
NorduGrid ARC 7 Documentation

Release ARC7

NorduGrid Collaboration

Oct 07, 2025

RELEASE NOTES

1	ARC tutorial	3
2	ARC Overview	5
3	Obtaining the software	7
4	Support and Community	15
5	General overview	17
6	Documentation for Infrastructure Admins	19
7	Technical Documents Describing ARC Components	283
8	Documentation for Developers	363
9	Documentation for Infrastructure Users	379
10	ARC Miscellaneous Pages	431

The *Advanced Resource Connector (ARC)* middleware, developed by the [NorduGrid Collaboration](#), is an open source software solution enabling e-Science computing infrastructures with emphasis on processing of large data volumes. ARC is being used to enable national and international e-infrastructures since its first release in 2002.

This document is dedicated to the ARC Version 7 collecting all relevant information in one place. You should be able to find information regarding the code, documentation, testing activities and support channels here. The information is refreshed daily, a snapshot of the development version can be found [here](#).

If you are new to ARC start reading the [Quickstart guide](#) quickstart guide to get an overview of main operations in the simple test case.

For a production *Computing Element* deployment follow the [Installation and Configuration Guide](#) that contains the structure and pointers to precise configuration of every ARC subsystem.

In case you are migrating to ARC 7 from an ARC 6 installation read the [Migration Guide](#). For an overview of the main changes compared to ARC 6, please visit [Main changes in ARC 7 compared to ARC 6](#).

The ultimate description of the ARC 7 configuration can be found in the [ARC Configuration Reference Document](#).

ARC TUTORIAL

In addition to the excellent [Quickstart guide](#) and [installation and configuration guide](#) you can find a new [tutorial](#) which will guide you through all necessary steps to set up a production ready system. This tutorial aims at showing some basic but essential configuration settings necessary for a common type of cluster configuration.

The tutorial was held first time at the EGI 2023 conference.

ARC OVERVIEW

A birds-eye overview of the ARC services, with an architecture figure of the ARC CE can be found in the [ARC CE components](#) overview.

OBTAINING THE SOFTWARE

ARC is available for variety of GNU/Linux flavors via stable *Reposiores* or *Nightly Builds* if you want to test the latest development release.

The source code is hosted in [NeIC's Coderefinery GitLab repository](#).

3.1 Release notes

3.1.1 Release Notes for NorduGrid ARC 7.1.1 03.10.2025

This is a bugfix release that addresses 4 bugs discovered in 7.1.0:

- arccp segfault (Bug 4271)
- Slurm 25 deprecation of options to the get-user-env was not fully fixed in 7.1.0 (get-user-env no longer supports any options at all) (Bug 4254)
- a glob pattern mismatch in submit_common.sh caused jobs to be aborted early (Bug 4270)
- A previous change to arc-job-cgroup causes get_cgroup_mount_root() to fail since the input arguments to fscanf() changed from 6 to 5 without changing the test value in the code from 6 to 5.

Note that GFAL2 is no longer supported in EL10. We therefore do not provide packages for GFAL for EL10 derivatives.

The following Bugzilla tickets have been addressed or solved

- 4271 arccp segfaults after upgrade to 7.1.0
- 4254 Job submission to slurm 25 is failing
- 4270 submit_common.sh: move_files_to_node and move_files_to_frontend glob patterns fail to match .* in Bash
- 4210 Tidy up arc-job-cgroup.c

Documentation

The ARC 7 documentation can be found at <https://www.nordugrid.org/arc/arc7>

If you miss something or have questions, please contact us!

Installing ARC 7

We recommend to install ARC release using the Nordugrid repository.

- Start with the basics: Token-version: <https://www.nordugrid.org/arc/arc7/admins/quickstart.html>
- Get production ready: https://www.nordugrid.org/arc/arc7/admins/arc_install_guide.html#arc-install-guide

There is also a nice ARC tutorial which you can follow if you are new to ARC: <https://www.nordugrid.org/arc/arc7/admins/tutorial/tutorial.html>

- We provide binary builds for all supported versions of Debian, Ubuntu, Fedora, Rocky (compatible with AlmaLinux), and CentOS Stream.

Contributing

You can submit pull requests to our public repo at <https://github.com/nordugrid/arc> which mirrors our Gitlab repo <https://source.coderefinery.org/nordugrid/arc>.

The pull requests will be automatically copied over to our Gitlab repo where they will be merged into the source-code if accepted.

Getting in touch

If you need help with the configuration file or have any other questions related to ARC, please contact us via * Email: nordugrid-discuss@nordugrid.org or wlcg-arc-ce-discuss@cern.ch * We have moved to Matrix for ARC support - you can always reach us there: https://matrix.to/#/#arc_support:matrix.org

Future support of ARC 6-series

Now that ARC 7.0.0 is released, we expect one final ARC 6 release, but after that we will only provide security updates of ARC 6. In particular: 1) No new feature development is planned or going on for ARC6 and no bug-fixing development

will happen on ARC6 code base in the future except for security issues.

- 2) Security fixes for ARC6 will be provided till end of March 2026.
- 3) Production Sites already running ARC 6 will be able to get deployment and configuration troubleshooting help via GGUS Helpdesk till end March 2027. This we call “operational site support”.
- 4) In EPEL ARC packages have their major release number as part of the package name, arc6 and arc7 packages can not be confused.
- 5) In the Nordugrid repository, you have to specifically enable the ARC 7 repo to get the ARC 7 update. See: <https://www.nordugrid.org/arc/arc7/common/repos/repository.html>

Known issues

- GFAL2 not supported in EL10

3.1.2 Release Notes for Nordugrid ARC 7.1.0 11.09.2025

We are happy to release ARC 7.1.0 which comes with a couple of important bugfixes and a few other minor improvements.

Highlights in this release

As of Slurm 25, there is no longer support for the timeout option to the `--get-user-env` parameter that the ARC Slurm submission scripts have been using. In ARC 7.1.0 this option has been removed, and the Slurm script now just uses `#SBATCH --get-user-env` with no custom timeout included. (Bugzilla ticket 4254).

During summer ARGO changed their REST API related to authentication, causing ARC jobs to fail accounting publishing. A patch was submitted to ARC which has been included in ARC 7.1.0 (but will not be included in ARC 6, see policy for support below). ARGO currently supports both old and new authentication API, so currently both ARC 7 and ARC 6 are functioning. However, we encourage upgrade to ARC 7 the sooner the better, to stay on the safe side. (Bugzilla ticket 4269).

The JSON structure of the published public ARC information has gradually been cleaned to follow the agreed schemas more strictly. One such change caused job submission failures from aCT. This has now been fixed (Bugzilla ticket 4256).

The following Bugzilla tickets have been addressed or solved

- 4269 ARGO REST header changes causing accounting problems
- 4258 Missing user.proxy filename
- 4256 Job submission from aCT to ARC fails with nightlies from 24.06.25
- 4254 Job submission to slurm 25 is failing
- 4253 Parsing authtokensgen expression
- 4251 Python3 bug in JURA parsing *endfrom*
- 4248 ARC 7 - Failed to start transfer request: SOAP: Failed processing request
- 4247 CEInfo REST Endpoint produces malformed JSON output if ARC CE doesn't have any RTEs enabled
- 4246 The command "`arcctl datastaging dtr state`" does not show all states
- 4243 Spelling of some messages from arcctl
- 4242 arcctl datastaging summary files crashes
- 4240 arc/utis/config.py is not using raw string literals for regular expressions
- 4239 arcls asks for password unnecessarily
- 4237 arc-datadelivery-service crashes
- 4222 Tripple warning when running arcctl as client

Documentation

The ARC 7 documentation can be found at <https://www.nordugrid.org/arc/arc7>

If you miss something or have questions, please contact us!

Installing ARC 7

We recommend to install ARC release using the Nordugrid repository.

- Start with the basics: Token-version: <https://www.nordugrid.org/arc/arc7/admins/quickstart.html>
- Get production ready: https://www.nordugrid.org/arc/arc7/admins/arc_install_guide.html#arc-install-guide

There is also a nice ARC tutorial which you can follow if you are new to ARC: <https://www.nordugrid.org/arc/arc7/admins/tutorial/tutorial.html>

- We provide binary builds for all supported versions of Debian, Ubuntu, Fedora, Rocky (compatible with AlmaLinux), and CentOS Stream.

Contributing

You can submit pull requests to our public repo at <https://github.com/nordugrid/arc> which mirrors our Gitlab repo <https://source.coderefinery.org/nordugrid/arc>.

The pull requests will be automatically copied over to our Gitlab repo where they will be merged into the source-code if accepted.

Getting in touch

If you need help with the configuration file or have any other questions related to ARC, please contact us via * Email: nordugrid-discuss@nordugrid.org or wlcg-arc-ce-discuss@cern.ch * We have moved to Matrix for ARC support - you can always reach us there: https://matrix.to/#/#arc_support:matrix.org

Future support of ARC 6-series

Now that ARC 7.0.0 is released, we expect one final ARC 6 release, but after that we will only provide security updates of ARC 6. In particular: 1) No new feature development is planned or going on for ARC6 and no bug-fixing development

will happen on ARC6 code base in the future except for security issues.

- 2) Security fixes for ARC6 will be provided till end of March 2026.
- 3) Production Sites already running ARC 6 will be able to get deployment and configuration troubleshooting help via GGUS Helpdesk till end March 2027. This we call “operational site support”.
- 4) In EPEL ARC packages have their major release number as part of the package name, arc6 and arc7 packages can not be confused.
- 5) In the Nordugrid repository, you have to specifically enable the ARC 7 repo to get the ARC 7 update. See: <https://www.nordugrid.org/arc/arc7/common/repos/repository.html>

3.1.3 Release Notes for NorduGrid ARC 7.0.0 27.03.2023

We are happy to announce the release of ARC 7.0.0, with REST as the primary interface, and with token support in addition to x509.

ARC 7 is backwards compatible with ARC 6 except for the deprecated components. This means that your arc configuration file will not need major changes. But pay attention when you restart your services after update, so that you can remove any blocks and options that are no longer supported.

The larger internal change in ARC 7 is the restructuring of the job metadata folder, the ARC control directory.

The restructuring will be done seamlessly from your ARC 6 control directory to the new restructured ARC 7 control directory. If you have any custom scripts that use per-job metadata files, make sure to update these to the

new structure: https://www.nordugrid.org/arc/arc7/tech/arex_internals/internal_files.html You can use the `arcctl` tool to fetch the directory path: `arcctl job path <jobid>`.

For more details about the changes in ARC 7: https://www.nordugrid.org/arc/arc7/common/changelog/arc6_to_arc7_changes.html

To upgrade from ARC 6, consult the migration guide: <https://www.nordugrid.org/arc/arc7/admins/migrating6to7.html#migration-6-to-7>

Highlights in this release

- Tokens (http://www.nordugrid.org/arc/arc7/tech/oidc_tokens/oidc_tokens.html)
 - Major improvements and added features compared to the token preview in ARC 6
 - Extended support for using token information (scopes) for authorization and identity mapping
 - Token based authentication is now enabled by default together with x509
 - Token information is part of the accounting records
 - Token claims can be used to link a token to a VO for accounting - either automatically if the claim `wlwg.group` is present, or manually in `arc.conf` with the `vomsless_vo` option (see below).
- New or changed options in `arc.conf`
 - `usetokenforvoms`: whether claims from WLCG compliant tokens should be used as VOMS attributes - currently `wlwg.group` claim is supported
 - `tokenscopes`: assigns JWT token scopes to perform specific actions
 - `authtokenmap`: map your token claims to some custom chosen attribute name for accounting
 - `authtokensgen`: advanced possibility using logical expression in order to map claims in a token to an authgroup.
 - `tlsserverorder`: force priority order of ciphers for TLS connections to be decided on server side
 - `vomsless_vo`: Added additional option to `vomsless_vo` configuration command. It is now possible to assign VO to jobs submitted without VOMS proxy differently for users matching different authgroups. This means that you can use this to assign an authgroup to a certain VO. For the token use-case you would first match a token to an authgroup using any of the available methods, and then use the `vomsless_vo` to assign the authgroup to the appropriate VO.
- New blocks in `arc.conf` - `infosys/accesscontrol`: AREX allows to control access to public information for non-authorized users. If this block contains no entries, public information is available to anyone. The `allowaccess` and `denyaccess` options are used to control access.
- The `ldap` interface is still present in ARC 7, but marked as legacy.
- The ARC server can use operating system provided CA bundles instead of the traditional IGTF CAs via the configurable `x509_cert_policy` in the common block in `arc.conf`. - An ARC-CE must currently use *either* the default grid option *or* the system option.
- REST 1.1 (<http://www.nordugrid.org/arc/arc7/tech/rest/rest.html>) - A new optional delegation parameter type extension to support tokens - A new optional queue parameter to define the default queue of the compute element - Support for bulk job submission of identical jobs
- A tidy up has been done to gain consistent handling of the benchmark parameters over the various `arc.conf` blocks
- A restructured `controldir` for increased scalability (https://www.nordugrid.org/arc/arc7/tech/arex_internals/internal_files.html) - When upgrading from ARC 6 to ARC 7, the `controldir` will seamlessly restructure itself.
- The `arc` control tool (`arcctl`) has several new capabilities. - `test token issuer (test-jwt)` - similar to the `test CA functionality (test-ca)` - `cleanup` functionality, to cleanup job metadata in control dir, `test-jwt` and `test-ca`

related files and accounting database - `arcctl config verify` behaviour has changed, it now only verifies `arc.conf` - `arcctl service verify` behaves as `arcctl config verify` earlier did, checks `arc.conf`, time sync and certificates

- The accounting subsystem in ARC7 is updated to improve operational efficiency - the database has undergone a redesign. Please see https://www.nordugrid.org/arc/arc7/common/changelog/arc6_to_arc7_changes.html for details.

Removed and obsoleted components:

- Server-side gridftp-based ARC solutions both for job management and data storage
- Both server and client-side support for EMIES job management interface- ARC Cache Index (ACIX), that was a distributed system that maintained a catalog of locations of cached files stored in various A-REX caches.
- userlists generation and mapping support by `nordugridmap` utility
- Archiving of the JURA accounting service (got removed already in the 6.8.0 release)
- GLUE1 and site-BDII support within the LDAP-based components of the ARC information system
- Remote batch system backend support via SSH
- Python-based LRMS backend framework (it never progressed beyond technology preview status)
- Support for ARGUS via AREX/WS component
- Support for python2
- LDAP information system support in the ARC7 clients
- EGIS resource discovery in the ARC 7 clients
- The following components and functionalities have been kept but labelled as DEPRECATED and target for near-future removal:
- OBSOLETE: LDAP-based server-side information system with GLUE2 and `nordugrid` schema rendering
- OBSOLETE: gridftp protocol for data staging (still available in ARC 7 and provided as a separate plugin package)
- OBSOLETE: Support for Loadleveler and LSF batch systems

Backwards compatibilities and incompatibilities

- ARC7 still supports x509 proxies in addition to tokens.
- Both authtokens and x509 authentication are enabled by default.
- Support for the EMIES interface and the gridftp job interface was completely removed.
- Publishing cluster information via REST interface is enabled by default [`arex/ws/publicinfo`] block is added to `zero.conf`
- No incompatible changes in the ARC7 accounting subsystem. APEL and SGAS supported.
- The internal structure of the `controldir` files have been changed, ARC internal scripts were updated but note that any custom scripts may break because of the changed file structure in the `controldir`.
- ARC 7 supports XRSL and ADL as job description languages, no backward incompatible changes were introduced, all ARC6 job files should work with ARC7.
- ARC 7 fully supports the SLURM, Condor and the fork batch systems and community efforts were made to provide some maintenance for OpenPBS as well. The other batch systems from ARC6 are kept in ARC7 but their status is unclear.

Documentation

The ARC 7 documentation can be found at <https://www.nordugrid.org/arc/arc7>

If you miss something or have questions, please contact us!

Installing ARC 7

We recommend to install ARC release using the Nordugrid repository.

- Start with the basics: Token-version: https://www.nordugrid.org/arc/arc7/admins/quickstart_jwt.html, x509-version: https://www.nordugrid.org/arc/arc7/admins/quickstart_x509.html
- Get production ready: https://www.nordugrid.org/arc/arc7/admins/arc_install_guide.html#arc-install-guide

There is also a nice ARC tutorial which you can follow if you are new to ARC: <https://www.nordugrid.org/arc/arc7/admins/tutorial/tutorial.html>

- We provide binary builds for all supported versions of Debian, Ubuntu, Fedora, Rocky (compatible with AlmaLinux), and CentOS Stream.

Contributing

You can submit pull requests to our public repo at <https://github.com/nordugrid/arc> which mirrors our Gitlab repo <https://source.coderefinery.org/nordugrid/arc>.

The pull requests will be automatically copied over to our Gitlab repo where they will be merged into the source-code if accepted.

Getting in touch

If you need help with the configuration file or have any other questions related to ARC, please contact us via * Email: nordugrid-discuss@nordugrid.org or wlcg-arc-ce-discuss@cern.ch * You can still reach us on our Skype channel dedicated to ARC support: <https://join.skype.com/dyf3A6Uutjy2>

- We will eventually be migrating to another service since Skype is shutting down in May, the new contact info will be announced later.

Future support of ARC 6-series

Now that ARC 7.0.0 is released, we expect one final ARC 6 release, but after that we will only provide security updates of ARC 6. In particular: 1) No new feature development is planned or going on for ARC6 and no bug-fixing development will happen on ARC6 code base in the future except for security issues. 2) Security fixes for ARC6 will be provided till end of March 2026. 3) Production Sites already running ARC 6 will be able to get deployment and configuration troubleshooting help via GGUS Helpdesk till end March 2027. This we call “operational site support”. 4) In EPEL ARC packages have their major release number as part of the package name, arc6 and arc7 packages can not be confused. 5) In the Nordugrid repository, you have to specifically enable the ARC 7 repo to get the ARC 7 update. See: <https://www.nordugrid.org/arc/arc7/common/repos/repository.html>

SUPPORT AND COMMUNITY

User support and site installation assistance is provided via the [nordugrid-discuss mailing list](#), and the [Nordugrid Bugzilla](#).

GENERAL OVERVIEW

5.1 ARC CE overview

5.1.1 The components and the infrastructure ecosystem around

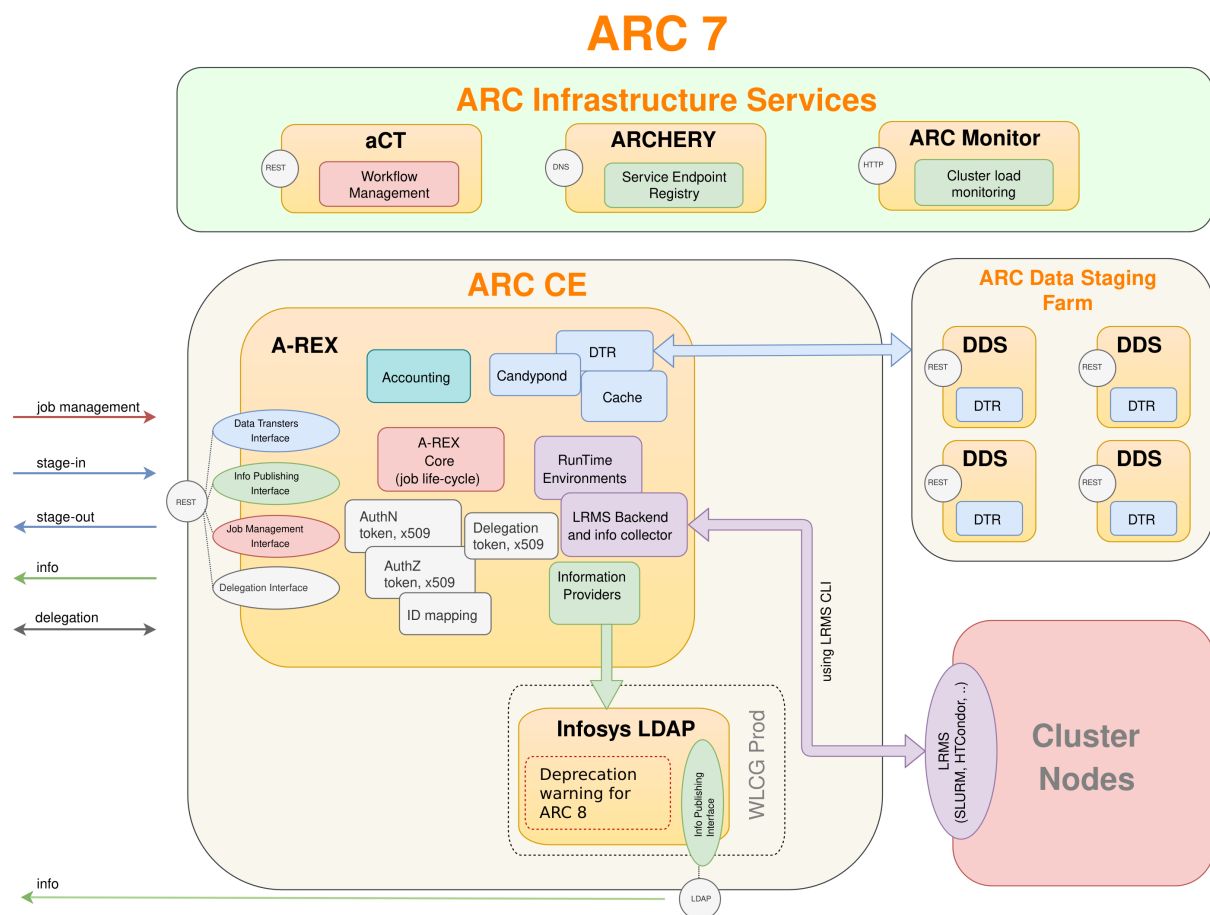


Fig. 5.1: ARC7 CE: internals, interfaces and the infrastructure ecosystem services around

DOCUMENTATION FOR INFRASTRUCTURE ADMINS

This section contains the full list of documentation related to ARC middleware services deployment, configuration and operations. If you are looking for the ARC Computing Element setup instruction or performance tuning parameters, you are in the right place.

6.1 ARC Configuration Reference Document

6.1.1 General configuration structure

This is the `arc.conf` REFERENCE DOCUMENT defining the configuration blocks and configuration options for the ARC services.

The `arc.conf` configuration file consists of the following blocks:

```
[common]
[authgroup:groupname]
[mapping]
[lrms]
[arex]
[arex/cache]
[arex/cache/cleaner]
[arex/data-staging]
[arex/ws]
[arex/ws/jobs]
[arex/ws/cache]
[arex/ws/candypond]
[arex/jura]
[arex/jura/sgas:targetname]
[arex/jura/apel:targetname]
[arex/ganglia]
[infosys]
[infosys/ldap]
[infosys/nordugrid]
[infosys/glue2]
[infosys/glue2/ldap]
[infosys/cluster]
[infosys/accesscontrol]
[queue:name]
[datadelivery-service]
[custom:blockname]
```

If `arc.conf.d` directory exists next to `arc.conf` file, all files in this directory ending with `.conf` will be read in alphabetical order and their content merged with the `arc.conf` file.

[block]

A block configures an ARC service, a service interface, a utility or a subsystem. Enabling (turning on) a functionality, a service or an interface requires the presence of the appropriate configuration block. To disable a service or an interface, simply delete or comment out the related `arc.conf` block (you may need to rerun the corresponding startup script).

The `[common]` block is mandatory even if not a single option is specified within. The presence of the block turns on the default values for the configuration options within the block.

As an example, in order to set up a minimalistic ARC CE offering no external interfaces you need to configure at least the `[common]`, `[mapping]`, `[arex]`, `[lrms]`, `[infosys]` and `[queue:name]` blocks.

As another example, an ARC-based data offloader would require the `[common]` and the `[datadelivery-service]` blocks.

A block is identified by its block header. A block header may consist of keywords and optionally block identifiers. Keywords may be separated by `/` and used to label subblocks (e.g. `[arex/jura]`), while block identifiers are separated by `:` from keywords. For example, in the `[queue:short]` block header `queue` is a keyword while `short` is an identifier, e.g. the name of the queue. Block headers must be UNIQUE.

A block starts with a unique `[keyword:identifier]` blockheader and ends where the next block starts, that is at the next `[blockheader]` directive.

A block may have sub-blocks e.g. the various interfaces of the AREX service are configured via sub-blocks (e.g. `[arex/ws]`). When a sub-block is enabled then the corresponding parent block must also appear in the `arc.conf` file.

Configuration blocks contain (config option, config value) pairs following the syntax in single line:

```
config_option=value element [optional value element]
```

Note

quotes around the configuration value(s) must NOT be used.

Note

the `arc.conf` is CASE-SENSITIVE!

Space handling syntax in `arc.conf` for configuration lines:

```
(stripped space)option(stripped space)=(stripped space)value(saved_
↪space)(value)(stripped space)
```

and for block headers:

```
[keyword:(stripped space)space is NOT allowed within identifier(stripped space)]
```

Detailed textual definition:

- All trailing and leading spaces on each configuration line are stripped and ignored. This applies both to block headers and block content.
- All spaces around the `=` sign in `option=value` kind of string (after 'a' is applied) are stripped and ignored. For example line `hostname = myhost.info` is treated as identical to `hostname=myhost.info`.
- In block headers of `[keyword]` kind (after 'a' is applied) no additional spaces are allowed around `keyword` and inside `keyword`.

- d) In block headers of `[keyword:identifier]` kind (after 'a' is applied) no additional spaces are allowed around `keyword` and inside both `keyword` and `identifier`. The spaces ARE allowed around `identifier` part and stripped and ignored.

Mandatory configuration options are indicated by an asterix prefix to the option name e.g: `*mandatory_configoption`. Mandatory options with undefined values will result in service stop during the startup process.

Each of the configuration options have well-defined default that is specified in this reference file. The default can take either a pre-set value, a special substitution or the keyword `undefined`. Configuration options within an enabled block take their default values in case they are missing (or commented out). Configuration parameters with `undefined` defaults takes no values. Furthermore, configuration options within disabled blocks takes no values either.

Configuration blocks are ORDER-DEPENDENT. The order dependency is also honoured within options inside a certain block. This means for instance that configuration blocks related to authorization **MUST** appear before used in the blocks such as `[mapping]`, `[arex/ws/jobs]` or `[gridftp/jobs]`. Order dependency within a block is for instance important when it comes to authorization decisions, as the first matching rule is used.

ARC configuration parser makes sure that blocks are **AUTOMATICALLY RE-ORDERED** in accordance to the defaults-defined order of block keywords. Several blocks with the same keyword but different identifiers (named blocks) will be sorted in according to read-out order.

When the same block is defined in the several files in `arc.conf.d`` directory, configuration options inside this block are appended in the read-out order. Complex logic (like overrides) is not supported. It is advised to use `arcctl config dump` to verify the desired running configuration after merge.

Below we give a detailed description of all the configuration options of the different configuration blocks. Every configuration option is described in a dedicated paragraph with the following reference syntax notation. This file is parsed at buildtime to assist in configuration default parsing and validation script and so it is important that it follows the agreed syntax: For each block or option please add explanatory text with two `##` followed by a space at the beginning of the line and then an example with a single `#` and no spaces at the beginning of the line.

example_config_option

Synopsis: `example_config_option = value [optional values]`

Description: Here comes the explanation of the config option. Mandatory configuration options are indicated by an asterix prefix to the option name e.g: `*mandatory_configoption` vs. `optional_configoption`. The explanation can be followed by the special keywords in a separate line:

- **multivalued** - used to indicate that config option can be specified multiple times. This forms a set of values for the same configuration option irrespective of lines order.
- **sequenced** - used to indicate that config option is a part of the sequence and its effect on configuration depends on the lines order. Sequenced option can be specified several times in the configuration sequence independently.

Missing such keywords means the config option can only occur once in the `arc.conf`. By default the `arc.conf` config options are optional and single-valued. For some config options only a fix set of values are allowed. These are listed in a separate line after the `allowedvalues` keyword. The default of every config option is explicitly given in the `default:` line. Default can be a pre-set value, a substitution or the `undefined` keyword. The last line of the paragraph is always a valid example preceded by a single `#`

This option in **multivalued**.

Allowed values: 12, 34, 56

Default: 34

Example:

```
example_config_option=56
```

6.1.2 Configuration blocks and options

[common] block

Common configuration affecting all ARC components, usually related to networking or security or service behaviour. This block is mandatory. The common block options may be overridden by the specific sections of the components later. The [common] always appears at the beginning of the config file. The config options set within this block are available for all the other blocks thus shared by the different components of ARC.

hostname

[common]

Synopsis: hostname = string

Description: The FQDN of the frontend on which the ARC services are deployed.

Default: \$EXEC{hostname -f}

Example:

```
hostname=myhost.org
```

http_proxy

[common]

Synopsis: http_proxy = url

Description: The http proxy server. This setting affects all client HTTP(s) requests that initiated by ARC core services, including data staging, SAML communications, and pushing SGAS accounting records. This variable is similar to setting the ARC_HTTP_PROXY environmental variable.

Default: undefined

Example:

```
http_proxy=proxy.mydomain.org:3128
```

x509_host_key

[common]

Synopsis: x509_host_key = path

Description: Server credential location. Sets the full path to the host private key. These variables are similar to the GSI environment variable X509_USER_KEY. If indicated, the variable can be set individually for each service/component in the corresponding block.

Default: /etc/grid-security/hostkey.pem

Example:

```
x509_host_key=/etc/grid-security/hostkey.pem
```

x509_host_cert

[common]

Synopsis: x509_host_cert = path

Description: Server credential location. Sets the full path to the host public certificate. These variables are similar to the GSI environment variable X509_USER_CERT. If indicated, the variable can be set individually for each service/component in the corresponding block.

Default: /etc/grid-security/hostcert.pem

Example:

```
x509_host_cert=/etc/grid-security/hostcert.pem
```

x509_cert_policy

[common]

Synopsis: x509_cert_policy = keyword

Description: layout of CA certificates. The following keywords are supported: grid, system, any. This variable defines if the server is going to use Globus (igtf) layout of CA certificates, just let OpenSSL handle that (system), or allow both options (any).

Default: grid

Example:

```
x509_cert_policy=grid
```

x509_cert_dir

[common]

Synopsis: x509_cert_dir = path

Description: Location of trusted CA certificates. This variable is similar to the GSI environment variable X509_CERT_DIR. If indicated, the variable can be set individually for each service/component in the corresponding block. If x509_cert_policy is set to 'system' this variable is ignored.

Default: /etc/grid-security/certificates

Example:

```
x509_cert_dir=/etc/grid-security/certificates
```

x509_voms_dir

[common]

Synopsis: x509_voms_dir = path

Description: the path to the directory containing *.lsc files needed for verification of VOMS service signature in the proxy-certificate.

Default: /etc/grid-security/vomsdir

Example:

```
x509_voms_dir=/etc/grid-security/vomsdir
```

voms_processing

[common]

Synopsis: voms_processing = keyword

Description: Defines how to behave if errors in VOMS AC processing detected. The following keywords are supported:

relaxed

use everything that passed validation.

standard

same as relaxed but fail if parsing errors took place and VOMS extension is marked as critical.
This is a default.

strict

fail if any parsing error was discovered

noerrors

fail if any parsing or validation error happened.

Allowed values: relaxed, standard, strict, noerrors

Default: standard

Example:

```
voms_processing=strict
```

[authgroup:groupname] block

These configuration blocks contain authorization rules. An [authgroup:groupname] block always defines a group of users where members of the group are those who satisfy the authorization rules. The rules within the block determine which user belong to the authgroup. Then, access control and identity mapping of ARC services are implemented via associating a authgroup with an interface, queue or a mapping rule using one of the allowaccess, denyaccess or [mapping] block parameters. The authgroup should not be mistaken for a virtual organisation (VO). An authgroup may match a single VO if only a single check (rule) on VO membership is performed.

IMPORTANT: Rules in an authgroup are processed in their order of appearance. The first matching rule decides the membership of the user to the authgroup being evaluated and the processing STOPS within that authgroup. This does not mean that the same user is not processed for the next authgroup: all [authgroup:groupname] blocks are evaluated, even if a user already has a match with one of the earlier groups.

All the objects used in the rules **MUST** be defined before it may be used. For example, to create group of authgroups you must first defined the child groups.

There are positively and negatively matching rules. If a rule is matched positively then the user tested is accepted into the respective group and further processing is stopped. Upon a negative match the user would be rejected for that group - processing stops too. The sign of rule is determined by prepending the rule with + (for positive) or - (for negative) signs. + is default and can be omitted. A rule may also be prepended with ! to invert result of rule, which will let the rule match the complement of users. That complement operator (!) may be combined with the operator for positive or negative matching.

subject

[authgroup:groupname]

Synopsis: subject = certificate_subject

Description: Rule to match specific subject of user's X.509 certificate. No masks, patterns and regular expressions are allowed.

This is **sequenced** option.

Default: undefined

Example:

```
subject=/O=Grid/O=Big VO/CN=Main Boss
subject=/O=Grid/O=Big VO/CN=Deputy Boss
```

file

[authgroup:groupname]

Synopsis: file = path

Description: Processes a list of DNs stored in an external file one per line in grid-mapfile format (see map_with_file from [mapping] block, unixname is ignored) and adds those to the authgroup.

This is **sequenced** option.

Default: undefined

Example:

```
file=/etc/grid-security/local_users
file=/etc/grid-security/atlas_users
```

voms

[authgroup:groupname]

Synopsis: voms = vo group role capabilities

Description: Match VOMS attribute in user's credential. Use * to match any value.

This is **sequenced** option.

Default: undefined

Example:

```
voms=nordugrid Guests * *
voms=atlas students prodman *
```

authgroup

[authgroup:groupname]

Synopsis: authgroup = group_name [group_name ...]

Description: Match user already belonging to one of specified authgroups. The authgroup referred here must be defined earlier in `arc.conf` configuration file. Multiple authgroup names may be specified for this rule. That allows creating hierarchical structure of authorization groups like `all-atlas` are those which are `atlas-users` and `atlas-admins`.

This is **sequenced** option.

Default: undefined

Example:

```
authgroup=local_admins
authgroup=local_admins remote_users
```

plugin

[authgroup:groupname]

Synopsis: plugin = timeout path [arg1 [arg2 [arg3...]]]

Description: Run external executable or function from shared library. Rule is matched if plugin returns 0. Any other return code or timeout are treated as rule not matched. In arguments following substitutions are supported:

- %D - subject of certificate
- %P - path to proxy

The environment variables passed to plugin contain basic information about user authentication. Following variables are set if corresponding information is available:

- X509_SUBJECT_NAME - common name of user's certificate.
- BEARER_TOKEN_#_SUBJECT - user's subject (identifier) extracted from JWT token (here # is tokens index, typically 0)
- BEARER_TOKEN_#_ISSUER - issuer of the token extracted from JWT token
- BEARER_TOKEN_#_AUDIENCE_# - designated audiences extracted from JWT token (here second # is audience's index starting from 0)
- BEARER_TOKEN_#_SCOPE_# - assigned scope extracted from JWT token (here second # is scope's index starting from 0)
- BEARER_TOKEN_#_GROUP_# - assigned WLCG group extracted from JWT token
- BEARER_TOKEN_#_CLAIM_<name>_# - raw claim values of the token of claim <name>

ARC ships with LCAS plugin that can be enabled with following plugin configuration. For more information about configuring LCAS refer to 'Using LCAS/LCMAPS' document.

This is **sequenced** option.

Default: undefined

Example:

```
plugin=10 /usr/libexec/arc/arc-lcas %D %P liblcas.so /usr/lib64 /etc/lcas/lcas.db
```

Warning

CHANGE: NEW environment variables in 7.0.0.

authtokens`[authgroup:groupname]`*Synopsis:* `authtokens = subject issuer audience scope group`*Description:* Match OIDC token claims. Use * to match any value.This is **sequenced** option.*Default:* undefined*Example:*`authtokens=e83eec5a-e2e3-43c6-bb67-df8f5ec3e8d0 https://wlcg.cloud.cnaf.infn.it/ * * *`**authtokensgen**`[authgroup:groupname]`*Synopsis:* `authtokensgen = logical expression`*Description:* Match OIDC token claims. Expression to match. Following operators are available:

- `=` - match token claim value (left part represents claim name) to specified string (right part), produces boolean result
- `~` - match token claim value (left part represents claim name) to regex expression (right part), produces boolean result
- `!` - boolean negation
- `&` - boolean AND
- `-` - boolean OR
- `^` - boolean XOR
- `()` - brackets are used to control priority of evaluation, without brackets all operators have same priority
- ```` - strings can be enclosed in quotes to allow special symbols in strings, to include `` in such string use ``

All empty spaces are optional. This functionality is experimental.

This is **sequenced** option.*Default:* undefined**Warning**

CHANGE: NEW in 7.0.0

all

[authgroup:groupname]

Synopsis: all = yes|no

Description: Matches any or none user identity. For yes argument this rule always returns positive match. For no it is always no match.

This is **sequenced** option.

Default: undefined

Example:

```
all=yes
```

[mapping] block

This block defines the grid-identity to local UNIX identity mapping rules used by various ARC components.

Rules in the [mapping] block are processed IN A SEQUENCE in line order of the configuration file (from top to bottom).

There are two kind of rules:

- mapping rules that defines how the particular authgroup members are mapped
- policy rules that modifies the mapping rules sequence processing

Default policy for mapping rules processing is:

- processing CONTINUES to the next rule if identity of user DO NOT match authgroup specified in the rule (can be redefined with policy_on_nogroup option)
- processing STOPS if identity of user match authgroup specified in the mapping rule. Depend on whether this mapping rule returns valid UNIX identity the processing can be redefined with policy_on_map and policy_on_nomap options.

Policy can be redefined at the any point of configuration sequence and affects all mapping rules defined after the policy rule.

Note

if mapping process STOPS and there is still no local UNIX identity identified, the user running A-REX will be used.

Note

when grid-identity is mapped to root account - request processing fails implicitly.

map_to_user

[mapping]

Synopsis: map_to_user = authgroup_name unixname[:unixgroup]

Description: the users that belongs to specified authgroup are mapped to unixname local UNIX account that may be optionally followed by a unixgroup UNIX group. In case of non-existing unixname account the mapping rule treated as a rule that did not returned mapped identity (nomap).

This is **sequenced** option.

Default: undefined

Example:

```
map_to_user=authgroupA nobody:nobody
```

map_to_pool

[mapping]

Synopsis: map_to_pool = authgroup_name directory

Description: the user that belong to specified authgroup is assigned one of the local UNIX accounts in the pool. Account names that are part of this pool are stored line-by-line in the pool file inside the directory. The directory also contains information about used account names stored in another files. If there are no more available accounts in the defined pool for mapping then accounts not used for a configurable time period may be reassigned. The pool behaviour, including account reuse, is configureable with the optional directory/config file that has INI syntax (line-by-line key=value). Possible keys of the config file are:

timeout

Define the timeout in days (default is 10) after which the UNIX account can be reassigned to another user if not used. The 0 value means no lease expiration.

This is **sequenced** option.

Default: undefined

Example:

```
map_to_pool=atlas /etc/grid-security/pool/atlas
```

map_with_file

[mapping]

Synopsis: map_with_file = authgroup_name file

Description: for users that belongs to specified authgroup the DN of certificate is matched against a list of DNs stored in the specified file, one per line followed by a local UNIX account name. The DN must be quoted if it contains blank spaces. This rule can be used to implement legacy grid-mapfile aproach.

This is **sequenced** option.

Default: undefined

Example:

```
map_with_file=authgroupB /etc/grid-security/grid-mapfile
```

map_with_plugin

[mapping]

Synopsis: `map_with_plugin = authgroup_name timeout plugin [arg1 [arg2 [...]]]`

Description: run external `plugin` executable with specified arguments to find the UNIX account name to which users that belongs to specified `authgroup` will be mapped to. A rule matches if the exit code is 0 and there is a UNIX account name printed on stdout (optionally followed by a UNIX group name separated by colon). The exit code 1 designates failed mapping. Any other code or timeout means fatal failure and will abort any further mapping processing. That will also cause rejection of corresponding connection. Plugin execution time is limited to `timeout` seconds. The environment variables passed to plugin contain basic information about user authentication. For description of those variables see ‘plugin’ command from [authgroup] section.

In the arguments the following substitutions are applied before the plugin is started:

- %D - subject of user’s certificate,
- %P - path to credentials’ proxy file.

The environment variables passed to the plugin contain basic information about user authentication. Following variables are set if corresponding information is available:

- X509_SUBJECT_NAME - common name of user’s certificate.
- BEARER_TOKEN_#_SUBJECT - user’s subject (identifier) extracted from JWT token (here # is tokens index, typically 0)
- BEARER_TOKEN_#_ISSUER - issuer of the token extracted from JWT token
- BEARER_TOKEN_#_AUDIENCE_# - designated audiences extracted from JWT token (here second # is audience’s index starting from 0)
- BEARER_TOKEN_#_SCOPE_# - assigned scope extracted from JWT token (here second # is scope’s index starting from 0)
- BEARER_TOKEN_#_GROUP_# - assigned WLCG group extracted from JWT token
- BEARER_TOKEN_#_CLAIM_<name>_# - raw claim values of the token of claim <name>

ARC ships with LCMAPS plugin that can be enabled with the corresponding configuration. For more information about configuring LCMAPS refer to ‘Using LCAS/LCMAPS’ document.

This is **sequenced** option.

Default: undefined

Example:

```
map_with_plugin=authgroupC 30 /usr/libexec/arc/arc-lcmaps %D %P liblcmaps.so /usr/
lib64 /etc/lcmaps/lcmaps.db arc
```

policy_on_nomap

[mapping]

Synopsis: `policy_on_nomap = continue/stop`

Description: redefines mapping rules sequence processing policy in case identity of user match `authgroup` specified in the mapping rule and mapping rule DO NOT return valid UNIX identity. Default policy is `stop` processing the further rules. For example this policy will be triggered if pool is depleted, certificate subject is missing in the map file used for defined `authgroup` or plugin execution failed.

This is **sequenced** option.

Default: undefined

Allowed values: continue, stop

Example:

```
policy_on_nomap=continue
```

policy_on_map

[mapping]

Synopsis: policy_on_map = continue/stop

Description: redefines mapping rules sequence processing policy in case identity of user match authgroup specified in the mapping rule and mapping rule return valid UNIX identity. Default policy is stop processing the further rules. This policy will be triggered if rule successfully returns the result (allocated in pool, matched in map file, plugin call was successful).

This is **sequenced** option.

Default: undefined

Allowed values: continue, stop

Example:

```
policy_on_map=stop
```

policy_on_nogroup

[mapping]

Synopsis: policy_on_nogroup = continue/stop

Description: redefines mapping rules sequence processing policy in case identity of user DO NOT match authgroup specified in the mapping rule. Default policy is continue processing the further rules.

This is **sequenced** option.

Default: undefined

Allowed values: continue, stop

Example:

```
policy_on_nogroup=stop
```

[lrms] block

This block specifies the characteristics of the Local Resource Manager System (batch system) underneath the ARC CE. This block contains all the lrms-specific parameters and information. Configuration values in this block are available for A-REX, the backends, accounting and infosys ARC subsystems.

ARC primarily supports SLURM, Condor and its own fork LRMS flavours, as listed below. Please note that the various PBS flavours have best effort community support only.

lrms

[lrms]

Synopsis: `*lrms = lrmstype [defaultqueue]`

Description: Sets the type of the LRMS (queue system) and optionally the default queue name. ONLY ONE LRMS IS ALLOWED. MULTIPLE LRMS ENTRIES WILL TRIGGER UNEXPECTED BEHAVIOUR.

For lrmstype, the following values can be chosen:

- fork - simple forking of jobs to the same node as the server
- condor - Condor
- pbs - PBS (covers Torque and other old PBS flavours e.g. OpenPBS, older PBSPro, etc)
- pbspro - Altair PBS Professional
- slurm - SLURM

The optional defaultqueue parameter specifies the name of an existing LRMS queue in the cluster that will be used by AREX as the default queue to submit grid jobs in case the client does not specify a queue name during the job submission process. This queue name must match one of the [queue:queue_name] blocks.

Allowed values: fork, sge, condor, pbs, pbspro, lsf, ll, slurm, boinc

Default: undefined

mandatory

Example:

```
lrms=pbspro gridlong
lrms=slurm
```

lrmsconfig

[lrms]

Synopsis: `lrmsconfig = text`

Description: An optional free text field to describe the configuration of your Local Resource Management System (batch system). The value is published in the infosys, and is not used otherwise.

Default: undefined

Example:

```
lrmsconfig=single job per processor
```

benchmark

[lrms]

Synopsis: `benchmark = string`

Description: Defines the default benchmark specification to store in the accounting AAR records if per-job data is missing. It is advised to set it to cluster-wide defaults in case of reporting to APEL to avoid records diversity for failed jobs or buggy backends.

Default: HEPSPEC 1.0

Example:

```
benchmark=HEPSPEC 12.26
```

Warning

CHANGE: MODIFIED in 7.0.0

defaultmemory

[lrms]

Synopsis: defaultmemory = number

Description: The LRMS memory request of job to be set by the LRMS backend scripts, if a user submits a job without specifying how much memory should be used. The order of precedence is: job description -> defaultmemory. This is the amount of memory (specified in MB) that a job will request.

Default: undefined

Example:

```
defaultmemory=512
```

nodename

[lrms]

Synopsis: nodename = path

Description: Redefine the command to obtain hostname of LRMS worker node. By default the value is defined on buildtime and depend on the OS. In most cases `/bin/hostname -f` will be used.

Note

this way of getting WN hostname will be used only in case of particular LRMS backend had no native LRMS-defined way.

Default: undefined

Example:

```
nodename=/bin/hostname -s
```

gnu_time

[lrms]

Synopsis: gnu_time = path

Description: Path to the GNU time command on the LRMS worker nodes. If time command exists on the node, jobscript will write additional diagnostic information.

Default: /usr/bin/time

Example:

```
gnu_time=/usr/bin/time
```

movetool

[lrms]

Synopsis: `movetool = comand`

Description: Redefine the command used to move files during jobscript execution on LRMS worker node (the command should be available on WNs). This in particular applies to files movement from sessiondir to scratchdir in the shared sessiondir case.

Default: `mv`

Example:

```
movetool=rsync -av
```

pbs_bin_path

[lrms]

Synopsis: `pbs_bin_path = path`

Description: The path to the qstat,pbsnodes,qmgr etc PBS binaries, no need to set if PBS is not used

Default: `/usr/bin`

Example:

```
pbs_bin_path=/usr/bin
```

pbs_log_path

[lrms]

Synopsis: `pbs_log_path = path`

Description: The path of the PBS server logfiles which are used by A-REX to determine whether a PBS job is completed. If not specified, A-REX will use qstat for that.

Default: `/var/spool/pbs/server_logs`

Example:

```
pbs_log_path=/var/spool/pbs/server_logs
```

pbs_dedicated_node_string

[lrms]

Synopsis: `pbs_dedicated_node_string = string`

Description: The string which is used in the PBS node config to distinguish the grid nodes from the rest. Suppose only a subset of nodes are available for grid jobs, and these nodes have a common `node property` string, in this case the string should be set to this value and only the nodes with the corresponding `pbs node property` are counted as grid enabled nodes. Setting the `dedicated_node_string` to the value of the `pbs node property` of the grid-enabled nodes will influence how the `totalcpus`, `user freecpus` is calculated. You don't need to set this attribute if your cluster is fully available for the grid and your cluster's PBS config does not use the `node property` method to assign certain nodes to grid queues. You shouldn't use this config option unless you make sure your PBS config makes use of the above described setup.

Default: `undefined`

Example:

```
pbs_dedicated_node_string=gridnode
```

condor_bin_path

[lrms]

Synopsis: condor_bin_path = path

Description: Path to Condor binaries. Must be set if Condor is used.

Default: /usr/bin

Example:

```
condor_bin_path=/opt/condor/bin
```

condor_config

[lrms]

Synopsis: condor_config = path

Description: Full path to Condor config file. Must be set if Condor is used and the config file is not in its default location (/etc/condor/condor_config or ~/condor/condor_config). The full path to the file should be given.

Default: /etc/condor/condor_config

Example:

```
condor_config=/opt/condor/etc/condor_config
```

condor_rank

[lrms]

Synopsis: condor_rank = ClassAd_float_expression

Description: If you are not happy with the way Condor picks nodes when running jobs, you can define your own ranking algorithm by optionally setting the condor_rank attribute. condor_rank should be set to a ClassAd float expression that you could use in the Rank attribute in a Condor job description.

Default: undefined

Example:

```
condor_rank=(1-LoadAvg/2)*(1-LoadAvg/2)*Memory/1000*KFlops/1000000
```

condor_requirements

[lrms]

Synopsis: condor_requirements = string

Description: Specify additional constraints for Condor resources. The value of condor_requirements must be a valid constraints string which is recognized by a condor_status -constraint ... command. It can reference pre-defined ClassAd attributes (like Memory, Opsys, Arch, HasJava, etc) but also custom ClassAd attributes. To define a custom attribute on a condor node, just add two lines like the ones below in the \$(hostname).local config file on the node:

```
NORDUGRID_RESOURCE=TRUE
STARTD_EXPRS = NORDUGRID_RESOURCE, $(STARTD_EXPRS)
```

A job submitted to this resource is allowed to run on any node which satisfies the `condor_requirements` constraint. If `condor_requirements` is not set, jobs will be allowed to run on any of the nodes in the pool. When configuring multiple queues, you can differentiate them based on memory size or disk space, for example.

Default: undefined

Example:

```
condor_requirements=(OpSys == "linux" && NORDUGRID_RESOURCE && Memory >= 1000 &&
↪Memory < 2000)
```

slurm_bin_path

[lrms]

Synopsis: `slurm_bin_path = path`

Description: Path to SLURM binaries, must be set if installed outside of normal PATH

Default: `/usr/bin`

Example:

```
slurm_bin_path=/usr/bin
```

slurm_wakeupperiod

[lrms]

Synopsis: `slurm_wakeupperiod = numsec`

Description: How long should infosys wait before querying SLURM for new data (seconds)

Default: `30`

Example:

```
slurm_wakeupperiod=15
```

slurm_use_sacct

[lrms]

Synopsis: `slurm_use_sacct = yes/no`

Description: Indicates whether ARC should use `sacct` instead of `scontrol` to obtain information about finished jobs

Allowed values: `yes, no`

Default: `yes`

Example:

```
slurm_use_sacct=yes
```

slurm_requirements

[lrms]

Synopsis: `slurm_requirements = string`

Description: Use this option to specify extra SLURM-specific parameters.

Default: undefined

Example:

```
slurm_requirements=mincpus=8
```

slurm_query_retries

[lrms]

Synopsis: `slurm_query_retries = number`

Description: Number of sacct/scontrol retries performed in scan-SLURM-job. If slurm is overloaded the sacct/scontrol command call may fail. If retries > 1 sacct/scontrol is retried after some seconds for that(those) particular job(s). If all retry attempts fail, the next scan-SLURM-job initiation will pick up the job(s) from last time.

Default: 1

Example:

```
slurm_query_retries=3
```

[arex] block

The [arex] block, together with its various subblocks, configures the A-REX service hosted in arched. A-REX takes care of various middleware tasks on the frontend such as job creation and management, stagein/stageout, LRMS job submission, data caching, etc...

user

[arex]

Synopsis: `user = user[:group]`

Description: Switch to a non root user/group after startup. Use with caution because of limited functionality when arex is not run under root.

Default: root

Example:

```
user=grid:grid
```

norootpower

[arex]

Synopsis: norootpower = yes|no

Description: If set to yes, all job management processes will switch to mapped user's identity while accessing session directory. This is useful if session directory is on NFS with root squashing turned on.

Allowed values: yes, no

Default: no

Example:

```
norootpower=yes
```

delegationdb

[arex]

Synopsis: delegationdb = db_name

Description: specify which DB to use to store delegations. Currently supported db_names are bdb and sqlite

Default: sqlite

Example:

```
delegationdb=sqlite
```

watchdog

[arex]

Synopsis: watchdog = yes/no

Description: Specifies if additional watchdog processes is spawned to restart main process if it is stuck or dies.

Allowed values: yes, no

Default: no

Example:

```
watchdog=no
```

loglevel

[arex]

Synopsis: loglevel = level

Description: Set loglevel of the arched daemon hosting A-REX service between 0 (FATAL) and 5 (DEBUG). Defaults to 3 (INFO).

Allowed values: 0, 1, 2, 3, 4, 5, FATAL, ERROR, WARNING, INFO, VERBOSE, DEBUG

Default: 3

Example:

```
loglevel=3
```

logfile

[arex]

Synopsis: logfile = path

Description: Specify A-REX log file location. If using an external log rotation tool be careful to make sure it matches the path specified here.

Default: /var/log/arc/arex.log

Example:

```
logfile=/var/log/arc/arex.log
```

joblog

[arex]

Synopsis: joblog = path

Description: Specifies where to store specialized log about started and finished jobs. If path is empty log is NOT written. Controlled by logrotate if default name is kept. This log is not used by any other part of ARC so can be safely disabled if you are not interested in storing jobs log.

Default: /var/log/arc/arex-jobs.log

Example:

```
joblog=
```

fixdirectories

[arex]

Synopsis: fixdirectories = yes/missing/no

Description: Specifies during startup A-REX should create all directories needed for its operation and set suitable default permissions. If no is specified then A-REX does nothing to prepare its operational environment. In case of missing A-REX only creates and sets permissions for directories which are not present yet. For yes all directories are created and permissions for all used directories are set to default safe values.

Allowed values: yes, missing, no

Default: yes

Example:

```
fixdirectories=yes
```

controldir

[arex]

Synopsis: controldir = path

Description: The directory of the A-REX's internal job metadata files. For a heavily loaded computing element you can consider to locate controldir on a dedicated partition optimized for small random reads and writes and for storing many small files. The directory is not needed on the nodes.

Default: /var/spool/arc/jobstatus

Example:

```
controldir=/var/spool/arc/jobstatus
```

sessiondir

[arex]

Synopsis: sessiondir = path [drain]

Description: the directory which holds the sessiondirs of the grid jobs. Multiple session directories may be specified. In this case jobs are spread evenly over the session directories. If sessiondir=* is set, the session directory will be spread over the \${HOME}/.jobs directories of every locally mapped unix user. It is preferred to use common session directories. The path may be followed by drain, in which case no new jobs will be assigned to that sessiondir, but current jobs will still be processed and accessible.

This option is **multivalued**.

Default: /var/spool/arc/sessiondir

Example:

```
sessiondir=/scratch/arcsessions drain
sessiondir=*
```

defaultttl

[arex]

Synopsis: defaultttl = [ttl [ttr]]

Description: The ttl parameter sets the time in seconds for how long a job session directory will survive after job execution has finished. If not specified the default is 1 week. The ttr parameter sets how long information about a job will be kept after the session directory is deleted. If not specified, the ttr default is one month.

Default: 604800 2592000

Example:

```
defaultttl=2592000
```

shared_filesystem

[arex]

Synopsis: shared_filesystem = yes/no

Description: Specifies if computing nodes can access folders mounted with protocols like NFS with the same pathnames as the frontend.

Note

the default 'yes' assumes that the paths to the session directories are the same on both frontend and nodes. If these paths are not the same, then one should set the scratchdir option.

The option changes the RUNTIME_NODE_SEES_FRONTEND variable in the submission scripts.

Allowed values: yes, no

Default: yes

Example:

```
shared_filesystem=yes
```

scratchdir

[arex]

Synopsis: scratchdir = path

Description: The path on computing node to move session directory to before execution. If defined should contain the path to the directory on the computing node which can be used to store a jobs' files during execution. Sets the environment variable `RUNTIME_LOCAL_SCRATCH_DIR`. If the variable is not set, then the session dir is not moved before execution. Don't set this parameter unless you want to move the sessiondir to scratchdir on the node.

Default: undefined

Example:

```
scratchdir=/local/scratch/
```

shared_scratch

[arex]

Synopsis: shared_scratch = path

Description: The path on frontend where scratchdir can be found. If defined should contain the path corresponding to that set in scratchdir as seen on the frontend machine. Sets the environment variable `RUNTIME_FRONTEND_SEES_NODE`.

Default: undefined

Example:

```
shared_scratch=/mnt/scratch
```

tmpdir

[arex]

Synopsis: tmpdir = path

Description: A temporary directory used by A-REX.

Default: /tmp

Example:

```
tmpdir=/tmp
```

runtimedir

[arex]

Synopsis: runtimedir = path

Description: The directory which holds the additional runtimeenvironment scripts, added by system administrator. Several directories can be specified. To enable RTEs to be advertised in the information system and used during submission the arcctl tool should be used.

This option is **multivalued**.

Default: undefined

Example:

```
runtimedir=/var/spool/arc/extraruntimes
runtimedir=/cvmfs/vo/arcruntime
```

maxjobs

[arex]

Synopsis: maxjobs = number1 number2 number3 number4 number5

Description: specifies maximum allowed number of jobs:

- number1 - jobs which are not in FINISHED state (jobs tracked in RAM)
- number2 - jobs being run (SUBMITTING, INLRMS states)
- number3 - jobs being processed per DN
- number4 - jobs in whole system
- number5 - LRMS scripts limit (jobs in SUBMITTING and CANCELING)

A parameter set to -1 means no limit.

Default: -1 -1 -1 -1 -1

Example:

```
maxjobs=10000 10 2000 -1 -1
```

maxrerun

[arex]

Synopsis: maxrerun = number

Description: Specifies how many times job can be rerun if it failed in LRMS. This is only an upper limit, the actual rerun value is set by the user in his xrsl.

Default: 5

Example:

```
maxrerun=5
```

statecallout

[arex]

Synopsis: statecallout = state options plugin_path [plugin_arguments]

Description: Enables a callout feature of A-REX: every time job goes to state A-REX will run plugin_path executable. The following states are allowed: ACCEPTED, PREPARING, SUBMIT, INLRMS, FINISHING, FINISHED and DELETED. Options consist of key=value pairs separated by comma. Possible keys are:

timeout

defines the timeout in seconds to wait for plugin execution (timeout= can be omitted).

onsuccess, onfailure, ontimeout

defines the action that A-REX should take on successful execution (exit code 0), failed execution (exit code is not 0) or execution timeout respectively.

Possible actions are:

- pass - continue executing job,
- fail - cancel job,
- log - write to log about the failure and continue executing job.

It is possible to use following substitutions to construct plugin command line:

- %R - session root (value of sessiondir in [arex] block)
- %C - controldir path
- %U - username of mapped UNIX account
- %u - numeric UID of mapped UNIX account
- %g - numeric GID of mapped UNIX account
- %H - home directory of mapped UNIX account as specified in /etc/passwd
- %Q - default queue (see lrms configuration option in [lrms] block)
- %L - LRMS name (see lrms configuration option in [lrms] block)
- %W - ARC installation path (corresponds to the ARC_LOCATION environmental variable)
- %F - path to configuration file for this instance
- %I - job ID (substituted in runtime)
- %S - job state (substituted in runtime)

Plugins included into ARC distribution:

- arc-blahp-logger - write accounting log for every finished job in BLAH format

This option is **multivalued**.

Default: undefined

Example:

```
statecallout=FINISHED timeout=10,onfailure=pass /usr/libexec/arc/arc-blahp-logger -I
→%I -U %u -L %C/job.%I.local -P %C/job.%I.proxy
```

wakeupperiod

[arex]

Synopsis: wakeupperiod = time

Description: Specifies how often A-REX checks for new jobs arrived, job state change requests, etc. That is responsiveness of A-REX. time is time period in seconds. Default is 3 minutes. Usually no need to change this parameter because important state changes are also triggering out-of-schedule checks.

Note

this parameter does not affect responsiveness of backend scripts - especially scan-<LRMS>-job. That means that upper estimation of time for

detecting job finished executing is sum of responsiveness of backend script + wakeupperiod.

Default: 180

Example:

```
wakeupperiod=180
```

infoproviders_timelimit

[arex]

Synopsis: infoproviders_timelimit = seconds

Description: Sets the execution time limit of the infoprovider scripts started by the A-REX. Infoprovider scripts running longer than the specified timelimit are gracefully handled by the A-REX (the behaviour depends on the state of the system) Increase this value if you have many jobs in the controldir and infoproviders need more time to process.

Default: 10800

Example:

```
infoproviders_timelimit=10800
```

pidfile

[arex]

Synopsis: pidfile = path

Description: Specify location of file containing PID of daemon process.

Default: /run/arched-arex.pid

Example:

```
pidfile=/run/arched-arex.pid
```

mail

[arex]

Synopsis: mail = email_address

Description: Specifies the email address from where the notification mails are sent

Default: \$VAR{user}@\$VAR{[common]hostname}

Example:

```
mail=grid.support@somewhere.org
```

helper

[arex]

Synopsis: helper = user executable arguments

Description: By enabling this parameter A-REX will run an external helper program under the user useraccount. The program will be kept running, every time the executable finishes it will be started again. As a limitation, currently only '.' is supported as username, which corresponds to the user running A-REX.

Default: undefined

Example:

```
helper=. /usr/local/bin/myutility
```

helperlog

[arex]

Synopsis: helperlog = path

Description: Configuration option to specify the location of log for helpers.

Default: /var/log/arc/job.helper.errors

Example:

```
helperlog=/var/log/arc/job.helper.errors
```

forcedefaultvoms

[arex]

Synopsis: forcedefaultvoms = VOMS_FQAN

Description: specify VOMS FQAN which user will be assigned if his/her credentials contain no VOMS attributes. To assign different values to different queues put this command into [queue] block.

Default: undefined

Example:

```
forcedefaultvoms=/vo/group/subgroup
```

usetokenforvoms

[arex]

Synopsis: usetokenforvoms = yes/no

Description: Whether claims from WLCG complaint tokens to be used as VOMS attributes.

Allowed values: yes, no

Default: no

Example:

```
usetokenforvoms=yes
```

Warning

CHANGE: NEW in 7.0.0

tokenscopes

[arex]

Synopsis: tokenscopes = action=scope[,action=scope[...]]

Description: assigns JWT token scopes required to perform specific actions. Multiple tokenscopes entries are allowed. Following actions are supported:

- info - information about server
- jobinfo - information about jobs
- jobcreate - create new job or restart existing
- jobcancel - cancel active jobs
- jobdelete - remove jobs from server
- datainfo - information about files in session directory
- datawrite - create new or modify files in session directory
- dataread - read files in session directory

The action=scope pairs can be replaced with identifier which works as shortcut for multiple actions and scopes. Only currently supported shortcut identifier is wlcg (see below)

Default: undefined

Following example assigns scopes according to WLCG profile and alternatively can be defined by tokenscopes=wlcg.

Example:

```
tokenscopes=jobinfo=compute.read,jobcreate=compute.create,jobcancel=compute.cancel,
↪ jobdelete=compute.cancel
tokenscopes=datainfo=compute.read,datawrite=compute.modify,dataread=compute.read
```

Warning

CHANGE: NEW in 7.0.0

authtokenmap

[arex]

Synopsis: authtokenmap = claim:attrname[,...]

Description: map your token claim key to a custom attribute name to be used in accounting

This option in **multivalued**.

Default: name:name,email:email,preferred_username:username,wlcg_groups:group

Example:

```
authtokenmap=sub:user
```

Warning

CHANGE: NEW in 7.0.0

[arex/cache] block

This subblock enables and configures the cache functionality of A-REX. A-REX can cache input files downloaded as part of the stage-in process of grid jobs so that subsequent jobs requiring the same file don't have to download it again. The cached file will be symlinked (or copied) into the session directory of the job. To disable to cache functionality simply comment out the [arex/cache] config block. It is a good idea to have the cache on its own separate file system that is shared with the nodes. For more information about the cache functionality of A-REX consult the Data Cache technical description in the online documentation.

cachedir

[arex/cache]

Synopsis: *cachedir = cache_path [link_path]

Description: Specifies a directory to store cached data. Multiple cache directories may be specified. Cached data will be distributed evenly over the caches. Optional link_path specifies the path at which the cache_path is accessible on computing nodes, if it is different from the path on the A-REX host. If link_path is set to . files are not soft-linked, but copied to session directory. If a cache directory needs to be drained, then link_path should specify drain, in which case no new files will be added to the cache and files in the cache will no longer be used. Setting link_path to readonly ensures that no new files are written to this cache, but existing files can still be used. Draining and read-only caches are not cleaned by the A-REX cache cleaner. A restart of A-REX is required when changing cache options.

This option in **multivalued**.

Default: undefined

Example:

```
cachedir=/scratch/cache
cachedir=/shared/cache /frontend/jobcache
cachedir=/fs1/cache drain
```

[arex/cache/cleaner] block

This subblock enables the cleaning functionality of the cache. If this block is not enabled then the cache will not be cleaned by A-REX. Either `cachesize` or `cachelifetime` should also be set to enable cleaning.

logfile

[arex/cache/cleaner]

Synopsis: `logfile = path`

Description: sets the filename where output of the cache-clean tool should be logged. Defaults to `/var/log/arc/cache-clean.log`.

Default: `/var/log/arc/cache-cleaner.log`

Example:

```
logfile=/tmp/cache-clean.log
```

loglevel

[arex/cache/cleaner]

Synopsis: `loglevel = level`

Description: specifies the level of logging by the cache-clean tool, between 0 (FATAL) and 5 (DEBUG). Defaults to 3 (INFO).

Allowed values: 0, 1, 2, 3, 4, 5, FATAL, ERROR, WARNING, INFO, VERBOSE, DEBUG

Default: 3

Example:

```
loglevel=4
```

cachesize

[arex/cache/cleaner]

Synopsis: `cachesize = max min`

Description: Specifies high and low watermarks for space used by cache, as a percentage of the space on the file system on which the cache directory is located. When the max is exceeded, files will be deleted to bring the used space down to the min level. It is a good idea to have the cache on its own separate file system.

Default: 100 100

Example:

```
cachesize=50 20
```

calculatesize

[arex/cache/cleaner]

Synopsis: calculatesize = filesystem/cachedir

Description: specifies the way the space occupied by the cache will be calculated. If set to cachedir then cache-clean calculates the size of the cache instead of using filesystem used space.

Allowed values: filesystem, cachedir

Default: filesystem

Example:

```
calculatesize=cachedir
```

cachelifetime

[arex/cache/cleaner]

Synopsis: cachelifetime = time

Description: Turns on time-based file cleaning. Files accessed less recently than the given time period will be deleted. Example values of this option are 1800, 90s, 24h, 30d. When no suffix is given the unit is seconds.

Default: undefined

Example:

```
cachelifetime=30d
```

cachespacetool

[arex/cache/cleaner]

Synopsis: cachespacetool = path [options]

Description: specifies an alternative tool to df that cache-clean should use to obtain space information on the cache file system. The output of this command must be total_bytes used_bytes. The cache directory is passed as the last argument to this command.

Default: undefined

Example:

```
cachespacetool=/etc/getspace.sh
```

cacheleantimeout

[arex/cache/cleaner]

Synopsis: cacheleantimeout = time

Description: the timeout in seconds for running the cache-clean tool. If using a large cache or slow file system this value can be increased to allow the cleaning to complete. Defaults to 3600 (1 hour).

Default: 3600

Example:

```
cacheleantimeout=10000
```

[arex/data-staging] block

This subblock enables and configures the data staging capabilities of A-REX. A subsystem called DTR (Data Transfer Reloaded) is responsible for collecting input data for a job before submission to the LRMS, and for staging out data after the job has finished. Automagic data staging of A-REX is a very powerful feature, disabling this functionality (by commenting out the subblock) is not recommended.

loglevel

[arex/data-staging]

Synopsis: `loglevel = number`

Description: Sets the log level for transfer logging in job.id.errors files, between 0 (FATAL) and 5 (DEBUG). Default is to use value set by loglevel option in [arex] section.

Allowed values: 0, 1, 2, 3, 4, 5, FATAL, ERROR, WARNING, INFO, VERBOSE, DEBUG

Default: `$VAR{[arex]loglevel}`

Example:

```
loglevel=4
```

logfile

[arex/data-staging]

Synopsis: `logfile = path`

Description: A central file in which all data staging messages from every job will be collected and logged in addition to their job.id.errors files. If this option is not present or the path is empty the log file is not created. This file is not automatically controlled by logrotate unless you name it as /var/log/arc/datastaging.log.

Default: undefined

Example:

```
logfile=/var/log/arc/datastaging.log
```

statefile

[arex/data-staging]

Synopsis: `statefile = path`

Description: A file in which data staging state information (for monitoring and recovery purposes) is periodically dumped.

Default: `$VAR{[arex]controldir}/dtr.state`

Example:

```
statefile=/tmp/dtr.state
```

usehostcert

[arex/data-staging]

Synopsis: usehostcert = yes/no

Description: Whether the A-REX host certificate should be used for communication with remote hosts instead of the users' proxies.

Allowed values: yes, no

Default: no

Example:

```
usehostcert=yes
```

maxtransfertries

[arex/data-staging]

Synopsis: maxtransfertries = number

Description: the maximum number of times download and upload will be attempted per job (retries are only performed if an error is judged to be temporary)

Default: 10

Example:

```
maxtransfertries=20
```

passivetransfer

[arex/data-staging]

Synopsis: passivetransfer = yes/no

Description: If yes, gridftp transfers are passive. Setting this option to yes can solve transfer problems caused by firewalls.

Allowed values: yes, no

Default: yes

Example:

```
passivetransfer=yes
```

globus_tcp_port_range

[arex/data-staging]

Synopsis: globus_tcp_port_range = port_range

Description: In a firewalled environment the software which uses GSI needs to know what ports are available. This parameter is only needed if passivetransfer=no was set. These variable are similar to the Globus environment variables GLOBUS_TCP_PORT_RANGE and GLOBUS_UDP_PORT_RANGE.

Default: 9000,9300

Example:

```
globus_tcp_port_range=9000,12000
```

globus_udp_port_range

[arex/data-staging]

Synopsis: globus_udp_port_range = port_range

Description: In a firewalled environment the software which uses GSI needs to know what ports are available. This parameter is only needed if `passivetransfer=no` was set. These variable are similar to the Globus enviroment variables `GLOBUS_TCP_PORT_RANGE` and `GLOBUS_UDP_PORT_RANGE`.

Default: 9000,9300

Example:

```
globus_udp_port_range=9000,12000
```

httpgetpartial

[arex/data-staging]

Synopsis: httpgetpartial = yes/no

Description: If yes, HTTP GET transfers may transfer data in chunks/parts. If no - data is always transfered in one piece.

Allowed values: yes, no

Default: no

Example:

```
httpgetpartial=no
```

speedcontrol

[arex/data-staging]

Synopsis: speedcontrol = min_speed min_time min_average_speed max_inactivity

Description: specifies how slow data transfer must be to trigger error. Transfer is cancelled if speed is below min_speed bytes per second for at least min_time seconds, or if average rate is below min_average_speed bytes per second, or no data was transferred for longer than max_inactivity seconds. Value of zero turns feature off.

Default: 0 300 0 300

Example:

```
speedcontrol=0 300 100 300
speedcontrol=
```

maxdelivery

[arex/data-staging]

Synopsis: maxdelivery = number

Description: Maximum number of concurrent file transfers, i.e. active transfers using network bandwidth. This is the total number for the whole system including any remote staging hosts.

Default: 10

Example:

```
maxdelivery=40
```

maxprocessor

[arex/data-staging]

Synopsis: maxprocessor = number

Description: Maximum number of concurrent files in each of the DTR internal pre- and post-processing states, eg cache check or replica resolution.

Default: 10

Example:

```
maxprocessor=20
```

maxemergency

[arex/data-staging]

Synopsis: maxemergency = number

Description: Maximum emergency slots which can be assigned to transfer shares when all slots up to the limits configured by the above two options are used by other shares. This ensures shares cannot be blocked by others.

Default: 1

Example:

```
maxemergency=5
```

maxprepared

[arex/data-staging]

Synopsis: maxprepared = number

Description: Maximum number of files in a prepared state, i.e. pinned on a remote storage such as SRM for transfer. A good value is a small multiple of maxdelivery.

Default: 200

Example:

```
maxprepared=250
```

sharepolicy

[arex/data-staging]

Synopsis: sharepolicy = grouping

Description: Defines the mechanism to be used for the grouping of the job transfers. DTR assigns the transfers to shares, so that those shares can be assigned to different priorities. Possible values for grouping are dn, voms:vo, voms:role and voms:group:

dn

each job is assigned to a share based on the DN of the user submitting the job.

voms:vo

each job is assigned to a share based on the VO specified in the proxy.

voms:role

each job is assigned to a share based on the role specified in the first attribute found in the proxy.

voms:group

each job is assigned to a share based on the group specified in the first attribute found in the proxy.

In case of the voms schemes, if the proxy is not a VOMS proxy, then a default share is used. If sharepolicy is not set then the client-defined priority is applied.

Default: undefined

Example:

```
sharepolicy=voms:role
```

sharepriority

[arex/data-staging]

Synopsis: sharepriority = share priority

Description: Defines a share with a fixed priority, different from the default (50). Priority is an integer between 1 (lowest) and 100 (highest).

This option is **multivalued**.

Default: undefined

Example:

```
sharepriority=myvo:students 20
sharepriority=myvo:production 80
```

copyurl

[arex/data-staging]

Synopsis: copyurl = url_head local_path

Description: Configures a mapping of URLs to locally- accessible paths. If a URL starts with url_head, the local_path will be substituted for the actual transfer. Applies to both input and output files.

Note

local_path can also be of URL type.

This option is **multivalued**.

Default: undefined

Example:

```
copyurl=gsiftp://example.org:2811/data/ /data/
copyurl=gsiftp://example2.org:2811/data/ /data/
```

Warning

CHANGE: MODIFIED in 7.0.0 - applies also to output files

linkurl

[arex/data-staging]

Synopsis: linkurl = url_head local_path [node_path]

Description: Identical to copyurl, configures DTR so that for certain URLs files won't be downloaded or copied (in case of copyurl), but soft-link will be created. The local_path specifies the way to access the file from the frontend, and is used to check permissions. The node_path specifies how the file can be accessed from computing nodes, and will be used for soft-link creation. If node_path is missing - local_path will be used. This option applies only to input files.

This option is **multivalued**.

Default: undefined

Example:

```
linkurl=gsiftp://somewhere.org/data /data
linkurl=gsiftp://example.org:2811/data/ /scratch/data/
```

preferredpattern

[arex/data-staging]

Synopsis: preferredpattern = pattern

Description: specifies a preferred pattern on which to sort multiple replicas of an input file. It consists of one or more patterns separated by a pipe character (|) listed in order of preference. Replicas will be ordered by the earliest match. If the dollar character (\$) is used at the end of a pattern, the pattern will be matched to the end of the hostname of the replica. If an exclamation mark (!) is used at the beginning of a pattern, any replicas matching the pattern will be excluded from the sorted replicas.

Default: undefined

Example:

```
preferredpattern=srm://myhost.ac.uk|.uk$|ndgf.org$|badhost.org$
```

The following options are used to configure multi-host data staging deployment scenario. In that setup a couple of additional data staging boxes are enabled to off-load data transfers.

deliveryservice

[arex/data-staging]

Synopsis: deliveryservice = URL

Description: The URL to a remote data delivery service which can perform remote data staging.

This option is **multivalued**.

Default: undefined

Example:

```
deliveryservice=https://myhost.org:443/datadeliveryservice
```

localdelivery

[arex/data-staging]

Synopsis: localdelivery = yes/no

Description: If any deliveryservice is defined, this option determines whether local data transfer is also performed.

Allowed values: yes, no

Default: no

Example:

```
localdelivery=yes
```

remotesizelimit

[arex/data-staging]

Synopsis: remotesizelimit = size

Description: Lower limit on file size (in bytes) of files that remote hosts should transfer. Can be used to increase performance by transferring small files using local processes.

Default: undefined

Example:

```
remotesizelimit=100000
```

[arex/ws] block

A-REX exposes a set of Web Service interfaces that can be used to create and manage jobs, obtain information about the CE and the jobs, handle delegations, access cache information, so on. Comment out this block if you don't want to provide WS-interfaces for various A-REX functionalities.

wsurl

[arex/ws]

Synopsis: wsurl = url

Description: Specifies the base URL under which the web service interfaces will be available. The URL argument must be a full URL consisting of protocol+host+port+path: e.g. https://<hostname>:<port>/<path> Make sure the chosen port is not blocked by firewall or other security rules.

Default: https://\$VAR{[common]hostname}:443/arex

Example:

```
wsurl=https://piff.hep.lu.se:443/arex
```

logfile

[arex/ws]

Synopsis: logfile = path

Description: Specify log file location for WS-interface operations.

Default: /var/log/arc/ws-interface.log

Example:

```
logfile=/var/log/arc/ws-interface.log
```

pidfile

[arex/ws]

Synopsis: pidfile = path

Description: Specify location of file containing PID of daemon process.

Default: /run/arched-arex-ws.pid

Example:

```
pidfile=/run/arched-arex-ws.pid
```

max_job_control_requests

[arex/ws]

Synopsis: max_job_control_requests = number

Description: The max number of simultaneously processed job management requests over WS interface - like job submission, cancel, status check etc.

Default: 100

Example:

```
max_job_control_requests=100
```

max_infosys_requests

[arex/ws]

Synopsis: max_infosys_requests = number

Description: The max number of simultaneously processed info requests over WS interface.

Default: 1

Example:

```
max_infosys_requests=1
```

max_data_transfer_requests

[arex/ws]

Synopsis: max_data_transfer_requests = number

Description: The max number of simultaneously processed data transfer requests over WS interface - like data staging.

Default: 100

Example:

```
max_data_transfer_requests=100
```

tlsciphers

[arex/ws]

Synopsis: tlsciphers = ciphers_list

Description: Override OpenSSL ciphers list enabled on server

Default: HIGH:!eNULL:!aNULL

Example:

```
tlsciphers=HIGH:!eNULL:!aNULL
```

tlsserverorder

[arex/ws]

Synopsis: tlsserverorder = yes

Description: Force priority order of ciphers for TLS connection to be decided on server sid

Default: no

Example:

```
tlsserverorder=yes
```

Warning

CHANGE: NEW in 7.0.0

tlsprotocols

[arex/ws]

Synopsis: `tlsprotocols = SSL/TLS protocols`

Description: Specify which protocols to enable This is space separated list of values - SSLv2 SSLv3 TLSv1.0 TLSv1.1 TLSv1.2 TLSv1.3

Default: TLSv1.2 TLSv1.3

Example:

```
tlsprotocols=TLSv1.2 TLSv1.3
```

tlscurve

[arex/ws]

Synopsis: `tlscurve = curve`

Description: Specify SSL/TLS ECDH curve name (SN)

Default: secp521r1

Example:

```
tlscurve=secp521r1
```

[arex/ws/jobs] block

This block enables the job management, info query, delegation protocols through REST interface. Read <http://www.nordugrid.org/arc/arc7/tech/rest/rest.html> for the REST interface specification.

allownew

[arex/ws/jobs]

Synopsis: `allownew = yes/no`

Description: The ‘allownew’ config parameter sets if the Computing Element accepts submission of new jobs via the WS-interface. This parameter can be used to close down the CE.

Allowed values: yes, no

Default: yes

Example:

```
allownew=yes
```

allownew_override

[arex/ws/jobs]

Synopsis: allownew_override = [authgroup ...]

Description: Defines which authorization groups are allowed to submit new jobs via the WS-interfaces when the CE is closed with allownew=no

Note

it requires the allownew=no to be set.

This option in **multivalued**.

Default: undefined

Example:

```
allownew_override=biouers atlasusers
allownew_override=yourauthgroup
```

allowaccess

[arex/ws/jobs]

Synopsis: allowaccess = authgroup

Description: Defines that the specified authgroup members are authorized to access the ARC-CE via this interface. A related config option the denyaccess (see below) can be used to reject access. Multiple allowaccess and denyaccess authorization statements are allowed within a configuration block. These statements are processed sequentially in the order they are specified in the config block. The processing stops on first allowaccess or denyaccess statement matching the authgroup membership. If there are no authorization statements specified, then no additional restrictions are applied for authorizing user access and the interface is open to everybody authenticated.

Default: undefined

This option in **multivalued**.

Example:

```
allowaccess=biouers
allowaccess=atlasusers
```

denyaccess

[arex/ws/jobs]

Synopsis: denyaccess = authgroup

Description: Defines that the specified authgroup members are REJECTED, not authorized to access the ARC-CE via this interface.

Note

a related config option the allowaccess (see above) can be used to grant access.

Multiple `denyaccess` and `allowaccess` authorization statements are allowed within a configuration block. These statements are processed sequentially in the order they are specified in the config block. The processing stops on first `allowaccess` or `denyaccess` statement matching the authgroup membership. If there are no authorization statements specified, then no additional restrictions are applied for authorizing user access and the interface is open to everybody authenticated.

Default: undefined

This option in **multivalued**.

Example:

```
denyaccess=blacklisted-users
```

maxjobdesc

[arex/ws/jobs]

Synopsis: maxjobdesc = size

Description: specifies maximal allowed size of job description in bytes. Default value is 5MB. Use 0 to set unlimited size.

Default: 5242880

Example:

```
maxjobdesc=0
```

[arex/ws/cache] block

The content of the A-REX cache can be accessed via a WS-interface. Configuring this block will allow reading cache files through a special URL. For example, if the remote file `gsiftp://remotehost/file1` is stored in the cache and the WS interfaces (configured above) are available via `wsurl` of `https://hostname:443/arex/`, then the cached copy of the file can be access via the following special URL: `https://hostname:443/arex/cache/gsiftp://remotehost/file1`. Comment out this block if you don't want to expose the cache content via WS-interface.

cacheaccess

[arex/ws/cache]

Synopsis: cacheaccess = rule

Description: This parameter defines the access control rules for the cache wsinterface, the rules for allowing access to files in the cache remotely through the A-REX web interface. If not set, then noone can access anything. The default is not set that means complete denial. A rule has three parts:

1. Regular expression defining a URL pattern
2. Credential attribute to match against a client's credential
3. Regular expression defining a credential value to match against a client's credential

A client is allowed to access the cached file if a URL pattern matches the cached file URL and the client's credential has the attribute and matches the value required for that pattern. Possible values for credential attribute are `dn`, `voms:vo`, `voms:role` and `voms:group`.

This option in **multivalued**.

Default: undefined

Example:

```
cacheaccess=gsiftp://host.org/private/data/. * voms:vo myvo:production
cacheaccess=gsiftp://host.org/private/data/bob/. * dn /O=Grid/O=NorduGrid/. *
```

[arex/ws/candypond] block

The CandyPond (Cache and deliver your pilot on-demand data) A-REX Web Service exposes various useful data-staging related operations for the pilot job model where input data for jobs is not known until the job is running on the worker node. This service is intended to be used by A-REX managed jobs. This service requires the [arex/data-staging] functionality. To use service from the job context enable EVN/CANDYPOND RTE.

The CandyPond service is available via the wsurl/candypond URL (e.g. <https://hostname:443/arex/candypond>)

[arex/jura] block

A-REX is responsible for collecting accounting measurements from various ARC subsystems, including batch system backends and DTR data staging.

A-REX writes all accounting data into the local accounting database that can be queried with `arcctl accounting`.

JURA is the accounting record generating and reporting ARC CE module. A-REX periodically executes JURA to create usage records based on the accounting target configuration and accounting database data.

Enable and configure this block if you want to send accounting records to accounting services.

Note

a dedicated `accounting` target subblock is needed for every accounting destination. The target subblocks are either of a type `apel` or `sgas`: `[arex/jura/apel:targetname]` or `[arex/jura/sgas:targetname]`.

logfile

[arex/jura]

Synopsis: logfile = path

Description: The name of the jura logfile.

Default: /var/log/arc/jura.log

Example:

```
logfile=/var/log/arc/jura.log
```

loglevel

[arex/jura]

Synopsis: loglevel = number

Description: Log level for the JURA accounting module.

Allowed values: 0, 1, 2, 3, 4, 5, FATAL, ERROR, WARNING, INFO, VERBOSE, DEBUG

Default: 3

Example:

```
loglevel=3
```

vomsless_vo

[arex/jura]

Synopsis: vomsless_vo = [authgroup]voname[#voissuer]

Description: This parameter allows the sysadmin to manually assign VOs during publishing to jobs that were submitted with VOMS-less grid proxies. voname is the VO name to be used in the generated records (the same as expected in voms-proxy) optional voissuer (relevant to SGAS only) value is the VOMS server identity (certificate DN). If authgroup is specified then this parameter is applied only if user belongs to specified authgroup.

Default: undefined

Example:

```
vomsless_vo=atlas
vomsless_vo=atlas#/DC=ch/DC=cern/OU=computers/CN=lcg-voms.cern.ch
vomsless_vo=atlasgroup atlas
```

vo_group

[arex/jura]

Synopsis: vo_group = group

Description: Adds an additional VO group attribute(s) to the usage records.

This option in **multivalued**.

Default: undefined

Example:

```
vo_group=/atlas/production
```

urdelivery_frequency

[arex/jura]

Synopsis: urdelivery_frequency = seconds

Description: Specifies the frequency of JURA process regular execution by the A-REX. The actual threshold of records reporting frequency can be defined on per-target basis.

Default: 3600

Example:

```
urdelivery_frequency=3600
```

x509_host_key

[arex/jura]

Synopsis: x509_host_key = path

Description: Optional parameter to overwrite [common] block values.

Default: \$VAR{[common]x509_host_key}

Example:

```
x509_host_key=/etc/grid-security/hostkey.pem
```

x509_host_cert

[arex/jura]

Synopsis: x509_host_cert = path

Description: Optional parameter to overwrite [common] block values.

Default: \$VAR{[common]x509_host_cert}

Example:

```
x509_host_cert=/etc/grid-security/hostcert.pem
```

x509_cert_dir

[arex/jura]

Synopsis: x509_cert_dir = path

Description: Optional parameter to overwrite [common] block values.

Default: \$VAR{[common]x509_cert_dir}

Example:

```
x509_cert_dir=/etc/grid-security/certificates
```

[arex/jura/sgas:targetname] block

An SGAS sub-block of [arex/jura] enables and configures an SGAS accounting server as a target destination to which JURA will send properly formatted usage records. You need to define a separate block with a unique targetname for every SGAS target server.

Note that the block name will be used by JURA to track that latest records sent to this targeted. Be aware that if you rename the block, target will be handled as a new one. However `targeturl` change will not trigger a new target handling.

targeturl

[arex/jura/sgas:targetname]

Synopsis: *targeturl = url

Description: The service endpoint URL of SGAS server.

Default: undefined

Example:

```
targeturl=https://grid.uio.no:8001/logger
```

localid_prefix

[arex/jura/sgas:targetname]

Synopsis: localid_prefix = prefix_string

Description: Sets a prefix value for the LocalJobID ur parameter for the SGAS usage records.

Default: undefined

Example:

```
localid_prefix=some_text_for_SGAS
```

vofilter

[arex/jura/sgas:targetname]

Synopsis: vofilter = vo

Description: Configures a job record filtering mechanism based on the VO attribute of the jobs. Only the matching job records, which was one of VO that you set here, will be sent to the target accounting service.

This option is **multivalued**.

Default: undefined

Example:

```
vofilter=atlas  
vofilter=fgi.csc.fi
```

urbatchsize

[arex/jura/sgas:targetname]

Synopsis: urbatchsize = number

Description: JURA sends usage records not one-by-one, but in batches. This options sets the size of a batch. Zero value means unlimited batch size.

Default: 50

Example:

```
urbatchsize=80
```

urdelivery_frequency

[arex/jura/sgas:targetname]

Synopsis: urdelivery_frequency = seconds

Description: Add optional minimal treshold of the interval between subsequent records publishing to this target.

Note

the actual delivery interval is the value divisible by urdelivery_frequency defined in [arex/jura] block that define the entire JURA process invocation frequency.

Default: undefined

Example:

```
urdelivery_frequency=3600
```

[arex/jura/apel:targetname] block

An APEL sub-block of [arex/jura] enables and configures an APEL accounting server as a target destination to which JURA will send properly formatted usage records. You need to define a separate block with a unique targetname for every APEL target server.

Note that the block name will be used by JURA to track that latest records sent to this targed. Be aware that if you rename the block, target will be handled as a new one. However targeturl change will not trigger a new target handling.

targeturl

[arex/jura/apel:targetname]

Synopsis: *targeturl = url

Description: The service endpoint URL of the APEL accounting server.

Default: undefined

Example:

```
targeturl=https://msg.argo.grnet.gr
```

topic

[arex/jura/apel:targetname]

Synopsis: topic = topic_name

Description: Sets the name of the APEL topic to which JURA will publish the accounting records. AMS destination topic for compute element is 'gLite-APEL'

Default: gLite-APEL

Example:

```
topic=/queue/global.accounting.test.cpu.central
```

project

[arex/jura/apel:targetname]

Synopsis: project = project_name

Description: Sets the name of the APEL project to use.

Default: accounting

Example:

```
project=accounting-nl
```

gocdb_name

[arex/jura/apel:targetname]

Synopsis: *gocdb_name = name

Description: Can be used to specify the GOCDB name of the resource. This value would be seen as Site attribute in the generated APEL records.

Default: undefined

Example:

```
gocdb_name=GRID_UIO_NO
```

apel_messages

[arex/jura/apel:targetname]

Synopsis: apel_messages = type

Description: Define what kind of records JURA will send to APEL services during regular publishing process. Possible cases are: per-job EMI CAR records (**urs**), APEL summary records (**summaries**) or APEL summary records v0.4 (**summaries-v04**) APEL Sync messages are always generated.

Allowed values: urs, summaries, summaries-v04

Default: summaries

Example:

```
apel_messages=urs
```

vofilter

[arex/jura/apel:targetname]

Synopsis: vofilter = vo

Description: Configures a job record filtering mechanism based on the VO attribute of the jobs. Only the matching job records, which was one of VO that you set here, will be sent to the target accounting service.

This option is **multivalued**.

Default: undefined

Example:

```
vofilter=atlas
vofilter=fgi.csc.fi
```

urbatchsize

[arex/jura/apel:targetname]

Synopsis: urbatchsize = number

Description: JURA sends usage records not one-by-one, but in batches. This options sets the size of a batch. Zero value means unlimited batch size. 500 is recommended to avoid too large messages using AMS

Default: 500

Example:

```
urbatchsize=500
```

urdelivery_frequency

[arex/jura/apel:targetname]

Synopsis: urdelivery_frequency = seconds

Description: Add optional minimal treshold of the interval between subsequent records publishing to this target.

Note

the actual delivery interval is the value divisible by urdelivery_frequency defined in [arex/jura] block that define the entire JURA process invocation frequency.

APEL recommended value is once per day for summaries. Use smaller values for urs.

Default: 86000

Example:

```
urdelivery_frequency=14000
```

[arex/ganglia] block

This block enables the monitoring of ARC-specific metrics. Earlier versions (ARC < 6.0) relied only on the standalone tool gangliarc, ganglia is now instead integrated into ARC, and gangliarc is obsolete.

Note

AREX ganglia (as gangliarc did) depends on an existing ganglia installation, as it sends its metrics to a running gmond process.

gmetric_bin_path

[arex/ganglia]

Synopsis: gmetric_bin_path = path

Description: The path to gmetric executable.

Default: /usr/bin/gmetric

Example:

```
gmetric_bin_path=/usr/local/bin/gmetric
```

metrics

[arex/ganglia]

Synopsis: metrics = name_of_the_metrics

Description: the metrics to be monitored. metrics takes a comma-separated list of one or more of the following metrics:

- staging – number of tasks in different data staging states - not yet implemented
- cache – free cache space
- session – free session directory space
- heartbeat – last modification time of A-REX heartbeat
- failedjobs – the number of failed jobs per last 100 finished
- jobstates – number of jobs in different A-REX stages
- all – all of the above metrics

Default: all

Allowed values: staging, cache, session, heartbeat, failedjobs, jobstates, all

Example:

```
metrics=all
```

frequency

[arex/ganglia]

Synopsis: frequency = seconds

Description: The period between each information gathering cycle, in seconds.

Default: 60

Example:

```
frequency=300
```

[infosys] block

This block enables and configures the core part of the information system. Enables the information collection to be used by other ARC components, including interfaces. Parameters in this block applies to all the infosys subsystems.

logfile

[infosys]

Synopsis: logfile = path

Description: Specifies log file location for the information provider scripts.

Default: /var/log/arc/infoprovider.log

Example:

```
logfile=/var/log/arc/infoprovider.log
```

loglevel

[infosys]

Synopsis: loglevel = number

Description: The loglevel for the infoprovider scripts (0-5). Each value corresponds to the following verbosity levels: FATAL => 0, ERROR => 1, WARNING => 2, INFO => 3, VERBOSE => 4, DEBUG => 5

Allowed values: 0, 1, 2, 3, 4, 5, FATAL, ERROR, WARNING, INFO, VERBOSE, DEBUG

Default: 3

Example:

```
loglevel=3
```

validity_ttl

[infosys]

Synopsis: validity_ttl = seconds

Description: The published infosys records advertise their validity e.g. how long the info should be considered up-to-date by the clients. Use this parameter to set the published validity value.

Note

different schemas may render this information differently.

Default: 10800

Example:

```
validity_ttl=10800
```

[infosys/ldap] block

This infosys subblock enables and configures the ldap hosting service for the infosys functionality. Using an LDAP server with some schema is one way to publish information about your Computing Element. Comment out this block if you don't want to run an LDAP-based information system.

hostname

[infosys/ldap]

Synopsis: hostname = FQDN

Description: the hostname of the machine running the slapd service will be the bind for slapd. If not present, will be taken from the [common]

Default: \$VAR{[common]hostname}

Example:

```
hostname=my.testbox
```

slapd_hostnamebind

[infosys/ldap]

Synopsis: slapd_hostnamebind = string

Description: May be used to set the hostname part of the network interface to which the slapd process will bind. Most of the cases no need to set since the hostname parameter is already sufficient. The example below will bind the slapd process to all the network interfaces available on the server.

Default: undefined

Example:

```
slapd_hostnamebind=*
```

port

[infosys/ldap]

Synopsis: port = port_number

Description: The port on which the slapd service runs. The default infosys port is assumed to be 2135 by many clients, therefore think twice before you change it because 3rd party clients assume 2135 to be the ldap infosys port.

Default: 2135

Example:

```
port=2135
```

user

[infosys/ldap]

Synopsis: `user = unix_user`

Description: overwrites the unix user running the slapd. By default the startup scripts search for well-known ldap-users like ldap or openldap than fall-back to root if not found.

Default: undefined

Example:

```
user=slapd
```

slapd

[infosys/ldap]

Synopsis: `slapd = path`

Description: explicitly define the path to slapd command. By default the startup scripts search for slapd binary in the system PATH.

Default: undefined

Example:

```
slapd=/usr/sbin/slapd
```

slapd_loglevel

[infosys/ldap]

Synopsis: `slapd_loglevel = number`

Description: Sets the native slapd loglevel (see man slapd). Slapd logs via syslog. The default is set to no-logging (0) and it is RECOMMENDED not to be changed in a production environment. Non-zero slap_loglevel value causes serious performance decrease.

Default: 0

Example:

```
slapd_loglevel=0
```

threads

[infosys/ldap]

Synopsis: `threads = number`

Description: The native slapd threads parameter, default is 32.

Default: 32

Example:

```
threads=128
```

timelimit

[infosys/ldap]

Synopsis: `timelimit = seconds`

Description: The native slapd `timelimit` parameter. Maximum number of seconds the slapd server will spend answering a search request. Default is 3600. You probably want a much lower value.

Default: 3600

Example:

```
timelimit=1800
```

idletimeout

[infosys/ldap]

Synopsis: `idletimeout = seconds`

Description: The native slapd `idletimeout` parameter. Maximum number of seconds the slapd server will wait before forcibly closing idle client connections. Its value must be larger than the value of `timelimit` option. If not set, it defaults to `timelimit + 1`.

Default: `$EVAL{$VAR{timelimit} + 1}`

Example:

```
idletimeout=1801
```

infosys_ldap_run_dir

[infosys/ldap]

Synopsis: `infosys_ldap_run_dir = path`

Description: The location where Nordugrid/GLUE2 LDAP ldif file will be generated, and where the fifo to sync between infoproviders and BDII will be generated.

Default: `/run/arc/infosys`

Example:

```
infosys_ldap_run_dir=/run/arc/infosys
```

ldap_schema_dir

[infosys/ldap]

Synopsis: `ldap_schema_dir = path`

Description: Allows to explicitly specify an additional path to the schema files. Note that this doesn't override standard location, but adds the specified path to the standard locations `/etc/ldap` and `/etc/openldap`. Normally it is sufficient to use only standard schema file locations, therefore not to set this parameter.

Default: undefined

Example:

```
ldap_schema_dir=/nfs/ldap/schema/
```

Note

the following options configure the third-party bdii ldap parameters. In 99% of cases no need to change anything and use the defaults. These variables are usually automatically set by ARC, and are here mostly for debug purposes and to tweak exotic BDII installations.

bdii_debug_level

[infosys/ldap]

Synopsis: `bdii_debug_level = level`

Description: Set this parameter to DEBUG to check bdii errors in bdii-update.log At the same time don't enable slapd logs this way reducing performance issues.

Default: WARNING

Example:

```
bdii_debug_level=ERROR
```

bdii_provider_timeout

[infosys/ldap]

Synopsis: `bdii_provider_timeout = seconds`

Description: This variable allows a system administrator to modify the behaviour of bdii-update. This is the time BDII waits for the bdii provider scripts generated by A-REX infosys to produce their output.

Default: 10800

Example:

```
bdii_provider_timeout=10800
```

Note

BDII5 uses these variables. These might change depending on BDII version. ARC sets them by inspecting distributed bdii configuration files. DO NOT change unless YOU KNOW WHAT YOU'RE DOING

bdii_location

[infosys/ldap]

Synopsis: `bdii_location = path`

Description: The installation directory for the BDII.

Default: /usr

Example:

```
bdii_location=/usr
```

bdii_run_dir

[infosys/ldap]

Synopsis: bdii_run_dir = path

Description: Contains BDII pid files and slapd pid files

Default: /run/arc/bdii

Example:

```
bdii_run_dir=/run/arc/bdii
```

bdii_log_dir

[infosys/ldap]

Synopsis: bdii_log_dir = path

Description: Contains infosys logs

Default: /var/log/arc/bdii

Example:

```
bdii_log_dir=/var/log/arc/bdii
```

bdii_tmp_dir

[infosys/ldap]

Synopsis: bdii_tmp_dir = path

Description: Contains provider scripts

Default: /var/tmp/arc/bdii

Example:

```
bdii_tmp_dir=/var/tmp/arc/bdii
```

bdii_var_dir

[infosys/ldap]

Synopsis: bdii_var_dir = path

Description: Contains slapd databases

Default: /var/lib/arc/bdii

Example:

```
bdii_var_dir=/var/lib/arc/bdii
```

bdii_update_pid_file

[infosys/ldap]

Synopsis: `bdii_update_pid_file = path`

Description: Allows to change bdii-update pidfiles filename and location

Default: `$VAR{bdii_run_dir}/bdii-update.pid`

Example:

```
bdii_update_pid_file=/run/arc/bdii/bdii-update.pid
```

bdii_database

[infosys/ldap]

Synopsis: `bdii_database = backend_type`

Description: Configure what ldap database backend should be used. If left `undefined` it will default to `hdb` for openldap versions up to 2.4 and to `mdb` for openldap versions 2.5 and later.

Default: `undefined`

Example:

```
bdii_database=hdb
```

bdii_conf

[infosys/ldap]

Synopsis: `bdii_conf = path`

Description: Location of the bdii config file generated by ARC.

Default: `$VAR{[infosys/ldap]infosys_ldap_run_dir}/bdii.conf`

Example:

```
bdii_conf=/run/arc/infosys/bdii.conf
```

bdii_update_cmd

[infosys/ldap]

Synopsis: `bdii_update_cmd = path`

Description: path to bdii-update script

Default: `$VAR{bdii_location}/sbin/bdii-update`

Example:

```
bdii_update_cmd=/usr/sbin/bdii-update
```

bdii_db_config

[infosys/ldap]

Synopsis: `bdii_db_config = path`

Description: path to slapd database configuration file

Default: `/etc/bdii/DB_CONFIG`

Example:

```
bdii_db_config=/etc/bdii/DB_CONFIG
```

bdii_archive_size

[infosys/ldap]

Synopsis: `bdii_archive_size = number`

Description: Sets BDII_ARCHIVE_SIZE in bdii configuration file

Default: `0`

Example:

```
bdii_archive_size=0
```

bdii_breathe_time

[infosys/ldap]

Synopsis: `bdii_breathe_time = number`

Description: Sets BDII_BREATHE_TIME in bdii configuration file

Default: `10`

Example:

```
bdii_breathe_time=10
```

bdii_delete_delay

[infosys/ldap]

Synopsis: `bdii_delete_delay = number`

Description: Sets BDII_DELETE_DELAY in bdii configuration file

Default: `0`

Example:

```
bdii_delete_delay=0
```

bdii_read_timeout

[infosys/ldap]

Synopsis: bdii_read_timeout = number

Description: Sets BDII_READ_TIMEOUT in bdii configuration file

Default: \$EVAL{\$VAR{bdii_provider_timeout} + \$VAR{[arex]infoproviders_timelimit} + \$VAR{[arex]wakeupperiod}}

Example:

```
bdii_read_timeout=300
```

Infosys Schema sub-blocks: The following infosys sub-blocks enable information publishing according to various information schema. In order to publish information in a certain schema, the corresponding sub-block must be defined in addition to the schema-neutral [infosys/cluster] and [queue:name] blocks! Comment out a specific schema block if you don't want to publish a specific information schema representation. Currently available information model (schema) sub-blocks:

- [infosys/nordugrid] - The native ARC info representation of a cluster and its queues
- [infosys/glue2] - The GLUE2 information model, both LDAP and XML (the latter is for WS-interface)
- [infosys/glue2/ldap] - The LDAP rendering of the GLUE2 model

[infosys/nordugrid] block

Enables the publication of the NorduGrid information model in the LDAP-based infosys. See the NORDUGRID-TECH-4 for schema definition. The configuration block does not contain any parameter. The information tree is populated based on the contents of the schema-neutral [infosys/cluster] and [queue:name] blocks.

[infosys/glue2] block

Enables the publication of the GLUE2 information model both in the LDAP and XML rendering. The information tree is populated based on the contents of the schema-neutral [infosys/cluster] and [queue:name] blocks and the GLUE2 specific schema sub-blocks.

admindomain_name

[infosys/glue2]

Synopsis: admindomain_name = string

Description: The Name attribute for the admindomain. This will show in top-BDII to group the resources belonging to this cluster. To group a bunch of clusters under the same AdminDomain, just use the same name. If not specified, will default to UNDEFINEDVALUE.

Default: UNDEFINEDVALUE

Example:

```
admindomain_name=ARC-TESTDOMAIN
```

admindomain_description

[infosys/glue2]

Synopsis: admindomain_description = text

Description: The free-form description of this domain.

Default: undefined

Example:

```
admindomain_description=ARC test Domain
```

admindomain_www

[infosys/glue2]

Synopsis: admindomain_www = url

Description: The URL pointing at a site holding information about the AdminDomain.

Default: undefined

Example:

```
admindomain_www=http://www.nordugrid.org/
```

admindomain_distributed

[infosys/glue2]

Synopsis: admindomain_distributed = yes/no

Description: Set this to yes if the domain is distributed that means, if the resources belonging to the domain are considered geographically distributed.

Allowed values: yes, no

Default: no

Example:

```
admindomain_distributed=yes
```

admindomain_owner

[infosys/glue2]

Synopsis: admindomain_owner = email

Description: The contact email of a responsible person for the domain

Default: undefined

Example:

```
admindomain_owner=admin@nordugrid.org
```

admindomain_otherinfo

[infosys/glue2]

Synopsis: `admindomain_otherinfo = text`

Description: Free-form text that fills the OtherInfo GLUE2 field. no need to set, used only for future development.

Default: undefined

Example:

```
admindomain_otherinfo=Test Other info
```

computingservice_qualitylevel

[infosys/glue2]

Synopsis: `computingservice_qualitylevel = qllevel`

Description: Allows a sysadmin to define different GLUE2 QualityLevel values for A-REX. Refer to GLUE2 documentation for the qualitylevel definitions.

Allowed values: production, pre-production, testing, development

Default: production

Example:

```
computingservice_qualitylevel=production
```

[infosys/glue2/ldap] block

Enables the publication of the LDAP-rendering of the GLUE2 infomodel.

showactivities

[infosys/glue2/ldap]

Synopsis: `showactivities = yes/no`

Description: Enables GLUE2 ComputingActivities in the LDAP rendering

Allowed values: yes, no

Default: no

Example:

```
showactivities=no
```

[infosys/cluster] block

Information schema-neutral blocks [infosys/cluster] and [queue:NAME] contain attributes that describe the computing cluster together with its queues. The parameters are available for every information model/schema representation.

This block describes the cluster characteristics of a Computing Element. The information specified here is mostly used by the Infosys ARC component.

alias

[infosys/cluster]

Synopsis: alias = text

Description: An arbitrary alias name of the cluster, optional.

Default: undefined

Example:

```
alias=Big Blue Cluster in Nowhere
```

hostname

[infosys/cluster]

Synopsis: hostname = fqdn

Description: Set the FQDN of the frontend.

Default: \$VAR{[common]hostname}

Example:

```
hostname=myhost.org
```

interactive_contactstring

[infosys/cluster]

Synopsis: interactive_contactstring = url

Description: the contact URL for interactive logins, set this if the cluster supports some sort of grid-enabled interactive login (gsi-ssh),

This option is **multivalued**.

Default: undefined

Example:

```
interactive_contactstring=gsissh://frontend.cluster:2200
```

Warning

CHANGE: REMOVED obsoleted in ARC 7.0.0

comment

[infosys/cluster]

Synopsis: comment = text

Description: Free text field for additional comments on the cluster in a single line, no newline character is allowed!

Default: undefined

Example:

```
comment=This cluster is specially designed for XYZ applications: www.xyz.org
```

cluster_location

[infosys/cluster]

Synopsis: cluster_location = formatted_string

Description: The geographical location of the cluster, preferably specified as a postal code with a two letter country prefix

Default: undefined

Example:

```
cluster_location=DK-2100
```

cluster_owner

[infosys/cluster]

Synopsis: cluster_owner = text

Description: It can be used to indicate the owner of a resource, multiple entries can be used

This option is **multivalued**.

Default: undefined

Example:

```
cluster_owner=World Grid Project
cluster_owner=University of NeverLand
```

advertisedvo

[infosys/cluster]

Synopsis: advertisedvo = vo_name

Description: This attribute is used to advertise which VOs are authorized on the cluster. Add only one VO for each advertisedvo entry. Multiple VOs in the same line will cause errors. These entries will be shown in all GLUE2 AccessPolicy and MappingPolicy objects, that is, they will apply for all Endpoints(Interfaces) and all Shares(currently queues). You can override the advertisedvos per queue. The information is also published in the NorduGrid schema.

Note

it is IMPORTANT to understand that this parameter is NOT enforcing any access control, it is just for information publishing!

This option in **multivalued**.

Default: undefined

Example:

```
advertisedvo=atlas
advertisedvo=community.nordugrid.org
```

clustersupport

[infosys/cluster]

Synopsis: clustersupport = email

Description: This is the support email address of the resource.

This option in **multivalued**.

Default: undefined

Example:

```
clustersupport=arc.support@mysite.org
clustersupport=arc.support@myproject.org
```

homogeneity

[infosys/cluster]

Synopsis: homogeneity = True/False

Description: Determines whether the cluster consists of identical NODES with respect to cputype, memory, installed software (opsys). The frontend is NOT needed to be homogeneous with the nodes. In case of inhomogeneous nodes, try to arrange the nodes into homogeneous groups assigned to a queue and use queue-level attributes. False may trigger multiple GLUE2 ExecutionEnvironments to be published if applicable.

Allowed values: True, False

Default: True

Example:

```
homogeneity=True
```

architecture

[infosys/cluster]

Synopsis: `architecture = string`

Description: Sets the hardware architecture of the NODES. The `architecture` is defined as the output of the `uname -m` (e.g. `i686`). Use this cluster attribute if only the NODES are homogeneous with respect to the architecture. Otherwise the queue-level attribute may be used for inhomogeneous nodes. If the frontend's architecture agrees to the nodes, the `adotf` (Automatically Determine On The Frontend) can be used to request automatic determination.

Default: `adotf`

Example:

```
architecture=adotf
```

opsys

[infosys/cluster]

Synopsis: `opsys = formatted_string`

Description: This multivalued attribute is meant to describe the operating system of the computing NODES. Set it to the `opsys` distribution of the NODES and not the frontend! `opsys` can also be used to describe the kernel or libc version in case those differ from the originally shipped ones. The distribution name should be given as `distro-name-version.number`, where spaces are not allowed. Kernel version should come in the form `kernelname-version.number`. If the NODES are inhomogeneous with respect to this attribute do NOT set it on cluster level, arrange your nodes into homogeneous groups assigned to a queue and use queue-level attributes. If `opsys=adotf`, will result in Automatic Determination of the Operating System On The Frontend, which should only be used if the frontend has the same OS as the nodes. The `adotf` discovered values will be used to fill GLUE2 `OSName`, `OSVersion` and `OSFamily` unless these values are explicitly defined for each queue. See the `[queue:queue-name]` block for their usage.

Note

any custom value other than `adotf` does NOT affect values in the GLUE2 schema.

This option is **multivalued**.

Default: `adotf`

Example:

```
opsys=Linux-2.6.18
opsys=glibc-2.5.58
opsys=CentOS-5.6
```

nodecpu

[infosys/cluster]

Synopsis: nodecpu = formatted_string

Description: This is the cputype of the homogeneous nodes. The string is constructed from the /proc/cpuinfo as the value of model name and @ and value of cpu MHz. Do NOT set this attribute on cluster level if the NODES are inhomogeneous with respect to cputype, instead arrange the nodes into homogeneous groups assigned to a queue and use queue-level attributes. Setting the nodecpu=adotf will result in Automatic Determination On The Frontend, which should only be used if the frontend has the same cputype as the homogeneous nodes.

Default: adotf

Example:

```
nodecpu=AMD Duron(tm) Processor @ 700 MHz
```

nodememory

[infosys/cluster]

Synopsis: nodememory = number

Description: This is the amount of memory (specified in MB) on the node which can be guaranteed to be available for the application. Please note in most cases it is less than the physical memory installed in the nodes. Do NOT set this attribute on cluster level if the NODES are inhomogeneous with respect to their memories, instead arrange the nodes into homogeneous groups assigned to a queue and use queue-level attributes.

Default: undefined

Example:

```
nodememory=64000
```

middleware

[infosys/cluster]

Synopsis: middleware = string

Description: The multivalued attribute shows the installed grid software on the cluster. Nordugrid-ARC is automatically set, no need to specify

This option in **multivalued**.

Default: undefined

Example:

```
middleware=my software
```

nodeaccess

[infosys/cluster]

Synopsis: nodeaccess = inbound/outbound

Description: Determines how the nodes can connect to the internet. Not setting anything means the nodes are sitting on a private isolated network. outbound access means the nodes can connect to the outside world while inbound access means the nodes can be connected from outside. inbound & outbound access together means the nodes are sitting on a fully open network.

This option is **multivalued**.

Default: undefined

Allowed values: inbound, outbound

Example:

```
nodeaccess=inbound
nodeaccess=outbound
```

localse

[infosys/cluster]

Synopsis: localse = url

Description: This multivalued parameter tells the BROKER that certain URLs (and locations below that) should be considered locally available to the cluster.

This option is **multivalued**.

Default: undefined

Example:

```
localse=gsiftp://my.storage/data1/
localse=gsiftp://my.storage/data2/
```

cpudistribution

[infosys/cluster]

Synopsis: cpudistribution = formatted_string

Description: This is the CPU distribution over nodes given in the form ncpu:m where:

n is the number of CPUs per machine m is the number of such machines

Example: 1cpu:3,2cpu:4,4cpu:1 represents a cluster with 3 single CPU machines, 4 dual CPU machines and one machine with 4 CPUs.

Default: undefined

Example:

```
cpudistribution=1cpu:3,2cpu:4,4cpu:1
```

maxcputime

[infosys/cluster]

Synopsis: maxcputime = number

Description: This is the maximum CPU time specified in seconds that the LRMS can allocate for the job. The default if not defined is that infoproviders get this value automatically from the LRMS. The purpose of this option is to tweak and override discovered value, or publish this value in case the LRMS module do not support automatic detection.

Default: undefined

Example:

```
maxcputime=3000000
```

mincputime

[infosys/cluster]

Synopsis: mincputime = number

Description: This is the minimum CPU time specified in seconds that the LRMS can allocate for the job. The default if not defined is that infoproviders get this value automatically from the LRMS. The purpose of this option is to tweak and override discovered value, or publish this value in case the LRMS module do not support automatic detection.

Default: undefined

Example:

```
mincputime=1200
```

maxwalltime

[infosys/cluster]

Synopsis: maxwalltime = number

Description: This is the maximum Wall time specified in seconds that the LRMS can allocate for the job. The default if not defined is that infoproviders get this value automatically from the LRMS. The purpose of this option is to tweak and override discovered value, or publish this value in case the LRMS module do not support automatic detection.

Default: undefined

Example:

```
maxwalltime=6000000
```

minwalltime

[infosys/cluster]

Synopsis: minwalltime = number

Description: This is the minimum Wall time specified in seconds that the LRMS can allocate for the job. The default if not defined is that infoproviders get this value automatically from the LRMS. The purpose of this option is to tweak and override discovered value, or publish this value in case the LRMS module do not support automatic detection.

Default: undefined

Example:

```
maxwalltime=1800
```

[infosys/accesscontrol] block

AREX allows to control access to public information for non-authorized users. If this block contains no entries public information is available to anyone.

Warning

CHANGE: NEW in 7.0.0

allowaccess

[infosys/accesscontrol]

Synopsis: allowaccess = authgroup

Description: Defines that only the specified authgroup members are authorized to access public information. For more information see similar configuration option in [arex/ws/jobs] block.

Default: undefined

This option is **multivalued**.

Example:

```
allowaccess=monitors
```

denyaccess

[infosys/accesscontrol]

Synopsis: denyaccess = authgroup

Description: Defines that the specified authgroup members are REJECTED, not authorized to access public information. For more information see similar configuration option in [arex/ws/jobs] block.

Default: undefined

This option is **multivalued**.

Example:

```
denyaccess=badactors
```

[queue:name] block

Each grid-enabled queue on the cluster should be represented and described by a separate queue block. The `queue_name` should be used as a label in the block name. In case of fork, or other LRMSes with no queue names, just use any unique string. A queue can represent a PBS/LSF/SGE/SLURM/LL queue, a SGE pool, a Condor pool or a single machine in case 'fork' type of LRMS. This block describes the queue characteristics.

homogeneity

[queue:name]

Synopsis: `homogeneity = True/False`

Description: determines whether the queue consists of identical NODES with respect to cputype, memory, installed software (opsys). In case of inhomogeneous nodes, try to arrange the nodes into homogeneous groups and assigned them to a queue. Possible values: True,False, the default is True.

Allowed values: True, False

Default: `$VAR{[infosys/cluster]homogeneity}`

Example:

```
homogeneity=True
```

comment

[queue:name]

Synopsis: `comment = text`

Description: A free-form text field for additional comments on the queue in a single line, no newline character is allowed!

Default: undefined

Example:

```
comment=This queue is nothing more than a condor pool
```

pbs_queue_node

[queue:name]

Synopsis: `pbs_queue_node = string`

Description: In PBS you can assign nodes to a queue (or a queue to nodes) by using the `node` property mark in PBS config.

Essentially, `pbs_queue_node` value is used to construct `nodes= string` in PBS script, such as `nodes=count:pbs_queue_node` where `count` is taken from the job description (1 if not specified).

This corresponds to setting the following parameter in PBS for this queue:

```
resources_default.neednodes = cpu_topology[:pbs_queue_node]
```

Setting the `pbs_queue_node` changes how the `queue-totalcpus`, `user freecpus` are determined for this queue.

You shouldn't use this option unless you are sure that your PBS configuration makes use of the above configuration. Read Nordugrid PBS instructions for more information: <http://www.nordugrid.org/documents/pbs-config.html>

Default: undefined

Example:

```
pbs_queue_node=gridlong_nodes
pbs_queue_node=ppn=4:ib
```

sges_jobopts

[queue:name]

Synopsis: sges_jobopts = string

Description: Per-queue override of additional SGE options to be used when submitting jobs to SGE to this queue

Default: undefined

Example:

```
sges_jobopts=-P atlas -r yes
```

condor_requirements

[queue:name]

Synopsis: condor_requirements = string

Description: It may be defined for each Condor queue. Use this option to determine which nodes belong to the current queue. The value of condor_requirements must be a valid constraints string which is recognized by a condor_status -constraint ... command. It can reference pre-defined ClassAd attributes (like Memory, Opsys, Arch, HasJava, etc) but also custom ClassAd attributes. To define a custom attribute on a condor node, just add two lines like the ones below in the \$(hostname).local config file on the node:

```
NORDUGRID_RESOURCE=TRUE
STARTD_EXPRS = NORDUGRID_RESOURCE, $(STARTD_EXPRS)
```

A job submitted to this queue is allowed to run on any node which satisfies the condor_requirements constraint. If condor_requirements is not set, jobs will be allowed to run on any of the nodes in the pool. When configuring multiple queues, you can differentiate them based on memory size or disk space, for example.

Default: \$VAR{[lrms]condor_requirements}

Example:

```
condor_requirements=(OpSys == "linux" && NORDUGRID_RESOURCE && Memory >= 1000 &&
↪Memory < 2000)
```

slurm_requirements

[queue:name]

Synopsis: slurm_requirements = string

Description: Use this option to specify extra SLURM-specific parameters.

Default: undefined

Example:

```
slurm_requirements=memory on node >> 200
```

totalcpus

[queue:name]

Synopsis: totalcpus = number

Description: Manually sets the number of cpus assigned to the queue. No need to specify the parameter in case the queue_node_string method was used to assign nodes to the queue (this case it is dynamically calculated and the static value is overwritten) or when the queue have access to the entire cluster (this case the cluster level totalcpus is the relevant parameter).

Default: undefined

Example:

```
totalcpus=32
```

queue-level configuration parameters: nodecpu, nodememory, architecture, opsys should be set if they are homogeneous over the nodes assigned to the queue AND they are different from the cluster-level value. Their meanings are described in the [infosys/cluster] block. Usage: this queue collects nodes with nodememory=512 while another queue has nodes with nodememory=256 -> don't set the cluster attributes but use the queue-level attributes. When the frontend's architecture or cputype agrees with the queue nodes, the `adotf` (Automatically Determine On The Frontend) can be used to request automatic determination of architecture or nodecpu. For GLUE2, fine tune configuration of ExecutionEnvironments' OSName, OSVersion, OSFamily is allowed with dedicated options `osname,osversion,osfamily`.

nodecpu

[queue:name]

Synopsis: nodecpu = formatted_string

Description: see description at [infosys/cluster] block

Default: \$VAR{[infosys/cluster]nodecpu}

Example:

```
nodecpu=AMD Duron(tm) Processor @ 700 MHz
```

nodememory

[queue:name]

Synopsis: nodememory = number

Description: see description at [infosys/cluster] block

Default: \$VAR{[infosys/cluster]nodememory}

Example:

```
nodememory=512
```

defaultmemory

[queue:name]

Synopsis: defaultmemory = number

Description: The LRMS memory request of job to be set by the LRMS backend scripts, if a user submits a job without specifying how much memory should be used. The order of precedence is: job description -> [lrms-defaultmemory] -> [queue-defaultmemory]. This is the amount of memory (specified in MB) that a job will request.

Default: undefined

Example:

```
defaultmemory=512
```

architecture

[queue:name]

Synopsis: architecture = string

Description: see description at [infosys/cluster] block

Default: \$VAR{[infosys/cluster]architecture}

Example:

```
architecture=adotf
```

opsys

[queue:name]

Synopsis: opsys = formatted_string

Description: see description at [infosys/cluster] block If osname, osversion are present, the values in opsys are ignored.

This option is **multivalued**.

Default: \$VAR{[infosys/cluster]opsys}

Example:

```
opsys=Linux-2.6.18
opsys=glibc-2.5.58
```

osname

[queue:name]

Synopsis: osname = string

Description: Only for GLUE2 overrides values defined in opsys for a single ExecutionEnvironment. Configuration of multiple ExecutionEnvironment for the same queue is not supported. Create a different queue for that.

Default: undefined

Example:

```
osname=Ubuntu
```

osversion

[queue:name]

Synopsis: `osversion = string`

Description: Only for GLUE2 overrides values defined in opsys for a single ExecutionEnvironment. Configuration of multiple ExecutionEnvironment for the same queue is not supported. Create a different queue for that.

Default: undefined

Example:

```
osversion=12.04
```

osfamily

[queue:name]

Synopsis: `osfamily = string`

Description: Only for GLUE2 overrides values defined in opsys for a single ExecutionEnvironment. Configuration of multiple ExecutionEnvironment for the same queue is not supported. Create a different queue for that.

Default: undefined

Example:

```
osfamily=linux
```

benchmark

[queue:name]

Synopsis: `benchmark = name value`

Description: Defines resource benchmark results for accounting and information publishing. The nodes in the same queue are assumed to be homogeneous with respect to the benchmark performance. In case of multiple benchmarks are specified:

- Accounting subsystem will use ONLY THE FIRST defined benchmark.
- Infosys will publish all defined benchmark values.

The values represent per-core CPU performance.

Note

APEL accounting services supports HEPscore23, HEPSPEC or Si2k benchmark types only.

This option is **multivalued**.

Default: HEPSPEC 1.0

Example:

```
benchmark=HEPscore23 16.5
benchmark=HEPSPEC 12.26
benchmark=Si2k 3065
```

allowaccess

[queue:name]

Synopsis: allowaccess = authgroup

Description: Defines that the specified authgroup members are authorized to submit jobs to this queue of ARC-CE after the user already granted access to the CE via one of the interfaces. A related config option the denyaccess (see below) can be used to deny submission to the queue. Multiple allowaccess and denyaccess authorization statements are allowed within a configuration block. These statements are processed sequentially in the order they are specified in the config block. The processing stops on first allowaccess or denyaccess statement matching the authgroup membership. If there are no authorization statements specified, then the queue is accessible by everyone already authorized.

Default: undefined

This option is **multivalued**.

Example:

```
allowaccess=biouers
allowaccess=atlasusers
```

denyaccess

[queue:name]

Synopsis: denyaccess = authgroup

Description: Defines that the specified authgroup members are NOT allowed to submit jobs to this queue of ARC-CE after despite the user is already granted access to the CE via one of the interfaces. A related config option the allowaccess (see below) can be used to grant job submission to the queue. Multiple allowaccess and denyaccess authorization statements are allowed within a configuration block. These statements are processed sequentially in the order they are specified in the config block. The processing stops on first allowaccess or denyaccess statement matching the authgroup membership. If there are no authorization statements specified, then the queue is accessible by everyone already authorized.

Default: undefined

This option is **multivalued**.

Example:

```
denyaccess=blacklisted-for-the-queue
```

advertisedvo

[queue:name]

Synopsis: advertisedvo = vo_name

Description: This attribute is used to advertise which VOs are authorized on the [queue:name] of the cluster. Add only one VO for each advertisedvo entry. Multiple VOs in the same line will cause errors. These entries will be shown in the MappingPolicy objects, that is, they will apply for the Shares that corresponds to the queue. The information is also published in the NorduGrid schema.

Note

if you have also configured advertisedvo in the [infosys/cluster] block, the result advertised VOs per queue will override whatever is defined in [infosys/cluster] block!

Note

it is IMPORTANT to understand that this parameter is NOT enforcing any access control, it is just for information publishing!

This option in **multivalued**.

Default: \$VAR{[infosys/cluster]advertisedvo}

Example:

```
advertisedvo=atlas
advertisedvo=community.nordugrid.org
```

maxslotsperjob

[queue:name]

Synopsis: maxslotsperjob = number

Description: This GLUE2 specific parameter configures the MaxSlotsPerJob value on a particular queue. This value is usually generated by LRMS infocollectors, but there are cases in which a system administrator might like to tweak it. Default is to publish what is returned by the LRMS, and if nothing is returned, NOT to publish the MaxSlotsPerJob attribute. If a system administrator sets the value here, that value will be published instead, regardless of what the LRMS returns. Each LRMS might have a different meaning for this value.

Default: undefined

Example:

```
maxslotsperjob=5
```

forcedefaultvoms

[queue:name]

Synopsis: forcedefaultvoms = VOMS_FQAN

Description: specify VOMS FQAN which user will be assigned if his/her credentials contain no VOMS attributes.

Default: \$VAR{[arex]forcedefaultvoms}

Example:

```
forcedefaultvoms=/vo/group/subgroup
```

maxcputime

[queue:name]

Synopsis: maxcputime = number

Description: This value overrides the one defined in the [infosys/cluster] block. See description in that block.

Default: undefined

Example:

```
maxcputime=3000000
```

mincputime

[queue:name]

Synopsis: mincputime = number

Description: This value overrides the one defined in the [infosys/cluster] block. See description in that block.

Default: undefined

Example:

```
mincputime=1200
```

maxwalltime

[queue:name]

Synopsis: maxwalltime = number

Description: This value overrides the one defined in the [infosys/cluster] block. See description in that block.

Default: undefined

Example:

```
maxwalltime=6000000
```

minwalltime

[queue:name]

Synopsis: minwalltime = number

Description: This value overrides the one defined in the [infosys/cluster] block. See description in that block.

Default: undefined

Example:

```
minwalltime=1800
```

[datadelivery-service] block

This block configures and enables the data delivery service. This service is intended to off-load data-staging from A-REX and usually deployed on one or more separate machines.

This service can also act as an independent data transfers service that case it would require an intelligent data manager that could replace A-REX's intelligence.

transfer_dir

[datadelivery-service]

Synopsis: *transfer_dir = path

Description: The directori(es) on the DDS host in which the service is allowed to read and write. When DDS is used as a remote transfer service assisting A-REX then this is usually one or more cache and/or session directories shared as a common mount with A-REX.

This option is **multivalued**.

Default: undefined

Example:

```
transfer_dir=/shared/arc/cache
transfer_dir=/shared/arc/session
```

hostname

[datadelivery-service]

Synopsis: hostname = FQDN

Description: The hostname of the machine on which DDS service runs.

Default: \$EXEC{hostname -f}

Example:

```
hostname=localhost
```

port

[datadelivery-service]

Synopsis: port = port

Description: Port on which service listens

Default: 443

Example:

```
port=8443
```

pidfile

[datadelivery-service]

Synopsis: pidfile = path

Description: pid file of the daemon

Default: /run/arched-datadelivery-service.pid

Example:

```
pidfile=/run/arched-datadelivery-service.pid
```

logfile

[datadelivery-service]

Synopsis: logfile = path

Description: log file of the daemon

Default: /var/log/arc/datadelivery-service.log

Example:

```
logfile=/tmp/delivery.log
```

loglevel

[datadelivery-service]

Synopsis: loglevel = level

Description: set loglevel of the data delivery service between 0 (FATAL) and 5 (DEBUG). Defaults to 3 (INFO).

Allowed values: 0, 1, 2, 3, 4, 5

Default: 3

Example:

```
loglevel=4
```

user

[datadelivery-service]

Synopsis: user = username

Description: Overwrites the user under which the service runs. The default is the user starting the service. DDS is very limited if not run as root.

Default: undefined

Example:

```
user=ddsuser
```

secure

[datadelivery-service]

Synopsis: secure = yes/no

Description: Set to no if the service should run without a host certificate. In this case the corresponding deliveryservice option in the [arex/data-staging] A-REX configuration block should use http rather than https URLs.

Allowed values: yes, no

Default: yes

Example:

```
secure=no
```

allowed_ip

[datadelivery-service]

Synopsis: *allowed_ip = ip

Description: IP address authorized to access service. Normally this is the A-REX host IP. By default the delivery service listens on all available interfaces, so if both IPv4 and IPv6 are enabled on this and the A-REX host, remember to add both A-REX host IPs here.

This option is **multivalued**.

Default: undefined

Example:

```
allowed_ip=192.0.2.1
allowed_ip=2001:db8:85a3::8a2e:370:7334
```

allowed_dn

[datadelivery-service]

Synopsis: allowed_dn = DN

Description: DN authorized to access service. This option restricts access to specified DNs (of the users who submit jobs to A-REX). It is only effective if secure=yes.

This option is **multivalued**.

Default: undefined

Example:

```
allowed_dn=/O=Grid/O=Big VO/CN=Main Boss
```

x509_host_key

[datadelivery-service]

Synopsis: x509_host_key = path

Description: Optional parameter to overwrite [common] block values.

Default: \$VAR{[common]x509_host_key}

Example:

```
x509_host_key=/etc/grid-security/hostkey.pem
```

x509_host_cert

[datadelivery-service]

Synopsis: x509_host_cert = path

Description: Optional parameter to overwrite [common] block values.

Default: \$VAR{[common]x509_host_cert}

Example:

```
x509_host_cert=/etc/grid-security/hostcert.pem
```

x509_cert_dir

[datadelivery-service]

Synopsis: x509_cert_dir = path

Description: Optional parameter to overwrite [common] block values.

Default: \$VAR{[common]x509_cert_dir}

Example:

```
x509_cert_dir=/etc/grid-security/certificates
```

[custom:name] block

This optional block is for those who wish to include non-ARC configuration in `arc.conf`. Custom blocks will be ignored by ARC components including the configuration validator. Any non-ARC configuration which is not in a custom block will be flagged as an error by the validator and A-REX will not start.

6.1.3 Removed blocks and options

This is the `arc.conf DELETED` file that contains all the configuration blocks and options that have been DELETED in ARC version 7.0.0 and later

[deleted:blocks] block

Following blocks and corresponding functionality are removed complete from ARC7 release and should be cleaned up from previous ARC6 configuration:

```
[authtokens] (always enabled in ARC7)
[lrms/ssh]
[arex/ws/publicinfo] (always enabled in ARC7)
[arex/ws/argus]
[gridftpd]
[gridftpd/jobs]
[gridftpd/filedir]
[infosys/glue1]
[infosys/glue1/site-bdii]
[acix-scanner]
[acix-index]
[userlist:name]
[nordugridmap]
```

Note

Options marked DELETED without stating a version were deleted in version 7.0.0 compared to the latest ARC6 supported configuration.

[authgroup:groupname] block

userlist

[authgroup:groupname]

Synopsis: `userlist = ulist_name [ulist_name ...]`

Description: Match user belonging to `ulist_name` defined in an earlier `[userlist:ulist_name]` block. Multiple `userlist` names are allowed for this rule.

This is **sequenced** option.

Default: undefined

Example:

```
userlist=biouers
```

Warning

CHANGE: DELETED

[arex/data-staging] block**use_remote_acix**

[arex/data-staging]

Synopsis: `use_remote_acix` = URL

Description: If configured then the ARC Cache Index, available at the URL, will be queried for every input file specified in a job description and any replicas found in sites with accessible caches will be added to the replica list of the input file. The replicas will be tried in the order specified by `preferredpattern` variable.

Default: undefined*Example:*

```
use_remote_acix=https://cacheindex.ndgf.org:6443/data/index
```

Warning

CHANGE: DELETED

6.2 ARC CE Deployment and Operation

6.2.1 Quickstart ARC: towards distributed computing in a few minutes

A few minutes is all it takes to set up an ARC *Computing Element* and try common distributed computing workflows.

ARC7 comes with so-called *zero configuration* (`/etc/arc.conf` and `/etc/arc.conf.d`) and works out of the box without any manual configuration. It has a pre-installed x509 host certificate signed by a Test-CA.

You can try ARC by using the legacy x509 user certificate, or with the newer Jason Web Token (JWT) capability. The procedure below splits into x509 versus token at [Step 6](#). The two require slightly different configuration options on the ARC server, and different procedures to acquire the authentication document (certificate or token).

The easiest and quickest way to get started is to install the ARC client on the same server as the ARC-CE.

If you want to test job submission on a remote ARC client, the client must trust the ARC-CE host certificate which is issued by the Test-CA, and you must therefore apply some extra steps for both the token and the x509 user case. These are described in [Step 6](#) in the “JWT Remote client” tab for the token case or “x509 Remote client” tab for the x509 case respectively.

Step 0. Prerequisites

- The zero configured A-REX comes with the REST interface enabled. It runs by default on port 443, so make sure it is not firewalled.
- If you are testing ARC with a remote client: Either register your ARC-CE on a DNS server, or map the ARC-CE host name to the `/etc/hosts` file on the client host.

Step 1. Enable Nordugrid ARC7 repos

Repository security

The Nordugrid RPM packages and DEB repositories are signed, and in order for the repository tools APT and YUM to verify them you must install the Nordugrid GPG key:

RHEL

Ubuntu

Debian

For rpm based distributions like Red Hat Enterprise Linux and Fedora:

```
[root ~]# rpm --import http://download.nordugrid.org/RPM-GPG-KEY-nordugrid-7
```

For Ubuntu distributions with sudo:

```
[user ~]$ wget -q http://download.nordugrid.org/DEB-GPG-KEY-nordugrid-7.asc -O- | \
↪ sudo apt-key add -
```

For Debian without sudo:

```
[root ~]# wget -q http://download.nordugrid.org/DEB-GPG-KEY-nordugrid-7.asc -O- | apt-
↪ key add -
```

Repository configuration

RHEL

Ubuntu

Debian

The Nordugrid ARC repositories for RedHat Enterprise Linux / CentOS packaging utility `dnf` can be configured through: `/etc/yum.repos.d/nordugrid.repo`

The repository configuration can be set up automatically with `dnf` by installing the `nordugrid-release` package or creating the configuration file manually.

Install with `dnf`

Install manually

The easiest way to configure DNF to use the Nordugrid repository for Red Hat Enterprise Linux, CentOS and similar distributions is to install the `nordugrid-release` package which can be found in the Nordugrid package repository for the appropriate RHEL/EPEL release.

Links to the release packages:

Rocky Linux/AlmaLinux

CentOS Stream

Fedora

Rocky Linux: 10 9 8

CentOS Stream: 10 9

Fedora: 41 40

Install with dnf (Fedora, CentOS Stream, Rocky Linux, CentOS Linux 8+9) by copying the appropriate link from above

```
[root ~]# dnf install <rhel-repo link>
```

This creates the appropriate repo files in `/etc/yum.repos.d/`.

For manual YUM repository setup, create a file `/etc/yum.repos.d/nordugrid.repo` with the following contents (here using Rocky as example, if you are on Fedora, replace `rocky` with `fedora`)

If you are installing an alpha, beta or release candiate, please set the `nordugrid-testing` to `enabled=1`.

```
[nordugrid]
name=NorduGrid - $basearch - base
baseurl=http://download.nordugrid.org/repos/7/rocky/$releasever/$basearch/
→base
enabled=1
gpgcheck=1
gpgkey=http://download.nordugrid.org/RPM-GPG-KEY-nordugrid-7

[nordugrid-updates]
name=NorduGrid - $basearch - updates
baseurl=http://download.nordugrid.org/repos/7/rocky/$releasever/$basearch/
→updates
enabled=1
gpgcheck=1
gpgkey=http://download.nordugrid.org/RPM-GPG-KEY-nordugrid-7

[nordugrid-testing]
name=NorduGrid - $basearch - testing
baseurl=http://download.nordugrid.org/repos/7/rocky/$releasever/$basearch/
→testing
enabled=0
gpgcheck=1
gpgkey=http://download.nordugrid.org/RPM-GPG-KEY-nordugrid-7
```

The NorduGrid ARC repositories for Ubuntu packaging utility APT can be configured through: `/etc/apt/sources.list` or when supported through a repo specific file: `/etc/apt/sources.list.d/nordugrid.list`.

The repository configuration can be set up automatically by means of installing `nordugrid-release` package or creating the configuration file manually.

Ubuntu version names:

Ubuntu release	Code name
25.04	plucky
24.04	noble
22.04	jammy

Install with dpkg

Install manually

The examples below give you the link for most recent Ubuntu releases. Packages are shown below for amd64 architecture. Replace `amd64` for `i386` if required for your architecture.

Ubuntu: 25.04 24.04 22.04

Install the source file with dpkg, example shown for Ubuntu 22.04:

```
[root ~]# wget -q https://download.nordugrid.org/packages/nordugrid-release/releases/
↪7/ubuntu/24.04/amd64/nordugrid-release_7~noble1_all.deb
[root ~]# dpkg -i nordugrid-release_7~noble1_all.deb
```

For a different version of Ubuntu, change the version names appropriately.

For manual APT repository setup for Ubuntu, the APT sources file should contain the following (here shown for Ubuntu 22.04 Jammy).

The configurations for the various APT based distributions can be found in the following sections. To enable a specific repository, remove the “#” from the beginning of the line, before the “deb” as shown for the Base Channel.

```
# Base channel - must be enabled
deb http://download.nordugrid.org/repos/7/ubuntu/ jammy main
deb-src http://download.nordugrid.org/repos/7/ubuntu/ jammy main

# Updates to the base release - should be enabled
deb http://download.nordugrid.org/repos/7/ubuntu/ jammy-updates main
deb-src http://download.nordugrid.org/repos/7/ubuntu/ jammy-updates main

# Scheduled package updates - optional
#deb http://download.nordugrid.org/repos/7/ubuntu/ jammy-experimental main
#deb-src http://download.nordugrid.org/repos/7/ubuntu/ jammy-experimental main
```

For a different release version, change the version name accordingly.

The NorduGrid ARC repositories for Ubuntu packaging utility APT can be configured through: /etc/apt/sources.list or when supported through a repo specific file: /etc/apt/sources.list.d/nordugrid.list.

The repository configuration can be set up automatically by means of installing nordugrid-release package or creating the configuration file manually.

Debian version names:

Debian release	Code name
13	trixie
12	bookworm
11	bullseye

Install with dpkg

Install manually

The examples below give you the link for most recent Debian releases. Packages are shown below for amd64 architecture. Replace amd64 for i386 if required for your architecture.

Debian: 13 12 11

Install the source file with dpkg, example shown for Debian 12:

```
[root ~]# wget -q https://download.nordugrid.org/packages/nordugrid-release/releases/
↪7/debian/12/amd64/nordugrid-release_7~bpo12+1_all.deb
[root ~]# dpkg -i nordugrid-release_7~bpo12+1_all.deb
```

For a different version of Debian, change the version names appropriately.

For manual APT repository setup for Debian, the APT sources file should contain the following (here shown for Debian 12 Bookworm).

The configurations for the various APT based distributions can be found in the following sections. To enable a specific repository, remove the “#” from the beginning of the line, before the “deb” as shown for the Base Channel.

```
# Base channel - must be enabled
deb http://download.nordugrid.org/repos/7/debian/ bookworm main
deb-src http://download.nordugrid.org/repos/7/debian/ bookworm main

# Updates to the base release - should be enabled
deb http://download.nordugrid.org/repos/7/debian/ bookworm-updates main
deb-src http://download.nordugrid.org/repos/7/debian/ bookworm-updates main

# Scheduled package updates - optional
#deb http://download.nordugrid.org/repos/7/debian/ bookworm-experimental main
#deb-src http://download.nordugrid.org/repos/7/debian/ bookworm-experimental main
```

For a different release version, change the version name accordingly.

Set up dependency repositories

RHEL

Ubuntu

Debian

RHEL8

RHEL9

```
[root ~]# dnf config-manager --set-enabled powertools && dnf makecache
```

```
[root ~]# dnf config-manager --set-enabled crb && dnf makecache
```

```
[root~]# apt-get update
```

```
[root~]# apt-get update
```

Step 2. Install A-REX

ARC Resource-coupled EXecution service (A-REX) is a core component that manages authentication, authorization and job life cycle. It is enough to have A-REX installed to have a minimal computing element:

RHEL

Ubuntu

Debian

```
[root ~]# dnf -y install nordugrid-arc-arex
```

```
[root ~]# apt-get -y install nordugrid-arc-arex
```

```
[root ~]# apt-get -y install nordugrid-arc-arex
```

Step 3. Run A-REX

To start ARC services just run:

```
[root ~]# arcctl service start --as-configured
```

You can check if A-REX is running with:

```
[root ~]# arcctl service list
arc-arex                (Installed, Disabled, Running)
arc-arex-ws             (Installed, Disabled, Running)
arc-datadelivery-service (Not installed, Disabled, Stopped)
arc-infosys-ldap        (Not installed, Disabled, Stopped)
```

Step 4. Install the ARC client

Install ARC client tools on the **client host**

Note

In the zero-conf setup - we install the client and the ARC control client tool on the same server as the ARC-CE, so client and host is the same machine. Typically you would install the client on another (remote) machine.

If you chose to use a remote machine, please repeat Step 1 on that machine.

RHEL

Ubuntu

Debian

```
[root@server]# dnf -y install nordugrid-arc-client nordugrid-arc-arcctl
```

```
[root@server]# apt-get -y install nordugrid-arc-client nordugrid-arc-arcctl
```

```
[root@server]# apt-get -y install nordugrid-arc-client nordugrid-arc-arcctl
```

Step 5. Install and enable autocompletion (optional)

arcctl tool automates many ARC CE operations and is designed with bash-completion in mind. If you would like to use ARC in production it is advised to have completion enabled.

Step 5a. Install and activate bash-completion

RHEL

Ubuntu

Debian

```
[root ~]# dnf install -y bash-completion python-argcomplete && activate-global-python-argcomplete
```

```
[root ~]# apt-get install -y bash-completion python-argcomplete && activate-global-python-argcomplete
```

```
[root ~]# apt-get install -y bash-completion python-argcomplete && activate-global-
↳python-argcomplete
```

Step 5b. Enable bash-completion

Open a new shell to enable bash-completion.

Step 6. Set up trust between ARC-CE and ARC client

JWT

x509

ARC comes with a test-JWT issuer that you can use to create user tokens for job submission in the zero-conf testing setup.

The issuer must be initialized before use, and in addition the ARC-CE must be configured to trust the JWT issuer. If the client is remote, there are additional steps to set up trust between the ARC-CE and the remote client.

Local client quick

Local client details

Remote client

If your ARC-CE and ARC client are the same machine you can apply the compressed command below that both initializes the test-jwt issuer and sets up the trust in one step.

```
[root@server]# $(arcctl test-jwt init --force | tail -n 1)
```

Step 6a.

For the zero-conf setup we will use ARC's inbuilt test-token issuer to submit a token. This is run on the ARC client machine. In our case this is the same machine as the ARC-CE server.

```
[user@server] arcctl test-jwt init
```

This will output an `arcctl deploy` command that needs to be issued on the ARC-CE server. Copy this to your clipboard.

Example output from `arcctl test-jwt init`:

```
Issuer URL: https://arc.example.org/arc/testjwt/8b7baf79
JWKS:
{
  "keys": [
    {
      "e": "AQAB",
      "kid": "testjwt",
      "kty": "RSA",
      "n": "r0nMfmRfhJFiyCPRUc8m9K7yl0qksmIRIQeiMNEi3_Und6WVNHlpERrzw6jTHu5wr_
↳Tk408ve-igludpqEZ5PUcV6K25MohYu1b6ifryDo6go-bQ0cEaEyZRYGm1scOUb_gWCAY0Le-
↳hv7hZGnQ3rojLZ2BJwUwBV0j5Hp_
↳ROPUdbifKfNkBiuJhGPJAegrPrKgsskQNA2GkXWACeS85WPKIQ54bkUiASsmz3_b0Ik9jQaQnHsU0znM3G-
↳EpjnLB-
↳1PS7FC1tIMaXcJ2BJZuFfkDyIv1Ymn8vKf9WeQjQ80L08k78pzTGOerZLcc5BQ2ZWEUhADWRWzkqHmEDymIw
↳",
      "use": "sig"
    }
  ]
}
```

(continues on next page)

(continued from previous page)

```
}
```

Run the following command on the ARC CE to trust the Test JWT issuer:

```
arcctl deploy jwt-issuer --deploy-conf test-jwt://H4sIAOzQRGYC/
↪ 73T3XKiMBQA4HfheqmgRbB3qFTx7CUys40AyFIiCQ0CTLY8d03brezT9C9yjlJznwnmeRDqaFI81SkytOHgjhvIV0elFKIhj8
↪ yB0Hvh4D5ZdWJgZWaWfUZE+aEIiFJIMkbioj4Lmjwh5FcyyFDpKDfL35JEpUMo7yBQCD6H876j5N22oqSMnRN/
↪ 4/
↪ 9xUFJVx3mScvQt2F3QL19QvfHimEvx58fCpSk7dtT2QRGuUz+lt5z0cs8ONgyJjJiGtKwDVCUq2fUz7wgBFY9WZv9WXvHvHYD1
↪ hRsy5QRSe9zZt3Jza8ELy010NjS8tjq2djVLDjnI5PVM18DTip08fBcVHrHOzDLD1FM/
↪ u430C1vJhlvCD+inFqEw+nqy7spq/7ylg2SbD3wjxDxbrY4Slqq3LhrWx4Yh5bnzjH/s4eLvBbZM/
↪ gwTtib+36xmOGQ2QfeH0dJZnm4knlpz5Z8lC7ku1ooTpNRTZTVfc05vNMF+42fQMrqcbtc4HnvXvRjzWxLutiEkG/
↪ 8i1to1nYtJrry2IPWbwBwJj6wzhywtKeR0F0xe/L2pn3tdt9/jR5vRydlNuv2+03J4Gqk1YEAAA=
```

Step 6b. On the A-REX server - setup trust of the client's Test JWT issuer

Run the `arcctl deploy jwt-issuer` command from Step 6a on the A-REX server.

Example:

```
[root@server]# arcctl deploy jwt-issuer --deploy-conf test-jwt://H4sIAOzQRGYC/
↪ 73T3XKiMBQA4HfheqmgRbB3qFTx7CUys40AyFIiCQ0CTLY8d03brezT9C9yjlJznwnmeRDqaFI81SkytOHgjhvIV0elFKIhj8
↪ yB0Hvh4D5ZdWJgZWaWfUZE+aEIiFJIMkbioj4Lmjwh5FcyyFDpKDfL35JEpUMo7yBQCD6H876j5N22oqSMnRN/
↪ 4/
↪ 9xUFJVx3mScvQt2F3QL19QvfHimEvx58fCpSk7dtT2QRGuUz+lt5z0cs8ONgyJjJiGtKwDVCUq2fUz7wgBFY9WZv9WXvHvHYD1
↪ hRsy5QRSe9zZt3Jza8ELy010NjS8tjq2djVLDjnI5PVM18DTip08fBcVHrHOzDLD1FM/
↪ u430C1vJhlvCD+inFqEw+nqy7spq/7ylg2SbD3wjxDxbrY4Slqq3LhrWx4Yh5bnzjH/s4eLvBbZM/
↪ gwTtib+36xmOGQ2QfeH0dJZnm4knlpz5Z8lC7ku1ooTpNRTZTVfc05vNMF+42fQMrqcbtc4HnvXvRjzWxLutiEkG/
↪ 8i1to1nYtJrry2IPWbwBwJj6wzhywtKeR0F0xe/L2pn3tdt9/jR5vRydlNuv2+03J4Gqk1YEAAA=
ARC CE now trust JWT signatures of https://arc.example.org/arc/testjwt/8b7baf79_
↪ issuer.
```

Auth configuration for issuer tokens has been written to `/etc/arc.conf.d/10-jwt-a7374e17.conf`

ARC restart is needed to apply configuration.

This command does two things

- 1) creates a `tokenissuers` folder in your control directory

Our control directory now contains a folder `tokenissuers/a7374e17` where the issuer url, key and metadata is stored:

```
[root@server]# ls /var/spool/arc/jobstatus/tokenissuers/a7374e17
issuer      keys      metadata
```

- 2) automatically sets up the arc configuration token authentication for tokens issued by this test-jwt issuer.

The token authentication file `10-jwt-a7374e17.conf` produced by the `arcctl deploy jwt-issuer` command looks like this in our example:

```
[root@server]# cat /etc/arc.conf.d/10-jwt-a7374e17.conf
[authgroup:testjwt-a7374e17]
authtokens = * https://arc.example.org/arc/testjwt/8b7baf79 arc * *

[mapping]
map_to_user = testjwt-a7374e17 nobody:nobody

[arex/ws/jobs]
allowaccess = testjwt-a7374e17
```

Here we see that a separate authgroup has been automatically created for tokens issued by this test-jwt issuer. Mapping of this authgroup is done to the nobody user (which is the only user we assume for zero-conf), and access is enabled for job submission by jobs issued with a token from this test-jwt issuer.

With the zero-conf setup A-REX is pre-installed with a x509 host certificate issued by the Test-CA. A remote client will need to trust this Test-CA. In addition the ARC-CE must trust the token issued by the Test-JWT issuer.

4 steps are needed to set up trust between the ARC-CE and the remote client:

Step 6a. On the A-REX server - printout Test-CA certificate

```
[user@server] arcctl test-ca info -o ca-cert
-----BEGIN CERTIFICATE-----
MIIFyTCCA7GgAwIBAgIUeLkSbksS9r3raPvkT2rR0ep06X8wDQYJKoZIhvcNAQEM
BQAwDETMBEGCgmSjomT8ixkArkWA29yZzEZMBcGCgmSjomT8ixkArkWCW5vcmR1
<output omitted>
TJ9f0I8ktHAClVlVJE9SIDWs2zPo8o4cmvLBAtxe+jaijn22THtpLLUSXt1ozexS
ZHGFtsUBuIoNzXoRXxJwkGBA1ZpLBbOpjyp6PzNcTPYFG51+EHTUMPKbfyQ5
-----END CERTIFICATE-----
```

Copy this output to your clipboard.

Step 6b. On the ARC client deploy the ca-cert

On your ARC client, issue the command below, and paste the CA certificate output from Step 6a when prompted.

```
[user@client]$ arcctl deploy ca-cert --x509-cert-dir ~/.globus
[2024-10-25 21:40:50,327] [ARCCTL.ThirdParty.Deploy] [INFO] [726706] [Deploying CA
↪Certificate to /etc/grid-security/certificates/ARCTestCAfdb0a5e3.pem]
[2024-10-25 21:40:50,328] [ARCCTL.ThirdParty.Deploy] [INFO] [726706] [Reading CA
↪Certificate PEM data from stdin]
-----BEGIN CERTIFICATE-----
MIIFyTCCA7GgAwIBAgIUeLkSbksS9r3raPvkT2rR0ep06X8wDQYJKoZIhvcNAQEM
BQAwDETMBEGCgmSjomT8ixkArkWA29yZzEZMBcGCgmSjomT8ixkArkWCW5vcmR1
<output omitted>
TJ9f0I8ktHAClVlVJE9SIDWs2zPo8o4cmvLBAtxe+jaijn22THtpLLUSXt1ozexS
ZHGFtsUBuIoNzXoRXxJwkGBA1ZpLBbOpjyp6PzNcTPYFG51+EHTUMPKbfyQ5
-----END CERTIFICATE-----
[2024-10-25 21:40:54,173] [ARCCTL.ThirdParty.Deploy] [INFO] [726706] [CA Certificate
↪for /DC=org/DC=nordugrid/DC=ARC/O=TestCA/CN=ARC TestCA fdb0a5e3 is deployed
↪successfully to /etc/grid-security/certificates/ARCTestCAfdb0a5e3.pem]
```

This will create all necessary files in your x509_cert_dir and allow your remote client to trust the host certificate of the ARC-CE server.

Next you must set up the Test-JWT in order to obtain a user token for job submission.

Step 6c. On the ARC client initialize the Test-JWT issuer

```
[user@client]$ arcctl test-jwt init
```

This will output an arcctl deploy command that needs to be issued on the ARC-CE server. Copy this to your clipboard. For explanations, see the tab “Local client details”.

Step 6d. On the ARC-CE server - setup trust of the client's Test JWT issuer

Copy the deploy command (last line of output) from Step 6c and run it on the ARC-CE server.

Example:

```
[root@server]# arcctl deploy jwt-issuer --deploy-conf test-jwt://H4sIAOzQRGYC/
↪73T3XXiMBQA4HfheqmgRbB3qFTxF7<...removed...output...>/jR5vRyd1Nuv2+03J4Gqk1YEAAA=
```

(continues on next page)

(continued from previous page)

```
ARC CE now trust JWT signatures of https://arc.example.org/arc/testjwt/8b7baf79_
↪issuer.

Auth configuration for issuer tokens has been written to /etc/arc.conf.d/10-jwt-
↪a7374e17.conf
ARC restart is needed to apply configuration.
```

You are now finally set, and trust between the ARC-CE and the remote client is established.

Note

When you set up your production ready service later on, you will remove the test-jwt related files, and configure authentication following instructions in [Authorization, Mapping and Queue selection rules](#).

During the test-user certificate generation, `arcctl test-ca` will automatically add the issued certificate subject to the `testCA.allowed-subjects` file, opening the job submission possibility to the test-user transparently. The `testCA.allowed-subjects` can be found in your `/etc/grid-security` folder.

No other subject will be able to submit to your system before you change the `authgroup` settings in `arc.conf`, which you will do once you configure ARC for production use.

You can test submission from the host running A-REX or from any other machine by following the instructions below.

Local client

Remote client

To generate a test certificate/key and install it to standard location inside a local user's home directory, run:

```
[root@server]# arcctl test-ca usercert --install-user user01
User certificate and key are installed to default /home/user01/.globus location for_
↪user user01.
```

Note

Replace `user01` with the actual username you want to submit jobs by. While it is technically possible to submit jobs as the root user, we strongly discourage that.

To generate a test certificate/key for a remote client, two steps are needed.

First create a tar-ball containing `usercert`, `ca-certs` and `setup` script using the TestCA automatically setup on your ARC-CE server:

```
[root@server]# arcctl test-ca usercert -t
[2025-08-06 18:44:20,340] [ARCCTL.TestCA] [INFO] [53878] [Certificate and key for_
↪user Test User 80404674 are exported to usercert-Test-User-80404674.tar.gz]
[2025-08-06 18:44:20,341] [ARCCTL.TestCA] [INFO] [53878] [Printing usage instructions_
↪for tarball]
tar xzf usercert-Test-User-80404674.tar.gz
source arc-testca-usercert/setenv.sh
```

In addition it adds the newly created x509 user certificate subject in the file `/etc/grid-security/testCA.allowed-subjects` which will ensure that jobs and other requests issued by a user with this certificate is accepted.

Next copy the tar-ball over to your client machine by a method of your choice (e.g. `scp`), and run the two commands as per instructions above:

```
[user@client]$ tar -xzf usercert-Test-User-80404674.tar.gz
[user@client]$ source arc-testca-usercert/setenv.sh
```

This sets up the necessary env vars and trust between the ARC-CE and ARC client machine.

Note

When you set up your production ready service later on, you will remove the test-CA related files, and configure authentication following instructions in *Authorization, Mapping and Queue selection rules*.

Step 7. Get token or x509 user certificate

JWT

x509

To submit jobs or perform any other action towards the ARC-CE you must authenticate yourself. We will do this using a token issued from the test-jwt issuer.

To generate a token do:

```
[user ~]$ export BEARER_TOKEN=$(arcctl test-jwt token)
```

To submit jobs or perform any other action towards the ARC-CE you must authenticate yourself. We will do this with a *proxy-certificate* which is a Single Sign-On token for distributed grid-infrastructure.

To generate a proxy certificate do:

```
[user ~]$ arcproxy
Your identity: /DC=org/DC=nordugrid/DC=ARC/O=TestCA/CN=Test User 50350053
Proxy generation succeeded
Your proxy is valid until: 2023-06-03 01:10:38
```

Step 8. Restart A-REX

On the ARC-CE, restart A-REX services to activate the configuration changes

```
[root ~]# arcctl service restart -a
```

Step 9. Check all is ok

You can run the client commands (arcinfo, arcsub etc) from the host running A-REX or from any other machine by installing the ARC client in steps 4 and 5.

You can start with the information query about your newly installed ARC computing element:

```
[user ~]$ arcinfo -C https://arc.example.org/arex
Computing service:
  Information endpoint: https://arc.example.org:443/arex
  Submission endpoint: https://arc.example.org:443/arex (status: ok, interface: org.
  ↪nordugrid.arcrest)
```

This means that all is ok, and the ARC client got back information from the ARC-CE that the information and service endpoints are available and ok.

Note

The examples use `arc.example.org` as a domain name for A-REX host. Step 0. Prerequisites for more information.

Tip: You can use `$(hostname)` instead of typing the hostname for these tests in your zero-conf setup and have a local client. For example:

```
arcinfo -C $(hostname)
```

Warning

It can take some minutes after the setup for everything to be fine, so if you see `status: critical` wait ca 1 minute and check again.

Step 10. Submit a job and check that it is running

A simple job can be submitted with the `arctest` tool:

```
[user ~]$ arctest -J 2 -C https://arc.example.org/arex
Job submitted with jobid: https://arc.example.org:443/arex/rest/1.0/jobs/f77b3d1b1efb
```

The job status can be checked with the `arcstat` tool:

```
[user ~]$ arcstat https://arc.example.org:443/arex/rest/1.0/jobs/f77b3d1b1efb
Job: https://arc.example.org:443/arex/rest/1.0/jobs/f77b3d1b1efb
Name: arctest2
State: Running

Status of 1 jobs was queried, 1 jobs returned information
```

To fetch the job's stdout run `arccat` tool:

```
[user ~]$ arccat https://arc.example.org:443/arex/rest/1.0/jobs/f77b3d1b1efb
HOSTNAME=arc.example.org
GRID_GLOBAL_JOBURL=https://arc.example.org:443/arex/f77b3d1b1efb
MALLOCARENA_MAX=2
PWD=/var/spool/arc/sessiondir/f77b3d1b1efb
SYSTEMD_EXEC_PID=374194
<output_omitted>
```

Step 11. Play more with the ARC Computing Element

As an admin you might frequently need to extract information from the logs and directories that ARC computing element uses. The brief list of the relevant paths can be obtained from:

```
[root ~]# arcctl config brief
ARC Storage Areas:
  Control directory:
    /var/spool/arc/jobstatus
  Session directories:
    /var/spool/arc/sessiondir
  Scratch directory on Worker Node:
    Not configured
  Additional user-defined RTE directories:
```

(continues on next page)

(continued from previous page)

```

    Not configured
ARC Log Files:
  A-REX Service log:
    /var/log/arc/arex.log
  A-REX Jobs log:
    /var/log/arc/arex-jobs.log
  A-REX Helpers log:
    /var/log/arc/job.helper.errors
  A-REX WS Interface log:
    /var/log/arc/ws-interface.log
  Infosys Infoproviders log:
    /var/log/arc/infoprovider.log

```

To get information and manage jobs on A-REX server, the `arcctl job` is useful. Operations include but are not limited to the following examples.

Listing jobs

```

[root ~]# arcctl job list
f5ab040cdc51
f617259d58ec
<output omitted>

```

```

[root ~]# arcctl job list --long
f5ab040cdc51      FINISHED   arctest2      https://wlcg.
↳cloud.cnaf.infn.it//b9f1e5e2-a8f0-4332-bd9d-58bd63898cc6
f617259d58ec      FINISHED   arctest2      https://wlcg.
↳cloud.cnaf.infn.it//b9f1e5e2-a8f0-4332-bd9d-58bd63898cc6
<output omitted>

```

Job general information

```

[root ~]# arcctl job info f77b3d1b1efb
Name           : arctest2
Owner          : https://wlcg.cloud.cnaf.infn.it//b9f1e5e2-a8f0-4332-bd9d-58bd63898cc6
State          : FINISHED
LRMS ID        : 376176
Modified       : 2023-06-02 16:07:05

```

Job log

```

[root ~]# arcctl job log f77b3d1b1efb
2023-06-02T14:06:51Z Job state change UNDEFINED -> ACCEPTED   Reason: (Re)Accepting_
↳new job
2023-06-02T14:06:51Z Job state change ACCEPTED -> PREPARING   Reason: Starting job_
↳processing
2023-06-02T14:06:51Z Job state change PREPARING -> SUBMIT      Reason: Pre-staging_
↳finished, passing job to LRMS
----- exiting submit_fork_job -----
2023-06-02T14:06:53Z Job state change SUBMIT -> INLRMS        Reason: Job is passed to LRMS

```

(continues on next page)

(continued from previous page)

```

----- Output of the job wrapper script -----
Detecting resource accounting method available for the job.
Looking for /usr/bin/time tool for accounting measurements
GNU time found and will be used for job accounting.
----- End of output -----
2023-06-02T14:07:05Z Job state change INLRMS -> FINISHING Reason: Job finished
↳executing in LRMS
2023-06-02T14:07:05Z Job state change FINISHING -> FINISHED Reason: Stage-out
↳finished.

```

A-REX logs that mentions the job

```

[root ~]# arcctl job log f77b3d1b1efb --service
### /var/log/arc/arex.log:
[2023-06-02 16:06:51] [Arc] [INFO] [374270/3] f77b3d1b1efb: State: ACCEPTED: parsing
↳job description
[2023-06-02 16:06:51] [Arc] [INFO] [374270/3] f77b3d1b1efb: State: ACCEPTED: moving
↳to PREPARING
[2023-06-02 16:06:51] [Arc] [INFO] [374270/3] f77b3d1b1efb: State: PREPARING from
↳ACCEPTED
[2023-06-02 16:06:51] [Arc] [INFO] [374270/3] f77b3d1b1efb: State: SUBMIT from
↳PREPARING
[2023-06-02 16:06:51] [Arc] [INFO] [374270/3] f77b3d1b1efb: state SUBMIT: starting
↳child: /usr/share/arc/submit-fork-job
[2023-06-02 16:06:53] [Arc] [INFO] [374270/3] f77b3d1b1efb: state SUBMIT: child
↳exited with code 0
[2023-06-02 16:06:53] [Arc] [INFO] [374270/3] f77b3d1b1efb: State: INLRMS from SUBMIT
[2023-06-02 16:07:05] [Arc] [INFO] [374270/3] f77b3d1b1efb: Job finished
[2023-06-02 16:07:05] [Arc] [INFO] [374270/3] f77b3d1b1efb: State: FINISHING from
↳INLRMS
[2023-06-02 16:07:05] [Arc] [INFO] [374270/3] f77b3d1b1efb: State: FINISHED from
↳FINISHING
### /var/log/arc/ws-interface.log:

```

Getting job attributes

```

[root ~]# arcctl job attr f77b3d1b1efb jobname
arctest2

```

Get production ready

Now you are ready to *Install production ARC7 Computing Element!*

6.2.2 ARC Computing Element Installation and Configuration Guide

Prerequisites

Choosing the host

It is assumed that ARC CE is installed on top of an existing Linux computing cluster. Many Linux distributions are supported. ARC works well also on a complete virtual computing cluster environment in a cloud.

ARC is non-intrusive towards existing systems. We suggest to deploy ARC CE on a dedicated (virtual) machine connected to the cluster network and filesystem.

ARC software is very lightweight and does not require powerful machines to run, however if ARC CE will perform data transfers the requirements are higher. As a minimum, a production CE with 4 cores and 8GB of RAM should be capable of handling up to 10,000 concurrent jobs without problems. One CE can easily handle the load of a single cluster, however multiple CEs may be deployed in parallel for redundancy.

Plan for storage areas

Several storage areas are necessary for job submission, execution and data storing. You should mount/export following directories:

- *session directory*
- *data staging cache directory* (if planned)
- decide to what extent to use NOT cross-mounted *scratch directory* on the worker nodes

Session directory (and the cache directory if used) is typically cross-mounted NFS share. Please note, that in the typical setup when A-REX is running as root NFS share need to be exported with `no_root_squash`.

Local resource management system (LRMS)

Install and configure your LRMS (batch system). ARC supports a variety of LRMS back-ends:

- fork - fork jobs on the ARC CE host node, not a cluster. Targeted for testing and development but not for real production workloads.
- condor - uses HTCondor-powered HTC resource
- slurm - for SLURM clusters
- pbs - any flavor of PBS batch system, including Torque and PBSPro
- pbspro - dedicated Altair PBS Professional backend (from 6.1 release)
- ll - Load Leveler batch system
- lsf - Load Sharing Facility batch system
- sge - Oragel Grid Engine (formerly Sun Grid Engine)
- boinc - works as a gateway to BOINC volunteer computing resources

Start by checking if you are able to submit jobs to the chosen LRMS from the ARC CE host.

You may consider setting up dedicated queues to use with ARC CE (e.g. per-VO queues).

Please also NOTICE that in some cases (*depending on LRMS*) you need to share the batch system log directories with ARC CE.

Configure OS accounts

Plan for local account(s) (or account pools) that will be used to execute jobs on the worker nodes.

These accounts should be also available on the ARC CE node.

Please note that ARC services are ran as root on the ARC CE node and switch to an appropriate local account when processing job data staging and job execution. This process is called *mapping*.

Installation

This section assumes you have already enabled the *Nordugrid repositories* for your package utility (dnf/apt).

Note

If you are using RHEL-based operating systems, ARC can be directly installed from the [EPEL repository](#). Please note that in EPEL the packages are named *nordugrid-arc7*.

Warning

If you are on an EL9 type server (CentOS-Stream 9, AlmaLinux 9, Rocky 9, RHEL9) you need to allow legacy crypto policies to be compatible with IGTF. On the command line of the ARC-CE server, issue:

```
update-crypto-policies --set LEGACY
```

Install ARC CE core packages from repositories as root:

```
dnf -y install nordugrid-arc-arex
```

or for Debian based systems as root:

```
apt-get install nordugrid-arc-arex
```

Any extra packages will be installed based on the ARC configuration file with *ARC Control Tool* as described [below](#). Full list of packages to install manually (especially additional plugins) can be found [here](#).

Grid security heavily relies on PKI and all actions requires certificates/keys for ARC CE as a service and users:

- for testing purposes, the ARC Test-CA and host certificate signed by the Test-CA are generated during A-REX installation.
- for production use please obtain a certificate signed by one of the [IGTF accredited CAs](#) and remove Test-CA files with `arcctl test-ca cleanup`.

In production ARC CE needs IGTF CA certificates deployed to authenticate users and other services, such as storage elements. To deploy IGTF CA certificates to ARC CE host, run as root¹:

```
arcctl deploy igtf-ca classic
```

¹ Use `--installrepo` argument to enable repositories with IGTF CA certificates if ARC is not installed from the Nordugrid repos.

Configuration

Configuration of ARC CE can be done by means of modifying the pre-shipped *zero configuration* available at `/etc/arc.conf`.

The purpose of this *zero configuration* is to offer a minimalistic working computing element out-of-the box right after package installation with zero additional configuration needed.

For production deployment you will need to customize the configuration in accordance to your actual setup and operations mode.

Note

ARC services must be restarted when changes have been made to `arc.conf`.

The ultimate information about available configuration options can be found in the [ARC Configuration Reference Document](#) which is also available locally as `/usr/share/doc/nordugrid-arc-*/arc.conf.reference`.

The most common configuration steps are explained below.

Configure authorization and mapping rules

Authorization rules define who can access the computing element (execute jobs, query info, etc). *Mapping rules* define which grid-users are mapped to which system accounts.

Both authorization and mapping rules in ARC rely on the concept of **authgroups**. Each authgroup represents a set of users, whose identities are matched to configured rules.

Once defined, authgroups can be applied to filter access to the CE per interface (*[arex/ws/jobs]* and/or per *queue*.

The `allowaccess` and/or `denyaccess` options in the corresponding block define which authgroups are allowed to access the interface or submit to the queue.

The *[mapping] block* used to configure the rules that defines how the particular authgroup members are mapped to OS accounts.

In the shipped *zero configuration* the `[authgroup: zero]` is defined and applied to A-REX WS interface, the effect of which is to deny any access unless user is listed in the `testCA.allowed-subjects` file. The mapping is configured with the `map_to_user` rule that assigns the same `nobody` account to everyone in the `zero` authgroup.

The typical configuration looks like this:

```
[authgroup: atlas]
voms = atlas * * *

[mapping]
map_to_pool = atlas /etc/grid-security/pool/atlas

[arex/ws/jobs]
allowaccess = atlas

[queue: qatlas]
allowaccess = atlas
```

Please read the *Authorization, Mapping and Queue selection rules* document to get familiar with all aspects of this important configuration step.

Provide LRMS-specific information

One more critical configuration step is to supply the ARC CE with relevant information regarding your LRMS specifics.

Specify your LRMS type

In the `arc.conf` there is a dedicated `[lrms]` block that defines the type of your LRMS, as well as several options related to the LRMS behaviour. For example, to instruct ARC to use SLURM, use the following configuration:

```
[lrms]
lrms = slurm
slurm_use_sacct = yes
```

Specify queues

In addition to specifying LRMS itself, it is necessary to list all the queues that will be exposed via the ARC CE, by using `[queue: name]` blocks.

```
[queue: atlas]
comment = Queue for ATLAS jobs
```

More information about configuring particular LRMS to work with ARC can be found in *Batch systems support* document.

Configure A-REX Subsystems

The ARC Resource-coupled EXecution service (A-REX) is a core service handling execution and entire life cycle of compute jobs.

Enable job management interfaces

In ARC 7 A-REX has only 1 main job management interface available, namely the ARC REST interface. To enable it, the block must be configured:

WS Interface ARC REST

`[arex/ws/jobs]`

In addition there is an experimental internal job submission interface. To use it, you must install the `nordugrid-arc-plugins-internal` package.

Enable data services

ARC has a built-in data transfer framework called DTR. It was designed to be used in environments in which data transfer was not possible or not desirable on the worker nodes such as HPC centres or sites without local storage.

DTR relies on users submitting jobs with pre-defined input and output files. When A-REX receives a job, it takes care of downloading the specified input files to the job's session directory, then submits the job to the batch system. After the batch job finishes, A-REX takes care of uploading any output files to grid storage.

Define the `[arex/data-staging]` block to enable data-staging capabilities. Data transfers can be scaled out using *multi-host data-staging*.

DTR also includes a cacheing capability. If cacheing is enabled then A-REX will download all input files to the cache, and create symlinks from the session directory for each file. If a job requests a file that is already cached,

A-REX will not download it again, but simply link from the existing cache file. Define the `[arex/cache]` block to enable cacheing.

More detailed technical documentation on ARC data features and advanced features such as CandyPond can be found in the [data overview pages](#).

RunTime Environments

RunTime Environments can modify the job execution cycle and are used for advertising software or features offered by the computing facility.

ARC ships several RTEs that are ready to be used and classified as system-defined.

One can add ones own directories with so-called user-defined RTEs using the `runtimeidir` configuration option in the `[arex]` block.

In ARC, both system- and user-defined directories are local to the ARC CE node and *SHOULD NOT* be shared to worker nodes.

To use an installed RTE, one should additionally **enable** this RTE with *ARC Control Tool*. For example, to enable the system-defined ENV/PROXY RTE, run:

```
[root ~]# arcctl rte enable ENV/PROXY
```

More details on operating RunTime Environments can be found in *RunTime Environments in ARC*.

Information system

ARC CE information system aims to collect and publish information to be used by special clients for matchmaking and/or monitoring the state and stats of the resource.

It is mandatory to configure the information system for production cases, like those of the WLCG computing infrastructure.

Defining general information

There are many information schemas and renderings of data available to comply to existing standards. There are several blocks that are used to define published information depending on schemas:

[infosys]

The most common block that enables internal information collection from ARC CE host and LRMS.

[infosys/cluster]

The common information about the whole cluster, including e.g. its capacity.

[queue: name]

For heterogeneous clusters, most of the information in the `[infosys/cluster]` block can be re-defined on per-queue basis.

[infosys/glue2]

Configures the GLUE2-specific values and enables internal glue2 rendering.

[infosys/ldap]

Enables LDAP/BDII dedicated services to publish information via LDAP protocol.

[infosys/glue2/ldap]

Enables GLUE2-schema LDAP rendering of the information.

[infosys/nordugrid]

Enables LDAP rendering of the information according to the NorduGrid schema.

Accounting

ARC CE has built-in functionality to measure job's resource usage metrics that can be used for analyses and publishing to the SGAS and APEL centralized accounting services.

If you need to configure accounting follow the [accounting guide](#).

Configure Firewall

Different ARC CE services open a set of ports that should be allowed in the firewall configuration.

To generate iptables configuration based on `arc.conf`, run as root:

```
arcctl deploy iptables-config
```

Enable and Run Services

To enable and run all services as configured in `arc.conf`, run as root:

```
arcctl service enable --as-configured --now
```

Instead of using *ARC Control Tool* to manage *ARC services*, you can always use your OS native tools.

Test Basic Functionality

To test some basic job submission to the configured ARC CE, follow the instructions provided in the `:doc:quickstart` guide.

6.2.3 ARC6 to ARC7 Migration Guide

The upgrade from ARC 6 to ARC 7 is just a matter of removing some deprecated packages (Step 1), enabling the relevant repository and updating the packages as per ordinary update procedures for the system you are on (Step 2), and in some cases you might want to migrate records in your accounting database (Step 3). Finally, check that the services have started properly and clean up your `arc.conf` if not (Step 4).

You will notice that the control directory will be restructured containing a new *jobs* directory where all the ARC job metadata files are stored. Your configuration should work without any changes for ARC 7, but please pay extra attention to any warnings or errors when you start ARC up after upgrading, and fix accordingly. Some configuration warnings in ARC 6 are now errors in ARC 7 and need to be attended.

ARC 7 is REST only, except for the information system which still has ldap possibility. If you rely on custom scripts using the legacy GridFTP or EMI-ES endpoints (jobs or info), they will have to be rewritten for REST.

If you are using an ARC 6 client to submit to an ARC 7 site, be sure to download the `nordugrid-arc-plugins-arcrest` package. For the ARC 7 package, this is already included.

Step 1. Remove deprecated packages

To be able to update from ARC 6 to ARC 7 you will have to remove the deprecated packages first (showing rhel flavour example):

```
arcctl service stop -a
dnf erase nordugrid-arc-plugins-gridftpjob nordugrid-arc-gridftpd nordugrid-release
```

Step 2. Update ARC

You can update either from the nordugrid repositories or from the EPEL or Debian repositories

From the nordugrid repositories

Follow these instructions to enable the nordugrid ARC 7 repo: [NorduGrid repository information](#). Next update all nordugrid packages with your usual procedure.

As an example, this is how to do the above on a rocky 9 system:

```
rpm --import http://download.nordugrid.org/RPM-GPG-KEY-nordugrid-7
dnf install https://download.nordugrid.org/packages/nordugrid-release/releases/7/
↪ rocky/9/x86_64/nordugrid-release-7-1.el9.noarch.rpm
dnf update --enablerepo=nordugrid "*nordugrid*"
```

Or from EPEL or Debian repositories

Follow instructions in [NorduGrid repository information](#).

Note that packages names in the official EPEL or Debian repositories are called `nordugrid-arc7-xxxx` instead of `nordugrid-arc-xxxx`.

Step 3. Migrate accounting from v1 to v2

ARC 7.0.0 comes with a new record format, and stores new files in this format automatically. The new record filename is: `<controldir>/accounting/accounting_v2.db`.

If you have some jobs in a non-terminal states during the update, those jobs will be missing in the new database and you will see accounting ERRORS in your ARC logs. You can migrate these jobs if you need to, to fix this.

If you need accounting records for old jobs to be updated or you are publishing the old accounting data, you need to migrate those recent ARC 6 records from old database to new database. This is done with the `arcctl accounting database migrate` [\(<https://www.nordugrid.org/arc/arc7/admins/commands/arcctl.html#migrate> command](https://www.nordugrid.org/arc/arc7/admins/commands/arcctl.html#migrate).

For clarity, an example is shown here, where we are migrating records from month of March and our controldirectory is `/var/spool/arc/jobstatus`:

```
arcctl accounting database migrate -b 2025-03-01 -e 2025-03-31 -f /var/spool/arc/
↪ jobstatus/accounting/accounting.db
```

Be careful not to migrate too large chunks at a time. One month is a good limit to go with.

Step 4. Post-install

On RHEL-flavour systems: Check your `arc.conf` and if it has been moved to `arc.conf.rpmsave` just copy it back.

```
arcctl service restart -a
```

Check for validation errors upon ARC startup - fix any errors shown, for instance, remove deprecated blocks.

Once validation passes - you are good to go!

Common errors

Outdated RTE's

Many sites have custom RTE's - if those contain paths to files in the controldir the RTE will fail if the path is one that has changed with the ARC 7 control dir restructuring.

One example is *ENV/PROXY* where we have encountered at least two sites with user-RTE's with the ARC 6 controldir structure assumed.

The error these sites see is in the jobs errors file - something like this - pointing to *job.<jobid>.proxy* - this is not the correct path.

```
` [2025-05-12 07:27:28] [INFO] [173/6] DTR 6030...50e5: a2de03bd1f29: All downloads
finished successfully 2025-05-12T07:27:28Z Job state change PREPARING -> SUBMIT
Reason: Pre-staging finished, passing job to LRMS ----- starting submit_slurm_job
----- cat: /var/spool/arc/jobstatus/job.a2de03bd1f29.proxy: No such file or
directory ERROR: Runtime script ENV/PROXY stage 0 execution failed with exit code
1 2025-05-12T07:27:28Z Job state change SUBMIT -> FINISHING Reason: Job failure
detected 2025-05-12T07:27:28Z Job state change FINISHING -> FINISHED Reason: Job
failure detected `
```

How to fix

1. First check your RTE's with

```
`arcctl rte list --long`
```

Check the path of your user-installed *ENV/PROXY*

2. Disable the user-installed *ENV/PROXY*

```
` arcctl disable ENV/PROXY `
```

3. Then remove the file e.g. if your runtime folder is configured to be *<controldir>/runtime*:

```
` rm /var/spool/arc/jobstatus/runtime/ENV/PROXY `
```

4. Enable and default (if all your jobs should use this RTE by default) the system *ENV/PROX*

```
` arcctl enable ENV/PROXY arcctl default ENV/PROXY `
```

6.2.4 ARC CE Deployment Scenarios

WLCG Deployment with Data Capabilities

All steps below are performed as the *root* user unless explicitly stated.

Prerequisites

First follow the installation and configuration guide: [ARC installation guide](#).

Install necessary packages

CA certificates

Install necessary CA certificates needed for WLCG (you may have already installed the `igtcf-ca` classic from the prerequisite step)

```
arcctl deploy igtcf-ca classic
```

Install and set up `fetch-crl`

To keep your CA's revocation lists up-to-date you need to install and set up `fetch-crl` tool. Required for a WLCG site.

Install the package:

```
dnf install fetch-crl
```

Enable and start `fetch-crl`:

```
systemctl start fetch-crl
```

Install `xrootd` plugin for data fetching

Your jobs may require to access data through the `xrootd` protocol. If so, install the plugin:

```
dnf install -y nordugrid-arc-plugins-xrootd
```

Prepare for voms service signature authentication

If your `arc.conf` contains authgroups using voms like

```
[authgroup:atlas-jobs]
voms = atlas * * * *

[authgroup:ops]
voms=ops * * *
voms=dteam * * *
```

Then you must install the corresponding voms directory. Here we show the installation for `rhel9` flavour CE, and for a site running jobs for the ATLAS experiment. Replace `atlas` with the experiment you serve, for instance `cms` or `lhcb` if needed. In addition, a WLCG site will need the `lcg` and `dteam` vomsdirs:

```
dnf install -y https://linuxsoft.cern.ch/wlcg/el9/x86_64/wlcg-iam-lsc-atlas-2.0.0-1.
↪el9.noarch.rpm
dnf install -y https://linuxsoft.cern.ch/wlcg/el9/x86_64/wlcg-voms-atlas-2.0.0-1.el9.
↪noarch.rpm
dnf install -y https://linuxsoft.cern.ch/wlcg/el9/x86_64/wlcg-iam-lsc-dteam-1.0.0-1.
↪el9.noarch.rpm
dnf install -y https://linuxsoft.cern.ch/wlcg/el9/x86_64/wlcg-iam-lsc-ops-2.0.0-1.el9.
↪noarch.rpm
```

To inspect what the installation does, you can issue the `rpm -ql` command on each rpm, which will show you what files were installed in what folders in your `/etc/grid-security/vomsdir` folder:

```
rpm -ql https://linuxsoft.cern.ch/wlcg/el9/x86_64/wlcg-iam-lsc-atlas-2.0.0-1.el9.
↳noarch.rpm
/etc/grid-security/vomsdir/atlas/voms-atlas-auth.app.cern.ch.lsc
/etc/grid-security/vomsdir/atlas/voms-atlas-auth.cern.ch.lsc
```

Up-to-date information about the VOMS services can be found [here](#).

Configure ARC datastaging and cache

The datastaging service is enabled by adding the `[arex/data-staging]` block to `arc.conf`

A minimal configuration with a central datastaging logfile enabled could look like:

```
[arex/data-staging]
logfile=/var/log/arc/datastaging.log
```

See `[arex/data-staging]` for other options. Especially the `preferredpattern` and `deliveryservice` in case you have one or more remote delivery service machine(s) set up.

You should also enable the cache and cache-cleaner by adding the blocks `[arex/cache]` and `[arex/cache/cleaner]`. An example configuration is shown below:

Example configuration:

```
[arex/cache]
cachedir=/grid/cache01
cachedir=/grid/cache02

[arex/cache/cleaner]
cachesize=90 80
cachelifetime=50d
calculatesize=filesystem
```

Runtime environments

ENV/PROXY

This runtimeenvironment will ensure the copying of the users proxy file from the ARC-CE to the compute node. In addition to other optional certificate/token related actions, see [ENV/PROXY](#) for more details.

```
[root ~]# arcctl rte enable ENV/PROXY
```

ATLAS RTE

If you are an ATLAS site, ATLAS requires a dummy ATLAS-SITE RTE. Create one and enable it like this:

```
[root ~]# arcctl rte enable --dummy APPS/HEP/ATLAS-SITE
```

ARC-CE only accepting non-igtf tokens

ARC Deployment beyond the WLCG scope

This scenario is for ARC-CE sites that only accept client requests using tokens issued by any of the CAs included in your systems CA bundle, or any other non-igtf CA. In other words: non-igtf tokens and no x509 user certificates. This could be relevant for distributed computing sites set up for user communities other than WLCG. In this scenario you would not do the `arcctl deploy igtf-ca` step which is included in the typical setup of a production ready ARC-CE.

Currently this requires configuration changes both on the ARC-CE and on the ARC client.

Note

The ARC-CE can not be configured to use *both* system CA's and igtf CA's, as openssl only looks in *one* location for authentication. If you also need an ARC-CE that accepts the traditional igtf certificates you must set up a separate ARC-CE for this.

ARC-CE config settings

The host certificate should be a plain ssl host certificate (not grid certificate), and installed in the default cert location which is e.g. `/etc/pki/tls/certs` for the certificate and e.g. `/etc/pki/tls/private` for the key for a rocky 9 server. The paths will depend on your OS.

To tell your ARC-CE to use system certificates instead of igtf, add the following in the `[common]` block in `arc.conf`

```
[common]
x509_cert_policy = system
```

ARC client settings

When interacting with an ARC-CE which is set up in system mode, you need to use the `--systemca` option. Example when submitting a job with `arcsub`:

```
arcsub --systemca -C arc.example.org myjob.xrls
```

How to install new CA

To accept a token from a CA not included in the default system CA's the admin must manually add the root CA's to the OS's default CA directory.

You can use the `arcctl` tool to accomplish this, as in this example case letsencrypt isrgrootx1:

```
arcctl deploy ca-cert --crt https://letsencrypt.org/certs/isrgrootx1.pem
```

`arcctl` takes the certificate, installs it to `x509_cert_dir` (configured in `arc.conf`), makes necessary hash-links, creates the signing policy file.

In the same fashion you will need to install root-CA's of any other token issuers you want to trust, if they are not already included in your CA bundle.

How to install an ssl host certificate

In the scenario show-cased here it would be natural that your ARC-CE has a normal ssl host certificate instead of the traditional igtf one. The example below shows how to install this certificate correctly for ARC. You will also need to install the corresponding root CA of this certificate as per instructions above, if not already included in your CA bundle.

Order an ordinary ssl host certificate (i.e. not a grid one). Once obtained, create a hostcert cocatenated with the intermediate cert (if no bundle is provided). We are assuming the default locations `/etc/pki/tls/certs` and `/etc/pki/tls/private` in the following example.

```
cat example.com.crt > hostcert.pem
cat intermediate.crt >> hostcert.pem
sudo cp hostcert.pem /etc/pki/tls/certs/
sudo cp hostkey.pem /etc/pki/tls/private/
```

Create hash link of hostcert:

```
cd /etc/pki/tls/certs
sudo ln -s hostcert.pem $(openssl x509 -hash -noout -in hostcert.pem).0
```

Verify the host certificate (assumes the root CA is installed in previous step):

```
openssl verify /etc/pki/tls/certs/hostcert.pem
/etc/pki/tls/certs/hostcert.pem: OK
```

You should now be all set to use your ordinary ssl host certificate for the ARC-CE.

Map to local user with token

In the WLCG it is common to map to a system user, for instance a pool of users for a certain VO membership. In non-WLCG use-cases it may be appropriate to instead map to the actual username of the ARC submitter on that particular compute site.

This can be achieved using tokens, when the condition is met that the site has an associated IAM and the user can fetch a token that includes the username claim, or some other claim that holds this username.

Here we will show an example of how such mapping can be achieved using the `map_with_plugin` option in the mapping block.

Step 1. Create auth group

In your `arc.conf`, create an authblock for tokens submitted by the token-issuer, in this example for the University of Oslo (uio).

```
[authgroup:uio-hpc-users]
authtokens = * https://uio-hpc-iam.uio.no/ * * *
```

Step 2. Map the authgroup

In your `arc.conf`, map the authgroup with your custom plugin

```
[mapping]
map_with_plugin = uio-hpc-users 10 /opt/arc/scripts/token-map.sh
```

Step 3. Create an executable script that returns the username

The `map_with_plugin` option has the `BEARER_TOKEN` env vars available as documented in the [map_with_plugin](#) option. In this example our token contains the claim `user` which holds the real username on the ARC compute site, which is what we are after. As usual, the username must already exist on the ARC-CE and on the compute nodes.

For our example we create a script `/opt/arc/scripts/token-map.sh` as configured in `arc.conf` that contains:

```
#!/bin/bash
# map by userid from token
echo "${BEARER_TOKEN_0_CLAIM_user_0}"
```

The script with output the username in this case, which is then used in the mapping as instructed in the `arc.conf` in Step 2.

JWT-only Infrastructure for small community based on ARC

This document describes the example setup of the infrastructure for a small research community based on ARC support for OAuth2 JWT tokens.

Intro: JWT tokens, IdPs and ARC Test-JWT

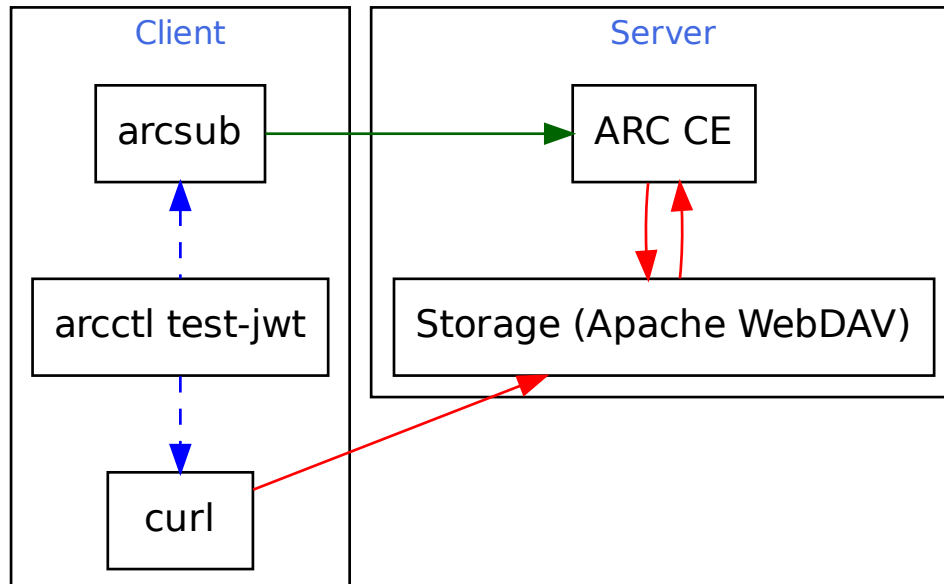
JWT Tokens are signed by the identity provider (IdP). Production infrastructure usually relies on the dedicated IdP hosted and managed by community. One of the common options is the [Keycloak](#).

In this example, we will use the *test-jwt* subsystem of the *ARC Control Tool* instead of the centrally managed IdP.

The *Test JWT* can be seen as a **private IdP** on the local client computer. Each client in the small community has its own private IdP.

In the x509 world it can be compared to the *own CA* for every client, trusted by servers. Every generated JWT token can be compared to a short-lived client certificate signed with such dedicated CA.

Prelude: Infrastructure actors



Client: user on the machine (laptop, VM, etc) with the ARC Client and *ARC Control Tool* installed. Each client has its own Test-JWT Issuer trusted by both the ARC CE and Storage.

ARC CE: tokens-only REST-only ARC 7 installation. Trusts the client(s) Test-JWTs.

Storage: the storage element providing WebDAV access to files with OAuth2 authentication. Trusts the client(s) Test-JWTs. This tutorial provides the example configuration for Apache.

For sake of simplicity we will install Apache and ARC CE on the same host.

Both client and server host needs to enable the *Nordugrid Repositories* before moving forward.

Verse: Test-JWT IdP

Setting up the *private IdP(s)* with ARC Test-JWT is very straight-forward.

It is private to the user, so it should be run as a normal user account by design. All Test-JWT files (keys) resides in the user \$HOME directory.

Install the client packages:

```
[root@token-clinet ~]# dnf install nordugrid-arc-client nordugrid-arc-arcctl
```

Init the Test-JWT Issuer¹:

```
[andrii@token-client ~]$ arcctl test-jwt init
[2024-10-23 09:19:59,164] [ARCCTL.JWTIssuer] [INFO] [502] [Showing the JWKS for JWT_
->Issuer https://token-client2/arc/testjwt/169d01c/1001]
{
```

(continues on next page)

¹ Here and in the rest of the document the long strings in the output, like keys, are truncated for readability. A real run uses longer keys.

(continued from previous page)

```
"keys": [
  {
    "e": "AQAB",
    "kid": "testjwt",
    "kty": "RSA",
    "n": "8cPfW5Mla0vPWT7m6Mbhx54lNid65BFYS02uwuXF...",
    "use": "sig"
  }
]
}
```

```
[2024-10-23 09:19:59,164] [ARCCTL.TestJWT] [INFO] [502] [Generating deployment_
↪command to be executed on ARC CE to trust the Test JWT issuer https://token-client2/
↪arc/testjwt/169d01c/1001]
arcctl deploy jwt-issuer --deploy-conf test-jwt://H4sIAJ+jGGcC/6XSXY+iMBQG4P/
↪C9RKggulc8SGojGYQUIfNhiAUrEBBKAY//
↪vUdSd7PfGq5030edKc9sYUkIRxSElm7cagpmlhzbwxR0Kq5o3jSj1BzEY5gpgALqwjj...
```

Note

Notice that `test-jwt init` prints the command to run on ARC CE. You can output it again with `arcctl test-jwt export`.

For the Apache configuration the JWK key will be needed as an escaped string. You can automate typing backslashes with `jq`:

```
[andrii@token-client ~]$ arcctl test-jwt info | jq '.keys[] | tostring'
[2024-10-23 09:23:36,291] [ARCCTL.JWTIssuer] [INFO] [507] [Showing the JWKS for JWT_
↪Issuer https://token-client2/arc/testjwt/169d01c/1001]
"{\"e\":\"AQAB\", \"kid\":\"testjwt\", \"kty\":\"RSA\", \"n\":\
↪\"8cPfW5Mla0vPWT7m6Mbhx54lNid65BFYS02uwuXFazHCDNLe1f1N...\", \"use\":\"sig\"}"
```

Before using the client, the new short-lived (12 hours by defaults) JWT access token needs to be generated. The ARC client expects it in the `BEARER_TOKEN` variable, so the shortcut is:

```
[andrii@token-client ~]$ export BEARER_TOKEN=$(arcctl test-jwt token)
```

Note

HINT: The Claims in the JWT token are fully customizable. For example, to include the meaningful username you can run `arcctl test-jwt token -n "UiO aCT"`. Find more info in [Test JWT Howto using ARC Control Tool](#).

Bridge: ARC CE

The ARC CE installation is as easy as:

```
[root@arc-jwt ~]# dnf -y install nordugrid-arc-arex
```

During the package installation the zero-conf (for testing only) is installed and the Test-CA hostkeys are generated and available at `/etc/grid-security/testCA-hostkey.pem` and `/etc/grid-security/testCA-hostcert.pem`.

To make the ARC CE trust the client's private IdP, copy the `arcctl deploy` command printed out by Test-JWT on client host, run it on the ARC CE and then restart the services:

```
[root@arc-jwt ~]# arcctl deploy jwt-issuer --deploy-conf test-jwt://H4sIAJ+jGGcC/
↪6XSXY+iMBQG4P/C9RKggulc8SGojGYQUIfNhiAUrEBBKAY....
[2024-10-23 09:20:09,364] [ARCCTL.ThirdParty.Deploy] [INFO] [41083] [ARC CE now trust
↪JWT signatures of https://token-client2/arc/testjwt/169d01c/1001 issuer.]
[2024-10-23 09:20:09,364] [ARCCTL.JWTIssuer] [INFO] [41083] [Auth configuration for
↪JWT issuer https://token-client2/arc/testjwt/169d01c/1001 has been written to /etc/
↪arc.conf.d/10-testjwt-62641c8b.conf]
[2024-10-23 09:20:09,364] [ARCCTL.JWTIssuer] [WARNING] [41083] [ARC services restart
↪is needed to apply configuration changes.]

