_dynamic(MPI_INFO_NULL, inter_comm, &win);
MPI_Win_get_attr(win, MPI_WIN_CREATE_FLAVOR, &create_flavor, &flag);
assert(*create_flavor == MPIX_WIN_FLAVOR_INTERCOMM);
MPI_Win_fence(0, win);

/* Check the accessibility and the content of the remote regions: */
for(i=0; i<remote_size; i++) {
    char hello_string[HELLO_STR_LEN];
    MPI_Get(hello_string, HELLO_STR_LEN, MPI_CHAR, i, 0, HELLO_STR_LEN, MPI_CHAR, win);
    MPI_Win_fence(0, win);
    printf("[%d] Get from %d: %s\n", world_rank, i, hello_string);
}

```

Usage Example on the DEEP-ER SDV

On the DEEP-ER SDV, there is already a special version of ParaStation MPI installed that features all the introduced API extensions. It is accessible via the module system:

```
> module load parastation/5.2.1-1-mt-wp6
```

When allocating a session with N nodes, one can run an MPI session (let's say with n processes distributed across the N nodes) where each of the processes is contributing its local and persistent memory region to an MPI window:

```

> salloc --partition=sdv --nodes=4 --time=01:00:00
salloc: Granted job allocation 2514
> srun -n4 -N4 ./hello 'Have fun!'
(0) Running on deeper-sdv13
(1) Running on deeper-sdv14

```

```

(2) Running on deeper-sdv15
(3) Running on deeper-sdv16
(0) The Window's port name is: shmid:347897856:92010569
(0) Calling finalize...
(1) Calling finalize...
(2) Calling finalize...
(3) Calling finalize...
(0) Calling finalize...
(0) Finalize done!
(1) Finalize done!
(2) Finalize done!
(3) Finalize done!

```

Afterwards, on all the nodes involved (and later on the NAM) one persistent memory region has been created by each of the MPI processes. The ?port name? for accessing the persistent window again is in this example:

```
shmid:347897856:92010569
```

By means of this port name (here to be passes as a command line argument), all the processes of a subsequent MPI session can access the persistent window provided that there is again at least one MPI processes running locally to each of the persistent but distributed regions:

```

> srun -n4 -N4 ./world shmid:347897856:92010569
[0] Running on deeper-sdv13
[1] Running on deeper-sdv14
[2] Running on deeper-sdv15
[3] Running on deeper-sdv16
[0] The Window's port name is: shmid:347897856:92010569
[1] Connection to persistent memory region established!
[3] Connection to persistent memory region established!
[0] Connection to persistent memory region established!
[2] Connection to persistent memory region established!
[0] Number of remote regions: 4
[0] Get from 0: Hello World from rank 0! Have fun!
[1] Get from 0: Hello World from rank 0! Have fun!
[2] Get from 0: Hello World from rank 0! Have fun!
[3] Get from 0: Hello World from rank 0! Have fun!
[0] Get from 1: Hello World from rank 1! Have fun!
[1] Get from 1: Hello World from rank 1! Have fun!
[2] Get from 1: Hello World from rank 1! Have fun!
[3] Get from 1: Hello World from rank 1! Have fun!
[0] Get from 2: Hello World from rank 2! Have fun!
[1] Get from 2: Hello World from rank 2! Have fun!
[2] Get from 2: Hello World from rank 2! Have fun!
[3] Get from 2: Hello World from rank 2! Have fun!
[0] Get from 3: Hello World from rank 3! Have fun!
[1] Get from 3: Hello World from rank 3! Have fun!
[2] Get from 3: Hello World from rank 3! Have fun!
[3] Get from 3: Hello World from rank 3! Have fun!
[0] Calling finalize...
[1] Calling finalize...
[2] Calling finalize...
[3] Calling finalize...
[0] Finalize done!
[1] Finalize done!
[2] Finalize done!
[3] Finalize done!

```

Advice to users

Please note that if not all persistent memory regions are covered by the subsequent session, the ?connection establishment? to the remote RMA window fails:

```
> srun -n4 -N2 ./world shmid: 347897856:92010569
[3] Running on deeper-sdv14
[1] Running on deeper-sdv13
[2] Running on deeper-sdv14
[0] Running on deeper-sdv13
[0] The Window's port name is: shmid:347930624:92010569
[0] ERROR: Could not connect to persistent memory region!
application called MPI_Abort(MPI_COMM_WORLD, -1)
?
```

Cleaning up of persistent memory regions

If the connection to a persistent memory region succeeds, the window and all of its memory regions will eventually be removed by the MPI_Win_free call of the subsequent MPI session (here by world.c) at least if not deep_mem_persistent is passed again as an Info argument. However, if a connection attempt fails, the persistent memory regions still persist.

For explicitly cleaning up those artefacts, one can use a simple batch script:

```
#!/bin/bash
keys=`ipcs | grep 777 | cut -d' ' -f2`
for key in $keys ; do
    ipcrm -m $key
done
```

Advice to administrators

Obviously, a good idea would be the integration of such an automated cleaning-up procedure as a default into the epilogue scripts for the jobs.