

Table of Contents

Information about the batch system (SLURM)	2
Overview	2
Available Partitions	2
Remark about environment	2
An introductory example	2
From a shell on a node	2
Running directly from the front ends	3
Batch script	3
Heterogeneous jobs	4
Heterogeneous jobs with MPI communication across modules	5
Workflows	6
slurm_workflow Library	10
Submitting jobs to alternative modules	11
Information on past jobs and accounting	11
FAQ	11
Is there a cheat sheet for all main Slurm commands?	11
Why's my job not running?	11
How can I check which jobs are running in the machine?	12
How do I do chain jobs with dependencies?	12
How can check the status of partitions and nodes?	12
Can I join stderr and stdout like it was done with -joe in Torque?	13

Information about the batch system (SLURM)

The DEEP prototype systems are running SLURM for resource management. Documentation of Slurm can be found [?here](#).

Overview

Slurm offers interactive and batch jobs (scripts submitted into the system). The relevant commands are `srun` and `sbatch`. The `srun` command can be used to spawn processes (**please do not use `mpiexec`**), both from the frontend and from within a batch script. You can also get a shell on a node to work locally there (e.g. to compile your application natively for a special platform or module).

Available Partitions

Please note that there is no default partition configured. In order to run a job, you have to specify one of the following partitions, using the `--partition=...` switch:

Name	Description
dp-cn	dp-cn[01-50], DEEP-EST Cluster nodes (Xeon Skylake)
dp-dam	dp-dam[01-16], DEEP-EST Dam nodes (Xeon Cascadelake + 1 V100 + 1 Stratix 10)
dp-dam-ext	dp-dam[09-16], DEEP-EST Dam nodes connected with Extoll Tourmalet
dp-esb-ib	dp-esb[26-75], DEEP-EST ESB nodes connected with IB EDR (Xeon Cascadelake + 1 V100)
dp-esb-ext	dp-esb[01-25], DEEP-EST ESB nodes connected with Extoll Fabri3)
dp-sdv-esb	dp-sdv-esb[01-02], DEEP-EST ESB Test nodes (Xeon Cascadelake + 1 V100)
ml-gpu	ml-gpu[01-03], GPU test nodes for ML applications (4 V100 cards)
knl	knl[01,04-06], KNL nodes
knl256	knl[01,05], KNL nodes with 64 cores
knl272	knl[04,06], KNL nodes with 68 cores
snc4	knl[05], KNL node in snc4 memory mode
debug	all compute nodes (no gateways)

Anytime, you can list the state of the partitions with the `sinfo` command. The properties of a partition can be seen using

```
scontrol show partition <partition>
```

Remark about environment

By default, Slurm passes the environment from your job submission session directly to the execution environment. Please be aware of this when running jobs with `srun` or when submitting scripts with `sbatch`. This behavior can be controlled via the `--export` option. Please refer to the [?Slurm documentation](#) to get more information about this.

In particular, when submitting job scripts, **it is recommended to load the necessary modules within the script and submit the script from a clean environment.**

An introductory example

Suppose you have an mpi executable named `hello_mpi`. There are three ways to start the binary.

From a shell on a node

First, start a shell on a node. Assume you would like to run your mpi task on 4 cluster nodes with 2 tasks per node:

```
[kreutz1@deepv /p/project/cdeep/kreutz1/Temp]$ srun -A deep -p dp-cn -N 4 -n 8 -t 00:30:00 --pty /bin/bash -i
[kreutz1@dp-cn01 /p/project/cdeep/kreutz1/Temp]$
```

The environment is transported to the remote shell, no `.profile`, `.bashrc`, ... are sourced (especially not the modules default from `/etc/profile.d/modules.sh`). As of March 2020, an account has to be specified using the `--account` (short `-A`) option, which is "deep" for the project members. For people not included in the DEEP-EST project, please use the "Budget" name you received along with your account creation.

Once you get to the compute node, start your application using `srun`. Note that the number of tasks used is the same as specified in the initial `srun` command above (4 nodes with two tasks each):

```
[kreutz1@deepv Temp]$ srun -A deep -p dp-cn -N 4 -n 8 -t 00:30:00 --pty /bin/bash -i
[kreutz1@dp-cn01 Temp]$ srun ./MPI_HelloWorld
Hello World from rank 3 of 8 on dp-cn02
Hello World from rank 7 of 8 on dp-cn04
Hello World from rank 2 of 8 on dp-cn02
Hello World from rank 6 of 8 on dp-cn04
Hello World from rank 0 of 8 on dp-cn01
Hello World from rank 4 of 8 on dp-cn03
Hello World from rank 1 of 8 on dp-cn01
Hello World from rank 5 of 8 on dp-cn03
```

You can ignore potential warnings about the cpu binding. ParaStation will pin your processes.