[root@arc-jwt ~]# arcctl service restart -a
```

If there are multiple clients - run the deploy command for every client used in the infrastructure.

All users are mapped to nobody in zero-conf. For a production ARC CE this requires adjustments via editing the config (/etc/arc.conf.d/10-testjwt-62641c8b.conf in this example). See [Configure authorization and mapping rules](#) for examples of production authgroups and mapping.

Note

As we have Test-CA hostkeys, the client need to trust the ARC CE Test CA. To do it:

1. Run `arcctl test-ca info -o ca-cert` on the ARC CE and copy the CA certificate
2. Run `arcctl deploy ca-cert` on the client and paste the certificate

Note

The JWT-only case needs certificates only for hosts, like any regular web-service. As users are excluded, a small community can easily rely on e.g. [LetsEncrypt](#) to manage certificates in production. Use [system CA bundle](#) in this case.

Chorus: Apache + WebDAV + mod_oauth2 = Storage Element

For an Apache based WebDAV storage elemnt with OAuth2, first we need to install the necessary packages.

Online the OIDC module, `mod_oauth2` is not included in EPEL. Package installation is done from GitHub releases.

```
[root@arc-jwt ~]# dnf -y install httpd mod_ssl
[root@arc-jwt ~]# dnf -y install https://github.com/OpenIDC/mod_oauth2/releases/
↪download/v4.0.0/mod_oauth2-4.0.0-1.el9.x86_64.rpm \
https://github.com/OpenIDC/liboauth2/releases/
↪download/v2.0.0/liboauth2-2.0.0-1.el9.x86_64.rpm \
https://github.com/OpenIDC/liboauth2/releases/
↪download/v2.0.0/liboauth2-apache-2.0.0-1.el9.x86_64.rpm
```

Start the Apache configuration with setting up TLS. For this example setup, we just link the ARC CE Test-CA generated hostcert/key to the location defined in default /etc/httpd/conf.d/ssl.conf:

```
[root@arc-jwt ~]# ln -s /etc/grid-security/testCA-hostkey.pem /etc/pki/tls/private/
↪localhost.key
[root@arc-jwt ~]# ln -s /etc/grid-security/testCA-hostcert.pem /etc/pki/tls/certs/
↪localhost.crt
```

As Apache is running on the same host as the ARC CE that also listens on 443, change the port to 8443 in /etc/httpd/conf.d/ssl.conf. It is advised to disable port 80 as well as a welcome page.

Once TLS is sorted out, WebDAV and OAuth2 configuration can be added like this:

```
[root@arc-jwt ~]# mkdir /srv/webdav
[root@arc-jwt ~]# echo "test" > /srv/webdav/test
[root@arc-jwt ~]# chown -R apache:apache /srv/webdav/

[root@arc-jwt ~]# cat <<EOF > /etc/httpd/conf.d/webdav.conf
Alias /webdav /srv/webdav

DavLockDB "/var/lib/httpd/DavLockDB"

<Directory /srv/webdav>
    DAV On
    AuthType oauth2
    OAuth2TokenVerify jwk "{\"e\":\"AQAB\",\"kid\":\"testjwt\",\"kty\":\"RSA\",\"n\":\"\\
→oy24GkRH3vs7DjGcp55wis2k-pET4...\", \"use\":\"sig\"}"
    OAuth2TokenVerify jwk "{\"e\":\"AQAB\",\"kid\":\"testjwt\",\"kty\":\"RSA\",\"n\":\"\\
→8cPfW5Mla0vPWT7m6Mbhx54lNid65...\", \"use\":\"sig\"}"
    <RequireAll>
        Require valid-user
    </RequireAll>
</Directory>
EOF

[root@arc-jwt ~]# systemctl restart httpd
```

The `OAuth2TokenVerify` contains the JWK key of the trusted Test-JWT Issuer. It can be specified multiple times, if there are multiple clients (multiple issuers).

You can verify that Apache WebDAV is working via `curl`:

```
[andrii@token-client ~]$ export BEARER_TOKEN=$(arcctl test-jwt token)
[andrii@token-client ~]$ curl -k -H "Authorization: bearer ${BEARER_TOKEN}" https://
→arc-jwt.local:8443/webdav/test
test
```

Coda: Submitting jobs with data staging

The example `xRSL` is inspired by the `arctest` job, but adding stage-in and stage-out to our WebDAV storage:

```
[andrii@token-client ~]$ cat << EOF > davtest.xrsl
&( executable = "/usr/bin/env" )
( jobname = "davtest" )
( stdout = "stdout" )
( join = "yes" )
( inputFiles = ("test" "https://arc-jwt.local:8443/webdav/test") )
( outputFiles = ("stdout" "https://arc-jwt.local:8443/webdav/test-output"
→"overwrite=yes"))
EOF
```

The ARC 7 client is delegating the JWT credentials for data-staging automatically, so the submission itself is just a generic `arcsub` call, nothing specific is needed.

```
[andrii@token-client ~]$ export BEARER_TOKEN=$(arcctl test-jwt token)

[andrii@token-client ~]$ arcsub -C https://arc-jwt.local/arex davtest.xrsl
Job submitted with jobid: https://arc-jwt.local:443/arex/rest/1.0/jobs/6ab04a5f9869
```

(continues on next page)

(continued from previous page)

```
[andrii@token-client ~]$ arcstat https://arc-jwt.local:443/arex/rest/1.0/jobs/
↳ 6ab04a5f9869
Job: https://arc-jwt.local:443/arex/rest/1.0/jobs/6ab04a5f9869
Name: davtest
State: Finished
Exit Code: 0
```

Once the job is finished, the ARC CE uploaded the output of env to WebDAV. We can get the file using curl:

```
[andrii@token-client ~]$ curl -k -H "Authorization: bearer ${BEARER_TOKEN}" https://
↳ arc-jwt.local:8443/webdav/test-output
HOSTNAME=arc-jwt.local
GRID_GLOBAL_JOBURL=https://arc-jwt.local:443/arex/6ab04a5f9869
...
```

Additionally, let's do some server-side checks related to job data transfers.

Starting with simple storage location check:

```
[root@arc-jwt ~]# ls -l /srv/webdav/
total 12
-rw-r--r-- 1 apache apache 944 Oct 23 23:29 test-output
-rw-r--r-- 1 apache apache 5 Oct 23 23:23 test
```

Apache logs, showing the WebDav requests (note username from JWT token in the logs):

```
[root@arc-jwt ~]# cat /var/log/httpd/ssl_access_log
fe80::d937:cdd1:ce2d:806e - Test User 81681464 [23/Oct/2024:23:29:20 +0200]
↳ "PROPFIND https://arc-jwt.local:8443/webdav/test HTTP/1.1" 207 816
fe80::d937:cdd1:ce2d:806e - Test User 81681464 [23/Oct/2024:23:29:21 +0200]
↳ "GET https://arc-jwt.local:8443/webdav/test HTTP/1.1" 200 5
fe80::d937:cdd1:ce2d:806e - Test User 81681464 [23/Oct/2024:23:29:34 +0200]
↳ "DELETE https://arc-jwt.local:8443/webdav/test-output HTTP/1.1" 404 196
fe80::d937:cdd1:ce2d:806e - Test User 81681464 [23/Oct/2024:23:29:34 +0200]
↳ "PUT https://arc-jwt.local:8443/webdav/test-output HTTP/1.1" 201 232
fe80::d937:cdd1:ce2d:806e - Test User 81681464 [23/Oct/2024:23:29:34 +0200]
↳ "PROPFIND https://arc-jwt.local:8443/webdav/test-output HTTP/1.1" 207 825
```

ARC CE job logs, showing data transfers (redacted to show only relevant information):

```
[root@arc-jwt ~]# arcctl job log https://arc-jwt.local:443/arex/rest/1.0/
↳ jobs/6ab04a5f9869
2024-10-23T21:29:20Z Job state change UNDEFINED -> ACCEPTED Reason:
↳ (Re)Accepting new job
2024-10-23T21:29:20Z Job state change ACCEPTED -> PREPARING Reason:
↳ Starting job processing
[2024-10-23 23:29:20] [INFO] [41237/5] DTR 7c07...ce0d: Scheduler received
↳ new DTR 7c076250-e5d6-4d04-8d45-9fe866dcce0d with source: https://arc-jwt.
↳ local:8443/webdav/test, destination: file:/var/spool/arc/sessiondir/
↳ 6ab04a5f9869/test, assigned to transfer share _default-download with
↳ priority 25
[2024-10-23 23:29:21] [INFO] [41237/4] DTR 7c07...ce0d: Transfer finished:
↳ 5 bytes transferred
[2024-10-23 23:29:21] [INFO] [41237/6] DTR 7c07...ce0d: 6ab04a5f9869: All
↳ downloads finished successfully
2024-10-23T21:29:21Z Job state change PREPARING -> SUBMIT Reason: Pre-
```

(continues on next page)

(continued from previous page)

```

→staging finished, passing job to LRMS
2024-10-23T21:29:22Z Job state change SUBMIT -> INLRMS Reason: Job is
→passed to LRMS
2024-10-23T21:29:33Z Job state change INLRMS -> FINISHING Reason: Job
→finished executing in LRMS
[2024-10-23 23:29:34] [INFO] [41237/5] DTR 93d8...b3bb: Scheduler received
→new DTR 93d85ed3-e121-4a97-9dd6-7a4f3d60b3bb with source: file:/var/spool/
→arc/sessiondir/6ab04a5f9869/stdout, destination: https://arc-jwt.
→local:8443/webdav/test-output, assigned to transfer share _default-upload
→with priority 25
[2024-10-23 23:29:34] [INFO] [41237/4] DTR 93d8...b3bb: Transfer finished:
→946 bytes transferred : checksum Adler32:49ad28a1
[2024-10-23 23:29:34] [INFO] [41237/6] DTR 93d8...b3bb: 6ab04a5f9869: All
→uploads finished successfully
2024-10-23T21:29:34Z Job state change FINISHING -> FINISHED Reason: Stage-
→out finished.

```

Accounting data related to data transfers:

```

[root@arc-jwt ~]# arcctl accounting job transfers https://arc-jwt.local:443/
→arex/6ab04a5f9869
Data transfers (downloads) performed during A-REX stage-in:
https://arc-jwt.local:8443/webdav/test:
Size: 5
Download timeframe: 2024-10-23 21:29:20 - 2024-10-23 21:29:21
Data transfers (uploads) performed during A-REX stage-out:
https://arc-jwt.local:8443/webdav/test-output:
Size: 946
Upload timeframe: 2024-10-23 21:29:33 - 2024-10-23 21:29:34

```

WLCG Deployment for pilot sites

6.2.5 Tuning tips for ARC and FAQ

Note

Work-in-progress!

CA certificates

The ENV/PROXY RTE by default makes a copy of CA certificates in each job's session directory. Copying many small files to a shared file system is potentially expensive so it is preferred to have CA certificates locally installed on each worker node, and turn off the copying in ENV/PROXY with `arcctl rte params-set ENV/PROXY COPY_CACERT_DIR No`.

The CA certificate packages can be found in the repositories for ARC (<http://www.nordugrid.org/documents/arc6/common/repos/repository.html>)

You should also install the `fetch-crl` package and enable the services as you would on the ARC-CE.

File-system tuning

If you have your session directories and/or cache directories on NFS make sure you have increased the number of nfs kernel threads. The default is 8 which is too low. The `RPCNFSDCOUNT` variable should be increased to 32 or maybe even 64 or even larger depending on your system (number of worker-nodes, number of cores on the nfs servers).

How to extract backtrace of core dump file

How to extract backtrace of ARC core dump file using gdb and save the output to a file.

```
gdb <path-to-arc-binary> <path-to-core-file>
(gdb) set logging file gdb.log
(gdb) set logging on
(gdb) thread apply all bt
```

Example: Assuming ARC is installed in default location: .. code-block:: console

```
gdb /usr/sbin/arched /var/log/arc/arccore/core.30547 (gdb) set logging file gdb.log (gdb) set logging
on (gdb) thread apply all bt
```

How to extract backtraces using a-rex-backtrace-collect

`a-rex-backtrace-collect` processes core file(s) collected in `ARC_LOGS_DIR/arccore` folder and produces their backtraces. The backtrace(s) are stored in files `<core_name>.backtrace`. The ARC installation location can be adjusted using `ARC_LOCATION` environment variable. The location of configuration file can be specified using `ARC_CONFIG` environment variable.

```
... code-block::console
a-rex-backtrace-collect
```

6.2.6 Operating ARC CE Subsystems

ARC Packages

Table 6.1: List of ARC binary packages

Package info	ARC package name	Block name Page 136, 1
Base package holding common files	<code>nordugrid-arc</code>	N/A
ARC Hosting Environment Daemon	<code>nordugrid-arc-hed</code>	N/A
ARC Resource-coupled EXecution service (A-REX)	<code>nordugrid-arc-arex</code>	[arex]
ARC Candypond Service	part of <code>nordugrid-arc-arex</code>	[arex/ws/candypond]
ARC LDAP-based Information Services ²	<code>nordugrid-arc-infosys-ldap</code>	[infosys/ldap]
ARC Data Delivery Service (DDS)	<code>nordugrid-arc-datadelivery-s</code>	[datadelivery-service]
ARC Control Admin tool for ARC-CE	<code>nordugrid-arc-arcctl-service</code>	N/A
ARC Control Admin tool for ARC client	<code>nordugrid-arc-arcctl</code>	N/A
ARC development files	<code>nordugrid-arc-devel</code>	N/A
Python 3 bindings for ARC	<code>python3-nordugrid-arc</code>	N/A
ARC command line clients	<code>nordugrid-arc-client</code>	own config
ARC test tools	<code>nordugrid-arc-test-utils</code>	N/A
ARC LDAP monitor web application	<code>nordugrid-arc-monitor</code>	N/A

continues on next page

Table 6.1 – continued from previous page

Package info	ARC package name	Block name ¹
ARC prometheus exporter	nordugrid-arc-arc-exporter	N/A
ARC base plugins (MCCs and DMCs)	nordugrid-arc-plugins-needed	N/A
ARC Globus plugins	nordugrid-arc-plugins-globus	N/A
ARC Globus common plugins	nordugrid-arc-plugins-globus	N/A
ARC xrootd plugins	nordugrid-arc-plugins-xrootd	N/A
ARC GFAL2 plugins ³	nordugrid-arc-plugins-gfal	N/A
ARC S3 plugins	nordugrid-arc-plugins-s3	N/A
ARC Internal plugin	nordugrid-arc-plugins-interr	N/A
ARC plugin for gridftp file transfers	nordugrid-arc-plugins-gridft	N/A
ARC plugin for lcas lcms	nordugrid-arc-plugins-lcas-l	N/A
ARC plugin for	nordugrid-arc-plugins-pythor	N/A
ARCHERY administration tool	nordugrid-arc-archery-manage	N/A
A-REX Python LRMS backends	nordugrid-arc-python-lrms	N/A
ARC optional worker nodes components	nordugrid-arc-wn	N/A
ARC python REST client	python3-arcrest	N/A
...	nordugrid-arc-arex-lrms-cont	N/A
Community developed runtime environments	nordugrid-arc-community-rtes	N/A

New in ARC 7

- python3-arcrest
- nordugrid-arc-arex-lrms-contrib

List of packages deprecated in ARC 7:

- nordugrid-arc-gridftpd
- nordugrid-arc-python-lrms
- nordugrid-arc-plugins-arcrest
- nordugrid-arc-plugins-gridftpjob
- nordugrid-arc-nordugridmap

List of debuginfo packages

- nordugrid-arc-debuginfo
- nordugrid-arc-debugsource
- nordugrid-arc-arex-debuginfo
- nordugrid-arc-client-debuginfo
- nordugrid-arc-datadelivery-service-debuginfo
- nordugrid-arc-devel-debuginfo
- nordugrid-arc-plugins-gfal-debuginfo
- nordugrid-arc-plugins-globus-common-debuginfo
- nordugrid-arc-plugins-gridftp-debuginfo
- nordugrid-arc-plugins-internal-debuginfo
- nordugrid-arc-plugins-lcas-debuginfo
- nordugrid-arc-plugins-needed-debuginfo

¹ Block names are used by `arcctl service enable --as-configured` command to fetch necessary packages automatically based on the blocks configured in `arc.conf`

² Package define LDAP/BDII/Glue-Schema dependencies and contains wrappers to start all this LDAP world. Infoproviders are in the A-REX package.

³ Support for specific protocols is provided by separate 3rd-party GFAL2 plugin packages. No longer supported for EL10 and derivatives.

- nordugrid-arc-plugins-python-debuginfo
- nordugrid-arc-plugins-s3-debuginfo
- nordugrid-arc-plugins-xrootd-debuginfo
- nordugrid-arc-test-utils-debuginfo
- nordugrid-arc-wn-debuginfo
- python3-nordugrid-arc-debuginfo

ARC Services

Table 6.2: List of ARC services

Block name ¹	ARC service name	Main process
[arex]	arc-arex	arched
[arex/ws]	arc-arex-ws	arched
[arex/ws/candypond]	started by arc-arex	arched
[infosys/ldap]	arc-infosys-ldap arc-infosys-ldap-slapd	slapd, bdii-update
[datadelivery-service]	arc-datadelivery-service	arched

In addition there is a stand-alone service for the Prometheus ARC exporter

Table 6.3: Additional services

Description	ARC service name	Main process
Prometheus exporter for job and datastaging metrics	arc-exporter	arc-exporter

Authorization, Mapping and Queue selection rules

Overview

ARC CE authorization and mapping rules rely on the concept of **authgroups** (configured by `[authgroup]` blocks). During the connection establishment and authentication process on any of the ARC CE interfaces (e.g. job submission on the figure), the **Authgroup matching** happened. Each authgroup represents a set of users, whose identities are matched to configured rules. The same user can match (belongs to) several authgroups.

If configured by `allowaccess` and `denyaccess` options, **per-interface authorization** rules will be enforced. This rules defines the list of authgroups that granted or denied access to the particular ARC CE interface. Without per-interface authorization rules, configuration access to the interface is granted to anyone who passed authentication process.

Warning

User will be successfully authenticated only in case of certificate and proxy-certificate has been passed validation.

In case of time synchronization problems, missing or invalid CA certificates and CRLs, missing VOMS LSC files, etc - the connection *WILL NOT* be established.

¹ Block names are used by `arcctl service start --as-configured` command to start/stop necessary services automatically based on the blocks configured in `arc.conf`

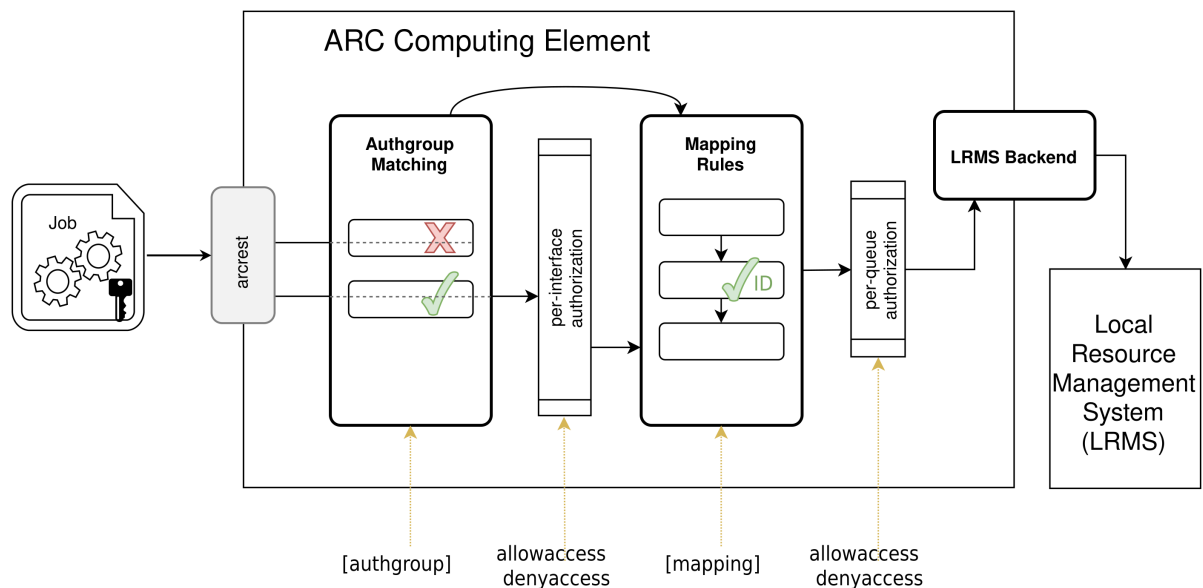


Fig. 6.1: General overview of authorization and mapping concepts in ARC

Mapping rules in ARC also relies on *authgroup* membership. Configured *mapping rules* are processed sequentially and define the OS account **ID** that will be assigned to authgroup members. Several mapping methods are available, depending on requirements.

Notice that in ARC authorization rules can also be **enforced per-queue**. It uses the same *allowaccess* and *denyaccess* options syntax as for *per-interface* configuration but can *additionally restrict access* to the particular *queue*. By default access to the queue granted to all users authorized on interface level.

Defining authgroups

Authgroups should be defined on the top of `arc.conf` using the `[authgroup:groupname]` block before it will be referenced in the other parts of configuration.

Each authgroup is a named object, that will be referenced by its name during authorization and mapping rules configuration. There are no special restrictions to the authgroup names except the absence of spaces, so you can even define `* authgroup` to blow the mind of other `arc.conf` readers.

Each config line in the `[authgroup:groupname]` block represent a matching rule that are processed sequentially.

When the matchig criteria of the rule has been satisfied by user identity - the processing stops within this authgroup.

Whether user belongs to this authgroup or not is defined by the type of rule that was matched: there are **positively** and **negatively** matching rules. By default all rules are positive (user IS a member of authgroup when matched) unless prefixed with `-` sign.

Matching one of the authgroups does not mean that the same user is not processed for the next authgroup. All authgroups are evaluated, even if a user already has a match with one of the earlier processed authgroups.

Note

Notice that:

- authgroup blocks should be defined before referencing!
- authgroup rules within blocks are order-dependent!
- all authgroup blocks are evaluated!

Complete list of authgroup rules can be found in the *ARC Configuration Reference* document. Some examples are:

```
[authgroup: norduguest]
-file = /etc/grid-security/banned.dns
voms = nordugrid Guests * *

[authgroup: admin]
subject = /O=Grid/O=Big VO/CN=Main Admin

[authgroup: lcas]
plugin = 10 /usr/libexec/arc/arc-lcas %D %P liblcas.so /usr/lib64 /etc/lcas/lcas.db

[authgroup: any]
authgroup = norduguest
authgroup = admin
authgroup = lcas

[authgroup: *]
all = yes
```

Applying authorization rules

You can enforce authentication restrictions on the REST interface which is configured with *[arex/ws/jobs]* block

In addition to interface level authorization, queue-level authorization can be configured using the same configuration approach,

The allowaccess configuration option defines that the *specified authgroup* members are authorized to access the ARC-CE via this interface or access particular queue. A related config option denyaccess can in turns be used to reject access.

Multiple allowaccess and denyaccess authorization statements are allowed within a configuration block. These statements are processed *sequentially, in the order they are specified* in the config block.

The processing stops on first allowaccess or denyaccess statement matching the authgroup membership. If there are no authorization statements specified, then no additional restrictions are applied for authorizing user access and the interface or queue is open to everybody authenticated.

Note

Default deny authorization approach can be configured using the empty authgroup.

Example: ARC zero configuration

In the shipped *zero configuration* the `[authgroup: zero]` is defined and applied to A-REX WS interface. The effect of this configuration is to allow access to CE only to the subjects stored in the `testCA.allowed-subjects` file. This file is empty by default and close down CE access until subjects are added by `arcctl test-ca usercert`.

```
[authgroup: zero]
file = /etc/grid-security/testCA.allowed-subjects

[arex/ws/jobs]
allowaccess = zero
```

Example: subject-based authorization

To authorize users based on certificate subject the *subject* or *file* rules can be used.

The *file* option support both:

- plain list of subjects (each line contains only a subject name),
- grid-mapfile format, when subject name followed by mapped account ID.

In both cases subject name should be enquoted if it contains spaces.

```
[authgroup: banana]
subject = /O=Grid/O=Bad Users/CN=The Worst

[authgroup: boss]
subject = /O=Grid/O=Big VO/CN=Main Boss

[authgroup: dnfromfile]
file = /etc/grid-security/local_users

[gridftpd/jobs]
denyaccess = banana
allowaccess = boss
allowaccess = dnfromfile
```

Example: VOMS-based authorization

To filter access based on VOMS certificate attributes, define one or more *[authgroup]* blocks using the *voms* keyword.

To verify VO membership signatures, ARC CE needs the so-called list of certificates (LSC) files that can be installed by *arcctl*.

Example configuration for *atlas*, *alice* and *swegrid.se* VOs¹:

1. Deploy LSC files:

```
[root ~]# arcctl deploy voms-lsc atlas --egi-vo
[root ~]# arcctl deploy voms-lsc alice --egi-vo
[root ~]# arcctl deploy voms-lsc swegrid.se --voms vomss://voms.ndgf.
→org:8443
```

2. Create authorization group and apply access restriction to interface and/or queue in *arc.conf*:

```
[authgroup: atlas]
voms = atlas * * *

[authgroup: alice]
voms = atlas * * *

[authgroup: swegrid]
voms = swegrid.se * * *

[authgroup: all]
authgroup = atlas
authgroup = alice
```

(continues on next page)

¹ In this example and in what follows, a simplified configuration is shown. An actual configuration will in most cases include different authgroups for different VO groups and roles.

(continued from previous page)

```
authgroup = swegrid

[gridftpd/jobs]
allowaccess = all

[arex/ws/jobs]
allowaccess = all

[queue: qlhc]
allowaccess = alice
allowaccess = atlas

[queue: qswegrid]
allowaccess = swegrid
```

Configure mapping

Any grid user should be mapped to a local account to *start processes* and *access files*.

Mapping rules configured in *[mapping] block* define which grid-users (*specified by authgroup*) are mapped to which system accounts (several mapping methods available).

Rules in the *[mapping] block* are processed *in a sequence* in line order of the configuration file (from top to bottom).

There are two kind of rules available:

- mapping rules (started with *map_*) that defines how the particular authgroup members are mapped,
- policy rules (started with *policy_*) that modifies the mapping rules sequence processing.

Default policy for mapping rules processing is:

- processing *CONTINUES* to the next rule if identity of user *DO NOT* match authgroup specified in the rule (can be redefined with *policy_on_nogroup* option)
- processing *STOPS* if identity of user matched the authgroup specified in the mapping rule. Depend on whether this mapping rule returns *valid UNIX identity* the processing can be redefined with *policy_on_map* and *policy_on_nomap* options.

Policy can be redefined at the any point of configuration sequence and affects all mapping rules defined *after* the policy rule.

Warning

If mapping process *STOPS* and there is still no local UNIX identity identified, the user running A-REX will be used (typically *root* unless redefined by *user* option for specific deployment case).

When grid-identity is mapped to *root* account - request processing fails implicitly!

Example: mapping to the same account

The `map_to_user` option allows to map all authgroup members to the same account specified as an argument.

For example in shipped *zero configuration* all users that are matched to authgroup `zero` are mapped to the same `nobody` account (and `nobody` group) that will work with local job forking:

```
[mapping]
map_to_user = zero nobody:nobody
```

Example: mapping to the accounts pool

The most secure and flexible way is to map authgroup members to account pools (so-called *map_to_pool* method). It is recommended to use pools mapping when the resource is under the use of different communities.

In this approach, every member of specified *authgroup* will be dynamically mapped to one of the available accounts in the configured pool.

Available pool account names are stored one per line in the `pool` file inside the dedicated directory. Accounts from pool are assigned by means of leasing approach. All leased accounts are stored in the other files placed in the same directory. They can be reassigned to other users after 10 days of inactivity.

Example configuration for atlas:

1. Create necessary number of accounts to be used on ARC CE and Worked Nodes of the cluster.
2. Define ARC accounts pool:

```
[root ~]# mkdir -p /etc/grid-security/pool/atlas
[root ~]# for u in atlas{001..100}; do echo $u >> /etc/grid-security/pool/
↪atlas/pool; done
```

2. Configure mapping in `arc.conf`²:

```
[mapping]
map_to_pool = atlas /etc/grid-security/pool/atlas
```

Example: mapping with external LCMAPS rules

ARC can run an external plugin to map users that can be configured with the *map_with_plugin* option.

To support several production loads, ARC ships with the built-in LCMAPS plugin included in A-REX package:

```
[authgroup:all]
all = yes

[mapping]
map_with_plugin = all 30 /usr/libexec/arc/arc-lcmaps %D %P liblcmaps.so /usr/lib64 /
↪etc/lcmaps/lcmaps-arc-argus.db arc
```

LCMAPS itself is a third-party tool that should be installed and configured separately, which is beyond the scope of this guide.

² `atlas` is the name used in `[authgroup: atlas]`

Queue selection rules and queue configuration

Note

Behaviour new in ARC 6.11

Once a submitter is authorized to submit to one of the submission interfaces defined in `arc.conf`, the next step for ARC is to select a queue and check authorisation at queue-level.

The logic is as follows

1. if (queue defined in job description `xrsl`)
 - submit to that queue w/o modification if authorised in the queue block
2. elseif (no queue defined in job description `xrsl` but a default queue is defined in `arc.conf`)
 - substitute the default queue into `xrsl` if authorised in the queue block
3. elseif (no queue in defined in job description `xrsl` and no default queue is defined in `arc.conf` and the VO is authorised in one of the `arc.conf` queues*)
 - substitute into `xrsl` the first of the queues where the VO is authorised according to `arc.conf`
4. else
 - reject
5. * any VO is authorised in a queue that has no authorisation defined at all.
 - This means that you must ensure that the authgroups you wish to allow submission from, are authorized in the queue you intend the authgroup to submit to. An example of an `arc.conf` setup that previously could result in the LRMS selecting the queue, but where submission today will be rejected could be the following:

```
[authgroup:exp1]
...

[authgroup:exp2]
...

[mapping]
..

[queue:queue1]
allowaccess= exp1

[queue:queue2]
allowaccess= exp1
```

Let's look at the `exp2` authgroup. According to this `arc.conf` we assume the authorisation and mapping is ok.

Old behaviour: If there is no queue information in the `xrsl`, the old behaviour would be to pass the job to the LRMS directly and let LRMS select the queue. This means that any queue in the system, even a queue not defined in `arc.conf` could be selected, e.g. a `queue3` defined by the LRMS, but not in `arc.conf`.

New behaviour: From ARC version 6.11 only `queue1` and `queue2` will be possible queues for `exp2` (and any authgroups), as they are the ones defined in `arc.conf`. Therefore in the example, none of the queues are allowed for `exp2` and the job will be rejected.

A possible rewrite of this `arc.conf` example-snippet, to allow `exp2` to submit to `queue2` could be:

```
[queue:queue1]
allowaccess= exp1

[queue:queue2]
<empty>
```

This is an example of rule 3: No queue is defined in the job description `xrsl` and no default queue is defined in `arc.conf` and the VO is authorised in (at least) one of the `arc.conf` queues. In this example all VO's are authorised for using `queue2`.

Batch systems support

Overview

The A-REX has to interfaced to the LRMS in order to be able to submit jobs and query their information. The A-REX supports several Local Resource Management Systems (LRMS), with which it interacts by several **backend scripts**.

The A-REX assumes that the LRMS has one or more **queues**, which is a couple of (usually homogeneous) worker nodes grouped together. The different LRMSes have different concepts of queues (or have no queues at all).

Nevertheless, in the A-REX configuration, the machines of the LRMS should be mapped to A-REX queues. The client side job submission tools query the information system for possible places to submit the jobs, where each queue on a CE is represented as an **execution target**, and treated separately.

Configuring A-REX to use one of these LRMS backends typically involves the following steps:

1. Sharing directories between A-REX, the LRMS frontend and its working nodes. It might involve setup of shared filesystems such as NFS or similar.
2. Configuring `[lrms] block` and `[queue]` blocks in `arc.conf` in respect to LRMS setup.
3. Configuring the A-REX in respect to the `shared scratch` directories configuration.

General LRMS configuration

In the `[lrms] block` the name of the LRMS has to be specified with the `lrms` option.

The supported LRMS are:

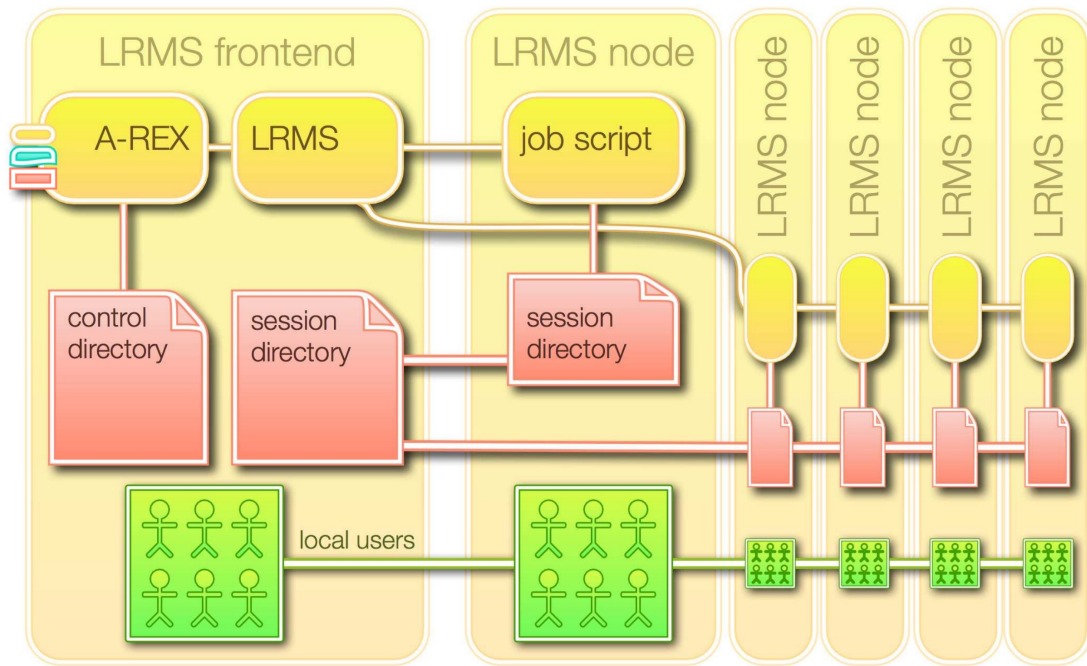
- `fork` - fork jobs on the ARC CE host node, not a cluster. Targeted for testing and development but not for real production workloads.
- `condor` - uses HTCondor-powered HTC resource
- `slurm` - for SLURM clusters

The community supported LRMSs are: * `pbs` - any flavor of PBS batch system, including Torque and PBSPRO * `ll` - Load Leveler batch system * `lsf` - Load Sharing Facility batch system * `sge` - Oracle Grid Engine (formerly Sun Grid Engine) * `boinc` - works as a gateway to BOINC volunteer computing resources

Each LRMS has its own specific configuration options that are prefixed with the LRMS name in `[lrms] block`.

Besides these specific options, the behaviour of the LRMS backend is affected by storage areas and limits setup, in particular:

- `tmpdir` - defines the path to the directory for temporary files on the worker nodes
- `shared_filesystem`, `scratchdir` and `shared_scratch` - changes the way A-REX will store the jobs' data during the processing. More details can be found in [Job scratch area](#) document.
- `defaultmemory` and `req-queue defaultmemory` - set the memory limit values for jobs that have no explicit requirements in the job description.



LRMS nodes

Fig. 6.2: The LRMS frontend and the nodes sharing the session directory and the local users

Accounting considerations

A-REX has several approaches to collect accounting information:

- using cgroups measurements of memory and CPU on the WNs
- using measurement from the GNU Time utility that wraps the job executable invocation inside the job script
- using data provided by LRMS

Depending on LRMS type in use there are different kinds of information available and/or missing in the LRMS accounting subsystem.

It is recommended to use cgroups or GNU Time methods to have reliable resources measurements in all cases. You can find more details in the *Measuring accounting metrics of the job* document.

Fork Backend

The Fork back-end is a simple back-end that interfaces to the local machine, i.e.: there is no batch system underneath. It simply forks the job, hence the name. The back-end then uses standard posix commands (e.g. `ps` or `kill`) to manage the job.

Fork is the default backend used in the ARC pre-shipped zero configuration.

Recommended batch system configuration

Since fork is a simple back-end and does not use any batch system, there is no specific configuration needed for the underlying system.

It still requires queue definition and the queue should be named `fork`.

Example:

```
[lrms]
lrms = fork

[queue:fork]
```

Known limitations

Since Fork is not a batch system, many of the queue specific attributes or detailed job information is not available. The support for the *Fork batch system* was introduced so that quick deployments and testing of the middleware can be possible without dealing with deployment of a real batch system since fork is available on every UNIX box.

The *Fork back-end* must not be used in production. The back-end by its nature, has lots of limitations, for example it does not support parallel jobs.

SLURM

SLURM is an open source resource manager designed for heterogeneous Linux clusters of all sizes. A large amount of HPC centres used SLURM.

Recommended batch system configuration

The backend should work with a normal installation of SLURM. Nodes with different specifications such as amount of memory or cores must be organized in different queues.

For production use-cases it is recommended to enable `slurm_use_sacct` option.

Example:

```
[lrms]
lrms=slurm
slurm_use_sacct=yes
defaultmemory=4096

[queue:normal]
comment=Queue for grid jobs
architecture=x86_64
totalcpus=1500

[queue:hmem]
comment=Queue for highmem grid jobs
architecture=x86_64
nodememory=2048
```

HTCondor

The **HTCondor system**, is another widely used and popular open source resource manager.

Recommended batch system configuration

Install HTCondor on the A-REX node and configure it as a submit machine. Next, add the following to the node's Condor configuration (or define CONDOR_IDS as an environment variable):

```
CONDOR_IDS = 0.0
```

CONDOR_IDS has to be 0.0, so that Condor will be run as root and can then access the Grid job's session directories (needed to extract various information from the job log).

Make sure that no normal users are allowed to submit Condor jobs from this node. If normal user logins are not allowed on the A-REX machine, then nothing needs to be done. If for some reason users are allowed to log into the A-REX machine, simply don't allow them to execute the `condor_submit` program. This can be done by putting

all local Unix users allocated to the grid in a single group, e.g. `griduser`, and then setting the file ownership and permissions on `condor_submit` like this:

```
[root ~]# chgrp griduser $condor_bin_path/condor_submit
[root ~]# chmod 750 $condor_bin_path/condor_submit
```

Example:

```
[lrms]
lrms = condor
defaultmemory = 2000

[queue:EL7]
comment = EL7 queue
defaultmemory = 3000
nodememory = 16384
condor_requirements = (Opsys == "linux") && (OpSysMajorVer == 66)
```

Known limitations

Only Vanilla universe is supported. MPI universe (for multi-CPU jobs) is not supported. Neither is Java universe (for running Java executables). ARC can only send jobs to Linux machines in the Condor pool, therefore excluding other unixes and Windows destinations.

Portable Batch System (PBS)

The Portable Batch System (PBS) is one of the most popular batch systems for small clusters. PBS comes in many flavours such as OpenPBS (unsupported), Terascale Open-Source Resource and QUEue Manager (TORQUE) and PBSPro (currently owned by Altair Engineering). ARC supports all the flavours and versions of PBS.

Recommended batch system configuration

PBS is a very powerful LRMS with dozens of configurable options. Server, queue and node attributes can be used to configure the cluster's behaviour. In order to correctly interface PBS to ARC (mainly the information provider scripts) there are a couple of configuration *REQUIREMENTS* asked to be implemented by the local system administrator:

1. The computing nodes **MUST** be declared as cluster nodes (job-exclusive), at the moment time-shared nodes are not supported by the ARC setup. If you intend to run more than one job on a single processor then you can use the virtual processor feature of PBS.
2. For each queue, one of the `max_user_run` or `max_running` attributes **MUST** be set and its value **SHOULD BE IN AGREEMENT** with the number of available resources (i.e. don't set the `max_running` = 10 if there are only six (virtual) processors in the system). If both `max_running` and `max_user_run` are set then obviously `max_user_run` has to be less or equal to `max_running`.
3. For the time being, do **NOT** set server limits like `max_running`, please use queue-based limits instead.
4. Avoid using the `max_load` and the `ideal_load` directives. The Node Manager (MOM) configuration file (`<PBS home on the node>/mom_priv/config`) should not contain any `max_load` or `ideal_load` directives. PBS closes down a node (no jobs are allocated to it) when the load on the node reaches the `max_load` value. The `max_load` value is meant for controlling time-shared nodes. In case of job-exclusive nodes there is no need for setting these directives, moreover incorrectly set values can close down a node.
5. Routing queues are now supported in a simple setup were a routing queue has a single queue behind it. This leverages MAUI work in most cases. Other setups (i.e. two or more execution queues behind a routing queue) cannot be used within ARC correctly.

6. PBS server logs SHOULD BE shared with ARC CE to allow backend scripts to check the job status and collect information needed for accounting. The path to logs on the ARC CE is defined with `pbs_log_path` option.

Additional useful configuration hints:

- If possible, please use queue-based attributes instead of server level ones.
- The `acl_user_enable = True` attribute may be used with the `acl_users = user1,user2` attribute to enable user access control for the queue.
- It is advisory to set the `max_queuable` attribute in order to avoid a painfully long dead queue.
- Node properties from the `<PBS home on the server>/server_priv/nodes` file together with the `resources_default.needsnodes` can be used to assign a queue to a certain type of node.

Checking the PBS configuration:

- The node definition can be checked by `pbsnodes -a`. All the nodes MUST have `ntype=cluster`.
- The required queue attributes can be checked as `qstat -f -Q queueName`. There MUST be a `max_user_run` or a `max_running` attribute listed with a REASONABLE value.

Example:

```
[lrms]
lrms = pbs
defaultmemory = 512
pbs_log_path = /net/bs/var/log/torque/server_logs

[queue:grid_rt]
comment = Realtime queue for infrastructure testing
allowaccess = ops
advertisedvo = ops

[queue:alien]
comment = Dedicated queue for ALICE
allowaccess = alice
advertisedvo = alice
defaultmemory = 3500
```

Known limitations

Some of the limitations are already mentioned under the PBS deployment requirements. No support for routing queues, difficulty of treating overlapping queues, the complexity of node string specifications for parallel jobs are the main shortcomings.

LoadLeveler

LoadLeveler(LL), or Tivoli Workload Scheduler LoadLeveler in full, is a parallel job scheduling system developed by IBM.

Recommended batch system configuration

The back-end should work fine with a standard installation of LoadLeveler. For the back-end to report the correct memory usage and cputime spent, while running. LoadLeveler has to be set-up to show this data in the `llq` command. Normally this is turned off for performance reasons. It is up to the cluster administrator to decide whether or not to publish this information. The back-end will work whether or not this is turned on.

Known limitations

There is at the moment no support for parallel jobs on the LoadLeveler back-end.

LSF

Load Sharing Facility (or simply LSF) is a commercial computer software job scheduler sold by Platform Computing. It can be used to execute batch jobs on networked Unix and Windows systems on many different architectures.

Recommended batch system configuration

Set up one or more LSF queues dedicated for access by grid users. All nodes in these queues should have a resource type which corresponds to the one of the the frontend and which is reported to the outside. The resource type needs to be set properly in the `lsb.queues` configuration file.

Be aware that LSF distinguishes between 32 and 64 bit for Linux. For a homogeneous cluster, the `type==any` option is a convenient alternative.

In `lsb.queues` set one of the following:

```
RES_REQ = type==X86_64
RES_REQ = type==any
```

See the `-R` option of the `bsub` command man page for more explanation.

The `reference_lrms_lsf_profile_path` option must be set to the filename of the LSF profile that the back-end should use.

Furthermore it is very important to specify the correct architecture for a given queue in `arc.conf`. Because the architecture flag is rarely set in the `xRSL` file the LSF back-end will automatically set the architecture to match the chosen queue.

LSF's standard behaviour is to assume the same architecture as the frontend. This will fail for instance if the frontend is a 32 bit machine and all the cluster resources are 64 bit. If this is not done the result will be jobs being rejected by LSF because LSF believes there are no useful resources available.

Known limitations

Parallel jobs have not been tested on the LSF back-end.

The back-end does not at present support reporting different number of free CPUs per user.

SGE

Sun Grid Engine (SGE, Oracle Grid Engine, Codine) is an open source batch system maintained by Sun (Oracle). It is supported on Linux, and Solaris in addition to numerous other systems.

Recommended batch system configuration

Set up one or more SGE queues for access by grid users. Queues can be shared by normal and grid users. In case it is desired to set up more than one ARC queue, make sure that the corresponding SGE queues have no shared nodes among them. Otherwise the counts of free and occupied CPUs might be wrong. Only SGE versions 6 and above are supported. You must also make sure that the ARC CE can run qacct, as this is used to supply accounting information.

Example:

```
[lrms]
lrms = sge
sge_root = /opt/nlge6
sge_bin_path = /opt/nlge6/bin/lx24-x86

[queue: long]
sge_jobopts= -P atlas -r yes
```

Known limitations

Multi-CPU support is not well tested. All users are shown with the same quotas in the information system, even if they are mapped to different local users. The requirement that one ARC queue maps to one SGE queue is too restrictive, as the SGE's notion of a queue differs widely from ARC's definition. The flexibility available in SGE for defining policies is difficult to accurately translate into Nordugrid's information schema. The closest equivalent of nordugrid-queue-maxqueueable is a per-cluster limit in SGE, and the value of nordugrid-queue-localqueued is not well defined if pending jobs can have multiple destination queues.

BOINC

BOINC is an open-source software platform for computing using volunteered resources. Support for BOINC in ARC is currently at the development level and to use it may require editing of the source code files to fit with each specific project.

Recommended batch system configuration

The BOINC database can be local to the ARC CE or remote. Read-access is required from the ARC CE to check for finished jobs and gather information on available resources. The ARC CE must be able to run commands in the project's bin/ directory.

Project-specific variables can be set up in an *RTE* which must be used for each job. The following example shows the variables which must be defined to allow job submission to BOINC for the project "example" to work:

```
export PROJECT_ROOT="/home/boinc/project/example" # project directory
export BOINC_APP="example"                       # app name
export WU_TEMPLATE="templates/example_IN"        # input file template
export RESULT_TEMPLATE="templates/example_OUT"   # output file template
export RTE_LOCATION="$PROJECT_ROOT/Input/RTE.tar.gz" # RTEs, see below
```

The last variable is a tarball of runtime environments required by the job.

Known limitations

The BOINC back-end was designed around projects that use virtualisation. The prototype implementation in the current ARC version may not be generic enough to suit all BOINC projects.

When preparing a BOINC job, the ARC CE copies a tarball of the session directory to the BOINC project download area. Once the job is completed and the output uploaded to the BOINC project upload area, a modified assimilator daemon must be used to copy the result back to the session directory so that it can be retrieved by ARC clients or uploaded to Grid storage by the ARC CE.

RunTime Environments in ARC

Understanding RunTime Environments

ARC Computing Element is a front-end to the various heterogeneous resource providers. To run jobs on the particular resource provider there are always a set of software or workflow-specific paths, tools, libraries, environmental variables or even dynamic content that should be recreated in the job content.

To provide a flexible way of job runtime environment tuning, ARC enforces the concept of the RunTime Environment (RTE).

ARC RunTime Environments (RTEs) provide two features:

Advertising

indicate the available environment to be requested by end-users

Modifying job environment

flexibly contextualize job execution environment

Advertising RTEs

Advertising RTEs provides user interfaces to application software and other resources in a way that is independent of the details of the local installation of the application and computing platform (OS, hardware, etc.).

It addresses setups typically required by large research groups or user bases, dealing with a common set of software. The actual implementation of a particular RTE may differ from site to site as necessary.

However, it should be designed so that resource providers with different accounting, licence or other site-specific implementation details can advertise the same application interface (RTE) for all users.

Despite possibly different parameters or implementation, the same software addressed by the same RTE name should be known by community. It is also supported to add RTE versioning at the end of the RTE name (after a dash). The RTE version will be used for resource matchmaking along with the RTE name.

For example to request ENV/PROXY and APPS/HEP/ATLAS with version 20.1.0.1 or greater in the *xRSL* job description:

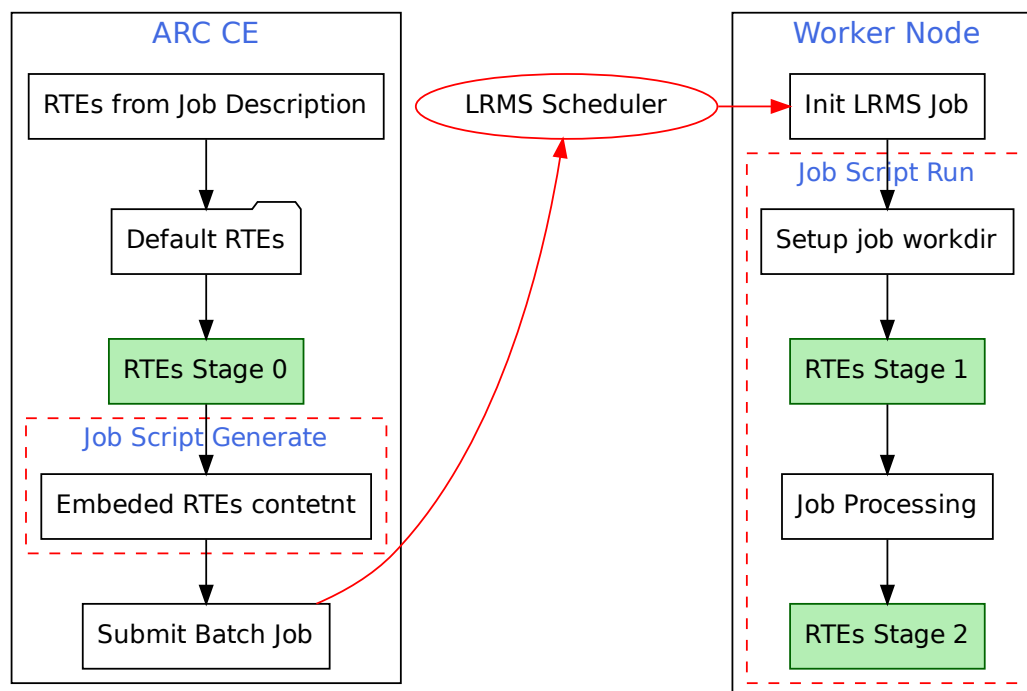
```
(runTimeEnvironment="ENV/PROXY")
(runTimeEnvironment>="APPS/HEP/ATLAS-20.1.0.1")
```

It is always up to the local system administrators to take a decision whether to *install* and *enable* a particular RTE or not.

Modifying job environment

The RTE content itself is a BASH script that is aimed to run any arbitrary code during the job life cycle.

The first argument of the RTE script indicates the so-called RTE stage. If the job description specifies additional arguments for corresponding RTE's those are appended starting at the second position.



There are 3 stages of an RTE execution:

Stage 0

RTE script sourced before the creation of the job's LRMS submission script. In this case the scripts are run by A-REX on the frontend (ARC CE), before the job is sent to the LRMS. Some environment variables are defined in this case, and can be changed to influence the job's execution later.

Stage 1

The Embedded RTE function runs before the main job processing on the Worker Node under the LRMS. Such stage can prepare the environment for some third-party software package. The current directory in this case is the one which would be used for execution of the job. The variable `$HOME` also points to this directory.

Stage 2

The embedded RTE function runs after the main job processing on the Worker Node under the LRMS. The main purpose is to clean possible changes done by *Stage 1* (like removing temporary files).

You can use [this template](#) to start writing custom RTE script that fulfill your needs.

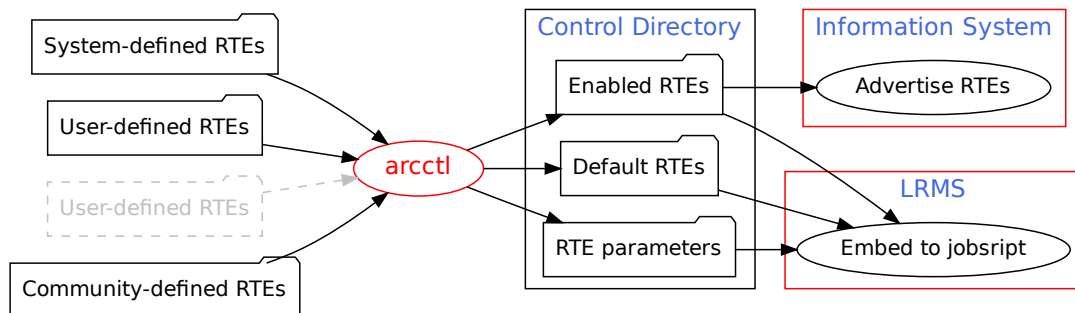
RunTime Environment script template

You can start with the following template to write custom RTE script:

```
#
# description: My RTE description ('# description:' is a keyword)
#
#####
### RTE Parameters      ###
#####
# NOTE! RTE parameters is advanced feature to keep the same code,
# but make it customizable for different clusters.
# Most probably you DO NOT need params for custom site-specific RTE.
# RTE parameters requires:
# 1. Parameters description headers ('# param:' is a keyword):
# param:PARAM_1:string:DEFAULT_VALUE:Description of parameter. Predefined 'string'
#   ↳ type means it can have arbitrary string value. Default is 'DEFAULT_VALUE'.
# param:PARAM_2:v1,v2,v3:v3:Description of parameter. This parameter can be set to 'v1
#   ↳ ', 'v2' or 'v3' value only. Default is 'v3'.
# These headers used by arcctl to operate parameters.
# Should be defined within first 20 lines of RTE script.
#
# 2. Definiton of parameters default values for shell
PARAM_1="${PARAM_1:-DEFAULT_VALUE}"
PARAM_2="${PARAM_2:-v3}"

#####
### RTE Code           ###
#####
if [ "x$1" = "x0" ]; then
    # RTE Stage 0
    # You can do something on ARC host as this stage.
    # Here is the right place to modify joboption_* variables.
    # Note that this code runs under the mapped user account already!
    # No root priveledges!
elif [ "x$1" = "x1" ]; then
    # RTE Stage 1
    # Do something on WN before job execution.
elif [ "x$1" = "x2" ]; then
    # RTE Stage 2
    # Do something on WN after job execution.
fi
```

Operating RunTime Environments



Installing RTE scripts

There are set of *System-defined RTEs* pre-installed with the ARC CE packages that aim to fulfill common workflows.

An ARC CE administrator can add additional RTE directories (so-called *User-defined RTEs*). These additional places should be specified in `arc.conf` using the `runtime_dir` configuration option in the `[arex]` block. Custom RTE scripts can be developed using [this template](#) as a starting point.

The *Community-defined RTEs* are additional RTEs created by research communities. These RTEs can be provisioned to the ARC CE from the trusted registries (including required software bundles) with the *ARC Control Tool*.

Note

In ARC directories with RTE script are local to ARC CE and **SHOULD NOT** be shared with worker nodes

The RTE names used for *advertising* are implied by directory structure, e.g. in the ENV/PROXY the ENV is the directory inside the *System RTEs* location and PROXY is the name of the file.

Enabling RTEs

The *Installed RTEs* must be enabled to be *advertised* and used during the job submission.

By name

To enable a particular RTE by name run the following command:

```
arcctl rte enable ENV/PROXY
```

By path

Especially if you have several RTEs with the same name installed, you can choose the exact one by specifying the filesystem path:

```
arcctl rte enable /usr/share/arc/rte/ENV/PROXY
```

Using wildcards

To enable several RTEs you can pass as many names as you want to the `arcctl` command. Additionally you can use glob (man 7 glob) wildcards in RTE names.

The following command will enable all APPS/HEP/ATLAS RTEs for SLC7 builds:

```
arcctl rte enable APPS/HEP/ATLAS-*X86_64-SLC7-*
```

Dummy RTEs

In case you need RTEs only for *advertising* but no need for script content, you can create a *Dummy RTE* for the specified name. The following command enables an APPS/MYAPP RTE with empty contents:

```
arcctl rte enable APPS/MYAPP --dummy
```

An example of dummy RTEs could be APPS/ATLAS-SITE used by the ATLAS experiment for sites to advertise that this indeed is an ATLAS-SITE.

Default RTEs

Default RTEs aim to address the workflows when *advertising* and implicit request in the job description is not needed, however *modification* of every submitted job (adjusting memory, setting LRMS scratch, etc) is required on the resource provider.

Installed RTEs can be selected for default inclusion to the job lifecycle with the following *ARC Control Tool* command:

```
arcctl rte default ENV/LRMS-SCRATCH
```

This will transparently add ENV/LRMS-SCRATCH to each job and will be executed the same way as *Enabled RTEs*. A default RTE does not need to be **enabled** to be executed, but it **can be enabled** if one wants to publish it in addition to executing it.

Note

You can use the same *by-name*, *by-path* and *wildcard* techniques as for *enabling*

RTE Parameters

To handle heterogeneity of resource providers, some RTEs can be parametrized.

For example, the system-defined ENV/PROXY RTE that transfers the delegated proxy-certificate to the worker node can optionally transfer CA certificate directories. This optional part is controlled by COPY_CACERT_DIR parameter.

To check if an RTE contains parameters and their default values, run the:

```
[root ~]# arcctl rte params-get ENV/PROXY
COPY_CACERT_DIR=Yes
```

You can also see the description and allowed values adding the `--long` keyword.

To set an RTE parameter value, the following command is used:

```
arcctl rte params-set ENV/PROXY COPY_CACERT_DIR No
```

List available RTEs and their status

To view the summary of all *installed*, *enabled* and *default* RTEs run:

```
[root ~]# arcctl rte list
<output omitted>
APPS/HEP/ATLAS-20.8.0-X86_64-SLC7-GCC48-OPT (user, enabled)
APPS/HEP/ATLAS-20.8.1-X86_64-SLC7-GCC48-OPT (user, enabled)
APPS/HEP/ATLAS-20.8.2-X86_64-SLC7-GCC49-OPT (user, enabled)
<output omitted>
ENV/LRMS-SCRATCH (system, default)
ENV/PROXY (system, masked, disabled)
ENV/PROXY (user, enabled)
ENV/RTE (system, disabled)
ENV/RUNTIME/ALIEN-2.17 (user, enabled)
VO-biomed-CVMFS (dummy, enabled)
```

The first tag describe the RTE origin (system, user or *dummy*). The following tags shows the status.

The special masked keyword indicates that the RTE name is used more that once and *by-name* operations will apply to another RTE script. For example ENV/PROXY will be enabled from the user-defined location as the system-defined is masked. However it is possible to enable a masked RTE *by specifying its path*.

Listing the particular kind of RTEs (e.g. *enabled*) is possible with the appropriate argument (see *ARC Control Tool* for all available options):

```
[root ~]# arcctl rte list --enabled
<output omitted>
APPS/HEP/ATLAS-20.8.2-X86_64-SLC7-GCC49-OPT
ENV/PROXY
ENV/RUNTIME/ALIEN-2.17
VO-biomed-CVMFS
```

The *long* listing allows to get the detailed pointers to RTEs locations and descriptions:

```
[root ~]# arcctl rte list --long
System pre-defined RTEs in /usr/share/arc/rte:
  ENV/PROXY # copy proxy certificate to the job session
↪directory
  ENV/RTE # copy RunTimeEnvironment scripts to the job
↪session directory
```

(continues on next page)

(continued from previous page)

ENV/LRMS-SCRATCH	# enables the usage of local to WN scratch_
↪directory defined by LRMS	
User-defined RTEs in /etc/arc/rte:	
ENV/RUNTIME/ALIEN-2.17	# RTE Description is Not Available
ENV/PROXY	# RTE Description is Not Available
Enabled RTEs:	
ENV/RUNTIME/ALIEN-2.17	-> /etc/arc/rte/ENV/RUNTIME/ALIEN-2.17
ENV/PROXY	-> /etc/arc/rte/ENV/PROXY
Default RTEs:	
ENV/LRMS-SCRATCH	-> /usr/share/arc/rte/ENV/LRMS-SCRATCH

View RTE content

Dumping the content of an RTE that will be embedded to job script is possible with the cat action:

```
[root ~]# arcctl rte cat ENV/LRMS-SCRATCH
SCRATCH_VAR="LOCALTMP"
# description: enables the usage of local to WN scratch directory defined by LRMS
# param:SCRATCH_VAR:string:WORKDIR:Variable name that holds the path to job-specific_
↪WN scratch directory

SCRATCH_VAR="${SCRATCH_VAR:-WORKDIR}"

if [ "x$1" = "x0" ]; then
    RUNTIME_LOCAL_SCRATCH_DIR="\${${SCRATCH_VAR}}"
fi
```

Disable and Undefault RTEs

Enabled RTEs can be disabled running:

```
arcctl rte disable ENV/PROXY
```

The similar operation for *default RTEs* is called undefault:

```
arcctl rte undefault ENV/LRMS-SCRATCH
```

Note

You can use the same *by-name*, *by-path* and *wildcard* techniques as for *enabling*

System-defined RunTime Environments shipped with ARC

ENV/PROXY

Export delegated credentials (proxy certificate) to the job's session directory. Optionally copies CA certificates directory from ARC CE to session directory.

Sets the X509_USER_PROXY, X509_USER_CERT and X509_CERT_DIR to make files instantly available to client tools.

Parameters:

- `COPY_CACERT_DIR` = Yes/No - If set to Yes, CA certificate directory will be copied to the session directory along with proxy certificate. Default is Yes.
- `USE_DELEGATION_DB` = Yes/No - If set to Yes RTE will try to extract proxy certificate from A-REX delegation DB (works in limited number of cases). Default is No.

ENV/RTE

Copy RunTime Environment scripts to the job session directory for some workloads that require the files themselves instead of embedding the RTEs to the jobscript.

Designed to be used as *default* RTE.

Has no parameters.

ENV/LRMS-SCRATCH

Many resource providers use *scratchdir* to move files to the local worker node disk before running the job.

When the local scratch directory is created dynamically by LRMS (e.g. in the job prologue) and then cleaned up automatically after the job completion, the ENV/LRMS-SCRATCH is needed. The scratch place should be indicated by some environmental variable that holds a path to such LRMS-defined scratch directory.

This RTE is designed to be used as a *default* RTE to enable this optional functionality.

Parameters:

- `SCRATCH_VAR` = name - Variable name that holds the path to job-specific WN scratch directory. Default is `WORKDIR`.
- `TMPDIR_LOCATION` = path - Define the `TMPDIR` path on WN. Variable names can be used as a part of the path, e.g. `'$WORKDIR/tmp'`.

Note

The ENV/LRMS-SCRATCH is **not needed** if the scratch directory used and created on the worker node is of type `</arc-conf-scratchdir-var>/<arc-job-id>`. It is only needed if the folder should be of type `<path defined by LRMS in SCRATCH_VAR>/<arc-job-id>`, i.e. a LRMS defined path is included.

ENV/CONDOR/DOCKER

ARC HTCondor backend supports submission to the Docker universe. This RTE enables this feature on-demand.

The RTE can be used by end-users when *enabled*. The RTE argument defines the Docker image name to be used, e.g:

```
(runtimeenvironment="ENV/DOCKER" "debian")
```

The RTE can also be used as a *default* RTE to enforce Docker universe submission for any job. The Docker image should be set with the RTE parameter.

Parameters:

- `DOCKER_IMAGE` = name - Docker image to use for submitted jobs by default.

ENV/SINGULARITY

A general-purpose RTE that allows to run the submitted jobscript inside a defined singularity container image.

This RTE is designed to be used as both a *default* and *enabled* RTE.

When *enabled*, a user can specify the singularity image as an RTE argument:

```
(runtimeenvironment="ENV/SINGULARITY" "mysoftwareimage.sif")
```

One can also specify docker images from central repos assuming the compute node has outbound access and can download the image:

```
(runtimeenvironment="ENV/SINGULARITY" "docker://rockylinux:9")
```

Parameters:

- **SINGULARITY_IMAGE** = *images* - Define singularity images to be used with a job. It accepts comma separated *vo:path* pairs, where *vo* defines the virtual organization with image located at *path*. The special default value for *vo* defines the image used by default (if specific image for VO is not defined). The special NULL value for *path* skips singularity usage.
- **SINGULARITY_OPTIONS** = *options* - Define additional options for singularity executable. In particular additional storage areas (e.g. CVMFS, CA certificates) can be specified to be mounted here.

ENV/CANDYPOND (experimental)

Makes available the arccandypond tool for use inside the job script on the Worker Nodes (including necessary environmental variables for its operation).

Note

The CandyPond service itself should be enabled (defining the *[arex/ws/candypond]* block) on ARC CE as well.

Parameters:

- **CANDYPOND_URL** = *url* - Redefine the URL of CandyPond service (default is auto – ARC CE URL used for job submission will be used automatically)

Working with community-defined RTEs

Community-defined RTEs are RTEs that created, managed and distributed by a particular community.

ARC provides the software solution to automate Community-defined RTEs discovery and software environment provisioning for distributed computing e-Infrastructures.

Typically community-defined RTEs describes software packages used for computations. The community itself is responsible for building, testing and verifying a particular version of software packages and the RTE scripts that will prepare such runtime environment to be used on the computing cluster worker nodes behind the computing element.

In general *RunTime Environments in ARC* are very flexible, so in addition to defining software packages such RTEs can be used to transfer data or communicate with community services to fetch/register some data. It is up to community to define what they need.

This document describes how to work with community RTEs from the ARC CE admin point of view.

To establish a community-defined RTEs registry refer to *this guide*.

1. Enable Community-defined RTEs support

Community RTEs support is added as a technology preview in the ARC 6.5 and turned off by default.

If you need to support Community RTEs deployment on the ARC CE, please install `nordugrid-arc-community-rtes` package:

```
[root ~]# dnf -y install nordugrid-arc-community-rtes
```

2. Establish trust chain with community

The trust-chain between community and site-admin is based on the digital signatures. All Community-defined RTEs are supposed to be signed using OpenPGP standard for signatures. Technical implementation relies on the GNU Privacy Guard (GPG) software.

To add a new community to the trusted list run:

```
[root ~]# arcctl rte community add example.org
The imported community public key data is:
pub 2048R/AA56A775 2020-01-30 [expires: 2022-01-29]
    Key fingerprint = 3A47 F0D4 E406 D854 EDAA ADB5 8FD6 DD57 AA56 A775
uid                               Example Computations Lab <support@example.org>
sub 2048R/3F914B9D 2020-01-30 [expires: 2022-01-29]

Is the community key fingerprint correct? (yes/NO): yes
```

Warning

Check the community key fingerprint matches the one provided to you by community authorities!

Note

Alternatively you can pass expected fingerprint value to `--fingerprint` option

In the provided example the `example.org` is the name of the community to add and at the same time it is interpreted as a domain name of ARCHERY community-defined RTEs registry.

If the ARCHERY domain name is different add the `--archery <DOMAIN>` option.

It is also possible to establish trust with community using OpenPGP compatible keyserver or web-based RTEs registry¹ as an alternative to ARCHERY.

3. Discover RTEs in the registry

Note

Examples below show `APPS/EXAMPLE.ORG/SIMULATION-3.0.1` is already deployed on ARC CE. That is result of the next step execution.

You can list all available community-defined RTEs with `rte-list` command:

¹ TODO: Document this alternative options. For now you can follow `--help` to find some info

```
[root ~]# arcctl rte community rte-list example.org
APPS/EXAMPLE.ORG/SIMULATION-3.0.1    (deployed, registry)
APPS/EXAMPLE.ORG/ANALYSIS-1.7.0      (registry)
ENV/EXAMPLE.ORG/SENDSTATS-1.0.0      (registry)
```

RTEs optionally provide description string that can be viewer with long listing:

```
[root ~]# arcctl rte community rte-list example.org --long
Community deployed RTEs:
  APPS/EXAMPLE.ORG/SIMULATION-3.0.1    # Example Simulation Software
Additional community RTEs available in the registry:
  APPS/EXAMPLE.ORG/ANALYSIS-1.7.0      # Example Analysis Software
  ENV/EXAMPLE.ORG/SENDSTATS-1.0.0      # Send stats to central example.org services
```

Before deployment it can be useful to look inside the RTE script. The `rte-cat` will show you the content:

```
[root ~]# arcctl rte community rte-cat example.org APPS/EXAMPLE.ORG/ANALYSIS-1.7.0
# description: Example Analysis Software
# download: url:https://example.org/soft/analysis.sif
↪checksum:md5:63490ad38190a6f172a9020c0c5615f4

if [ "x$1" = "x1" ]; then
  mkdir ${RUNTIME_JOB_DIR}/bin
  cat > ${RUNTIME_JOB_DIR}/bin/example-analysis <<EOF
#!/bin/bash
exec singularity run ${RUNTIME_JOB_SWDIR}/analysis.sif "$@"
EOF
  chmod +x ${RUNTIME_JOB_DIR}/bin/example-analysis
  export PATH=${RUNTIME_JOB_DIR}/bin:${PATH}
fi
```

4. Deploy community-defined RTE

Deploying a community-defined RTE from the registry requires nothing more than passing the RTE name to `rte-deploy`:

```
[root ~]# arcctl rte community rte-deploy example.org APPS/EXAMPLE.ORG/SIMULATION-3.0.
↪1
```

This command will:

- fetch the RTE script signed by community
- verify the signature using trusted community public keys (installed during step 2)
- deploy the RTE script itself to be used further with `arcctl rte`
- download files specified in the community-defined RTE script to the community software location
- verify the checksum data of downloaded files

Note

HINT: It is useful to increase the debug level to at least INFO level during the deployment phase to monitor the progress.

5. Enable community-defined RTE

After deployment of community-defined RTEs, operating can be done *as usual* - the same was as for other RTE types.

In particular you can list, *enable* or *default* RTEs including deployed from community registry:

```
[root ~]# arcctl rte list
ENV/CANDYPOND          (system, disabled)
ENV/PROXY               (system, enabled)
ENV/RTE                 (system, disabled)
ENV/SINGULARITY         (system, disabled)
APPS/EXAMPLE.ORG/SIMULATION-3.0.1 (community, disabled)

[root ~]# arcctl rte enable APPS/EXAMPLE.ORG/SIMULATION-3.0.1
```

Additional information and hints

This section provide information how to customize the community-defined RTEs operations.

Location of deployed community software

By default the location for the deployed community software is picked up automatically based on `arc.conf` and relies on the session directory in particular.

You can discover and change the location with `arcctl`:

```
[root ~]# arcctl rte community config-get example.org
SOFTWARE_DIR=/shared/session/_software/example.org
SOFTWARE_SHARED=Yes

[root ~]# arcctl rte community config-set example.org SOFTWARE_DIR /opt/community/
↳ example.org
```

Operating without the registry

It is possible to deploy community-defined RTEs using the same `arcctl` automations without registry.

During the deployment phase you can provide the URL to the signed RTE file with the `--url` option:

```
[root ~]# arcctl rte community deploy example.org ENV/URLDEPLOYED-1.0.0 --url http://
↳ example.org/rte.signed
```

Or it can be even an RTE script without a signature if you trust the contents:

```
[root ~]# arcctl rte community deploy example.org APPS/SIM-DEVEL --url file:///home/
↳ example/dev/myrte.sh --insecure
```

Removing RTEs and communities

If you want to remove a deployed RTE or an entire community, there are `rte-remove` and `remove` actions respectively:

```
[root ~]# arcctl rte community rte-remove example.org APPS/EXAMPLE.ORG/SIMULATION-3.0.
↪ 1
[2020-02-06 18:11:03,653] [ARCCTL.RunTimeEnvironment.Community] [ERROR] [32505]↪
↪ [Community
RTE APPS/EXAMPLE.ORG/SIMULATION-3.0.1 is enabled. Please disable it first or use "--
↪ force"
to disable and undefault automatically]
[root ~]# arcctl rte community rte-remove example.org APPS/EXAMPLE.ORG/SIMULATION-3.0.
↪ 1 --force
[root ~]# arcctl rte community remove example.org
```

Measuring accounting metrics of the job

ARC CE has built-in capabilities to collect information about per-job resource consumption. This includes both ARC CE resources (e.g. data transfers, software environments) and worker nodes resources (e.g. CPU and memory usage). The full list of attributes stored in the A-REX Accounting Records (AAR) can be found in [this document](#).

A-REX can use different methods (described below) to measure memory and CPU usage on the worker nodes, depending on their availability in the particular deployment case.

Measuring memory and CPU usage on the WN with cgroups

When recent versions of GNU/Linux OS are used on the worker nodes the most precise and transparent way to measure all job workload is to rely on the `cgroups` kernel subsystem. Any `systemd`-based Linux distribution relies on `cgroups` heavily and they are already used.

Note

Some older operating systems may require mounting the `cgroups` tree explicitly. For example in RHEL6 it can be easily done with `libcgroup`:

```
[root ~]# dnf install libcgroup
[root ~]# service cgconfig start
```

The benefit of using `cgroups` is that everything will be accounted. Even if several payloads are executed (e.g. in pilot mode) or extra helper processes are spawned - the resource accounting will be accurate for the all workload done.

Enabling cgroups usage

To be able to use `cgroups` for accounting ARC needs an extra tool installed on the worker nodes – the `arc-job-cgroup`. Based on the tool availability, the job script will or will NOT use `cgroups` for measuring accounting metrics automatically.

The `arc-job-cgroup` tool is available for the majority of OSes as a packaged binary build as a part of ARC distribution (`nordugrid-arc-wn` package). So the easiest way to install it is to use your package manager on the worker nodes, e.g.:

```
[root ~]# dnf install nordugrid-arc-wn
```

If it is not possible to install the packaged version for some reason, it is easy to compile the pure C source code with standard C library calls only.

```
[root ~]# wget https://source.coderefinery.org/nordugrid/arc/raw/master/src/wn/arc-
↪ job-cgroup.c
[root ~]# cc -o arc-job-cgroup arc-job-cgroup.c
[root ~]# mv arc-job-cgroup /usr/local/bin/
[root ~]# chmod 4755 /usr/local/bin/arc-job-cgroup
```

How ARC operates cgroups

The idea behind LRMS-independent cgroup-based resource usage measurements in ARC is to:

- create child cgroups for memory and cpuacct controllers
- put the jobscript process into created cgroups (this will automatically catch all child processes)
- collect the accounting data at the end of the jobscript
- remove the child cgroup created at the beginning (moving all processes to parent cgroup)

If cgroups are used in the Kernel, the process already belongs to some cgroup. It can be either a root cgroup (used for all processes) or some dedicated cgroup created by LRMS with cgroups support, container management system, etc.

All resources used by the child cgroup are accounted in the parent cgroup. Moreover all parent-defined limits are inherited and enforced as well. So creating another child cgroup in hierarchy is safe from all points of view.

Warning

Creating a child cgroup and put a task into it requires root privileges. This is the reason behind the SUID bit for arc-job-cgroup. However the code itself is as simple as the mkdir. You can review [these 333 lines](#) to reassure any possible fears.

If the arc-job-cgroup tool is not available, the cgroups tree is not mounted, or there are any other issues with cgroups creation, the job script code falls back to the GNU time measurement method.

Measuring memory and CPU usage on the WN with GNU time

The GNU time utility is capable of measuring and displaying information about the resources used by the executable it runs.

It is used as a part of the ARC-generated job script if found on the worker node.

Note

Changed in version 6.2.

In case of successful cgroups usage, GNU time will NOT be used by job script.

Warning

The GNU time is a separate binary typically installed by dedicated package. Do not mix it up with built-in version of time in your shell (e.g. bash, zsh).

Typically you can install it with e.g. `dnf install time` or similar package management command.

For a non-standard location of GNU Time the `gnu_time` configuration option can be used to define it.

If the GNU `time` utility is not available the job will run as it is and only LRMS-provided metrics will be accounted.

Using LRMS-provided metrics

After the job has finished execution in the LRMS, the batch system backend scan-script extracts accounting information about the job from the LRMS, either executing command line clients, parsing logs or using API.

The exact data measurements and the method of these data collection completely depends on the LRMS backend implementation and differs from one backend to another.

Common metrics include `LRMSStartTime` and `LRMSEndTime`. There are also typically some memory and CPU usage metrics available.

Accounting Subsystem

Warning

Information in this chapter is relevant only for 6.4+ ARC releases.

Moreover ARC 6.12 get accounting changes to address the APEL move to ARGO messaging service protocol. If you are publishing to APEL you must update to 6.12+ ARC release.

Note

If you are looking for the information about legacy accounting subsystem in 6.0-6.3 ARC releases please read *Accounting with legacy JURA* but it is highly recommended to update to recent release.

The next generation ARC Accounting Subsystem aimed to improve scalability, eliminate bottlenecks of the *legacy implementation* and provide a local on-site ARC CE accounting database to be queried and analyzed by CE admins.

Overview

Figure shows the overview of next generation accounting subsystem in ARC. More details can be found in *ARC Accounting Technical Details*.

The central point of ARC next generation Accounting Subsystem is a local SQLite accounting database that stores all the *A-Rex Accounting Record (AAR)* information. AAR defines all accounting information stored about a single ARC CE job.

The database is populated directly by A-REX based on per-job files inside the *control directory*. In particular the `.diag` file is a main source of *resource usage* information, `.statistics` holds data transfers measurements and `.local` is a general datastore about the job properties like ID, owner, etc.

The record about the job is created in the accounting database as soon as the job enters ARC CE. Than each state change of the job is recorded. Once the job reached terminal state (both successfull completion or failed) all resource usage data is written becomes available via local queries (`arcctl accounting`) or for reporting to SGAS/APEL (`jura-ng`).

Accounting subsystem is integral part of the A-REX and allways enabled.

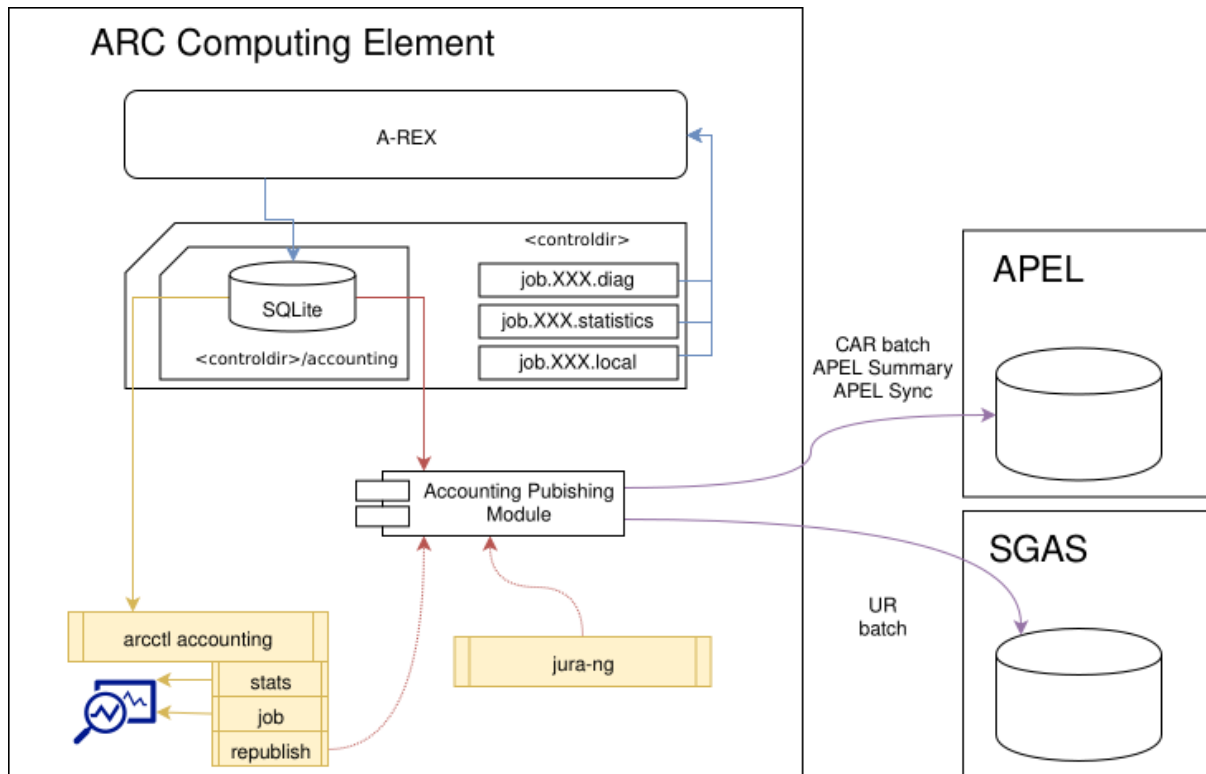


Fig. 6.3: Next generation ARC accounting subsystem overview: creating AAR records, getting access to the data and publishing to external SGAS/APEL services

Configuration

In the typical use-case accounting subsystem just works and does not requires additional configuration.

However in some rare cases there are several things you can consider to tune.

WLCG VOs

The recommended and typically used way to authorize WLCG VOs is to use `[authgroup]` block `voms` configuration option that exactly match attributes in VOMS extension of proxy certificate.

If you have authorization rules configured this way NOTHING should be configured.

In case you do not have at least one `voms` option in the any of the defined `[authgroup]` blocks A-REX will not trigger the VOMS attributes parsing, that consequently leads to no VO info in the accounting!

To have VO info in the accounting in this particular case you can add standalone authroup with `voms` option that have no further usage in `arc.conf`:

```
[authgroup: parsevoms]
voms = * * * *
```

There is also `forcedefaultvoms` configuration option (can be defined on per-queue basis) that can define the accounted VO for jobs that have no VOMS extension in owner's certificate.

Enabling accounting records reporting

It is possible to send resource usage reports to the centralized SGAS and/or APEL accounting databases.

Based on the local accounting database data the *Accounting Publishing Module* is capable of generating the:

- OGF Usage Record 1.0 (UR) XML format to be sent to the SGAS LUTS (Logging and Usage Tracking Service)
- EMI Compute Accounting Record (CAR) v1.2 XML format for individual job records to be sent to APEL
- APEL Summaries (to reduce traffic) and APEL Sync messages for full-featured integration with APEL services

Regular publishing sequence is handled by `jura-ng` helper tool and can be enabled with `[arex/jura]` block.

A-REX periodically runs `jura-ng` with default and minimum *hourly period*. It can be increased with *urdelivery_frequency* option. Furthermore it can be increased per-target, using the same option inside a particular target block.

JURA has dedicated log file defined by *logfile* option. Log rotation has been set for default `/var/log/arc/jura.log` location.

Accounting services for sending the records are configured with dedicated sub-blocks. You need to define a separate block with a unique targetname for every target server used.

Note

The target block name will be used by `jura-ng` to track that latest records sent to this target. Be aware that if you rename the block, target will be handled as a new one. However `targeturl` change will not trigger a new target handling and records will continue publishing using the latest recorded timestamp

The AARs data will be reported to all of the defined destinations, unless `vofiler` option configured for some of them to filter records by VO name.

Warning

There were several issues in the codebase and misunderstanding from the operational point of view how the *benchmarks* are propagated to the central accounting services with ARC accounting. Please read *About benchmarks and accounting publishing* for more details.

Configuring reporting to SGAS

The *SGAS sub-block* enables and configures an SGAS accounting server as a target destination to which ARC CE will send properly formatted OGF.98 Usage Record 1.0 (UR) XML records.

The *targeturl* option is the only mandatory parameter to configure SGAS target. In the specific setup cases you can also apply *VO filtering* and *set prefix* for local job IDs.

Example:

```
[arex/jura/sgas: NeIC]
targeturl = https://grid.uio.no:8001/logger
urbatchsize = 80
```

Configuring reporting to APEL

Changed in version 6.12.

The *APEL sub-block* enables and configures an APEL accounting server as a target destination to which ARC will send the data.

The *targeturl* option defines the APEL broker URL to send records to. The currently known APEL AMS endpoint is provided in the *targeturl example* but you should refer to APEL for up to date information.

The *apel_messages* option allows you to choose between per-job EMI Compute Accounting Record (CAR) XML records publishing and APEL Summaries publishing. Sending summaries is a default and recommended by APEL behaviour that allows to save resources and traffic.

APEL Sync records that synchronize the total job counters per-month are always sent.

You also need *GOCDB name* of the resource. Since move to AMS publishing the *APEL topic* is always gLite-APEL that is a default value.

For correct production accounting setup it is recommended to specify resource *benchmarking results* in the *[queue:name] block*. ARC assumes that the nodes in the same queue are homogeneous with respect to the benchmark performance and benchmark values are specified per-queue.

Example:

```
[arex/jura/apel: EGI]
targeturl = https://msg.argo.grnet.gr
gocdb_name = RAL-LCG2

[queue: grid]
benchmark = HEPSPEC 8.73
```

Lookup local accounting data

Data in ARC CE accounting database can be viewed with *ARC Control Tool*. Timeframe of interest and many other filters can be specified, e.g.:

```
[root ~]# arcctl accounting stats --filter-vo ops --start-from 2019-07-10
A-REX Accounting Statistics:
  Number of Jobs: 1317
  Execution timeframe: 2019-07-10 00:01:45 - 2019-07-26 12:48:13
  Total WallTime: 4:40:28
  Total CPUtime: 0:22:32 (including 0:00:00 of kernel time)
  Data staged in: 2.3M
  Data staged out: 9.7K

[root ~]# arcctl accounting job events
↳ UuFLDmmS5unf5481mks8bjnABFKDmABFKDmmSMKDmNBFKDm9nRnVo
2019-07-10 17:30:00  ACCEPTED
2019-07-10 17:30:00  PREPARING
2019-07-10 17:30:00  DTRDOWNLOADSTART
2019-07-10 17:30:08  SUBMIT
2019-07-10 17:30:08  DTRDOWNLOADEND
2019-07-10 17:30:11  INLRMS
2019-07-10 17:30:16  LRMSSTART
2019-07-10 17:30:26  LRMSSEND
2019-07-10 17:30:47  FINISHING
2019-07-10 17:30:47  FINISHED
2019-07-14 15:01:16  DELETED
```

(continues on next page)

(continued from previous page)

```
[root ~]# arcctl accounting job transfers_
↳UuflDmmS5unf5481mks8bjnABFKDmABFKDmmSMKDmNBFKDm9nRnVo
Data transfers (downloads) performed during A-REX stage-in:
  http://download.nordugrid.org:80/packages/nordugrid-arc/releases/6.1.0/src/
↳nordugrid-arc-6.1.0.tar.gz:
  Size: 5.2M
  Download timeframe: 2019-07-10 17:30:00 - 2019-07-10 17:30:04
No stage-out data transfers (uploads) performed by A-REX.
```

More queries examples can be found in [this document](#).

Republishing records

When something goes wrong with accounting services, network, etc there is possible need of republishing local records again.

In the current implementation of accounting subsystem, there is no difference between publishing and re-publishing. The same *Accounting Publishing Module* will be used to generate the records to be sent to target service for defined timeframe.

Rebuplishing is triggered by ARC CE administrator using *ARC Control Tool*.

The most streamlined way is to republish data to the target that is already configured in `arc.conf` for regular publishing:

```
[root ~]# arcctl accounting republish -b 2019-06-01 -e 2019-07-01 -t EGI
```

However it is also possible to define all target options from the command line, without the defined target in `arc.conf`:

```
[root ~]# arcctl accounting republish --end-from 2019-06-01 --end-till 2019-07-01 \
> --apel-url https://msg.argo.grnet.gr --gocdb-name "UA-KNU" \
> --apel-messages summaries --apel-topic gLite-APEL
```

Clean up of the <controldir>/logs folder

With the accounting subsystem change in ARC 6.4 legacy archive files are no longer used or written to the <controldir>/logs directory. It is recommended to manually wipe this directory.

Job scratch area

ARC allows to configure different approaches to manage the job scratch area during the job life cycle. In most cases this is achieved by generating a wrapper submission script to the batch system that carries on tasks relative to directory creation and data movement.

The key elements are:

job session directory

directory on the ARC CE where the job files are located.

job scratch directory

directory on LRMS-managed worker nodes (WNs) where all I/O during computation is performed.

<arc-job-id>

a unique identifier for the job, assigned by ARC.

Example: 869c20be3993

job files

- input files, those coming from the client and data-staging framework that are needed for the actual job processing, copied inside a folder named after *<arc-job-id>*, that will also contain the job *stdout* and *stderr* files.
- some metadata files used by ARC, mainly *<arc-job-id>.comment* and *<arc-job-id>.diag*

Job session directory is configured with *sessiondir* configuration option. It is possible to configure several session directory root paths. A-REX will then select one of the available directories and append the *<arc-job-id>* to the path.

Example:

```
/nfs/sessiondir5/
├── 869c20be3993
│   ├── script.sh
│   ├── outfile
│   ├── stderr
│   └── stdout
├── 869c20be3993.comment
└── 869c20be3993.diag
```

There are several configuration options that affect the selection of *job scratch directory*:

shared_filesystem

defines if the *job session directory* is shared between ARC CE and WNs (by means of e.g. NFS). Sets the environment variable `RUNTIME_NODE_SEES_FRONTEND`.

scratchdir

defines the path to *job scratch directory* on the WN. Sets the environment variable `RUNTIME_LOCAL_SCRATCH_DIR`.

movetool

defines what tool the job wrapper will use to move data from the session directory to the local WNs scratch directory. Sets the environment variable `RUNTIME_LOCAL_SCRATCH_MOVE_TOOL`.

shared_scratch

defines that the WNs scratch directory can be accessible from ARC CE (by means of e.g. NFS) using the configured path. Sets the environment variable `RUNTIME_FRONTEND_SEES_NODE`.

Currently this legacy option is not used much, as it is uncommon that nodes share a single scratch directory. If you require such functionality please contact the ARC team.

Note

Described environment variables can be redefined by *RunTime Environments* dynamically. For example *ENV/LRMS-SCRATCH* can be used to utilize local scratch that created dynamically by LRMS.

Compute inside shared session directory

In the simplest case, where the *job session directory* is **shared** between ARC CE and WNs (e.g. via NFS) and accessible on the same path - the *job scratch directory* is the *job session directory*.

Configuration:

```
[arex]
shared_filesystem = yes
```

Variables set in the wrapper script:

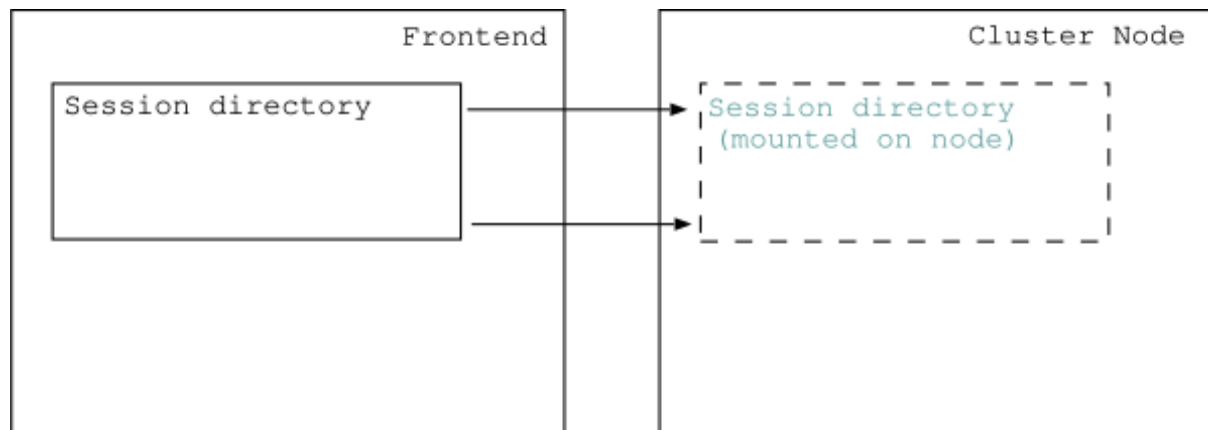


Fig. 6.4: Sessiondir is shared between ARC CE and WNs. No local scratchdir defined.