Running directly from the front ends

You can run the application directly from the frontend, bypassing the shell. Do not forget to set the correct environment for running your executable on the login node as this will be used for execution with `srun`.

```
[kreutz1@deepv Temp]$ ml GCC/10.3.0 ParaStationMPI/5.4.9-1
[kreutz1@deepv Temp]$ srun -A deep -p dp-cn -N 4 -n 8 -t 00:30:00 ./MPI_HelloWorld
Hello World from rank 7 of 8 on dp-cn04
Hello World from rank 3 of 8 on dp-cn02
Hello World from rank 6 of 8 on dp-cn04
Hello World from rank 2 of 8 on dp-cn02
Hello World from rank 4 of 8 on dp-cn03
Hello World from rank 0 of 8 on dp-cn01
Hello World from rank 1 of 8 on dp-cn01
Hello World from rank 5 of 8 on dp-cn03
```

It can be useful to create an allocation which you can be used for several runs of your job:

```
[kreutz1@deepv Temp]$ salloc -A deep -p dp-cn -N 4 -n 8 -t 00:30:00
salloc: Granted job allocation 69263
[kreutz1@deepv Temp]$ srun ./MPI_HelloWorld
Hello World from rank 7 of 8 on dp-cn04
Hello World from rank 3 of 8 on dp-cn02
Hello World from rank 6 of 8 on dp-cn04
Hello World from rank 2 of 8 on dp-cn02
Hello World from rank 5 of 8 on dp-cn03
Hello World from rank 1 of 8 on dp-cn01
Hello World from rank 4 of 8 on dp-cn03
Hello World from rank 0 of 8 on dp-cn01
...
# several more runs
...
[kreutz1@deepv Temp]$ exit
exit
salloc: Relinquishing job allocation 69263
```

Batch script

As stated above, it is recommended to load the necessary modules within the script and submit the script from a clean environment.

The following script `hello_cluster.sh` will unload all modules and load the modules required for executing the given binary:

```
#!/bin/bash

#SBATCH --partition=dp-esb
#SBATCH -A deep
#SBATCH -N 4
#SBATCH -n 8
#SBATCH -o /p/project/cdeep/kreutz1/hello_cluster-%j.out
#SBATCH -e /p/project/cdeep/kreutz1/hello_cluster-%j.err
#SBATCH --time=00:10:00

ml purge
ml GCC ParaStationMPI
srun ./MPI_HelloWorld
```

This script requests 4 nodes of the ESB module with 8 tasks, specifies the stdout and stderr files, and asks for 10 minutes of walltime. You can submit the job script as follows:

```
[kreutz1@deepv Temp]$ sbatch hello_cluster.sh
Submitted batch job 69264
```

... and check what it's doing:

```
[kreutz1@deepv Temp]$ squeue -u $USER
```

JOBID	PARTITION	NAME	USER	ST	TIME	NODES	NODELIST(REASON)
69264	dp-cn	hello_cl	kreutz1	CG	0:04	4	dp-cn[01-04]

Once finished, you can check the result (and the error file if needed)

```
[kreutz1@deepv Temp]$ cat /p/project/cdeep/kreutz1/hello_cluster-69264.out
Hello World from rank 7 of 8 on dp-esb37
Hello World from rank 3 of 8 on dp-esb35
Hello World from rank 5 of 8 on dp-esb36
Hello World from rank 1 of 8 on dp-esb34
Hello World from rank 6 of 8 on dp-esb37
Hello World from rank 2 of 8 on dp-esb35
Hello World from rank 4 of 8 on dp-esb36
Hello World from rank 0 of 8 on dp-esb34
```

Heterogeneous jobs

As of version 17.11 of Slurm, heterogeneous jobs are supported. For example, the user can run:

```
srun --account=deep --partition=dp-cn -N 1 -n 1 hostname : --partition=dp-dam -N 1 -n 1 hostname
dp-cn01
dp-dam01
```

Please notice the `:` separating the definitions for each sub-job of the heterogeneous job. Also, please be aware that it is possible to have more than two sub-jobs in a heterogeneous job.

The user can also request several sets of nodes in a heterogeneous allocation using `salloc`. For example:

```
salloc --partition=dp-cn -N 2 : --partition=dp-dam -N 4
```

In order to submit a heterogeneous job via `sbatch`, the user needs to set the batch script similar to the following one:

```
#!/bin/bash

#SBATCH --job-name=imb_execute_1
#SBATCH --account=deep
#SBATCH --mail-user=
#SBATCH --mail-type=ALL
#SBATCH --output=job.out
#SBATCH --error=job.err
#SBATCH --time=00:02:00

#SBATCH --partition=dp-cn
#SBATCH --nodes=1
#SBATCH --ntasks=12
#SBATCH --ntasks-per-node=12
#SBATCH --cpus-per-task=1

#SBATCH packjob

#SBATCH --partition=dp-dam
#SBATCH --constraint=
#SBATCH --nodes=1
#SBATCH --ntasks=12
#SBATCH --ntasks-per-node=12
#SBATCH --cpus-per-task=1

srun ./app_cn : ./app_dam
```

Here the `packjob` keyword allows to define Slurm parameters for each sub-job of the heterogeneous job. Some Slurm options can be defined once at the beginning of the script and are automatically propagated to all sub-jobs of the heterogeneous job, while some others (i.e. `--nodes` or `--ntasks`) must be defined for each sub-job. You can find a list of the propagated options on the [?Slurm documentation](#).

When submitting a heterogeneous job with this colon notation using ParaStationMPI, a unique `MPI_COMM_WORLD` is created, spanning across the two partitions. If this is not desired, one can use the `--pack-group` key to submit independent job steps to the different node-groups of a heterogeneous allocation:

```
srun --pack-group=0 ./app_cn ; srun --pack-group=1 ./app_dam
```

Using this configuration implies that inter-communication must be established manually by the applications during run time, if needed.

For more information about heterogeneous jobs please refer to the [?relevant page](#) of the Slurm documentation.

Heterogeneous jobs with MPI communication across modules

In order to establish MPI communication across modules using different interconnect technologies, some special Gateway nodes must be used. On the DEEP-EST system, MPI communication across gateways is needed only between Infiniband and Extoll interconnects.

Attention: Only ParaStation MPI supports MPI communication across gateway nodes.

This is an example job script for setting up an Intel MPI benchmark between a Cluster and a DAM node using a IB ↔ Extoll gateway for MPI communication:

```
#!/bin/bash

# Script to launch IMB PingPong between DAM-CN using 1 Gateway
# Use the gateway allocation provided by SLURM
# Use the packjob feature to launch separately CM and DAM executable

# General configuration of the job
#SBATCH --job-name=modular-imb
#SBATCH --account=deep
```

```

#SBATCH --time=00:10:00
#SBATCH --output=modular-imb-%j.out
#SBATCH --error=modular-imb-%j.err

# Configure the gateway daemon
#SBATCH --gw_num=1
#SBATCH --gw_psgwd_per_node=1

# Configure node and process count on the CM
#SBATCH --partition=dp-cn
#SBATCH --nodes=1
#SBATCH --ntasks-per-node=1

#SBATCH packjob

# Configure node and process count on the DAM
#SBATCH --partition=dp-dam-ext
#SBATCH --nodes=1
#SBATCH --ntasks-per-node=1

# Echo job configuration
echo "DEBUG: SLURM_JOB_NODELIST=$SLURM_JOB_NODELIST"
echo "DEBUG: SLURM_NNODES=$SLURM_NNODES"
echo "DEBUG: SLURM_TASKS_PER_NODE=$SLURM_TASKS_PER_NODE"

# Set the environment to use PS-MPI
module --force purge
module use $OTHERSTAGES
module load Stages/Devel-2019a
module load Intel
module load ParaStationMPI

# Show the hosts we are running on
srun hostname : hostname

# Execute
APP="./IMB-MPI1 Uniband"
srun ${APP} : ${APP}

```

Attention: During the first part of 2020, only the DAM nodes will have Extoll interconnect (and only the nodes belonging to the deep-dam-ext partition will have Extoll active), while the CM and the ESB nodes will be connected via Infiniband. This will change later during the course of the project (expected end of Summer 2020), when the ESB will be equipped with Extoll connectivity (Infiniband will be removed from the ESB and left only for the CM).