```
RUNTIME_NODE_SEES_FRONTEND='yes'
```

Compute inside a local WN scratch directory

Session directory is shared, WN scratch directory is not shared

For I/O performance reasons it is possible to perform computations inside a local directory on the WN.

In this case *job scratch directory* is created in the configured local directory and files from the *job session directory* are **moved** to job scratch before execution starts (a).

Only files representing the job's *stdout* and *stderr* are placed in the original *job session directory* and soft-linked in scratch (b).

After the execution has completed all output files are moved to the *job session directory* (b) and are then available for clients to download (c).

The move is performed by default using the *mv* command. The admin can change this behaviour using the *movetool* option.

Configuration:

```
[lrms]
movetool = rsync -av

[arex]
shared_filesystem = yes
scratchdir = /mnt/scratch/arc
```

Variables set in the wrapper script:

```
RUNTIME_NODE_SEES_FRONTEND='yes'
RUNTIME_LOCAL_SCRATCH_DIR='/mnt/scratch/arc'
RUNTIME_LOCAL_SCRATCH_MOVE_TOOL='rsync -av'
```

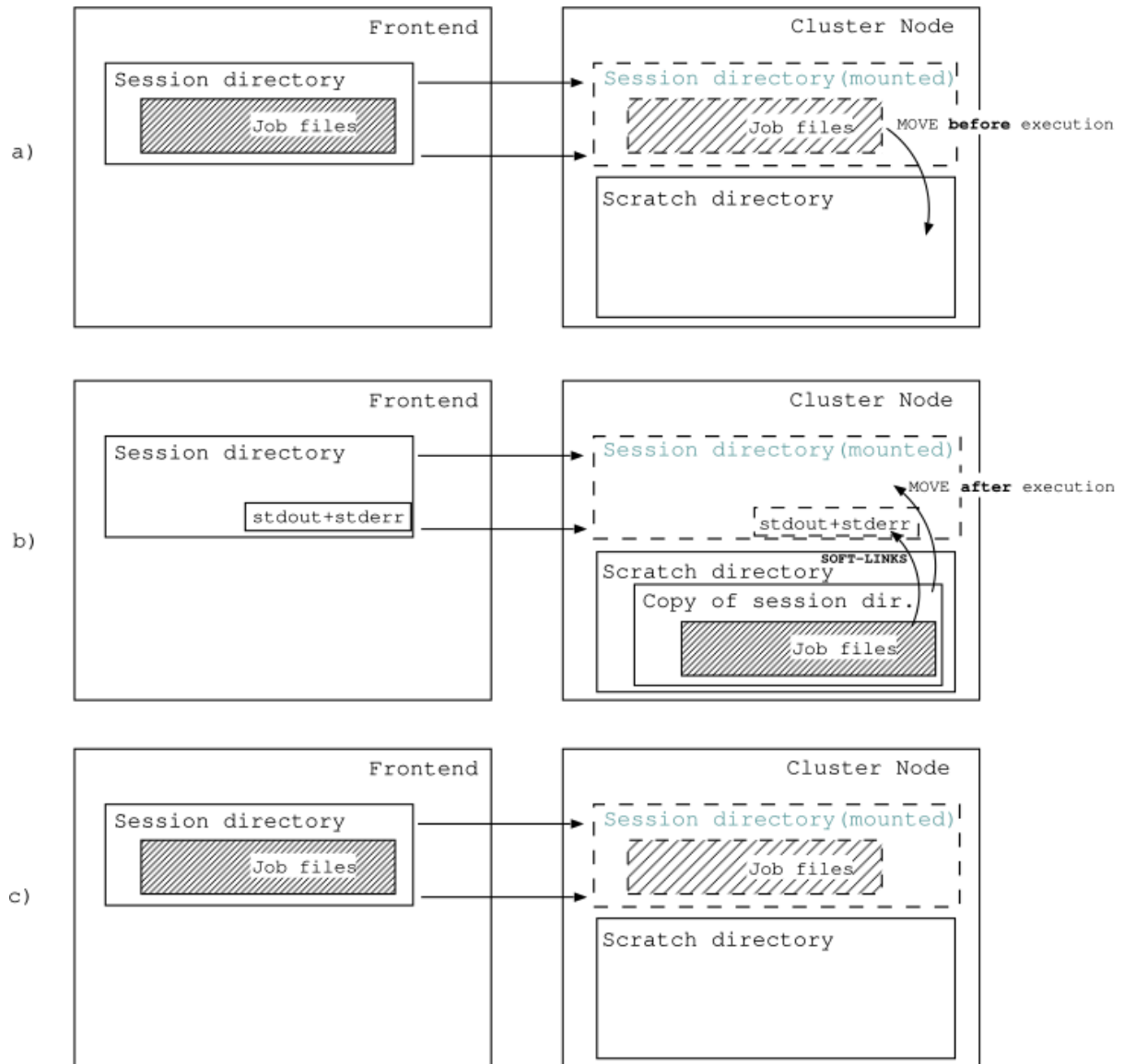


Fig. 6.5: Sessiondir is shared between ARC CE and WNs, scratchdir is defined and available only on nodes.

Session directory is NOT shared, WN scratch directory is NOT shared

If the session directory is not shared the data movement between ARC CE and WN can be done by one or a combination of these two:

- **means of LRMS backend**

How to implement the data movement between the ARC CE and WN depends on the particular batch system backend used.

This is not described here, please refer to your batch system manuals.

For example, in *PBS* this corresponds to `#PBS -W stagein` options. In *SLURM* it could be custom *prolog* and *epilog* scripts.

- **A custom made RTE**

This topic is currently not part of this guide, but it is implemented in some ARC setup, so please ask for support from the community.

Warning

This setup is not recommended, and is not actively supported by ARC. You can as always reach out to the ARC community for help and support if you need this setup.

Two known use cases of setting up ARC to work with such setup are described below.

Use-case 1: Static path to WN scratch directory

In this use case the path of the scratch directory on every WN is a static path that does not change depending on the job.

ARC will create a wrapper script that expects all the job files to be on the WN in the path `$RUNTIME_LOCAL_SCRATCH_DIR/<arc-job-id>`.

Warning

In this use case, when the *job session directory* is not shared (`shared_filesystem = no`) *job scratch directory* MUST BE defined (`scratchdir = path`) to instruct LRMS where to find the files. Otherwise job submission to LRMS will fail.

After all input files are gathered in the *job session directory* on ARC CE, the LRMS or the custom RTE copies files to the *job scratch directory* on WN (Figure a)).

The job performs all I/O using local *job scratch directory*. After execution all declared output files (including *stdout* and *stderr*) must be staged out to the *job session directory* on ARC CE by LRMS or custom RTE (Figure b)).

When the output job files are available in the *job session directory* on ARC CE they are ready to be uploaded to external storage elements or be downloaded by the user (Figure c)).

Configuration:

```
[arex]
shared_filesystem = no
scratchdir = /mnt/scratch/arc
```

Variables set in the wrapper script:

```
RUNTIME_NODE_SEES_FRONTEND=' '
RUNTIME_LOCAL_SCRATCH_DIR='/mnt/scratch/arc'
```

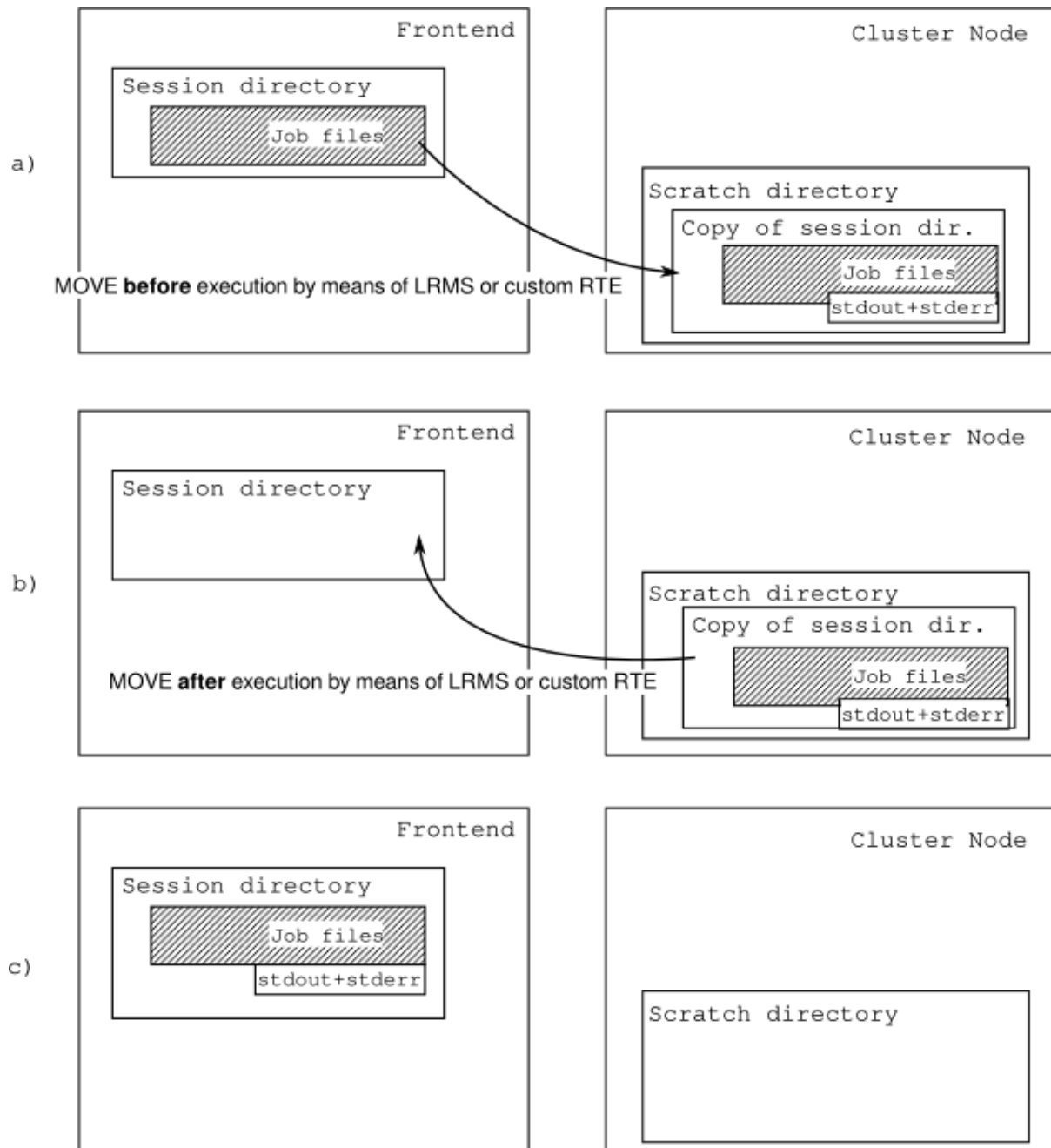


Fig. 6.6: NEITHER sessiondir NOR WN scratchdir are shared between ARC CE and WNs.

Use-case 2: Dynamic job path to the WN scratch directory

In this use case the LRMS generates a dynamic path for each submitted job. This path is not known to ARC during the creation time of the wrapper script, therefore it must be obtained in some way at runtime, after the job is submitted.

A sysadmin may already have defined such variable in their prolog/epilog scripts.

The workflow is very similar to the one described in the picture of the static path use case, with the difference that the **name** of the LRMS specific environment variable that contains the dynamic path can be defined and configured in a special *RTE* called *ENV/LRMS-SCRATCH*, so that it will be used at runtime to generate a path of the kind `<lrms-job-id>/<arc-job-id>`.

The sysadmin can now create a custom RTE with LRMS specific commands and make use of the variable above when referencing the dynamic dir.

In this scenario the `arc.conf` variable *scratchdir* can be defined but it is not mandatory, as the LRMS may have a completely different base path for each node/job.

However, if configured, the `scratchdir` path will be prefixed to the dynamic path, as in: `<arc-conf-scratchdir-var>/<lrms-job-id>/<arc-job-id>`

Please refer to the documentation of the bundled *ENV/LRMS-SCRATCH* runtime environment for more details.

Configuration:

```
[arex]
shared_filesystem = no
```

Variables set in the wrapper script:

```
RUNTIME_NODE_SEES_FRONTEND=' '
RUNTIME_LOCAL_SCRATCH_DIR=' '
```

Accounting with legacy JURA

WARNING: This component was deprecated in ARC 6.4 and completely removed in ARC 6.8!

Warning

Information in **this chapter is relevant only for 6.0-6.3 ARC releases.**

Starting from ARC 6.4 release the *Accounting Subsystem* with local accounting database will be introduced. Make sure you are reading the documentation that match your ARC CE release version.

The *Job Usage Reporter of ARC* (JURA) is a component which is capable to create standard-compliant usage records from job usage information provided by the A-REX (*Job Log* files) and send the records to remote accounting services.

JURA is capable of creating two types of usage records from the job log files:

- Usage Record 1.0 (UR) XML format to be sent to an SGAS LUTS (Logging and Usage Tracking Service)
- Compute Accounting Record (CAR) XML format to be sent to APEL

Overview

Figure shows the overview of accounting workflow in ARC. More details can be found in [Legacy JURA Accounting Technical Details](#).

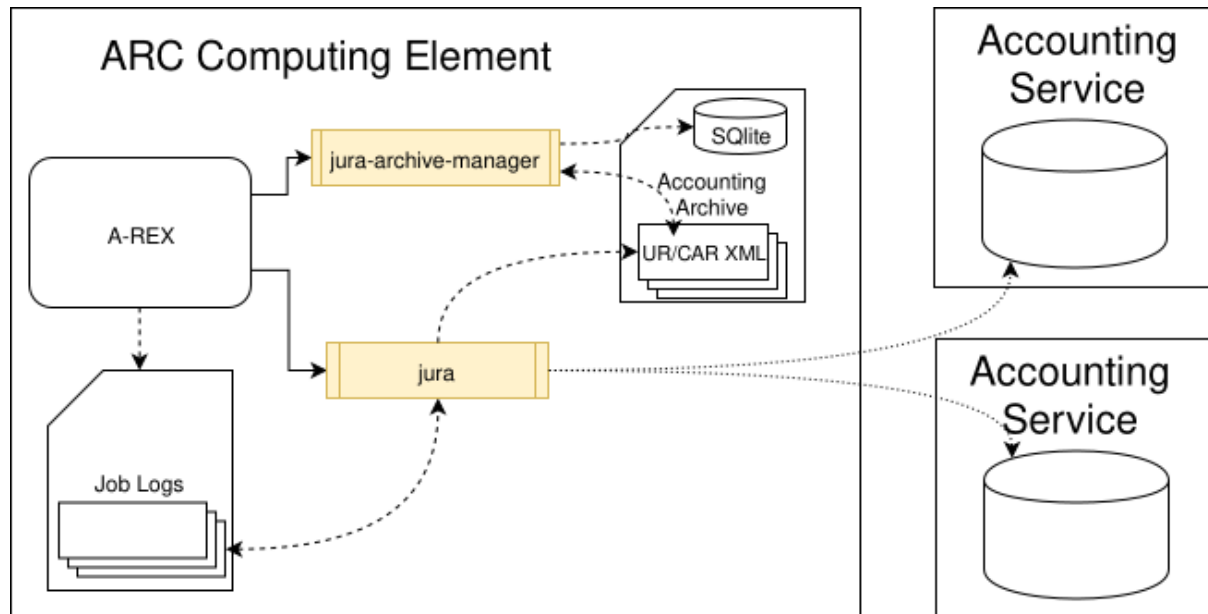


Fig. 6.7: Accounting records processing in ARC CE: creating, processing, publishing and archiving

When accounting is enabled, the sequence of events is the following:

1. A-REX writes the per-job accounting records to *Job Logs* directory (logs directory inside the *controldir*).
2. A-REX periodically runs *jura*. Default and minimum period is hourly but can be increased with *urdelivery_frequency* option.
3. Jura reads jobs accounting information from *Job Logs* and depending on the targets configured in *arc.conf* creates accounting records according to required format.
4. Jura publishes accounting records to the configured accounting services.
5. If records archiving is enabled, JURA also puts the records into Accounting Archive
6. When archiving functionality is used, A-REX also runs the *jura-archive-manager* process which manages the archive structure and indexes the records in the local database. The local accounting database can then be used to *lookup statistics* or to *republish data*.

Enabling accounting records reporing

JURA is a stand-alone binary application that is a part of the A-REX package and can be enabled with *[arex/jura]* block.

JURA has a dedicated log file as defined by the *logfile* option. Log rotation is set for the default */var/log/arc/jura.log* location. Accounting logs can be also viewed with *arcctl*:

Changed in version 6.4: Extra legacy argument is required to access legacy accounting subsystem functionality. Use *arcctl accounting legacy logs*.

```
[root ~]# arcctl accounting logs
```

Accounting services for sending the records are configured with dedicated sub-blocks. You need to define a separate block with an unique targetname for every target server used.

The usage record of each job will be reported to all of the destinations, unless the `vofiler` option is configured for some of them to filter records by VO name.

Configuring reporting to SGAS

The *SGAS sub-block* enables and configures an SGAS accounting server as a target destination to which JURA will send properly formatted Usage Record 1.0 (UR) XML records.

The *targeturl* option is the only mandatory parameter to configure an SGAS target. In the specific setup cases you can also apply *VO filtering* and *set prefix* for local job IDs.

Example:

```
[arex/jura/sgas: NeIC]
targeturl = https://grid.uio.no:8001/logger
urbatchsize = 80
```

Configuring reporting to APEL

The *APEL sub-block* enables and configures an APEL accounting server as a target destination to which JURA will send properly formatted Compute Accounting Record (CAR) XML records.

The *targeturl* option defines the APEL broker URL to send records to. Unlike the native APEL client that relies on Top-BDII infrastructure in runtime to find the blockers, ARC uses a dedicated target URL in the configuration. However, it is possible to look up available targets with `arcctl`:

```
[root ~]# arcctl accounting apel-brokers
http://mq.cro-ngi.hr:6163/
http://broker-prod1.argo.grnet.gr:6163/
[root ~]# arcctl accounting apel-brokers --ssl
https://mq.cro-ngi.hr:6162/
https://broker-prod1.argo.grnet.gr:6162/
```

SSL or non-SSL connecton should be defined accordingly with the `reference_arex_jura_apel_use_ssl` option.

You also need the *GOCDB name* of the resource and the *APEL topic* to which JURA will publish accounting records. For a correct production accounting setup it is recommended to specify resource benchmarking results.

Example:

```
[arex/jura/apel: EGI]
targeturl = https://mq.cro-ngi.hr:6162
topic = /queue/global.accounting.cpu.central
gocdb_name = RAL-LCG2
benchmark_type = Si2k
benchmark_value = 2625.00
use_ssl = yes
```

Note

Publishing to APEL relies on SSM framework that has own logfile viewable with `arcctl accounting logs --ssm`.

Records archiving

The archiving functionality allows to store generated usage records in a specified directory on the disk.

The `reference_removed_arex_jura_archiving` enables this functionality. Directory path for the archived Jura usage records is defined with the `reference_arex_jura_archiving_archivedir` option.

Once archive is enabled, Jura writes generated usage records to files named `usagerecord[CAR].<jobid>.<random>`. If a A-REX job log file is processed repeatedly – for example because of temporary connection failures to the remote accounting service – and a respective usage record archive file already exists, then the usage record is not generated again. Instead, the content of the archive file is used without change (the creation time stamp is also retained).

Note

Records archive can grow very rapidly on heavy loaded sites with huge amount of jobs coming daily. To limit the archiving time, use the `reference_arex_jura_archiving_archivettl` option.

To better organize archived records and index the accounting data in the local database, the `jura-archive-manager` helper process is run periodically by A-REX when archive is enabled.

The archive manager:

- reads the resource usage data and stores it to an indexed SQLite database file `accounting.db`.
- moves files according to the following directory structure: `YYYY-MM/DD/<jj>/<kkkk>.<random>.{UR|CAR}`, where `jj` is first 2 characters of `<jobid>` and `kkkk` is the rest of `<jobid>`

Lookup local accounting data

Data in the local accounting database can be viewed with `arcctl` for each type of destination. Timeframe of interest and many other filters can be specified, e.g.:

Changed in version 6.3: Extra legacy argument is required to access legacy accounting subsystem functionality. Use `arcctl accounting legacy stats`.

```
[root ~]# arcctl accounting stats -t apel -b 2018-11-01 -e 2018-11-30
Statistics for APEL jobs from 2018-11-01 00:00:02 till 2018-11-29 23:59:13:
  Number of jobs:          355168
  Total WallTime: 48146 days, 18:06:50
  Total CPUTime:  71927 days, 10:33:33
[root ~]# arcctl accounting stats -t apel -v
ops
belle
atlas
ops.ndgff.org
[root ~]# arcctl accounting stats -t apel --filter-vo belle
Statistics for APEL jobs from 2015-04-13 07:28:52 till 2019-04-02 00:25:57:
  Number of jobs:          869400
  Total WallTime: 216866 days, 3:33:33
  Total CPUTime:  193604 days, 19:40:58
```

Republishing records

When something goes wrong with accounting services, network etc, it may be needed to republish local records again.

Warning

Republishing is ONLY possible when *Records archiving* is enabled and records still exist for the required time period.

The *ARC Control Tool* should be used to republish data. It is tightly integrated with the accounting archive database: `arcctl` will look up needed records in the database for the defined republishing period, prepare them for republishing and then invoke `jura` with correct configuration settings.

Warning

It is strongly recommended to AVOID running republishing in ARC5-way, executing the `jura` command manually. Even if you manage to get it working, it can cause unpredictable results.

Define the timeframe of interest and run the command to republish data:

Changed in version 6.4: Extra legacy argument is required to access legacy accounting subsystem functionality. Use `arcctl accounting legacy republish`.

```
[root ~]# arcctl accounting republish -b 2018-12-01 -e 2018-12-31 -s https://sgas.
↳grid.org.ua:8443/sgas
[root ~]# arcctl -d DEBUG accounting republish -b 2019-04-01 -e 2019-04-10 -a https://
↳mq.cro-ngi.hr:6162
```

Next generation ARC accounting: arcctl example queries

General statistic queries examples

```
[root arc-dev]# arcctl accounting stats
A-REX Accounting Statistics:
  Number of Jobs: 238
  Execution timeframe: 2019-07-09 18:42:18 - 2019-07-12 15:16:56
  Total WallTime: 6 days, 1:35:37
  Total CPUTime: 4:19:45 (including 0:00:00 of kernel time)
  Data staged in: 62.1M
  Data staged out: 1.7K

[root arc-dev]# arcctl accounting stats -o
brief          data-staged-in  jobcount        json            walltime
cputime        data-staged-out jobids          users           wlcgvos
[root arc-dev]# arcctl accounting stats -o wlcgvos
moldynggrid
ops
matmoden

[root arc-dev]# arcctl accounting stats --filter-vo ops -o jobcount
222

[root arc-dev]# arcctl accounting stats --filter-vo ops --start-from 2019-07-11
```

(continues on next page)

(continued from previous page)

```
A-REX Accounting Statistics:
  Number of Jobs: 131
  Execution timeframe: 2019-07-11 00:16:08 - 2019-07-12 15:33:26
  Total WallTime: 0:20:38
  Total CPUTime: 0:02:02 (including 0:00:00 of kernel time)
  Data staged in: 219.3K
  Data staged out: 1.0K

[root arc-dev]# arcctl accounting stats --filter-queue grid \
> --filter-user /DC\=org/DC\=ugrid/O\=people/O\=KNU/CN\=Andrii\ Salnikov
A-REX Accounting Statistics:
  Number of Jobs: 8
  Execution timeframe: 2019-07-09 18:42:18 - 2019-07-10 17:31:59
  Total WallTime: 0:01:20
  Total CPUTime: 0:00:00 (including 0:00:00 of kernel time)
  Data staged in: 61.7M
  Data staged out: 0

[root arc-dev]# arcctl accounting stats --filter-endpoint org.nordugrid.arcrest \
> --end-till 2019-07-10 -o data-staged-out
75

[root arc-dev]# arcctl accounting stats --filter-state failed --filter-vo moldyngrid_
↪ -o data-staged-in
21574172

[root arc-dev]# arcctl accounting stats --filter-extra rte ENV/PROXY -o json | jq .
{
  "stageout": 0,
  "rangeend": 1562938441,
  "count": 86,
  "cpukerneltime": 0,
  "users": [
    "/DC=EU/DC=EGI/C=HR/O=Robots/O=SRCE/CN=Robot:argo-egi@cro-ngi.hr",
    "/DC=EU/DC=EGI/C=GR/O=Robots/O=Greek Research and Technology Network/
↪ CN=Robot:argo-egi@grnet.gr"
  ],
  "wlcgvos": [
    "ops"
  ],
  "cpuusertime": 0,
  "cputime": 0,
  "rangestart": 1562688395,
  "stagein": 398180,
  "walltime": 172
}

[root arc-dev]# arcctl accounting stats --filter-extra jobname "test 04"
A-REX Accounting Statistics:
  Number of Jobs: 4
  Execution timeframe: 2019-07-10 13:40:00 - 2019-07-10 17:31:59
  Total WallTime: 0:00:40
  Total CPUTime: 0:00:00 (including 0:00:00 of kernel time)
  Data staged in: 30.9M
  Data staged out: 0
```

(continues on next page)

(continued from previous page)

```
[root arc-dev]# arcctl accounting stats --filter-extra jobname "test 04" -o jobids
g3AMDmCCP5unf5481mks8bjnABFKDmABFKDmN9IKDmIBFKDmnjU0ym
nEhMDmvqQ5unf5481mks8bjnABFKDmABFKDmN9IKDmZBFKDtXEYGn
UufLDmnmS5unf5481mks8bjnABFKDmABFKDmmSMKDmNBFKDm9nRnVo
bdoLDmCnS5unf5481mks8bjnABFKDmABFKDmmSMKDmaBFKdMfKldDn
```

Job-specific accounting queries

```
[root arc-dev]# arcctl accounting job events
↳ bdoLDmCnS5unf5481mks8bjnABFKDmABFKDmmSMKDmaBFKdMfKldDn
2019-07-10 17:30:26 ACCEPTED
2019-07-10 17:30:26 PREPARING
2019-07-10 17:30:26 DTRDOWNLOADSTART
2019-07-10 17:30:27 SUBMIT
2019-07-10 17:30:27 DTRDOWNLOADEND
2019-07-10 17:30:29 INLRMS
2019-07-10 17:30:47 LRMSSTART
2019-07-10 17:30:57 LRMSSEND
2019-07-10 17:31:59 FINISHING
2019-07-10 17:31:59 FINISHED

[root arc-dev]# arcctl accounting job transfers
↳ bdoLDmCnS5unf5481mks8bjnABFKDmABFKDmmSMKDmaBFKdMfKldDn
Data transfers (downloads) performed during A-REX stage-in:
  http://download.nordugrid.org:80/packages/nordugrid-arc/releases/6.1.0/src/
↳ nordugrid-arc-6.1.0.tar.gz (from cache):
  Size: 5.2M
  Download timeframe: 2019-07-10 17:30:26 - 2019-07-10 17:30:27
  http://download.nordugrid.org:80/packages/nordugrid-arc/releases/6.0.0/src/
↳ nordugrid-arc-6.0.0.tar.gz (from cache):
  Size: 5.1M
  Download timeframe: 2019-07-10 17:30:26 - 2019-07-10 17:30:27
No stage-out data transfers (uploads) performed by A-REX.

[root arc-dev]# arcctl accounting job info
↳ bdoLDmCnS5unf5481mks8bjnABFKDmABFKDmmSMKDmaBFKdMfKldDn
Job bdoLDmCnS5unf5481mks8bjnABFKDmABFKDmmSMKDmaBFKdMfKldDn accounting info:
=====
Job description:
  Job was submitted at 2019-07-10 17:30:26 via "org.ogf.glue.emies.activitycreation"
↳ interface using "https://arc.matmoden.kiev.ua:443/arex" endpoint.
  Job owned by "/DC=org/DC=ugrid/O=people/O=KNU/CN=Andrii Salnikov" as a member of
↳ "moldyngrid" WLCG VO.
  It was targeted to the "grid" queue with "381309.head.energrid.ipme.kiev.ua" LRMS
↳ ID.
  Job completed with exit code 0 at 2019-07-10 17:31:59.
  Following job properties are recorded:
    Localuser: prdmdg24
    Clienthost: 130.235.185.244:58458
    Lrms: pbs
    Nodenames: n2
    Jobname: test 04
Resource usage:
  Execution timeframe: 2019-07-10 17:30:26 - 2019-07-10 17:31:59
```

(continues on next page)

(continued from previous page)

```

Used WallTime: 10
Used CPUTime: 0 (including 0 of kernel time)
Used WN Scratch: 0
Max physical memory: 528
Max virtual memory: 0
Used CPUs: 1 on 1 node(s)
Data staged in: 0
Data staged out: 0
Used RunTime Environments:
  There are no RTEs used by the job.
Auth token attributes provided:
  VOMS FQAN: /moldyngrid
  VOMS FQAN: /moldyngrid/Role=VO-Admin
  VOMS FQAN: /moldyngrid/Role=production
  VOMS FQAN: /moldyngrid/edu

```

About benchmarks and accounting publishing

After the changes in the accounting system in ARC 6.4.0 there were some issues related to missing benchmark values in the accounting records. Some of them are related to a bug that unfortunately snuck into the codebase, however sites can have issues with benchmarks for other reasons too.

This page is aimed to clarify how benchmarks are recorded and propagated, in what situations problems can occur, and how to fix them.

If you see `HEPSPEC 1.0` is being used in the `jura.log` - (some of) your job records are missing benchmark values.

Follow the questions to clarify your case.

Which version of ARC you have?

Depending on the ARC version, there are several issues related to the benchmark values processing:

- **ARC < 6.4.0:**
 - completely different accounting codebase is in use, information in this document is irrelevant. General advise is definitely to update to the recent version, bugs in the old codebase will not be fixed.
- **ARC < 6.5.0:**
 - bug with handling benchmark values in the publishing code
 - HTCondor backend with non-shared filesystem have missing benchmarks
 - APEL summaries query performance is degrading with increasing amount of stored records
- **ARC < 6.8.0:**
 - HTCondor backend with non-shared filesystem have missing benchmarks
 - APEL summaries query performance is degrading with increasing amount of stored records
- **ARC >= 6.8.0:**
 - all known benchmark related issues are fixed
 - if you see `HEPSPEC 1.0` is being used there is some valid reason for it, including missconfiguration

Note

It is important to understand that benchmark values **are part of the job accounting record**. Benchmark data in the job accounting record is **defined on the job start time** and stored when job is just finished.

If you have jobs started before update of ARC or configuration fix (depending on your case) - you need to manually fix already stored records. There is no way changes will be applied retrospectively to already stored records automatically.

What are the reasons for missed benchmark value in the job records?

There are several valid reasons when you will see `HEPSPEC 1.0 is being used` message:

1. The job was started when ARC was at version < 6.5
2. The job was started when the `[queue:name] block` in `arc.conf` had no proper benchmark defined
3. The permissions or other issues (including HTCondor backend bug with non-shared filesystem in ARC < 6.8.0) prevents the writing of `.diag` files on the worker nodes
4. The job failed in LRMS before even execution of initial jobscript wrapper part (node failure, etc).

The last issue is simply can happen very rarely and nothing to do with it, but such jobs has zero cputime, so benchmark is really irrelevant.

Nevertheless, to eliminate the log message that annoys admins and avoid additional type of summary records during the publishing, ARC 6.8.0 introduced `benchmark` option in `[lrms] block` that will be used as a fallback if the benchmark metric is missing in the job data.

Warning

Again! The `benchmark` option in `[lrms] block` has no influence on already stored records. It is during **storing** time but **NOT publishing** time.

How to fix missing benchmark values manually?

Already stored accounting records that has no benchmark values can be fixed by issuing an `sqlite` query that adds benchmark value. Following example assumes the controldir is the default `/var/spool/arc/jobstatus`:

```
[root ~]# sqlite3 /var/spool/arc/jobstatus/accounting/accounting.db "insert into
↪ JobExtraInfo
(RecordID, InfoKey, InfoValue)select distinct RecordID, 'benchmark', 'HEPSPEC:12.1'
from JobExtraInfo where RecordID not in
(select RecordID from JobExtraInfo where InfoKey='benchmark');"
```

If you discover that some records use the default benchmark of `HEPSPEC 1.0` instead of your desired benchmark value in `arc.conf` (e.g. you had added benchmark values after job start) you can update the values as well:

```
[root ~]# sqlite3 /var/spool/arc/jobstatus/accounting/accounting.db "update
↪ JobExtraInfo
set InfoValue = 'HEPSPEC:12.1' where InfoKey = 'benchmark' and InfoValue =
↪ 'HEPSPEC:1.0';"
```

What should I know to avoid running into bechmark issues?

To understand how the HEPSPEC 1.0 is being used occurs in the `jura.log` there are 3 points to understand:

1. JURA is only the publisher and it sends the data about the jobs stored in the local ARC accounting database. NO values from `arc.conf` (apart from where to publish records) are used during publishing.
2. Info about benchmarks is part of the job accounting data stored in the ARC local accounting database when the job is in the finishing state. Moreover, the static data, including the benchmark defined in `arc.conf` are defined during the jobscript generation (job start time). Any update to `arc.conf` AFTER the job start HAS NO EFFECT on already stored records.
3. In case of publishing to APEL, the default method to use is *APEL summaries*. This means that `jura` will send (update) the total counters about last 2 month of data that aggregated per VO, DN, Endpoint (include queue) and Benchmark! CONSEQUENTLY if any single job within 2 month timeframe is missing the benchmark data - this warning about using HEPSPEC 1.0 will be there!

Warning

For ARC < 6.8.0 the APEL summary query includes grouping by benchmark which was out of scope of the initial ARC accounting database design. The extra tables join is harmful to performance on heavy loaded sites! The recommended mitigation to save ARC CE CPU cycles is to go back to individual usage records publishing with `apel_messages = urs` option.

In the ARC 6.8.0 the APEL summary querying were improved and performance hit is not that valuable. You can use summaries on the heavy loaded sites as well.

It is also important to understand the **chain of benchmark propagation** for the issues troubleshooting:

1. The value of *benchmark* defined in the `[queue:name] block` in `arc.conf` is written to the `.diag` file as it is on the frontend (controldir).
2. The `.diag` file from the control directory is copied next to the job's session directory and either shared to worker node (shared sessiondir case) or moved by LRMS. See more details about shared vs non-shared sessiondir in the *Job scratch area* document.
3. During job execution jobscript writed data to `.diag` on the worker node. This includes benchmark that can be redefined in runtime (e.g. by *RunTime Environments in ARC*)
4. After job completion the `.diag` from worker node is moved to the frontend's session directory if sessiondir is not shared.
5. On the frontend `.diag` from session directory merged with `.diag` in the control directory and more information from the LRMS accounting is added to it.
6. A-REX parse the `.diag` in the control directory and store data to the database. From ARC 6.8.0 at this stage the default fallback *benchmark* is added to the data from `arc.conf` if missing in the `.diag`.

So, should I do something if I see "HEPSPEC 1.0 is being used" message?

If this is a rare single job that just failed in LRMS before writing the accounting data - nothing to worry about.

But if it annoys you, you can fix even single job data manually as describer above. Or starting from ARC 6.8.0 you can define the fallback *benchmark* to aviod it completely.

To identify how many jobs are missing benchmark data in the database, run the following query:

```
[root ~]# sqlite3 /var/spool/arc/jobstatus/accounting/accounting.db "select JobID_
↪from AAR
  where RecordID not in ( select RecordID from JobExtraInfo where InfoKey='benchmark
↪');"

```

This returns list of the job IDs with missing benchmark data. Than you can use:

```
[root ~]# arcctl accounting job info <JobID>
```

to find what are those jobs.

If there are many, than something definitely goes wrong and you should:

1. Check if you are facing the known issues if you are running ARC < 6.8.0. ARC update + manual records fix will solve your problem in this case.
2. Check the `arc.conf` syntax in respect to benchmark. It should be defined in the `[queue:name] block` and use either HEPSPC or si2k. Manual records fix for already stored records is needed anyway.
3. Check the `.diag` file contains information, use `arcctl` to check stored data, check A-REX logs for any hints.
4. Open a [bugzilla ticket](#) if nothing helps.

ARC CE Data Staging and Caching

ARC Information System

TBD

6.2.7 ARC tutorial

In this tutorial we will take an existing HPC cluster, install and configure ARC to become a production grid site, and test that it is working by submitting some test-jobs from a remote ARC client server.

Note

This is a tutorial with the not yet released ARC 7 and with using tokens. ARC still works as before with x509 user certificates, the differences will be pointed out.

Prerequisites

There are many ways to set up ARC for production depending on the infrastructure you are installing ARC in front of. Sites will need different configuration options as a result. For a detailed overview of getting a production ready site the [ARC installation and configuration guide](#) is a good source of reference, in addition to the [ARC configuration reference document](#).

In this tutorial we will set up a site with a set of recommended and commonly used configuration options.

To follow this tutorial you should have a functional cluster with the following required components:

Requirements

- Already installed and working batch system (in this tutorial we will focus on SLURM or HTCondor, but other supported batch systems are also ok)
- One server for ARC installation - this server could also host the batch system (be the SLURM master or HTCondor CM)
 - The ARC server must be able to submit jobs to the batch system
 - The ARC server must have an x509 host certificate
- Minimally 1 compute node

- Shared filesystem between the ARC server and the compute nodes
- Optional: a separate server for the SLURM master/HTCondor CM

Requirements for tutorial part 3

- Minimally 1 server for remote datadelivery service

Tutorial overview

- Part 0: *Install ARC nightly build repo*
- Part 1: *Zero-conf test-setup* (from Step 2)
- Part 2: *Production ready setup with LOCAL ARC datadelivery*
- Part 3: *Production ready setup with REMOTE ARC datadelivery service*
- Part 4: Brief mention of the powerful *statecallout* and *RunTimeEnvironment* functionality that comes with ARC.
- Part 5: *Create a small community grid with ARCHERY*

Nordugrid ARC nightly build repo setup for ARC tutorial

If you are using one of the test-infrastructures set up in connection to the tutorial, you can follow these steps in this document to set up the nightly repo.

If you are using some other infrastructure, please go to the [official nightly repo page](#) for instructions.

Enable ARC7 nightly builds for almalinux (rocky)

Setup a CRON job (e.g. `/etc/cron.daily/update-arc-nightly-time.sh`) to fetch the latest nightlies date. As root user do:

```
cat << EOF > /etc/cron.daily/update-arc-nightly-time.sh
#!/bin/bash
curl -s http://builds.nordugrid.org/nightlies/nordugrid-arc/next/ | sed -n 's/^\.*<a.*>
↪\(..*\)\</a>.*$/\1/p' | sort | tail -1 > /etc/yum/vars/arcnightly
EOF
```

Run the CRON script once manually* to have the `arcnightly` variable initialized. Also, make sure the script has executable permissions:

```
sudo chmod +x /etc/cron.daily/update-arc-nightly-time.sh
sudo /bin/bash /etc/cron.daily/update-arc-nightly-time.sh
```

To add the nightlies repository to your RHEL-based system, create a `/etc/yum.repos.d/nordugrid-nightly.repo` with the following content. As root user do:

```
cat << EOF > /etc/yum.repos.d/nordugrid-nightly.repo
[nordugrid-nightly]
name=Nordugrid ARC Next Nightly Builds - \${basearch}
baseurl=http://builds.nordugrid.org/nightlies/nordugrid-arc/next/\${arcnightly}/rocky/9/
↪\${basearch}
enabled=1
gpgcheck=0
EOF
```

Enable powertools/crb for dependencies

For RHEL8 flavour

```
sudo dnf install -y epel-release
sudo dnf config-manager --set-enabled powertools
```

For RHEL9 flavour:

```
sudo dnf install -y epel-release
sudo dnf config-manager --set-enabled crb
```

Navigation

- Go back to the *Overview*

Tutorial Part 2 - Production setup with local ARC datadelivery

Tutorial configuration values

In the “zero-conf” setup in the [previous step](#) no other user than your self can submit jobs to ARC, and the job will only run as a fork. There are minimally 5 steps to configure ARC for a production ready server. These are:

- Step 1. *Tell ARC what LRMS system you are using*
- Step 2. *Tell ARC what are the paths to the session directory, control directory and compute scratch directory*
- Step 3. *Tell ARC what LRMS queues to use*
- Step 4. *Enable ARC data-staging service and tell ARC what cache directory to use*
- Step 5. *Set up the authentication and authorization including external to local user mapping rules*

That will be it. And we will finally test the setup by *submitting a couple of test-jobs*.

We will be using some example ARC configuration parameter values suited to fit with the test-infrastructure set up in relation to the EGI 2023 ARC tutorial. You will naturally have to replace these values with the actual values for your specific system. Here is an overview of these for convenience:

- ARC session directory: `/grid/session1`
- ARC control directory: `/grid/control`
- ARC cache directory: `/grid/cache1`
- batch system (LRMS): `slurm` or `condor`
- name of batch system queues: `arc-tutorial` and `wlcg`
- compute node scratch directory: `/scratch`

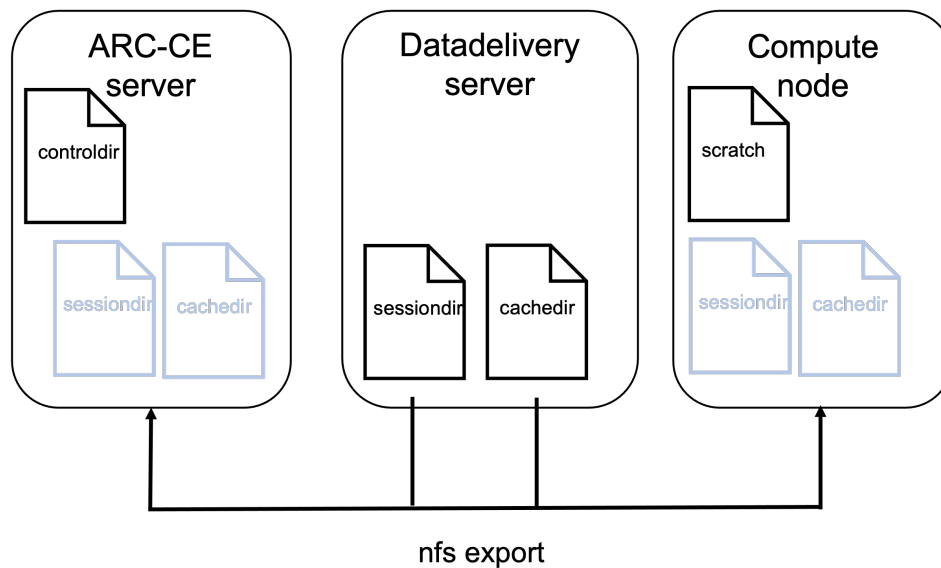


Fig. 6.8: Overview of server and file system for ARC example cluster

Step 1. Tell ARC what LRMS system you are using

In this tutorial we either use slurm or the htcondor batch systems.

Simply replace the “zero-conf” fork lrms with `slurm` or `condor` like this:

```
[lrms]
lrms = slurm
```

Or if you are on HTCondor:

```
[lrms]
lrms = condor
```

ARC supports many other batch-systems - for details see `arc.conf` reference: `[lrms]`

Step 2. A-REX directories - `[arex]` block

In a production system we would want to decide where the ARC related files and folders should live. The relevant A-REX directories are at least the control-directory, the session-directory and the cache directory for data-delivery service. We will address the two first ones here, and the cache directory in [Step 4](#)

Controldir

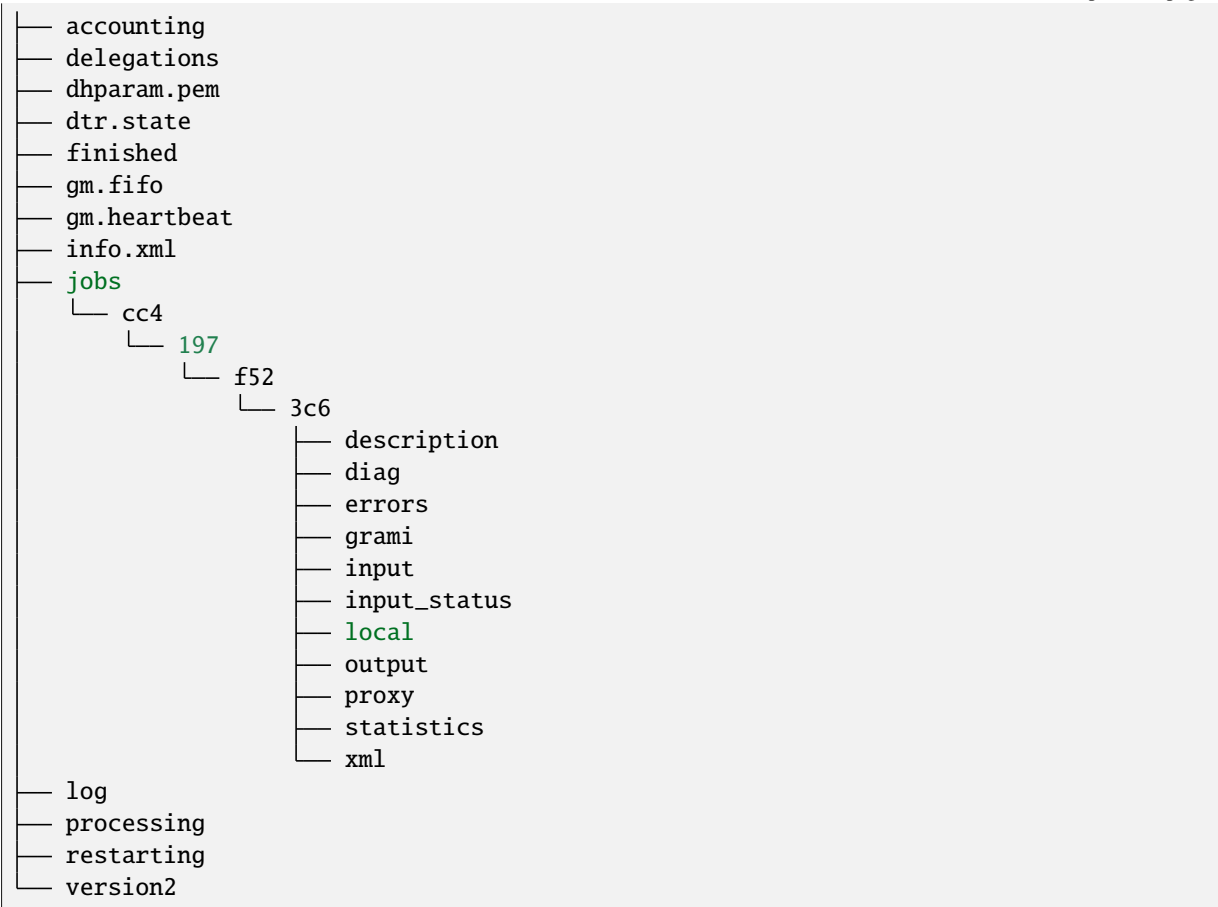
The control directory holds A-REX internal job metadata files and is local to the ARC-CE (i.e. not shared with the compute nodes). The default value is `/var/spool/arc/jobstatus`. We will instead point to `/grid/control` in this example, which correspond to the folders prepared on the test-servers if you are following along as part of an ARC workshop. If you are setting up your clusters you use the dedicated path for the ARC control directory that you have prepared (or just keep the default). ARC will automatically create the directory if it does not already exist. There can only be one single control directory.

Example of the control directory file tree that gets created by ARC-CE:

```
/grid/control/
├─ accepting
```

(continues on next page)

(continued from previous page)

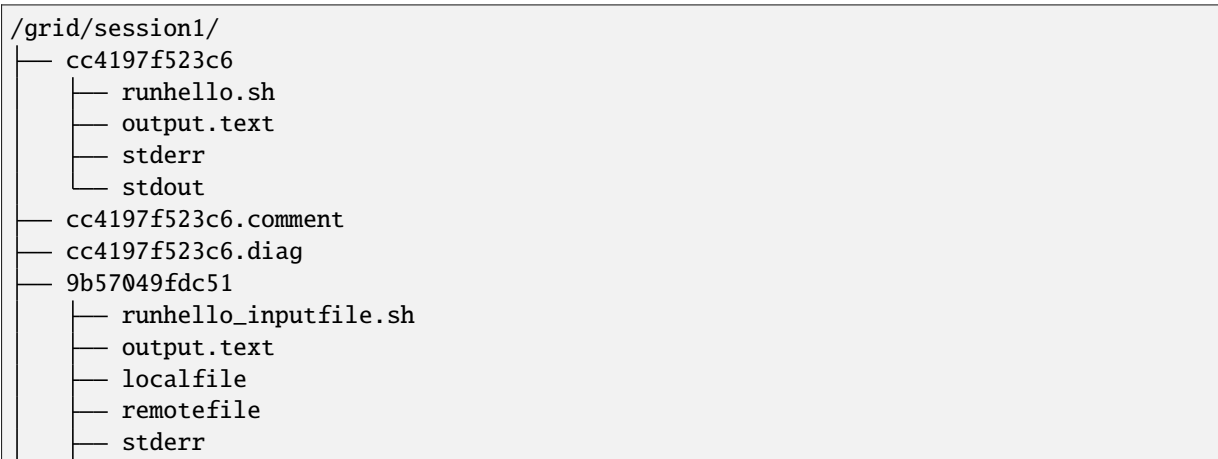


Visit [Internal files of A-REX](#) for more details.

Sessiondir

The jobs working directories are created as subdirectories of the ARC session directories. Each job has its dedicated subfolder session directory identified with the ARC job id. They hold all the input, output, executables and log-file for the job. There can be multiple ARC session directories. The default value is `/var/spool/arc/sessiondir` but we will change it to `/grid/session1` in this example to demonstrate a non-default value. ARC will automatically create the directory if it does not already exist.

Example of the file tree of the session directory where you see two jobs present with their corresponding working directories and files.



(continues on next page)

(continued from previous page)

```
|
|   └─ stdout
└─ 9b57049fdc51.comment
```

This is how the arex block now looks after adding our systems directories:

```
[arex]
controldir = /grid/control
sessiondir = /grid/session1
```

For more details, see `arc.conf` reference: *controldir* and *sessiondir*.

Configure a local compute scratch directory

To avoid excessive traffic on the shared filesystem it is advised to configure a local scratch directory on the compute nodes. If a scratch directory is configured, A-REX will move the jobs files from the shared session directory onto the local scratch directory on the compute node once the job starts. The only file that is not moved is the jobs log-file, this is simply soft-linked to the compute-node. Note that if the ARC cache is used, the input-file links are simply moved to the scratch directory. The outputfiles are created in the compute nodes scratch directory while the job is running.

Once the job is finished, the files including the output-files are moved back into the jobs session directory.

We configure ARC for using a local scratch directory on the compute nodes by adding the `scratchdir` directory path to the `[arex]` block. In this example the path used is `/scratch`. The `[arex]` block now look like this:

```
[arex]
controldir = /grid/control
sessiondir = /grid/session1
scratchdir = /scratch
```

For more details about the configuration option, see `arc.conf` reference: *scratchdir*. In addition, for details on how the job scratch area works, see: *Job scratch area*.

Step 3. Tell ARC what LRMS queues to use - [queues] block

Each queue on the cluster that should receive jobs from ARC must be represented and described by a separate queue block. The `queue_name` should be used as a label in the block name. In case of fork, or other LRMSes with no queue names, just use any unique string. In the case of SLURM the `queue_name` must match an existing SLURM partition.

There can be several queues, for instance a queue for the ordinary grid jobs, and a separate queue for low priority preemptable grid jobs, or for instance a separate queue for different experiments.

The queues can also be configured to accept jobs from certain *authgroups*, which we will see in Step 5-3. In the example shown here, all *authgroups* can use all queues as no *authgroups* are specified.

```
[queue: arc-tutorial]
comment = Queue for ARC test jobs

[queue:wlcg]
comment = Queue for WLCG jobs
```

Consult the `arc.conf` reference *[queue]* for more details.

Step 4. Enable ARC data-staging service and tell ARC what cache directory to use

With the ARC data-staging enabled, ARC will make sure to fetch any remote input files to the ARC server, and the job will be forwarded to the LRMS system only once all files are downloaded.

The ARC cache is a must when you have data-staging enabled. The datadelivery service will check the ARC cache to see if the file is already downloaded, and if it is, will not perform a new download, but simply link the file to the current job using it.

Example of the file tree of the cache directory is shown below. We see the `data` folder holding the actual data files, in addition to the `joblinks` folder identified by the ARC id where the hardlinks to the cached data files are stored.

```
/grid/cache1/
├── data
│   ├── 15
│   │   ├── ed291ec90b64fe7c153c74f65af09f0db2b70e
│   │   └── ed291ec90b64fe7c153c74f65af09f0db2b70e.meta
│   └── 31
│       ├── b9f10eb0968c50246534619de8c78e42a46156
│       └── b9f10eb0968c50246534619de8c78e42a46156.meta
└── joblinks
    ├── cc9197f523c6
    └── remotefile
```

To configure the cache and the data-staging, right below the `[arex]` block add the follow (close to minimal) configuration:

```
[arex/cache]
cachedir = /grid/cache1

[arex/data-staging]
logfile=/var/log/arc/datastaging.log
```

Note

We enable the central `datastaging.log` here, otherwise, the datataging logs will only reside on any remote data delivery servers, and it is convenient to have them accessible from the ARC-CE itself.

It is important also to set up cache-cleaning by enabling the see [\[arex/cache/cleaning\]](#) block. However, for the sake of simplicity, we will not do that in this tutorial.

Check the [\[arex/cache\]](#) and [\[arex/data-staging\]](#) for details.

Step 5. Set up authentication and authorization

The authentication and authorization consists of 3 steps

- Step 5-1: Define authentication groups so-called authgroups
- Step 5-2: Define mapping rules
- Step 5-3: Enable authorization of authgroups to services and queues

Note

Rules in an authgroup are processed in order of appearance. The first matching rule decides the membership of the user to the authgroup being evaluated and the processing STOPS WITHIN that authgroup. A user can

be matched to more than one authgroup. For the mapping rules, the first successful match is used by default, but this can be configured.

Confer with the `arc.conf` reference for details on the `[authgroup]` and the `[mapping]` configuration blocks.

Step 5-1: Define authentication groups so-called authgroups - `[authgroups]` block

Now that the configuration matches the cluster setup with batch system type, folder structure and datastaging, it is time to set up authentication and authorization, in order for the right users to be allowed to submit jobs to the cluster. At the moment only you are allowed to submit jobs according to ARC as result of the “zero-conf” setup.

The authentication and authorization functionality in ARC is very powerful and you can define the authgroups as fine- or coarse-grained as needed. Here we will show a few examples of some possibilities.

Up until now we only have the authgroup connected with the test-jwt with corresponding access rules in the `/etc/arc.conf.d/10-testjwt-<hash>.conf` file, our example was:

```
[root ~]# cat /etc/arc.conf.d/10-testjwt-a7374e17.conf
[authgroup:testjwt-a7374e17]
authtokens = * https://arc.example.org/arc/testjwt/8b7baf79 * * *

[mapping]
map_to_user = testjwt-a7374e17 nobody:nobody

[arex/ws/jobs]
allowaccess = testjwt-a7374e17
```

For the purpose of this tutorial we will keep using the access configuration produced by test-jwt, and also the host certificate issued by the test-CA. But we will make a change to the mapping: change the mapped user from the nobody to arcuser (or any other user that exists on the whole compute cluster):

```
[mapping]
map_to_user = testjwt-93397e33 arcuser:arcuser
```

Note

In real production you will remove any configuration settings related to the test-CA or test-jwt with `sudo arcctl test-jwt cleanup` and `sudo arcctl test-ca cleanup`

We can add some relevant OIDC token claims (add any other of relevance to you):

```
[authgroup: wlcg_iam]
authtokens = * https://wlcg.cloud.cnaf.infn.it/ * * *

[authgroup: egi_aai]
authtokens = * https://aai.egi.eu/oidc/ * * *
```

In addition to some other relevant authgroups based on x509 and voms (add any other of relevance to you):

```
[authgroup:atlas-admin]
subject=/DC=ch/DC=cern/OU=Organic Units/OU=Users/CN=ddmadmin/CN=531497/CN=Robot:↵
↵ATLAS Data Management

[authgroup:cms]
voms=cms * * *
```

Note

Each authgroup can contain multiple authorization rules and any combination of authorization methods. For instance can the `wlcg_iam` authgroup used as an example here, contain both *authtokens* rule and *voms* rules.

It can be convenient to group the authgroups, an example is shown here:

```
[authgroup:wlcg]
authgroup = wlcg_iam
authgroup = egi_aai
authgroup = atlas-admin
authgroup = cms
```

Another example could be:

```
[authgroup:x509]
authgroup = atlas-admin
authgroup = cms

[authgroup:tokens]
authgroup = wlcg_iam
authgroup = egi_aai
```

In other words - group the authgroups in the way you find convenient for the mapping and the authorisation to services.

Now that we have defined some authgroups which authenticate users coming with either tokens from WLCG or EGI, or the x509 robot user certificate of ATLAS Data Management, or an x509 user certificate that has encoded belonging to the cms VO - we are ready to define the mapping for these different user groups.

Step 5-2: Define mapping rules

With the needed authgroups defined, next we need to map the authgroups to a user on the system. ARC will not allow any job to run as root.

In this tutorial we will map all authgroups to the same user for simplicity. This user must exist on the ARC compute element as well as on the compute nodes (which it does for the test infrastructure if you are attending in the context of an ARC workshop). In this example the linux user and groupname is `arcuser`.

```
[mapping]
map_to_user = wlcg arcuser:arcuser
```

Other popular mapping methods are the *map_to_pool* and *map_to_plugin* methods.

The job will be submitted to the batch system by this user, and hence run on the compute node as this user.

Step 5-3: Enable authorization of authgroups to services and queues

The final step in the authentication and authorization setup is to tell ARC what authgroups are allowed to do what. By using `allowaccess` and `denyaccess` you can control what interfaces and queues a certain authgroup is allowed to use.

In the case no `allowaccess` is added, ARC allows anyone authenticated to access this service/queue.

We will set up some very simple authorization in this tutorial.

```
[arex/ws/jobs]
allowaccess = wlcg
```

We can also specify access of authgroups to particular queues e.g.

```
[queue: arc-tutorial]
comment = Queue for ARC test jobs

[queue:wlcg]
comment = Queue for WLCG jobs
allowaccess = wlcg
```

Note

In the case of SLURM - the SLURM partitions must correspond to the queue names. If a job is submitted without specifying a queue (more on that below), the default queue of the batch system will be used.

Our full arc.conf now looks like this (merging arc.conf with files in arc.conf.d folder)

```
[common]
x509_host_key = /etc/grid-security/testCA-hostkey.pem
x509_host_cert = /etc/grid-security/testCA-hostcert.pem

[authgroup:testjwt-a7374e17]
authtokens = * https://arc.example.org/arc/testjwt/8b7baf79 * * *

[authgroup: wlcg_iam]
authtokens = * https://wlcg.cloud.cnaf.infn.it/ * * *

[authgroup: egi_aai]
authtokens = * https://aai.egi.eu/oidc/ * * *

[authgroup:atlas-admin]
subject=/DC=ch/DC=cern/OU=Organic Units/OU=Users/CN=ddmadmin/CN=531497/CN=Robot:↵
↵ATLAS Data Management
voms = atlas * lcgadmin *

[authgroup:cms]
voms=cms * * *

[authgroup:wlcg]
authgroup = wlcg_iam
authgroup = egi_aai
authgroup = atlas-admin
authgroup = cms

[authgroup:x509]
authgroup = atlas-admin
authgroup = cms

[authgroup:tokens]
authgroup = wlcg_iam
authgroup = egi_aai

[mapping]
map_to_user = wlcg arcuser:arcuser
map_to_user = testjwt-93397e33 arcuser:arcuser
```

(continues on next page)

(continued from previous page)

```
[lrms]
lrms=slurm

[arex]
controldir=/grid/control
sessiondir=/grid/session1
scratchdir=/scratch

[arex/cache]
cachedir=/grid/cache1

[arex/data-staging]
logfile=/var/log/arc/datastaging.log

[arex/ws]
[arex/ws/jobs]
allowaccess=wlcg
allowaccess = testjwt-a7374e17

[infosys]
[infosys/glue2]

[infosys/cluster]
nodeaccess=outbound

[queue: arc-tutorial]
comment = Queue for ARC test jobs

[queue:wlcg]
comment = Queue for WLCG jobs
allowaccess = wlcg
```

Restart ARC after configuration changes

For any configuration changes, ARC must be restarted. Before proceeding, do

```
$ arcctl service restart -a
```

Submit a job using tokens

Get a token

Generate your token, and store it as an environment variable on the ARC client-machine. We will use the ARC test-jwt token in this tutorial. But if you have a real token, you can use that as long as you also configure ARC to accept it - see *step5_auth*.

```
export BEARER_TOKEN=$(arcctl test-jwt token)
```

In this tutorial we will submit two test-jobs, one without any datastaging, and one with datastaging. If you are using the ARC client machine or one of the ARC test-sites set up in context with an ARC workshop, the test-job descriptions are already present in the */home/almalinux/arcjob_tutorial* folder.

Note

If you are on your own client server (and not a client in connection to a live tutorial workshop), please download and untar the files (job description, input-files, executable) - or use your own custom test-jobs. Note that to use your own remote ARC client you must ensure that the ARC-CE accepts the token that you will use, and that the tokens issuer CA is installed on both server and client. If you are using a test-jwt on the client Follow the steps here to set up the required trust between the ARC-CE and ARC client. If you have a real IGTF compliant host certificate on the ARC client, the previous steps deploying the igtf-ca's on the ARC-Ce should be all you need to do: `arcctl deploy igtf-ca classic`.

Submit a job with no inputdata

In this example we are assuming you are submitting the job on the client installed on the ARC-CE server itself. Replace `$(hostname)` with your ARC-CE endpoint if this is not the case.

The executable looks like this:

```
#!/bin/sh
echo "Am sleeping for 60 seconds"
sleep 60
exit
```

And the job-description looks like this:

```
&(executable=runhello.sh)
(jobname=hello_test)
(stdout=std.out)
(stderr=std.err)
(cpuTime="1 hours")
(memory="1000")
(outputfiles=("output.txt" ""))
```

Submit the job (make sure your token is still valid, if not, generate a new one - see [token generation](#)).

```
arcsub -d DEBUG -C $(hostname) hello.xrsl
```

Submit a job with local and remote inputdata

Now let's submit a job that requires ARC to stage data for us. We will use this executable:

```
#!/bin/sh
echo "Computing md5sum for localfile";md5sum localfile>>output.txt;
echo "Computing md5sum for remotefile";md5sum remotefile>>output.txt;
echo "Am sleeping for 60 seconds"
sleep 60
exit
```

And this job-description file:

```
&(executable=runhello_inputfile.sh)
(jobname=hello_test)
(stdout=std.out)
(stderr=std.err)
(cpuTime="1 hours")
(memory="1000")
```

(continues on next page)

(continued from previous page)

```
(inputfiles=
("localfile" "file.txt")
("remotefile" "http://download.nordugrid.org/repos/6/rocky/9/source/updates/repodata/
↪repomd.xml"))
(outputfiles=("output.txt" ""))
```

Submit the job (make sure your token is still valid, if not, generate a new one - see [token generation](#))

```
arcsub -d DEBUG -C $(hostname) hello_inputfile.xrsl
```

Inspecting the job

Let us inspect what happens when the job runs in the production setup

The controldir, sessiondir, cachedir are used as configured - we see files populated as the jobs are running. We also see that the owner of the files in the sessiondir and cachedir is the arcuser as configured.

```
[almalinux@arc.example.org ~]$ ls -lhrt /grid/session1/
total 28K
-rw-----. 1 arcuser arcuser 274 May 24 15:53 9be514eedc51.comment
drwx-----. 2 arcuser arcuser 54 May 24 15:54 9be514eedc51
```

```
[almalinux@arc.example.org ~]$ ls -lhrt /grid/cache1/
total 0
drwxr-xr-x. 3 root root 16 May 24 15:24 data
drwxr-xr-x. 2 root root 6 May 24 16:04 joblinks
```

```
[almalinux@arc.example.org ~]$ ls -lhrt /grid/cache1/joblinks/
total 0
drwx-----. 2 arcuser arcuser 24 May 24 16:12 9be514eedc51
```

```
[almalinux@arc.example.org ~]$ ls -lhrt /grid/control/jobs/
total 0
drwxr-xr-x. 4 root root 28 May 24 15:53 9be
```

A few seconds after the job is submitted with arcsub - we can check Slurm to see if the job is running:

```
[almalinux@arc.example.org ~]$ squeue
      JOBID PARTITION   NAME     USER ST       TIME  NODES NODELIST(REASON)
       8 arc-tutor hello_te  arcuser  R      0:02      1 c1
```

Indeed, running as with slurmid 8 in the arc-tutorial partition and as arcuser as expected, and on the compute node c1.

Logging into the compute node c1 we can check the scratch folder for this job:

```
[almalinux@c1 ~]$ sudo ls -lhrt /scratch/9be514eedc51/
total 12K
-rw-----. 1 arcuser arcuser 33 May 24 15:53 localfile
-rwx-----. 1 arcuser arcuser 197 May 24 15:53 runhello_inputfile.sh
lrwxrwxrwx. 1 arcuser arcuser 45 May 24 15:53 remotefile -> /grid/cache1/joblinks/
↪9be514eedc51/remotefile
lrwxrwxrwx. 1 arcuser arcuser 35 May 24 15:53 std.out -> /grid/session1/9be514eedc51/
↪std.out
lrwxrwxrwx. 1 arcuser arcuser 35 May 24 15:53 std.err -> /grid/session1/9be514eedc51/
```

(continues on next page)

(continued from previous page)

```
↪std.err
-rw-----. 1 arcuser arcuser 89 May 24 15:53 output.txt
```

- On the compute node, the job is running in the /scratch directory, and files are owned by the arcuser as configured
- We see that the file that ARC staged from a remote source is saved in the ARC cachedir and is linked in /scratch
- We see that the jobs output/error files std.out and std.err still live in the ARC-CE sessiondir as expected, ARC does not move these files
- We see that the outputfile produced by the job executable: output.txt is created in the scratch directory
- Once the job is finished all files will be moved back to the ARC-CE sessiondirectory.

Navigation

- Go back to the [Overview](#)
- Go to the next step - configure remote data delivery service: [Next](#)

Tutorial Part 3 - Production ready setup with remote ARC datadelivery service

For this part of the tutorial you need at least one remote datadelivery server set up. It makes sense that the session directory and the cache directory is hosted on this server and shared to the ARC-CE and the compute nodes.

Install ARC Datadelivery service

Step 1. Setup nightly repo

On the remote data delivery server, install the nordugrid nightly repo following the same procedure as for the ARC-CE, namely [this one](#).

Step 2. Install ARC packages and CA's

```
[root ~]# dnf install -y nordugrid-arc-datadelivery-service
```

Create the log-directory

```
mkdir /var/log/arc
```

Step 3. Configure ARC Datadelivery service

On the remote data delivery server you must create the following minimal arc.conf:

```
[common]
hostname = THIS_IP

[datadelivery-service]
allowed_ip = IP_OF_ARC_CE
transfer_dir = /grid/cache1
transfer_dir = /grid/session1
```

(continues on next page)

(continued from previous page)

```
secure=no
```

In our case we have made it simple, we do not require a host certificate (`secure=no`) and we have not set any FQDN either for this remote datadelivery server, so we use the ip-adress as hostname. The `allowed_ip` is the ip-adress of the ARC-CE(s) that we will allow requests from. Make sure to add the actual IP's for your servers.

Step 4. Install CA's

To trust the host certificate of the ARC-CE install CA's.

Step 4a. Install igtf bundle

```
[root ~]# arcctl deploy igtf-ca classic --installrepo igtf
```

Step 4b. Install test-CA certificates from ARC-CE

Note

This step is only relevant for the tutorial if you are using the test-host certificates on the ARC-CE. Otherwise the igtf bundle should be enough.

Since we are still running with the test host certificate issued by the test-CA on the ARC-CE we must also transfer these from the ARC-CE to the datadelivery server.

On the ARC-CE

On the **A-REX host** print out the Test-CA certificate:

```
[user@arc-ce] arcctl test-ca info -o ca-cert
-----BEGIN CERTIFICATE-----
MIIFyTCCA7GgAwIBAgIUeLkSbksS9r3raPvkT2rR0ep06X8wDQYJKoZIhvcNAQEM
BQAwDETMBEgCGmSJomT8ixkARkWA29yZzEZMBcGCmSJomT8ixkARkWCW5vcmR1
<output omitted>
TJ9f0I8ktHACLvLvJE9SIDWs2zPo8o4cmvLBAtxe+jaijn22THtpLLUSXt1ozexS
ZHGFtsUBuIoNzXoRXxJwkGBA1ZpLBbOpjyp6PzNcTPYFG51+EHTUMPkbfyQ5
-----END CERTIFICATE-----
```

Copy this output to your clipboard

On the ARC datadelivery server

On the remote ARC datadelivery server paste in the output from the step below as shown in the example here:

```
[root@data-delivery ~]# arcctl deploy ca-cert
[2024-10-25 21:40:50,328] [ARCCTL.ThirdParty.Deploy] [INFO] [726706] [Reading CA_
↳Certificate PEM data from stdin]
-----BEGIN CERTIFICATE-----
MIIIFyTCCA7GgAwIBAgIUeLkSbksS9r3raPvkT2rR0ep06X8wDQYJKoZIhvcNAQEM
BQAwDETMBEGCgmSJomT8ixkARkWA29yZzEZMBcGCgmSJomT8ixkARkWCW5vcnR1
<output omitted>
TJ9f0I8ktHAClVlvJE9SIDWs2zPo8o4cmvLBAtxe+jaijn22THtpLLUSXt1ozexS
ZHGFtsUBuIoNzXoRXxJwkGBA1ZpLBbOpjyp6PzNcTPYFG51+EHTUMPkbfyQ5
-----END CERTIFICATE-----
[2024-10-25 21:40:54,173] [ARCCTL.ThirdParty.Deploy] [INFO] [726706] [CA Certificate_
↳for /DC=org/DC=nordugrid/DC=ARC/O=TestCA/CN=ARC TestCA fdb0a5e3 is deployed_
↳successfully to /etc/grid-security/certificates/ARCTestCAfdb0a5e3.pem]
```

This will create all necessary files in your `x509_cert_dir` and allow your remote client to trust the ARC-CE.

Step 5. Start the remote Datadelivery service

```
[root ~]# arcctl service start -a
```

Step 6. Configure remote datadelivery service On the ARC-CE

For the ARC-CE we must add the following option to the `[arex/data-staging]` block - using the FQDN is you have that, or the ip-adress if not. Since we have set `secure=no` we use `http` otherwise `https` is used.

```
deliveryservice=http://<ip-remote-datadelivery-server>:<configured-port>/
↳datadeliveryservice
```

For instance:

```
deliveryservice=http://158.37.63.219:443/datadeliveryservice
```

Step 7. Restart the remote ARC-CE

```
[root ~]# arcctl restart service -a
```

That is it. Now you have a remote data delivery server set up, and the ARC-CE only takes care of the high level logic of the transfers, and not the actual transfers themselves.

For more details on the A-REX data transfer framework, see [here](#), and for the remote data deliver service, see [here](#).

- Go back to the [Overview](#)

Setting up ARCHERY for small research community

This tutorial demonstrates the registry setup for a small community.

The following scenario defines the use-case: a research community is using a set of ARC-enabled resources (HPC/grid/cloud) to run computations. There is no central job-server, but researchers want to submit jobs to all available ARC CEs, without the need to choose and specify a particular CE for each job.

With ARCHERY we can achieve this without introducing an extra service, relying solely on DNS.

Please refer to the [ARCHERY documentation](#) for more details beyond the scope of this tutorial.

Prerequisites

- Basic knowledge of DNS records
- Have a set of *ARC-CE's already configured*

Example infrastructure

Our example infrastructure topology is flat and consist of 3 ARC CEs:

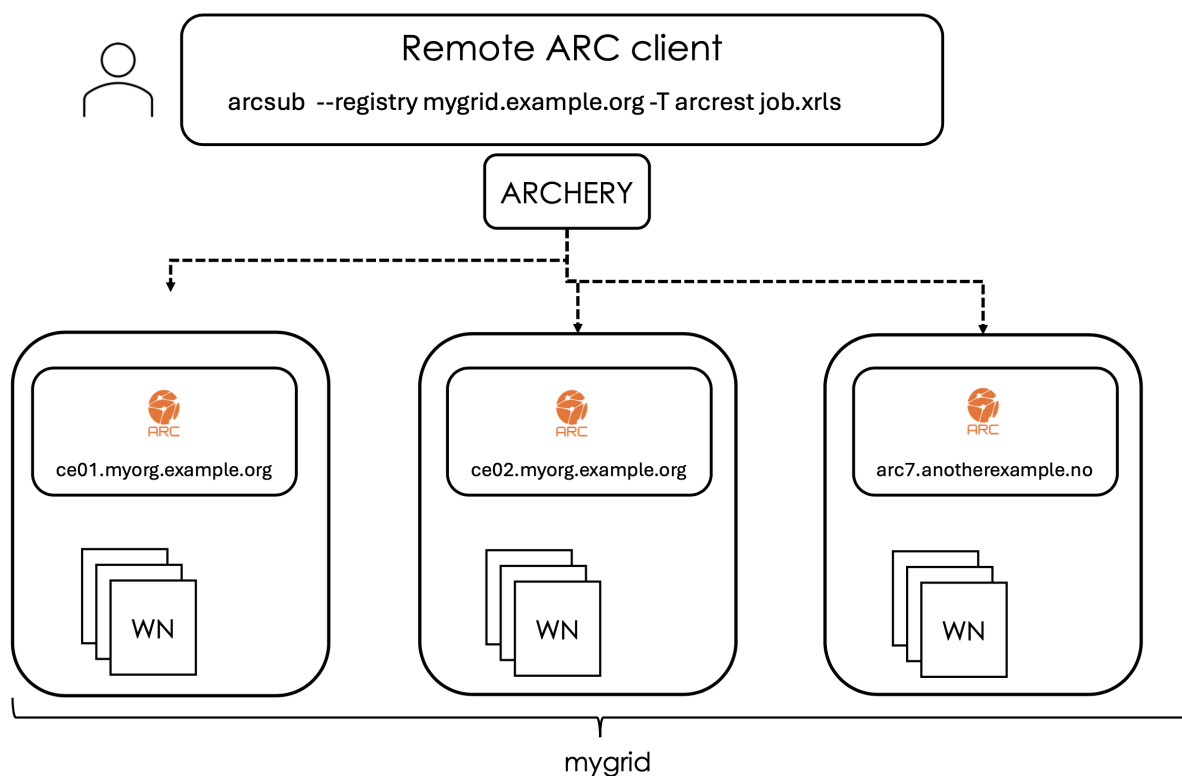


Fig. 6.9: Small research community e-Infrastructure topology - a set of ARC-CE's forming a community grid "myorg".

A researcher submitting a job from the remote client machine will use ARCHERY as a unified endpoint to submit jobs to any of the ARC-CE's in the grid (3 in our example), without the need of knowing what resources are available and choosing a particular resource.

Define ARCHERY in DNS

This tutorial uses `myorg.example.org` as the DNS domain name of the registry.

Step 1 - Collect ARC CE endpoints information

We will assume ARC 7, and are therefore only interested in the `arcrest` endpoints. This reduces the required endpoints to just 1 per ARC CE and makes the task of defining ARCHERY records simpler.

You might know the endpoint for ARC CE configuration already, but otherwise you can use `arcinfo` to get it:

```
[user@client ~]# arcinfo ce01.myorg.example.org
Computing service: ce01.myorg.example.org (production)
  Information endpoint: https://ce01.myorg.example.org:443/arex (org.ogf.glue.emies.
↪resourceinfo)
  Information endpoint: https://ce01.myorg.example.org:443/arex (org.nordugrid.
↪arcrest)
  Submission endpoint: https://ce01.myorg.example.org:443/arex (status: ok,↪
↪interface: org.ogf.glue.emies.activitycreation)
  Submission endpoint: https://ce01.myorg.example.org:443/arex (status: ok,↪
↪interface: org.nordugrid.arcrest)
```

We are interested in the `org.nordugrid.arcrest` endpoint. So our endpoint list for the example topology is:

- `https://ce01.myorg.example.org:443/arex`
- `https://ce02.myorg.example.org:443/arex`
- `https://arc7.anotherexample.no:443/arex`

Step 2 - Construct DNS records

According to the ARCHERY data model, this topology can be described with:

1. An Entry point Group Object (representing our set of CE's). DNS name should be `_archery.myorg.example.org`. It will contain pointers to Service Objects for every ARC CE.
2. Service Objects for each ARC CE. The DNS name is arbitrary and for simplicity we just name them `ce01.myorg.example.org`, `ce02.myorg.example.org` and `ce03.myorg.example.org`.

DNS records for the example topology are therefore as follows:

Name	Record-type	Record
↪-----		
↪_archery.myorg.example.org.	TXT	"u=dns://ce01.myorg.example.org.↪
↪t=archery.service"		
↪_archery.myorg.example.org.	TXT	"u=dns://ce02.myorg.example.org.↪
↪t=archery.service"		
↪_archery.myorg.example.org.	TXT	"u=dns://ce03.myorg.example.org.↪
↪t=archery.service"		
ce01.myorg.example.org.	TXT	"o=service t=org.nordugrid.arex id=ce01.↪
↪myorg.example.org"		
ce01.myorg.example.org.	TXT	"u=https://ce01.myorg.example.org:443/↪
↪arex t=org.nordugrid.arcrest"		
ce02.myorg.example.org.	TXT	"o=service t=org.nordugrid.arex↪
↪id=ce02.myorg.example.org"		

(continues on next page)

(continued from previous page)

```
ce02.myorg.example.org.      TXT      "u=https://ce02.myorg.example.org:443/
↪arex t=org.nordugrid.arcrest"

ce03.myorg.example.org.      TXT      "o=service t=org.nordugrid.arex
↪id=arc7.anotherexample.no"
ce03.myorg.example.org.      TXT      "u=https://arc7.anotherexample.no:443/
↪arex t=org.nordugrid.arcrest"
```

Step 3 - Add records to DNS

In this example we will add DNS records through an openstack DNS API. Depending on your DNS management interface available, you need to adjust accordingly.

To add DNS records for the myorg.example.org entry point _archery group:

```
[user ~]$ openstack recordset create --type TXT myorg.example.org. _archery.myorg.
↪example.org. \
--record "\"u=dns://ce01.myorg.example.org. t=archery.service\"\"" \
--record "\"u=dns://ce02.myorg.example.org. t=archery.service\"\"" \
--record "\"u=dns://ce03.myorg.example.org. t=archery.service\"\""
```

Similarly add DNS records for each of the three ARC CE services, defining the REST endpoints:

```
[user ~]$ openstack recordset create --type TXT myorg.example.org. ce01.myorg.
↪example.org. \
--record "\"o=service t=org.nordugrid.arex id=ce01.myorg.example.org\"\"" \
--record "\"u=https://ce01.myorg.example.org:443/arex t=org.nordugrid.arcrest\"\""

[user ~]$ openstack recordset create --type TXT myorg.example.org. ce02.myorg.
↪example.org. \
--record "\"o=service t=org.nordugrid.arex id=ce02.myorg.example.org\"\"" \
--record "\"u=https://ce02.myorg.example.org:443/arex t=org.nordugrid.arcrest\"\""

[user ~]$ openstack recordset create --type TXT myorg.example.org. ce03.myorg.
↪example.org. \
--record "\"o=service t=org.nordugrid.arex id=arc7.anotherexample.no\"\"" \
--record "\"u=https://arc7.anotherexample.no:443/arex t=org.nordugrid.arcrest\"\""
```

That is all that is needed to define small community registry to use distributed ARC CE resources.

Test submission using ARCHERY endpoint

Now you can use the --registry option in arctsub to submit your jobs to any of the ARC-CE from your small community grid:

```
[user ~]$ arctsub --registry myorg.example.org myjob.xmls
Job submitted with jobid: https://ce02.myorg.example.org:443/arex/rest/1.0/jobs/
↪361a038bdc51

[user ~]$ arctsub --registry myorg.example.org myjob.xmls
Job submitted with jobid: https://ce01.myorg.example.org:443/arex/rest/1.0/jobs/
↪36171e49dc51

[user ~]$ arctsub --registry myorg.example.org myjob.xmls
```

(continues on next page)

(continued from previous page)

```
Job submitted with jobid: https://arc7.anotherexample.no:443/arex/rest/1.0/jobs/
↪361db27023c6
```

(Optional alternative) Use archery-manage instead of Step 1 and 2

The *archery-manage* tool aims to automate DNS management of ARCHERY records on scale. It supports dynamic DNS updates and addresses more complex cases beyond the small community infrastructure with just a couple of ARC CEs.

In this example we are using *archery-manage* only as a help-tool to discover service endpoints and output the necessary DNS records that we in turn will feed into our DNS service. It is up to you to decide whether the manual or *archery-manage* way is a better fit for you.

The tool is distributed along other ARC packages, so you just need to install it with (for example on RHEL):

```
[root ~]# dnf install -y nordugrid-arc-archery-manage
```

First you need to define the list of ARC CEs you want to include, one CE per line. In our example it is:

```
[user ~]# cat celist.txt
ce01.myorg.example.org
ce02.myorg.example.org
arc7.anotherexample.no
```

Then you simply need to run *archery-manage* in records output mode.

Note

Currently *archery-manage* uses the LDAP service on the ARC-CE for endpoints discovery, so it has to be enabled to use this method. In the future, *archery-manage* will use the REST interface as well without the need for LDAP.

Exact command for our example use-case is:

```
[root ~]# archery-manage -s arcce-list:celist.txt -o zonefile -f type:org.nordugrid.
↪arcrest --domain myorg.example.org
2b40751af1.myorg.example.org.    TXT "o=service t=org.nordugrid.arex id=ce02.myorg.
↪example.org"
2b40751af1.myorg.example.org.    TXT "u=https://ce02:443/arex t=org.nordugrid.arcrest
↪"
2e670cff3b.myorg.example.org.    TXT "o=service t=org.nordugrid.arex id=arc7.
↪anotherexample.no"
2e670cff3b.myorg.example.org.    TXT "u=https://arc7.anotherexample.no:443/arex
↪t=org.nordugrid.arcrest"
53f2071f0a.myorg.example.org.    TXT "o=service t=org.nordugrid.arex id=ce01.myorg.
↪example.org"
53f2071f0a.myorg.example.org.    TXT "u=https://ce01:443/arex t=org.nordugrid.arcrest
↪"
_archery.myorg.example.org.      TXT "u=dns://2b40751af1.myorg.example.org.
↪t=archery.service"
_archery.myorg.example.org.      TXT "u=dns://2e670cff3b.myorg.example.org.
↪t=archery.service"
_archery.myorg.example.org.      TXT "u=dns://53f2071f0a.myorg.example.org.
↪t=archery.service"
```

Notice we use the option `-f` in the example to filter for only REST endpoints. You can add other endpoints as well if you like.

This gives us the constructed DNS records we need to add to DNS service. And we can now either add these manually as in the Step 3 above, or user `archery-manage` with the `--ddns-update` option to *update the DNS*, if you have access to do so.

Navigation

- Go back to the [Overview](#)

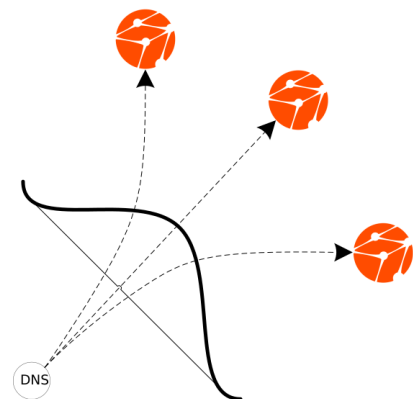
All commands in an bash file

All commands that are executed in this tutorial are aggregated in the two scripts belows. You can download these on the ARC-CE and execute them as root.

Warning: this will only work for the servers prepared in the context of an actual hands-on workshop, as configuration values for other sites will be different.

- Script with all commands to set up the zero-conf ARC-CE (slurm).
- Script with commands to set up production ARC-CE (slurm)

6.3 ARCHERY



The **ARC Hierarchical Endpoints Registry (ARCHERY)** is a novel DNS-based registry for e-Science infrastructures to store service endpoints and execution environment definitions.

ARCHERY inherits all the benefits of the DNS infrastructure since deployment, management and operation of the system is fully relying on the DNS protocol.

ARCHERY implements a minimalistic *data model* representing services and their endpoints within e-Infrastructures and embeds the service endpoint information directly into the DNS database.

All the services available within an e-Infrastructure are indexed by means of an ARCHERY instance.

Nordugrid Collaboration, the organization behind the ARC software, operates a top-level ARCHERY hosted within `nordugrid.org` DNS zone.

Similar hierarchical service catalogues can be deployed by any organization or project under a dedicated DNS zone. Instructions for howto create such registry is given [here](#).

Note

Sites are NOT required to run any reristration service on an ARC CE.

Read [this document](#) to understand what needs to be done so that an ARC CE would appear in an ARCHERY registry.

Previously a different technology was used as a service registry that was based on LDAP (see details [here](#)).

Following documents cover the generic aspects of the ARCHERY deployment and operation, including the very specific case of Nordugrid infrastructure:

6.3.1 ARCHERY deployment for Nordugrid

The Nordugrid topology for ARCHERY

The root element of the hierarchical service registry tree is the `nordugrid.org` DNS zone hosted by Niels Bohr Institutet.

The root zone holds references to:

- *per-country registries* in the dedicated sub-zones under `<country>.archery.nordugrid.org` domains or
- *projects-based registries* under project-owned specific DNS zones (e.g. `archery.myproject.org`)

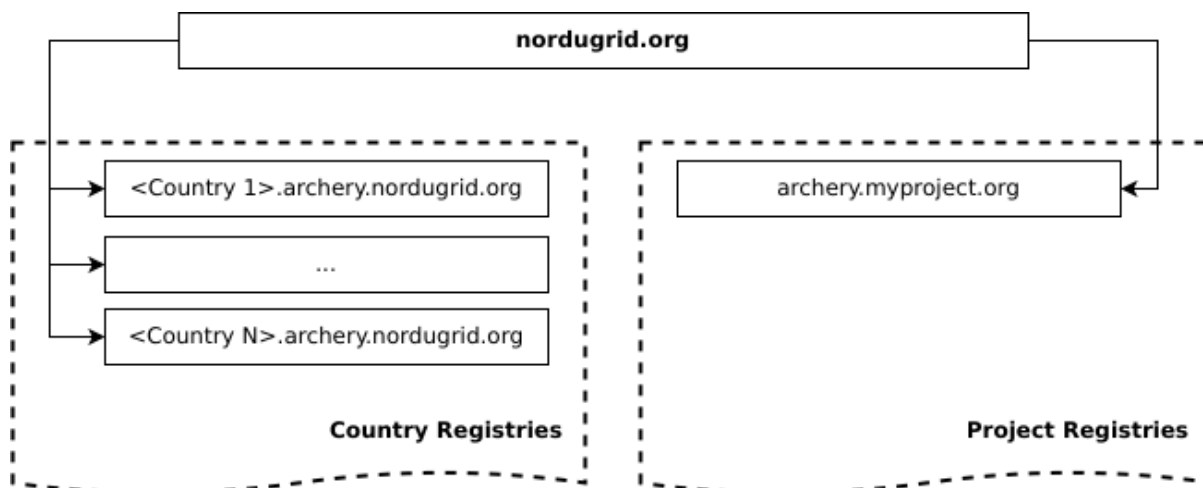


Fig. 6.10: The Nordugrid ARCHERY topology model

Per-Country ARCHERY DNS zones of the Nordugrid

Services indexed by Nordugrid are naturally belonging to a country. Therefore the deployed Nordugrid ARCHERY topology contains a set of contry-level DNS zones.

These per-country ARCHERY zones are maintained in a collaborative manner by different organizations responsible for specific country services.

A country-level DNS subzone was created for every country that previously run an ARC CE (services were taken from the monitor and GOCDB)

Note

Zones whose maintenance is *not delegated yet* are served by Niels Bohr Institutet along with the root `nordugrid.org` zone.

The summary table shows the country-level ARCHERY entry points together with the administrating contact details.

Country	Main entry point	Additional points	entry	Administration contact
Denmark	dk.archery.nordugrid.org			Petter Urkedal, urkedal@nbi.dk
Estonia	ee.archery.nordugrid.org			Zone is not delegated yet
Finland	fi.archery.nordugrid.org			Zone is not delegated yet
Hungary	hu.archery.nordugrid.org			Zone is not delegated yet
Portugal and Spain	pt.archery.nordugrid.org			Zone is not delegated yet
Lithuania	lt.archery.nordugrid.org			Zone is not delegated yet
Norway	no.archery.nordugrid.org			Zone is not delegated yet
Slovakia	sk.archery.nordugrid.org			Zone is not delegated yet
Slovenia	si.archery.nordugrid.org			Zone is not delegated yet
Sweden	se.archery.nordugrid.org			Zone is not delegated yet
Switzerland	ch.archery.nordugrid.org			Zone is not delegated yet
Ukraine	ua.archery.nordugrid.org	grid.org.ua		Andrii Salnikov, archery@grid.org.ua
UK	uk.archery.nordugrid.org			Zone is not delegated yet

Contact us if you want to take over the administration of a contry sub-zone.

Project-based ARCHERY Registries of NorduGrid

Other way to organize services is to group them by the host project. ARCHERY allows services to be affiliated with multiple projects.

It is assumed that projects are administrating their own list of services and that information is stored under a dedicated ARCHERY project-owned DNS zone.

The summary table collects ARCHERY entry points and administration contact information for projects indexed by the NorduGrid.

Registry	Entry point	Administration contact
NeIC Services for ATLAS	atlas.archery.ndgf.org	Maiken Pedersen, maikenp@uio.no

ARCHERY DNS zone administration tasks

As an ARCHERY DNS zone administrator you are expected to do following administration tasks:

- one-time initial *DNS zone setup*,
- defining and keeping up to date ARCHERY instance *services topology* in the configuration file,
- *operate regular runs* of *archery-manage* to keep DNS zone info up to date (e.g. cron-job).

Registering site in the ARCHERY

In ARCHERY sites are indexed top-to-bottom and the list of ARC CEs a maintained on the country-level and/or project-level.

As a site administrator you should technically do nothing for site registration.

Warning

In contrast to previous indexing system, ARCHERY does not requires bottop-to-top registration.

Static list of CEs is managed by ARCHERY administrator. In case your CE is not yet in the registry¹ contact your country registry manager to find the details about the established registration procedures.

Contact information for known ARCHERY managers can be found below: they will provide details for the organizational procedures (if any) upon request.

You can also [contact us](#) in case of any questions.

6.3.2 Initial setup of ARCHERY instance

This document describes how to setup an ARCHERY instance in the DNS zone, including both DNS aspects and initial service endpoint information embedding.

The described procedures apply to either *NorduGrid deployment* case or any other ARCHERY instance setup flavors, e.g. country-level, project/experiment, virtual organization, etc.

Note

Administrative access to the DNS server is required during the ARCHERY initial setup to create DNS zone and configure remote access to this zone.

Further *operating* of ARCHERY instance DOES NOT require administrative access to the DNS.

Choose DNS zone name

ARCHERY registry instance is accessible by DNS name, that is used as an entry point (e.g. to *submit jobs* to the infrastructure).

This guide uses `example.org` as an organization-owned domain name that will contain a zone for ARCHERY setup.

Despite it is possible to put records directly into the `example.org` zone, for manageability and security¹ reasons it is strongly advised to have dedicated DNS sub-zone configured for ARCHERY records.

In the *ARCHERY deployment for NorduGrid* per-country sub-zone names are *pre-defined*.

¹ You can check [NorduGrid Monitor](#) that takes ARC CE list from ARCHERY

¹ The zone content will be updated dynamically. Setting up a dedicated sub-zone is the easiest way to isolate ARCHERY data and completely eliminate the risk of changing other records that are not related to ARCHERY.

Any other setup can use arbitrary sub-zone name. This guide uses `index.example.org` DNS zone for ARCHERY setup.

Generate transaction signature key to manage ARCHERY

ARCHERY implies the usage of Dynamic DNS (DDNS) updates to manage the data inside DNS zone.

The `archery-manage` tool, that is part of NorduGrid ARC middleware, will do the DDNS updates for you as described below.

DDNS eliminates typos, allows to maintain up-to-date information and keeps it simple operate ARCHERY without administrative access to the DNS server itself.

Use the following command to generate the key:

```
[user ~]$ dnssec-keygen -a HMAC-MD5 -b 256 -n USER archery
```

From the generated files you need a *secret* part to be included in both BIND and `archery-manage` configuration.

Define key in BIND

Note

Configuration examples in this guide are provided for BIND, however you can use any name server implementation configured in a similar way.

Create the `/etc/named/archery.key` and put the generated secret key inside:

```
key archery_key {
    algorithm hmac-md5;
    secret "S0Me+SecRet+keYgener@tedwithdnssec==";
};
```

Include the key definition into `/etc/named.conf` using the following config line:

```
include "/etc/named/archery.key";
```

Keyfile for archery-manage

Create a file `archery-manage.key` and put the generated key in the following format:

```
archery_key:S0Me+SecRet+keYgener@tedwithdnssec==
```

Configure DNS zone for ARCHERY

It is generally required by the worldwide DNS infrastructure that at least one slave DNS server should be configured for every DNS zone for reliability reasons.

In this guide the following addressed will be used:

- Primary (master) DNS: `ns1.example.org` (`192.0.2.100`)
- Secondary (slave) DNS: `ns2.example.org` (`192.0.2.200`)

Define zone in BIND

Add zone definition to master DNS `/etc/named.conf`:

```
zone "index.example.org." IN {
    type master;
    file "master/index.example.org.db";
    notify yes;
    also-notify {
        # slave DNS IP address
        192.0.2.200;
    };
    allow-transfer {
        # slave DNS IP address
        192.0.2.200;
    };
    allow-update {
        key archery_key;
    };
};
```

Please observe the `allow-update` directive that authorize DDNS update requests signed by `archery-manage` key.

The secondary DNS should be configured without any special options:

```
zone "index.example.org." IN {
    type slave;
    file "slave/index.example.org.db";
    masters {
        192.0.2.100;
    };
    allow-transfer {
        192.0.2.100;
    };
};
```

Create zonefile with a basic zone info

Basic zonefile requires only SOA record. It will be filled with data by `archery-manage` later.

You can use following zonefile template (timers are subject to arrange depending on the planned update frequency):

```
$ORIGIN example.org.
$TTL 3600
index IN SOA ns1.example.org. hostmaster.example.org. (
    2018082401 ; serial
    1200      ; refresh (20 minutes)
    180       ; retry (3 minutes)
    604800    ; expire (1 week)
    60        ; minimum (1 minute)
)
NS ns1.example.org.
NS ns2.example.org.
```

Define records in parent zone

Note

If you setup a country-level index for the Nordugrid infrastructure such records are defined in the parent `nordugrid.org` zone.

Please provide your DNS setup information to us instead of following this section.

Define NS records² to refer to defined subzone:

```
$ORIGIN example.org.  
# dedicated ARCHERY zone  
index    NS ns1.example.org.  
index    NS ns2.example.org.
```

To create an ARCHERY entry point in the parent zone you can:

- define CNAME record to use `example.org` as an entry point:

```
$ORIGIN example.org.  
# ARCHERY entry point  
_archery CNAME _archery.index
```

- OR define TXT resource record with ARCHERY data pointing to group:

```
$ORIGIN example.org.  
# ARCHERY entry point  
_archery TXT "u=index.example.org t=archery.group"
```

The same technique can be used to define any other DNS aliases for an entry point (even in a complete different domain).

CNAME is recommended if you referring only one ARCHERY group.

Populate ARCHERY DNS zone with initial data

The *archery-manage* is a dedicated tool to manage information in the ARCHERY DNS zone. It is available as `nordugrid-arc-archery-manage` package in the Nordugrid repositories and EPEL.

The tool uses the configuration file in the plain text of JSON format that define services topology. Configuration file syntax is very simple and described in details in the *operations guide*.

For initial data provisioning you should *run archery-manage once* supplying config, DNS zone and transaction signature key path *generated in the previous steps*.

After zone is populated with data *Operating ARCHERY instance* comes down to keeping it up to date, running periodic updates.

² If you plan to use a different out-of-scope domain names in NS don't forget to add glue A records.

6.3.3 Operating ARCHERY instance

JSON topology configuration file for ARCHERY

JSON configuration file for archery-manage is a most flexible and customizable topology source for provisioning ARCHERY.

Following sections explains JSON syntax with the example use-cases.

Simple ARC services group

The simplest JSON config that describes single group of ARC services (matching the `arcce-list` plain-text source) can be written as follows. The `arc-services` keyword will trigger automatic endpoints discovery from ARC information system.

```
{
  "arc-services": [
    "arc1.example.org",
    "arc2.example.org"
  ]
}
```

Defining nested groups and arbitrary services

The following JSON configures the 2 subgroups - ARC and Storage.

The ARC subgroup contains 2 ARC CE. ARC CE endpoints will be discovered automatically by means of querying information system.

The Storage subgroup contains 2 manually defined services. All service endpoints specification is included into JSON file.

```
{
  "groups": [
    {
      "id": "ARC",
      "arc-services": [ "arc1.example.org", "arc2.example.org" ]
    },
    {
      "id": "Storage",
      "services": [
        {
          "id": "se1.example.org",
          "type": "DPM",
          "endpoints": [ { "httpg://se1.example.org:8446/srm/managerv2": "SRM" } ]
        },
        {
          "id": "se2.exmple.org",
          "type": "dCache",
          "endpoints": [
            { "gsiftp://ccsrm.ihep.ac.cn:2811": "gsiftp" },
            { "httpg://ccsrm.ihep.ac.cn:8446/srm/managerv2": "SRM" },
            { "xroot://ccsrm.ihep.ac.cn:1094": "xroot" }
          ]
        }
      ]
    }
  ]
}
```

(continues on next page)

(continued from previous page)

```
]
}
```

Integration with BDII

The archery-manage is able to fetch services and their endpoints from Site-BDII.

The `external-source` keyword in JSON configuration allows to inject discovered services to the desired branch of the ARCHERY registry tree.

```
{
  "groups": [
    {
      "id": "NGI_XX",
      "type": "ngi.type",
      "groups": [
        {
          "id": "XX-Site1",
          "type": "site.type",
          "external-source": {
            "sitebdii": "ldap://bdii.site1.example.org:2170/GLUE2DomainID=XX-Site1,
↪o=glue",
          }
        },
        {
          "id": "XX-Site2",
          "type": "site.type",
          "external-source": {
            "sitebdii": "ldap://bdii.site2.example.org:2170/GLUE2DomainID=XX-Site2,
↪o=glue",
          }
        }
      ]
    },
    {
      "id": "NGI_YY",
      "type": "ngi.type",
      "groups": [
        {
          "id": "YY-WLCG-Site",
          "type": "site.type",
          "external-source": {
            "sitebdii": "ldap://bdii.wlcgsite.example.org:2170/GLUE2DomainID=YY-WLCG-
↪Site,o=glue",
            "filters": [ "vo:atlas,cms", "portscan" ]
          }
        }
      ]
    }
  ],
}
```

Please notice that for YY-WLCG-Site site, during the information fetching from the Site-BDII, only services (and their endpoints) that match defined filters will be added to ARCHERY registry. In particular, it is services that according to published AccessPolicy allows to use endpoints by atlas or cms VOs and pass the network reachability test.

ARCHERY as a community trusted software registry

The archery-manage is able to provision community trusted software registry objects into the ARCHERY zone.

The JSON configuration should define the software and points to *RunTimeEnvironment scripts* location:

```
{
  "software": {
    "rtes_dir": "/home/community/rtesroot"
  }
}
```

Read more details in *this document*.

Referencing existing ARCHERY objects

ARCHERY allows to embedd overlapped trees into the DNS.

For example if particular service or group is already represented in DNS, it can be “linked” to another ARCHERY tree.

To specify such linking to existing objects the `external-archery-object` keyword should be defined in JSON config.