A general description of how the user can request and use gateway nodes is provided at [this section](#) of the JURECA documentation.

Attention: some information provided on the JURECA documentation do not apply for the DEEP system. In particular:

- as of 31/03/2020, the DEEP system has 2 gateway nodes.
- As of 09/01/2020 the gateway nodes are exclusive to the job requesting them. Given the limited number of gateway nodes available on the system, this may change in the future.
- As of 09/04/2020 the xenv utility (necessary on JURECA to load modules for different architectures - Haswell and KNL) is not needed any more on DEEP when using the latest version of ParaStationMPI (currently available in the Devel-2019a stage and soon available on the default production stage).

Workflows

The version of Slurm installed on the system enables workflows (chains of jobs) with the possibility of having some overlap between the dependent jobs. This allows them to exchange data over the network rather than writing and reading it using a common file system.

Workflows can be submitted in two ways:

- using the new `--delay` option provided in `sbatch` command, which allows to start a job with a fixed delay from the start of the previous job;
- submitting separate jobs using an `afterok` dependency and later requesting a change in dependency type from `afterok` to `after` (using our provided shared library), which allows the second job to start if resources are available.

An example project that uses all the features discussed is provided [?here](#).

The following simple example script helps to understand the mechanism of new `delay` switch for workflows.

```
[huda1@deepv scripts]$ cat test.sh
#!/bin/sh

NAME=$(hostname)
echo "$NAME: Going to sleep for $1 seconds"
sleep $1
echo "$NAME: Awake"

[huda1@deepv scripts]$ cat batch_workflow.sh
#!/bin/bash
#SBATCH -p sdv -N2 -t3

#SBATCH packjob

#SBATCH -p sdv -N1 -t3 --delay 2

srun test.sh 175

[huda1@deepv scripts]$
```

In the above `sbatch` script, the usage of `--delay` can be seen. The option takes values in minutes and allows us to delay the subsequent job of by a user-defined number of minutes from the start of the first job in the job pack. After submission of this job pack (which uses the same syntax as a heterogeneous job), Slurm divides it into separate jobs. Also, Slurm ensures that the delay is respected by using reservations, rather than the usual scheduler.

Here is the example execution of this script.

```
[huda1@deepv scripts]$ sbatch batch_workflow.sh
Submitted batch job 81458
[huda1@deepv scripts]$ squeue -u huda1
```

JOBID	PARTITION	NAME	USER	ST	TIME	NODES	NODELIST(Reason)
81458	sdv	batch_wo	huda1	CF	0:01	2	deeper-sdv[02-03]
81459	sdv	batch_wo	huda1	PD	0:00	1	(Reservation)

```
[huda1@deepv scripts]$
```

Here the second job (81459) will start 2 minutes after the start of the first job (81458), and it is listed as PD (Pending) with reason Reservation because it will start as soon as its reservation will begin.

Similarly, the output files will be different for each separated job in the job pack. the final outputs are:

```
[huda1@deepv scripts]$ cat slurm-81458.out
deeper-sdv02: Going to sleep for 175 seconds
deeper-sdv03: Going to sleep for 175 seconds
deeper-sdv02: Awake
deeper-sdv03: Awake

[huda1@deepv scripts]$ cat slurm-81459.out
deeper-sdv01: Going to sleep for 175 seconds
deeper-sdv01: Awake

[huda1@deepv scripts]$
```

Another feature to note is that if there are multiple jobs in a job pack and any number of consecutive jobs have the same delay values, they are combined into a new heterogeneous job. This allows to have heterogeneous jobs within workflows. Here is an example of such a script:

```
[huda1@deepv scripts]$ cat batch_workflow_complex.sh
#!/bin/bash

#SBATCH -p sdv -N 2 -t 3
#SBATCH -J first

#SBATCH packjob

#SBATCH -p sdv -N 1 -t 3 --delay 2
#SBATCH -J second

#SBATCH packjob

#SBATCH -p sdv -N 1 -t 2 --delay 2
#SBATCH -J second

#SBATCH packjob

#SBATCH -p sdv -N 2 -t 3 --delay 4
#SBATCH -J third

if [ "$SLURM_JOB_NAME" == "first" ]
then
    srun ./test.sh 150

elif [ "$SLURM_JOB_NAME" == "second" ]
then
    srun ./test.sh 150 : ./test.sh 115

elif [ "$SLURM_JOB_NAME" == "third" ]
then
    srun ./test.sh 155

fi

[huda1@deepv scripts]$
```

Note the delay values for the second and third job in the script are equal.