```
{
  "groups": [
    {
      "id": "Sweden",
      "type": "country",
      "arc-services": [
        "arc1.example.org",
        "arc2.example.org",
        "arc3.example.org"
      ]
    },
    {
      "external-archery-object": "dns://_archery.norway.nordugrid.org"
    },
    {
      "id": "Infrastructure Services",
      "services": [
        {
          "external-archery-object": "dns://voms.services.cern.ch"
        },
        {
          "id": "voms.ndgf.org",
          "type": "org.glite.voms",
          "endpoints": [
            { "voms://voms.ndgf.org:15015/nordugrid.org": "org.glite.voms" },
            { "https://voms.ndgf.org:8443/voms/nordugrid.org": "org.glite.voms-admin" }
          ]
        }
      ]
    }
  ]
}
```

Example JSON configuration above defines 3 subgroups in the ARCHERY:

- Sweden ARC services defined with `arc-services`
- Reference to existing country-level ARCHERY deployment under `norway.nordugrid.org` domain
- Common Infrastructure Services group that includes manual specification of `voms.ndgf.org` service endpoints and external VOMS service defined in another ARCHERY instance under `voms.services.cern.ch` domain.

Custom DNS data in ARCHERY zone controlled by `archery-manage`

It is possible to add arbitrary records into the ARCHERY DNS zone.

In particular it is useful to handle subzone delegation by declaring the NS and A records in the same configuration file.

Use `raw-dns` keyword to define array of the DNS records to be managed in the zone. Each record is represented by object that have `name`, DNS record type and `rdata` field that contains:

- string for a single record
- list of strings for set of records
- `null` if this record should be removed from the DNS

```
{
  "groups": [
    {
      "id": "si",
      "arc-service": [
        "meja.arnes.si",
        "jost.arnes.si"
      ]
    },
    {
      "external-archery-object": "dns://_archery.ua.archery.nordugrid.org"
    }
  ],
  "raw-dns": [
    { "name": "ua", "type": "NS", "rdata": [
      "ns1.ua.archery.nordugrid.org.",
      "ns2.ua.archery.nordugrid.org."
    ] },
    { "name": "ns1.ua", "type": "A", "rdata": "194.44.249.94" },
    { "name": "ns2.ua", "type": "A", "rdata": "194.44.249.10" }
  ]
}
```

Example JSON configuration above defines:

- 2 subgroups: one for Slovenia (`si`) and one delegated for Ukraine
- NS and 2 glue A records to define subzone delegation to the different server

Managing ARCHERY data in DNS

ARCHERY reuses the existing DNS infrastructure services thus eliminating the need to develop, deploy and operate new set of custom dedicated services.

To simplify the process of *rendering* ARCHERY records and injecting those into the DNS, ARC6 comes with the *archery-manage* information management tool.

The *archery-manage* tool had been designed to simplify common operations with ARCHERY, including registry initial bootstrap, data migration from the other service registries and keeping dynamic information up to date. It discover service endpoints, querying the resources in the defined topology configuration, then based on this information generate DNS records suitable for ARCHERY operations.

Relying on the dynamic DNS updates feature, the *archery-manage* had been designed to modify data in the DNS zone remotely, eliminating the need to interact with DNS zone configuration itself after the initial setup.

This approach also makes the setup and access rights delegation simple, fully separating the DNS hosting itself and ARCHERY data management machine.

Example of the DNS zone configuration can be found in *Initial setup of ARCHERY instance* guide.

The archery-manage data processing

It is important to understand the *archery-manage* data processing chain to efficiently maintain the ARCHERY instance:

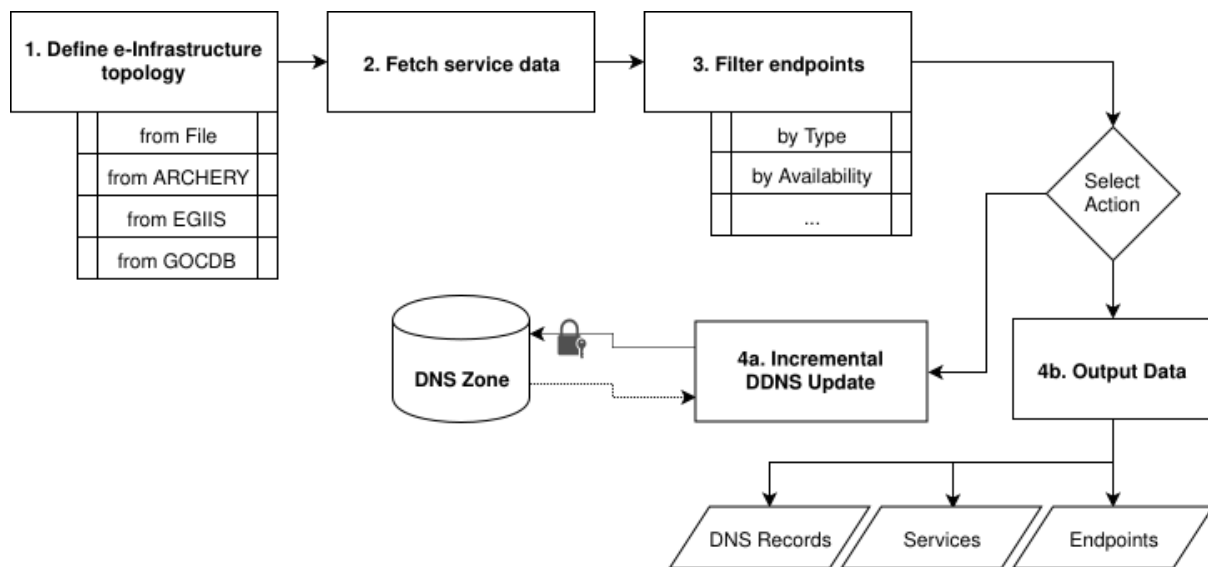


Fig. 6.11: The archery-manage data processing chain

Step 1. Define e-Infrastructure topology

Topology data defines how services are grouped within the e-Infrastructure. It comes either from a configuration file or from other databases that holds such information (including another ARCHERY instance). Interaction with already established databases (e.g. GOCDB) simplifies the integration and/or migration process.

Static list of ARC CE hostnames defined line-by-line in the plain text file is a trivial topology source that can form a group of computing elements that represent country or organization.

Step 2. Fetch service data

Topology database provides the pointers to information services that can be used to query service data. During this step the `archery-manage` tool discovers available endpoints and fetches service information.

For ARC CE hosts, the infosys LDAP GLUE2 is used to discover available endpoints with a fall-back to legacy LDAP Nordugrid schema querying.

Step 3. Filter endpoints

Set of discovered endpoints later passed to the filtering process. Based on the endpoint data (e.g. endpoint type), or additional testing (e.g. endpoint network availability check) endpoints that does not pass the filters are excluded.

Filters are extensible by design. The up to date list of currently supported filters can be obtained passing `-f help` to `archery-manage`. The overview of mostly used filters can be found *in operations guide*.

Step 4a. Incremental DDNS Update

The target automation use-case is to push the discovered data to the DNS database. This is done automatically with Dynamic DNS updates over the network. Comparing the data already available in the DNS with discovered information, `archery-manage` constructs the incremental update that only applies the difference.

Step 4b. Output data

In addition to automatic updating of the DNS database, it is also possible to automate manual operation use-case. For this the tool following the same processing chain, can be used to print out endpoint or service lists with their types or the ARCHERY DNS records that can be manually added to DNS zone configurations.

The first and main task of the administrator of ARCHERY DNS zone is to keep services topology up to date.

Services topology is a static list of e-Infrastructure services (and optionally their grouping) defined in the simple configuration file (see below).

The *archery-manage* tool than uses the topology configuration file, fetch necessary information about the defined services and update data in *ARCHERY DNS zone*.

As an ARCHERY DNS zone administrator you should establish regular updates (e.g. cron-job or similar) to keep the information up to date.

The following sections will provide details about this operations tasks.

Define infrastructure services topology

In the simplest way, when a topology is a flat list of ARC services - the text file with hostnames specified line-by-line can be used.

More advanced topologies (e.g. grouping or non-ARC services) can be described in the *JSON topology configuration file for ARCHERY*.

Simple topology: flat list of ARC CEs

CE hostnames can be defined line-by-line in a plain text file:

```
ce01.example.org
ce02.example.org
arc6.example.org
```

Alternatively you can use *JSON config* syntax as well:

```
{
  "arc-services": [
    "ce01.example.org",
    "ce02.example.org",
    "arc6.example.org"
  ]
}
```

The *archery-manage* will automatically detect ARC service endpoints using information services on ARC CE and push it to DNS zone with an incremental DDNS update.

Hierarchical topology and non-ARC services

ARCHERY is NOT restricted to the ARC services only and can hold any e-Infrastructure service endpoints information.

To define such services the JSON configuration file, as a most flexible source of topology definition for *archery-manage* should be used.

Please consult the *JSON topology configuration file for ARCHERY* document for such topologies configuration examples.

Run archery-manage to update DNS zone information

To populate DNS zone with endpoint information based on the configured services topology you should run *archery-manage*:

```
[user ~]$ archery-manage -s arcce-list:ces.list --ddns-update \
    --domain index.example.org --ddns-master-ip 192.0.2.100 \
    --ddns-tsig-keyfile archery-manage.key
```

The *-s* key that define the *type* of configuration file to use and *path* to configuration file separated by colon. The *arcce-list* type corresponds to the plain text list of ARC services. The *json* type should be used for *JSON configuration file*.

The *--ddns-tsig-keyfile* points to the location of *generated transaction signature key*.

Options *--domain* and *--ddns-master-ip* defined the dedicated ARCHERY DNS zone name and master DNS server IP respectively.

Note

Updates are performed over the network, so you can run *archery-manage* on any host. It SHOULD NOT be the DNS server itself.

The *archery-manage* writes a logs to stderr that indicate the status of each operations performed during the run, including fetching, filtering and DDNS updates. You can increase logs verbosity with *-d* option.

Consider to add filters to *archery-manage*:

Filter: Port connectivity

Check network connectivity to endpoint TCP port and filter endpoints that do not pass this test.

It is advised to have the `portscan` filter enabled during the regular operations.

```
[user ~]$ archery-manage -s arcce-list:ce.list -f portscan ...
```

Filter: Endpoint type

By default any service endpoints types that are defined or discovered will be published. If you want to filter endpoints based on type there is a dedicated filter to apply.

For example, if you want only ARC REST endpoints in the ARCHERY:

```
[user ~]$ archery-manage -s json:/etc/archery.json -f type:org.nordugrid.arcrest ...
```

Both ARC REST and LDAP GLUE2 endpoints:

```
[user ~]$ archery-manage -s json:/etc/archery.json -f type:org.nordugrid.arcrest,org.
↪nordugrid.ldapglue2 ...
```

Filter: VO

Note

VO filtering is only available when service endpoints are automatically discovered based on information services (e.g. ARC services, Site-BDII, etc)

For project-based ARCHERY deployment it is also useful to filter endpoints based on VO access policy. Only endpoints that advertise specified VO support will be added¹:

```
[user ~]$ archery-manage -s arcce-list:ce.list -f vo:exampleVO ...
```

Check the data is embedded to DNS

Once you populate ARCHERY DNS zone with data, you can try to *query the ARCHERY data from the DNS* to verify the update is working. This includes manual DNS queries or *job submission*.

Setup regular updates to ARCHERY

To keep information about endpoints up to date setup a CRON job (or Systemd Timer) to run `archery-manage` regularly.

In combination with at least *port filtering* this allows to eliminate stale endpoints and actualize information in registry.

¹ Resource information (GLUE2PolicyRule in GLUE2 and nordugrid-cluster-acl in Nordugrid LDAP) will be used as a source of supported VOs.

ARCHERY operations hints

Optimize information fetching frequency

Available services endpoints itself is not subject of rapid change. But the endpoint availability status is.

The *archery-manage* automatic endpoints fetching can be done less frequent as actual endpoint availability scan using the *JSON configuration* output:

1. Setup endpoints fetching to expanded JSON config.

Run *archery-manage* with defined topology and desired filters but without DDNS update. This job can be run less frequent, e.g. once daily.

```
[user ~]$ archery-manage -s json:/etc/archery.json -f vo:exampleVO -
↳o json > /etc/archery-expanded.json
```

2. Update information in the ARCHERY DNS zone using expanded JSON config.

For actual ARCHERY DNS zone updates use obtained expanded JSON config. This job should be run more frequently (e.g. every 5 minutes) with a *portscan filter* to eliminate stale endpoints.

```
[user ~]$ archery-manage -s json:/etc/archery-expanded.json -f
↳portscan \
    --ddns-update --ddns-tsig-keyfile archery-manage.key \
    --domain index.example.org --ddns-master-ip 192.0.2.100
```

Configure LDAP-monitor to use ARCHERY

The popular ARC CE monitoring tool - the LDAP monitor, can be used with ARCHERY out-of-the box.

It support fetching endpoints² from ARCHERY with the following configuration in *settings.inc*:

```
<?php
$archery_list = array (
    array (
        "endpoint" => "example.org",
    )
);
?>
```

6.3.4 Querying ARCHERY registry data

On the registry client level ARCHERY benefits from the distributed DNS caching on many levels, making the registry scalable.

Since DNS client is an integral part of any operating system, obtaining and processing service endpoint information from ARCHERY comes down to parsing and interpreting the data obtained from the DNS in accordance to the ARCHERY *data model rendering*.

There are several tools and options available to fetch ARCHERY data:

² Nordugrid monitor only works with *org.nordugrid.ldapng* data (and LDAP GLUE2 is experimental), so Nordugrid LDAP schema publishing is mandatory to use monitor.

The archery-manage tool as registry client

Using ARCHERY as a topology source for *archery-manage* allows to use tool output capabilities as a registry client. For example:

```
[user ~]$ archery-manage -s archery:moldyngrid.org -o arc-CEs
golowood.mao.kiev.ua
arc.imbg.org.ua
arc.univ.kiev.ua
uagrid.org.ua
grid.isma.kharkov.ua

[user ~]$ archery-manage -s archery:egi.grid.org.ua -o services
glue:sehn02.atlas.ualberta.ca/data : org.dcache.storage
glue:lcg-se1.sfu.computecanada.ca/data : org.dcache.storage
urn:ogf:ComputingService:lcg-ce1.sfu.computecanada.ca:arex : org.nordugrid.arex
urn:ogf:ComputingService:lcg-ce2.sfu.computecanada.ca:arex : org.nordugrid.arex
neutsrv1.triumf.ca_bdii-site_497425383 : bdii_site
<output omitted>

[user ~]$ archery-manage -s archery:nordugrid.org -o endpoints --json | jq .
[
{
  "t": "org.nordugrid.gridftpjob",
  "u": "gsiftp://deckard.dcs.cu.dk:2811/jobs",
  "rr_data": "u=gsiftp://deckard.dcs.cu.dk:2811/jobs t=org.nordugrid.gridftpjob"
},
{
  "t": "org.nordugrid.ldapglue2",
  "u": "ldap://deckard.dcs.cu.dk:2135/o=glue",
  "rr_data": "u=ldap://deckard.dcs.cu.dk:2135/o=glue t=org.nordugrid.ldapglue2"
},
<output omitted>
```

Any DNS client for manual data retrieval

Any DNS client or library can be used to obtain ARCHERY resource record sets. The *rendering* is designed to be simple and raw data is human readable:

```
[user ~]$ host -t TXT _archery.nordugrid.org
_archery.nordugrid.org descriptive text "u=uk.archery.nordugrid.org t=archery.group"
_archery.nordugrid.org descriptive text "u=ua.archery.nordugrid.org t=archery.group"
_archery.nordugrid.org descriptive text "u=fi.archery.nordugrid.org t=archery.group"
_archery.nordugrid.org descriptive text "u=lt.archery.nordugrid.org t=archery.group"
_archery.nordugrid.org descriptive text "u=si.archery.nordugrid.org t=archery.group"
_archery.nordugrid.org descriptive text "u=hu.archery.nordugrid.org t=archery.group"
_archery.nordugrid.org descriptive text "u=ch.archery.nordugrid.org t=archery.group"
_archery.nordugrid.org descriptive text "u=atlas.archery.ndgf.org t=archery.group"
_archery.nordugrid.org descriptive text "u=pt.archery.nordugrid.org t=archery.group"
_archery.nordugrid.org descriptive text "u=se.archery.nordugrid.org t=archery.group"
_archery.nordugrid.org descriptive text "u=no.archery.nordugrid.org t=archery.group"
_archery.nordugrid.org descriptive text "u=sk.archery.nordugrid.org t=archery.group"
_archery.nordugrid.org descriptive text "u=ee.archery.nordugrid.org t=archery.group"
_archery.nordugrid.org descriptive text "u=dk.archery.nordugrid.org t=archery.group"

[user ~]$ host -t TXT _archery.dk.archery.nordugrid.org
_archery.dk.archery.nordugrid.org descriptive text "u=dns://e06b294c0d._archery.dk.
```

(continues on next page)

(continued from previous page)

```
↪ archery.nordugrid.org. t=archery.service"

[user ~]$ host -t TXT e06b294c0d._archery.dk.archery.nordugrid.org
e06b294c0d._archery.dk.archery.nordugrid.org descriptive text "o=service t=org.
↪ nordugrid.ares id=deckard.dcsc.ku.dk"
e06b294c0d._archery.dk.archery.nordugrid.org descriptive text "u=ldap://deckard.dcsc.
↪ ku.dk:2135/o=glue t=org.nordugrid.ldapglue2"
e06b294c0d._archery.dk.archery.nordugrid.org descriptive text "u=ldap://deckard.dcsc.
↪ ku.dk:2135/Mds-Vo-Name=local,o=grid t=org.nordugrid.ldapng"
e06b294c0d._archery.dk.archery.nordugrid.org descriptive text "u=gsift://deckard.
↪ dcsc.ku.dk:2811/jobs t=org.nordugrid.gridftpjob"
```

ARC Client/SDK endpoint retrieval plugin

ARCHERY endpoint retrieval plugin for ARC client/SDK is a part of release starting from ARC 5.4.4.

It is installed within core plugins set and can be used transparently with -g option:

```
[manf@arc-client ~]$ arctest -J 2 -g moldyngrid.org -d INFO
INFO: Configuration (/etc/arc/client.conf) loaded
INFO: Configuration (/home/manf/.arc/client.conf) loaded
INFO: Using proxy file: /home/manf/.globus/x509_user_proxy
INFO: Using certificate file: /home/manf/.globus/usercert.pem
INFO: Using key file: /home/manf/.globus/userkey.pem
INFO: Using CA certificate directory: /etc/grid-security/certificates
INFO: Broker Random loaded
INFO: Found service endpoint index.moldyngrid.org (type org.nordugrid.archery)
INFO: Status for service endpoint "dns://d9b05b0008._archery.index.moldyngrid.org."
↪ is set to inactive in ARCHERY. Skipping.
INFO: Status for service endpoint "dns://ac525b5fea._archery.index.moldyngrid.org."
↪ is set to inactive in ARCHERY. Skipping.
INFO: Found service endpoint dns://7c99fb635c._archery.index.moldyngrid.org. (type
↪ archery.service)
INFO: Found service endpoint dns://616f87be8e._archery.index.moldyngrid.org. (type
↪ archery.service)
INFO: Found service endpoint dns://46e36be5ec._archery.index.moldyngrid.org. (type
↪ archery.service)
INFO: Found service endpoint ldap://grid.isma.kharkov.ua:2135/Mds-Vo-Name=local,
↪ o=grid (type org.nordugrid.ldapng)
INFO: Found service endpoint https://grid.isma.kharkov.ua:60443/ares (type org.ogf.
↪ glue.emies.activitymanagement)
<output omitted>
INFO: Found service endpoint gsift://arc6.univ.kiev.ua:2811/jobs (type org.nordugrid.
↪ gridftpjob)
INFO: Found service endpoint ldap://arc6.univ.kiev.ua:2135/o=glue (type org.nordugrid.
↪ ldapglue2)
INFO: Found service endpoint https://arc6.univ.kiev.ua:443/ares (type org.nordugrid.
↪ arcrest)
Submitting test-job 2:
&( executable = "/usr/bin/env" )( stdout = "stdout" )( stderr = "stdout" )( gmlog =
↪ "gmlog" )( jobname = "arctest2" )( clientxrs1 = "&( executable = "/usr/bin/env""
↪ )( jobname = "arctest2" )( stdout = "stdout" )( join = "yes" )( gmlog = "
↪ "gmlog"" )" )
Client version: nordugrid-arc-6.2.0
Test submitted with jobid: gsift://arc6.univ.kiev.ua:2811/jobs/
↪ RHZKDmPeedvnjwO5upha6lOqABFKDmABFKDmTeLKDmABFKDmItfmy
```

NorduGrid monitor is fetching ARCHERY data

NorduGrid infrastructure monitoring web application (aka [monitor](#)) is able to visualize realtime ARC CE information obtained using ARCHERY service endpoint discovery.

If you have own monitor setup, you can easily [configure](#) ARCHERY as an information source.

Community instance of ARCHERY capable to hold execution environment definitions (known as *RunTime Environments*). This turns ARCHERY into all-in-one community registry that cover both *resources available to community* and *execution software environments* verified and used by community.

Following document covers the ARCHERY usage as RunTime Environments registry:

6.3.5 ARCHERY as a community-defined RTEs registry

Community-defined RTEs is a new concept that aimed to automate software environment provisioning for distributed computing infrastructures.

Community-defined RTEs are prepared and digitally signed by community. All RTEs are then indexed in the registry to be discoverable. The ARC CE admin can deploy such RTEs with a single *ARC Control Tool* command after establishing a trust chain with community.

This document describes how to use ARCHERY as a Community-defined software environment registry.

1. Signing keys

The trust-chain between community and site-admin is based on the digital signatures. All Community-defined RTEs are supposed to be signed using OpenPGP standard for signatures. Technical implementation relies on the GNU Privacy Guard (GPG) software.

If you already have GPG keypair it can be used to sign community RTEs.

If not, generate a new keypair with:

```
[user ~]$ gpg --gen-key
```

2. Preparing RTEs

Community-defined RTE script is any script that follows *RunTime Environments in ARC* design for both naming scheme and the content.

In many production cases community RTEs require additional files that need to be deployed on ARC CE.

Such files can be listed within RTE script metadata (at the top of file), using the `download` keyword, for example:

```
# download: url:https://example.org/mySW.sif
↪checksum:sha256:1c0176d901301be17f86df3330f121d0ae603652eced38e8dbae413f75fad670
```

This will instruct the `arcctl` to download the `mySW.sif` singularity image from `https://example.org/`, verify the file checksum and place it into the software deployment directory for this community.

During the job execution, these downloaded files are available to payload.

Path to the directory containing these files stored in the `RUNTIME_JOB_SWDIR` variable in the job context.

Note

If RTE requires complete software package bundle it is advised to have Singularity container or at least a single tarball to be deployed on the ARC CE. Then write a simple commands for *RTE stage 1* to extract the files or define paths to the container image.

3. Signing and publishing RTEs

The *ARCHERY Manage Tool* completely automates the RTEs signing and publishing process.

Just do it

In case your RTEs are organized in the typical directory-based structure, all you need to do is to add following configuration into the *JSON topology configuration file for ARCHERY*:

```
{
  "software": {
    "rtes_dir": "/home/community/rtesroot"
  }
}
```

This configuration will instruct archery-manage to:

- export public key from GPG database (default location) to be published
- index all RTEs available in the specified directory
- extract descriptions from RTE scripts if available (to be published)
- sign all RTE scripts with GPG
- embeds public key, RTEs index and signed RTEs content into the ARCHERY zone

Customize the process

GPG location and key

In case your GPG database is not in the standard location and/or you have several keypairs in the database use the following options to point archery-manage to the right place:

```
{
  "software": {
    "gpg_home": "/srv/community-gpg",
    "gpg_keyid": "community@example.org",
    "rtes_dir": "/home/community/rtesroot"
  }
}
```

Manually define public key

In case you are not going to sign RTEs on the same machine you probably have to define the community public key manually.

One way to do it - provide URL that points to the key location:

```
{
  "software": {
    "pubkey_url": "https://example.org/gpg.public.key",
    "rtes_dir": "/home/community/rtesroot"
  }
}
```

Or you can put base64-encoded key value¹ directly into the configuration:

¹ Use `gpg --export | base64 -w0` to get value

```
{
  "software": {
    "pubkey":
    ↪ "mQENBF4y4ssBCADN9317J1KiEZSkDX2T00sCjtRLcL46XcN21lxHPFJm43ziG2GVqFHdoLQEhoN4ozvfNKxndTFkNTh6aP2Cd
    ↪ hLD5/M12Iq4ayYdBBwLJFdZDg/PSclX3kRFS4CDDvVt1XfQ1Jc/XLaV0k7nFDdbTbFEvfenMgcen/
    ↪ cxmlj1GI7IplNoiFoBjLP72NwbPW/
    ↪ WdxGkPzHO+9ypW1WSYHNTL9442zXypcOYyfRhxae5D+pSIEmptu7279Hr9u70lhYBRQW4uGI7CTiuqWxBYUKVK1s0UyjB79EJ
    ↪ VBE5jYENAA7uRe3hxeBqvlm0wmohBQ2QTDSpgwuZqp7M5kiK/
    ↪ FPf8q0QtyZGoYlCGUQc0dsBXyRjnH0sUiTKl0M9/
    ↪ PDC+u+Y2teR2XibdQ8ja8CedU09NT5Lbhx+MSv2azLSxLjB7zAVEzPCfxgcXX3G5zfNFs04+5FoRlGsoGirEN9dkZAxGa71hMh
    ↪ RBDV85aiE4RLeCZ+LQL5IND6ftezQuw3GfRzy8w5EycGcpfWoAr+D5WBhibkBDQReMuLLAQgAr9onAVIzORKjbCapUGVCiTsNp
    ↪ ",
    "rtes_dir": "/home/community/rtesroot"
  }
}
```

Host signed RTEs instead of embedding into DNS

The RTE scripts are typically small enough to be emedded directly into the ARCHERY DNS zone. But if you want to publish it on the web instead, just use `signed_rtes_url` option.

```
{
  "software": {
    "signed_rtes_url": "https://example.org/rtes/",
    "rtes_dir": "/home/community/rtesroot"
  }
}
```

In this case ARCHERY will store only the references to signed RTE scripts in `https://example.org/rtes/` location.

Signed RTEs scripts will be saved locally in the `signed` directory² and NEEDS TO BE manually uploaded to the web-hosting to be actually accessible in that location.

Manually define RTEs

Instead of organizing RTEs in the directory-like structure, they can be explicetely described in the configuration one-by-one. This allows more flexibility, as you can point to already existing and signed RTEs published anywhere.

To manully describe community RTEs define `rtes` array in the configuration:

```
{
  "software": {
    "rtes": [
      {
        "name": "APPS/COMMUNITY/SW1-1.0.0",
        "description": "Community SW1",
        "url": "https://example.org/rtes/SW1-1.0.0.signed",
      },
      {
        "name": "APPS/COMMUNITY/SW2-1.7.0",
        "description": "Community SW2",
        "data":
        ↪ "owGbWmVwMHYf+1u+Kqw5aWMaxmPJ3EEG+ka6xnoGcRzrbnrJCSWpxclFlQkpmfZ6Xgm5/

```

(continues on next page)

² can be also redefined with additional `signed_rtes_dir` option

(continued from previous page)

```

↪ jUpnnXpSZohBspBCcX15SnliUqqCRklmkyQVUm1+el50fmKJQWpRjlVFSUlBspa+fm5+TUpmXDtSjl1+Ur19UklqsX2ykV5yZp
↪ FwA=",
    },
    {
      "name": "APPS/COMMUNITY/SW3-0.2.1",
      "path": "/home/community/rtes/APPS/COMMUNITY/SW3-0.2.1"
    }
  ]
}
}

```

Array contains objects describing RTEs. Each object defines the name, optionally description and points to the content of community-defined RTE script, using one of the following options:

- `url` - provide URL to signed RTE script
- `data` - provide base64-encoded signed RTE script content
- `path` - local path to plain (unsigned) RTE script location

To sign RTE script without `archery-manage`, you can use following command:

```
[console ~]# gpg --output rtescript.signed --sign rtescript.sh
```

It is possible to use both `rtes` array and `rtes_dir` simultaneously.

Note

If you are looking for more details behind the ARCHERY idea - read [this paper](#).

6.4 ARC Tools Reference

This section holds an online version of the command line options reference for ARC command line tools.

6.4.1 ARC Tools How-Tos

Example use-cases and highlight for ARC command line tools usage patterns.

Test JWT Howto using ARC Control Tool

The *Test JWT* can be seen as a **personal Identity Provider (IdP)** on the local client computer.

It is part of the main `arcctl` package available both for clients and server installations. Server brings it as dependency, but on the client, it should be installed separately:

```
dnf install nordugrid-arc-arcctl
```

Initialize IdP

TestJWT is supposed to be run on the clint computer by regular user:

```
[user01@arc ~]$ arcctl test-jwt init
[2024-10-25 22:54:34,846] [ARCCTL.JWTIssuer] [INFO] [684] [Showing the JWKS for JWT
↳Issuer https://arc.local/arc/testjwt/89685a6f/1002]
{
  "keys": [
    {
      "e": "AQAB",
      "kid": "testjwt",
      "kty": "RSA",
      "n": "9JutYVvLUeeTqxzyBMhgwnqexVKxTb0HNrC6iZElqeyFSEK7CBrkUhSwk8Fp-
↳0qDIbtLHnXUJNZqzkcPhKFyEVN3dbXX7_3LhhVrfM9QnE0MZS2NTF5x7VgoT-NrddZ7Tm2_
↳4yD81xIs66QmJpRvPqEYT8Ds76rq5GnQP4xUitI0zicX8a7jIQk5dWRfHx8ilHdY0FtuYDxgvwm4COUaKxkhPW5FbAtcVS-
↳cJ00B1gIPdoceRUql0EkSgIIC8d7EM8plmN7G9o0W3L9wn4_Z4rn6J-
↳75nLg5KTww37SiF7ELu76EIuHCjewFBRmva10If9xTu0bzwywb3V0f8TRYw",
      "use": "sig"
    }
  ]
}
[2024-10-25 22:54:34,846] [ARCCTL.TestJWT] [INFO] [684] [Generating deployment
↳command to be executed on ARC CE to trust the Test JWT issuer https://arc.local/arc/
↳testjwt/89685a6f/1002]
arcctl deploy jwt-issuer --deploy-conf test-jwt://H4sIAIoFHGcC/6WSyZKiQBRF/
↳4V12+AAie4coEQRlcmho8NgSCBlhsQEDf+90q6u6HVvR/
↳LexTnxIt97MCnEju9gh5k8GFTXDayYCRNhXNQTLnUq72eSe07ySiyGNb4SzIpjQeQdIWD7HDdgfjA4j2F2gZlf5CjDX8PZPzCV
↳V88lRF4SqVC+hhlH97rn8SanQaHOUVujv/Y/yUQCq8kri+NBX6ouKFMc8P/LWwGHb0/
↳fVgIBVN99MZVcfIp+Uv/
↳+q4o103pjRnNI1XDT7ZN9WC0CzbezfbRCHJStja69Z0uaVWzQV0lpISdrIhrcF8VsVWZJBYlIseVy4UF6vL7GittHN5j71dtJY
↳DMx0cBl1C7HfKrUg7NNVod92pXQyxUUNhKrk37L9btRaCCvcHXlH0QFXZR/z/kEPlq2IkqV/
↳4mTcnBZteCPpaL6lnHUBR7sDL7tT7NlGz1ttt7N+q0z83IO6VSacFBuhosxFH0gbsUhsDbyNc+4wVMckG130oyoTVj3AZ2rIr0
```

Trust JWT Issuer on ARC CE

Copy the arcctl deploy command from the output of init or run arcctl test-jwt export to show the command again.

```
[root@arc-ce ~]# arcctl deploy jwt-issuer --deploy-conf test-jwt://H4sIAIoFHGcC/
↳6WSyZKiQBRF/4V12+AAie4coEQRlcmho8NgSCBlhsQEDf+90q6u6HVvR/
↳LexTnxIt97MCnEju9gh5k8GFTXDayYCRNhXNQTLnUq72eSe07ySiyGNb4SzIpjQeQdIWD7HDdgfjA4j2F2gZlf5CjDX8PZPzCV
↳V88lRF4SqVC+hhlH97rn8SanQaHOUVujv/Y/yUQCq8kri+NBX6ouKFMc8P/LWwGHb0/
↳fVgIBVN99MZVcfIp+Uv/
↳+q4o103pjRnNI1XDT7ZN9WC0CzbezfbRCHJStja69Z0uaVWzQV0lpISdrIhrcF8VsVWZJBYlIseVy4UF6vL7GittHN5j71dtJY
↳DMx0cBl1C7HfKrUg7NNVod92pXQyxUUNhKrk37L9btRaCCvcHXlH0QFXZR/z/kEPlq2IkqV/
↳4mTcnBZteCPpaL6lnHUBR7sDL7tT7NlGz1ttt7N+q0z83IO6VSacFBuhosxFH0gbsUhsDbyNc+4wVMckG130oyoTVj3AZ2rIr0
[2024-10-25 22:58:12,626] [ARCCTL.ThirdParty.Deploy] [INFO] [266740] [ARC CE now
↳trust JWT signatures of https://arc.local/arc/testjwt/89685a6f/1002 issuer.]
[2024-10-25 22:58:12,626] [ARCCTL.JWTIssuer] [INFO] [266740] [Auth configuration for
↳JWT issuer https://arc.local/arc/testjwt/89685a6f/1002 has been written to /etc/arc.
↳conf.d/10-testjwt-8475ff6c.conf]
[2024-10-25 22:58:12,626] [ARCCTL.JWTIssuer] [WARNING] [266740] [ARC services restart
↳is needed to apply configuration changes.]

[root@arc-ce ~]# arcctl service restart -a
```

Issue JWT token

To generate token signed by TestJWT instance run (print to stdout):

```
[user01@arc ~]$ arcctl test-jwt token
eyJhbGciOiJSUzI1NiIsImtpZCI6InRlc3Rqd3QifQ.
eyJhdWQiOiJhcmMiLCJhenAiOiJhcmMiLCJleHAiOjE3Mjk5MzMyMDMsImldCI6MTcyOTg5MDAwMywiaXNzIjoiaHR0cHM6Ly9R7Q15g5jVHdyESmKzNB68-j32iB6-yxz4v00n0xA7RXURKMB-JjT6xc20Hh9McTXrenhxW3o_Z-Fo4ReEgU2VCeKCswRQavxXPejP9v0rnpTJM3oKyHsu07GnuEp8nf-7DaKgbJwK3Y1Baidd-6ygXjWCS9M0zeJLP2BgerPMPItjR8KiYw6eLDp7UAs3LYxSM1Ktajt2bvhCGn2qG_8t0zaM0rOelWbL50TafnlSw8xlkaHmUXb-cQwoL6c17zKcBYySwmX8Ed7QJLNAUUAavyqwpP16J54nKbGbERPhH__iupSRhHMqCUf_yh-y_vuQI13upZH9QaHGYsv2B60Q
```

ARC client relies on the BEARER_TOKEN variable, so the convenient way is to run:

```
[user01@arc ~]$ export BEARER_TOKEN=$( arcctl test-jwt token )
```

Fine-tune token claims

Token claims can be modified via command line arguments to `arcctl test-jwt token`.

Command help is a best friend here:

```
[user01@arc ~]$ arcctl test-jwt token -h
usage: arcctl test-jwt token [-h] [-p PROFILE] [-n USERNAME] [-v VALIDITY] [-s SCOPES] [-c CLAIMS]

optional arguments:
  -h, --help            show this help message and exit
  -p PROFILE, --profile PROFILE
                        Generate using token named profile (default is default)
  -n USERNAME, --username USERNAME
                        Use specified username instead of automatically generated
  -v VALIDITY, --validity VALIDITY
                        Validity of the token in hours (default is 12)
  -s SCOPES, --scopes SCOPES
                        Additional scopes to include into the token
  -c CLAIMS, --claims CLAIMS
                        Additional claims (JSON) to include into the token
```

For example:

```
arcctl test-jwt token -n 'Andrii' -s "storage.read storage.write" -v 72`
```

will generate a token with a subject “Andrii”, adds stuff to list of scopes and increase validity from default 12 to 72 hours.

Using token profiles

You can create a profile with customized claims for re-use in tokens generation instead of using command line arguments:

```
[user01@arc ~]$ arcctl test-jwt config-set -p wlcg username "Atlas User"
[user01@arc ~]$ arcctl test-jwt config-set -p wlcg scopes "compute.create:/"
[user01@arc ~]$ arcctl test-jwt config-set -p wlcg claims '{"wlcg.ver": "1.0", "wlcg.
→groups": ["/atlas", "/atlas/production"]}'
[user01@arc ~]$ arcctl test-jwt config-set -p wlcg validity 24

[user01@arc ~]$ arcctl test-jwt config-get -p wlcg
{
  "username": "Atlas User",
  "scopes": "compute.create:/",
  "claims": "{\"wlcg.ver\": \"1.0\", \"wlcg.groups\": [\"/atlas\", \"/atlas/
→production\"]}",
  "validity": "24"
}

[user01@arc ~]$ arcctl test-jwt token -p wlcg
```

Security Operations using ARC Control Tool

This page collects general information about security operations.

Finding jobs submitted by DN

In case it is necessary to find all the jobs belonging to a specific user, the following command may be run by the local CE admin:

```
[root ~]# arcctl job list --owner <DN>
```

Killing jobs submitted by DN

In case you need to kill all jobs owned by a specific user, identified by a certain certificate DN:

```
[root ~]# arcctl job killall --owner <DN>
```

Removing jobs submitted by DN

Warning

Do not do this before any needed forensics have been performed

In case you want to remove all files related to a jobs run by a specific user:

```
[root ~]# arcctl job cleanall --owner <DN>
```

6.4.2 ARC Control Tool

NorduGrid ARC Control Tool

```
usage: arcctl [-h] [-c CONFIG] [-d {CRITICAL,ERROR,WARNING,INFO,DEBUG}]
            COMPONENT ...
```

Named Arguments

-c, --config	config file location (default is /etc/arc.conf)
-d, --debug	Possible choices: CRITICAL, ERROR, WARNING, INFO, DEBUG verbosity level (default is “WARNING”) Default: “WARNING”

ARC Components

COMPONENT	Possible choices: deploy, test-ca, test-jwt, config, service, rte, accounting, job, cache, datastaging, cleanup DESCRIPTION
------------------	--

Sub-commands:

deploy

Third party components deployment

```
arcctl deploy [-h] ACTION ...
```

Deployment Actions

ACTION	Possible choices: igtf-ca, ca-cert, vomses, voms-lsc, jwt-issuer, iptables-config DESCRIPTION
---------------	--

Sub-commands:

igtf-ca

Deploy IGTF CA certificates

```
arcctl deploy igtf-ca [-h] [-i {igtf,egi-trustanchors}]
                    {classic,iota,mics,slcs} [{classic,iota,mics,slcs} ...]
```

Positional Arguments

bundle Possible choices: classic, iota, mics, slcs
IGTF CA bundle name

Named Arguments

-i, --installrepo Possible choices: igtf, egi-trustanchors
Add specified repository that contains IGTF CA certificates

ca-cert

Deploy trusted CA certificate

```
arcctl deploy ca-cert [-h] [-c CRT] [-f] [--x509-cert-dir X509_CERT_DIR]
                        [--url-ca {system,grid,insecure}]
```

Named Arguments

-c, --crt CA Certificate PEM file location: HTTPS URL or path to file (default is to read from stdin)

-f, --force Overwrite if already exists
Default: False

--x509-cert-dir Redefine CA Certificates dir location

--url-ca Possible choices: system, grid, insecure
PKI CA Bundle for URL locations (default is “system”)
Default: “system”

vomses

Deploy VOMS client configuration files

```
arcctl deploy vomses [-h] (-v VOMS | -e) [-u] [-c] vo
```

Positional Arguments

vo VO Name

Named Arguments

-v, --voms	VOMS-Admin URL
-e, --egi-vo	Print URL to EGI VOs database Default: False
-u, --user	Install to user's home instead of /etc Default: False
-c, --use-client-cert	Use client certificate to contact VOMS-Admin Default: False

voms-lsc

Deploy VOMS server-side list-of-certificates files

```
arcctl deploy voms-lsc [-h] (-v VOMS | -e) [--pythonssl] vo
```

Positional Arguments

vo	VO Name
-----------	---------

Named Arguments

-v, --voms	VOMS-Admin URL
-e, --egi-vo	Print URL to EGI VOs database Default: False
--pythonssl	Use Python SSL module to establish TLS connection (default is to call external OpenSSL binary) Default: False

jwt-issuer

Deploy JWT Issuer trust data

```
arcctl deploy jwt-issuer [-h] [--url-ca {system,grid,insecure}] [-i] url
```

Positional Arguments

url	Issuer metadata URL (.well-known/configuration or test-jwt://)
------------	--

Named Arguments

--url-ca	Possible choices: system, grid, insecure PKI CA Bundle (default is “system”) Default: “system”
-i, --deploy-conf	Automatically add atuh configuration snipped to arc.conf.d Default: False

iptables-config

Generate iptables config to allow ARC CE configured services

```
arcctl deploy iptables-config [-h] [--any-state] [--multiport]
```

Named Arguments

--any-state	Do not add ‘-state NEW’ to filter configuration Default: False
--multiport	Use one-line multiport filter instead of per-service entries Default: False

test-ca

ARC Test CA control

```
arcctl test-ca [-h] [--ca-id CA_ID] [--ca-dir CA_DIR] ACTION ...
```

Named Arguments

--ca-id	Define CA ID to work with (default is to use hostname-based hash)
--ca-dir	Redefine path to CA files directory

Test CA Actions

ACTION	Possible choices: init, info, cleanup, hostcert, usercert DESCRIPTION
---------------	--

Sub-commands:

init

Generate self-signed TestCA files

```
arcctl test-ca init [-h] [-d {sha224,sha256,sha384,sha512}] [-v VALIDITY] [-f]
```

Named Arguments

-d, --digest	Possible choices: sha224, sha256, sha384, sha512 Digest to use (default is “sha384”) Default: “sha384”
-v, --validity	Validity of certificate in days (default is 90) Default: 90
-f, --force	Overwrite files if exist Default: False

info

Show information about TestCA

```
arcctl test-ca info [-h] [-o {dn,files,ca-cert}]
```

Named Arguments

-o, --output	Possible choices: dn, files, ca-cert Specify what information to show (default is “dn”) Default: “dn”
---------------------	---

cleanup

Cleanup TestCA files

```
arcctl test-ca cleanup [-h]
```

hostcert

Generate and sign testing host certificate

```
arcctl test-ca hostcert [-h] [-d {sha224,sha256,sha384,sha512}] [-v VALIDITY]
                        [-n HOSTNAME] [-f] [-t]
```

Named Arguments

-d, --digest	Possible choices: sha224, sha256, sha384, sha512 Digest to use (default is “sha384”) Default: “sha384”
-v, --validity	Validity of certificate in days (default is 30) Default: 30
-n, --hostname	Generate certificate for specified hostname instead of this host
-f, --force	Overwrite files if exist Default: False
-t, --export-tar	Export tar archive to use from another host Default: False

usercert

Generate and sign testing user certificate

```
arcctl test-ca usercert [-h] [-d {sha224,sha256,sha384,sha512}] [-v VALIDITY]
                        [-n USERNAME] [-i INSTALL_USER] [-t] [-f] [--no-auth]
```

Named Arguments

-d, --digest	Possible choices: sha224, sha256, sha384, sha512 Digest to use (default is “sha384”) Default: “sha384”
-v, --validity	Validity of certificate in days (default is 30) Default: 30
-n, --username	Use specified username instead of automatically generated
-i, --install-user	Install certificates to \$HOME/.globus for specified user instead of workdir
-t, --export-tar	Export tar archive to use from another host Default: False
-f, --force	Overwrite files if exist Default: False
--no-auth	Do not add user subject to allowed list Default: False

test-jwt

ARC Test JWT control

```
arcctl test-jwt [-h] [--iss-id ISS_ID] [--jwk-dir JWK_DIR] ACTION ...
```

Named Arguments

--iss-id	Define arcctl token Issuer ID to work with (default is hostname)
--jwk-dir	Redefine path to JWK files directory (default is “/etc/grid-security/jwt”) Default: “/etc/grid-security/jwt”

Test JWT Actions

ACTION	Possible choices: init, info, export, cleanup, config-get, config-set, token
DESCRIPTION	

Sub-commands:

init

Generate RSA key-pair for JWT signing

```
arcctl test-jwt init [-h] [-f]
```

Named Arguments

-f, --force	Overwrite files if exist Default: False
--------------------	--

info

Show information about Test JWT issuer

```
arcctl test-jwt info [-h]
```

export

Export JWT issuer information to be imported to ARC CE

```
arcctl test-jwt export [-h]
```

cleanup

Cleanup TestJWT Issuer files

```
arcctl test-jwt cleanup [-h]
```

config-get

Get JWT token generation config

```
arcctl test-jwt config-get [-h] [-p PROFILE] [key]
```

Positional Arguments

key	Config key
------------	------------

Named Arguments

-p, --profile	Config named profile (default is “default” Default: “default”
----------------------	--

config-set

Set JWT token generation config

```
arcctl test-jwt config-set [-h] [-p PROFILE]
                           {username,validity,scopes,claims} value
```

Positional Arguments

key	Possible choices: username, validity, scopes, claims Config key as in token options
value	Config value

Named Arguments

-p, --profile	Config named profile (default is “default” Default: “default”
----------------------	--

token

Issue JWT token

```
arcctl test-jwt token [-h] [-p PROFILE] [-n USERNAME] [-v VALIDITY]
                      [-s SCOPES] [-c CLAIMS]
```

Named Arguments

-p, --profile	Generate using token named profile (default is “default” Default: “default”
-n, --username	Use specified username instead of automatically generated
-v, --validity	Validity of the token in hours (default is 12)
-s, --scopes	Additional scopes to include into the token
-c, --claims	Additional claims (JSON) to include into the token

config

ARC CE configuration control

```
arcctl config [-h] ACTION ...
```

Config Actions

ACTION	Possible choices: dump, get, describe, brief, verify
DESCRIPTION	

Sub-commands:

dump

Dump ARC CE running configuration

```
arcctl config dump [-h]
```

get

Print configuration option value

```
arcctl config get [-h] block option
```

Positional Arguments

block	Name of configuration block (without square brackets)
option	Configuration option name

describe

Describe configuration option

```
arcctl config describe [-h] [-r REFERENCE] block option
```

Positional Arguments

block	Name of configuration block (without square brackets)
option	Configuration option name

Named Arguments

-r, --reference	Redefine arc.conf.reference location (default is “/usr/share/doc/nordugrid-arc/arc.conf.reference”) Default: “/usr/share/doc/nordugrid-arc/arc.conf.reference”
------------------------	---

brief

Print configuration brief points

```
arcctl config brief [-h] [-t {storage,logs}]
```

Named Arguments

-t, --type	Possible choices: storage, logs Show brief only for provided options type
-------------------	--

verify

Verify ARC CE configuration syntax

```
arcctl config verify [-h] [-r REFERENCE]
```

Named Arguments

-r, --reference Redefine arc.conf.reference location (default is “/usr/share/doc/nordugrid-arc/arc.conf.reference”)
 Default: “/usr/share/doc/nordugrid-arc/arc.conf.reference”

service

ARC CE services control

```
arcctl service [-h] ACTION ...
```

Services Actions

ACTION Possible choices: enable, disable, start, restart, stop, verify, list
 DESCRIPTION

Sub-commands:

enable

Enable ARC CE services

```
arcctl service enable [-h] [--now] (-a | -s SERVICE)
```

Named Arguments

--now Start the services just after enable
 Default: False
-a, --as-configured Use information from arc.conf to get services list
 Default: False
-s, --service Service name

disable

Disable ARC CE services

```
arcctl service disable [-h] [--now] (-a | -s SERVICE)
```

Named Arguments

--now	Stop the services just after disable Default: False
-a, --as-configured	Use information from arc.conf to get services list Default: False
-s, --service	Service name

start

Start ARC CE services

```
arcctl service start [-h] (-a | -s SERVICE)
```

Named Arguments

-a, --as-configured	Use information from arc.conf to get services list Default: False
-s, --service	Service name

restart

Restart ARC CE services

```
arcctl service restart [-h] (-a | -s SERVICE)
```

Named Arguments

-a, --as-configured	Use information from arc.conf to get services list Default: False
-s, --service	Service name

stop

Stop ARC CE services

```
arcctl service stop [-h] (-a | -s SERVICE)
```

Named Arguments

- a, --as-configured** Use information from arc.conf to get services list
Default: False
- s, --service** Service name

verify

Verify ARC CE

```
arcctl service verify [-h] [-r REFERENCE]
```

Named Arguments

- r, --reference** Redefine arc.conf.reference location (default is “/usr/share/doc/nordugrid-arc/arc.conf.reference”)
Default: “/usr/share/doc/nordugrid-arc/arc.conf.reference”

list

List ARC CE services and their states

```
arcctl service list [-h] [-i | -e | -a]
```

Named Arguments

- i, --installed** Show only installed services
Default: False
- e, --enabled** Show only enabled services
Default: False
- a, --active** Show only running services
Default: False

rte

RunTime Environments

```
arcctl rte [-h] ACTION ...
```

RunTime Environments Actions

ACTION	Possible choices: enable, disable, list, default, undefault, cat, params-get, params-set, params-unset, community
	DESCRIPTION

Sub-commands:

enable

Enable RTE to be used by A-REX

```
arcctl rte enable [-h] [-f] [-d] rte [rte ...]
```

Positional Arguments

rte	RTE name
------------	----------

Named Arguments

-f, --force	Force RTE enabling Default: False
-d, --dummy	Enable dummy RTE that do nothing but published in the infosys Default: False

disable

Disable RTE to be used by A-REX

```
arcctl rte disable [-h] rte [rte ...]
```

Positional Arguments

rte	RTE name
------------	----------

list

List RunTime Environments

```
arcctl rte list [-h] [-l] [-e | -d | -a | -s | -u | -n | -c]
```

Named Arguments

-l, --long	Detailed listing of RTEs Default: False
-e, --enabled	List enabled RTEs Default: False
-d, --default	List default RTEs Default: False
-a, --available	List available RTEs Default: False
-s, --system	List available system RTEs Default: False
-u, --user	List available user-defined RTEs Default: False
-n, --dummy	List dummy enabled RTEs Default: False
-c, --community	List deployed community RTEs Default: False

default

Transparently use RTE for every A-REX job

```
arcctl rte default [-h] [-f] rte [rte ...]
```

Positional Arguments

rte	RTE name
------------	----------

Named Arguments

-f, --force	Force RTE enabling Default: False
--------------------	--------------------------------------

undefault

Remove RTE from transparent A-REX usage

```
arcctl rte undefault [-h] rte [rte ...]
```

Positional Arguments

rte RTE name

cat

Print the content of RTE file

```
arcctl rte cat [-h] rte
```

Positional Arguments

rte RTE name

params-get

List configurable RTE parameters

```
arcctl rte params-get [-h] [-l] rte
```

Positional Arguments

rte RTE name

Named Arguments

-l, --long Detailed listing of parameters
Default: False

params-set

Set configurable RTE parameter

```
arcctl rte params-set [-h] rte parameter value
```

Positional Arguments

rte RTE name
parameter RTE parameter to configure
value RTE parameter value to set

params-unset

Use default value for RTE parameter

```
arcctl rte params-unset [-h] rte parameter
```

Positional Arguments

rte	RTE name
parameter	RTE parameter to unset

community

Operating community-defined RunTimeEnvironments

```
arcctl rte community [-h] ACTION ...
```

Community RTE Actions

ACTION	Possible choices: add, remove, list, config-get, config-set, rte-list, rte-cat, rte-deploy, rte-remove
	DESCRIPTION

Sub-commands:

add

Add new trusted community to ARC CE

```
arcctl rte community add [-h] [-f FINGERPRINT]
                        [-a ARCHERY | -u URL | --pubkey PUBKEY | --keyserver_
↳KEYSERVER]
                        community
```

Positional Arguments

community	Trusted community name
------------------	------------------------

Named Arguments

-f, --fingerprint	Fingerprint of the community key
-a, --archery	Use ARCHERY domain name (this is the default with community name as a domain)
-u, --url	Use JSON URL
--pubkey	Manually defined location (URL) of the public key
--keyserver	Manually defined location of PGP keyserver

remove

Remove trusted community from ARC CE

```
arcctl rte community remove [-h] [-f] community
```

Positional Arguments

community	Trusted community name
------------------	------------------------

Named Arguments

-f, --force	Disable and undefault all community RTEs automatically Default: False
--------------------	--

list

List trusted communities

```
arcctl rte community list [-h] [-l]
```

Named Arguments

-l, --long	Print more information Default: False
-------------------	--

config-get

Get config variables for trusted community

```
arcctl rte community config-get [-h] [-l] community [option ...]
```

Positional Arguments

community	Trusted community name
option	Configuration option name

Named Arguments

-l, --long	Print more information Default: False
-------------------	--

config-set

Set config variable for trusted community

```
arcctl rte community config-set [-h] community option value
```

Positional Arguments

community	Trusted community name
option	Configuration option name
value	Configuration option value

rte-list

List RTEs provided by community

```
arcctl rte community rte-list [-h] [-l | -a | -d] community
```

Positional Arguments

community	Trusted community name
------------------	------------------------

Named Arguments

-l, --long	Print more information Default: False
-a, --available	List RTEs available in the software registry Default: False
-d, --deployed	List deployed community RTEs Default: False

rte-cat

Print the content of RTEs provided by community

```
arcctl rte community rte-cat [-h] community rtename
```

Positional Arguments

community	Trusted community name
rtename	RunTimeEnvironment name

rte-deploy

Deploy RTE provided by community

```
arcctl rte community rte-deploy [-h] [-u URL] [-f] [--insecure]
                                community rtename
```

Positional Arguments

community	Trusted community name
rtename	RunTimeEnvironment name

Named Arguments

-u, --url	Explicitly define URL to signed RTE file
-f, --force	Force RTE files redeployment if already exists Default: False
--insecure	Do not validate community signature for URL-based deployment Default: False

rte-remove

Remove deployed community RTE

```
arcctl rte community rte-remove [-h] [-f] community rtename
```

Positional Arguments

community	Trusted community name
rtename	RunTimeEnvironment name

Named Arguments

-f, --force	Disable and undefault RTE automatically Default: False
--------------------	---

accounting

A-REX Accounting records management

```
arcctl accounting [-h] [--database-file DATABASE_FILE] ACTION ...
```

Named Arguments

--database-file Set custom location of accounting database to work with. Path or file name inside controldir.

Accounting Actions

ACTION Possible choices: stats, job, database, republish
DESCRIPTION

Sub-commands:

stats

Show A-REX AAR statistics

```
arcctl accounting stats [-h] [-b END_FROM] [-e END_TILL] [-s START_FROM] [-a]
                        [--filter-vo FILTER_VO] [--filter-fqan FILTER_FQAN]
                        [--filter-user FILTER_USER]
                        [--filter-state FILTER_STATE]
                        [--filter-queue FILTER_QUEUE]
                        [--filter-endpoint FILTER_ENDPOINT]
                        [--filter-extra ATTRIBUTE VALUE]
                        [-o {brief,jobcount,walltime,cputime,data-staged-in,data-
→staged-out,wlcgvos,fqans,users,jobids,json}]
```

Named Arguments

-b, --end-from Define the job completion time range beginning (YYYY-MM-DD [HH:mm[:ss]])

-e, --end-till Define the job completion time range end (YYYY-MM-DD [HH:mm[:ss]])

-s, --start-from Define the job start time constraint (YYYY-MM-DD [HH:mm[:ss]])

-a, --show-active Include active jobs (not in terminal states) to the stats
Default: False

--filter-vo Account jobs owned by specified WLCG VO(s)

--filter-fqan Account jobs owned by users with specified FQAN(s)

--filter-user Account jobs owned by specified user(s)

--filter-state Account jobs in the defined state(s)

--filter-queue Account jobs submitted to the defined queue(s)

--filter-endpoint Account jobs submitted via defined endpoint type(s)

--filter-extra Filter extra attributes (e.g. jobname, project, vomsfqan, rte, dtrurl, etc)

-o, --output Possible choices: brief, jobcount, walltime, cputime, data-staged-in, data-staged-out, wlcgvos, fqans, users, jobids, json
Define what kind of stats you want to output (default is “brief”)
Default: “brief”

job

Show job accounting data

```
arcctl accounting job [-h] ACTION ...
```

Job Accounting Actions

ACTION	Possible choices: info, events, transfers
	DESCRIPTION

Sub-commands:

info

Show job accounting data

```
arcctl accounting job info [-h]
                             [-o {all,description,resources,rtes,authtokens,json}]
                             jobid
```

Positional Arguments

jobid	Job ID
--------------	--------

Named Arguments

-o, --output	Possible choices: all, description, resources, rtes, authtokens, json
	Define what kind of job information you want to output (default is “all”)
	Default: “all”

events

Show job event history

```
arcctl accounting job events [-h] jobid
```

Positional Arguments

jobid	Job ID
--------------	--------

transfers

Show job data transfers statistics

```
arcctl accounting job transfers [-h] jobid
```

Positional Arguments

jobid	Job ID
--------------	--------

database

Accounting database operations

```
arcctl accounting database [-h] ACTION ...
```

Accounting Database Actions

ACTION	Possible choices: backup, optimize, migrate, cleanup, rotate
DESCRIPTION	

Sub-commands:

backup

Run online backup of accounting database

```
arcctl accounting database backup [-h] --location LOCATION
```

Named Arguments

--location	Database backup file path
-------------------	---------------------------

optimize

Optimize database query performance

```
arcctl accounting database optimize [-h] [-b END_FROM] [-e END_TILL]
```

Named Arguments

-b, --end-from	Define time range start for ANALYSE queries (YYYY-MM-DD [HH:mm[:ss]]). Default is 2 weeks ago.
-e, --end-till	Optionally define to which point in time migrate records (YYYY-MM-DD [HH:mm[:ss]]).

migrate

Migrate records from legacy ARC6 database

```
arcctl accounting database migrate [-h] -f FROM_DB -b END_FROM [-e END_TILL]
```

Named Arguments

-f, --from-db	Legacy database file (default is “accounting.db”) Default: “accounting.db”
-b, --end-from	Define from which point in time migrate records (YYYY-MM-DD [HH:mm[:ss]])
-e, --end-till	Optionally define to which point in time migrate records (YYYY-MM-DD [HH:mm[:ss]]).

cleanup

Delete old records from accounting database

```
arcctl accounting database cleanup [-h] -e END_TILL [-b END_FROM] [--yes]
                                   [--vacuum]
```

Named Arguments

-e, --end-till	Cleanup records older than specified completion time (YYYY-MM-DD [HH:mm[:ss]])
-b, --end-from	Optionally define from which point in time cleanup records (YYYY-MM-DD [HH:mm[:ss]]).
--yes	Answer yes to all questions Default: False
--vacuum	Vacuum database (this is blocking action) Default: False

rotate

Rotate accounting database

```
arcctl accounting database rotate [-h] --date DATE [--rotate-to ROTATE_TO]
                                   [--yes] [--vacuum]
```

Named Arguments

--date	Rotate database at specified date (YYYY-MM-DD)
--rotate-to	Rotated database file. Path or file name inside controldir. Default is to add date suffix.
--yes	Answer yes to all questions Default: False
--vacuum	Vacuum database (this is blocking action) Default: False

republish

Republish accounting records to defined target

```
arcctl accounting republish [-h] [--dry-run] -b END_FROM -e END_TILL
                             (-t TARGET_NAME | -a APEL_URL | -s SGAS_URL)
                             [--apel-topic {gLite-APEL,/queue/global.accounting.test.
↪cpu.central}]
                             [--apel-messages {urs,summaries,summaries-v04}]
                             [--apel-project APEL_PROJECT]
                             [--gocdb-name GOCDB_NAME]
                             [--localid-prefix LOCALID_PREFIX]
                             [--vofilter VOFILTER] [--urbatchsize URBATCHSIZE]
```

Named Arguments

--dry-run	Print out messages instead of sending them to target Default: False
-b, --end-from	Define republishing timeframe start (YYYY-MM-DD [HH:mm[:ss]])
-e, --end-till	Define republishing timeframe end (YYYY-MM-DD [HH:mm[:ss]])
-t, --target-name	Specify configured accounting target name from arc.conf (e.g. neic_sgas).
-a, --apel-url	Specify APEL server URL (e.g. https://msg.argo.grnet.gr)
-s, --sgas-url	Specify SGAS server URL (e.g. https://grid.uio.no:8001/logger)

APEL

Options to be used when target is specified using --apel-url

--apel-topic	Possible choices: gLite-APEL, /queue/global.accounting.test.cpu.central Define APEL topic (default is gLite-APEL)
--apel-messages	Possible choices: urs, summaries, summaries-v04 Define APEL messages (default is summaries)
--apel-project	Define APEL project (default is “accounting”)
--gocdb-name	(Re)define GOCDB site name

SGAS

Options to be used when target is specified using `--sgas-url`

--localid-prefix Define optional SGAS localid prefix

Other options

Works for both APEL and SGAS targets

--vofilter Republish only jobs owned by these VOs

--urbatchsize Size of records batch to be send (default is 50 for SGAS, 500 for APEL)

job

A-REX Jobs

```
arcctl job [-h] [-t CACHETTL] ACTION ...
```

Named Arguments

-t, --cachettl GM-Jobs output caching validity in seconds (default is 30)
Default: 30

Jobs Control Actions

ACTION Possible choices: list, script, log, info, stdout, stderr, attr, path, kill, killall, clean, cleanall, stats, accounting, datastaging
DESCRIPTION

Sub-commands:

list

List available A-REX jobs

```
arcctl job list [-h] [-l]
                  [-s {ACCEPTED,PREPARING,SUBMIT,INLRMS,FINISHING,FINISHED,DELETED,
↪ CANCELING}]
                  [-o OWNER]
```

Named Arguments

-l, --long	Detailed listing of jobs Default: False
-s, --state	Possible choices: ACCEPTED, PREPARING, SUBMIT, INLRMS, FINISH- ING, FINISHED, DELETED, CANCELING
	Filter jobs by state
-o, --owner	Filter jobs by owner

script

Display job script submitted to LRMS

```
arcctl job script [-h] jobid
```

Positional Arguments

jobid	Job ID
--------------	--------

log

Display job log

```
arcctl job log [-h] [-f] [-r] [-s] jobid
```

Positional Arguments

jobid	Job ID
--------------	--------

Named Arguments

-f, --follow	Follow the job log output Default: False
-r, --raw	Show raw logfile content as it is (including jobscrip) Default: False
-s, --service	Show ARC CE logs containing the jobID instead of job log Default: False

info

Show job main info

```
arcctl job info [-h] jobid
```

Positional Arguments

jobid	Job ID
--------------	--------

stdout

Show job executable stdout

```
arcctl job stdout [-h] [-f] jobid
```

Positional Arguments

jobid	Job ID
--------------	--------

Named Arguments

-f, --follow	Follow the job log output
	Default: False

stderr

Show job executable stderr

```
arcctl job stderr [-h] [-f] jobid
```

Positional Arguments

jobid	Job ID
--------------	--------

Named Arguments

-f, --follow	Follow the job log output
	Default: False

attr

Get job attribute

```
arcctl job attr [-h] jobid [attr]
```

Positional Arguments

jobid	Job ID
attr	Attribute name

path

Print control directory path for a job

```
arcctl job path [-h] jobid
```

Positional Arguments

jobid	Job ID
--------------	--------

kill

Cancel job

```
arcctl job kill [-h] jobid [jobid ...]
```

Positional Arguments

jobid	Job ID
--------------	--------

killall

Cancel all jobs

```
arcctl job killall [-h]
                    [-s {ACCEPTED,PREPARING,SUBMIT,INLRMS,FINISHING,FINISHED,DELETED,
↪ CANCELING}]
                    [-o OWNER]
```

Named Arguments

-s, --state	Possible choices: ACCEPTED, PREPARING, SUBMIT, INLRMS, FINISH- ING, FINISHED, DELETED, CANCELING
	Filter jobs by state
-o, --owner	Filter jobs by owner

clean

Clean job

```
arcctl job clean [-h] jobid [jobid ...]
```

Positional Arguments

jobid	Job ID
--------------	--------

cleanall

Clean all jobs

```
arcctl job cleanall [-h]
                    [-s {ACCEPTED,PREPARING,SUBMIT,INLRMS,FINISHING,FINISHED,DELETED,
↪CANCELING}]
                    [-o OWNER]
```

Named Arguments

-s, --state	Possible choices: ACCEPTED, PREPARING, SUBMIT, INLRMS, FINISH- ING, FINISHED, DELETED, CANCELING
	Filter jobs by state
-o, --owner	Filter jobs by owner

stats

Show jobs statistics

```
arcctl job stats [-h] [-l] [-t | -d]
```

Named Arguments

-l, --long	Detailed output of stats
	Default: False
-t, --total	Show server total stats
	Default: False
-d, --data-staging	Show server datastaging stats
	Default: False

accounting

Show job accounting data

```
arcctl job accounting [-h] ACTION ...
```

Job Accounting Actions

ACTION	Possible choices: info, events, transfers
	DESCRIPTION

Sub-commands:

info

Show job accounting data

```
arcctl job accounting info [-h]
                           [-o {all,description,resources,rtes,authtokens,json}]
                           jobid
```

Positional Arguments

jobid	Job ID
--------------	--------

Named Arguments

-o, --output	Possible choices: all, description, resources, rtes, authtokens, json
	Define what kind of job information you want to output (default is “all”)
	Default: “all”

events

Show job event history

```
arcctl job accounting events [-h] jobid
```

Positional Arguments

jobid	Job ID
--------------	--------

transfers

Show job data transfers statistics

```
arcctl job accounting transfers [-h] jobid
```

Positional Arguments

jobid	Job ID
--------------	--------

datastaging

Job Datastaging Information for jobs preparing or running.

```
arcctl job datastaging [-h] ACTION ...
```

Job Datastaging Menu

ACTION	Possible choices: get-totaltime, get-details, get-all
DESCRIPTION	

Sub-commands:

get-totaltime

Show the total time spent in the preparation stage for the selected job

```
arcctl job datastaging get-totaltime [-h] jobid
```

Positional Arguments

jobid	Job ID
--------------	--------

get-details

Show details related to the files downloaded for the selected job

```
arcctl job datastaging get-details [-h] jobid
```

Positional Arguments

jobid	Job ID
--------------	--------

get-all

Show details related to datastaging for all preparing jobs

```
arcctl job datastaging get-all [-h]
```

cache

ARC A-REX Cache control

```
arcctl cache [-h] ACTION ...
```

A-REX Cache Actions

ACTION	Possible choices: stats, list, is-cached
DESCRIPTION	

Sub-commands:

stats

Show cache usage statistics

```
arcctl cache stats [-h]
```

list

List cached URLs

```
arcctl cache list [-h] [-l]
```

Named Arguments

-l, --long	Output paths to cached files
	Default: False

is-cached

Checks is the URL already in A-REX cache

```
arcctl cache is-cached [-h] [-q] url
```

Positional Arguments

url URL to check

Named Arguments

-q, --quiet Do not output path to cached file
Default: False

datastaging

DataStaging info

```
arcctl datastaging [-h] ACTION ...
```

DataStaging Control Actions

ACTION Possible choices: job, summary, remote_dds, dtr
DESCRIPTION

Sub-commands:

job

Job Datastaging Information for a preparing or running job.

```
arcctl datastaging job [-h] ACTION ...
```

Job Datastaging Menu

ACTION Possible choices: get-totaltime, get-details, get-all
DESCRIPTION

Sub-commands:

get-totaltime

Show the total time spent in the preparation stage for the selected job

```
arcctl datastaging job get-totaltime [-h] jobid
```

Positional Arguments

jobid Job ID

get-details

Show details related to the files downloaded for the selected job

```
arcctl datastaging job get-details [-h] jobid
```

Positional Arguments

jobid Job ID

get-all

Show details related to datastaging for all preparing jobs

```
arcctl datastaging job get-all [-h]
```

summary

Job Datastaging Summary Information for jobs preparing or running.

```
arcctl datastaging summary [-h] ACTION ...
```

Job Datastaging Summary Menu

ACTION	Possible choices: jobs, files
	DESCRIPTION

Sub-commands:

jobs

Show overview of the duration of datastaging for all jobs INLRMS or PREPARING status. If no timewindow is specified, all active jobs will be accounted for. If a timewindow is specified, only jobs that have had any activity in the errors file in the specified timewindow are included.

```
arcctl datastaging summary jobs [-h] [-d DAYS] [-hr HOURS] [-m MINUTES]
                                [-s SECONDS] [-sh]
```

Named Arguments

-d, --days	Timewindow in days.
-hr, --hours	Timewindow in hours.
-m, --minutes	Timewindow in minutes.
-s, --seconds	Timewindow in seconds.
-sh, --showhist	Print out ASCII histograms of the datastaging summary. Default: False

files

Show the total number file and total file-size downloaded for all jobs INLRMS or PREPARING status. If no timewindow is specified, all active jobs will be accounted for. If a timewindow is specified, only jobs that have had any activity in the errors file in the specified timewindow are included.

```
arcctl datastaging summary files [-h] [-d DAYS] [-hr HOURS] [-m MINUTES]
                                [-s SECONDS]
```

Named Arguments

-d, --days	Timewindow in days.
-hr, --hours	Timewindow in hours.
-m, --minutes	Timewindow in minutes.
-s, --seconds	Timewindow in seconds.

remote_dds

Remote Datastaging file information for preparing jobs.

```
arcctl datastaging remote_dds [-h]
```

dtr

Data-delivery transfer (DTR) information

```
arcctl datastaging dtr [-h] ACTION ...
```

DTR info menu

ACTION	Possible choices: state
DESCRIPTION	

Sub-commands:

state

Show summary of DTR state info

```
arcctl datastaging dtr state [-h] state
```

Positional Arguments

state

cleanup

ARC CE housekeeping

```
arcctl cleanup [-h] ACTION ...
```

Cleanup Actions

ACTION	Possible choices: controldir, jwt-issuer, accounting, test-ca
DESCRIPTION	

Sub-commands:

controldir

Cleanup jobs metadata in controldir

```
arcctl cleanup controldir [-h] [--force] [--yes] [--dry-run]
                           (-g | -w | -f FINISHED_BEFORE)
```

Named Arguments

--force	Override any safety checks Default: False
--yes	Answer yes to all questions Default: False
--dry-run	Do not perform any actions, print the intended actions instead Default: False
-g, --gc	Cleanup leftovers in control directory Default: False
-w, --wipe	Cleanup all job records from control directory. Default: False
-f, --finished-before	Cleanup jobs finished before specified date (YYYY-MM-DD [HH:mm[:ss]])

jwt-issuer

Cleanup JWT issuer trust settings

```
arcctl cleanup jwt-issuer [-h] (-l | -i ISSUER)
```

Named Arguments

-l, --list	List JWT issuers trusted by ARC CE
	Default: False
-i, --issuer	JWT Issuer ID to cleanup

accounting

Cleanup accounting database

```
arcctl cleanup accounting [-h]
```

test-ca

Cleanup TestCA files

```
arcctl cleanup test-ca [-h] [--ca-id CA_ID] [--ca-dir CA_DIR]
```

Named Arguments

--ca-id	Define CA ID to work with (default is to use hostname-based hash)
--ca-dir	Redefine path to CA files directory

6.4.3 ARCHERY Manage Tool

The archery-manage tool used to simplify common operations with ARCHERY, including registry initial bootstrap, integration with topology databases and keeping dynamic information up to date.

```
usage: archery-manage [-h] [-d {CRITICAL,ERROR,WARNING,INFO,DEBUG}] -s SOURCE
                        [-f FILTER]
                        [-o {arc-CEs,services,endpoints,zonefile,json,_debug}]
                        [--json] [--output-all] [-u] [--domain DOMAIN]
                        [--ddns-master-ip DDNS_MASTER_IP]
                        [--ddns-tsig-keyfile DDNS_TSIG_KEYFILE]
                        [--ddns-tsig-algorithm {HMAC-MD5,HMAC-SHA1,HMAC-SHA224,HMAC-
↪SHA256,HMAC-SHA384,HMAC-SHA512}]
                        [--ttl TTL] [--threads THREADS] [--timeout TIMEOUT]
```

Named Arguments

-d, --debug	Possible choices: CRITICAL, ERROR, WARNING, INFO, DEBUG Default: "INFO"
-s, --source	Services topology source (use 'help' value to print available sources)
-f, --filter	Add endpoints filter (use 'help' value to print available filters)
-o, --output	Possible choices: arc-CEs, services, endpoints, zonefile, json, _debug Write requested data to stdout
--json	Change output format from plaintext to JSON Default: False
--output-all	Output all services/endpoints including inactive (filters are still applied) Default: False
-u, --ddns-update	Invoke DNS zone incremental DDNS update secured by TSIG key Default: False
--domain	Domain name of the ARCHERY endpoint to use (required for DDNS update)
--ddns-master-ip	Master DNS IP address (required for DDNS update)
--ddns-tsig-keyfile	TSIG keyfile (required for DDNS update)
--ddns-tsig-algorithm	Possible choices: HMAC-MD5, HMAC-SHA1, HMAC-SHA224, HMAC-SHA256, HMAC-SHA384, HMAC-SHA512 Cryptographic algorithm for TSIG Default: "HMAC-MD5"
--ttl	DNS resource records TTL value to use (default is 3600) Default: 3600
--threads	Number of threads to fetch information in parallel (default is 8) Default: 8
--timeout	Per-source information fetching timeout (default is 10 seconds) Default: 10

6.4.4 ARC Configuration Parser

Nordugrid ARC configuration parser

```
usage: arconfig-parser [-h] [--debug {CRITICAL,ERROR,WARNING,INFO,DEBUG}]
                    [--load] [--save] [-r RUNCONFIG] [-c CONFIG]
                    [-d DEFAULTS] [-b BLOCK] [-o OPTION] [-s]
                    [-e {bash,json}] [-f EXPORT_FILTER]
```

Named Arguments

--debug Possible choices: CRITICAL, ERROR, WARNING, INFO, DEBUG
 verbosity level (default is “WARNING”)
 Default: “WARNING”

Runtime configuration

Work with runtime configuration that includes default values

--load load ARC runtime configuration
 Default: False

--save save ARC runtime configuration
 Default: False

-r, --runconfig runtime config file location (default is “/run/arc/arc.runtime.conf”)
 Default: “/run/arc/arc.runtime.conf”

Configuration files

Initial ARC configuration files

-c, --config config file location (default is “/etc/arc.conf”)
 Default: “/etc/arc.conf”

-d, --defaults defaults file location (default is “/usr/share/arc/arc.parser.defaults”)
 Default: “/usr/share/arc/arc.parser.defaults”

Getting values

Get blocks and configuration option values

-b, --block block name (can be specified several times)

-o, --option option name

-s, --subblocks match subblocks against supplied block name(s)
 Default: False

-e, --export Possible choices: bash, json
 export configuration to the defined format

-f, --export-filter limit bash export to specified options only

6.5 NorduGrid repository information

Note

If you are looking for testing repo or nightly repo information, please see

- *Testing repo*
- *Nightlies repo*

The NorduGrid ARC packages are available through YUM and APT repositories for several systems. We have release-based repositories that you can follow. This will keep your install to a particular release of NorduGrid ARC with only minor and bug-fixing updates. You can also choose to follow the repository “latest” which will always point to the latest stable release. For each distribution there are 3 channels (repositories) available:

- *base* - Base packages (mandatory)
- *updates* - Updates to the base release (strongly recommended)
- *testing* - Packages almost ready to go into the updates repository (optional). The alpha, beta and release candidate releases can be found here.

6.5.1 Repository security

The NorduGrid RPM packages and DEB repositories are signed, and in order for the repository tools APT and YUM to verify them you must install the NorduGrid GPG key:

RHEL

Ubuntu

Debian

For rpm based distributions like Red Hat Enterprise Linux and Fedora:

```
[root ~]# rpm --import http://download.nordugrid.org/RPM-GPG-KEY-nordugrid-7
```

For Ubuntu distributions with sudo:

```
[user ~]$ wget -q http://download.nordugrid.org/DEB-GPG-KEY-nordugrid-7.asc -O- | 
↪ sudo apt-key add -
```

For Debian without sudo:

```
[root ~]# wget -q http://download.nordugrid.org/DEB-GPG-KEY-nordugrid-7.asc -O- | apt-  
↪ key add -
```

6.5.2 Repository configuration

RHEL

Ubuntu

Debian

The NorduGrid ARC repositories for RedHat Enterprise Linux / CentOS packaging utility dnf can be configured through: `/etc/yum.repos.d/nordugrid.repo`

The repository configuration can be set up automatically with dnf by installing the `nordugrid-release` package or creating the configuration file manually.

Install with dnf

Install manually

The easiest way to configure DNF to use the NorduGrid repository for Red Hat Enterprise Linux, CentOS and similar distributions is to install the `nordugrid-release` package which can be found in the NorduGrid package repository for the appropriate RHEL/EPEL release.

Links to the release packages:

Rocky Linux/AlmaLinux

CentOS Stream

Fedora

Rocky Linux: 10 9 8

CentOS Stream: 10 9

Fedora: 41 40

Install with `dnf` (Fedora, CentOS Stream, Rocky Linux, CentOS Linux 8+9) by copying the appropriate link from above

```
[root ~]# dnf install <rhel-repo link>
```

This creates the appropriate repo files in `/etc/yum.repos.d/`.

For manual YUM repository setup, create a file `/etc/yum.repos.d/nordugrid.repo` with the following contents (here using Rocky as example, if you are on Fedora, replace `rocky` with `fedora`)

If you are installing an alpha, beta or release candidate, please set the `nordugrid-testing` to `enabled=1`.

```
[nordugrid]
name=NorduGrid - $basearch - base
baseurl=http://download.nordugrid.org/repos/7/rocky/$releasever/$basearch/
↳base
enabled=1
gpgcheck=1
gpgkey=http://download.nordugrid.org/RPM-GPG-KEY-nordugrid-7

[nordugrid-updates]
name=NorduGrid - $basearch - updates
baseurl=http://download.nordugrid.org/repos/7/rocky/$releasever/$basearch/
↳updates
enabled=1
gpgcheck=1
gpgkey=http://download.nordugrid.org/RPM-GPG-KEY-nordugrid-7

[nordugrid-testing]
name=NorduGrid - $basearch - testing
baseurl=http://download.nordugrid.org/repos/7/rocky/$releasever/$basearch/
↳testing
enabled=0
gpgcheck=1
gpgkey=http://download.nordugrid.org/RPM-GPG-KEY-nordugrid-7
```

The NorduGrid ARC repositories for Ubuntu packaging utility APT can be configured through: `/etc/apt/sources.list` or when supported through a repo specific file: `/etc/apt/sources.list.d/nordugrid.list`.

The repository configuration can be set up automatically by means of installing `nordugrid-release` package or creating the configuration file manually.

Ubuntu version names:

Ubuntu release	Code name
25.04	plucky
24.04	noble
22.04	jammy

Install with `dpkg`

Install manually

The examples below give you the link for most recent Ubuntu releases. Packages are shown below for amd64 architecture. Replace amd64 for i386 if required for your architecture.

Ubuntu: 25.04 24.04 22.04

Install the source file with dpkg, example shown for Ubuntu 22.04:

```
[root ~]# wget -q https://download.nordugrid.org/packages/nordugrid-release/releases/
↪7/ubuntu/24.04/amd64/nordugrid-release_7~noble1_all.deb
[root ~]# dpkg -i nordugrid-release_7~noble1_all.deb
```

For a different version of Ubuntu, change the version names appropriately.

For manual APT repository setup for Ubuntu, the APT sources file should contain the following (here shown for Ubuntu 22.04 Jammy).

The configurations for the various APT based distributions can be found in the following sections. To enable a specific repository, remove the “#” from the beginning of the line, before the “deb” as shown for the Base Channel.

```
# Base channel - must be enabled
deb http://download.nordugrid.org/repos/7/ubuntu/ jammy main
deb-src http://download.nordugrid.org/repos/7/ubuntu/ jammy main

# Updates to the base release - should be enabled
deb http://download.nordugrid.org/repos/7/ubuntu/ jammy-updates main
deb-src http://download.nordugrid.org/repos/7/ubuntu/ jammy-updates main

# Scheduled package updates - optional
#deb http://download.nordugrid.org/repos/7/ubuntu/ jammy-experimental main
#deb-src http://download.nordugrid.org/repos/7/ubuntu/ jammy-experimental main
```

For a different release version, change the version name accordingly.

The NorduGrid ARC repositories for Ubuntu packaging utility APT can be configured through: /etc/apt/sources.list or when supported through a repo specific file: /etc/apt/sources.list.d/nordugrid.list.

The repository configuration can be set up automatically by means of installing nordugrid-release package or creating the configuration file manually.

Debian version names:

Debian release	Code name
13	trixie
12	bookworm
11	bullseye

Install with dpkg

Install manually

The examples below give you the link for most recent Debian releases. Packages are shown below for amd64 architecture. Replace amd64 for i386 if required for your architecture.

Debian: 13 12 11

Install the source file with dpkg, example shown for Debian 12:

```
[root ~]# wget -q https://download.nordugrid.org/packages/nordugrid-release/releases/
↪7/debian/12/amd64/nordugrid-release_7~bpo12+1_all.deb
[root ~]# dpkg -i nordugrid-release_7~bpo12+1_all.deb
```

For a different version of Debian, change the version names appropriately.

For manual APT repository setup for Debian, the APT sources file should contain the following (here shown for Debian 12 Bookworm).

The configurations for the various APT based distributions can be found in the following sections. To enable a specific repository, remove the “#” from the beginning of the line, before the “deb” as shown for the Base Channel.

```
# Base channel - must be enabled
deb http://download.nordugrid.org/repos/7/debian/ bookworm main
deb-src http://download.nordugrid.org/repos/7/debian/ bookworm main

# Updates to the base release - should be enabled
deb http://download.nordugrid.org/repos/7/debian/ bookworm-updates main
deb-src http://download.nordugrid.org/repos/7/debian/ bookworm-updates main

# Scheduled package updates - optional
#deb http://download.nordugrid.org/repos/7/debian/ bookworm-experimental main
#deb-src http://download.nordugrid.org/repos/7/debian/ bookworm-experimental main
```

For a different release version, change the version name accordingly.

6.5.3 Finish setup

RHEL

Ubuntu

Debian

RHEL8

RHEL9

```
[root ~]# dnf config-manager --set-enabled powertools && dnf makecache
```

```
[root ~]# dnf config-manager --set-enabled crb && dnf makecache
```

```
[root~]# apt-get update
```

```
[root~]# apt-get update
```

6.5.4 Install ARC packages

Please refer to the *ARC Computing Element Installation and Configuration Guide* for package selection and configuration.

6.6 ARC5 to ARC6 Migration Guide

ARC 6 comes with major re-engineering in the code base leading to improvements in the deployment, configuration, maintenance and scalability of ARC CE. It introduces new features while some legacy interfaces, components and features have been obsoleted and removed.

Despite all the code changes the supported ARC 6 CE interfaces are unchanged as well as the structure of the control directory, session directory, accounting logs, etc. This makes it possible to upgrade ARC5 to ARC6 without a complete re-installation “from scratch”.

Note

Despite the fact that we haven't observed any issues during normal package upgrades it is recommended to perform a clean install "from scratch" wherever possible in order to eliminate hidden bugs and ARC5 leftovers.

Following sections provides you with the necessary info about the major steps of an ARC5 to ARC6 migration.

6.6.1 Quick reference

Here is a quick **6 step** checklist for what to do:

1. Rewrite your `arc.conf`: you can *send your configuration to us* and we will help to rework it for ARC6.
2. Learn the ARC6 packages and services names.
3. Get familiar with *ARC Control Tool* that will simplify your life as an ARC6 administrator.
4. Don't forget to enable *RunTime Environments (RTE)* with `arcctl rte enable`.
5. Be aware of *accounting subsystem* changes.
6. Say "Good Bye!" to EGIIS, the ARC community is migrating over to a DNS-based ARCHERY system.

The details:

6.6.2 Configuration file

Because of the complete overhaul of the server-side ARC CE configuration, an existing ARC5 configuration can not be used with an ARC6 installation, it will immediately fail during the configuration validation process of the ARC service startup. Therefore you must create a new ARC6 configuration from scratch.

Note

There is no automatic converter from an old ARC5 configuration file to the new ARC6 syntax. The main reason for this is that there are many conceptual changes rather than simple option and block renaming changes.

We offer support in configuration rework via [nordugrid-discuss](#) mailing list or [ARC6 Support Skype Channel](#).

Just send your configuration to us!

When reworking the configuration take notice of these **syntax changes**:

- no more quotes in option values
- blocks are now not only grouping options, but have a enable/disable functionality
- in block names keywords are separated by `/` to represent subblocks (if such exist), while block identifiers are separated by `:` from keywords

Note

It is advised to read the *General configuration structure* section of ARC configuration reference document before starting any rewrite.

There are also **changes in the configuration approach** in several areas. The major highlights are:

- The way *authorization and user mapping* is handled has been greatly reworked and improved. In particular mapping via *gridmap files* is NO longer required nor recommended. If you still have *gridmap files* in your configuration consider whether you really want to keep this approach in ARC6.

- Jura *accounting configuration* is now configured within dedicated blocks.
- Infosys schemas are now also enabled/disabled with blocks. In particular the default ARC5 NorduGrid LDAP schema should be explicitly enabled with *[infosys/nordugrid] block*.

Warning

ARC6 configuration also comes with consistent defaults for everything. If a block is enabled all configuration options in this block that has no specific value in `arc.conf` get their default values.

Configuration with all defaults included can be printed with `arcctl config dump`.

Notice that any changes in configuration requires restart off all affected ARC services.

Note

Config validation is enforced during A-REX start or can be run manually with `arcctl config validate`. So don't worry if you had misspelled something, the validator will catch it!

6.6.3 Operating ARC services

ARC6 also comes with the a revised consistent packages and services naming scheme. We made the service and package names more clear and consistent with configuration block names in `arc.conf`.

Log file locations and internal format have been reviewed to provide consistency across the ARC subsystems. Log file locations can be printed with `arcctl config brief`.

ARC Control Tool

To simplify common ARC CE operations, ARC6 introduces *ARC Control Tool* - the single entry point to administrate and work with various ARC and third-party components. The `arcctl` tool helps ARC site admins by wrapping typical complex operations and unifying ARC CE administration under a single utility.

The `arcctl` tool was designed with BASH-completion in mind. It completes subsystem names, arguments, jobIDs, certificate DNs, RTE names, etc.

Note

To get familiar with some of the `arcctl` functionality it is advised to follow `/admins/try_arc6` tutorial. Using the `arcctl` tool to administrate an ARC6 site is recommended even for the most experienced ARC5 admins.

RunTime Environments

The new framework for *RunTime Environments (RTE)* has gotten a redesigned architecture and a largely extended functionality (default, installed, enabled RTEs, introduction of RTE parameters).

The major conceptual and behaviour differences are:

- in ARC6 RTE directories are local to the ARC CE head node and *SHOULD NOT* be shared to worker nodes
- in ARC6 the RTE concept got extended by the introduction of RTE types such as default, installed, enabled, dummy, etc.. An ordinary ARC5 RTE would correspond to an installed user-defined ARC6 RTE. Please read the documentation.
- All installed ARC5 RTEs (available in a directory) should be enabled explicitly (e.g. `arcctl rte enable ENV/PROXY`) to be used on an ARC6 CE.

Note

You can use wildcards to enable all RTEs you have in one shot: `arcctl rte enable '*'`

Accounting

In ARC 6.0 *Jura accounting* had redesigned archive structure that introduces an archived records database and comes with a redesigned record republishing to accounting services, in addition to the already mentioned configuration changes.

To make ARC 6.0-6.3 accounting compatible with ARC5 jura, an automatic archive conversion is done by an additional jura-archive-manager process run by A-REX.

Warning

The JURA archive conversion consumes CPU and I/O resources of your ARC CE. If you are upgrading and have a lot of records in archive already expect an additional load during conversion.

In ARC 6.4 the *next generation accounting subsystem* had been introduced. Instead of published records archiving it introduced local SQLite database that stores A-REX accounting records for each job for both on-site CE auditing and publishing.

Warning

If you want to keep old archived usage records from ARC 5 release it is advised to upgrade from ARC 5 to ARC 6.3 first, wait for records conversion completion, then update to ARC 6.4 or higher versions.

6.6.4 Retiring the EGIIS

All configuration blocks and services in ARC5 that correspond to registration to EGIIS are now gone in ARC6.

ARCHERY, the new DNS-based service endpoint catalogue for ARC is replacing EGIIS and it works *top-to-bottom* without explicit registration, i.e. nothing registration related needs to be configured in the `arc.conf` of an ARC6 CE.

If your CE had been previously registered in the country-level Nordugrid EGIIS your ARC CE registration is already migrated to the nordugrid.org ARCHERY instance.

If you find your CE missing (e.g. in the [nordugrid top-level monitor](#)) please contact the country-level ARCHERY instance admins as described in [this document](#).

6.7 Main changes in ARC 6 compared to ARC 5

There are many improvements and changes in the new ARC release, and with these we hope to have greatly enhanced the ease of setting up, configuring and managing an ARC production site, and to have improved the reliability and scalability of the ARC CE by the internal restructuring ARC has undergone.

Note

Despite all the new features and code changes, the supported ARC 6 CE interfaces are unchanged - the latest ARC 5 clients are compatible with an ARC 6 CE and vice versa.

- Complete overhaul of server-side *ARC CE configuration*:

- Configuration has been completely reworked, cleaned up, regrouped and made more streamlined.
 - In particular, the way authorization and user mapping is handled has been greatly improved. Worth mentioning is the support of higher granularity vo-based authorization, and that mapping via gridmap files is no longer required nor recommended.
 - ARC now comes with a so-called zero configuration - a preconfigured minimal ARC setup automatically installed with ARC, including test-CA and test-host certificate for immediate test submission.
 - Default values throughout ARC have had a big tidy up and are now handled consistently through the ARC runtime configuration.
 - Configuration blocks are now used to enable (turn on) or disable a functionality, a service or an interface. For example, the NorduGrid schema is no longer published unless it is enabled via the corresponding configuration block.
 - Validation of the ARC configuration is enforced by A-REX startup scripts. If you have configuration errors then A-REX will not start, and you will be pointed to the error.
- **Scalability and manageability improvements:**
 - The internal job-loop in A-Rex has been re-engineered to be event-driven
 - xrootd plugin has been re-written to improve data transfer performance
 - Consistent handling of server-side logging and improved logfile structure
 - Reworked startup scripts
 - Streamlined package name scheme, where e.g. ldap services have been separated out
 - *The RTE framework* has got a redesigned architecture and a largely extended functionality (default, installed and enabled RTEs, introduction of RTE parameters)
 - **A new framework for *RunTimeEnvironments (RTE)*:**
 - The RTE framework has gotten a redesigned architecture and largely extended functionality (default, installed, enabled RTEs, introduction of RTE parameters).

Note

RTEs must in ARC 6 be *explicitly enabled through the new `arcctl` tool* as a separate step after installation.

- **There are several system RTEs installed together with ARC which you can enable on demand. These are:**
 - * ENV/CANDYPOND: makes ARC Candypond (“Cache and deliver your pilot on-demand data”) client available on the Worker Nodes and ready to be used
 - * ENV/CONDOR/DOCKER: enables submission to Docker universe in HTCondor backend
 - * ENV/LRMS-SCRATCH: enables the usage of local WN scratch directory defined by LRMS
 - * ENV/PROXY: copies proxy certificate to the job session directory
 - * ENV/RTE: copies RunTimeEnvironment scripts to the job session directory
- ***ARCCTL*, a new server-side management and control tool for sysadmins has been developed.**
 - `arcctl` is meant to be the central one-stop-shop tool for administrating an ARC CE
 - With this tool you can handle RTEs, the jobs, accounting, ARC services and many other things related to an ARC CE.
 - `arcctl` also offers help with integrating 3rd party services and components with an ARC CE such as handling CAs, VOMS, etc..

- ***JURA Accounting* subsystem improvements:**
 - Accounting configuration, archive and operations has been restructured to improve typical operations.
 - Improved republishing via `arcctl`.
 - Better integration with APEL SSM.
 - Archive structure is relying on a database for republishing and gathering stats locally.
- ***ARCHERY*, the new DNS-based service endpoint catalogue for ARC**
 - `archery-manage` package to populate the ARCHERY registry with ARC CE endpoint info
 - Official top-level DNS registry under the `nordugrid.org` domain
 - ARCHERY-compatible clients to submit and monitor jobs (`arc cli`, `ldap monitor`)
 - ARCHERY replaces the OBSOLETE EGIIS service of the ARC 5 release series.
- **TECHNOLOGY PREVIEW components:**
 - The REST interface - enabled together with the A-Rex Web-Service.
 - The INTERNAL interface - a special interface aimed for restrictive HPC sites, to be used with a local installation of the ARC Control Tower.
 - *Candypond* - “Cache and deliver your pilot on-demand data” service to bring the power of ARC Cache to the pilot job universe.
 - Experimental Python LRMS for SLURM with SSH access support.
- **The ARC client and SDK:**
 - The client has undergone a major internal cleanup, large number of submission, target retrieval and job description plugins got removed as a result of OBSOLETE third-party interfaces and technologies in ARC.
 - New plugins for ARCHERY and the REST And INTERNAL interfaces were implemented.
 - The `arcstat` cli of the ARC6 client and the underlying SDK now handles walltime as $(wall-time) * (\text{number of cores})$ for multicore jobs.

6.7.1 REMOVED and obsoleted components

- Large number of configuration options of `arc.conf` got DELETED during the server-side configuration rework. A detailed list of those options are available in the `arc.conf.DELETED` file.
- The separate `gangliarc` component is discontinued, instead we now offer `ganglia` integrated into AREX.
- Nordugrid no longer distributes a source code bundle including documentation, `gangliarc` and `nagios`.
- CREAM related plugins and server-side components.
- UNICORE related plugins and server-side components.
- All the windows & solaris related ARC code and build.
- JAVA bindings for ARC SDK.
- Support for BES, WSRF and other non-EMIES WS flavours as job management interfaces.
- Support for JDL, ARCJSDL, RSL job description dialects: ARC keeps only XRSL and ADL as supported job description languages.
- EMIR indexing service including ARC CE registration to EMIR and ARC SDK plugins.
- Server-side EGIIS, including ARC CE registration to EGIIS. BUT: keep EGIIS plugins in the client.
- GLUE1 support from the ARC client SDK, server-side support for GLUE1 temporarily kept but labelled as OBSOLETE.

- ARC SDK support to obtain information from top-BDII.
- Publishing nordugrid-authuser objects in the Nordugrid LDAP schema

6.7.2 List of bugs fixed since ARC 5

A rather long (but not exhaustive) list of bugs fixed since ARC 5 can be found here: <http://www.nordugrid.org/arc/arc6/common/changelog/bugs-6.0.0.html>

6.8 List of bugs fixed since ARC 5.4.4

Table 6.4: bugs_600

Bug ID	Summary
BUGZ-895	ARC Administrative tools
BUGZ-1421	No log of backend script problems
BUGZ-3034	arcstat -s is case-sensitive
BUGZ-3106	arcproxy could be more expressive when it finds a problem
BUGZ-3360	SGE and LL backends fail to report correct node OS/system information in GLUE2 info-system
BUGZ-3384	Support for per-queue authorisation configuration and publishing
BUGZ-3451	Configure number of cpus manually
BUGZ-3476	Crashes in multiple ARC components
BUGZ-3545	Patch for Correct Cores Parsing
BUGZ-3557	undetected job submission in case of heavy filesystem load
BUGZ-3565	Allow setting default VO in ~/.arc/client.conf to be used in arcproxy
BUGZ-3566	Implement RTEs processing without shared directory
BUGZ-3569	Exporting CPU/Wall time limits to Glue2/BDII
BUGZ-3570	Arcsub used 100GB memory
BUGZ-3584	JURA: create one log file per job, not per submission
BUGZ-3624	Data delivery service can only listen to one network interface.
BUGZ-3626	Force GLUE2ComputingManagerTotalLogicalCPUs to be totalcpus when this value is defined in arc.conf
BUGZ-3632	arcproxy fails in Ubuntu 16.04, 16.10 and recent Debian systems
BUGZ-3637	arcget with multiple jobs crashes
BUGZ-3643	The watchdog crashes
BUGZ-3662	arcsub crashes
BUGZ-3667	JSON output for arcstat
BUGZ-3674	settings in client.conf ignored
BUGZ-3675	Problems retrieving jobs with arcget 1
BUGZ-3676	Problems retrieving jobs with arcget 2 (first byte missing)
BUGZ-3677	arcproxy fails retrieving attributes from voms
BUGZ-3682	Better error message when DN not in gridmap file
BUGZ-3690	Warnings about missing information on deleted jobs
BUGZ-3695	Slowness with arccp and the xrootd protocol
BUGZ-3700	ARC1ClusterInfo.pm uses netstat
BUGZ-3702	Cannot use arc data commands without certificates
BUGZ-3707	Seg fault Triolith - related or not to the data-staging problems on Triolith
BUGZ-3713	Malformed jobs.dat entries
BUGZ-3722	Xenial repo for the (nordugrid) ARC source is not correct
BUGZ-3756	package update failed to restart A-REX
BUGZ-3772	Can't use different credentials within one submission process when jobs require user input files
BUGZ-3773	Enabling arex-ganglia breaks controldir access

continues on next page

Table 6.4 – continued from previous page

Bug ID	Summary
BUGZ-3778	arcctl not working if arc code configured with <code>--disable-swig</code> flag and installed with <code>make install</code>
BUGZ-3788	Poor performance with <code>arcep</code> and HTTPS
BUGZ-3812	A-REX hangs inside XRootd after fork

TECHNICAL DOCUMENTS DESCRIBING ARC COMPONENTS

The following documents give a deep technical description of the various ARC components. If you are looking for architecture internals (how parts of ARC were designed) you can follow this section.

7.1 ARC Data Services Technical Description

7.1.1 A-REX Data Cache technical description

Structure of the cache directory

Cached files are stored in sub-directories under the data directory in each main cache directory. Filenames are constructed from an SHA-1 hash of the URL of the file and split into subdirectories based on the two initial characters of the hash. In the extremely unlikely event of a collision between two URLs having the same SHA-1 hash, caching will not be used for the second file.

When multiple caches are used, a new cache file goes to a randomly selected cache, where each cache is weighted according to the size of the file system on which it is located.

For example: if there are two caches of 1TB and 9TB then on average 10% of input files will go to the first cache and 90% will go to the second cache.

Some associated metadata including the corresponding URL and an expiry time, if available, are stored in a file with the same name as the cache file, with a `.meta` suffix.

For example, with a cache directory `/cache` the file `srn://srn.nordugrid.org/grid/atlas/file1:`

- is mapped to `/cache/data/37/b19acc950c37876a61d2de6e238d38c9e94c0`,
- the file `/cache/data/37/b19acc950c37876a61d2de6e238d38c9e94c0.meta` contains the original URL and an expiry time if one is available.

At the start of a file download, the cache file is locked, so that it cannot be deleted and so that another download process cannot write the same file simultaneously. This is done by creating a file with the same name as the cache filename but with a `.lock` suffix. This file contains the process ID of the process and the hostname of the host holding the lock. If this file is present, another process cannot do anything with the cache file and must wait until the cache file is unlocked (i.e. the `.lock` file no longer exists). The lock is continually updated during the transfer, and is considered stale if 15 minutes have passed since the last update. These stale locks, caused for example by a download process exiting abnormally, will therefore automatically be cleaned up. Also, if the process corresponding to the process ID stored inside the lock is no longer running on the host specified in the lock, it is safe to assume that the lock file can be deleted. If a file is requested which already exists in the cache (and is not locked), the cache file is not locked, but checks are done at the end of cache processing to ensure the file was not modified during the processing.

How the cache works

If a job requests an input file which can be cached or is allowed to be cached, it is stored in the selected cache directory, then a hard link is created in a per-job directory, under the `joblinks` subdirectory of the main cache directory. Then depending on the configuration, either the hard-link is copied or soft-linked to the SD. The former option is advised if the cache is on a file system which will suffer poor performance from a large number of jobs reading files on it, or the file system containing the cache is not accessible from worker nodes. The latter option is the default option. Files marked as executable in the job will be stored in the cache without executable permissions, but they will be copied to the SD and the appropriate permissions applied to the copy.

The per-job directory is only readable by the local user running the job, and the cache directory is readable only by the A-REX user. This means that the local user cannot access any other users' cache files. It also means that cache files can be removed without needing to know whether they are in use by a currently running job. However, as deleting a file which has hard links does not free space on the disk, cache files are not deleted until all per-job hard links are deleted.

Warning

If a cache is mounted from an NFS server and the A-REX is run by the root user, the server must have the `no_root_squash` option set for the A-REX host in the `/etc/exports` file, otherwise the A-REX will not be able to create the required directories.

Note

Note that when running A-REX under a non-privileged user account, all cache files will be owned and accessible by the same user, and therefore modifiable by running jobs. This is potentially dangerous and so caching should be used with caution in this case.

If the file system containing the cache is full and it is impossible to free any space, the download fails and is retried without using caching.

Before giving access to a file already in the cache, the A-REX contacts the initial file source to check if the user has read permission on the file. In order to prevent repeated checks on source files, this authentication information is cached for a limited time. On passing the check for a cached file, the user's DN is stored in the `.meta` file, with an expiry time equivalent to the lifetime remaining for the user's proxy certificate. This means that the permission check is not performed for this user for this file until this time is up (usually several hours). File creation and validity times from the original source are also checked to make sure the cached file is fresh enough. If the modification time of the source is later than that of the cached file, the file will be downloaded again. The file will also be downloaded again if the modification date of the source is not available, as it is assumed the cache file is out of date. These checks are not performed if the DN is cached and is still valid.

The A-REX checks the cache periodically if it is configured to do automatic cleaning. If the used space on the file system containing the cache exceeds the high water-mark given in the configuration file it tries to remove the least-recently accessed files to reduce size to the low water-mark.

Cache cleaning

When `[arex/cache/cleaner] block` is defined the cache is cleaned automatically periodically (every 5 minutes) by the A-REX to keep the size of each cache within the configured limits. Files are removed from the cache if the total size of the cache is greater than the configured limit. Files which are not locked are removed in order of access time, starting with the earliest, until the size is lower than the configured lower limit. If the lower limit cannot be reached (because too many files are locked, or other files outside the cache are taking up space on the file system), the cleaning will stop before the lower limit is reached.

Since the limits on cache size are given as a percentage of space used on the filesystem on which the cache is located, it is recommended that each cache has its own dedicated file system.

If the cache shares space with other data on a file system, the option `calculatesize=cachedir` should be set in `arc.conf` so that the cache limits are applied on the size of the cache rather than the file system.

With large caches mounted over NFS and an A-REX heavily loaded with data transfer processes, cache cleaning can become slow, leading to caches filling up beyond their configured limits. For performance reasons it may be advantageous to disable cache cleaning by the A-REX, and run the `cache-clean` tool (usually `/usr/libexec/arc/cache-clean`) independently on the machine hosting the file system.

Caches can be added to and removed from the configuration as required without affecting any cached data, but after changing the configuration file, the A-REX should be restarted. If a cache is to be removed and all data erased, it is recommended that the cache be put in a *draining* state until all currently running jobs possibly accessing files in this cache have finished. In this state the cache will not be used by any new jobs, but the hard links in the *joblinks* directory will be cleaned up as each job finishes. Once this directory is empty it is safe to delete the entire cache.

Caches may also be marked as *read-only*, so that data cached there can be used by new jobs, but no new data will be written there. Note that read-only caches are not cleaned by A-REX.

Exposing the Cache

Normally the ARC cache is internal to the CE and is not exposed to the outside. However it may be beneficial to allow reading cache files, if for example the file is lost from Grid storage or as a fallback when Grid storage is down. This can be done via HTTPS through the A-REX web services interface.

Specifying `[arex/ws/cache] block` opens remote read access to certain cache files for certain credential properties. When configured this allows cached files to be read from the A-REX WS endpoint, for example if file `gsiftp://my.host/file1` is cached at CE `a-rex.host` the file is accessible (if credentials allow) at:

```
https://a-rex.host/arex/cache/gsiftp://my.host/file1
```

Since remote reading can increase the load on A-REX, the number of concurrent requests should be limited. This can be done using the `max_data_transfer_requests` configuration option.

7.1.2 A-REX data transfer framework (DTR) technical description

This page describes the data staging framework for ARC, code-named DTR (Data Transfer Reloaded).

Overview

ARC's Computing Element (A-REX) performs the task of data transfer for jobs before and after the jobs run. The requirements and the design steps for the data staging framework are described in *DTR Design and Implementation Details*. The framework is called DTR (Data Transfer Reloaded) and uses a three-layer architecture, shown in the figure below:

The Generator uses user input of tasks to construct a Data Transfer Request (also DTR) per file that needs to be transferred. These DTRs are sent to the Scheduler for processing. The Scheduler sends DTRs to the Pre-processor for anything that needs to be done up until the physical transfer takes place (e.g. cache check, resolve replicas) and then to Delivery for the transfer itself. Once the transfer has finished the Post-processor handles any post-transfer operations (e.g. register replicas, release requests). The number of slots available for each component is limited, so the Scheduler controls queues and decides when to allocate slots to specific DTRs, based on the prioritisation algorithm implemented. See *DTR priority and shares system* for more information.

This layered architecture allows any implementation of a particular component to be easily substituted for another, for example a GUI with which users can enter DTRs (Generator) or an external point-to-point file transfer service (Delivery).

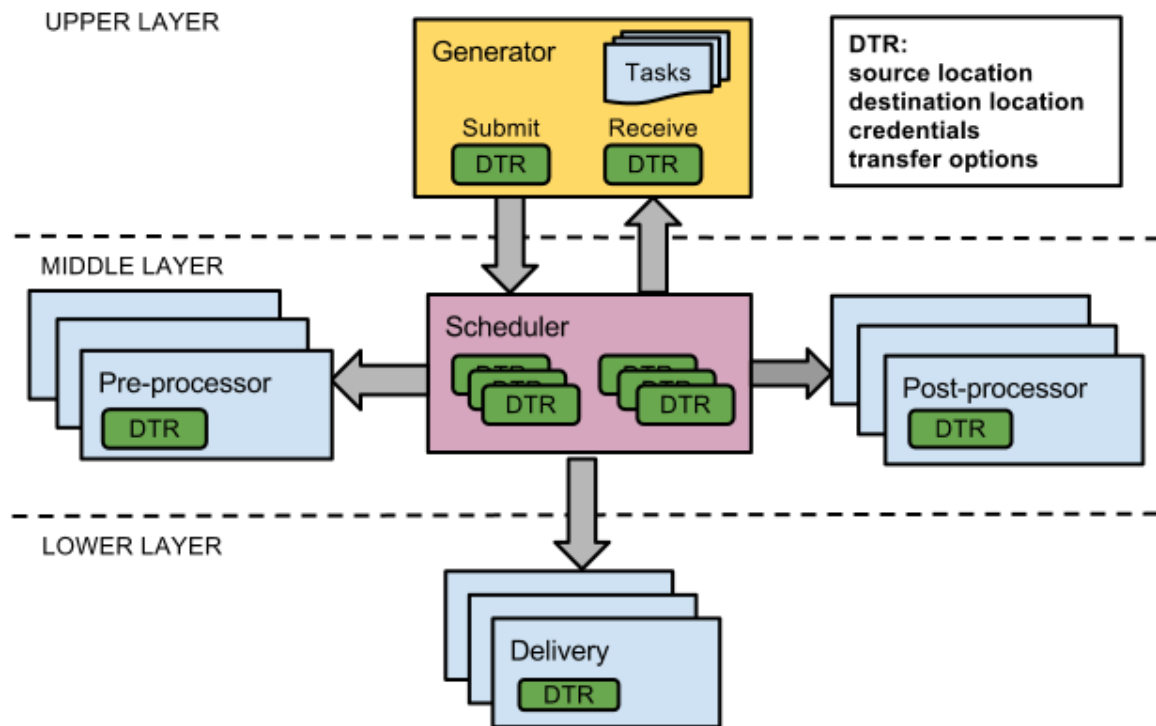


Fig. 7.1: DTR three-layer architecture

Implementation

The middle and lower layers of the architecture (Scheduler, Processor and Delivery) are implemented as a separate library **libarcdatastaging** (in `src/libs/data-staging` in the ARC source tree). This library is included in the **nordugrid-arc** common libraries package. It depends on some other common ARC libraries and the DMC modules (which enable various data access protocols and are included in **nordugrid-arc-plugins-*** packages) but is independent of other components such as A-REX or ARC clients. A simple Generator is included in this library for testing purposes. A Generator for A-REX is implemented in `src/services/a-rex/grid-manager/jobs/DTRGenerator.(h|cpp)`, which turns job descriptions into data transfer requests.

Configuration

Data staging is configured through the `[arex/data-staging]` block in `arc.conf`. Reasonable default values exist for all parameters but the `[arex/data-staging]` block can be used to tune the parameters, and also enable *multi-host data staging*. A selection of parameters are shown below:

Parameter	Explanation	Default Value
maxdelivery	Maximum delivery slots	10
maxprocessor	Maximum processor slots per state	10
maxemergency	Maximum emergency slots for delivery and processor	1
maxprepared	Maximum prepared files (for example pinned files using SRM)	200
sharetype	Transfer share scheme (dn, voms:vo, voms:group or voms:role)	None
definedshare	Defined share and priority	_default 50
Multi-host related parameters		
deliveryservice	URL of remote host which can perform data delivery	None
localdelivery	Whether local delivery should also be done	no
remotesizelimit	File size limit (in bytes) below which local transfer is always used	0
usehostcert	Whether the host certificate should be used in communication with remote delivery services instead of the user's proxy	no

Description of other data-staging parameters can be found in [\[arex/data-staging\] block](#). The multi-host parameters are explained in more detail in [ARC Data Delivery Service Technical Description](#)

Example:

```
[data-staging]
maxdelivery = 10
maxprocessor = 20
maxemergency = 2
maxprepared = 50
sharetype = voms:role
definedshare = myvo:production 80
deliveryservice = https://spare.host:60003/datadeliveryservice
localdelivery yes
remotesizelimit = 1000000
```

Client-side priorities

To specify the priority of jobs on the client side, the `priority` element can be added to an XRSL job description, eg:

```
("priority" = "80")
```

For a full explanation of how priorities work see [DTR priority and shares system](#).

gm-jobs -s

The command “gm-jobs -s” to show transfer shares information now shows the same information at the per-file level rather than per-job. The number in “Preparing” are the number of DTRs in TRANSFERRING state, i.e. doing physical transfer. Other DTR states count towards the “Pending” files. For example:

Preparing/Pending files	Transfer share
2/86	atlas:null-download
3/32	atlas:production-download

As before, per-job logging information is in the controldir/job.id.errors files, but A-REX can also be configured to log all DTR messages to a central log file in addition through the logfile parameter.

Supported Protocols

The following access and transfer protocols are supported. Note that third-party transfer is not supported.

- file
- HTTP(s/g)
- GridFTP
- SRM
- Xrootd
- LDAP
- Rucio
- S3
- RFIO/DCAP/LFC (through GFAL2 plugins)

Multi-host Data Staging

To increase overall bandwidth, multiple hosts can be used to perform the physical transfers. See [ARC Data Delivery Service Technical Description](#) for details.

Monitoring

In A-REX the state, priority and share of all DTRs is logged to the file controldir/dtr.state periodically (every second). This can then be used by the Gangliarc framework to show data staging information as ganglia metrics.

Advantages

DTR offers many advantages over the previous system, including:

High performance

When a transfer finishes in Delivery, there is always another prepared and ready, so the network is always fully used. A file stuck in a pre-processing step does not block others preparing or affect any physical transfers running or queued. Cached files are processed instantly rather than waiting behind those needing transferred. Bulk calls are implemented for some operations of indexing catalogs and SRM protocols.

Fast

All state is held in memory, which enables extremely fast queue processing. The system knows which files are writing to cache and so does not need to constantly poll the file system for lock files.

Clean

When a DTR is cancelled mid-transfer, the destination file is deleted and all resources such as SRM pins and cache locks are cleaned up before returning the DTR to the Generator. On A-REX shutdown all DTRs can be cleanly cancelled in this way.

Fault tolerance

The state of the system is frequently dumped to a file, so in the event of crash or power cut, this file can be read to recover the state of ongoing transfers. Transfers stopped mid-way are automatically restarted after cleaning up the half-finished attempt.

Intelligence

Error handling has vastly improved so that temporary errors caused by network glitches, timeouts, busy remote services etc are retried transparently.

Prioritisation

Both the server admins and users have control over which data transfers have which priority.

Monitoring

Admins can see at a glance the state of the system and using a standard framework like Ganglia means admins can monitor ARC in the same way as the rest of their system.

Scaleable

An arbitrary number of extra hosts can be easily added to the system to scale up the bandwidth available. The system has been tested with up to tens of thousands of concurrent DTRs.

Configurable

The system can run with no configuration changes, or many detailed options can be tweaked.

Generic flexible framework

The framework is not specific to ARC's Computing Element (A-REX) and can be used by any generic data transfer application.

Open Issues

- Provide a way for the infosys to obtain DTR status information
 - First basic implementation: when DTR changes state write current state to .input or .output file
- Decide whether or not to cancel all DTRs in a job when one fails
 - Current logic: if downloading, cancel all DTRs in job, if uploading don't cancel any
 - Should be configurable by user - also EMI execution service interface allows specifying per-file what to do in case of error
- Priorities: more sophisticated algorithms for handling priorities
- Advanced features such as pausing and resuming transfers

Related Documents

DTR Design and Implementation Details

This page documents the requirements and design stages of the new data staging framework which took place around mid-2010.

Issues with previous implementation

1. Queueing happens per job - that makes it impossible to use potentially more effective processing order
2. Slow data transfers block those which could be done faster.
3. Big data transfers block small ones.
4. Jobs waiting for already cached (or to be cached) files are blocked by other jobs in queue.
5. Special features of sophisticated protocols are not taken into account - like SRM's "try later".
6. No priorities aka flexible queues.
7. No support for different credentials for different files.
8. No bandwidth handling.
9. No handling of files with different access latency (eg tape vs disk)
10. No mechanism to choose a preferred replica for the LFC (catalog) inputs, for example: if replicas are in ndgf, swegrid, signet, unige, try with ndgf first.

Task Summary

The initial task was NOT to solve all these issues. The task was to create a framework which could be extended to solve them later, or to find/adopt such a framework.

Requirements

1. Effective usage of bandwidth. Whenever any transfer is paused due to any reason (tape stage in, retry later) for estimated time another transfer should use the available bandwidth.
2. Transfer negotiation (with protocols such as SRM) should be independent of physical data transfer.
3. Each transfer should be capable to use its own credentials.
4. Transfer should be capable of pausing (temporary cancel) and resuming (if protocol allows).
5. Automatic/dynamic redistribution of bandwidth is needed to allow short transfers to pass through even while big transfers are taking the whole bandwidth.
6. Transfer from multiple alternative locations.
7. Cache checks should happen independently of data transfer to avoid locks.
8. Jobs where all files are cached should be processed immediately.
9. Better description of file source/destination than just URL (options are difficult to handle, something nicer is needed)
10. Priorities at different levels: among user groups, inside groups. Any other levels? 3 possible levels: among VO, users/roles inside VO, inside user identity.
11. Ability for users to set relative priority of their own jobs, both before submission and while job is in queue.

Security Requirements

- It must be built into the design that no process has higher privileges than necessary
- Elevated privileges are required for:
 - Access to cache. Cache files are only read/writeable by root user so they cannot be modified by jobs
 - Access to session and control directories. Access to these directories should be performed under the uid of the job owner. The current method of running as root and chown'ing must not be used.
- Elevated privileges are not required for any other parts of the system such as scheduling

Performance Requirements

- Must scale higher than current highest workloads
- Must be able to handle up to 10000 active jobs (between ACCEPTED and FINISHED)
- Must be able to handle up to 1000 active physical transfers whilst ensuring all available bandwidth is used
- Must be able to handle transfers which have to wait for several hours before the physical file is ready to transfer

Possible solution to URL options problem

There has been a long standing problem with the format and syntax used to express URLs and associated metadata and options. While not directly related to data staging, it will be addressed as part of the data staging work.

Architecture Proposal

- **3 layers:**
 1. Higher layer processing data staging requirements of jobs, collecting requested attributes/properties, resolving priorities (flattening them), managing credentials.
 2. Middle layer schedules individual transfers, suspends and resumes them, distributes bandwidth, etc.
 3. Lower level handles individual transfer protocol, communicates with middle layer to acquire, release and pre-allocate resources (mostly bandwidth), caches connections (if possible).
- Any layer can be outsourced to external service, for example gLite FTS
- **Basic file staging stages:**
 1. Identify transfer options
 2. Check cache
 3. Evaluate authorization (may require bandwidth)
 4. Resolve location (meta-URL, if needed, may require bandwidth)
 5. If needed repeat all steps from beginning
 6. Do transfer
 7. Post-processing (eg register replicas, release cache locks)

DTR - Data Transfer Request

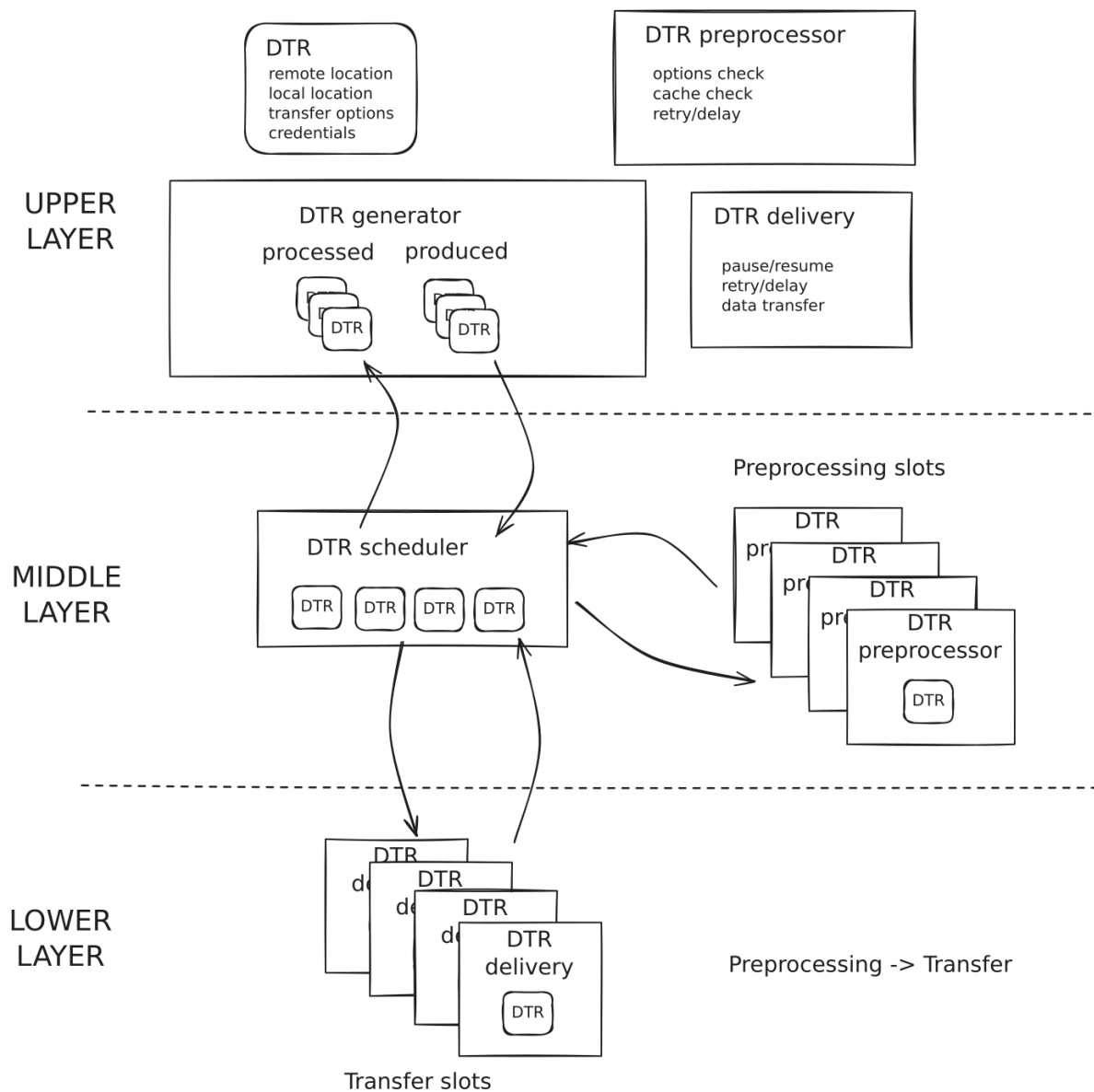


Fig. 7.2: Functional components and layers.

Requirements for components interfaces (based on protocol descriptions and architecture)

DTR Description

DTR stands for Data Transfer Request. This is the structure that contains several fields that fully describe the file transfer to be performed. One DTR is generated by the generator per each file transfer.

A detailed description and state diagrams of DTRs can be found in the [Detailed description of DTRs](#)

DTR Generator to DTR Scheduler

- Push DTR from Generator to Scheduler:
 - DTR contains the description of the single entity to be transferred - uploaded or downloaded. That includes endpoint, transfer parameters and limits, credentials - possibly multiple, etc.
 - Multiple DTRs may be affiliated together. Possible reasons and uses:
 - * Belong to the same job
 - * Belong to a bunch of jobs which user indicated as preferably processed together
 - * Belong to the same VO and assigned priorities to be applied within group
 - * Failure of one DTR in a group may cancel processing of other DTRs (not sure, may be implemented in Generator)
 - DTR may have assigned priorities levels. Probably related to groups.
- Receive DTR from Scheduler to Generator:
 - Returned DTR indicates outcome of processing, either positive or negative. In the latter case it includes description of encountered problems and level of severity.
- Cancel DTR in Scheduler
- Modify DTR properties in scheduler. Possible usage: - Manipulate priorities

DTR Scheduler to DTR Preprocessor

- Push DTR from Scheduler to Preprocessor
 - Because DTR preprocessing is supposed to take short time it may include processing timeout
- Receive DTR from Preprocessor to Scheduler
 - Returned DTR indicates outcome of processing
 - * Positive
 - DTR comes with information need for further processing either in Preprocessor or Delivery unit
 - DTR may contain multiple/alternative additional endpoints
 - Probably such DTR may be presented in a tree-like diagram
 - * Failure - includes description of encountered problems and level of severity
 - * Delayed processing
 - Includes retry time and possible margins.
 - Scheduler must ensure this DTR will go back to Preprocessor within specified time margins.
- Cancel DTR in Preprocessor

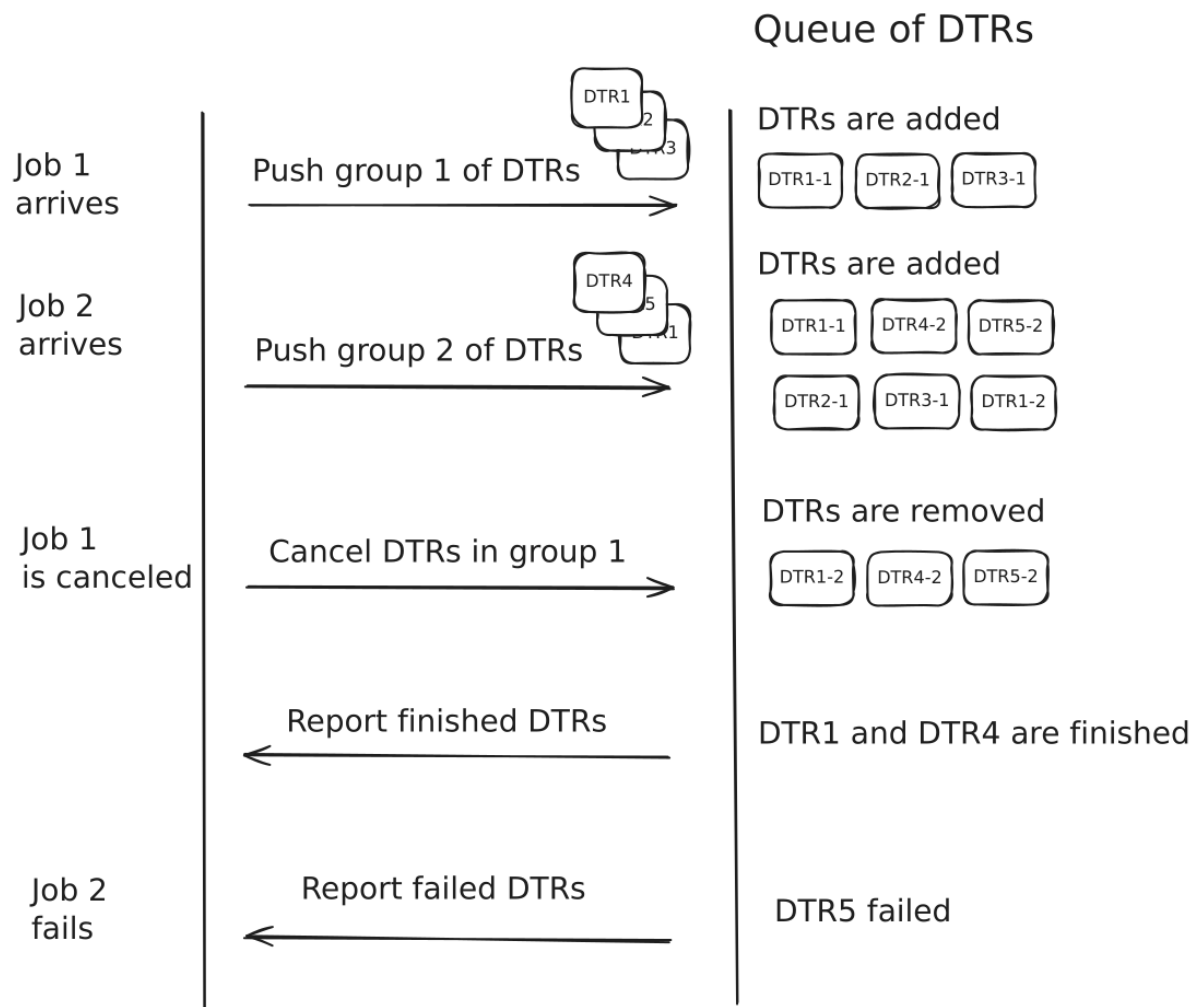


Fig. 7.3: Example of communication session between DTR Generator and DTR Scheduler.

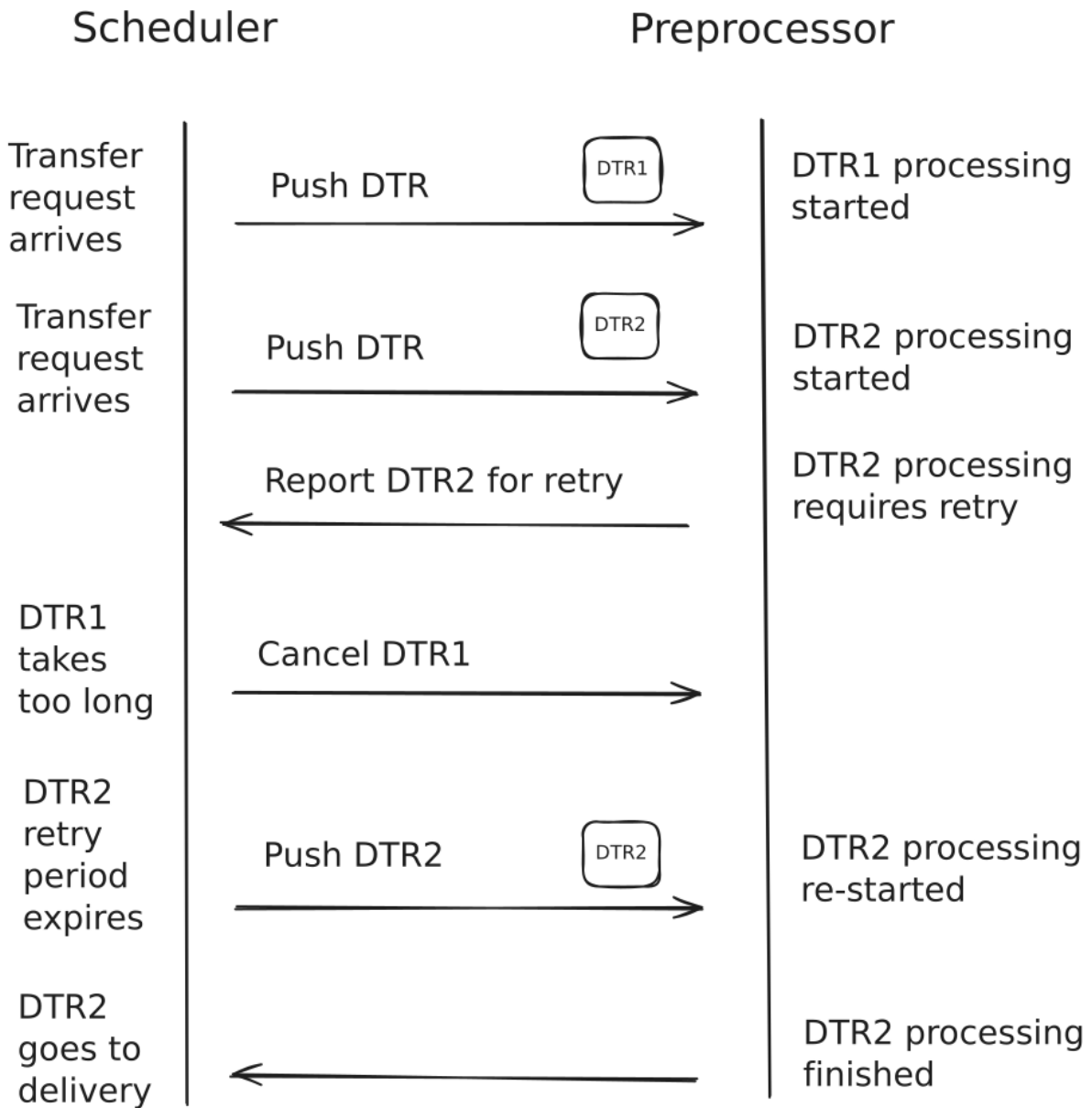


Fig. 7.4: Example of communication session between DTR Scheduler and DTR preprocessor.

DTR scheduler to DTR delivery

- Push DTR from Scheduler to Delivery
 - DTR may have bandwidth assigned
 - DTR may have timeout related parameters assigned - minimal transfer rate, maximal inactivity timeout, etc.
- Cancel DTR in Delivery
- Suspend DTR in Delivery - should DTR leave Delivery or should it stay there?
- Receive DTR from Delivery to Scheduler. Returned DTR indicates outcome of processing
 - Positive
 - Partially positive (partial data delivered)
 - Redirection
 - Failure - includes description of encountered problems and level of severity
- Get information about bandwidth currently used by Delivery
- Modify assigned bandwidth
 - May be used to free some bandwidth for urgent transfers

Component Workflows

Generator

The Generator is an integral part of the a-rex process. Internally it performs 4 tasks in the following order:

- Makes DTRs out of existing job descriptions also assigning priorities and grouping information and makes them available to the Scheduler.
- Communicates immediate requests like DTR cancel or suspend to the Scheduler.
- Monitors DTR states as reported by the Scheduler (and possibly by other modules) in order to provide feedback to client tools of A-REX asking for job state.
- Receives finished/failed DTRs from the Scheduler and initiates job processing continuation. Note: This part may be merged with previous one.

Scheduler

“Queues” are queues of DTRs waiting to enter pre/post-processing or delivery. They are kept internal to the scheduler. The scheduler is the only place with complete knowledge of the system and the only place where priorities are controlled and error conditions are handled. When an event is received and involves sending the DTR to a queue, the DTR is put at a certain position in the queue. Depending on the DTR priority, other DTRs in the queue may be moved, paused or stopped to allow higher priority DTRs to proceed quicker or consume more resources. On receiving an event, the relevant queue is examined and action taken (eg if a delivery finished start a new one).

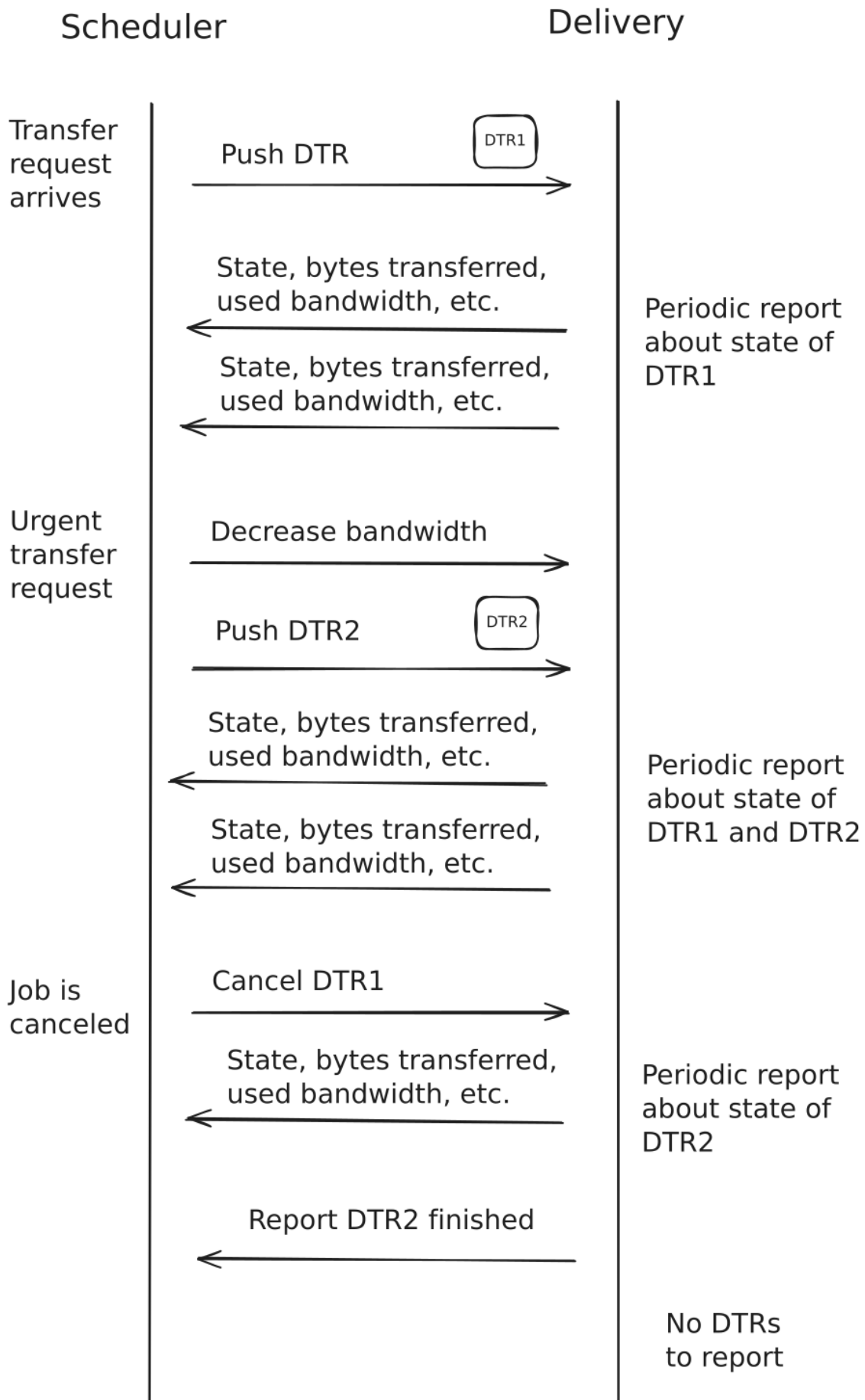


Fig. 7.5: Example of communication session between DTR Scheduler and DTR Delivery.

Reactions to new events

For simplicity error conditions are not included in the workflow here but described separately on the *Detailed description of DTRs*. They are examined by the scheduler, which will decide the workflow - do any necessary post-processing, decide whether to retry (immediately or after some delay) or to report back to the generator that the DTR definitively failed. DTR state transitions are also described in more detail and with diagrams on the *Detailed description of DTRs*.

- New DTR from generator
 - if cacheable or meta-protocol:
 - * add to pre-processor queue (for cache check, replica resolution etc)
 - else if ready to be delivered (base protocol for src and dest):
 - * add to delivery queue
- DTR returned from pre-processor
 - if cached:
 - * send to post-processor to be linked
 - else if need more pre-processing:
 - * add to pre-processing queue
 - else:
 - * add to delivery queue
- DTR returned from delivery
 - if post-processing required (index registration, cache linking, release request):
 - * put in post-processor queue
 - else:
 - * return to generator
- DTR returned from post-processor:
 - Return to generator
- DTR cancel notification from generator
 - if before pre-processing:
 - * return to generator
 - else if in post-processing:
 - * wait until finished and then send back to clean up
 - else:
 - * cancel immediately and add to post-processing queue for clean up
- DTR modify notification from generator
 - change request immediately, and modify queue if appropriate

Processor

The processor is divided into two logical parts: the pre-processor and the post-processor. Either part is invoked by the scheduler as a process/thread and has the DTR to process. Therefore, the pre- or post-processor can be a straightforward function, performing the next steps:

The pre-processor:

- check endpoint for its presence in cache
 - if successful, mark DTR as BYPASS_TRANSFER, return the DTR to the scheduler
 - if file is not in cache, construct cache endpoint for DTR destination and return to scheduler
- resolve the replicas of the file, if needed
 - return the DTR with a failure if no locations have been found
 - return the list of replicas found to scheduler
 - Note: The pre-processor doesn't care if the resolved locations represent meta-protocols themselves, it's the scheduler's job to determine it and possibly send this DTR for the pre-processing once again.
- query an endpoint
 - supply information on size, checksum, creation date, access latency as available according to protocol
- If an asynchronous request needs to be performed, for example SRM prepareToGet
 - Start request, mark DTR as STAGING_PREPARING_WAIT, return to the scheduler. An estimated wait time may be set by the pre-processor from information supplied from the remote service.
- If a polling request arrives from the scheduler
 - Check state of asynchronous request and report back to scheduler success, wait more, or error.
- If the cancellation request arrives from the scheduler
 - interrupt the operation, mark the DTR as PREPROCESSING_CANCELLED, return to the scheduler
- If the preprocessor hits the timeout during performing these tasks
 - interrupt the operation, mark the DTR as PREPROCESSING_NOT_FINISHED, return to the scheduler.

Possible features of the pre-processor in the future

- for each of resolved locations:
 - check options one by one (we still have to define transfer options, below is the example), mark in the list to which extent it satisfies the option
 - * can the location provide a required bandwidth
 - * other options, specified for the user
 - * keep processing all the options even if some of them are not satisfied – the scheduler may later review the options and start the processing again
 - request the file size (if possible) and compute estimated transfer time, mark it in the list for this location
- if there are in the list locations that satisfied all the checks
 - return the DTR as TRANSFER_READY to the scheduler with the list of these locations and their estimated transfer times
- if there are no locations in the list that satisfied all the checks
 - return the DTR with the least severe encountered failure (from most severe to least severe, Location is dead - Location is not authorized - Location doesn't satisfy the options), so the scheduler can either drop the DTR (in case of dead/non-authorized storages) or review transfer options and try again (in case of unsatisfied options)

The Post-processor:

- release stage requests;
- register a replica in an index service;
- release cache locks;

Delivery

Delivery is a component where the transfer orders are handled. Delivery process is listening to the Scheduler to receive the ready DTRs for transferring or any other events such as decreasing the bandwidth or suspending and canceling the ongoing jobs. Delivery process, reports the status of the DTRs and events to the scheduler periodically.

- Transfer request delivery
 - Pick up a DTR from delivery queue.
 - Check the source and destination, set the bandwidth, timeout and etc.
 - Start transferring by placing the received data in a buffer.
 - Periodically report the status of DTR such as transfered bytes, used bandwidth to Scheduler.
 - By finishing the data transfer do a checksum
 - * if checksum is correct return a SUCCESS status to the Scheduler
 - * else retry transfer again
- High priority transfer request event
 - Change the DTR to use max bandwidth n
- Job suspension
 - Suspend a transfer request
 - Keep the transfered data information in the delivery process
 - Release the used bandwidth
 - Report the status of suspended DTR to Scheduler
- Transferring cancel
 - Stop transferring
 - Release the resources if any is in use
 - Clean up the buffer
- Status report to Scheduler
 - Periodically collect the information of the DTRs in progress.
 - Calculate used bandwidth, transfered data bytes, status
 - Report the status to the Scheduler

Protocol Interfaces

Current Interface

The following interface is defined for each protocol through the DataPoint class:

- StartReading - Start reading data from the source
- StopReading - Stop reading data from the source
- StartWriting - Start writing data to destination
- StopWriting - Stop writing data to destination
- Resolve - Find physical replicas from an indexing service
- Check - Verify that the current credentials have access
- ListFiles - Find metadata
- Remove - Delete
- (Pre)Register and (Pre)Unregister - add and delete in indexing service

These do not fulfil all the requirements of the DTR interfaces described above.

New Interface

The main limitation of the current interface is that it does not handle efficiently those protocols such as SRM which involve asynchronous preparation steps. In the new framework, having to wait for a file to be prepared should not block other activities. Therefore the proposal is to split protocols as now into meta and direct protocols, but introduce a third class of stageable protocols and appropriate extra methods to handle them. A stageable protocol could also be a meta protocol or a direct protocol. Extra methods are also needed to handle pausing and cancellation of transfers.

For Meta Protocols (eg LFC, RLS)

- Resolve - Resolve replicas in an indexing service
- (Pre)Register and (Pre)Unregister - add and delete in indexing service

For Stageable Protocols (eg SRM, Chelonia)

- PrepareReading - Prepare the storage service for reading. This may involve preparation of Transport URLs (TURLs) which should then be used for physical data transfer, or operations like reading data from tape to disk. If the protocol's implementation is asynchronous then this method may return a status that tells the caller to wait and poll later. The caller should call this method again after some period of time (the remote service may provide an estimate of the preparation time).
- PrepareWriting - Prepare the storage service for writing. Works similarly to PrepareReading.
- FinishReading - Release or abort requests made during PrepareReading, usually called after physical transfer has completed.
- FinishWriting - Release or abort requests made during PrepareWriting, usually called after physical transfer has completed.
- Suspend - Pause current preparation of transfer
- Resume - Resume suspended preparation of transfer (Note: depending on the protocol, suspend and resume may be implemented as stop current request and start new request)
- Cancel - Stop preparation of transfer

For Direct Protocols (eg FTP, HTTP)

- StartReading - Start reading data from the source
- StopReading - Stop reading data from the source
- StartWriting - Start writing data to destination
- StopWriting - Stop writing data to destination
- Modify - Change parameters of transfer such as bandwidth limits
- Suspend - Pause current transfer
- Resume - Resume suspended transfer
- Cancel - Stop current transfer

For All Protocols

- Check - Verify that the current credentials have access
- List - Find metadata
- Remove - Delete

Existing ARC Code

We should aim to re-use as much as possible of the existing ARC code. The arc1 code base will be used for all developments. Code for job description handling and interfacing with the LRMS can remain unchanged, all that concerns us is the code handling the PREPARING and FINISHING states. Job state handling is done in states.cpp. This code can remain largely unchanged but in the ActJobPreparing/Finishing methods the job enters the new system at the upper layer.

The lower level code for data transfer for each protocol is handled in the DMCs. These can largely remain unchanged except for the extra methods in the meta-protocols above.

Caching code can remain unchanged.

PREPARING/FINISHING State Semantics

The semantics of these states may need to be changed - at the moment PREPARING/FINISHING means that the job is transferring data and PENDING those states means a job cannot enter the state due to some limit. In our new system there is less of a distinction between jobs waiting and jobs transferring data, also some files within the job may be transferring while some are waiting. Once the job enters the upper layer of the new system it will be in a staging state even though it may have to wait a long time before any data is transferred.

Processes and Threads

In the current architecture a persistent A-REX thread is spawned by HED. A new data staging process (downloader/uploader) is forked for each job to stage its input or output files. Within each process a thread is created per file. The status of each thread is reported back to the loader process through callbacks and the process exit code tells the A-REX the status of all transfers. There are some problems with this approach:

- The A-REX knows nothing about what happens inside the loader until the process exits and cannot communicate with the loader
- It is not possible to change the uid of each thread individually so they all run under the uid of the downloader/uploader. The uid of the downloader/uploader depends on configuration. If possible uid of mapped user is used. But if there is cache shared among users that uid will be root or uid of A-REX.

To solve this last problem, processes in the delivery layer writing to or from the session dir must run under the locally mapped uid. Writing to the cache must always be done as root and so cache and non-cache downloads must be done in different processes. This leads to the conclusion that the delivery layer must be separate processes from the scheduling layer. Then to solve the first problem there needs to be a method of communication between the delivery and scheduling layers, which must be two way, so that the scheduler can modify on-going transfers and the transfer can report its status back to the scheduler.

Notes on running processes under mapping uid:

- It is only open() operation which needs to be run under special uid. If we follow convention that all open() operations in an executable are called through single wrapper function and put global lock around it, then we can have filesystem access under selected uid inside multi-threaded application. Unfortunately for NFS the process needs to maintain the uid throughout the whole transfer.
- This may introduce performance issue if open() operation takes too long, like in case of remote file system.
- open() may be called by an external library which we have no control over
- according to ‘man open’ “UID mapping is performed by the server upon read and write requests” and hence suggested approach will fail on NFS.

Operations carried out by the pre- and post-processor require access to the filesystem in the following steps

- Cache preparation - locking the cache file and checking its existence. This must be done as root.
- Cache finalisation - copying or linking the cached file to the session directory. The hard link to the per-job dir must be done as root but the soft linking or copying to the session dir must be done as the mapped user.
 - Last is rather “must” than “may” to make it work on NFS with root_squash=on. The situation may exist where cache is on NFS and then process copying file must switch uid while accessing file and its copy.
- Access to proxy - the proxy is needed when contacting secure remote services. The proxy in the control dir is owned by the mapped user. Therefore either we have to:
 - Make a copy of the proxy owned by root - this does not fit our security requirements above
 - * Note1: This is how it is done in current implementation. There is no security hole here because this is same proxy which was obtained by A-REX and written to control directory under root uid. So this proxy already belonged to root and making belong to root again makes little difference.
 - * Note2: Because as a rule proxy is stored on local file system it is always accessible by root. Copying of proxy was needed in current implementation due to limitation of Globus libraries - those were able to accept proxy only from file and were checking if proxy belonged to current uid. Because ARC data library allows assigning credentials to communication directly (at least it should, but may be not implemented for some protocols, needs checking) and because proxy is readable by root such trick is not needed anymore.
 - Use processes rather than threads for the pre- and post-processor, changing the uid of the process to the mapped user
 - * This approach is most probably not needed for this purpose but may be very much desirable for fighting NFS.
 - As suggested above, use global open() function to open with mapped id then switch back to root - this suffers the same problems mentioned above
 - * There is no need to do that for proxy because root can always open files belonging to other users

The generator will be a thin layer between the main A-REX process which moves jobs from state to state and the scheduler which processes DTRs - when a job enters the PREPARING or FINISHING state the generator will create DTRs and send them to the scheduler. For the scheduler to be efficient it should run as a separate process (or thread). Each pre-processing step such as resolving replicas in a catalog or a cache permission check should be fast (< few seconds) but should be run asynchronously in order not to block the scheduler if something bad happens.

In summary we have the following processes corresponding to the three layers:

- Persistent main A-REX process (thread of HED) and persistent generator thread

- Persistent scheduler process (or thread) with temporary threads for pre-processing
- Temporary delivery processes created per file transfer

Using processes vs threads should add no more CPU load, however more memory is required. For increasing performance it should be possible to reuse processes in a way similar to how threads are reused.

Implementation Choices

The current system uses files in the control directory for all communication. We may need something more sophisticated for this more complex system, either internal to the new system or also to replace the current files in control dir method. Possibilities:

- Files
- Sockets
- Database
- RPC
- Libevent - <http://en.wikipedia.org/wiki/Libevent>
- Pipes
- Message passing
 - DBus - <http://www.freedesktop.org/wiki/Software/dbus>
 - ActiveMQ - <http://activemq.apache.org/>
- ...

Suggestion 1

- Use persistent object - file, database record - for storing DTRs. Each object includes:
 - Description of DTR
 - Owner of DTR
 - Last requested action
 - Current state inside owner (or maybe inside every module).
- Each DTR has own ID - may be just file name or record key
- Keep simple communication channels between modules - like 2 pipes for example
 - Whenever DTR is changed its ID is sent over communication channel to module which is supposed to react
 - As backup modules can scan modification timestamps of objects periodically

Pros:

- Simple
- Persistency

Cons:

- Monitoring of ever changing state - like bandwidth usage - would require constant modification of files/records.
 - This problem could be solved by providing information which needs no persistency through communication channel. But that would make communication more complex.
 - Another possible solution is to mmap (should work well with files) objects and flush them to persistent store only if persistent part of information is modified.

Suggestion 2 (used in current implementation)

- DTR objects can be rather complicated, so keep them only in memory
- In case of process failure all DTRs are reconstructed from control files into initial NEW state
- DTRs are passed as objects between threads
 - Separate threads run for A-REX (including the Generator), Scheduler, and Delivery components
 - Delivery thread starts new processes for each transfer
 - * These communicate simple status messages (“heartbeats”) through pipes to main delivery thread
- Communication between threads through callbacks

Pros:

- Simplifies development - no need for complex persistency layer or serialisation code
- Fast communication through thread callbacks

Cons:

- Having lots of threads increases risks of deadlocks and race conditions
- No persistency
- No way to communicate from scheduler to transfer processes
 - If the only communication required is to pause/resume/cancel a transfer it can be done through signals eg SIGSTOP/SIGCONT/SIGTERM

Implementation idea for Suggestion 2

- Generator, Scheduler, Delivery processes are singletons and run as persistent threads. The Processor is a singleton.
 - Note: C++-wise it would be probably more correct to use static methods instead of singletons. With such approach there would be no need to handle singletons outside of class itself hence simplifying interface. Although inside class there may be singleton.
- A-REX initiates a request to the data staging system by calling `Generator::receiveJob()`
 - Generator creates DTRs and after they finish it modifies the job state directly
 - A-REX can query job status from Scheduler via Generator
- Generator communicates with scheduler by calling `Scheduler::receive_dtr` method
 - There is also a `Scheduler::cancel_dtrs` method which is called by the Generator to cancel DTRs
- Scheduler has an instance of `DTRLIST` class, which is the storage for all the DTRs in the system.
- Scheduler communicates (passes a DTR) to any other process by calling `DTR::push()` method
- Pre-, post-processor and delivery communicate with scheduler by simply changing the status of DTR and changing the owner to “SCHEDULER”. Scheduler will pick up DTRs with changed statuses automatically during next iteration.
- When DTR goes from the scheduler to the pre or post-processor, it calls the `processDTR()` method of the singleton processor within `push()`. The processor then spawns a thread and returns.
- When DTR goes from the scheduler to delivery, within `push()` it calls the `processDTR()` method of the singleton delivery. The delivery then spawns a process and returns.
 - The delivery singleton receives messages through pipes from the transfer processes and reports information in the DTR object.

Detailed description of DTRs

DTR stands for Data Transfer Request. This is the structure that contains several fields that fully describe the file transfer to be performed. One DTR is generated by the generator per each file transfer.

Fields of the DTR

More or less required:

- DTR ID
- source endpoint
- destination endpoint
 - for source and destination, a list of metadata such as file size, checksum, creation date etc
 - for source and destination (if applicable) a list of replicas
 - for source and destination (if applicable) current replica
 - for source and destination (if applicable) TURL or delivery-level URL used for transfer
 - for source and destination (if applicable) request ID (in the case of asynchronous requests to remote storage services)
- credentials
- cache information
 - if the file is cacheable, the filename in cache
 - cache directories configuration
 - caching state (already in cache, cache currently locked etc)
- local user information (uid/gid)
- Job ID this transfer belongs to
- priority of the transfer - a number set by the generator which flattens priorities
- transfer share this DTR belongs to
- sub-share the DTR belongs to - may be set by the Generator
- tries left
- flags to handle properties and strategies when dealing with index servers
 - flag to say whether DTR is replicating inside the same logical filename
 - flag to say whether DTR should force registration to an existing logical filename, if the source is different
- mapping info - mapping information of local files to which remote files may be mapped to in the configuration (copyurl/linkurl)
- status of the DTR
- error status
 - type of error
 - location of error
 - text description of error detail
- number of bytes transferred/offset
- timing properties

- timeout - time which DTR is allowed to remain in current state
- creation time
- last modification time
- process time - wait until this time to do further processing
- cancel (set to true if request is to be cancelled)
- bulk operation flags to combine several DTRs in a bulk request
- delivery endpoint, whether Delivery is to be carried out by a local process or remote service
- current owner - who is in charge for this DTR right now
- logger object, so each DTR can have its own log
- lock, since DTRs can be modified by several processes, for avoiding writing collisions

Possible

- affiliation (if we use the affiliation of multiple DTRs, see right below).
- history of states

Multiple DTRs may be affiliated together. Possible reasons and uses:

- Belong to same job
- Belong to bunch of jobs which user indicated as preferably processed together
- Belong to same VO and assigned priorities to be applied within group
- Failure of one DTR in group may cancel processing of other DTRs (should be implemented in Generator)

State transitions of DTR

All possible states of a DTR, with arrows indicating the normal flow of DTRs between states. Each state is explained in detail below. Error conditions are not included here but are shown in another diagram further down.

Status codes

The following table describes all non-error status codes, and also the action taken in the event of a cancellation request being received while in that state. In general if all of the data transfer has been completed before receiving a cancellation request, the destination file is not deleted. The main reason for this is to preserve cache files, as the user may wish to run the same job soon after cancelling it.

Table 7.

Status Code	Text Description
Statuses set by the generator	
NEW	The DTR has just been built by the generator
CANCEL	A request has been made to cancel the DTR
Statuses set by the scheduler	
CHECK_CACHE	The DTR destination is cacheable and the cache should be checked for the file's existence
RESOLVE	The DTR source is a meta-protocol and should be resolved
QUERY_REPLICA	The DTR source should be queried to check existence, check file size, checksum etc.
PRE_CLEAN	The destination in the DTR should be deleted before writing
STAGE_PREPARE_SOURCE	The DTR source is a meta-protocol which must be prepared or staged
STAGE_PREPARE_DESTINATION	The DTR destination is a meta-protocol which must be prepared or staged
TRANSFER_WAIT	The DTR is ready to be sent to delivery but must wait due to transfer limits or priority
TRANSFER	The DTR should be transferred immediately
RELEASE_REQUEST	The DTR transfer has finished and any requests made on remote storage should be released

Table 7.1 – c

Status Code	Text Description
REGISTER_REPLICA	The DTR destination is a meta-protocol and the new replica should be registered
PROCESS_CACHE	The DTR destination is cacheable and the cached file should be unlocked and linked/copied
DONE	The DTR completed successfully
CANCELLED	The DTR has been cancelled successfully
ERROR	An error occurred with the DTR
Statuses set by the pre-processor	
CHECKING_CACHE	The pre-processor is checking the cache
CACHE_WAIT	The cache file is locked and the scheduler should wait before trying to obtain the lock
CACHE_CHECKED	The cache check is complete
RESOLVING	The pre-processor is resolving replicas
RESOLVED	The replica resolution is complete
QUERYING_REPLICA	The pre-processor is querying a replica
REPLICA_QUERIED	The replica querying is complete
PRE_CLEANNING	The pre-processor is deleting the destination file
PRE_CLEANNED	The destination file has been deleted
STAGING_PREPARING	The pre-processor is making a staging or preparing request
STAGING_PREPARING_WAIT	The staging or preparing request is not ready and the scheduler should wait before polling
STAGED_PREPARED	The staging or preparing request is complete
Statuses set by the delivery	
TRANSFERRING	The transfer of the DTR is on-going
TRANSFERRED	The transfer completed successfully
Statuses set by the post-processor	
RELEASING_REQUEST	The post-processor is releasing a stage or prepare request
REQUEST_RELEASED	The release of stage or prepare request is complete
REGISTERING_REPLICA	The post-processor is registering a replica in an index service
REPLICA_REGISTERED	Replica registration is complete
PROCESSING_CACHE	The post-processor is releasing locks and copying/linking the cached file to the session d
CACHE_PROCESSED	Cache processing is complete

Error Conditions of DTRs

The following diagram shows possible error conditions and actions taken. For simplicity and because all error handling logic takes place within the scheduler, the pre- and post-processor and the delivery layers are not shown.

Errors are categorised into the following types:

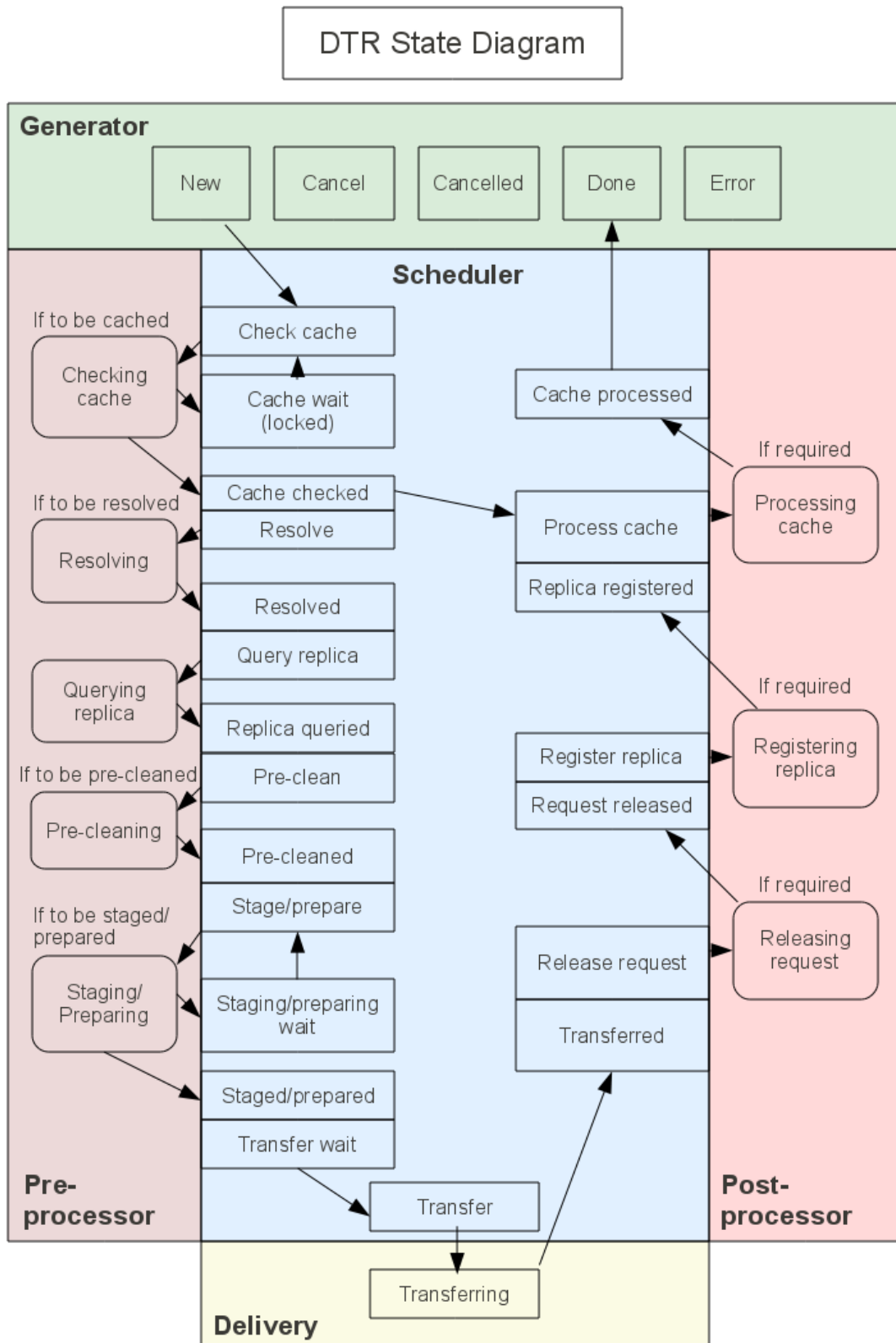


Fig. 7.6: DTR State Diagram

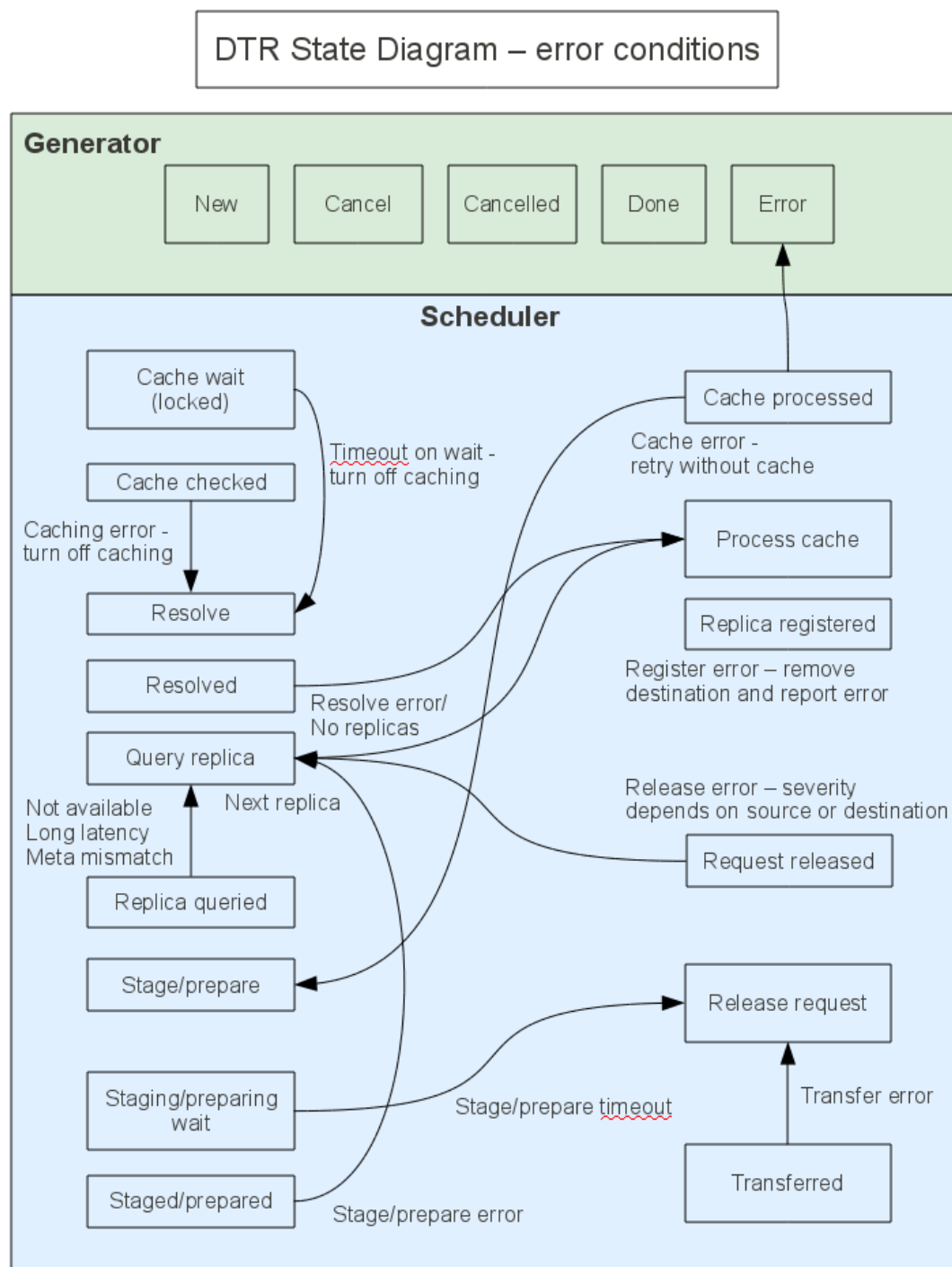


Fig. 7.7: DTR Error State Diagram

Error	Explanation	Retry	Action
INTERNAL_LOG_ERROR	Internal error in data staging logic	No	Stop processing and report back to generator
INTERNAL_PROS_ERROR	Internal error like losing contact with an external process	Yes	Clean if necessary and retry
SELF_REFERENCE_ERROR	Attempt to copy a file to itself	No	Return to generator
CACHE_ERROR	A problem occurred in cache handling	Yes	Retry without caching
TEMPORARY_REMOTE_ERROR	Error such as connection timeout on remote service	Yes	Retry with an increasing back-off
PERMANENT_REMOTE_ERROR	Error such as file not existing, permission denied etc on remote service	No	Follow cancellation steps and return failed DTR to generator
LOCAL_FILE_ERROR	Error with a local file	No	Follow cancellation steps and return to generator
TRANSFER_SPEED_ERROR	Transfer rate was below specified limits	Yes	Retry transfer. If all retries fail, report back to generator - it will make the decision on whether to cancel other related DTRs. (Future work: make decision on whether other transfers caused slow transfer and whether cancelling others would help or should be done)
STAGING_TIMEOUT_ERROR	The staging process took too long	No	Try a different replica - if none available, cancel and report back to generator

Methods of DTRs

DTR::push (DTR, receiver) – pass the DTR from one process to another, e.g. *DTR::push (dtr, preprocessor)*

Implementation

Within Data Staging framework there is a global list of DTRs. Pointers to the DTRs are passed around between components, which can modify them directly and push them between each other.

DTR priority and shares system

Here we describe the priority and shares system for the new data staging framework. During the design stage there were several ideas taken from other research in the field and the first implementation of the transfer shares model in ARC.

Ideas behind priorities and fair-share in data staging

The initial idea was giving every DTR that comes into the system a fixed priority, then sorting the queue according to priorities and launching the first N DTRs with the highest priorities. This scheme also allows easy incorporation of pre-emption: if the job with higher priority appears it just pushes other DTRs out of the front of the queue and then during the next scheduler loop we can start these DTRs and suspend the pushed ones.

However, this idea can potentially lead to the situation that demanded the implementation of transfer shares in ARC. If a user or VO with the highest priority submits bunch of jobs at once all the other will be blocked, because DTRs from this bunch will occupy the front of the queue for a long time.

The idea of transfer shares comes in handy now. The available transfer slots should be shared among different VOs/Users, so nobody would be blocked. VOs/Users with higher priority get more transfer slots than the other. However, strict limits on the number of slots per share are not flexible enough - if the transfer pattern changes then strict limits could cause problems, squeezing lots of users/jobs into one share with a few slots and blocking others. The User or VO must also be able to decide the relative priority of jobs within its own share.

Current Implementation

The ideas above led to the creation of two configurable properties: user-defined job priority and server-defined share priority. Users may define a priority for their jobs in the job description (“priority” attribute in xrsl), and this is a measure of the priority given to this job within the share it gets assigned to when submitted. On the server-side, it is possible to define a share type, and priorities of certain shares. The share priority is used to determine the number of transfer slots to assign to the share, taking into account which shares are currently active (active meaning the share has at least one DTR in the system).

When the Scheduler receives a new DTR, it is placed into a transfer share, which is defined by a User DN, VO, Group inside VO or role inside VO as it was in previous versions of ARC. Currently it’s possible to use only one sharing criteria in the configuration, i.e. it’s not possible to use simultaneously sharing by User and VO.

Priority is defined as a number between 1 and 100 inclusive - a higher number is a higher priority. In the A-REX configuration it is possible to specify a base priority for certain shares. If the DTR doesn’t belong to any of these specified shares, it is placed in a “_default” share with default base priority (50). The scheduler sets the priority of the DTR to the base priority of the share multiplied by the user priority from the job description (default 50) divided by 100, therefore default priority of a DTR is 25. In this system the priority set in the job description effectively defines a percentage of the base priority. Thus service administrators can set maximum priority limits for certain shares, but users or VOs have full control of their jobs’ priority within the share.

While revising the Delivery and Processor queues, the scheduler separates DTRs according to the shares they belong to. Inside every share DTRs are sorted according to their priorities. Then the scheduler determines the number of transfer slots that every active share can grab. The number is determined dynamically depending on priorities of active shares. Each share receives the number of slots which corresponds to the weight of its priority in the summed priority of all active shares. After the number of slots for each share is determined the scheduler just launches N[i] highest priority DTRs in each share, where N[i] is the number of transfer slots for i-th share.

The reason for weighting the DTR priority by the share priority is for occasions when the Scheduler considers the entire queue of DTRs, for example when allowing highest priority DTRs to pass certain limits.

Example: there are two active shares, one has base priority 60, the other 40. The summarized priority is 100 (60 + 40). The first share has a weight of 60%, the second 40%. So the first will grab 60% of configured transfer slots, and the second – 40%. If the system is configured with 5 Delivery slots, then the first share will take 3 slots and the second 2 slots. The 3 highest priority DTRs from the first share and 2 highest priority from the second share will be assigned to those slots.

Emergency Shares

To avoid the situation where a fixed limit of slots are used up by slow transfers and a new high priority transfer has to wait for a slot, we have “emergency” transfer slots. If there are transfers in the queue from a particular share, but all slots are filled with transfers from other shares, one emergency slot can be assigned to this share to allow transfers to start immediately. The share may use an emergency slot until any other transfer finishes, at which point the emergency slot becomes a regular slot and a new transfer does not start from the queue.

Sub-shares

The Generator can assign DTRs to “sub-shares” to give a higher granularity than the standard criteria and when assigning transfer slots. Sub-shares are treated as separate shares. In A-REX, different sub-shares are assigned to downloads and uploads, and in this case emergency transfer slots prove useful for preventing jobs not being able to finish because all transfer slots are taken by downloaders. If this happens then emergency slots can be used for uploads.

Potential Problems

- Within a share, high priority jobs block low priority jobs. Thus if there is a constant stream of high priority jobs in a share, then some low priority jobs in the same share may never run. Possible solutions:
 - Increasing the priority as the time spent in the queue increases (returning to previous priority after leaving the queue). This is currently implemented as increasing the priority by 1 every 5 minutes after the DTR’s timeout has passed.
 - Changing simple highest-priority-first to a random algorithm where higher priorities are weighted higher
 - Making a higher granularity of shares by splitting each priority or priority range into its own share - this is probably too complicated

A-REX Configuration

The configuration varies depending on the ARC version. In the examples below VO roles are used to assign shares, the atlas slow_prod role is assigned a low priority share and the atlas validation role is assigned a higher priority share.

```
[arex/data-staging]
sharepolicy=voms:role
sharepriority=atlas:slow-prod 20
sharepriority=atlas:validation 80
```

If both shares are active and there are 10 slots available, then DTRs in the slow-prod share will get 2 slots and those in the validation share get 8 slots, and so the jobs in the validation share will have a higher throughput (assuming similar numbers of files and file sizes in each type of job).

Example

A user wants their job to be high (but not top) priority and specifies (“priority” = “80”) in the job description. The user has a VOMS proxy with no role defined and submits the job to a site with the above configuration. The job is assigned to the default share and DTRs have priority 40 (50 x 80 / 100). The user then creates a VOMS proxy with the ATLAS validation role and submits another job with the same priority to the same site. This time the job goes to the configured atlas:validation share and the DTRs have priority 64 (80 x 80 / 100). Note that the priority of a DTR only affects the its position within a share and does not influence the distribution of slots between shares.

Dynamically modifying priorities

From version 6.2 ARC supports dynamically adjusting priorities of active DTRs for example to push certain transfers to the front of the queue.

This can be done by creating a file in the same location as the dtr.state file (by default the control dir) with the name dtr.state.prio. This contains lines of DTR ID and priority, eg:

```
c3e12bfa-ba9a-4a45-be73-0b3b2a206e1c 90  
ff2d5329-5440-49aa-83cb-72daeb5e3e59 95
```

This file is read by the DTR system before it sorts the queue by priority and it changes any DTRs found in the file to the given priority. Then it renames the file to dtr.state.prio.read.

7.1.3 ARC Data Delivery Service Technical Description

This page describes how to set up DTR over multiple hosts.

Introduction

In DTR the transfers can be split over multiple hosts to scale up the data transfer throughput. Remote hosts only perform simple point-to-point transfers between physical endpoints and all the logic is kept in the main A-REX host. This has many advantages:

- Keeping all the high-level logic in one place allows intelligent load balancing.
- The remote hosts do not do submission to the batch system and so all LRMS configuration only needs to be done in one place.
- A problem with a remote host does not mean a lost job - the transfer can simply be retried on another host.
- The setup and configuration of a remote host is very simple.

The datadelivery service has been developed within HED and can be deployed on one or many remote hosts. Once a DTR has passed all the pre-processing steps and is ready for transfer, the Scheduler can send it to a remote datadelivery service to execute the transfer. The Scheduler will then poll the request until it has completed. The Scheduler spreads randomly the total number of transfers (up to the configured limit and depending on the configured access rights of each service) between all the staging hosts.

The architecture is shown in the figure below. The datadelivery service on the remote host acts as a very simple Scheduler which sends all transfer requests it receives straight to Delivery.

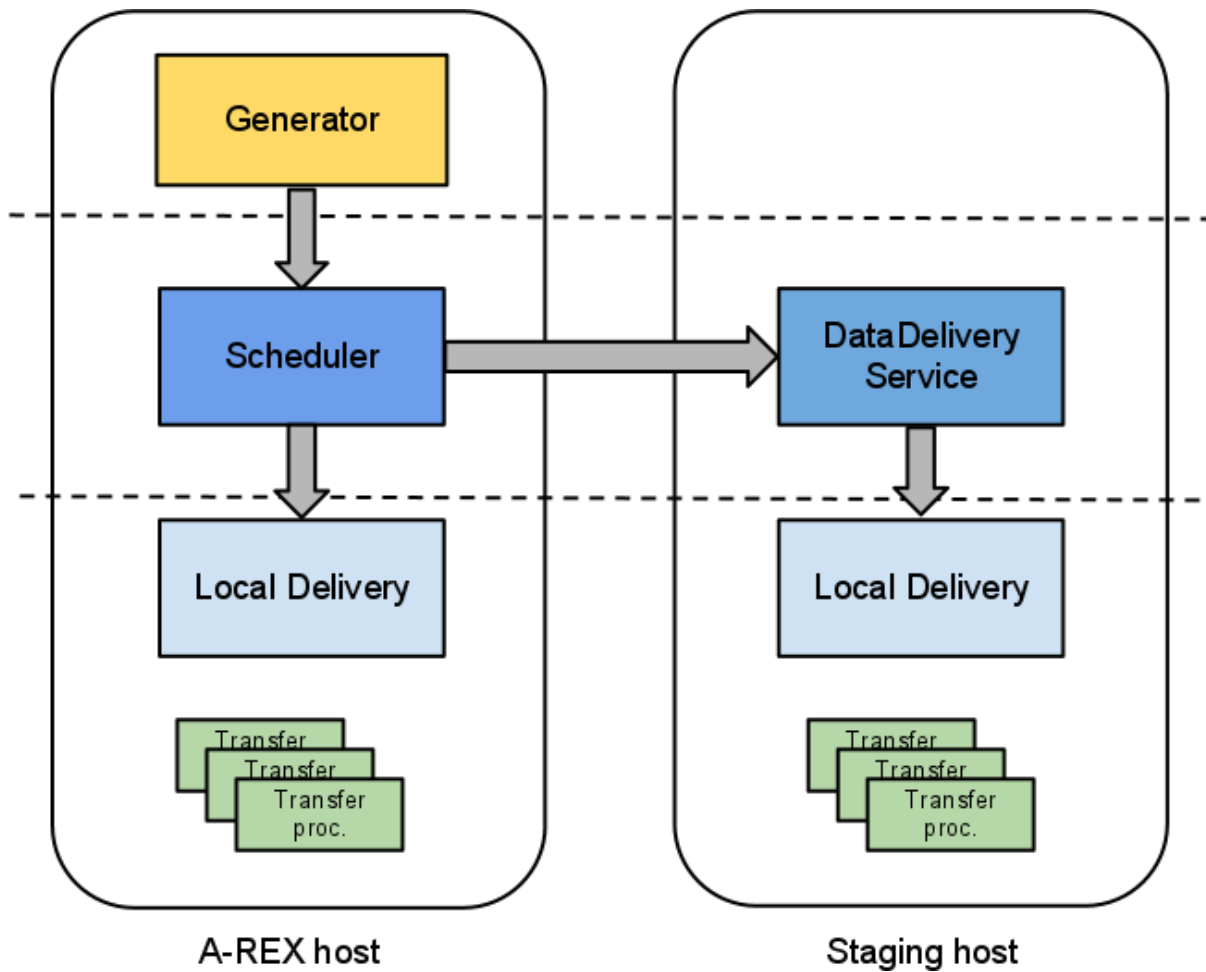


Fig. 7.8: Multi-host datastaging

Installation

No extra components are necessary on the A-REX host.

For the remote staging hosts, the package **nordugrid-arc-datadelivery-service** can be found in the usual EPEL or Nordugrid repositories, or is built automatically when building the source tree. This package should be installed on each remote host. It depends on some core ARC packages and the usual external packages that an ARC installation requires. Note that it is not necessary to install A-REX on the staging hosts.

IMPORTANT If support for transfer protocols that are not included in the **nordugrid-arc-plugins-base** package (e.g. GridFTP or Xrootd) is required then the corresponding plugins packages must also be installed.

CA certificates are required on each remote host to authenticate connections with storage services. Depending on the local set up, each host may also require a host certificate (more info below).

To start|stop the service (as superuser):

```
systemctl start|stop arc-datadelivery-service
```

The log file for the service can be found at `/var/log/arc/datadelivery-service.log` (configurable - see below). If logrotate is running, this log will be rotated every day.

Configuration

Remote Hosts

The datadelivery service uses the `[datadelivery-service]` section of the `arc.conf` configuration file.

By default the datadelivery service runs with TLS enabled and a host certificate is required for each host. The path to the host credentials may be specified by the usual x509 options. Host certificates are not a strict requirement and it is possible to run without TLS using the `secure=no` option. This means that A-REX can no longer verify the authenticity of the remote hosts and so the decision to run with or without host certificates should be based on local site policy.

The following configuration options are supported:

Parameter	Explanation	Default Value
<code>x509_host_k</code>	Path to host key	<code>/etc/grid-security/hostkey.pem</code>
<code>x509_host_c</code>	Path to host cert	<code>/etc/grid-security/hostcert.pem</code>
<code>x509_cert_dir</code>	Path to CA certificates	<code>/etc/grid-security/certificates</code>
<code>hostname</code>	Hostname of service host	<code>localhost</code>
<code>port</code>	Port on which service runs	<code>443</code>
<code>pidfile</code>	pid file	<code>/var/run/arched-datadelivery-service.pid</code>
<code>logfile</code>	Log file	<code>/var/log/arc/datadelivery-service.log</code>
<code>loglevel</code>	Logging level (0 (FATAL) to 5 (DEBUG))	<code>2 (WARNING)</code>
<code>user</code>	User under which service runs (should only be changed in special cases)	<code>root</code>
<code>secure</code>	Set to “no” if the service should run without a host certificate	<code>yes</code>
<code>allowed_ip</code>	IP address authorized to access service (can be specified multiple times)	No default, must be specified
<code>allowed_dn</code>	DN authorized to access service (can be specified multiple times)	No default
<code>transfer_dir</code>	Path the service is allowed to read/write to (can be specified multiple times)	No default, must be specified must be specified

At least one `allowed_ip` and at least one `transfer_dir` are the only mandatory parameters, but it can be useful to change interface, port and loglevel from the default.

Since the service can copy files to and from the service host, it is dangerous to allow open access to any clients. Usually `allowed_ip` is set to the IP address of the A-REX host since this is the only host which should have access to the service. The service can be further locked down by specifying authorized DNs so that only certain users are allowed to have their jobs' files staged by the service. Note that DN filtering is not possible if `secure=no` is specified.

It is also dangerous to allow requests which can copy to or from anywhere on the host filesystem, so filesystem access is restricted through the `transfer_dir` parameter, which specifies the path(s) that requests are allowed to use. The service is able to read and write to any path matching `transfer_dir`. `transfer_dir(s)` should be specified which match every cache and session dir. In some situations it may be desirable to set for example one cache per remote host where the cache is local to the host. A-REX checks on start-up which dirs are accessible by which remote hosts and uses that info to direct the DTRs to the right hosts.

Configuration example:

```
[datadelivery-service]
loglevel = 3
hostname = delivery.host.1
port = 60002
allowed_ip = 1.2.3.4
transfer_dir = /var/arc/session
transfer_dir = /var/arc/cache
```

If the site has multiple caches like `/var/arc/cache-01`, `/var/arc/cache-02` etc, then this configuration will allow the delivery service to write to all of them.

It is vital that the cache and session file systems are mounted on the same path on the A-REX and each remote host, as the system assumes that a transfer to the cache or session directory can be done using the same path on all hosts. It is also vital that user accounts are in sync across all hosts, so that a user represented by a uid on the A-REX host maps to the same account on all the remote hosts. No mapping of DN to local user is done on the remote hosts, and so no mapping infrastructure like gridmap files is required. The local user id of the user owning the transfer is passed in the request to the service and the transfer process is executed under that uid (except when writing to cache, when the root account is used).

The service is started and stopped by the `arc-datadelivery-service` systemd unit:

```
systemctl start arc-datadelivery-service
```

A-REX Host

The configuration for setting up A-REX to use remote datadelivery services uses options in the `[arex/data-staging]` section of the regular `arc.conf`.

Parameter	Explanation	Default Value
<code>delivery-service</code>	URL of remote host which can perform data delivery. Hostname and port must match those specified in the datadelivery service configuration.	None
<code>localdelivery</code>	Whether local delivery should also be done	no
<code>remote-sizelimit</code>	File size limit (in bytes) below which local transfer is always used	0
<code>use-hostcert</code>	Whether the host certificate should be used in communication with remote datadelivery services instead of the user's proxy	no

Example:

```
[arex/data-staging]
deliveryservice = https://delivery.host.1:60002/datadeliveryservice
deliveryservice = https://delivery.host.2:60002/datadeliveryservice
localdelivery = yes
remotesizelimit = 1000000
```

Multiple remote datadelivery services can be specified, and the Scheduler will randomly divide DTRs between those services where the transfer is allowed (according to transfer_dir configuration on each service). The presence of a datadeliveryservice option turns off the regular “local” delivery on the A-REX host so only the remote service will be used. If it is desired to also do transfers on the A-REX host, **localdelivery = yes** must be used. If no remote services are specified then local delivery is always enabled.

If the datadelivery service is running with secure=no then https should be replaced by http in the deliveryservice URLs.

Normally the credentials of the user who submitted the job are used for communication between A-REX and remote datadelivery services, however for extra security it is possible to use the A-REX host certificate for this instead, by specifying **usehostcert = yes**. In this case the host cert is only used for establishing the secure connection with the remote service, it is still the user’s credentials which are delegated and used for the transfer. The host cert must also be able to be used as a client certificate, in other words must have the X509 extension “X509v3 Extended Key Usage: TLS Web Client Authentication”. This option has no effect if the datadelivery service is running with secure=no.

Communication with remote services involves some degree of overhead such as the SSL handshake, delegating credentials etc, and when transferring small files this overhead can become a significant fraction of the transfer time. Therefore it is possible to specify a file size limit (in bytes) with **remotesizelimit**. Any files smaller than this limit will use local transfer, even if local transfer is disabled through **localdelivery = no**.

Deployment Scenarios

Several ways of deploying multi-host data staging are possible, using any number of remote hosts. Two examples are shown here.

Shared Storage

The cache and session directories are on a storage system such as Lustre or GPFS, which is mounted on all hosts. All hosts can access the cache and session directory and so DTRs will be split randomly between them. No transfer is done on the A-REX host.

Local Caches

Each remote host has its own local disk as a cache and these caches are mounted on the A-REX host. DTRs will be sent to the host corresponding to the cache chosen for the DTR so that all cache transfers are to local disks. The session directory is not available on the remote hosts so all uploads and non-cache downloads will run on the A-REX host.

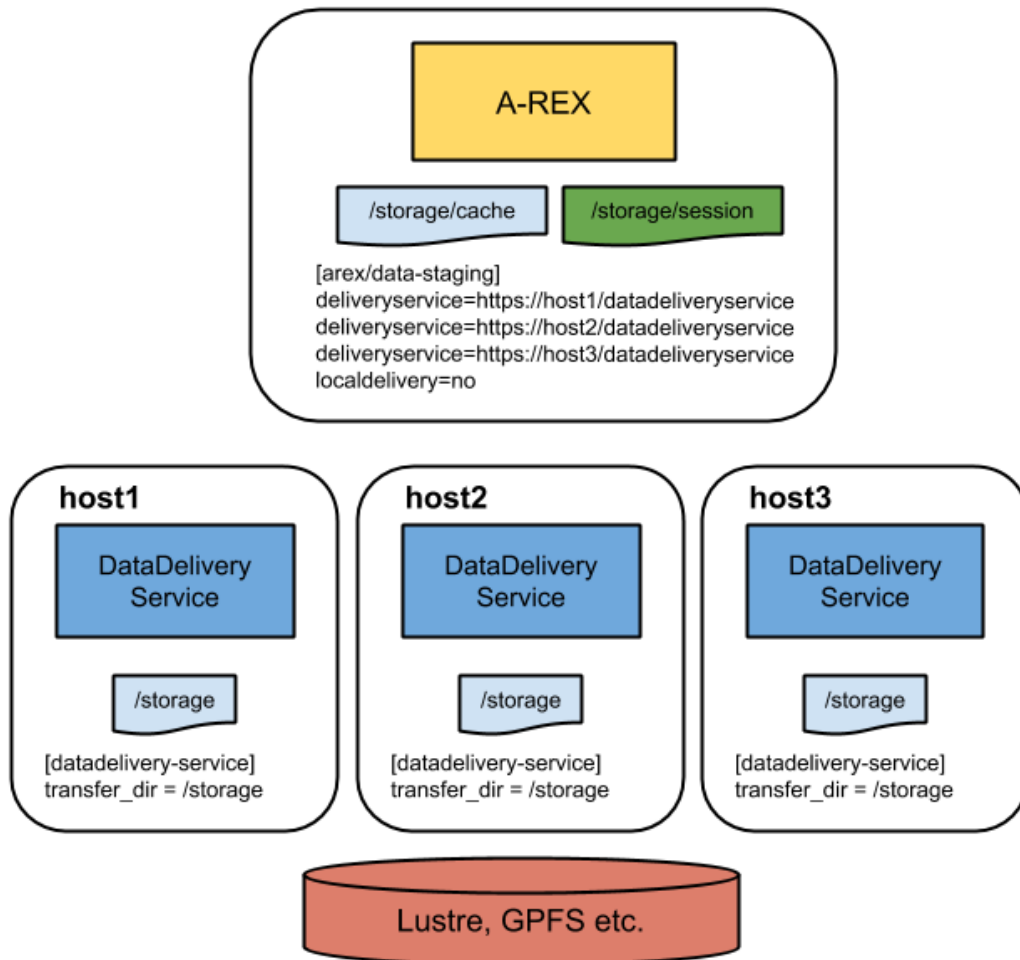


Fig. 7.9: DDS deployment: shared storage

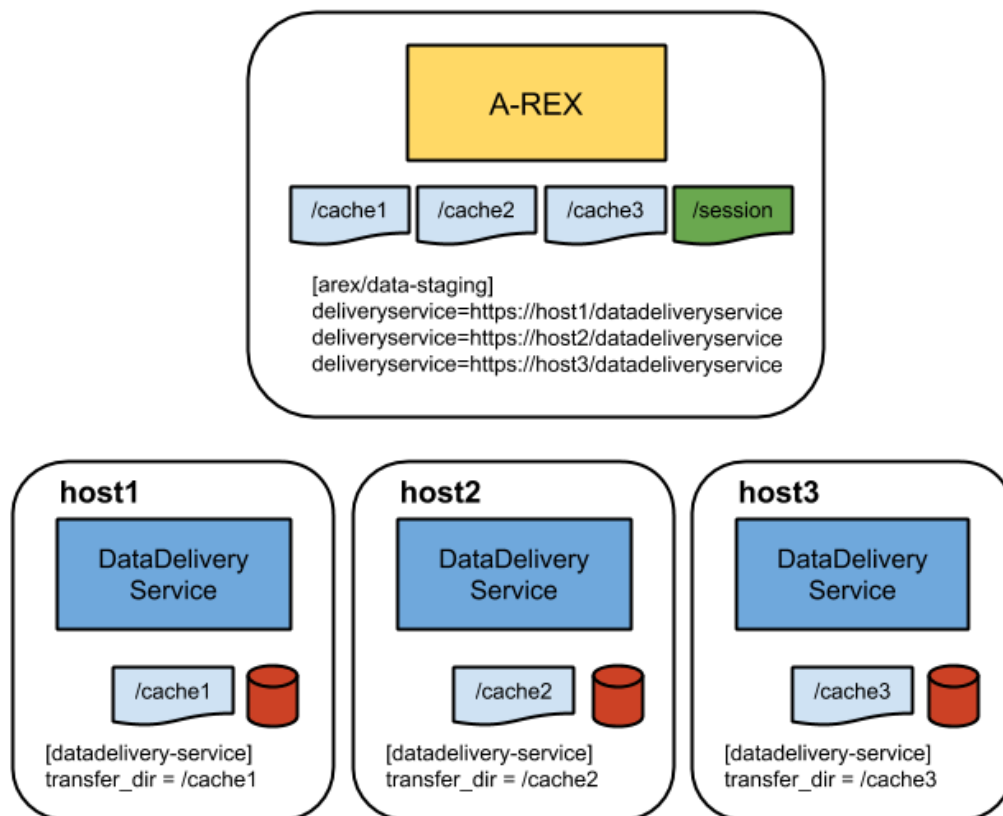


Fig. 7.10: DDS deployment: local disk for cache on DDS nodes

Security

The remote datadelivery service may require a host certificate (see above) and allow incoming connections from the A-REX host, but no other incoming connections are required. It must have outbound connectivity to the world to perform the transfers. As explained above, access is normally restricted to the A-REX host and transfers restricted to the cache and session directories. For each transfer A-REX delegates the credentials of the user who submitted the job to the remote service, which creates a temporary file containing the credentials. This file is deleted when the transfer finishes. Creating a file should become unnecessary if pure in-memory credentials are fully supported by ARC and all transfer protocols.

The transfer process itself is executed under the uid of the session directory owner, unless the transfer is to cache in which case it is executed by root. It is therefore important that user accounts are synchronised across all hosts. The file systems with the cache and session directories must be mounted on the same paths with the same user access rights on all hosts.

Proxies

HED services do not support legacy proxies such as those generated by default by voms-proxy-init. In order to use remote datadelivery services, the jobs must be submitted to A-REX using RFC proxies, which can be generated by arcproxy or giving the option -rfc to voms-proxy-init. If a legacy proxy is used, local transfer will be used even if it is disabled in the configuration.

Monitoring Remote Hosts

When the first DTR is received by the Scheduler, it pings all the configured remote datadelivery services to check that they are running and to get the list of allowed directories from each one. If a service is unreachable it will not be used. If all services are unreachable then local delivery will be used even if it is turned off in the configuration. The allowed directories information is used to direct DTRs to services where the transfer is allowed. This procedure is repeated every 5 minutes (for the first DTR received after the 5 minute limit), and only the successful services are used until the next check (unless none are successful, in which case the check is done for every DTR until one succeeds). This means that configuration changes in the remote hosts will be picked up automatically after some time. However any changes in A-REX's configuration, such as adding a new remote host, require an A-REX restart to be effective.

If remote datadelivery services are enabled, the number of DTRs assigned to each one can be seen on separate plots in Gangliarc.

7.1.4 CandyPond technical description

This page describes technical details of CandyPond, which stands for “Cache and deliver your pilot on-demand data”.

Description and Purpose

The ARC caching system automatically saves to local disk job input files for use with future jobs. The cache is completely internal to the computing element and cannot be accessed or manipulated from the outside. CandyPond exposes various operations of the cache to the outside and can be useful in a pilot job model where input data for jobs is not known until the job is running on the worker node. When the pilot picks its payload it can contact CandyPond to gain access to a file that is already cached, or if it is not cached ask CandyPond to download the file to cache.

Installation and Configuration

CandyPond is an integral part of A-REX and is available as part of the `nordugrid-arc-arex` package.

It is enabled in A-REX by adding the block `[arex/ws/candypond]` to `arc.conf`. The `[arex]`, `[arex/ws]` and `[arex/data-staging]` blocks are also required.

Runtime Environment Configuration

A runtime environment `ENV/CANDYPOND` exists to provide a convenient python module `arccandypond` to the job running on the worker node. This can either be used as a command line interface or as a python API.

Note that the `ENV/PROXY` runtime environment is also needed in order to have access to the proxy on the worker node.

`ENV/CANDYPOND` will automatically detect the correct URL of the CandyPond service, but if it is desired to use a different URL then it can be set with

```
arcctl rte params-set ENV/CANDYPOND CANDYPOND_URL <url>
```

Command Line Interface

`arccandypond get <url> <file>` can be used in place of whatever usual command the job would use to download input data. This command asks Candypond to download the url to cache if not already present, and link to the file specified in the job’s working directory.

`arccandypond check <url>` can be used to check if the given url already exists in the cache. It will exit with 0 if the file is present, 1 if not, or 2 if an error occurred.

Python API

The job can import the module and use the `CacheLink` and `CacheCheck` methods to perform the equivalent of `get` and `check` commands.

Example Use Case

In this example a job is submitted which uses `arccandypond` to download input data to cache and have it available to the job.

The `xrsl` file defines the required runtime environments. Note that no input files are specified.