Attention The delay value for the 4th job (-J third) is relative to the start of the first job and not from the start of middle 2 jobs. So it will start after 2 minutes of the start time of the middle jobs. Also, note the usage of the environment variable SLURM_JOB_NAME in the script to decide which command to run in which job. The example execution leads to the following:

The example execution leads to the following:

```
[huda1@deepv scripts]$ sbatch batch_workflow_complex.sh
Submitted batch job 81460

[huda1@deepv scripts]$ squeue -u huda1
```

JOBID	PARTITION	NAME	USER	ST	TIME	NODES	NODELIST(REASON)
81461+0	sdv	second	huda1	PD	0:00	1	(Resources)
81461+1	sdv	second	huda1	PD	0:00	1	(Resources)
81463	sdv	third	huda1	PD	0:00	2	(Resources)
81460	sdv	first	huda1	PD	0:00	2	(Resources)

```

[huda1@deepv scripts]$
```


Note that the submitted heterogeneous job has been divided into a single job (81460), a job pack (81461) and again a single job (81643). Similarly, three different output files will be generated, one for each new job.

```
[huda1@deepv scripts]$ cat slurm-81460.out
deeper-sdv03: Going to sleep for 150 seconds
deeper-sdv04: Going to sleep for 150 seconds
deeper-sdv03: Awake
deeper-sdv04: Awake
```

```
[huda1@deepv scripts]$ cat slurm-81461.out
deeper-sdv01: Going to sleep for 150 seconds
deeper-sdv02: Going to sleep for 115 seconds
deeper-sdv02: Awake
deeper-sdv01: Awake
```

```
[huda1@deepv scripts]$ cat slurm-81463.out
deeper-sdv01: Going to sleep for 155 seconds
deeper-sdv02: Going to sleep for 155 seconds
deeper-sdv01: Awake
deeper-sdv02: Awake
```

```
[huda1@deepv scripts]$
```

If a job exits earlier than the allocated time asked by the user, the corresponding reservation for this job is deleted 5 minutes after the end of the job, automatically and the resources become available for the other jobs. However, users should be careful with the requested time when submitting workflows as the larger time values can delay the scheduling of the workflows depending on the situation of the resources.

The workflows created using `delay switch` ensure overlap between the applications. Instead, using the alternative method (which uses Slurm job dependencies) does not ensure a time overlap between two consecutive jobs of a workflow. Though, in this case users do not need to guess the time a job will take and how much should the delay between jobs starting times should be.

Jobs can be chained in Slurm with the aid of the following script:

```
[huda1@deepv scripts]$ cat chain_jobs.sh
#!/usr/bin/env bash

if [ $# -lt 3 ]
then
    echo "$0: ERROR (MISSING ARGUMENTS)"
    exit 1
fi

LOCKFILE=$1
DEPENDENCY_TYPE=$2
shift 2
SUBMITSCRIPT=$*

if [ -f $LOCKFILE ]
then
    if [[ "$DEPENDENCY_TYPE" =~ ^(after|afterany|afterok|afternotok)$ ]]; then
        DEPEND_JOBID=`head -1 $LOCKFILE`
        echo "sbatch --dependency=${DEPENDENCY_TYPE}:${DEPEND_JOBID} $SUBMITSCRIPT"
        JOBID=`sbatch --dependency=${DEPENDENCY_TYPE}:${DEPEND_JOBID} $SUBMITSCRIPT`
    else
        echo "$0: ERROR (WRONG DEPENDENCY TYPE: choose among 'after', 'afterany', 'afterok' or 'afternotok')"
    fi
else
    echo "sbatch $SUBMITSCRIPT"
    JOBID=`sbatch $SUBMITSCRIPT`
fi
```

```

echo "RETURN: $JOBID"
# the JOBID is the last field of the output line
echo "${JOBID##* } > $LOCKFILE

exit 0

```

This is a modified version of the of the chainJobs.sh included in JUBE, which allows to select the desired dependency type between two consecutive jobs. Here is an example of submission of a workflow with Slurm dependencies using the previous script (here called chain_jobs.sh):

```

[huda1@deepv scripts]$ ./chain_jobs.sh lockfile afterok simple_job.sh
sbatch simple_job.sh
RETURN: Submitted batch job 98626
[huda1@deepv scripts]$ ./chain_jobs.sh lockfile afterok simple_job.sh
sbatch --dependency=afterok:98626 simple_job.sh
RETURN: Submitted batch job 98627
[huda1@deepv scripts]$ ./chain_jobs.sh lockfile afterok simple_job.sh
sbatch --dependency=afterok:98627 simple_job.sh
RETURN: Submitted batch job 98628
[huda1@deepv scripts]$ squeue -u huda1
      JOBID PARTITION   NAME   USER  ST       TIME  NODES NODELIST(REASON)
      98627         sdv simple_j  huda1  PD        0:00      2 (Dependency)
      98628         sdv simple_j  huda1  PD        0:00      2 (Dependency)
      98626         sdv simple_j  huda1   R        0:21      2 deeper-sdv[01-02]
[huda1@deepv scripts]$ scontrol show job 98628 | grep Dependency
  JobState=PENDING Reason=Dependency Dependency=afterok:98627
[huda1@deepv scripts]$ cat lockfile
98628

```

Please note that lockfile must not exist previous to the first submission. After the first job submission, that file will contain the id of last submitted job, which is later used by the subsequent call to the chain_job.sh script to set the dependency.

slurm_workflow Library

In order to improve the usability of workflows, a library has been developed and deployed on the system to allow users to interact with the scheduler from within applications involved in a workflow. The library is called slurm_workflow.

The library has two functions.

The first function is relevant to workflows created using the --delay switch and moves all the reservations of the remaining workflow jobs.

```

/*
IN:   number of minutes from now. The start time of the next reservation of the workflow is moved to this time if doable.
OUT:  0 successful, non zero unsuccessful.slurm_wf_error is set.
*/
int slurm_wf_move_all_res(uint32_t t);

```

The minimum value usable for the parameter is currently 2 (minutes).

The second function changes the dependencies type of all jobs dependent on the current job from afterok:job_id to after:job_id.

```

/*
OUT:  0 successful, error no otherwise.
*/

int slurm_change_dep();

```

This enables the jobs in workflow eligible for allocation by Slurm.

Both functions allow an application to notify the scheduler that it is ready for the start of the subsequent jobs of a workflow. This is particularly relevant in case a network connection must be established between the two applications, but only after a certain time from the start of the first job.

When using the library, the header file can be included using `#include <slurm/slurm_workflow.h>` and the library should be linked against using `-lslurm_workflow -lslurm`.

Submitting jobs to alternative modules

Users can submit batch jobs to multiple modules by using `--module-list` extension for Slurm `sbatch` command. This extension accepts two modules, a primary module, submitted with higher priority, and an alternative module that receives a lower priority in the job queue. In the below example the job is submitted to two modules: the primary module is `dp-cn`, while the secondary module is `dp-dam`.

```
sbatch --module-list=dp-cn,dp-dam job.batch
```

The parameters for the alternative module are automatically calculated by using an internal conversion model. `Module-list` is an alternative to `--partition`, which does not apply any conversion model, and submits the job to multiple partitions but with the same configuration. Available conversion models:

- CM to DAM: number of requested nodes / 2, number of tasks per node x2 (CPU cores ratio)
- ESB to CM: same number of nodes, time limit * 10 (due to GPU vs CPU performance), number of tasks per node / 3 (CPU cores ratio)
- DAM to CM: number of nodes * 2 (due to available memory), time limit * 10 (due to GPU vs CPU performance), number of tasks per node / 2 (CPU cores ratio)
- DAM to ESB: number of nodes * 4 (due to available memory), time limit / 4 (due to using more nodes), number of tasks per node / 6 (CPU cores ratio)

At submission time it is recommended to specify the number of nodes, the number of tasks per node, the number of CPUs per task, and the number of GPUs per node.

`Module-list` is currently not compatible with other dependencies specified with `--dependency` clause, and also with other partitions specified with `--partition`.