```
$ cat candypond.xrsl
&
("executable" = "candypond.sh")
("runtimeenvironment" = "ENV/CANDYPOND")
("runtimeenvironment" = "ENV/PROXY")
("jobname" = "candypond_test" )
("walltime" = "3600" )
("cputime" = "3600" )
("stderr" = "stderr")
("stdout" = "stdout")
("gmlog" = "gmlog")
("outputfiles" =
  ("stdout" "")
  ("stderr" "")
  ("gmlog" ""))
)
```

The job script uses `candypond` to download the input file to cache and link to the job's working directory:

```
$ cat candypond.sh
#!/bin/sh
arccandypond get http://www.nordugrid.org:80/data/run.sh run.sh
ls -lrt
echo
cat run.sh
```

Submit the job:

```
$ arccsub candypond.xrsl
Job submitted with jobid: https://...
```

Check the output:

```
$ arccat https://...
{'http://www.nordugrid.org:80/data/run.sh': ('0', 'Success')}
total 28
-rwx----- 1 dcameron dcameron 257 Apr 10 20:23 candypond.sh
drwx----- 4 dcameron dcameron 4096 Apr 10 20:23 arc
-rw----- 1 dcameron dcameron 10662 Apr 10 20:23 user.proxy
-rw----- 1 dcameron dcameron 0 Apr 10 20:23 stderr
lrwxrwxrwx 1 dcameron dcameron 89 Apr 10 20:23 run.sh -> /opt/var/arc/cache/
→joblinks/eOGODmV3WYunPSAtDmVmuSEmABFKDmABFKDmJSFKDmGBFKDm5lgA5m/run.sh
-rw----- 1 dcameron dcameron 62 Apr 10 20:23 stdout

#!/bin/sh

GCC=`which g++ 2>/dev/null`
echo $GCC
if [ -z $GCC ]; then
  echo "Could not find the g++-compiler!"
  exit 0
fi
```

(continues on next page)

(continued from previous page)

```
make
chmod 755 prime
./prime $1
```

Note the symbolic link to the cache.

Issues and Notes

- Calls to `arccandypond get` may block for a long time if the file needs to be downloaded to cache and A-REX is already busy with data staging or the file is very large. A timeout option will be added in the future.
- CandyPond (like the A-REX web service interface it is part of) does not accept legacy proxies. This type of proxy is created by default with older versions of `grid/voms-proxy-init`, but an RFC-compliant proxy can be generated using the `-rfc` option.
- CandyPond links files to the session dir. If a scratch directory is used for executing the job, the cache files are moved there from the session directory. This requires that the scratch dir is accessible from the CandyPond host, so it cannot be used in situations where the scratch directory can only be accessed by the underlying LRMS.

7.2 ARC Data Services Overview

Warning

Whiteboard sketches to be replaced by series of digital figures

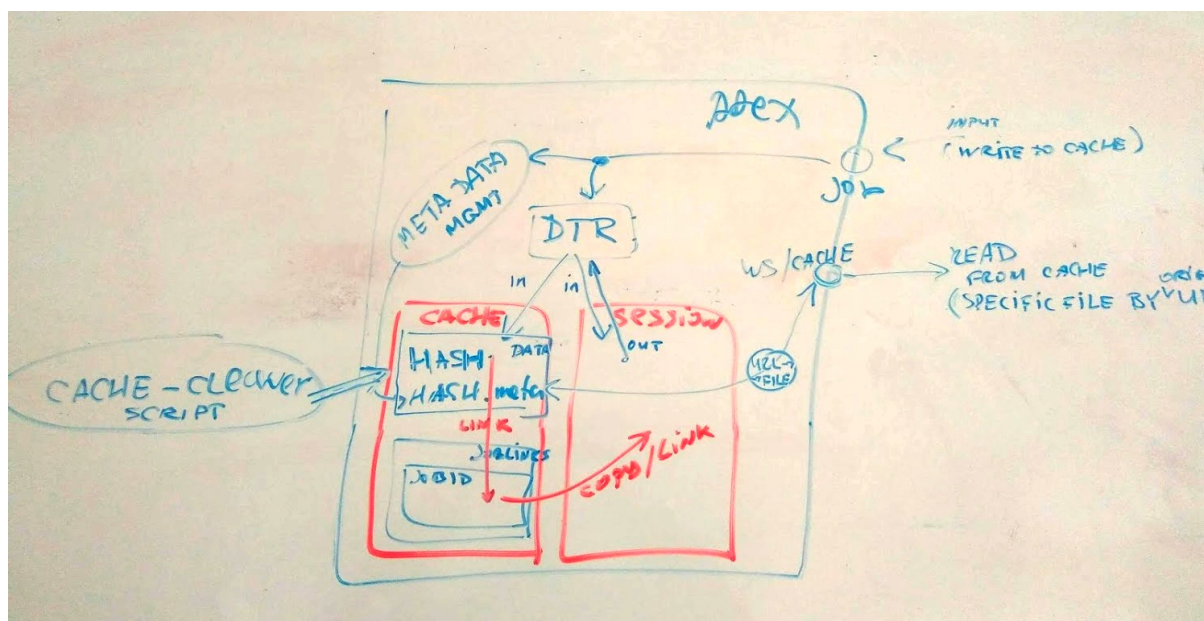


Fig. 7.11: Using A-REX cache (stand-alone A-REX installation)

Fig. 7.12: Available ARC data services that coupled with A-REX cache/job data management

General accounting configuration and operations flows are described in *Accounting Subsystem*. This section contains more technical details about implementation of each component of accounting subsystem.

Collecting the accounting information

- `.local` file contains general information associated with the job. All IDs, ownership, `authToken` attributes are taken from this file. The data in `.local` are written and updated by the A-REX *JobControl* modules.
- `.statistics` file is a dedicated file written by the *DTR data transfer framework* that contains data transfer measurements.
- `.diag` file is written by the *LRMS Scripts*: initially by `submit-<lrms>-job.sh`, then *JobScript* during the job execution on the Worker Node and finally by `scan-<lrms>-job.sh` that adds extract from LRMS accounting data. It contains but is not limited to *resource usage* and worker node environment data.

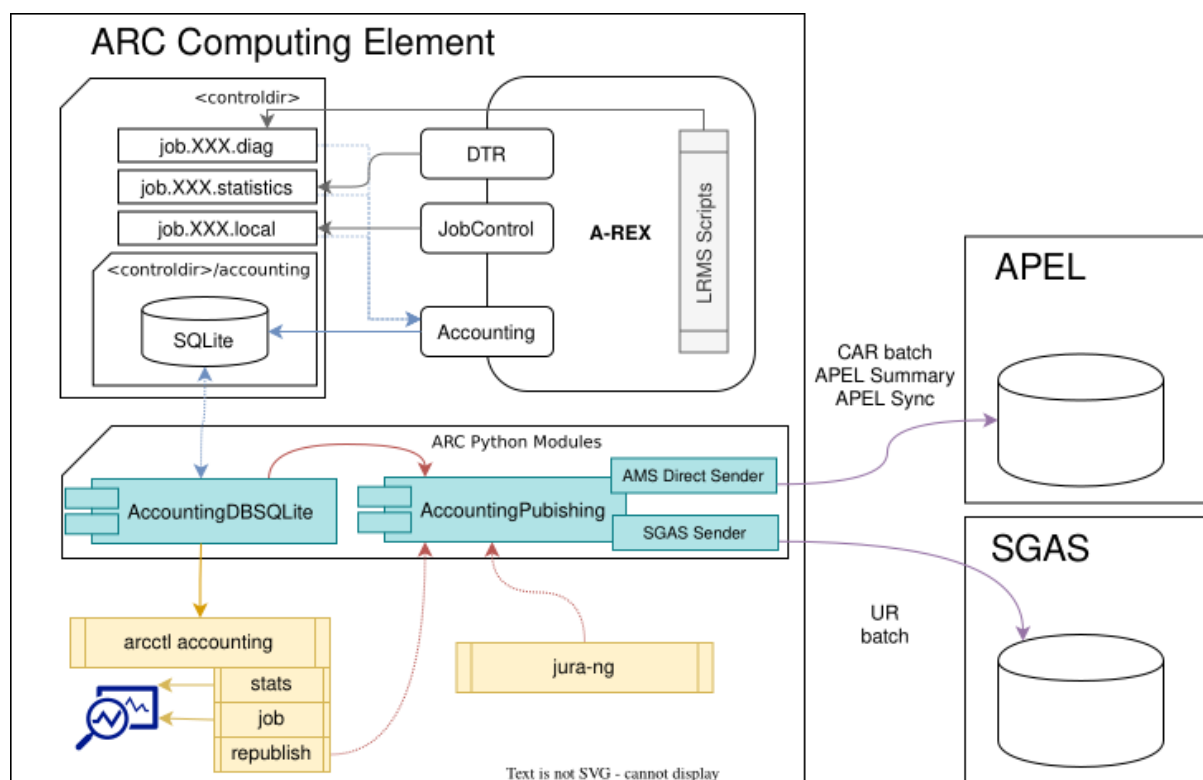


Fig. 7.13: Technical details of ARC CE accounting workflow: information collection, AAR creation, querying and publishing

The initial record about the job is created based on the first **ACCEPTED** job event. The ID, ownership and submission time is recorded during this step and accounting job status is marked as **in-progress**.

Any subsequent job events triggers event data recording in the database, and allow to track data staging time, lrms enqueueing time, etc.

When the **FINISHED** job event occurs (execution is completed) the A-REX *Accounting* subsystem updates all AAR metrics in the database, storing resource usage, endtime, etc. Such a state is indicated by the `status={completed|failed|aborted}`.

Using the local accounting database

Using the accounting data for statistics lookup and/or publishing to external services is accomplished via the developed `arc.control` Python modules.

The `AccountingDBSQLite` module is responsible for handling all low-level database operations and it hides SQL queries under the API needed for other workflows.

In particular the accounting subsystem of *ARC Control Tool* provides a command line interface to the *typical queries* that can get you the accounting data in a flexible manner.

The records publishing is carried out by the `AccountingPublishing` Python module that includes:

- classes for generating usage records in OGF.98 UR, EMI CAR 1.2, APEL Summaries and APEL Sync formats
- classes that handle the records POST-ing to SGAS endpoint
- classes that handle the records sending to APEL via AMS protocol
- general wrapping classes to handle regular publishing and republishing of the data

Both the `arccctl accounting republish` tool and the `jura-ng` tool (that runs regularly by A-REX) use the same `AccountingPublishing` Python module.

The regular publishing process stores the last published record endtime in the dedicated Publishing database. The next round of regular publishing queries the stored time and query the records since then.

7.3.2 Accounting data publishing details

Reporting to SGAS

SGAS has a simple custom web service interface loosely based on WS-ResourceProperties.

AccountingPublishing Python module uses the insertion method of this interface to report URs directly to the Python `httplib` library with SSL context wrapping.

To increase communication efficiency the `AccountingPublishing SGASSender` class sends URs in batches. SGAS accepts a batch of URs in a single request. The batch is an XML element called `UsageRecords`, containing elements representing URs. The maximal number of URs in a batch can be set as a `urbatchsize` configuration parameter of SGAS target.

Reporting to APEL

Changed in version 6.12.

APEL currently uses AMS REST protocol for records sending.

The `AccountingPublishing APELAMSDirectSender` class implements the AMS REST communication without external dependencies.

Communication code relies on the same Python `httplib` library with SSL context wrapping.

It connects to AMS endpoint with valid SSL context with client certificate authentication and obtains AMS authentication token.

Messages sent to APEL are S/MIME signed using `openssl` binary tool and then sent to the endpoint using AMS authentication token.

Reporting to APEL also works with sending records in batches. The default `urbatchsize` value is set to 500 according to APEL recommendations but can be lowered if you run into message size issues (e.g. sending large individual records).

Republishing

Republishing simply triggers the same `AccountingPublishing` classes for the defined timeframe that comes from the command line.

All records are regenerated from accounting database data and sent to the target.

7.3.3 Security

The accounting directory `<controldir>/acconting` is by default accessible only by the user running A-REX (root in most cases).

All usage records are submitted with use of the X.509 credentials specified by the value of `x509_` set of configuration options of `arc.conf`. No proxies are used for communication with accounting services.

The only access restriction made by a SGAS service is matching the Distinguished Name of the client (in this context ARC CE) with a set of trusted DNs. When access is granted, policies are then applied by SGAS, allowing either publishing and/or querying rights. Clients with publishing rights can insert any UR, regardless of content. By default, querying rights only allows retrieving URs pertaining to jobs submitted by the querying entity.

Publishing records to APEL requires `glite-APEL` endpoint defined for the grid-site in the GOCDB. The ARC CE certificate DN should be added to the `glite-APEL` endpoint.

7.3.4 Third-party accounting queries

The *ARC Control Tool* accounting stats interface is powerful enough to get custom information from the accounting database as shown in *examples*.

However if you want to get a specific report or integrate ARC accounting database with third-party software you can of course use SQLite directly.

The SQLite database file location is: <controldir>/accounting/accounting.db.

It is worth to be aware of the *ARC Accounting Database Schema* to develop third-party queries.

7.3.5 Definition of the A-REX Accounting Record including attribute mappings to SGAS and APEL

ARC CE is measuring and collecting a lot of accounting information needed but not limited to the data required by common aggregated accounting SGAS and APEL services.

All accounting information stored about a job is defined by what we called *A-REX Accounting Record (AAR)*.

AARs has a representation inside the local accounting database according to schema and representations inside A-REX and Python modules.

Local stats are generated based on the stored AARs information and provides the way for on-site CE operations analyses.

The following tables include a flat list of the properties (NOT the database rendering) included into the AAR:

Table 7.2: Attributes used in current implementation

A-REX Accounting Record (AAR)	SGAS OGF-UR	APEL CAR	Content description
jobid	JobIdentity. GlobalJobId, RecordIdentity is composed of jobid and hostname taken from the endpointurl.	JobIdentity. GlobalJobId, RecordIdentity is composed of jobid and hostname taken from the endpointurl.	The global unique jobid assigned by A-REX.
localid	JobIdentity. LocalJobId	JobIdentity. LocalJobId	LRMS job ID
jobname	JobName	JobName	User specified job name
endpointurl	MachineName	MachineName, SubmitHost, Site	The A-REX job submission endpoint URL used for this job
endpointtype	not used	not used	The A-REX job submission endpoint type used for this job
lrms	not used	Infrastructure (used as a part of it)	The LRMS behind A- REX
queue	Queue	Queue	The name of the LRMS queue of the job
nodename	Host	Host	WN name(s) as given by LRMS separated by :
clienthost	SubmitHost (port removed)	not used	Client connection socket from the client to A-REX
usersn	UserIdentity. GobalUserName	UserIdentity. GobalUserName	The global user identity, at the moment it is the SN from the certificate
localuser	UserIdentity. LocalUserId	UserIdentity. LocalUserId	The mapped local userid

continues on next page

Table 7.2 – continued from previous page

A-REX Accounting Record (AAR)	SGAS OGF-UR	APEL CAR	Content description
authtokenattributes	UserIdentity.VO and child structures	UserIdentity.Group and UserIdentity.GroupAttribute	contains the attributes of auth token (VOMS FAQNs in current implementation)
projectname	ProjectName	UserIdentity.GroupAttribute	User-defined name of the project the job belongs to
status	Status	Status	The terminal state of an A-REX job: aborted, failed, completed
exitcode	not used	ExitStatus	The exit code of the payload in the LRMS
submissiontime	StartTime	StartTime	The timestamp of job acceptance at A-REX
endtime	EndTime	EndTime	The timestamp when the job reached the terminal state in A-REX
nodecount	NodeCount	NodeCount	Number of allocated worker nodes
inputfile	FileTransfers	not used	Details of downloaded inputfile: url, size, transfer start, transfer end, downloaded from cache
outputfile	FileTransfers	not used	Details of uploaded outputfile: url, size, transfer start, transfer end
usedmemory	Memory	Memory	Maximum virtual memory used by the job
usedmaxresident	Memory	Memory	Maximum resident memory used by the job
usedaverageresident	Memory	Memory	To be dropped from the AAR schema
usedwalltime	WallDuration	WallDuration	The measured clocktime elapsed during the execution of the job in the LRMS. No matter on how many cores, processors, nodes the user job ran on.
usedcputime	CpuDuration	CpuDuration (with type all)	The total CPU time consumed by the job. If the job ran on many cores/processors/nodes, all separate consumptions shall be aggregated in this value.
usedusercputime	CpuDuration (with type user but should not be there)	CpuDuration (with type user)	The user part of the usedcputime
usedkernelcputime	CpuDuration (with type system but should not be there)	CpuDuration (with type system)	The kernel part of the usedcputime
cores	Processors	Processors	The number of cores allocated to the job

continues on next page

Table 7.2 – continued from previous page

A-REX Record (AAR)	Accounting	SGAS OGF-UR	APEL CAR	Content description
usedscratchspace		StorageUsageBlock		The used size of scratch dir at the end of the job termination in the LRMS.
systemsoftware				The type and version of the system software (i.e. opsys, glibc, compiler, or the entire container wrapping the system software)
wninstance		ServiceLevel		Coarse-grain characterization tag for the WorkerNode, e.g. BigMemory or t2.micro (aka Amazon instance type)
RTEs				List of used RTEs, including default ones as well.
data-stagein-volume		Network class has something similar		The total volume of downloaded job input data in GBs
data-stagein-time				The time spent by the DTR sysem to download input data for the job
data-stageout-volume		Network class has something similar		The total volume of uploaded job output data in GBs
data-stageout-time				The time spent by the DTR sysem to upload output data of the job
lrms-submission-time				The timestamp when the job was handed over to the LRMS system
lrmsstarttime				The timestamp when the payload starts in the LRMS
lrmsendtime				The timestamp when the payload completed in the LRMS
benchmark		Benchmark	ServiceLevel	The type and the corresponding benchmark value of the assigned WN

Table 7.3: NOT USED SGAS or APEL attributes

SGAS OGF-UR	APEL CAR
ProcessID	
Charge	Charge
Swap	Swap

7.4 ARC CE REST interface specification

Note

The current interface version is 1.1

7.4.1 The REST API endpoint

The various functionalities of the service are accessible through HTTP(S) URL built upon following pattern:

`<service endpoint URL>/rest/<version>/<functionality>`

- `<service endpoint URL>` represents mounting point of the service and may look like `https://arc.example.org:443/arex`.
- `<version>` is two parts number separated by dot. Current version is 1.1.
- `<functionality>` is one of keywords defined below.

Further the part `<service endpoint URL>/rest/<version>` is referred as `<base URL>`.

All parts of URL to the right of hostname are case-sensitive.

Depending on Accept header in HTTP request (Accept: application/json, Accept: text/xml or Accept: application/xml), information in the response rendered in either JSON or XML format. If not specified it defaults to text/html and output is compatible with ordinary web browser.

In the HTTP response headers the HTTP Status-Code (RFC7231) indicates the status of the overall request (e.g. 403 corresponds to the forbidden).

For the operations that support multiple requests per single API call, in addition to the Status-Code in HTTP header, the per-request Status-Codes are returned. They are included as a part of the response array in HTTP body using the same RFC2731 values following the syntax defined below.

7.4.2 Description of functionalities and operations

Requesting supported versions

GET `<service endpoint URL>/rest`

Operations:

- GET - returns list of supported REST API versions
- POST, PUT, DELETE - not supported

Example response:

The XML response is like:

```
<versions>
  <version>1.0</version>
  <version>1.1</version>
  <version>1.2</version>
</versions>
```

The JSON is:

```
{version: [ "1.0", "1.1", "1.2" ]}
```

or

```
{version: "1.0"}
```

Obtaining CE resource information

GET <base URL>/info[?schema=glue2]

Operations:

- GET - retrieve generic information about cluster properties. It accepts the optional `schema` parameter. The default and only supported value in the current ARC release is `glue2`. The `CRR` rendering might be added in the future ARC releases. XML or JSON returned according to request headers.
- HEAD - supported
- PUT, POST, DELETE - not supported.

Example QUERY:

```
GET https://host.domain.org:443/arex/rest/1.0/info?schema=glue2 HTTP/1.1
Accept: application/xml
```

The XML response is:

```
<InfoRoot>
  <Domains xmlns="http://schemas.ogf.org/glue/2009/03/spec_2.0_r1"
  →xmlns:xsi="http://www.w3.org/2001/XMLSchema-instance"  xsi:schemaLocation=
  →"https://raw.githubusercontent.com/OGF-GLUE/XSD/master/schema/GLUE2.xsd">
    <AdminDomain BaseType="Domain" CreationTime="2018-11-06T20:26:46Z"
    →Validity="10800">
      <ID>urn:ad:UNDEFINEDVALUE</ID>
      <Name>UNDEFINEDVALUE</Name>
      <Distributed>false</Distributed>
      <Services>
        <ComputingService BaseType="Service" CreationTime="2018-11-
        →06T20:26:46Z" Validity="10800">
          <ID>urn:ogf:ComputingService:arc.zero:arex</ID>
          <Capability>data.transfer.ceph.srm</Capability>
          <Capability>executionmanagement.jobmanager</Capability>
        ... output omitted ...
      </ComputingService>
    </AdminDomain>
  </Domains>
```

Operating jobs

GET <base URL>/jobs[?state=<state1>[,<state2>[...]]]

POST <base URL>/jobs?action=new[&queue=<name>][&delegation_id=<id>][&instances=<number>][&instances_min=<number>]

POST <base URL>/jobs?action={info|status|kill|clean|restart|delegations}

Operations:

- GET - get list of jobs
- HEAD - supported
- POST - job submission and management
- PUT, DELETE - not supported

Get list of jobs

GET <base URL>/jobs retrieves list of jobs belonging to authenticated user as application/xml or application/json. Returned document contains list of job IDs.

It accepts the optional state parameters. When defined the returned document contains only jobs in the *requested state(s)*.

Example QUERY:

```
GET https://host.domain.org:443/arex/rest/1.0/jobs HTTP/1.1
Accept: application/xml
```

The XML response is:

```
<jobs>
  <job>
    <id>1234567890abcdef</id>
  </job>
  <job>
    <id>fedcba0987654321</id>
  </job>
</jobs>
```

The JSON is:

```
{
  "job": [
    {"id": "1234567890abcdef"},
    {"id": "fedcba0987654321"}
  ]
}
```

Job submission (create a new job)

POST <base URL>/jobs?action=new initiates creation of a new job instance or multiple jobs.

Request body contains job description, in one of the supported formats: ADL as Content-type: application/xml or xRSL as Content-type: applicaton/rsl.

The optional queue parameter defines the default value for the computing element **queue**. The value has same effect as xRSL *queue* (or ADL QueueName) and applied to all job descriptions that does not have it specified. If xRSL/ADL already contains queue the value from xRSL/ADL is used instead.

The optional delegation_id parameter defines the default value for the *delegation* ID for data staging. The value has same effect as xRSL *delegationid* (or ADL DelegationID) and applied to all job descriptions that does not have it specified. If xRSL/ADL already contains delegationid the value from xRSL/ADL is used instead.

Bulk submission of array of job descriptions

It is possible to pass multiple job descriptions of the same type in document body to submit an array:

- ADL descriptions are enclosed in <ActivityDescriptions> element
- XRSLS uses + to *merge multiple jobs*.

Response contains 201 code. Response body contains an array of elements corresponding to the sequence of the job descriptions in the requests in the same order. The elements of the array in the response contains:

- **status-code**: a 3-digit integer result code of the attempt to understand and satisfy the request (according to RFC7231)

- **reason**: a short textual description of the Status-Code
- **id**: job UUID or None if not assigned (non-successfull submission)
- **state**: the job state according to *state model* or None if not available (non-successfull submission)

The XML response is:

```
<jobs>
  <job>
    <status-code>201</status-code>
    <reason>Created</reason>
    <id>1234567890abcdef</id>
    <state>ACCEPTING</state>
  </job>
  <job>
    <status-code>500</status-code>
    <reason>Requested RTE is missing</reason>
  </job>
</jobs>
```

The JSON is:

```
{
  "job": [
    {
      "status-code": "201",
      "reason": "Created",
      "id": "1234567890abcdef",
      "state": "ACCEPTING"
    },
    {
      "status-code": "500",
      "reason": "Requested RTE is missing",
    }
  ]
}
```

Bulk submission of identical job instances

The optional **instances** parameter can be used to request the creation of **<number>** of identical jobs defined by the **same** identical job description document in the request body.

Note

When the **instances** parameter is used, passing array of multiple job descriptions is not supported and REST 1.1 will return 403 if an array of jobs are submitted.

The server response will contain successfully created job elements with unique job IDs for each job instance. Note that due to internal server limits, it is possible that fewer job instances are created than requested.

The XML response is:

```
<jobs>
  <job>
    <status-code>201</status-code>
    <reason>Created</reason>
```

(continues on next page)

(continued from previous page)

```

    <id>1234567890abcdef</id>
    <state>ACCEPTING</state>
  </job>
  <job>
    <status-code>201</status-code>
    <reason>Created</reason>
    <id>1234567890abcdeg</id>
    <state>ACCEPTING</state>
  </job>
</jobs>

```

The JSON is:

```

{
  "job": [
    {
      "status-code": "201",
      "reason": "Created",
      "id": "1234567890abcdef",
      "state": "ACCEPTING"
    },
    {
      "status-code": "201",
      "reason": "Created",
      "id": "1234567890abcdeg",
      "state": "ACCEPTING"
    }
  ]
}

```

The optional `instances_min` parameter can be used to specify the *still acceptable lowest number* of identical jobs to be created in case the server has an internal limit and can not create all requested `instances`.

The default value of the `instances_min` parameter is 1 meaning any ammount of jobs up to *instances* is acceptable by a client. The actual number of created jobs will be determined by server configuration limits.

If server is not able to create the `instances_min` jobs, entire bulk operation is cancelled and REST API returns 403 error code. Use the same value for the `instances_min` and `instances` parameters if only the exact `<number>` of identical job creation is acceptable.

Jobs management

POST `<base URL>/jobs?action={info|status|kill|clean|restart|delegations}` - job management operations supporting arrays of jobs.

Request body contains list of jobids as JSON/XML (e.g. output of GET `<base URL>/jobs` can be reused).

Example of the body in XML:

```

<jobs>
  <job>
    <id>1234567890abcdef</id>
  </job>
  <job>
    <id>fedcba0987654321</id>
  </job>
</jobs>

```

And in JSON:

```
{
  "job": [
    {"id": "1234567890abcdef"},
    {"id": "fedcba0987654321"}
  ]
}
```

Response depends on the requested action:

Job info

POST <base URL>/jobs?action=info retrieves full information about job(s) according to the GLUE2 activity information XML document, or in JSON format.

Response contains 201 code. Response body contains an array of elements corresponding to the job IDs in the requests. The elements of the array in the response contains:

- **status-code:** a 3-digit integer result code of the attempt to understand and satisfy the request (according to RFC7231). The 200 is only positive response.
- **reason:** a short textual description of the Status-Code
- **id:** job UUID
- **info_document:** GLUE2 activity information about the job or empty documents if not available (request if not satisfiable)

Job status

POST <base URL>/jobs?action=status retrieves information about job(s) current state.

Response body contains an array of elements corresponding to the job IDs in the requests. The elements of the array in the response contains:

- **status-code:** a 3-digit integer result code of the attempt to understand and satisfy the request (according to RFC7231). The 200 is only positive response.
- **reason:** a short textual description of the Status-Code
- **id:** job UUID
- **state:** the job state according to *state model* or None if not available

Killing jobs

POST <base URL>/jobs?action=kill send a request to kill job(s).

Response body contains an array of elements corresponding to the job IDs in the requests. The elements of the array in the response contains:

- **status-code:** a 3-digit integer result code of the attempt to understand and satisfy the request (according to RFC7231). The response code is 202 to indicate request is queued for later execution and is only positive response.
- **reason:** a short textual description of the Status-Code
- **id:** job UUID

Clean job files

POST <base URL>/jobs?action=clean send a request to clean job(s) files.

Response body contains an array of elements corresponding to the job IDs in the requests. The elements of the array in the response contains:

- **status-code:** a 3-digit integer result code of the attempt to understand and satisfy the request (according to RFC7231). The response code is 202 to indicate request is queued for later execution and is only positive response.
- **reason:** a short textual description of the Status-Code

- id: job UUID

Restart job

POST <base URL>/jobs?action=restart send a request to restart job(s).

Response body contains an array of elements corresponding to the job IDs in the requests. The elements of the array in the response contains:

- status-code: a 3-digit integer result code of the attempt to understand and satisfy the request (according to RFC7231). The response code is 202 to indicate request is queued for later execution.
- reason: a short textual description of the Status-Code
- id: job UUID

Job delegations

POST <base URL>/jobs?action=delegations - retrieves list of delegations associated with the job.

Response body contains an array of elements corresponding to the job IDs in the requests. The elements of the array in the response contains:

- status-code: a 3-digit integer result code of the attempt to understand and satisfy the request (according to RFC7231), 200 is only positive response
- reason: a short textual description of the Status-Code
- id: job UUID
- delegation_id: an array of assigned delegation IDs

File operations

Files belonging to specific job are operated using <base URL>/jobs/<job id> URL.

Working with session directory

GET <base URL>/jobs/<job id>/session/<path>

DELETE <base URL>/jobs/<job id>/session/<path>

PUT <base URL>/jobs/<job id>/session/<path>

Operations:

- GET, HEAD, PUT, DELETE - supported for files stored in job's session directory and perform usual actions.
- GET, HEAD - for directories retrieves list of stored files (consider WebDAV for format)
- DELETE - for directories removes whole directory
- PUT - for directory not supported.
- POST - not supported.

Delegation functionality

GET <base URL>/delegations[?type={x509|jwt}]

POST <base URL>/delegations?action=new[&type={x509|jwt}]

Operations:

- GET - retrieves list of delegations belonging to authenticated user
- HEAD - supported
- POST - create new delegation

- PUT, DELETE - not supported

POST <base URL>/delegations/<delegation id>?action=get,renew,delete

PUT <base URL>/delegations/<delegation id>

Operations:

- GET, HEAD - not supported
- POST - manage particular delegation ID
- PUT - store x509 delegation public part for particular delegation ID

Get list of delegations

GET <base URL>/delegations[&type={x509|jwt}] - retrieves list of delegations belonging to authenticated user. It accepts the optional `type` parameter that allows to filter delegations based on the type. By default all types are returned. Supported values are: `x509` for proxy-certificate delegation and `jwt` for data staging token delegation.

QUERY:

```
GET https://host.domain.org:443/arex/rest/1.0/delegations HTTP/1.1
Accept: application/xml
```

The XML response is:

```
<delegations>
  <delegation>
    <id>1234567890abcdef</id>
    <type>x509</type>
  </delegation>
  <delegation>
    <id>fedcba0987654321</id>
    <type>jwt</type>
  </delegation>
</delegations>
```

The JSON formatted response (make consistent across specification):

```
{
  "delegation": [
    { "id": "1234567890abcdef", "type": "x509" },
    { "id": "fedcba0987654321", "type": "jwt" }
  ]
}
```

New delegation

Delegation protocol depends on the delegation type (`x509` or `jwt`) specified with an optional `type` parameter. If not explicitly specified the default delegation type is `x509` for backward compatibility with REST API 1.0.

X.509 delegation

X.509 delegation is a 2-step process:

1. Step 1 generates pair of private/public keys on server side and communicates X.509 certificate request to the client.
2. Client signs public key and stores delegated certificate to finish delegation procedure.

Corresponding REST API calls:

1 step

POST <base URL>/delegations?action=new&type=x509 starts a new delegation process. Response is 201 and contains certificate request of application/x-pem-file type and URL of delegation in Location HTTP header with assigned delegation id.

2 step

PUT <base URL>/delegations/<delegation id> stores public part (2nd step). Request body contains signed certificate (Content-type: application/x-pem-file). Response is 200 on success.

JWT delegation

JWT delegation is a single API request.

POST <base URL>/delegations?action=new&type=jwt stores provided JWT token in the ARC CE delegation database. Request should contain the X-Token-Delegation header that provides bearer token. No verification of the delegation token is performed on ARC CE side. It will be passed as it is via Authorization header to the endpoints supporting JWT. Response is 200 and contains URL of delegation in Location HTTP header with assigned delegation id.

QUERY:

```
POST https://host.domain.org:443/arex/rest/1.0/delegations?
→action=new&type=jwt HTTP/1.1
X-Token-Delegation: bearer_
→eyJhbGciOiJSUzI1NiIsInR5cCIgOiAiSldUIiwia2lkIiA6ICI3ZzZGSzBPam43YW.
→...
```

Delegations management

Delegations are managed one-by-one. The same delegation ID can be re-used for multiple jobs (submitted separately or in batch).

The delegation ID to be used in the job context required to be either explicitly specified as a part of the job description in a description language defined way (e.g. DelegationID in ADL and *delegationid* in xRSL) or passed as delegation_id parameter during job submission.

POST <base URL>/delegations/<delegation id>?action=get,renew,delete used to manage delegation.

Request body is empty and action is defined by action value.

Response is structured depending on the action:

Get delegation

POST <base URL>/delegations/<delegation id>?action=get depending on delegation type returns:

- for x509 public part of the stored delegation (application/x-pem-file content type)
- for JWT stored delegation token (application/jwt content type)

Delete delegation

POST <base URL>/delegations/<delegation id>?action=delete removes delegation. Response is 200 with no body expected.

Renew delegation

The process is similar to creation of the new delegation and depends on delegation type (x509 or jwt).

For x509 the POST <base URL>/delegations/<delegation id>?action=renew API call initiates renewal of delegation. Response is 201 with certificate request of application/x-pem-file type that should be followed by the PUT <base URL>/delegations/<delegation id> call for signed certificate upload.

For jwt the token stored in the ARC CE delegation database will be replaced by the new one supplied via X-Delegation header. Response is 200 on success.

A-REX control directory files access for debugging purposes

GET <base URL>/jobs/<job id>/diagnose/<file type>

Operations:

- GET - return the content of file in A-REX control directory for requested jobID
- HEAD - supported
- POST, PUT, DELETE - not supported

The <file type> matches the controldir file suffix and can be one of the following:

- failed
- local
- errors
- description
- diag
- comment
- status
- acl
- xml
- input
- output
- input_status
- output_status
- statistics

7.4.3 REST Interface Job States

Table 7.4: State identifiers used with ARC REST API

REST API State Name	Description	A-REX Internal State
ACCEPT-ING	This is the initial job state. The job has reached the cluster, a session directory was created, the submission client can optionally upload files to the sessiondir. The job waits to be detected by the A-REX, the job processing on the CE hasn't started yet	ACCEPTED
ACCEPTED	In the ACCEPTED state the newly created job has been detected by A-REX but can't go to the next state due to an internal A-REX limit. The submission client can optionally upload files to the sessiondir.	PEND-ING:ACCEPTED
PREPAR-ING	The job is undergoing the data stage-in process, input data is being gathered into the session directory (via external downloads or making cached copies available). During this state the submission client still can upload files to the session directory. <i>This is an I/O heavy job state.</i>	PREPARING
PREPARED	The job successfully completed the data stage-in process and is being held waiting in A-REX's internal queue before it can be passed over to the batch system	PEND-ING:PREPARING
SUBMIT-TING	The job environment (via using RTEs) and the job batch submission script is being prepared to be followed by the submission to the batch system via using the available batch submission client interface	SUBMIT
QUEUING	The job is under the control of the local batch system and is "queuing in the batch system", waiting for a node/available slot	INLRMS
RUNNING	The job is under the control of the local batch system and is "running in the batch system", executing on an allocated node under the control of the batch system	INLRMS
HELD	The job is under the control of the local batch system and is being put on hold or being suspended, for some reason the job is in a "pending state" of the batch system	INLRMS
EXIT-INGLRMS	The job is under the control of the local batch system and is finishing its execution on the worker node, the job is "exiting" from the batch system either because the job is completed or because it was terminated	INLRMS
OTHER	The job is under the control of the local batch system and is in some "other" native batch system state which can not be mapped to any of the previously described batch systems states.	INLRMS
EXECUTED	The job has successfully completed in the batch system. The job is waiting to be picked up by the A-REX for further processing or waiting for an available data stage-out slot.	PEND-ING:INLRMS
FINISHING	The job is undergoing the data stage-out process, A-REX is moving output data to the specified output file locations, the session directory is being cleaned up. Note that failed or terminated jobs can also undergo the FINISHING state. <i>This is an I/O heavy job state</i>	FINISHING
FINISHED	Successful completion of the job on the cluster. The job has finished ALL its activity on the cluster AND no errors occurred during the job's lifetime.	FINISHED
FAILED	Unsuccessful completion of the job. The job failed during one of the processing stages. The job has finished ALL its activity on the cluster and there occurred some problems during the lifetime of the job.	FINISHED
KILLING	The job was requested to be terminated by an authorized user and as a result it is being killed. A-REX is terminating any active process related to the job, e.g. it interacts with the LRMS by running the job-cancel script or stops data staging processes. Once the job has finished ALL its activity on the cluster it will be moved to the KILLED state.	CAN-CELLING
KILLED	The job was terminated as a result of an authorized user request. The job has finished ALL its activity on the cluster.	FINISHED
WIPED	The generated result of jobs are kept available in the session directory on the cluster for a while after the job reaches its final state (FINISHED , FAILED or KILLED). Later, the job's session directory and most of the job related data are going to be deleted from the cluster when an expiration time is exceeded. Jobs with expired session directory lifetime are "deleted" from the cluster in the sense that only a minimal set of info is kept about such a job and their state is changed to WIPED	DELETED

7.4.4 Status of This Document

This document provides normative specification for the ARC REST Interface version 1.1.

This specification was designed by the requirements listed below:

1. Support for versioning: via URL paths like <https://arc.zero:443/arex/rest/1.1/jobs>
2. Usable with simple tools (wget, curl)
3. Friendly to common HTTP REST frameworks
4. Interactive access to session directory content
5. Machine readable error/result codes/messages
6. No drastic changes to information representation and jobs handling
7. Support for different response formats: xml, json
8. Support for bulk operations

7.5 ARCHERY data model and DNS records rendering

ARCHERY implies a minimalistic data model that is targeting mostly static information to address the distributed e-Infrastructure resource grouping, software environments provisioning and service discovery needs.

The ARCHERY objects and their relations are embedded into the DNS infrastructure as a dedicated Resource Record sets according to the rendering described below.

7.5.1 ARCHERY objects

To represent the distributed e-infrastructure concept the ARCHERY data model defines three type of objects: the **Endpoint Object**, the **Service Object** and the **Group Object**.

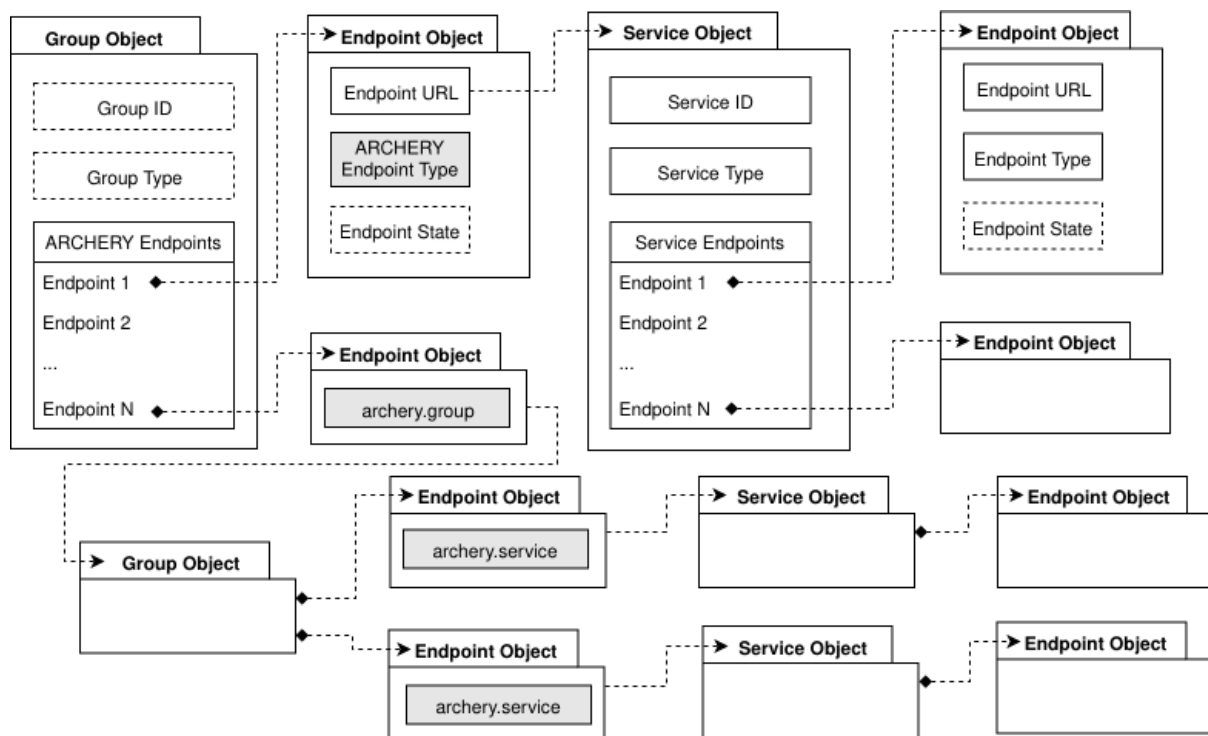


Fig. 7.14: ARCHERY e-Infrastructure topology objects, their attributes and relations in the data model.

To represent the *community software environments* in the e-infrastructure two additional type of objects are defined in the ARCHERY data model: the **Software Object** and the **RTE Object**.

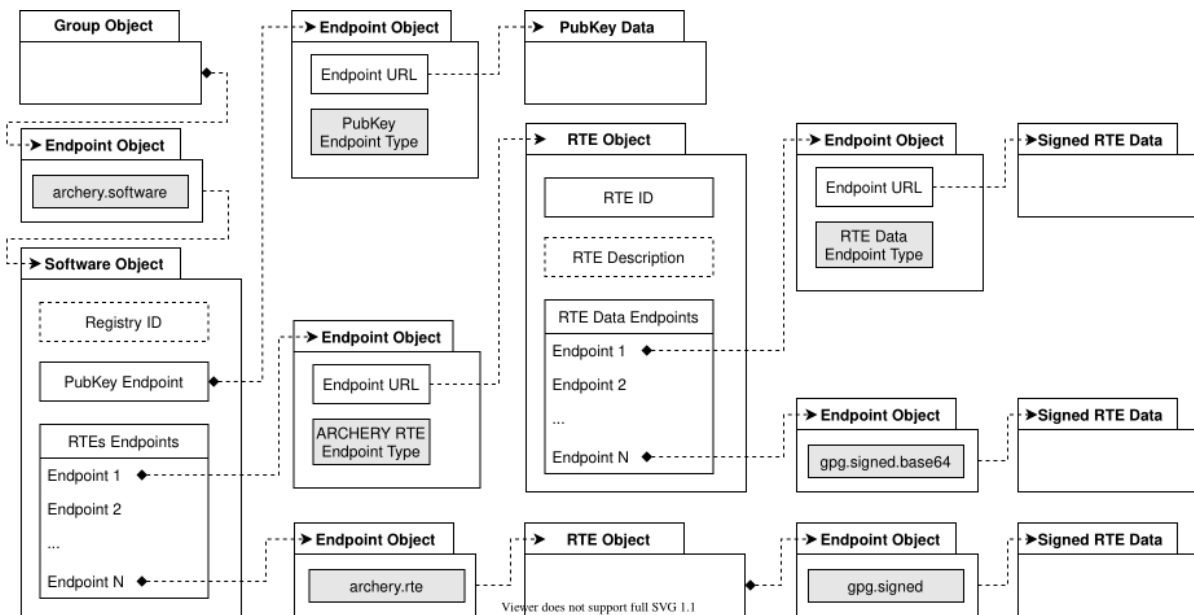


Fig. 7.15: ARCHERY software environment objects, their attributes and relations in the data model.

Endpoint object

Endpoint object is used to capture information about a network location that can be used to access specific service functionality, including accessing information within ARCHERY registry service itself. For the later special endpoint types were defined (see below). The Endpoint object is described with the following set of attributes:

- **Endpoint URL** - defines the network location by which the service functionality is accessible;
- **Endpoint Type** - contains the value from the Endpoint type enumeration defined by the e-Infrastructure operators¹. In addition to ordinary service endpoint types two special ARCHERY endpoint types `archery.group` and `archery.service` had been introduced for accessing the ARCHERY Group and Service objects within the registry. Another special types includes `archery.software` and `archery.rte` that used to access ARCHERY Software and RTE objects respectively when community software registry is embedded into the DNS.
- **Endpoint State (optional)** - Boolean value that indicates the endpoint availability. In case of an ARCHERY service endpoint, if the endpoint state is false (unavailable) it is still used to represent e-Infrastructure topology, but should not be used during service endpoints discovery. Missing attribute is interpreted as true i.e. the endpoint is available.

Service object

Service object represents an e-Infrastructure service, like Computing Element, Storage Element, etc. The Service object is described with the following set of attributes:

- **Service ID** - holds the service identifier as an arbitrary string. For example the ID can be derived from the hostname;
- **Service Type** - contains the value from the Service type enumeration defined by the e-Infrastructure operators;
- **Service Endpoints** - an array of Endpoint objects associated with the Service.

¹ The `org.ogf.glue.emies.activiticreation`, `org.nordugrid.ldapglue2`, `org.nordugrid.arcrest` are examples of endpoint types used in NorduGrid infrastructure.

Group object

Group object is used to organize other objects such as ARCHERY Service or Group. The grouping was introduced to the data model so that infrastructure topologies (i.e. hierarchies or federations) can be represented inside the registry in a flexible way. The Group object is described with the following set of attributes:

- **Group ID (optional)** - holds the group identifier as an arbitrary string.
- **Group Type (optional)** - defines the grouping type based on organizational structure; example values could be: Site, Country, Tier, etc.
- **ARCHERY Endpoints** - an array of Endpoint objects of special defined ARCHERY endpoint types (`archery.group`, `archery.service` or `archery.software`). These Endpoint objects describe the optional state and URL of the ARCHERY objects inside the registry.

Software object

Software objects are used to represent a *community-defined RTEs* registry instances. The Software object provides community public key data and enumerates the available RTEs. It is described with the following set of attributes:

- **Registry ID (optional)** - holds the software registry identifier as an arbitrary string.
- **PubKey Endpoint** - and Endpoint object that contains URL to the public key location. It can points to external URLs or raw data inside the DNS. Technical implementation uses `gpg.pubkey` and `gpg.pubkey.base64` endpoint types depending on Base64 encoding usage.
- **RTEs Endpoints** - and array of Endpoint objects of special defined ARCHERY endpoint type `archery.rte`. These Endpoint objects describe the optional state and URL of the ARCHERY RTE objects inside the registry.

RTE Object

RTE Object is used to represent a particular *community-defined RTE script*. It encapsulates the RTE ID (according to the naming and versioning convention), optional RTE description and contains Endpoints to signed RTE scripts. It is described with the following set of attributes:

- **RTE ID** - holds the RTE identifier according to naming and versioning scheme
- **RTE Description (optional)** - an arbitrary string that contains the human-readable description of RTE
- **RTE Data Endpoints** - and array of Endpoint objects that points to signed RTE data. It can points to external URLs or raw data inside the DNS. Technical implementation uses `gpg.signed` and `gpg.signed.base64` endpoint types depending on Base64 encoding usage.

Entry point group object

Following the data model the service endpoint discovery process can start from the **entry point** Group Object and recursively contacting all the ARCHERY Endpoints.

The Software objects can be attached to any parent Group Object and discovered using the same recursive traversal of the e-Infrastructure topology.

ARCHERY object attributes allow taking into account group, service and/or endpoint types as well as availability status during the recursive discovery process to query only the subset of objects.

7.5.2 DNS Resource Records rendering

Rendering ARCHERY objects

ARCHERY objects are rendered using TXT resource records (RRs). These TXT RRs contain the space separated key=value pairs where keys correspond to the data model object attributes. Boolean object attributes have values specified as 0 or 1.

The **Endpoint Object** is rendered with a single TXT RR where Endpoint Object attributes – *Endpoint URL*, *Endpoint Type* and optional *Endpoint State* are represented by the u, t, s keys in this specific order within the space separated key-value pairs.

The **Service Object** of the ARCHERY data model is rendered by a dedicated RRSet identified by unique domain name that can be used to access this object inside the registry.

The service RRSet is composed of single service object identity RR and one RR for every Endpoint objects associated with the service. The service object identity RR has the format of o=service followed by t=<Service Type> and id=<Service ID>:

```
<DNS name> TXT "o=service t=<Service Type> id=<Service ID>"
<DNS name> TXT "u=<Endpoint URL> t=<Endpoint Type> [s={0|1}]"
<DNS name> TXT ...
<DNS name> TXT "u=<Endpoint URL> t=<Endpoint Type> [s={0|1}]"
```

The **Group Object** of is rendered by a dedicated RRSet identified by unique domain name that can be used to access this object inside the registry.

The Group RRSet is composed of single Group object identity RR and one RR for every ARCHERY Endpoint objects pointing to other Service or Group objects that are part of this specific Group. The Group object identity RR has the format of o=group followed by optional t=<Group Type> and optional id=<Group ID>:

```
<DNS name> TXT "o=group [t=<Group Type>] [id=<Group ID>]"
<DNS name> TXT "u=<DNS URL> t=archery.{group|service} [s={0|1}]"
<DNS name> TXT ...
<DNS name> TXT "u=<DNS URL> t=archery.{group|service} [s={0|1}]"
```

The **Software Object** of the ARCHERY data model is rendered by a dedicated RRSet identified by unique domain name that can be used to access this object inside the registry.

The software object RRSet is composed of single object identity RR, PubKey Endpoint object RR and one RR for every RTEs Endpoint objects associated with the registry. The software object identity RR has the format of o=software followed by optional id=<Registry ID>:

```
<DNS name> TXT "o=software [id=<Registry ID>]"
<DNS name> TXT "u=<PubKey URL> t=gpg.pubkey[.base64]"
<DNS name> TXT "u=<DNS URL> t=archery.rte [s={0|1}]"
<DNS name> TXT ...
<DNS name> TXT "u=<DNS URL> t=archery.rte [s={0|1}]"
```

The **RTE Object** of the ARCHERY data model is rendered by a dedicated RRSet identified by unique domain name that can be used to access this object inside the registry.

The RTE object RRSet is composed of single object identity RR, and one RR for every RTE Data Endpoint object pointing to signed RTE data location. The RTE object identity RR has the format of o=rte followed by id=<RTE ID> and optional d=<RTE Descition>. The description should be encoded to eliminate spaces, using the URLEncode is advised. Typically there is only one RTE Data Endpoint object that point to either DNS or external URL. Several RTE Data Endpoints imply external URL redundancy:

```
<DNS name> TXT "o=rte id=<RTE ID> [d=<RTE Description>]"
<DNS name> TXT "u=<RTE Data URL> t=gpg.singed[.base64] [s={0|1}]"
```

Embedding ARCHERY object relations

In the registry rendering the entry point Group Object is distinguished by the predefined RRSet DNS name starting with the mandatory `_archery`. Any other objects may have arbitrary RRSet name in any DNS zone. Exact RRSet naming scheme does not affect service discovery process, but should be consistent and transparent from the operational point of view.

The child-parent Endpoint Object to Service Object relation is implicitly defined by means of grouping within the same RRSet. The same applies to “RTE Data Endpoint Object to RTE Object” and “PubKey Endpoint object to Software Object” relations.

The ARCHERY Data model allows grouping of Service, Group and Software Objects into Group objects via the intermediate Endpoint objects. This grouping in the DNS rendering is implemented by using dedicated ARCHERY Endpoint types that are contained as RR text inside the Group object’s RRSet. The same applies to “RTEs Endpoint Object to Software Object” relations.

Both Group Object to Group Object, Group Object to Service Object, Group Object to Software Object and Software Object to RTE Object relations rest upon the unique DNS name of the RRSet of the pointed object.

7.6 A-REX Technical Description

7.6.1 Internal files of the A-REX

A-REX stores information about jobs in files in the control directory. Information is stored in files to make it easier to recover in case of failure, but for faster processing job state is also held in memory while A-REX is running.

The files and sub-directories in the control directory and their formats are described below:

- `version2` - empty file serves as indication that information is stored in this folder according to second format version. This file is mandatory.
- `logs` – sub-directory with information prepared for reporting plugins.
- `dhparam.pem` - file with Diffie-Hellman parameters for establishing TLS connection with A-REX server. This file is generated first time A-REX is started and it may take some time till it is populated.
- `dtr.state` - file with current state of data statging functionality
- `gm-heartbeat` - modification time of this file is continuously updated by A-REX to indicate it’s main processing loop is running
- `gm.fifo` - FIFO communication channel to A-REX mostly used by babckend scripts to indicate jobs whose state have cahnged
- `info.xml` - file with current state of A-REX CE expressed in GLUE2 XML
- `accepting`, `restarting`, `processing`, `finished` - sub-directories which store current state of jobs being processed split accrding to state (see below). Following files are stored there:
 - `ID.status` – file with current state of the job. Here ID corresponds to arbitrary ASCII string assigned to submitted job. Currently ID is made of 12 lowercase hex symbols. But that may change without notice. This is a plain text file containing a single word representing the internal name of current state of the job. Possible values and corresponding external job states are:
 - * `ACCEPTED`
 - * `PREPARING`
 - * `SUBMIT`
 - * `INLRMS`
 - * `FINISHING`
 - * `FINISHED`

* CANCELING

* DELETED

See corresponding Section for a description of the various states. Additionally each value can be prepended the prefix “PENDING:” (like PENDING:ACCEPTED, see corresponding Section). This is used to show that a job is ready to be moved to the next state but it has to stay in it’s current state only because otherwise some limits set in the configuration would be exceeded.

These files are split through aforementioned sub-directories according to following rules:

- * accepting - for jobs in ACCEPTED state
 - * finished - for jobs in FINISHED and DELETED states
 - * processing - for other states
 - * restarting - temporary location for jobs being restarted on user request or after restart of A-REX
- Other per-job files are stored in hierarchy of subdirectories jobs/SUBID/SUBID/SUBID/SUBID. Here SUBID are arbitrary ASCII string which if put together form ID of the job. The files are stored inside set of sub-directories to reduce load on filesystems which typically suffer performance decrease when amount of files in directory increases. Currently each SUBID consists of 3 lowercase hex symbols. But that may change. Following files are stored there:
 - description – file contains the description of the job (JD).
 - local – information about the job used by the A-REX. It consists of lines of format “name = value”. Not all of them are always available. The following names are defined:
 - * globalid – job identifier as seen by user tools. Depending on used interface it is either BES ActivityIdentifier XML tree, GUID of EMI ES or GridFTP URL.
 - * headnode – URL of service interface used to submit this job.
 - * interface – name of interface used for jobs submission - org.nordugrid.arcrest, org.ogf.glue.emies.activitycreation or org.nordugrid.gridftpjob.
 - * lrms – name of the LRMS backend to be used for local submission
 - * queue – name of the queue to run the job at
 - * localid – job id in LRMS (appears only after the job reached state InLRMS)
 - * args – main executable name followed by a list of command-line arguments
 - * argscode – code which main executable returns in case of success
 - * pre – executable name followed by a list of command-line arguments for executable to run before main executable. There maybe few of them
 - * precode – code which pre-executable returns in case of success
 - * post – executable name followed by a list of command-line arguments for executable to run after main executable. There maybe few of them
 - * postcode – code which post-executable returns in case of success
 - * subject – user certificate’s subject, also known as the distinguished name (DN)
 - * starttime – GMT time when the job was accepted represented in the Generalized Time format of LDAP
 - * lifetime – time period to preserve the SD after the job has finished in seconds
 - * notify – email addresses and flags to send mail to about the job specified status changes
 - * processtime – GMT time when to start processing the job in Generalized Time format
 - * exectime – GMT time when to start job execution in Generalized Time format
 - * clientname – name (as provided by the user interface) and IP address:port of the submitting client machine

- * clientsoftware – version of software used to submit the job
 - * rerun – number of retries left to rerun the job
 - * priority – data staging priority (1 - 100)
 - * downloads – number of files to download into the SD before execution
 - * uploads – number of files to upload from the SD after execution
 - * jobname – name of the job as supplied by the user
 - * projectname – name of the project as supplied by the user. There may be few of them
 - * jobreport – URL of a user requested accounting service. The A-REX will also send job records to this service in addition to the default accounting service configured in the configuration. There may be few of them
 - * cleantime – GMT time when the job should be removed from the cluster and it's SD deleted in Generalized Time format
 - * expiretime – GMT time when the credentials delegated to the job expire in Generalized Time format
 - * gmlog – directory name which holds files containing information about the job when accessed through GridFTP interface
 - * sessiondir – the job's SD
 - * failedstate – state in which job failed (available only if it is possible to restart the job)
 - * failedcause – contains internal for jobs failed because of processing error and client if client requested job cancellation.
 - * credentialserver – URL of MyProxy server to use for renewing credentials.
 - * freestagein – yes if client is allowed to stage-in any file
 - * activityid – Job-id of previous job in case the job has been resubmitted or migrated. This value can appear multiple times if a job has been resubmitted or migrate more than once.
 - * migrateactivityid –
 - * forcemigration – This boolean is only used for migration of jobs. It determines whether the job should persist if the termination of the previous job fails.
 - * transfershare – name of share used in Preparing and Finishing states. This file is filled partially during job submission and fully when the job moves from the Accepted to the Preparing state.
- input – list of input files. Each line contains 3 values separated by a space. First value contains name of the file relative to the SD. Second value is a URL or a file description. Example:

input.dat gsiftp://grid.domain.org/dir/input 12378.dat

A URL represents a location from which a file can be downloaded. Each URL can contain additional options.

A file description refers to a file uploaded from the UI and consists of [size][.checksum] where

- * size - size of the file in bytes.
- * checksum - checksum of the file identical to the one produced by cksum (1).

These values are used to verify the transfer of the uploaded file. Both size and checksum can be left out. A special kind of file description . is used to specify files which are not required to exist. The third optional value is path to delegated credentials to be used for communication with remote server.

This file is used by the data staging subsystem of the A-REX. Files with URL will be downloaded to the SD or cache and files with 'file description' will simply be checked to exist. Each time a new valid file appears in the SD it is removed from the list and input file is updated.

- input_status – contains list of files uploaded by client to the SD.

- output – list of output files. Each line contains 1, 2 or 3 values separated by a space. First value is the name of the file relative to the SD. The second value, if present, is a URL. Supported URLs are the same as those supported by input file. Optional 3rd value is path to delegated credentials to be used while accessing remote server.

This file is used by the data staging subsystem of the A-REX. Files with URL will be uploaded to SE and remaining files will be left in the SD. Each time a file is uploaded it is removed from the list and output file is updated. Files not mentioned as output files are removed from the SD at the beginning of the Finishing state.

- output_status – list of output files successfully pushed to remote locations.
- failed – the existence of this file marks the failure of the job. It can also contain one or more lines of text describing the reason of failure. Failure includes the return code different from zero of the job itself.
- errors – this file contains the output produced by external utilities like data staging, script for job submission to LRMS, etc on their stderr handle. Those are not necessarily errors, but can be just useful information about actions taken during the job processing. In case of problem include content of that file while asking for help.
- diag – information about resources used during execution of job and other information suitable for diagnostics and statistics. It's format is similar to that of local file. The following names are at least defined:
 - * nodename – name of computing node which was used to execute job,
 - * runtimeenvironments – used runtime environments separated by ';',
 - * exitcode – numerical exit code of job,
 - * frontend distribution – name and version of operating system distribution on frontend computer,
 - * frontend system – name of operating on frontend computer,
 - * frontend subject – subject (DN) of certificate representing frontend computer,
 - * frontend ca – subject (DN) of issuer of certificate representing frontend computer, and other information provided by GNU time utility. Note that some implementations of time insert unrequested information in their output. Hence some lines can have broken format.
- proxy – delegated X509 credentials or chain of public certificates.
- proxy.tmp – temporary X509 credentials with different UNIX ownership used by processes run with effective user id different from job owner's id.
- statistics – statistics on input and output data transfer
- xml - job's current state expressed in GLUE2 XML rendering

There may be other files inside jobs/SUBID/SUBID/SUBID/SUBID sub-directories which are created and used by different parts of the A-REX. Their presence can not be guaranteed and can change depending on changes in the A-REX code.

- delegations – sub-directory containing collection of delegated credentials and sqlite database for associating them to submitted jobs.

- list - sqlite database which stores information about available delegations and their links to jobs.

Delegations are stored in structure of sub-directories similar to one used for jobs with delegation ids used to generate names of sub-directories. But last item in hierarchy is not a directory. It's a file containing delegation content. It's either X.509 proxy or OIDC token.

- accounting - sub-directory containing following accounting related information:
 - accounting.db - sqlite database with accounting records for each job processed and related information

7.7 ARC support for OIDC

7.7.1 Support level

Only tokens conforming to the WLCG profile have been tested. The current validation is not strict: the token is parsed and the signature is checked if present, but no additional requirements are imposed. Tokens are only accepted for client authentication during job submission through the REST interface.

7.7.2 Obtaining and using tokens

The suggested way for obtaining a token is through the oidc-agent utility - <https://indigo-dc.gitbook.io/oidc-agent/>. Install it following the instructions for your Linux distribution.

Point your browser to <https://wlcg.cloud.cnaf.infn.it/> and create an account.

Start the oidc-agent. It will print few lines of shell commands. Copy them to the command line and execute. This will set up environment variables for other oidc-* commands.

Start oidc-gen. It will guide you through the steps to register the OIDC client and create a profile for the oidc-agent. When asked to select the issuer, select <https://wlcg.cloud.cnaf.infn.it/>. When asked about the scope, write `openid profile wlcg wlcg.groups`. You only need to run the oidc-gen once. Next time when you use the oidc-agent, you can load an already created profile with the 'oidc-add NAME_YOU_CHOOSE' command.

When the oidc-gen command tells you "To continue and approve the registered client visit the following URL in a Browser of your choice:" point your browser (which must be running on the same machine as the oidc tools do) to the address given.

Obtain the token and store it in the BEARER_TOKEN variable:

```
export BEARER_TOKEN=`oidc-token NAME_YOU_CHOOSE`
```

If the oidc tools were installed on a different machine than your ARC client, then first obtain the token on the oidc tool machine:

```
oidc-token NAME_YOU_CHOOSE
```

Then copy the token string and on your ARC client machine do:

```
export BEARER_TOKEN=<token-string>
```

Now submit the job to an ARC CE with arcsub through the REST interface. For that, use the option '-T arcrest'. The token stored in the BEARER_TOKEN variable will be used instead of the X.509 certificate for authenticating the user to the ARC CE server. Note that arcsub still requires an X.509 proxy which will be delegated to the CE. Data staging currently will also still use the X.509 proxy credentials.

Note: You can use any other method for obtaining a WLCG compliant OIDC token. Just store it in the BEARER_TOKEN variable before calling arcsub.

7.7.3 Configuring authorization on server

Token processing is enabled by the presence of the [authtokens] configuration block.

The user can be authorized on the server by adding a dedicated command to the authgroup block:

```
authtokens=subject issuer audience scope group
```

The specified parameters must match those in the provided token. Parameters can be '*' to match any value, for example

```
authtokens=e83eec5a-e2e3-43c6-bb67-df8f5ec3e8d0 https://wlcg.cloud.cnaf.infn.it/ * * *
```

matches a user with subject e83eec5a-e2e3-43c6-bb67-df8f5ec3e8d0 in token issued by <https://wlcg.cloud.cnaf.infn.it/>.

Note

Until handling of authtokens is integrated with arcproxy you will have to find the subject of the token using a tool like e.g. <https://jwt.io/>. Alternatively you can install `faat` (<https://pypi.org/project/faat/>).

A full example configuration could look like

```
[authtokens]

[authgroup: wlcg_iam]
authtokens = * https://wlcg.cloud.cnaf.infn.it/ * compute.create /wlcg/pilots
authtokens = * https://wlcg.cloud.cnaf.infn.it/ * compute.read /wlcg/pilots
authtokens = * https://wlcg.cloud.cnaf.infn.it/ * compute.modify /wlcg/pilots
authtokens = * https://wlcg.cloud.cnaf.infn.it/ * compute.cancel /wlcg/pilots

[mapping]
map_to_user = wlcg_iam wlcg:wlcg
policy_on_nomap=stop

[arex/ws/jobs]
allowaccess=wlcg_iam
```

User mapping to a local account is implemented using a simulated X.509 user subject. The subject provided by an OIDC token is unique only in scope of the identity provider. To generate a globally unique user-identifier, the issuer and the subject are concatenated as “issuer/subject” to provide an identifier suitable for user mapping. For example, a user with subject e83eec5a-e2e3-43c6-bb67-df8f5ec3e8d0 in the token issued by <https://wlcg.cloud.cnaf.infn.it/> is represented by a simulated identifier <https://wlcg.cloud.cnaf.infn.it/e83eec5a-e2e3-43c6-bb67-df8f5ec3e8d0>

7.8 ARC Information System Technical Details

General Information System configuration items are described in the configuration items section *[infosys] block*

This page contains a description of the architecture of the ARC Information System.

7.8.1 Overview: Purpose of the Information System

The ARC Information System (short name: “infosys”) is a collection of scripts that generates information documents about the status of the A-REX service, the LRMS and the submitted jobs.

These documents can be retrieved using the *ARC CE REST interface* formatted as XML or JSON.

A group of legacy components provide the same information as LDIF documents via a LDAP server .

The information is presented according to two schemas, the default *GLUE2 schema* and the legacy *Nordugrid schema*.

Some of the content of the documents is static but there is also a dynamic component to provide minimal statistics and job progress.

It is important to understand the asynchronoyus nature of the content of these documents, as this is NOT a real time system. Every rendered item contains a validity field so that information consumers (clients) can determine how fresh is the presented information. See section *Static and dynamic info* for more details.

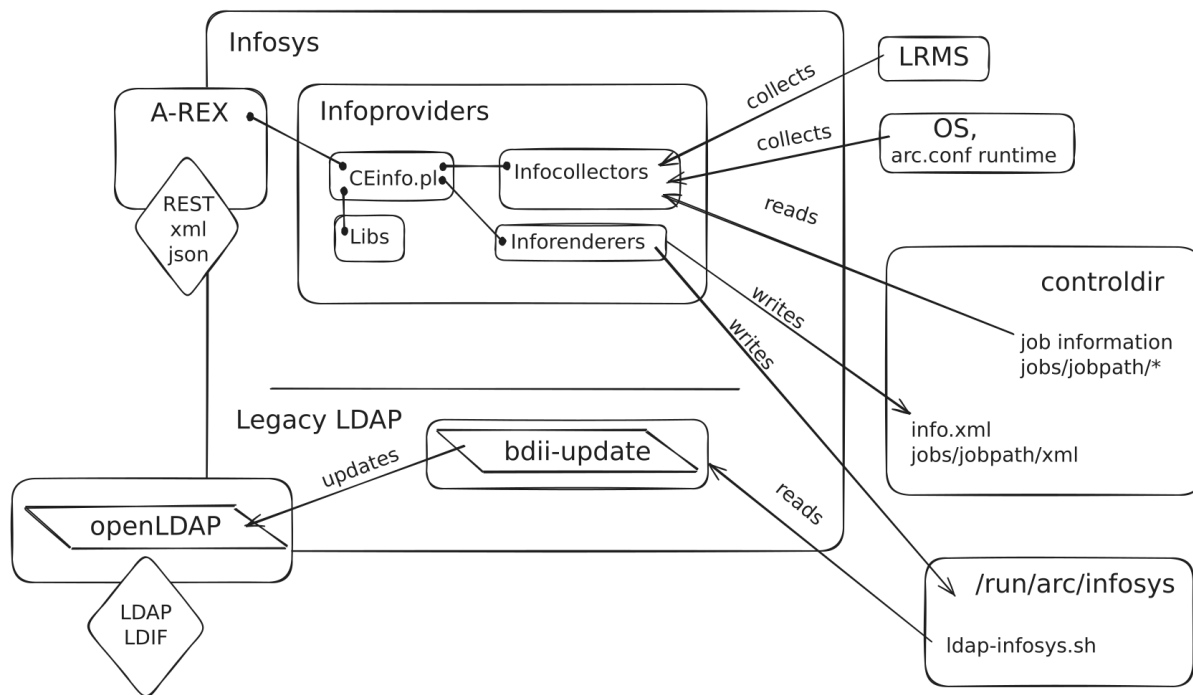


Fig. 7.16: Components and workflow of the ARC Information System

7.8.2 Inputs and Outputs

The main **inputs** to the information system are:

- the runtime `arc.conf` for most of the static data;
- the content of the control directory where both A-REX and the LRMS scan scripts write information about the status of the jobs;
- direct queries to the LRMS to try to get the most up-to-date information from it
- The output of commands used to collect info about load, storage and other system related matters.

The **outputs** of each “run” or “scan” is mainly three types of files:

- `info.xml`, that contains static data about the Computing Element and aggregated dynamic statistics about the status of the LRMS and the jobs, generated in the control directory. The output is formatted according to the [NorduGRID implementation of the GLUE2 schema](#).
- For each job, a file containing information about that job is generated in a `<controldir>/jobs/jobpath/xml` file. The output is formatted according to the [NorduGRID implementation of the GLUE2 schema](#) for *ComputingActivity* objects.
- Optionally, a BASH script `ldap-infosys.sh` is created when LDAP information is required, generated inside `/run/arc/infosys`. The script contains an LDIF representation of the NorduGRID and GLUE2 schemas.

7.8.3 Description of the components

This section described the functionalities and purposes of the Information System components. It is not meant to describe the code in detail but to give an understanding of the architecture.

The Infoproviders

The infoproviders are script or executables that output a formatted document containing information about the status of the ARC computing element. Currently there are two information providers, one of which is generated by the other for historical reasons:

- `CEinfo.pl`, a PERL script responsible for parsing the configuration, collecting all information from different data sources, aggregating it into proper data structures, writing it out in different formats, but mainly generating the `info.xml` documents and an `xml` document for each job. It also generates the legacy information provider `ldif-provider.sh` below when `[infosys/ldap] block` and `[infosys/nordugrid] block` are present in configuration.
- `ldif-provider.sh` is a bash script generated by `CEinfo.pl` and it's a legacy functionality for backward compatibility of previous ARC versions that outputs information about the cluster as a complete LDIF document representing the cluster information for both NorduGRID and GLUE2 schemas. GLUE2 information is only rendered if the `[infosys/glue2/ldap] block` is present in the configuration. For a more detailed description of the LDAP backend see *The legacy LDAP subsystem*.

The information collectors

The information collectors are a set of PERL modules responsible to collect and aggregate information of various kinds.

Their purposes are described in the table below.

Table 7.5: Information collectors :widths: 25 75 :header-rows: 1

Module name	Purpose
<code><LRMS>mod pm</code>	Modules responsible to query directly the named <code><LRMS></code> to obtain fresher information about queues/partitions and jobs. The <code>mod</code> prefix identifies the latest version of these modules that can automatically collect information about GLUE2 ExecutionEnvironments, that is, the hardware of nodes of each partition in the cluster.
<code><LRMS>. pm</code>	Legacy LRMS modules that cannot return proper hardware information about the nodes. Since some of the LRMS backends are maintained by communities, not all communities have interest in collecting detailed hardware information, so they were never developed in that direction.
<code>HostInfo. pm</code> and <code>Sysinfo. pm</code>	Collect information about the status of the frontend and the mount points, as well as other details such as hostname, load, available storage etc.
<code>RTEinfo. pm</code>	Collects information about Run Time Environments (See also <i>RunTime Environments in ARC</i>)
<code>ARC1Clust pm</code>	Aggregates information from all sources into a datastructure that can be used to generate GLUE2 documents.
<code>ARC0Clust pm</code>	Aggregates information from all sources into a datastructure that can be used to generate LDIF documents. Used for legacy LDAP information rendering.

The information renderers

Information renderers are responsible for transforming the datastructures generated by infocollectors into documents in various formats.

Below a description of the modules and their purposes.

Table 7.6: Information renderers

Module name	Purpose
<code>XmlPrinter.pm</code>	Generic library to render XML documents.
<code>GLUE2xmlPrinter.pm</code>	Specialized library to render XML documents according to the GLUE2 schema.
<code>LdifPrinter.pm</code>	Generic library to render LDIF documents.
<code>NgldifPrinter.pm</code>	Specialized library to render LDIF documents according to the NorduGRID schema.
<code>GLUE2ldifPrinter.pm</code>	Specialized library to render LDIF documents according to the GLUE2 schema.

Helper libraries

The infoproviders make use of intermediate libraries to perform various tasks required for collection, rendering and other support features. In the following table their description.

Table 7.7: Libraries of the information system

Module name	Purpose
<code>LogUtils.pm</code>	Logging facility to keep consistency across logging in various infoprovider modules.
<code>IniParser.pm</code>	Library to parse ARC configuration INI-formatted files.
<code>ConfigCentral.pm</code>	Main library for parsing configuration files and translating them into datastructures usable by the information system scripts.
<code>InfoChecker.pm</code>	Lint library to cleanup configuration entries that are not used by the information system. It also does some format correction and consistency checks.
<code>LRMSInfo.pm</code>	Interface to any kind of LRMS information collector. It takes care of selecting the proper LRMS information collector, and add compatibility fixes for legacy collectors.
<code>ARC0mod.pm</code>	Compatibility layer to LRMS information collectors to add hardware datastructures to the legacy modules, for consistency.
<code>InfosysHelper.pm</code>	Helper library to bridge A-REX infocollection and LDAP updates. It is used to synchronize the collection to the update of the database to avoid inconsistencies between the XML and LDIF documents.

The legacy LDAP subsystem

The legacy LDAP subsystem provides information about the cluster using the LDAP protocol in the form of LDIF files.

It supports two main custom GRID schemas:

- the NorduGRID schema http://www.nordugrid.org/documents/arc_infosys.pdf
- the GLUE2 schema <https://ogf.org/documents/GFD.218.pdf>.

It is constituted by 4 main elements:

- A LDAP server, *openLDAP* (*slapd*), that serves information on port 2135 to anonymous users.
- A third party Python script *bdi i -update* that periodically executes *ldapadd/modify/delete* against the LDAP database to keep the information up to date.

- A generated `ldif-provider.sh` script that outputs a LDIF document containing all the information about the cluster, used by `bdii-update`
- A set of startup scripts that configure the LDAP database so that it can work with legacy GRID information from ARC.

This subsystem is disabled by default in ARC7. It can be enabled in configuration by adding the following blocks:

- `[infosys/ldap] block` - configures the LDAP subsystem
- `[infosys/glue2/ldap] block` - enables LDAP GLUE2
- `[infosys/nordugrid] block` - enables classic NorduGRID schema

7.8.4 Static and dynamic info

It is important to understand the nature of the contents of the documents in the information system. It is not a real-time system: the information is collected asynchronously and serially from various sources and then published with a *validity timestamp* so that the clients can decide what to do with such info.

Most of the static information is taken from configuration files of various subsystems. The dynamic information updated regularly concerns storage availability, system load, status of jobs, statistics about jobs currently managed by the system.

The information system shows no historical records of any object, nor keeps a log of what happened before the latest scan.

Although the generated data is used by A-REX in some REST responses, A-REX has some in-memory information that is fresher than the one presented by the information system. See more details below.

A-REX itself has no means to check what is the state of a job inside the LRMS, therefore the LRMS backends and the information system are the only sources of information about what is happening with jobs inside the cluster batch system.

A-REX in-memory information

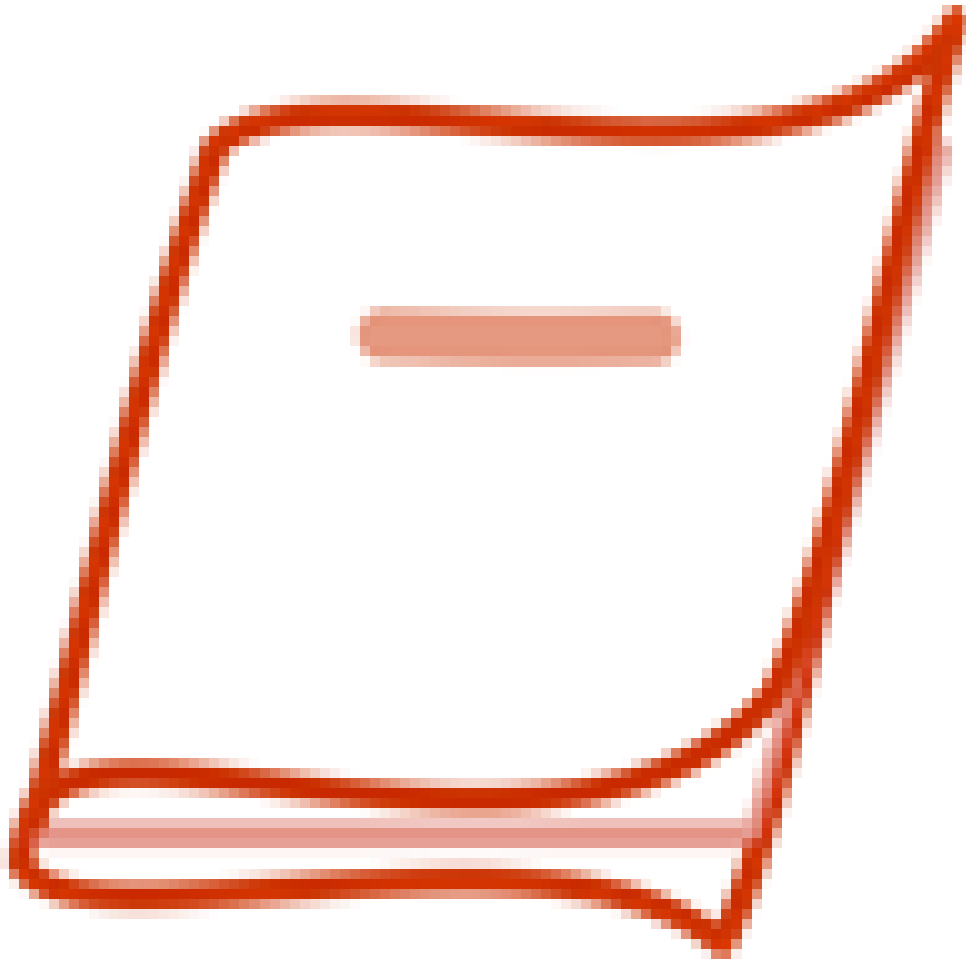
A-REX keeps in-memory coarse-grained information about important changes in the job states. The *job status* retrievable using the [REST API](#) is more reliable than the job state contained/generated by the infosys and therefore for tracking a job it is more reliable to retrieve this information via REST instead of LDAP.

7.9 Old Technical Documents

Note

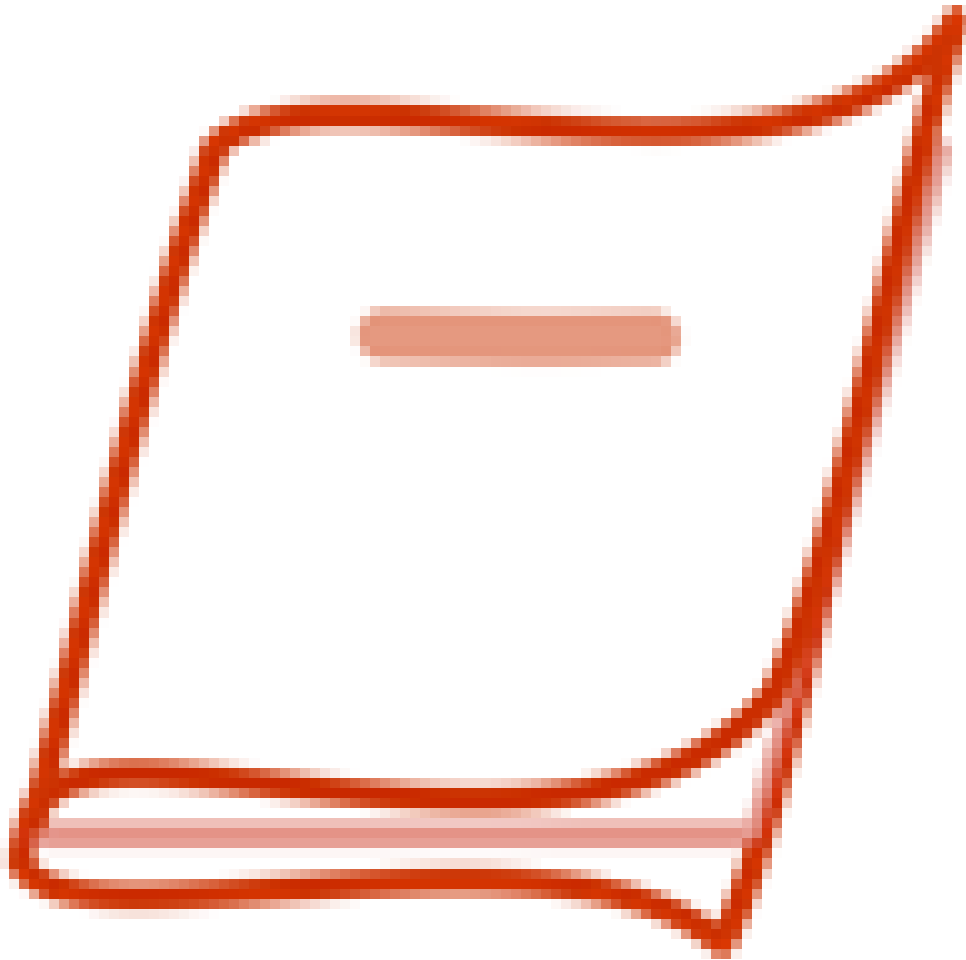
Many of the technical documents [exists](#) for ARC5 only. Those that are verified to be relevant for ARC6 will be listed below. For ARC 7 some parts of the documentation may still be relevant, but can not be guaranteed.

7.9.1 Hosting Environment of the Advanced Resource Connector middleware



Document gives a deep technical description of the HED service container.

7.9.2 A Client Library for ARC



Document describes from a technical viewpoint the plugin-based client library of ARC.

7.10 Legacy JURA Accounting Technical Details

WARNING: This component was deprecated in ARC 6.4 and completely removed in ARC 6.8!

Warning

Information in **this chapter is relevant only for 6.0-6.3 ARC releases.**

Starting from ARC 6.4 release the *next generation accounting subsystem* with local accounting database will be introduced. Make sure you are reading the documentation that match your ARC CE release version.

General accounting configuration and operations flows are described in *Accounting with legacy JURA*. This section contains more technical details about implementation of each component of accounting subsystem.

7.10.1 Records processing and publishing

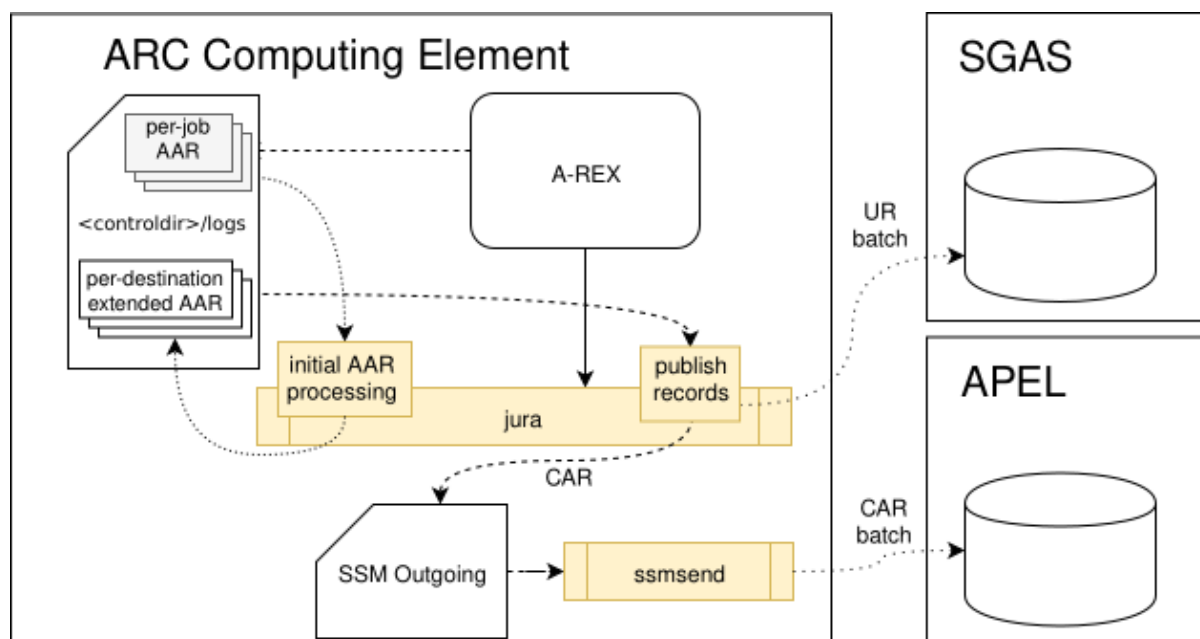


Fig. 7.17: ARC CE accounting: records creation, processing and publishing

AREX Accounting Records (AAR job log files)

The A-REX Accounting Records (AAR) are job log files generated by A-REX. AAR is the only source of accounting information for JURA. AARs are written by A-REX based on job data available, including `.diag` files that backend scripts creates based on batch system data and/or GNU `time` utility measurements. These job log files reside under the `<control_dir>/logs` directory. The name of the AAR job log files consist of the ID of the job and a random string to avoid collision of multiple job log files for the same job: `<jobid>.<random>`.

The AAR job log file consists of `name=value` lines, where `value` is either a job-related resource consumption data or a static info like name, ID or proxy certificate.

A-REX generates at least two job log files for each job: one at the time of job submission, another one after the job finishes, and possibly others at various job events. Please note JURA makes use only one of the AREX generated files belonging to the same job: the one that corresponds to the `FINISHED` job event (such state is indicated by the `status={completed|failed|aborted}`).

JURA initial AAR processing

A-REX periodically runs `jura` that loop over available A-REX job log records in the `<control_dir>/logs`. JURA opens all the files and processes only those that corresponds to a `FINISHED` job state.

JURA converts AARs to *per-job per-destination* extended AARs that contains the target information from `arc.conf` as well. One extended AAR is generated per accounting target. The extended AAR job log files named `<jobid>.<random>_<random2>` where first `<random>` is taken from original AAR.

The original AAR `<jobid>.<random>` file is deleted once *per-destination* extended AAR logs are created by JURA.

Note

JURA as part of the initial processing deletes all files corresponding to non-finished job states.

JURA Publishing loop

JURA publishing subsystem loop over *extended AAR* logs in the A-REX `<control_dir>/logs` directory. JURA generates records in the Usage Record (UR) format proposed by the Open Grid Forum (OGF) for SGAS or Compute Accounting Record (CAR) XML for APEL.

The *extended AAR* job log file is deleted once record is successfully submitted, thus preventing multiple insertion of same usage records. If submission to destination fails, the extended AAR log files are kept, so another attempt is made upon a subsequent run of JURA. This mechanism will be repeated until the expiration time passes at which point the next execution of JURA removes the file without processing.

Please note that the JURA publishing loop is backward compatible with ARC 5 implementation.

Reporting to SGAS

SGAS has a simple custom web service interface loosely based on WS-ResourceProperties. JURA uses the insertion method of this interface to report URs directly using ARC HTTP client implementation. The corresponding processed extended AAR job log files are deleted after receiving a non-fault response from the service.

To increase communication efficiency JURA sends URs in batches. SGAS accepts a batch of URs in a single request. The batch is an XML element called `UsageRecords`, containing elements representing URs.

The process of handling batches is the following: JURA does not send all usage records immediately after generation, but instead collects them in a batch until reaching the maximal number of records or until running out of job log files. The maximal number of URs in a batch can be set as a *urbatchsize* configuration parameter of SGAS target.

Reporting to APEL

APEL uses the SSM framework for communication.

JURA send records to APEL by means of invoking helper `ssmsend` process that uses SSM python librairies developed by APEL.

ARC ships minimal set of SSM libraries along with A-REX binary packages to allow SSM usage. If SSM binary packages from APEL are available for your OS (e.g. EL6), you can install this packages and they will be used instead of those shipped with ARC automatically.

JURA prepares the messages to be sent by `ssmsend` and puts them into *SSM Outgoing* directory located in the `/var/spool/arc/ssm/<destination hostname>/outgoing/00000000/`. Generated messages are XML based CAR records with file name format `<YYYYMMDDhhmmss>`.

Reporting to APEL also works with sending records in batches. The default *urbatchsize* value is set to 1000 according to APEL recommendations.

7.10.2 Accounting archive

After records are archived by `jura` a dedicated `jura-archive-manager` process manages the archive layout and archive database.

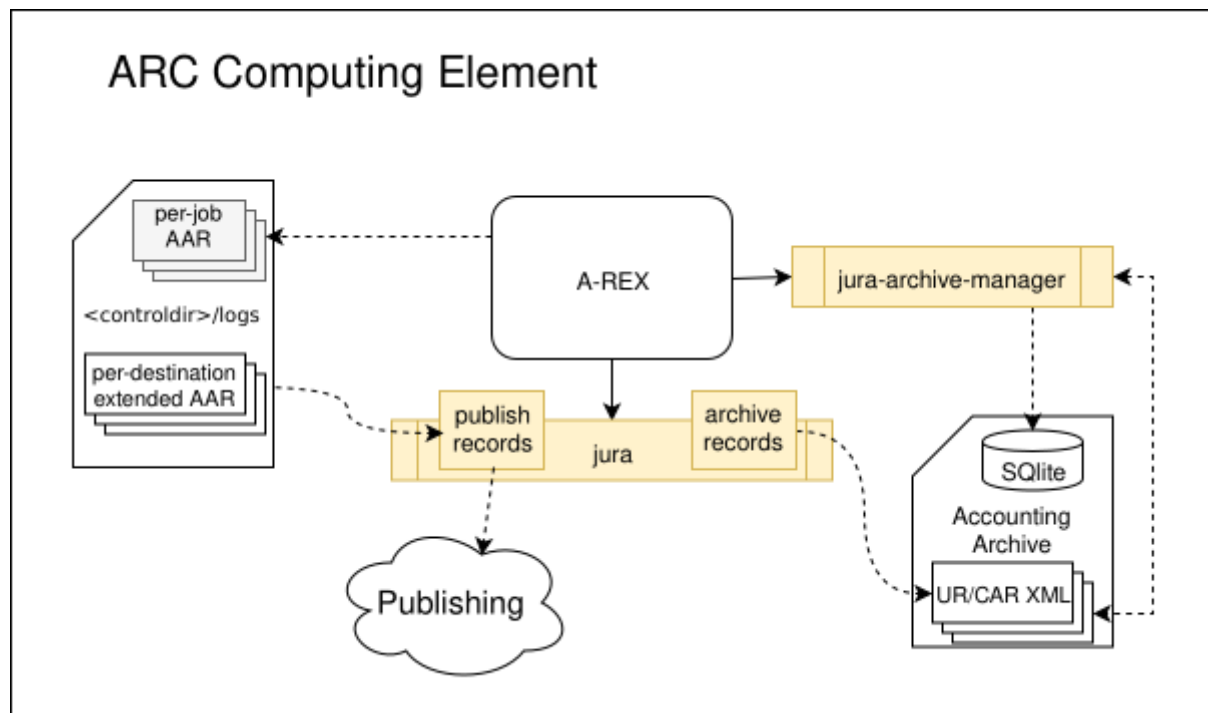


Fig. 7.18: ARC CE accounting: records archiving

7.10.3 Republishing process

Republishing is a complex workflow that involves:

- archive database to find out which subset of records should be republished
- exporting necessary records from archive structure to temporary flat directory
- triggering jura with corresponding configuration to grab records from temporary directory and send them to target
- cleanup after republishing

This chain is controlled from `arcctl` code.

7.10.4 Security

The JURA executable runs with the same user privileges as the A-REX. The owner of a job log file is the local user mapped for the submitter entity of the corresponding job. Since these files contain confidential data, A-REX restricts access to them allowing only read access for the job owner, thus when JURA is executed by A-REX it is allowed to read and delete job log files.

All usage records are submitted using the X.509 credentials specified by the value of `x509_` set of configuration options of `arc.conf`. No proxies are used for communication with accounting services.

The only access restriction made by a SGAS service is matching the Distinguished Name of the client (in this context JURA) with a set of trusted DNs. When access is granted, policies are then applied by SGAS, allowing either publishing and/or querying rights. Clients with publishing right can insert any UR, regardless of content. By default, querying right only allows retrieving URs pertaining to jobs submitted by the querying entity.

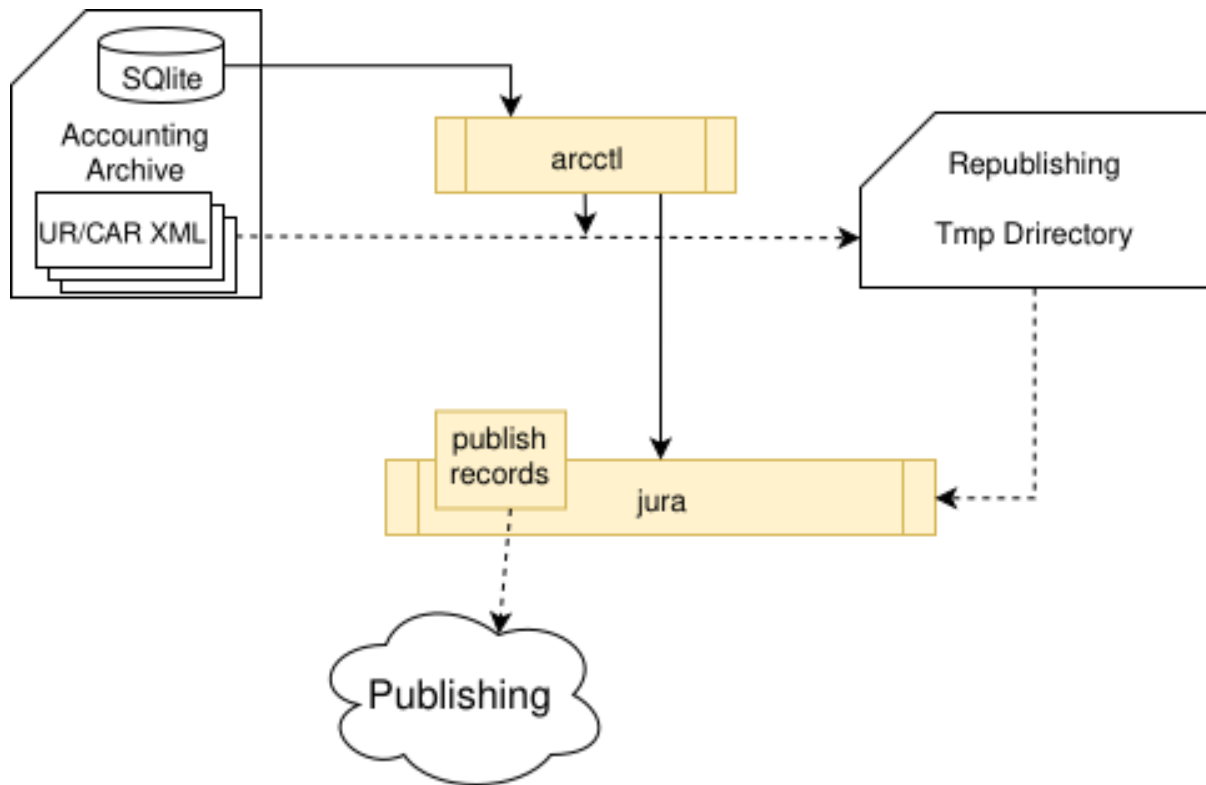


Fig. 7.19: ARC CE accounting: records republishing

7.10.5 Implementation and API

JURA as part of the ARC software stack is written in C++, and utilizes the functionality provided by the ARC libraries, including secure HTTPS communication provided by the ARC pluggable TLS and HTTP modules.

The modular design is also present in the usage reporting part of the JURA code, making it possible to extend JURAs support of accounting services. To create a JURA module one should simply write a C++ class which inherits from the abstract `Arc::Destination` class, and it must extend the two methods:

- `static Arc::Destination* Arc::Destination::createDestination(Arc::JobLogFile&)`
- `void Arc::Destination::report(Arc::JobLogFile&)`

The static `createDestination` method should initialize a object of the specialized class, using the configuration options specified in the passed `Arc::LogFile` object, and the memory allocated by the method should be freed by the caller. Then the `report` method should carry out the transfer of the UR, represented by the `JobLogfile` object, to the accounting service.

JURA archive manager is written in Python and share the classes for managing ARC components with *ARC Control Tool*.

7.10.6 Limitations

In the following list some issues which limits the functionality of JURA is described:

- The current implementation of JURA and A-REX supports only one expiration time for all the reporting destinations. Even though the configuration enables the specification of different expiration values per reporting destination, it is not taken into account by the system, the last value is used as the common expiration time value.
- It is not possible to use different credentials per destinations.
- If you are updating from ARC5 with an old jura accounting archive already containing records, the conversion process to index the archive structure will be initiated and will cause serious system load until finished. To avoid old archive conversion, you can move records before update.
- Some *optional UR properties* are not supported.
- Memory can be reported incorrectly with buggy GNU “time” results.

7.11 ARC Accounting Database Schema

DOCUMENTATION FOR DEVELOPERS

If you are looking for development-internal details of ARC (like how some stuff was coded) this part of documentation is for you. Mainly for those who want to contribute to the project development, advanced troubleshooters or just interested.

8.1 Implementation Details for Developers

8.1.1 General `arc.conf` python configuration parser

Parsing configuration

Initial configuration parsing

The latest version of `arcconfig-parser` is designed to operate with the defaults file that holds the default values for all possible parameters.

You find the `arcconfig-parser` in your `$ARC_LOCATION/libexec/arc` folder which depends on your OS and installation options, but an example is `/usr/libexec/arc/`.

At the time of initial parsing the following chain is executed:

- All blocks and options are parsed from `arc.conf`
- All blocks and options are parsed from `arc.conf.d` and merged with `arc.conf`
- For all blocks defined in the merged `arc.conf`, missing options are added from the defaults file
- Special constructs in values are substituted (see *Special constructs can be used as values*)

Parameters that are optional by design and do not have a default value (specified with `not set value`) are not included to the parsed configuration.

Runtime configuration

Configuration that includes both `arc.conf` and defaults config is called the *runtime configuration*.

In some cases it is useful to save and load the runtime configuration:

- To supply C++ services (`a-rex`, `gridftpd`) with configuration that includes defaults from common place
- For repetitive operations on config to eliminate full-chain processing of each invocation

To save the runtime configuration to the default location (`/var/run/arc/`):