Information on past jobs and accounting

The `sacct` command can be used to enquire the Slurm database about a past job.

```
[kreutz1@deepv Temp]$ sacct -j 69268
```

JobID	JobName	Partition	Account	AllocCPUS	State	ExitCode
69268+0	bash	dp-cn	deepest-a+	96	COMPLETED	0:0
69268+0.0	MPI_Hello+		deepest-a+	2	COMPLETED	0:0
69268+1	bash	dp-dam	deepest-a+	384	COMPLETED	0:0

On the Cluster (CM) nodes it's possible to query the consumed energy for a certain job:

```
[kreutz1@deepv kreutz1]$ sacct -o ConsumedEnergy,JobName,JobID,CPUTime,AllocNodes -j 69326
```

ConsumedEnergy	JobName	JobID	CPUTime	AllocNodes
496.70K	hpl_MKL_0+	69326	16:28:48	1
0	batch	69326.batch	16:28:48	1
496.70K	xlinpack_+	69326.0	08:10:24	1

This feature will also be for the ESB nodes.

FAQ

Is there a cheat sheet for all main Slurm commands?

Yes, it is available [?here](#).

Why's my job not running?

You can check the state of your job with

```
scontrol show job <job id>
```

In the output, look for the Reason field.

You can check the existing reservations using

```
scontrol show res
```

How can I check which jobs are running in the machine?

Please use the `squeue` command (the `"-u $USER"` option to only list jobs belonging to your user id).

How do I do chain jobs with dependencies?

Please confer the `sbatch/srun` man page, especially the

```
-d, --dependency=<dependency_list>
```

entry.

Also, jobs can be chained after they have been submitted using the `scontrol` command by updating their Dependency field.

How can check the status of partitions and nodes?

The main command to use is `sinfo`. By default, when called alone, `sinfo` will list the available partitions and the number of nodes in each partition in a given status. For example:

```
[deamicis1@deepv hybridhello]$ sinfo
PARTITION    AVAIL  TIMELIMIT  NODES  STATE NODELIST
sdv           up    20:00:00    11    idle deeper-sdv[06-16]
knl           up    20:00:00     1  drain knl01
knl           up    20:00:00     3    idle knl[04-06]
knl256        up    20:00:00     1  drain knl01
knl256        up    20:00:00     1    idle knl05
knl272        up    20:00:00     2    idle knl[04,06]
snc4          up    20:00:00     1    idle knl05
extoll        up    20:00:00    11    idle deeper-sdv[06-16]
ml-gpu        up    20:00:00     3    idle ml-gpu[01-03]
dp-cn         up    20:00:00     1  drain dp-cn33
dp-cn         up    20:00:00     5   resv dp-cn[09-10,25,49-50]
dp-cn         up    20:00:00    44    idle dp-cn[01-08,11-24,26-32,34-48]
dp-dam        up    20:00:00     1  drain* dp-dam08
dp-dam        up    20:00:00     2  drain dp-dam[03,07]
dp-dam        up    20:00:00     3   resv dp-dam[05,09-10]
dp-dam        up    20:00:00     2  alloc dp-dam[01,04]
dp-dam        up    20:00:00     8    idle dp-dam[02,06,11-16]
dp-dam-ext    up    20:00:00     2   resv dp-dam[09-10]
dp-dam-ext    up    20:00:00     6    idle dp-dam[11-16]
dp-esb        up    20:00:00    51  drain* dp-esb[11,26-75]
dp-esb        up    20:00:00     2  drain dp-esb[08,23]
dp-esb        up    20:00:00     2  alloc dp-esb[09-10]
dp-esb        up    20:00:00    20    idle dp-esb[01-07,12-22,24-25]
dp-sdv-esb    up    20:00:00     2   resv dp-sdv-esb[01-02]
psgw-cluster  up    20:00:00     1    idle nfgw01
psgw-booster  up    20:00:00     1    idle nfgw02
debug         up    20:00:00     1  drain* dp-dam08
debug         up    20:00:00     4  drain dp-cn33,dp-dam[03,07],knl01
debug         up    20:00:00    10   resv dp-cn[09-10,25,49-50],dp-dam[05,09-10],dp-sdv-esb[01-02]
debug         up    20:00:00     2  alloc dp-dam[01,04]
debug         up    20:00:00    69    idle deeper-sdv[06-16],dp-cn[01-08,11-24,26-32,34-48],dp-dam[02,06,11-16],knl[04-06]
```

Please refer to the man page for `sinfo` for more information.

Can I join stderr and stdout like it was done with -joe in Torque?

Not directly. In your batch script, redirect stdout and stderr to the same file:

```
...  
#SBATCH -o /point/to/the/common/logfile-%j.log  
#SBATCH -e /point/to/the/common/logfile-%j.log  
...
```

(The %j will place the job id in the output file). N.B. It might be more efficient to redirect the output of your script's commands to a dedicated file.