```
arcconfig-parser --save
```

To save the runtime configuration to a specified location:

```
arcconfig-parser --save -r /var/run/arc/arex.arc.conf
```

To load the runtime configuration instead of the full-chain processing and e.g. get the value of the `x509_host_key` in the `[common]` block:

```
arcconfig-parser --load -b common -o x509_host_key
```

Special constructs can be used as values

When defaults include references to another config part to be consistent with the implied `arc.conf` structure.

The most obvious example is if `x509_host_key` is not found in e.g. the `[arex/jura]` block, then it should be taken from the `[common]` block.

The config parser is following this logic (especially in respect to defaults) and use special constructs to accomplish this behaviour.

Command substitutions

Configuration option values can contain the construct `$EXEC{<command>}` that substituted to the stdout of `<command>`.

For example:

```
hostname=$EXEC{hostname -f}
```

Option values substitutions

The construct `$VAR{[block]option}` can be used to substitute the values of another option value.

If the option is in the same block as the referencing option, the block name can be omitted - `$VAR{option}`.

For example:

```
x509_host_key=$VAR{[common]x509_host_key}
bdii_update_cmd=$VAR{bdii_location}/sbin/bdii-update
```

Evaluation of simple code

For limited number of cases the `arc.conf` default values relies on arithmetic operations. For this purpose the `$EVAL{string}` special construct has been introduced.

For example:

```
bdii_read_timeout=$EVAL{$VAR{bdii_provider_timeout} + $VAR{[arex]infoproviders_
↪timelimit} + $VAR{[arex]wakeupperiod}}
```

Getting the configuration values

If the `--option` argument is passed to `arcconfig-parser`, it returns the value of the specified option to stdout.

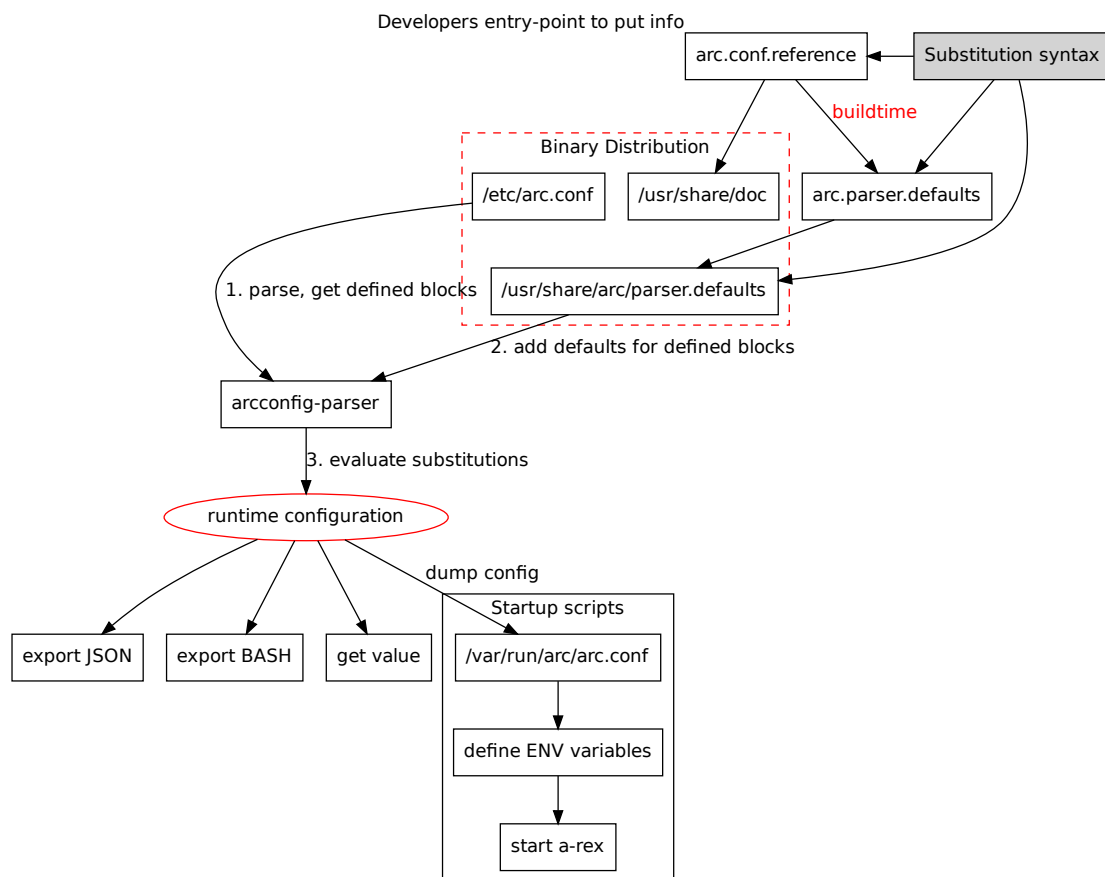
Without `--option` `arcconfig-parser` can be used to operate with configuration blocks:

- check blocks existence (exit code used to indicate the status of the check)
- return the list of subblocks

With the `--export` option `arcconfig-parser` allows to export config in the following formats:

- `json` - returns entire configuration or subset of blocks as-is in JSON to stdout
- `bash` - for `[common]` block or specified configuration subset returns `CONFIG_option_name=value` pairs to stdout. Block names ARE NOT included in the exports and option values precedence will be used in the order of passed blocks. If automatic subblocks expansion used with bash export, for every block in sequence - it's subblocks are processed first (in `arc.conf` defined order). It is possible to filter the options that will be exported with additional `--filter` option that can be specified several times.

Common configuration parsing sequence



Options reference

The full list of available options of `arconfig-parser` is available instantly using `--help` option. An online version is available [here](#).

Examples

Get value of an option in a block:

```
# arconfig-parser --block infosys --option logfile
/var/log/arc/infoprovider.log
```

Get value of an option in blocks in order they are specified (e.g. if not found in `[gridftpd]` look in the `[common]` block¹):

```
# arconfig-parser --block gridftpd --block common --option x509_host_key
/etc/grid-security/hostkey.pem
```

Export entire configuration to JSON²:

```
# arconfig-parser --export json
```

Export `[infosys]` block options to JSON (for Perl):

```
# arconfig-parser --block infosys --export json
{"infosys": {"loglevel": "5"},...
```

Export `[infosys]` block and all their subblocks options to JSON:

```
# arconfig-parser --block infosys --subblocks --export json
{"infosys/glue2/ldap": {"showactivities": "no"},...
```

Export for BASH (compatible with config representation in shell-based LRMS backends):

```
# arconfig-parser --block infosys --block arex --block common --export bash
CONFIG_controlldir="/var/spool/arc/jobstatus"
CONFIG_defaultttl="1210000"
CONFIG_delegationdb="sqlite"
CONFIG_hostname="ce01.example.org"
CONFIG_maaxrerun="5"
CONFIG_maxjobs="10000 -1"
CONFIG_runtimedir="/home/grid/arc/runtime"
CONFIG_sessiondir="__array__" # <= NEW define for multivalued values that indicate
↪ indexed vars
CONFIG_sessiondir_0="/mnt/scratch/grid/arc/session"
CONFIG_sessiondir_1="/home/grid/arc/session drain"
...
```

Export for BASH with exported options filtering:

```
# arconfig-parser -b common -f hostname -f x509_cert_dir -e bash
CONFIG_hostname="ce01.example.org"
CONFIG_x509_cert_dir="/etc/grid-security/certificates"
```

Using BASH export:

¹ Block dependencies are now implied by the defaults file, so for most cases it is enough to specify only the block in question

² *HINT:* use `arconfig-parser --export json | jq .` to view highlighted JSON structure in shell

```
# eval "$( arcconfig-parser --block infosys --block arex --block common --export bash_
↪)"
# echo "$CONFIG_gridmap"
```

Check block(s) exists ([common/perflog] does not exists in the example):

```
# arcconfig-parser --block common/perflog --block arex
# echo $?
1
```

List block subblocks:

```
# arcconfig-parser --block infosys --subblocks
infosys
infosys/ldap
infosys/nordugrid
infosys/glue2
infosys/glue2/ldap
infosys/cluster
```

Using the parser as a Python module:

```
from arc.utils import config

# initial parsing with defaults
config.parse_arc_conf('/etc/arc.conf', '/usr/share/arc/arc.parser.defaults')

# get parsed dictionary and list of blocks in the arc.conf order
>>> confdict = config.get_config_dict()
>>> confblocks = config.get_config_blocks()

# get list of all [queue] subblocks sorted by name
>>> sb = config.get_subblocks(['queue'], is_sorted=True)
>>> sb
['queue:grid', 'queue:grid_rt']

# get value of 'x509_host_key' from [arex] block and than from [common] if not found_
↪in [arex]
>>> a = config.get_value('x509_host_key', ['arex', 'common'])
>>> a
'/etc/grid-security/hostkey.pem'

# get value of 'allowactivedata' option from [gridftpd] block
>>> b = config.get_value('allowactivedata', 'gridftpd')
>>> b
'yes'

# get value of 'allowactivedata' option from [gridftpd] block (always return list)
>>> c = config.get_value('allowactivedata', 'gridftpd', force_list=True)
>>> c
['yes']

# get value of 'allowactivedata' option from [gridftpd] block (return boolean value)
>>> d = config.get_value('allowactivedata', 'gridftpd', bool_ynsno=True)
>>> d
True
```

8.1.2 LRMS shell-backends overview for developers

CONFIG variables used in LRMS shell-backend:

lrms_common.sh:

```
$CONFIG_runtimedir      [arex]
$CONFIG_shared_scratch  [arex]
$CONFIG_shared_filesystem [arex]
$CONFIG_scratchdir      [arex]
$CONFIG_gnu_time        [lrms]
$CONFIG_nodename        [lrms]
$CONFIG_enable_perflog_reporting [common]
$CONFIG_perflogdir      [common]
```

submit_common.sh:

```
$CONFIG_defaultmemory   [queue] [lrms]
$CONFIG_hostname        [common]
$CONFIG_controldir      [arex]
```

lrms=condor¹:

```
# $CONFIG_enable_perflog_reporting [common] not in reference
# $CONFIG_perflogdir               [common] not in reference
# $CONFIG_controldir               [arex] (for perflog)

$CONFIG_condor_requirements [queue] [lrms]
$CONFIG_condor_rank         [lrms]
# $CONFIG_shared_filesystem [arex]
$CONFIG_condor_bin_path     [lrms]
$CONFIG_condor_config       [lrms]
```

lrms=slurm:

```
# $CONFIG_enable_perflog_reporting [common] not in reference
# $CONFIG_perflogdir               [common] not in reference
# $CONFIG_controldir               [arex] (for perflog)

$CONFIG_slurm_wakeupperiod [lrms]
$CONFIG_slurm_use_sacct    [lrms]
$CONFIG_slurm_bin_path     [lrms]
# $CONFIG_shared_filesystem [arex]
```

lrms=fork:

no variables

lrms=boinc:

```
$CONFIG_boinc_app_id [lrms]
$CONFIG_boinc_db_host [lrms]
$CONFIG_boinc_db_port [lrms]
$CONFIG_boinc_db_user [lrms]
$CONFIG_boinc_db_pass [lrms]
$CONFIG_boinc_db_name [lrms]
```

lrms=ll:

¹ Here and following # prefix is for options and are used in *_common scripts and not unique to particular backend

```
# $CONFIG_enable_perflg_reporting [common] not in reference
# $CONFIG_perflgdir [common] not in reference
# $CONFIG_controldir [arex] (for perflg)

$CONFIG_ll_bin_path [lrms]
$CONFIG_ll_consumable_resources [lrms]
$CONFIG_ll_parallel_single_jobs *not in reference
# $CONFIG_scratchdir [arex]
```

lrms=lsf:

```
# $CONFIG_enable_perflg_reporting [common] not in reference
# $CONFIG_perflgdir [common] not in reference
# $CONFIG_controldir [arex] (for perflg)

$CONFIG_lsf_architecture [lrms]
$CONFIG_lsf_bin_path [lrms]
```

lrms=pbs:

```
# $CONFIG_enable_perflg_reporting [common] not in reference
# $CONFIG_perflgdir [common] not in reference
# $CONFIG_controldir [arex] (for perflg)

$CONFIG_pbs_queue_node [queue]
$CONFIG_pbs_bin_path [lrms]
$CONFIG_nodememory [queue] ([infosys/cluster] parser substitution,
↪ fallback only)
$CONFIG_pbs_log_path [lrms]
# $CONFIG_shared_filesystem [arex]
```

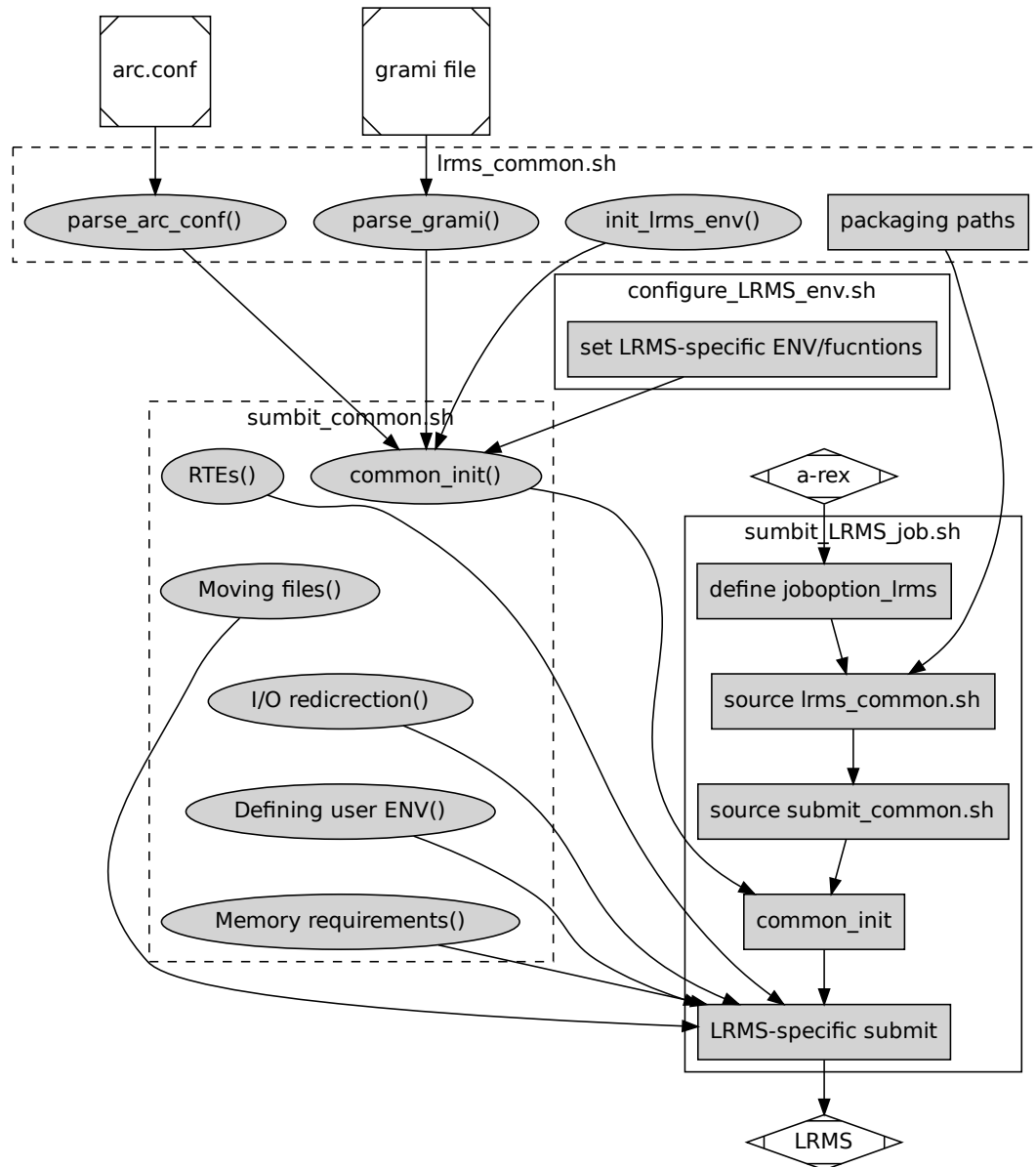
lrms=sge:

```
# $CONFIG_enable_perflg_reporting [common] not in reference
# $CONFIG_perflgdir [common] not in reference
# $CONFIG_controldir [arex] (for perflg)

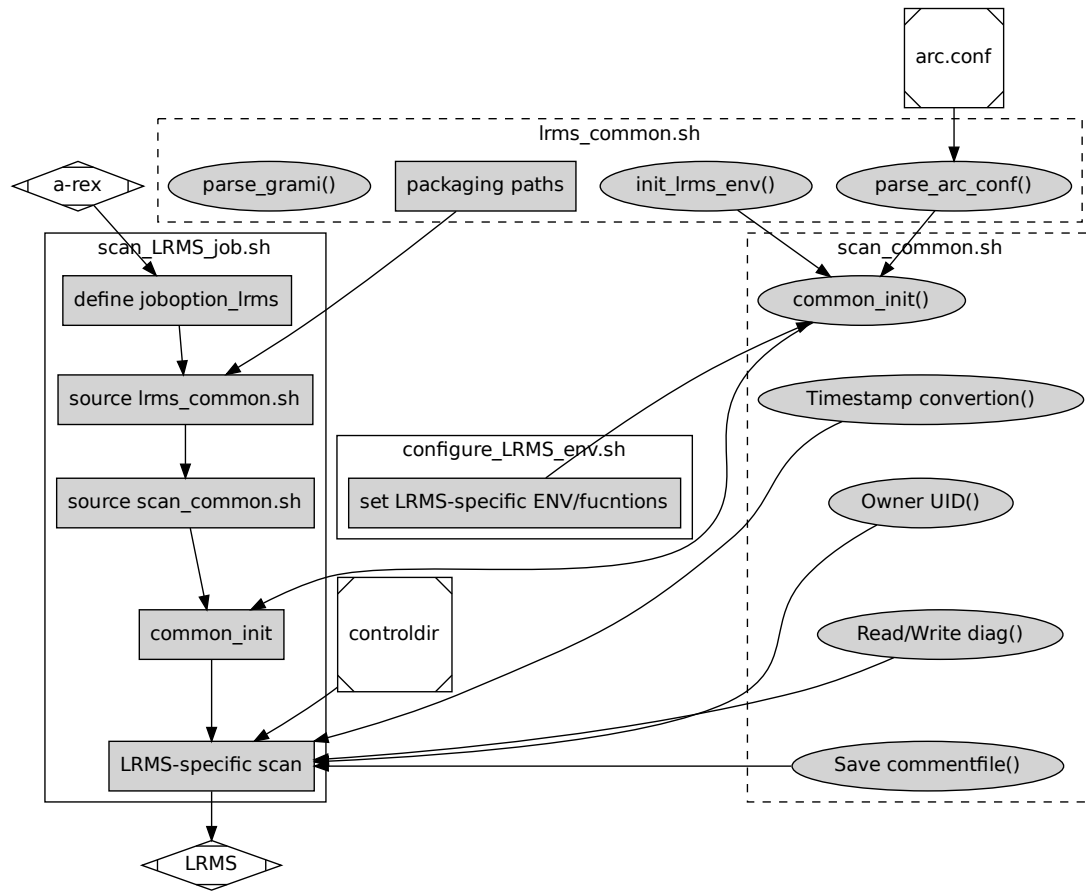
$CONFIG_sge_root [lrms]
$CONFIG_sge_cell [lrms]
$CONFIG_sge_qmaster_port [lrms]
$CONFIG_sge_execd_port [lrms]
$CONFIG_sge_bin_path [lrms]
$CONFIG_sge_jobopts [queue] [lrms]
# $CONFIG_scratchdir [arex]
```

Call graph

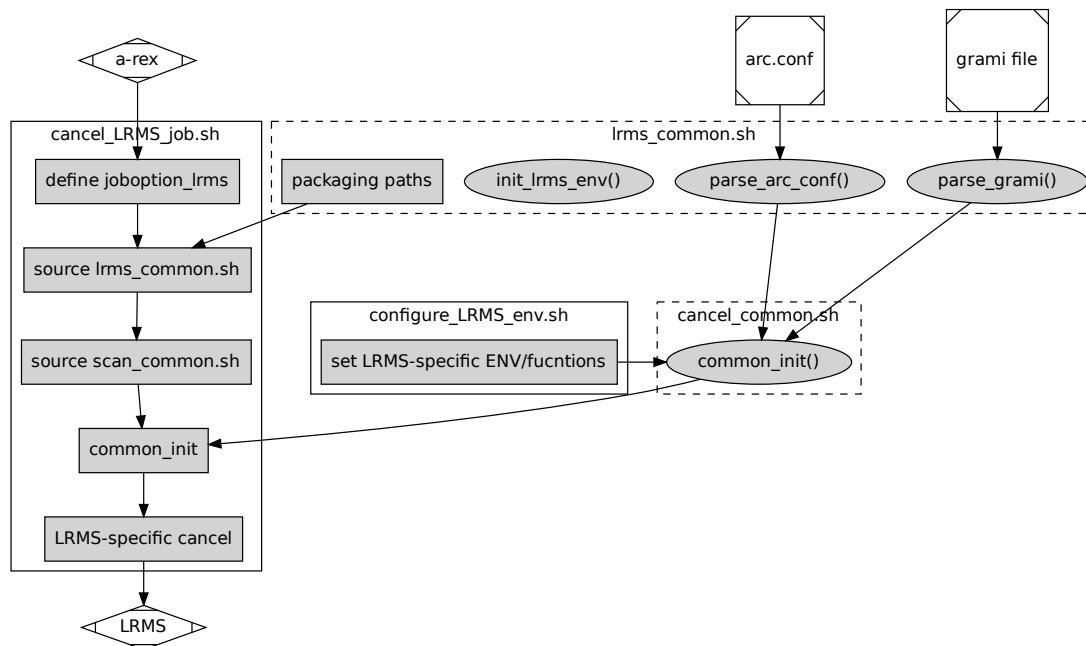
Submitting jobs



Scanning jobs



Canceling jobs



ARC memory limits processing:

Current logic of memory limits processing:

- `nodememory` - advertise memory for matchmaking: max memory on the nodes (in `[infosys/cluster]` block or per-queue)
- `defaultmemory` - enforce during submission if no memory limit specified in the job description (in `[lrms]` block or per-queue)

The ARC logic is *no enforcement = no limit*

Backends behaviour with no memory enforcement limit:

- boinc - set to hardcoded 2GB
- condor - no enforcement
- form - no memory handling at all
- ll - no enforcement
- lsf - no enforcement
- pbs - no enforcement²
- sge - no enforcement
- slurm - no enforcement

² `exclusivenode` is memory-based and `nodememory` value is used in this case

8.1.3 Contributing to Documentation

Nordugrid ARC documentation is mainly written in [reStructuredText](#) and build with [Sphinx](#) to HTML pages, LaTeX (for printable PDF versions) and ePub.

Recent source tree (master branch) build is available instantly via [Coderefinery GitLab Pages](#) and deployed to the [nordugrid.org](#) on nightly basis.

Everyone is welcome to contribute. For trivial fixes just click on the “Edit in GitLab” button and commit the changes.

You need to have an account in the [Coderefinery GitLab](#) and be a member of the project to be allowed to edit. [Contact us](#) if you want to contribute, but are not a member yet.

Committing and reviewing changes

Contribution is possible by direct push to the repo, no need for merge requests if you are fixing typos or other content that does not require review!

Note

In the rendered HTML version there is a link at the bottom of each page labeled *Edit page source on GitLab*. Using this link you can edit the page source in your web-browser without the need to checkout source tree

In case you want to add something new that requires the review process:

- create the new branch (in the same repository, no need to fork)
- open a merge request to master branch
- in-line and general discussion of the particular contribution is inside merge request.

Documentation structure

As reflected in the [main index](#) page, the documentation is logically divided into the following groups:

- *Documentation for Infrastructure Users*
- *General overview*
- *Documentation for Developers*
- *Technical Documents Describing ARC Components*

Admins documentation is the most developed part for ARC. We aimed to have dedicated documents for each particular topic and then bound them in the following way:

- The introductory *Quickstart ARC: towards distributed computing in a few minutes* contains a zero-configuration case hands-on instruction to help new users and admins getting started with ARC CE.
- Main *ARC Computing Element Installation and Configuration Guide* that contains general installation and configuration flow for production *Computing Element* deployment. The guide itself should be brief enough and holds only main examples. All detailed instruction for each particular subsystem configuration should be written in the dedicated document that linked to this guide.
- *ARC Configuration Reference Document* automatically rendered from text version in the code source tree. Targeted for web-search. The ultimate configuration options description we have.
- All other dedicated documents go to *Operating ARC CE Subsystems* section in case someone wants to access them directly instead of following installation guides links.

Source tree directory structure

Sources of documents in rST format are placed into the source directory that contains the following structure:

- `_static` - logo and CSS files¹
- `_templates` - Sphinx HTML theme layout tuning¹
- `_extensions` - Custom sphinx extensions developed for ARC¹
- **common** - directory used to store common for all ARC customers documentation, including admins, users and developes. This documents are hidden in the main TOC tree, instead they are linked from other various documents when needed.
 - `repos` - this directory holds repository configuration instructions for release, nightlies, etc.
 - `changelog` - documents that describe changes between ARC versions.
- **admins** - directory used to store *General overview*. The main documents that represent the entry points to *ARC Computing Element* installation and configuration resides directly in this directory. You also find the quickstart guide here: *Quickstart guide*.
 - `details` - this subdirectory designed to hold detailed configuration instructions for ARC CE sub-systems, like *RunTime Environments in ARC* that describe all aspects of RTEs in-depth. This documents intentionally moved deeper in the main TOC tree to prevent first levels flooding. If you can find a chapter in main *Quickstart guide* that can point to the in-depth document in question - this is a right place to store it.
 - `commands` - place for automatically generated command line option reference for ARC tools. At the time of writing the Python tools represented there.
 - `archery` - for the documents that not related to *Computing Element* another subdirectory should be created. At the time of writing only ARCHERY documentation of that kind is available.
- `developers` - directory used to store *Documentation for Developers*. As long as we don't have many of this kind of documents everything that contains implementation details (how stuff coded, which variable used, etc) should go there.
- `tech` - directory used to store *Technical Documents Describing ARC Components*.
- `users` - directory used to store *Documentation for Infrastructure Users*. If document is user-oriented - place it here.
- `sdk` - place for Doxygen SDK documentation integration to Sphinx build¹
- `testing` - directory used to store the documents for *ARC 7 Testing Area*.
- `wip` - “work in progress” area hidden in the TOC tree. You can use it for incomplete documents that should not be publicly advertised yet, but still will be built and available via `/wip/` URL in produced HTML files.

Index files

Each directory in the source tree contains `index.rst` file that is used to link other documents in the same TOC tree. Upper-level `index.rst` file contains references to `index.rst` files in the subdirectories, that in turns contains pointers to the other documents.

Note

When you add a new document, add a reference to the `index.rst` in the same directory

Just follow the `index.rst` chain (look for `toctree` keyword) starting from the source directory to get familiar with TOC linking structure.

¹ Do not touch unless you are modifying the Sphinx build itself. Nothing there affects documentation writing process.

Storing images

When you need to add images to your document you should upload the image file itself and refer to it from the `.rst`.

ARC6 documentation structure implies that image files are stored inside the `images` subdirectory in the document location. Then referencing is done by relative path, e.g.:

```
.. figure:: images/shared_sessiondir_yes.svg
```

Building the docs

The top directory in the source tree contains `build.sh` script that:

- checkout the ARC source code tree and configures it with packaging-like paths
- builds Doxygen SDK documentation
- converts `arc.conf.reference` to rST
- copies documentation parts from ARC source code tree (for *developers* section)
- prepare automatically generated documentation for CLI commands
- builds HTML
- builds PDF (with LaTeX)
- builds ePub

On commit the GitLab CI configured to automatically invoke the `build.sh` to produce rendered documentation. Documentation archive is available as CI job artifacts and for `master` branch is deployed to [GitLab Pages](#).

To build the docs on your local machine you should have at least Sphinx installed. Additionally you should be able to configure ARC source tree and ARC Python command dependencies for auto-generated parts.

Complete and up-to-date list of dependencies defined for CI build and can be found in `.gitlab-ci.yml`.

PDF builds with LaTeX is the most heavy part from the both time and needed additional packages perspective. For HTML rendering local debugging it is recommended to use the `html` script argument the skips PDF and ePub:

```
[user@localhost doc]$ ./build.sh html
```

Point your browser to `file:///path/to/doc/build/html/index.html` to view the HTML rendering locally.

Writing Documentation in reStructuredText

reStructuredText (reST) is the default plaintext markup language used by Sphinx. There is possible to render other markups but for consistency and better cross-referencing NorduGrid ARC6 documentation written solely in reStructuredText.

General Syntax

reStructuredText markup specification is well documented in the several sources and was designed to be a simple and readable in plain-text. Common text editors (including `vim` and `emacs`) recognize reST markup and provide syntax highlighting out of the box.

Start with [reStructuredText Primer](#) on Sphinx docs.

Complete reST Markup Syntax can be found on [Docutils](#) starting with [Quick reStructuredText](#) document.

Sphinx also uses interpreted text roles to insert semantic markup (cross-referencing, etc) into documents. To get familiar read [this document](#).

Just study the markup following the documentation and reading the already written documents just here.

Code snippets

There are different options for representing the code, starting with the simple literal blocks identified with `::`.

It is advised to use `code-block::` directive to enable syntax highlighted rendering:

```
.. code-block:: ini

[gridftpd/jobs]
allowaccess = staticdn dnfromfile
```

In this example, the `ini` tag represents the syntax highlighting lexer identifier for Pygments. Look for [available lexers](#) to serve your needs. The most common for our docs are: `console`, `bash`, `ini` and `cfg`.

References

Referencing another parts of documentation is necessary to achieve usability. The typical referencing cases are:

Using custom label

Create label just above any of the paragraph headers with. The following markup creates `my_label` label.

```
.. _my_label:

My Heading
=====
```

Referring to a label is possible from any other document with:

```
read the :ref:`my_label`
```

In this case paragraph heading will be used for hyperlink text. If you want some custom text for hyperlink text use the following syntax:

```
read this :ref:`text <my_label>`
```

Referring arc.conf.reference

Autogenerator reST rendering of `arc.conf.reference` already contains labels for all configuration options and blocks that can be used.

The label name has the following structure:

- `reference_<block name>[_<sub_block>...]<option name>` represents configuration option inside block
- `reference_<block name>[_<sub_block>...]` represents block itself

For example:

```
In the ``arc.conf`` there is a dedicated :ref:`[lrms] <reference_lrms>` block that
↳ defines the type of your LRMS.
*Job session directory* is configured with :ref:`reference_arex_sessiondir`
↳ configuration option.
```

Referring bugz

Custom Sphinx plugin for ARC documentation introduces new reST roles that can be used to mention bugs in NorduGrid Bugzilla or GitLab Issues:

More details can be found in `:bugz:`3802`` and `:issue:`57``.

Referencing docs

Referencing the the whole document is similar to using labels, but instead of label name the document name (filename without extension) is used with `:doc:` keyword.

Example 1: Refer to the `try_arc.rst` in the same source tree directory (relative path). Use the document header as hyperlink text:

```
:doc:`try_arc`
```

Example 2: Refer to the `repository.rst` by absolute path (starting from sources top directory). Use the custom hyperlink text:

```
:doc:`NorduGrid Repositories </common/repos/repository>`
```

Adding notes

To highlight statement visually use notes and warnings:

```
.. note::
    Zero configured A-REX comes with ARC REST and EMI-ES interfaces enabled.

.. warning::
    This information is valid for releases of ARC 6 starting from 6.0.
```

Adding images

In the ARC6 documentation images should be included with the caption. It is accomplished with the `figure` keyword.

After *storing image* inside `images` subdirectory, include it in reST document as:

```
.. figure:: images/shared_sessiondir_yes.svg
   :align: center
   :alt: Sessiondir is shared between ARC CE and WNs

   Sessiondir is shared between ARC CE and WNs. No local scratchdir defined.
```

Sphinx build had been configured with image format autocoverison feature. So you can use any image format, including vector graphics.

Graphviz

If you want to illustrate some process or structure that can be shown as graph, consider using built-in [Graphviz](#) functionality.

Simple usage examples can be found in `admins/details/rtes.rst` document within ARC documentation. Give it a try, it is easy.

Converting from other sources

Converting the old documentation parts from LaTeX or Mediawiki markdown to reST is not a magical process that do everything automatically unfortunately.

However you can get a good start with a [pandoc](#) tool that do the conversion.

E.g. to convert LaTeX source simply run:

```
[user ~]$ pandoc -f latex -t rst acix.tex > acix.rst
```

The ammount of efforts to edit the resulted `.rst` file and fix formatting issues is completely depends on the source itself (e.g. how the console ouput was formatted in the origin document). Images and references in most cases should be fixed separately after the pandoc.

DOCUMENTATION FOR INFRASTRUCTURE USERS

This part of the documentation is targeted to distributed computing infrastructure users either that use clients or SDK to run jobs and handle data transfers.

9.1 Installing ARC Client Tools

9.1.1 Step 1. Enable NorduGrid ARC repos

Prepare your system to install via the *NorduGrid Repositories*.

If you need the latest client tools that are not yet released, configure nightly builds repository following *Using ARC packages from nightly builds* instructions.

Note

On several operating systems you can install ARC without using the NorduGrid Repositories (e.g. using EPEL for CentOS/RHEL). But we recommend that you even on those systems install from the NorduGrid Repositories as they contains IGTF certificates and get updated faster on new releases.

9.1.2 Step 2. Install packages

Client tools

The main client tools package is `nordugrid-arc-client`. Install it using your OS packet manager¹:

```
[root ~]# dnf -y install nordugrid-arc-client
```

This package brings necessary plugins as dependency and allows you to submit jobs and perform data transfers in the ARC ecosystem, e.g.:

- submit jobs via ARC REST interface defined by xRSL or EMI-ES ADL descriptions
- query ARCHERY registry and ARC REST GLUE2 local information system
- support local and HTTP(S) data transfers
- support Rucio and SRM

If you need more features, additional plugins are available (see below).

¹ Examples are shown for dnf-based systems. You can use apt or any other packet managers in the similar way.

ARCCTL

ARC Control Tool is one-stop-shop for sysadmins and users running ARC that automate many operations.

For the client side it will allow you to easily deploy CA certificates and VOMS configuration:

```
[root ~]# dnf -y install nordugrid-arc-arcctl
```

Additional plugins

Note

In ARC 7, the *nordugrid-arc-plugins-arcrest* is included in the *nordugrid-arc-plugins-needed*, so you no longer need to install it specially as you had to in ARC 6.

- Data transfers:
 - *gridftp* - support for gridftp data transfers
 - *xrootd* - support for xroot protocol
 - *s3* - support for s3 protocol
 - *gfal* - add support for numerous data transfer protocols and file catalogues via installed GFAL2 plugins. No longer supported as of EL10.

```
[root ~]# dnf -y install nordugrid-arc-plugins-<PLUGIN NAME>
```

9.1.3 Step 3. Setting up credentials

e-Science distributed computing networks (grid) heavily relies on cryptography and uses either personal X.509 certificates/keys or OIDC tokens to identify entities and sets of dedicated Certification Authorities (CA).

In the most common workflow you are also required to be a member of some Virtual Organization (VO) to get access to infrastructure resources. On the technical level the VO membership is currently handled by infrastructure VOMS services.

ARC supports both x509 and OIDC tokens, and both can be used by the same ARC-CE (provided the ARC-CE is configured to) or ARC client.

A new feature in ARC 7 is to support grids that only rely on the default system CA bundle, or other non-igtf CA's. For details how the ARC-CE must be set up, and how to submit jobs to such a CE, please refer to: [ARC-CE only accepting non-igtf tokens](#).

CA certificates bundle

The CA certificates are used to verify entities (e.g. users as well as compute and storage services) and should be installed on the client host as well for infrastructure security reasons.

In the distributed grid environment a set of dedicated CAs is used as a part of [IGTF](#).

The IGTF CA certificates bundle can easily be installed with *ARC Control Tool*²:

```
[root ~]# arcctl deploy igtf-ca classic
```

² In the case that ARC is not installed from the Nordugrid repositories, use the `--installrepo` argument to enable third-party repositories with IGTF CA certificates

Personal authentication

As mentioned ARC 7 supports both OIDC tokens and x509 user certificates. To interact with an ARC-CE you must obtain one of these.

Personal X.509 certificate

For production infrastructure usage you need a personal certificate signed by one of the [IGTF accredited CAs](#).

Typically there is at least one IGTF accreditation CA in each country that you can find [on map](#).

Organizational and technical procedures varies from CA to CA, so you should read the instructions on a chosen CA web-site.

Once you get your certificate, install it to your client host and ensure the permissions are set correctly. For historical reasons the default location for certificate and key is `.globus` unless redefined in *client configuration file*:

```
[user@client ~]$ ls -l ~/.globus/
total 12
-rw-r--r-- 1 user user 6353 Oct  1 11:55 usercert.pem
-r----- 1 user user 1854 Oct  1 11:55 userkey.pem
```

Note

For testing and development purposes it is possible to use local testing CA. *ARC Control Tool* provides built-in test CA capabilities that allows you to bootstrap own test CA and generate host and user certificates.

OIDC token

To interact with an ARC-CE using tokens you must obtain a bearer token from a token issuer that your grid supports. In the case of WLCG this can be for instance the WLCG-IAM: <https://wlcg.cloud.cnaf.infn.it/login>, EGI Check-in: <https://aai.egi.eu/registry/> or SciTokens: <https://scitokens.org/>. In other grids, it will be the token issuer that your community trusts.

Once you have created an identity at the required Identity Provider, and you have obtained a bearer token, you store the token in the environment variable `BEARER_TOKEN` before issuing the client requests.

```
export BEARER_TOKEN=<token-string>
```

For more details on the OIDC token support, please see: [ARC support for OIDC](#).

Virtual Organization membership

In most cases you do have some implicit VO affiliation on your workplace. If not, you can search for VOs using e.g. [EGI VO\(s\) search tool](#).

Every Virtual Organisation (VO) has own procedures and policies. Please contact VO support team for membership instructions.

To prove that you are a member of the particular VO you need to obtain a special token from the VOMS server as a part of your authentication process.

Unfortunately, the EGI API is no longer publicly accessible. For information how to obtain the token, please visit the EGI Database for your experiment, for instance for CMS: https://operations-portal.egi.eu/vo/view/voname/cms#VOV_section. and consult your experiment about how to install the necessary vomses.

9.1.4 Step 4. Try it out

When all is set, you can try to submit a job to the grid infrastructure.

Create a proxy certificate

To submit a job, or perform any other action you need a so-called *proxy-certificate* which is a Single Sign-On token for distributed grid-infrastructure. It is generated in the following way:

```
[user ~]$ arcproxy
Your identity: /DC=org/DC=AccreditedCA/O=people/CN=John Smith
Proxy generation succeeded
Your proxy is valid until: 2019-11-28 02:06:57
```

If you need to include you VO membership confirmation token (attribute certificate), specify the VO name as well:

```
[user ~]$ arcproxy -S area51
Your identity: /DC=org/DC=AccreditedCA/O=people/CN=John Smith
Contacting VOMS server (named area51): voms.example.org on port: 15001
Proxy generation succeeded
Your proxy is valid until: 2019-11-28 02:08:27
```

Or obtain a bearer token from your token Issuer

```
export BEARER_TOKEN=<token-string>
```

Submit a test job

ARC client tools includes the `arctest` utility that comes with several test jobs on board. So you can try to submit a job without the need to write a job description.

You can use the Nordugrid top-level ARCHERY service registry (nordugrid.org) to find available CEs automatically³:

```
[user ~]$ arctest -J 2 --registry nordugrid.org
Submitting test-job 2:
&( executable = "/usr/bin/env" )( stdout = "stdout" )( stderr = "stdout" )( gmlog =
↪ "gmlog" )( jobname = "arctest2" )( clientxrsl = "&( executable = ""/usr/bin/env""
↪ )( jobname = ""arctest2"" )( stdout = ""stdout"" )( join = ""yes"" )( gmlog = "
↪ "gmlog"" )" )
Client version: nordugrid-arc-6.5.0
Test submitted with jobid: https://arc.example.org:443/arex/e932ad551df1
Computing service: KNU ARC
```

The job status can be then checked with the `arcstat` tool:

```
[user ~]$ arcstat https://arc.example.org:443/arex/e932ad551df1
Job: https://arc.example.org:443/arex/e932ad551df1
Name: arctest2
State: Running

Status of 1 jobs was queried, 1 jobs returned information
```

³ HINT: you can add `-d INFO` to get more details of CE selection process

To fetch the job's stdout run `arccat` tool:

```
[user ~]$ arccat https://arc.example.org:443/arex/e932ad551df1
GRIDMAP=/dev/null
HOSTNAME=arc.example.org
TMPDIR=/tmp
<output omitted>
```

9.2 Overview of ARC client tools

This document contains general overview of available ARC client tools. For more detailed information about possible tool usage options, please refer to the `--help` messages and corresponding man pages.

9.2.1 Submission endpoint types

In ARC 7 only the `arcrest` job submission interface is available, and therefore you do not have to specify the interface explicitly, as `arcrest` is the default interface.

Table 9.1: The submission types short and long

short	long
<code>arcrest</code>	<code>org.nordugrid.arcrest</code>
<code>internal</code>	<code>org.nordugrid.internal</code>

Example of use:

```
arcsb -T arcrest -C some-ce-endpoint
```

9.2.2 Credentials

`arcproxy`

The `arcproxy` command creates a proxy certificate from a key/certificate pair which can then be used to access resources.

It is all-in-one tool for all kind of operations with proxy certificates, including VOMS extension generation and communications with MyProxy services.

```
[user ~]$ arcproxy -S myVO
Enter pass phrase for private key:
Your identity: /DC=org/DC=AccreditedCA/O=people/CN=John Smith
Contacting VOMS server (named myVO): voms.example.org on port: 15004
Proxy generation succeeded
Your proxy is valid until: 2024-01-22 22:50:44

[user ~]$ arcproxy -I
Subject: /DC=org/DC=AccreditedCA/O=people/CN=John Smith/CN=2123720174
Issuer: /DC=org/DC=AccreditedCA/O=people/CN=John Smith
Identity: /DC=org/DC=AccreditedCA/O=people/CN=John Smith
Time left for proxy: 11 hours 56 minutes 45 seconds
Proxy path: /tmp/x509up_u1000
Proxy type: X.509 Proxy Certificate Profile RFC compliant impersonation proxy - RFC
->inheritAll proxy
```

(continues on next page)

(continued from previous page)

```

Proxy key length: 2048
Proxy signature: sha256
===== AC extension information for VO myVO =====
VO          : myVO
subject     : /DC=org/DC=AccreditedCA/O=people/CN=John Smith
issuer      : /DC=org/DC=AccreditedCA/O=hosts/CN=voms.example.org
uri         : voms.example.org:15004
attribute   : /myVO/Role=NULL/Capability=NULL
Time left for AC: 11 hours 56 minutes 47 seconds

[user ~]$ arcproxy -L myproxy.example.org -M put -S myVO -U myproxyuser
Enter pass phrase for private key:
Your identity: /DC=org/DC=AccreditedCA/O=people/CN=John Smith
Contacting VOMS server (named myVO): voms.example.org on port: 15004
Proxy generation succeeded
Your proxy is valid until: 2024-01-29 10:51:52
Enter pass phrase for MyProxy server (new):
Verifying - Enter pass phrase for MyProxy server (new):
Succeeded to put a proxy onto MyProxy server

```

arcnew

The arcnew command is used for renewing the proxy of jobs that have been submitted to computing resources.

```

[user ~]$ arcnew https://arc.example.org:443/arex/rest/1.0/jobs/f77b3d1b1efb
Jobs processed: 1, renewed: 1

```

9.2.3 Job submission and management

arcsb

The arcsb command is used for submitting jobs to computing resources.

Note

Submission endpoint selection options are changed in ARC 6.5.0 release

```

[user ~]$ arcsb -T arcrest -C arc.example.org myJob.xrsl
Job submitted with jobid: https://arc.example.org:443/arex/rest/1.0/jobs/f77b3d1b1efb

```

arcstat

The arcstat command is used for obtaining the status of jobs that have been submitted to computing resources.

```

[user ~]$ arcstat https://arc.example.org:443/arex/rest/1.0/jobs/f77b3d1b1efb
Job: https://arc.example.org:443/arex/rest/1.0/jobs/f77b3d1b1efb
Name: myJob
State: Running

Status of 1 jobs was queried, 1 jobs returned information

```

arccat

The `arccat` command performs the `cat` command on the stdout, stderr or A-REX error log of the job. It can also show the specified file from job's session directory.

```
[user ~]$ arccat https://arc.example.org:443/arex/rest/1.0/jobs/f77b3d1b1efb
SHELL=/bin/bash
HISTSIZE=1000
<output omitted>

[user ~]$ arccat -l https://arc.example.org:443/arex/rest/1.0/jobs/f77b3d1b1efb
2024-01-22T09:57:45Z Job state change UNDEFINED -> ACCEPTED Reason: (Re)Accepting
↪new job
2024-01-22T09:57:45Z Job state change ACCEPTED -> PREPARING Reason: Starting job
↪processing
2024-01-22T09:57:45Z Job state change PREPARING -> SUBMIT Reason: Pre-staging
↪finished, passing job to LRMS
----- starting submit_pbs_job -----
PBS jobname: myJob
<output omitted>
```

arckill

The `arckill` command is used to kill running jobs.

```
[user ~]$ arckill https://arc.example.org:443/arex/rest/1.0/jobs/f77b3d1b1efb
Jobs processed: 1, successfully killed: 1, successfully cleaned: 1
```

arcclean

The `arcclean` command removes a job from the computing resource.

```
[user ~]$ arckill --keep https://arc.example.org:443/arex/rest/1.0/jobs/f77b3d1b1efb
Jobs processed: 1, successfully killed: 1

[user ~]$ arcclean https://arc.example.org:443/arex/rest/1.0/jobs/f77b3d1b1efb
Jobs processed: 1, deleted: 1
```

arcresub

The `arcresub` allows to resubmit job to the same or other compute resources.

```
[user ~]$ arcresub --same https://arc.example.org:443/arex/rest/1.0/jobs/f77b3d1b1efb
Job submitted with jobid: https://arc.example.org:443/arex/rest/1.0/jobs/28c62c12a26b
```

arcresume

The `arcresume` command is used for resuming a job that was submitted to compute resources and then subsequently failed. The job will be resumed at the last ok state reported by the cluster.

```
[user ~]$ arcresume https://arc.example.org:443/arex/rest/1.0/jobs/28c62c12a26b
Jobs processed: 1, resumed: 1
```

arcget

The `arcget` command is used for retrieving the results from a job.

```
[user ~]$ arcget --keep https://arc.example.org:443/arex/rest/1.0/jobs/28c62c12a26b
Results stored at: 8c62c12a26b
Jobs processed: 1, successfully retrieved: 1
```

arctest

The `arctest` command is used for testing resources. It is able to:

- submit several test job
- prints info about installed user- and CA-certificates

Note

Submission endpoint selection options are changed in ARC 6.5.0 release

```
[user ~]$ arctest -J 2 -C arc.example.org
Job submitted with jobid: https://arc.example.org:443/arex/rest/1.0/jobs/28c62c12a26b

[user ~]$ arctest -E
Certificate information:
  Certificate: /home/john/.globus/usercert.pem
  Subject name: /DC=org/DC=AccreditedCA/O=people/CN=John Smith
  Valid until: 2024-04-21 18:48:15

Proxy certificate information:
  Proxy: /tmp/x509up_u1000
  Proxy-subject: /DC=org/DC=AccreditedCA/O=people/CN=John Smith/CN=2123720174
  Valid for: 11 hours 37 minutes 33 seconds

Certificate issuer: /DC=org/DC=AccreditedCA

CA-certificates installed:
  /DC=org/DC=AccreditedCA
<output omitted>
```

9.2.4 Data manipulation

arcls

The `arcls` command is used for listing files in storage elements and replica catalogs.

```
[user ~]$ arcls srm://se.example.org/dpm/example.org/home/myVO/
data.1
data.2

[user ~]$ arcls -lL rucio://rucio-lb-prod.cern.ch/replicas/mc15_13TeV/EVNT.19785567._
↪000729.pool.root.1
<Name>                <Type>  <Size>        <Modified>      <Checksum>
↪ <Latency>
EVNT.19785567._000729.pool.root.1  file 167468788      (n/a)
↪adler32:f7332aeb      (n/a)
    root://xrootd.lcg.triumf.ca:1094//atlas/atlasdatadisk/rucio/mc15_13TeV/ee/3b/
↪EVNT.19785567._000729.pool.root.1
    srm://srm.ndgf.org:8443/srm/managerv2?SFN=/atlas/disk/atlasdatadisk/rucio/mc15_
↪13TeV/ee/3b/EVNT.19785567._000729.pool.root.1
    gsiftp://gridftp.echo.stfc.ac.uk:2811/atlas:datadisk/rucio/mc15_13TeV/ee/3b/EVNT.
↪19785567._000729.pool.root.1
```

arccp

The `arccp` command copies files to, from and between storage services.

```
[user ~]$ arccp /mnt/data/data.123 srm://se.example.org/dpm/example.org/home/myVO/
↪data.123

[user ~]$ arccp rucio://rucio-lb-prod.cern.ch/replicas/mc15_13TeV/EVNT.19785567._
↪000729.pool.root.1 /tmp/
```

arcmkdir

The `arcmkdir` command creates directories on storage elements and catalogs.

```
[user ~]$ arcmkdir srm://se.example.org/dpm/example.org/home/myVO/myData
```

arcrename

The `arcrename` command renames files on storage elements.

```
[user ~]$ arcrename srm://se.example.org/dpm/example.org/home/myVO/data.123 srm://se.
↪example.org/dpm/example.org/home/myVO/data.124
```

arcrm

The `arcrm` command deletes files on storage elements.

```
[user ~]$ arcrm srm://se.example.org/dpm/example.org/home/myVO/data.124
```

9.2.5 Information services

arcinfo

The `arcinfo` command is used for obtaining the status of computing resources and detailed information about the resource according to published data in the information system.

```
[user ~]$ arcinfo arc.example.org
Computing service: (production)
Information endpoint: https://arc.example.org:443/arex (org.nordugrid.arcrest)
Submission endpoint: https://arc.example.org:443/arex (status: ok, interface: org.
↪nordugrid.arcrest)
```

```
[user ~]$ arcinfo -l arc.example.org | grep Implementation
Implementation name: nordugrid-arc
```

arcsync

The `arcsync` command synchronizes your local jobs database with the information at a given computing element(s).

```
[user ~]$ $ arcsync -C arc.example.org
Synchronizing the local list of active jobs with the information in the
information system can result in some inconsistencies. Very recently submitted
jobs might not yet be present, whereas jobs very recently scheduled for
deletion can still be present.
Are you sure you want to synchronize your local job list? [y/n] y
Total number of new jobs found: 7
```

9.3 How to submit the job?

The heterogeneous nature of the distributed computing requires the formal *job description* in order to match resources and execution environment of the particular CE to job needs.

The *client tools* then, based on the job description, select submission endpoint and pass job to chosen CE for processing.

9.3.1 Submission sequence

Make sure you have *client tools* installed and credentials configured first.

Then follow the sequence below.

Generate proxy certificate or JWT token

JWT

x509

To submit jobs or perform any other action towards the ARC-CE you must authenticate yourself. We will do this using a token issued from the test-jwt issuer.

To generate a token do:

```
[user ~]$ export BEARER_TOKEN=$(arcctl test-jwt token)
```

To submit jobs or perform any other action towards the ARC-CE you must authenticate yourself. We will do this with a *proxy-certificate* which is a Single Sign-On token for distributed grid-infrastructure.

Note

In most production cases you need to add VO affiliation information to proxy certificate with `-S <VO>` option.

To generate a proxy certificate do:

```
[user ~]$ arcproxy -S myVO
Your identity: /DC=org/DC=nordugrid/DC=ARC/O=TestCA/CN=My Cert
Contacting VOMS server (named myVO): voms.myvo.example.org on port: 15001
Proxy generation succeeded
Your proxy is valid until: 2025-02-29 06:31:54
```

Create job description

ARC supports the *Extended Resource Specification Language (xRSL)* and *EMI Activity Description Language (ADL)*¹ for specifying job descriptions.

According to the syntax you should write xRSL or ADL file, describing your job files and resource requirements.

For example, the xRSL job description for `arctest -J 2` job is:

```
&( executable = "/usr/bin/env" )( jobname = "arctest2" )( stdout = "stdout" )( join =
↪ "yes" )( gmlog = "gmlog" )
```

The same description in ADL:

```
<?xml version="1.0"?>
<ActivityDescription xmlns="http://www.eu-emi.eu/es/2010/12/adl" xmlns:emiypes=
↪ "http://www.eu-emi.eu/es/2010/12/types" xmlns:nordugrid-adl="http://www.nordugrid.
↪ org/es/2011/12/nordugrid-adl">
  <ActivityIdentification>
    <Name>arctest2</Name>
  </ActivityIdentification>
  <Application>
    <Executable>
```

(continues on next page)

¹ See section 9.3 of EMI-ES document for complete ADL specification

(continued from previous page)

```
<Path>/usr/bin/env</Path>
</Executable>
<Output>stdout</Output>
<Error>stdout</Error>
<LoggingDirectory>gmlog</LoggingDirectory>
</Application>
</ActivityDescription>
```

Run arctsub

The arctsub tool is responsible for job submission. This includes selecting (matchmaking) resources based on the job description and sending the job description and all job files to selected submission endpoint.

In the simplest case you can use a top-level *Nordugrid ARCHERY registry* to find CE in entire e-Infrastrucutre that fits your job and allows you to use it resources:

```
[user ~]$ arctsub --registry nordugrid.org myjobdescription.xml
Job submitted with jobid: https://arc.example.org:443/arex/e932ad551df1
```

To specify the CEs you are aiming to submit jobs to more precisely, please read the next section for details.

9.3.2 Specify CEs for job submission

The arctsub accepts a set of options that allow to select target CE for the job submission.

There CEs can be passed manually (using -C, --computing-element=ce option) or fetched from the registries (-Y, --registry=registry).

The submission endpoint type (actual job submission interface) and information endpoint type (interface to query the information for matchmaking and brokering) can be defined as well, if you want to use some interface explicitly.

Following is a detailed description of how options works together to select submission endpoint, including the examples of usage.

Computing Element

The Computing Element (-C, --computing-element=CE) used to specify computing element name or a a complete endpoint URL directly.

Supported options:

- CE name (given as FQDN of the cluster frontend)
- URL of info or jobsubmission endpoint on the CE

The URL is one of the valid info or jobsubmission URLs of the computing element. The usage of URL is needed when the services are ruining on non-default ports or path.

Except the case when *info-endpoint-type* is NONE, the URL is interpreted as info-endpoint URL. In case of --info-endpoint-type=NONE, the URL represents the jobsubmission endpoint URL.

If only the --computing-element is specified without any preferred info-endpoint-type or submission-endpoint-type then arctsub will fetch resource info from all the available info endpoints parallel AND select from the discovered and supported (via available plugin) jobsubmission interfaces.

Examples:

```
[user ~]$ arctsub --computing-element my.cluster.org
```

will trigger a submission to a randomly selected jobsubmission-endpoint-type advertised via local infosys and supported by installed submission plugins.

```
[user ~]$ arctsub --computing-element ldap://my.cluster.org:389/o=glue
```

will trigger a submission to a available jobsubmission endpoint discovered from the info interface running on the LDAP port 389.

```
[user ~]$ arctsub --submission-endpoint-type arcrest --computing-element https://my.
↳cluster.org/arex
```

will trigger a submission to the ARC REST jobsubmission endpoint after the infosys discovery completed against the specified ARC REST GLUE2 info endpoint URL. The URL above is interpreted as info URL.

```
[user ~]$ arctsub --submission-endpoint-type arcrest --info-endpoint-type NONE --
↳computing-element https://my.cluster.org/arex
```

will trigger a submission to the ARC REST jobsubmission endpoint specified via the given URL. Info discovery is turned off. The URL is interpreted as jobsubmission endpoint URL.

Registry

The -Y, --registry=registry option used to specify service endpoint URL registries.

It takes FQDN's of ARCHERY servers with optional specification of protocol.

Supported options:

- ARCHERY: `arctsub --registry nordugrid.org` or `arctsub --registry dns://nordugrid.org`

Note

Support for EGIIS, EMIR and BDII service catalogues are REMOVED and the corresponding plugins are not part of ARC7.

In case only a short hostname given with --registry without specifying the full protocol URL (ldap or dns), then arctsub will attempt to contact both `ldap://hostname` and `dns://hostname`.

Info endpoint type

Use the -Q, --info-endpoint-type=type for choosing the cluster/resource information retrieval interface type.

Supported options:

```
org.nordugrid.arcrest
org.nordugrid.ldap.nordugrid
org.nordugrid.ldap.glue2
org.nordugrid.internal
NONE
```

The special value NONE will disable the local infosys queries during job submission.

Note

Support for emies and gridftp information retrieval is REMOVED and the corresponding plugins deleted as of ARC7.

Note

The `org.nordugrid` prefix can be omitted, e.g. both `org.nordugrid.arcrest` and `arcrest` should be accepted.

Submission endpoint type

In ARC 7, only the *arcrest* submission endpoint type is available. If you are using an ARC 7 client to submit to an ARC 7 server it is therefore not necessary to specify submission endpoint. If you are submitting the job with an ARC 6 client to an ARC 7 server, you might want to specify the submission endpoint type:

Use `-T, --submission-endpoint-type=type` for choosing the jobsubmission endpoint type.

Supported options:

```
org.nordugrid.arcrest
org.nordugrid.internal
```

Note

`org.nordugrid` prefix can be omitted, e.g. both `org.nordugrid.arcrest` and `arcrest` should be accepted.

Note

support for EMI-ES and GridFTP submission interface types are REMOVED and the corresponding plugins deleted as of ARC7.

9.4 How to work with data?

ARC libraries are very flexible in terms of supported data transfer protocols. It is designed to be extendable via a set of pluggable Data Management Components (DMC).

The available protocols to work with depends on the DMCs *installed on the system* (look for *Additional plugins* in particular).

Note

- Notice the special plugin to integrate GFAL2 additional plugins for data transfer protocols in addition to ARC native DMCs. For example, to add support for legacy LFC file catalogue protocol to ARC you have to install `nordugrid-arc-plugins-gfal` and `gfal2-plugin-lfc`.
- GLAL2 is not supported in EL10 and derivatives, and we therefore do not provide packages for GFAL2 for this OS version.

This applies both to client tools and *ARC Data Staging*.

9.4.1 Data transfer URLs

File locations in ARC can be specified both as local file names, and as Internet standard *Uniform Resource Locators (URL)*. There are also some additional URL *options* that can be used.

Depending on the installed ARC components some or all of the following transfer protocols and metadata services are supported:

Table 9.2: List of main supported protocols

Protocol	Description
ftp	ordinary <i>File Transfer Protocol (FTP)</i>
gsiftp	GridFTP, the Globus -enhanced FTP protocol with security, encryption, etc. developed by The Globus Alliance
http	ordinary <i>Hyper-Text Transfer Protocol (HTTP)</i> with PUT and GET methods using multiple streams
https	HTTP with SSL
httpg	HTTP with Globus GSI
dav	WebDAV
davs	WebDAV with SSL
ldap	ordinary <i>Lightweight Data Access Protocol (LDAP)</i>
srm	Storage Resource Manager (SRM) service
root	Xrootd protocol
rucio	Rucio – a data management system used by ATLAS and other scientific experiments
s3	Amazon S3
file	local to the host file name with a full path

An URL can be used in a standard form, i.e.

```
protocol://[host[:port]]/file
```

Or, to enhance the performance or take advantage of various features, it can have additional options:

```
protocol://[host[:port]][;option[;option[...]]/file[:metadataoption[:metadataoption[...]]]
```

For a metadata service URL, construction is the following:

```
protocol://[url[|url[...]]@]host[:port][;option[;option[...]]]
/lfn[:metadataoption[:metadataoption[...]]]
```

where the nested URL(s) are physical replicas. Options are passed on to all replicas, but if it is desired to use the same option with a different value for all replicas, the option can be specified as a common option using the following syntax:

```
protocol://[;commonoption[;commonoption]] [url[|url[...]]@]host[:port]
[;option[;option[...]]/lfn[:metadataoption[:metadataoption[...]]]
```

In user-level tools, URLs may be expressed using this syntax, or there may be simpler ways to construct complex URLs. In particular, command line tools such as `arccp`, and the `xRSL` languages provide methods to express URLs and options in a simpler way.

For the SRM service, the syntax is:

```
srm://host[:port][;options]/[service_path?SFN=]file[:metadataoptions]
```

Versions 1.1 and 2.2 of the SRM protocol are supported. The default *service_path* is `srm/managerv2` when the server supports v2.2, `srm/managerv1` otherwise.

For Rucio the following URL is used to look up replicas of the given scope and name:

```
rucio://rucio-lb-prod.cern.ch/replicas/scope/name
```

The Rucio authorisation URL can be specified with the environment variable `$RUCIO_AUTH_URL`. The Rucio account to use can be specified either through the `rucioaccount` URL option or `$RUCIO_ACCOUNT` environment variable. If neither are specified the account is taken from the VOMS `nickname` attribute.

S3 authentication is done through keys which must be set by the environment variables `$S3_ACCESS_KEY` and `$S3_SECRET_KEY`.

The URL components are:

<code>host[:port]</code>	Hostname or IP address [and port] of a server
<code>lfn</code>	Logical File Name
<code>url</code>	URL of the file as registered in indexing service
<code>service_path</code>	End-point path of the web service
<code>file</code>	File name with full path
<code>option</code>	URL option
<code>commonoption</code>	URL option for all replicas
<code>metadataoption</code>	Metadata option

The following URL options are supported:

<code>threads=<number></code>	specifies number of parallel streams to be used by GridFTP or HTTP(s,g); default value is 1, maximal value is 10
<code>exec=yes no</code>	means the file should be treated as executable
<code>preserve=yes no</code>	specify if file must be uploaded to this destination even if job processing failed (default is no)
<code>cache=yes no renew copy c</code>	indicates whether the file should be cached; default for input files in A-REX is yes. <code>renew</code> forces a download of the file, even if the cached copy is still valid. <code>copy</code> forces the cached file to be copied (rather than linked) to the session directory, this is useful if for example the file is to be modified. <code>check</code> forces a check of the permission and modification time against the original source. <code>invariant</code> disables checking the original source modification time.
<code>readonly=yes no</code>	for transfers to <code>file://</code> destinations, specifies whether the file should be read-only (unmodifiable) or not; default is yes
<code>secure=yes no</code>	indicates whether the GridFTP data channel should be encrypted; default is no
<code>blocksize=<number></code>	specifies size of chunks/blocks/buffers used in GridFTP or HTTP(s,g) transactions; default is protocol dependent
<code>checksum=cksum md5 adler3</code>	specifies the algorithm for checksum to be computed (for transfer verification or provided to the indexing server). This is overridden by any metadata options specified (see below). If this option is not provided, the default for the protocol is used. <code>checksum=no</code> disables checksum calculation.
<code>overwrite=yes no</code>	makes software trying (or not) to overwrite existing file(s); if yes, the tool will try to remove any information/content associated with the specified URL before writing to the destination.
<code>protocol=gsi gssapi ssl t</code>	distinguishes between different kinds of HTTPS/HTTPG and SRM protocols. Here <code>gssapi</code> stands for HTTPG implementation using only GSSAPI functions to wrap data and <code>gsi</code> uses additional headers as implemented in Globus IO. The <code>ssl</code> and <code>tls</code> options stand for the usual HTTPS and are specifically usable only if used with the SRM protocol. The <code>ssl3</code> option is mostly the same as the <code>ssl</code> one but uses SSLv3 handshakes while establishing HTTPS connections. The default is <code>gssapi</code> for SRM connections, <code>tls</code> for HTTPS and <code>gssapi</code> for HTTPG. In the case of SRM, if default fails, <code>gsi</code> is tried.
<code>spacetoken=<pattern></code>	specifies a space token to be used for uploads to SRM storage elements supporting SRM version 2.2 or higher
<code>autodir=yes no</code>	specifies whether before writing to the specified location the software should try to create all directories mentioned in the specified URL. Currently this applies to FTP and GridFTP only. Default value for these protocols is yes
<code>tcpnodelay=yes no</code>	controls the use of the TCP_NODELAY socket option (which disables Nagle's algorithm). Applies to HTTP(S) only. Default is no (supported only in <code>arc1s</code> and other <code>arc*</code> tools)
<code>transferprotocol=protocol</code>	specifies transfer protocols for meta-URLs such as SRM. Multiple protocols can be specified as a comma-separated list in order of preference.
<code>rucioaccount=account</code>	specifies the Rucio account to use when authenticating with Rucio.
<code>httpputpartial=yes no</code>	while storing a file on a HTTP(S) server, the software will try to send it in chunks/parts. If the server reports error for the partial PUT command, the software will fall back to transferring the file in a single piece. This behavior is non-standard and not all servers report errors properly. Hence the default is a safer no.
<code>httpgetpartial=yes no</code>	while retrieving a file from a HTTP(S) server, the software will try to read it in chunks/parts. If the server does not support the partial GET command, it usually ignores requests for partial transfer range and the file is transferred in one piece. Default is yes.
<code>failureallowed=yes no</code>	if set to yes for a job input or output file, then a failure to transfer this file will not cause a failure of the job. Default is no.
<code>relativeuri=yes no</code>	if set to yes, HTTP operations will use the path instead of the full URL. Default is no.
<code>accesslatency=disk tape</code>	filter replicas returned from an index service based on their access latency. (available from version 6.12)
<code>tlscred=none x509</code>	specifies kind of authentication information to be used in TLS negotiation -
9.4. How to work with data?	x509 for X.509 proxy or key/certificate, none to force using nothing. Default is x509 .
<code>tokencred=none token</code>	specifies kind of authentication information to be put into HTTP header - token for bearer token, none to force using nothing. Default is token .

Local files are referred to by specifying either a location relative to the job submission working directory, or by an absolute path (the one that starts with /), preceded with a `file://` prefix.

URLs also support metadata options which can be used for registering additional metadata attributes or querying the service using metadata attributes. These options are specified at the end of the LFN and consist of name and value pairs separated by colons. The following attributes are supported:

<code>checksumtype</code>	Type of checksum. Supported values are cksum (default), md5 and Adler32
<code>checksumvalue</code>	The checksum of the file

The checksum attributes may also be used to validate files that were uploaded to remote storage.

Examples of URLs are:

- `http://grid.domain.org/dir/script.sh`
- `gsiftp://grid.domain.org:2811;threads=10;secure=yes/dir/input_12378.dat`
- `ldap://grid.domain.org:389/lc=collection1,rc=Nordugrid,dc=nordugrid,dc=org`
- `file:///home/ausser/griddir/steer.cra`
- `srm://srm.domain.org/griddir/user/file1:checksumtype=adler32:checksumvalue=123456781`
- `srm://srm.domain.org;transferprotocol=https/data/file22`
- `rucio://rucio-lb-prod.cern.ch/replicas/user.grid/data.root`

9.4.2 Stage-in during submission

During the job submission to computing element, data can be moved from the client machine along with the job description.

The *inputFiles* directive of the job description automatically activates this kind of data movement if `source` of the data is a local path or empty (`""`). An empty source value means that input filename is taken from the current working directory on the submission machine.

For example:

```
(inputFiles=("data1" "/mnt/data/analyses/data11.22.33")
("data2" ""))
```

During the job submission process the `/mnt/data/analyses/data11.22.33` file will be uploaded to the job session directory on the CE as `data1`. The `data2` file from the current working directory will be uploaded as `data2`.

9.4.3 Stage-in on ARC CE

Instead of copying data from submission machine ARC CE can download it from the available storage services specified by *URL* in the *inputFiles* source value.

In this case all the power of the *A-REX Data Caching* can be used as well.

For example:

```
(inputfiles = "data.root.1" "srm://srm.mystorage.example.org/atlas/disk/data11_7TeV/
↪data.root.1")
```

During job submission the data WILL NOT be uploaded. Instead, A-REX will analyze the job description and starts a Stage-In process. During Stage-in the `data.root.1` file will be downloaded from provided SRM URL.

¹ This is a destination URL. The file will be copied to `srm.domain.org` at the path `griddir/user/file1` and the checksum will be compared to what is reported by the SRM service after the transfer.

² This is a source or destination URL. When getting a transport URL from SRM, the HTTPS transfer protocol will be requested.

9.4.4 Manual data movement with arc* tools

There is a set of *data management tools* that can be used to manipulate data manually, out of jobs processing context.

arcls, arccp, arcrm and arcmkdir work similar to classic Unix data movement commands but accept local or remote URLs.

Again, *any URL* supported by installed ARC data management plugins can be passed to the tools as an argument.

```
[user ~]$ arccp srm://srm.mystorage.example.org/atlas/disk/data11_7TeV/data.root.1 /
↪mnt/data/data.root.1

[user ~]$ arcls http://download.nordugrid.org/repos/6
centos/
debian/
fedora/
ubuntu/
```

9.5 Job Description Language (xRSL)

In the distributed computing job are submitted to a large number of very different computing resources, which are often widely distributed geographically.

In order to abstract from the heterogeneous nature of these resources, a high-level job description language is necessary. This document introduces the xRSL job description language.

Note

The *EMI Activity Description Language (ADL)* is also supported by ARC (refer to section 9.3 of EMI-ES document for complete ADL specification)

The diversity of resources implies special requirements for a proper description of a job, introducing many new options as opposed to a conventional high-performance computing center use case.

The NorduGrid project introduced the **extended Resource Specification Language (xRSL)**. xRSL adopts the general syntax of RSL language developed by Globus Alliance but extend and redesign some attributes and introduce the two levels of job option specifications:

- **User-side xRSL**, i.e., the set of attributes specified by a user in a job-specific file. This file is interpreted by a *Client* and after the necessary modifications is passed to the ARC computing service – A-REX
- **Server-side xRSL**, i.e., the set of attributes pre-processed by a client, and ready to be interpreted by the A-REX. Effectively, this is an *internal* job representation of ARC.

A user only has to know the user-side part, and utilize it to describe the grid jobs. A-REX, however, uses slightly different notations, supplied by the client tools; therefore developers of such tools must take care of converting user-submitted xRSL to the internal job description (Server-side xRSL).

In what follows, description of the xRSL is given, using the following notations:

<xxxx>	parameter to be substituted with a corresponding string or a number
[xxxx]	optional parameter
xxx yyy zzz	list of possible values of a parameter
–"	same as above

9.5.1 xRSL syntax and rules

A job is described by means of xRSL attributes, which can be either passed via a command-line, or, more conveniently, be collected in a so-called xRSL-file (*.xrs1).

Such a file contains a plain list of attribute-value pairs and boolean operators & (for AND) and | (for OR). Attribute names are **case-insensitive**.

If the attributes are specified from the command line, the entire description must be enclosed either in single or in double quotes. Single quotes enclosure is completely analogous to xRSL-file usage, while double quotes imply standard shell expansion of the enclosed string. This has implications when strings contain local shell variables: they **will not be expanded** unless the task description is entered from the command line and is enclosed in double quotes.

Attribute value

An attribute-value pair is a key element of specification. It consists of an expression that assigns one or more values to an attribute, and is enclosed in round brackets:

```
(attribute="value")
```

For multi-valued attributes:

```
(attribute="value1" "value2")
```

List of values

Certain attributes do not have assigned value; instead, they contain a list of values that should be put in proper correspondence with each other:

```
(attribute=("value1" "value2")("value3" "value4"))
```

In the example above, value1 and value3 are put in correspondence to value2 and value4 respectively, according to the context of the attribute.

Quoting

Values should be enclosed in quotes if they contain blank spaces or special characters.

The special characters are:

```
+ & | ( ) = < > ! " ' ^ # $
```

To quote a string containing special characters, you can use either single or double quotes. If your string, however, contains both such quotes, you can define any character as an own delimiter, by preceding it with the “carat” (^) character: attribute=^*My "good" value^* makes use of a carat-escaped asterisk as a delimiter.

Implicit conjunction

An xRSL job description starts with an ampersand (&) , to indicate implicit conjunction of all the attributes:

```
&(attribute1=value1)(attribute2="value 2")...
```

Whenever a disjunct-request of two or more attributes is needed, the following construction can be used:

```
(|(attribute="value1")(attribute="value2")...)
```

Note

Only few selected attributes (indicated further in the document) can be requested by the user multiple times, like in the disjunct request example above. Most attributes **must be unique**, i.e., appear only once in the job description document.

Operators

In expressions, the following relational operators are allowed, in general:

```
=    !=    >    <    >=    <=
```

However, most attributes can *only be used with equality operator* =. For few attributes (as indicated in the document), some other operators can be used in client-side job description as well.

Blank spaces

The xRSL attributes can be written in a single string, or split in lines arbitrary; blank spaces between and inside (attribute="value") relations are ignored.

Comments

Commented lines should start with (*) and be closed with *):

```
(*attribute="value1"*)
```

Comments can not be nested.

Multiple jobs in one description

Multiple job description in one file can be specified via multi-request operator +, which should precede multiple job descriptions:

```
+(&( ... ))(&( ... ))(&( ... ))
```

9.5.2 User-side xRSL attributes

The following attributes can be specified in a user's xRSL script. Some have to be modified by the client tool before being passed to the A-REX.

Attribute names are case-insensitive, although assigned values may well be case-sensitive, if they represent file names, environment variables etc.

Note

It is possible to use unsupported attributes in job description. Standard ARC client submission commands (`arcsub` and `arcresub`) must be used with a command line option `-U` in order to accept unknown attributes. Without this command line option, the client tool will consider job description invalid if it contains unsupported attributes.

executable

Unique:	yes
Operators:	=
User xRSL:	(executable=<string>)
A-REX xRSL:	-"-
Example:	(executable="local_to_job.exe")

The executable to be submitted as a main task to a Local Resource Management System (LRMS).

- **string** - file name (including path), local to the computing element (CE)

Executable is a file that has to be executed as the main process of the task. It could be either a pre-compiled binary, or a script. Users may transfer their own executables, or use the ones known to be already installed on the remote system (CE).

If an executable has to be transferred to the destination site (CE) from some source, it has to be specified in the `inputFiles` list. If it is not specified in `inputFiles`, the source is expected to be local to the user (client) and will be added as such to the `inputFiles` list by the ARC Client.

If the file name starts with a leading slash (/), it is considered to be **the full path to the executable at the destination site (CE)**; otherwise the location of the file is **relative** to the session directory (where job input and files are stored).

If the xRSL string is entered from the command line and is enclosed in double quotes, standard shell expansion of variables takes place. That is, if the file name contains an environment variable (`$...`), the value of this variable is resolved locally, but if the name itself is also enclosed in double quotes, it will be resolved at the remote computing element:

(executable=`$ROOT_DIR/myprog.exe`) – `$ROOT_DIR` is resolved locally (*will cause errors if the path does not exist at the execution machine*)

(executable="`$ROOT_DIR/myprog.exe`") – `$ROOT_DIR` will be resolved remotely

To use executable at the destination site without specifying absolute path a trivial wrapper script can be used as the executable input file:

```
#!/bin/sh

# load local environment
module load openmpi/latest
```

(continues on next page)

(continued from previous page)

```
# pass-through all the arguments
mpirun "$@"
```

Other option is to use dummy environment variable:

(executable="\${PROG_PATH}\${PROG_PATH:+/}myprog.exe") – \$PROG_PATH will be resolved remotely and possibly provided by relevant RTE or be an empty string when left undefined. In the latter case myprog.exe will be searched in the \$PATH list on the destination site (CE).

The executable can be omitted. In that case the job can be used to copy files from sources in inputFiles to destinations in outputFiles by giving them matching on-CE filename`s, or it can be used to seed the CE's cache with the ``inputFiles. Both operations may be suboptimal: The copy may go via the CE and may not be streamed, and the presence of any particular file in cache at any given point in time can't be predicted.

arguments

Unique:	yes
Operators:	=
User xRSL:	(arguments=<string> [string] ...)
A-REX xRSL:	(arguments=<executable> <string> [string] ...)
Example:	(arguments="10000" \$(ATLAS)/input.dat)

List of the arguments for the executable.

- **string** - an argument
- **executable** - the executable to be run by LRMS, taken by the ARC Client from the user-specified executable attribute

inputFiles

Unique:	yes
Operators:	=
User xRSL:	(inputFiles=(<filename> <source> [option] ...) ...)
A-REX xRSL:	(inputFiles=(<filename> <URL> [option] ...) (<filename> [size][.checksum]) ...)
Example:	(inputFiles=("local_to_job" "gsiftp://sel.lu.se/p1/remote.1" "threads=5") ("local_to_job.dat" "/scratch/local_to_me.dat") ("same_name_as_in_my_current_dir" ""))

List of files to be copied to the computing element before job execution.

- **filename** - destination file name, local to the computing element and always relative to the session directory
- **source** - source of the file: (remote URLs, or a path, local to the submission node). If void ("", use the quotes!), the input file is taken from the submission directory.
- **option** - URL options for source. The ARC Client converts **source** and any options given here to a URL with the syntax described [here](#). Additionally **delegationid=<id>** can be specified here to assign X.509 delegation identifier to be used for corresponding URL.
- **URL** - URL of the file
- **size** - file size in bytes
- **checksum** - file checksum (as returned by cksum)

If the **inputFiles** list does not contain the standard input file (as specified by **stdin**) and/or the executable file (as specified by **executable**), an ARC client must append these files to the list. If the **<source>** is a URL, any options given by **option** are added to it, then it is passed by the ARC Client to the A-REX as shown in the example above. A-REX recognizes all URLs except **file:///**.

Internally, the client must forward the (**<filename>** **<source>** [**option**] ...) request to the execution service without changes, unless **<source>** is a local path, void ("") or **file:///**. In case **<source>** is a local path, void ("") or **file:///**, the client must extract file size and checksum, and substitute the **<source>** string with **[size].[checksum]**. In the unlikely case when it is impossible to extract file size, the **<source>** string must be substituted by a void one ("").

Note

Please note that the **inputFiles** attribute is not meant to operate with directories, for reasons of access control and checksum verifications. You must specify a pair ("**<local_to_job>**" "**<source>**") for each file.

executables

Unique:	yes
Operators:	=
User xRSL:	(executables=<string> [string] ...)
A-REX xRSL:	-"-
Example:	(executables="myscript.sh" "myjob.exe")

List of files from the **inputFiles** set, which will be given executable permissions.

- **string** - file name, local to the computing element and relative to the session directory

If the executable file (as specified in **executable** and if relative to the session directory) is not in the **executables** list, it will be added to the list by the ARC Client.

cache

Unique:	yes
Operators:	=
User xRSL:	(cache="yes" "no")
A-REX xRSL:	-"-
Example:	(cache="yes")

Specifies whether input files specified in the **inputFiles** should be placed by default in the cache or not. This affects all input files, even those described by **executables**.

If not specified, default value is **yes**.

Note

Cached files can not be modified by jobs by default. If your job has to modify input files, please use the (`readonly=no`) URL option for those files. This option does not affect whether or not the file is cached.

outputFiles

Unique:	yes
Operators:	=
User xRSL:	(outputFiles=(<code><string></code> <code><URL></code> [<code>option</code>] ...) ...)
A-REX xRSL:	(outputFiles=(<code><string></code> <code><URL></code>) ...)
Example:	(outputFiles=("local_to_job.dat" "gsiftp://se1.uio.no/stored.dat") ("local_to_job_dir/" ""))

List of files to be retrieved by the user or uploaded by the A-REX and optionally indexed (registered) in a data indexing service.

- **string** - file name, local to the *Computing Element (CE)*. If this string ends with a backslash / and `<URL>` is empty, the entire directory will be kept at the execution site. If however this string ends with a backslash / but the `<URL>` is a remote location, the contents of the directory are transferred to the destination.
- **URL** - destination URL of the remote file; if void ("", use the quotes!), the file is kept for manual retrieval. Note that this can not be a local file: // URL.
- **option** - URL options for destination URL. See [URL Options](#) for possible values. When the destination is an indexing service, a physical file location may be specified by the additional option `location`. This option can be given multiple times. The CE will attempt to upload the file to the specified locations in the order they are given until one succeeds. Options specified after a location option only affect that location. Another additional `delegationid=<id>` option can be specified here to assign X.509 delegation identifier to be used for corresponding URL. Before passing to the A-REX, the ARC Client adds to URL any options and locations given here, using the syntax described in [URL Options](#).

If the list does not contain standard output, standard error file names and A-REX log-files directory name (as specified by `stdout`, `stderr` and `gmlog`), the ARC Client appends these items to the `outputFiles` list. If the `<URL>` is not specified (void, "", use the quotes!), files will be kept on the CE and should be downloaded by the user via the ARC Client. If specified name of file ends with /, the entire directory is kept.

A convenient way to keep the entire job directory at the remote site for a manual retrieval is to specify (`outputfiles=(/ )`).

In some cases, the list of output files may only be known after the job has completed. ARC allows a user to specify a list of output files dynamically in a file or files in the session directory as part of their job. The file(s) containing the output file information can be defined in the xRSL script as the path to the file relative to the session directory preceeded by '@'. The format of these files is lines of 2 values separated by a space. The first value contains name of the output file relative to the session directory and the second value is a URL to which the file will be uploaded.

Example:

```
(outputFiles="@output.files" "")
output.files is generated by the user and contains
file1 gsiftp://grid.domain.org/file1
file2 gsiftp://grid.domain.org/file2
```

After the job completes, the file `output.files` in the session directory will be read and any files described within will be uploaded to the given URLs.

cpuTime

Unique:	yes
Operators:	=
User xRSL:	(cpuTime=<time>)
A-REX xRSL:	(cpuTime=<tttt>)
Example:	(cpuTime="240")

Maximal CPU time request for the job. For a multi-processor job, this is a sum over all requested processors.

- `time` - time (in minutes if no unit is specified)
- `tttt` - time converted by the ARC Client from `time` to seconds.

The client converts time specified in the user-side XRSL file to seconds. If no time unit is specified, the client assumes the time given in minutes. Otherwise, a text format is accepted, i.e., any of the following will be interpreted properly (make sure to enclose such strings in quotes!):

```
1 week
3 days
2 days, 12 hours
1 hour, 30 minutes
36 hours
9 days
240 minutes
```

If both `cpuTime` and `wallTime` are specified, the ARC Client converts them both. `cpuTime` can not be specified together with `gridTime` or `benchmarks`.

Note

This attribute should be used to direct the job to a system with sufficient CPU resources, typically, a batch queue with the sufficient upper time limit. Jobs exceeding this maximum most likely will be **terminated** by remote systems! If time limits are not specified, the limit is not set and jobs can run as long as the system settings allow (note that in this case you can not avoid queues with too short time limits).

wallTime

Unique:	yes
Operators:	=
User xRSL:	(wallTime=<time>)
A-REX xRSL:	(wallTime=<tttt>)
Example:	(wallTime="240")

Maximal wall clock time request for the job.

- `time` - time (in minutes if no unit is specified)
- `tttt` - time converted by the ARC Client to seconds

The client converts time specified in the user-side XRSL file seconds. If no time unit is specified, the client assumes the time given in minutes. Otherwise, a text format is accepted, i.e., any of the following will be interpreted properly (make sure to enclose such strings in quotes!):

1 week
 3 days
 2 days, 12 hours
 1 hour, 30 minutes
 36 hours
 9 days
 240 minutes

If both `cpuTime` and `wallTime` are specified, the ARC Client converts them both. `wallTime` can not be specified together with `gridTime` or `benchmarks`. If only `wallTime` is specified, but not `cpuTime`, the corresponding `cpuTime` value is evaluated by the ARC Client and added to the job description.

Note

This attribute should be used to direct the job to a system with sufficient CPU resources, typically, a batch queue with the sufficient upper time limit. Jobs exceeding this maximum most likely will be **terminated** by remote systems! If time limits are not specified, the limit is not set and jobs can run as long as the system settings allow (note that in this case you can not avoid queues with too short time limits).

gridTime

Unique:	yes
Operators:	=
User xRSL:	(gridTime=<time>)
A-REX xRSL:	none
Example:	(gridTime="2 h")

Maximal CPU time request for the job scaled to the 2.8 GHz Intel Pentium 4 processor.

- `time` - time (in minutes if no unit is specified)

The attribute is completely analogous to `cpuTime`, except that it will be recalculated to the actual CPU time request for each queue, depending on the published processor clock speed.

`gridTime` can not be specified together with `cpuTime` or `wallTime`. If only `gridTime` is specified, but not `cpuTime`, the corresponding `cpuTime` value is evaluated by the ARC Client and added to the job description.

benchmarks

Unique:	yes
Operators:	=
User xRSL:	(benchmarks=(<code><string></code> <code><value></code> <code><time></code>) ...)
A-REX xRSL:	
Example:	(benchmarks=("mybenchmark" "10" "1 hour, 30 minutes"))

Evaluate a job's `cpuTime` based on benchmark values.

- `string` - benchmark name
- `value` - benchmark value of reference machine
- `time` - the `cpuTime` the job requires on the reference machine

`benchmarks` can not be specified together with `cpuTime` or `wallTime`. If only `benchmarks` is specified, but not `cpuTime`, the corresponding `cpuTime` value is evaluated by the ARC Client and added to the job description.

memory

Unique:	yes
Operators:	=
User xRSL:	(memory= <code><integer></code>)
A-REX xRSL:	- "-"
Example:	(memory>="500")

Memory required for the job, per count for parallel jobs.

- `integer` - size (Mbytes)

Note

Similarly to `cpuTime`, this attribute should be used to direct a job to a resource with a sufficient capacity. Jobs exceeding this memory limit will most likely be **terminated** by the remote system.

disk

Unique:	no
Operators:	= != > < >= <=
User xRSL:	(disk= <code><integer></code>)
A-REX xRSL:	none
Example:	(disk="500")

Disk space required for the job.

- `integer` - disk space, Mbytes

Note

This attribute is used at the job submission time to find a system with sufficient disk space. However, it **does not guarantee** that this space will be available at the end of the job, as most known systems do not allow for disk space allocation. Eventually, a remote system can terminate a job that exceeds the requested disk space.

runTimeEnvironment

Unique:	no
Operators:	= != > < >= <=
User xRSL:	(runTimeEnvironment=<string> [argument] ...)(runTimeEnvironment=<string> [argument] ...)
A-REX xRSL:	only = is allowed
Example:	(runTimeEnvironment>="APPS/HEP/ATLAS-10.0.1" "find" "HIGGS")

Required runtime environment.

- `string` - environment name

The site to submit the job to will be chosen by the ARC Client among those advertising specified runtime environments. Before starting the job, the A-REX will set up environment variables and paths according to those requested. Runtime environment names are defined by Virtual Organizations, and tend to be organized in name spaces.

To request several environments, repeat the attribute string:

```
(runTimeEnvironment="ENV1")(runTimeEnvironment="ENV2") etc.
```

To make a disjunct-request, use a boolean expression:

```
(|(runTimeEnvironment="env1")(runTimeEnvironment="env2"))).
```

You can use `>=` or `<=` operators: job will be submitted to any suitable site that satisfies such requirements, and among the available at the sites runtime environments, the highest version satisfying a requirement will be requested in the pre-processed xRSL script.

Runtime environment string interpretation is case-insensitive. If a runtime environment string consists of a name and a version number, a partial specification is possible: it is sufficient to request only the name and use `>` or `>=` operators to select the highest version.

The optional arguments can be supplied as additional strings after runtime environment string. On the server side they are passed to software responsible for preparing corresponding infrastructure.

middleware

Unique:	no
Operators:	= != > < >= <=
User xRSL:	(middleware=<string>)
A-REX xRSL:	only = is allowed
Example:	(middleware="nordugrid-arc-0.5.99")

Required middleware version. Make sure to specify full name and version number.

```
string  Grid middleware name.
```

The site to submit the job to will be chosen by the ARC Client among those advertising specified middleware. Usage is identical to that of the `runTimeEnvironment`. Use the `>=` operator to request a version “equal or higher”.

opsys

Unique:	no
Operators:	= != > < >= <=
User xRSL:	(opsys=<string>)
A-REX xRSL:	only = is allowed
Example:	(opsys="FC3")

Required operating system.

string Operating system name and version.
--

The site to submit the job to will be chosen by the ARC Client among those advertising specified operating system. Usage is identical to that of `runTimeEnvironment` and `middleware`. Use the `>=` operator to request a version “equal or higher”.

stdin

Unique:	yes
Operators:	=
User xRSL:	(stdin=<string>)
A-REX xRSL:	-"-
Example:	(stdin="myinput.dat")

The standard input file.

string file name, local to the computing element.
--

The standard input file should be listed in the `inputFiles` attribute; otherwise it will be forced to that list by the ARC Client.

stdout

Unique:	yes
Operators:	=
User xRSL:	(stdout=<string>)
A-REX xRSL:	-"-
Example:	(stdout="myoutput.txt")

The standard output file.

string file name, local to the computing element and relative to the session directory.
--

The standard output file should be listed in the `outputFiles` attribute; otherwise it will be forced to that list by the ARC Client. If the standard output is not defined, ARC Client assigns a name.

stderr

Unique:	yes
Operators:	=
User xRSL:	(stderr=<string>)
A-REX xRSL:	-"-
Example:	(stderr="myjob.err")

The standard error file.

string file name, local to the computing element and relative to the session directory.

The standard error file should be listed as an `outputFiles` attribute; otherwise it will be forced to that list by the ARC Client. If the standard error is not defined, ARC Client assigns a name. If `join` is specified with value `yes`, ARC Client adds `stderr` to the pre-processed xRSL script with the same value as `stdout`.

join

Unique:	yes
Operators:	=
User xRSL:	(join="yes" "no")
A-REX xRSL:	none
Example:	(join="yes")

If yes, joins `stderr` and `stdout` files into the `stdout` one. Default is no.

gmlog

Unique:	yes
Operators:	=
User xRSL:	(gmlog=<string>)
A-REX xRSL:	-"-
Example:	(gmlog="myjob.log")

A name of the directory containing grid-specific diagnostics per job.

string a directory, local to the computing element and relative to the session directory

This directory is kept in the session directory to be available for retrieval (ARC Client forces it to the list if `outputFiles`)

jobName

Unique:	yes
Operators:	=
User xRSL:	(jobName=<string>)
A-REX xRSL:	- "-
Example:	(jobName="My Job nr. 1")

User-specified job name.

string	job name
--------	----------

This name is meant for convenience of the user. It can be used to select the job while using the ARC Client. It is also available through the Information System.

ftpThreads

Unique:	yes
Operators:	=
User xRSL:	(ftpThreads=<integer>)
A-REX xRSL:	- "-
Example:	(ftpThreads="4")

Defines how many parallel streams will be used by the A-REX during gsiftp and http(s|g) transfers of files.

integer	a number from 1 to 10
---------	-----------------------

If not specified, parallelism is not used.

acl

Unique:	no
Operators:	=
User xRSL:	(acl=<xml>)
A-REX xRSL:	- "-
Example:	(acl="<?xml version='\"1.0\"'?> <gac1 version='\"0.0. 1\"'><entry><any-user></any-user> <allow><write/><read/><list/><admin/></ allow></entry></gac1>")

Makes use of GACL rules to list users who are allowed to access and control job in addition to job's owner. Access and control levels are specified per user. any-user tag refers to any user authorized at the execution cluster. To get more information about GACL please refer to <http://www.gridsite.org>.

xml	a GACL-compliant XML string defining access control list
-----	--

Following job control levels can be specified via acl:

write	allows to modify contents of job data (job directory) and control job flow (cancel, clean, etc.)
read	allows to read content of job data (contents of job directory)
list	allows to list files available for the job (contents of job directory)
admin	allows to do everything – full equivalence to job ownership

queue

Unique:	yes
Operators:	= !=
User xRSL:	(queue=<string>)
A-REX xRSL:	only = is allowed
Example:	(queue="pclong")

The name of the remote batch queue.

Use only when you are sure that the queue by this name does exist.

string	known queue name
---------------	------------------

While users are not expected to specify **queue** in job descriptions, this attribute **must** be present in the Server-side xRSL. In fact, this is primarily an internal attribute, added to the job description by client tools after resource discovery and matchmaking. Still, users can specify this attribute to explicitly force job submission to a queue: when specified explicitly by the user, this value will not be overwritten by the ARC Client, and an attempt will be made to submit the job to the specified queue.

If for some reason (e.g. due to a client tool error) **queue** is absent from the Server-side xRSL, A-REX on the selected cluster will attempt to submit the job to the default queue if such is specified in the A-REX configuration.

startTime

Unique:	yes
Operators:	=
User xRSL:	(startTime=<time>)
A-REX xRSL:	(startTime=<tttt>)
Example:	(startTime="2002-05-25 21:30:00")

Time to start job processing by the Grid Manager, such as e.g. start downloading input files.

time	time string, YYYY-MM-DD hh:mm:ss
tttt	time string, YYYYMMDDhhmmss[Z] (converted by the ARC Client from time)

Actual job processing on a worker node starts depending on local scheduling mechanisms, but not sooner than **startTime**.

lifeTime

Unique:	yes
Operators:	=
User xRSL:	(lifeTime=<time>)
A-REX xRSL:	(lifeTime=<tttt>)
Example:	(lifeTime="2 weeks")

Maximal time to keep job files (the session directory) on the gatekeeper upon job completion.

time	time (in minutes if no unit is specified)
tttt	time (seconds, converted by the ARC Client from time)

Typical life time is 1 day (24 hours). Specified life time can not exceed local settings.

notify

Unique:	yes
Operators:	=
User xRSL:	(notify=<string> [string] ...)
A-REX xRSL:	-"-
Example:	(notify="be your.name@your.domain.com")

Request e-mail notifications on job status change.

string	string of the format: [b][q][f][e][c][d] user1@domain1 [user2@domain2] ...
---------------	--

Here flags indicating the job status are:

b	begin (PREPARING)
q	queued (INLRMS)
f	finalizing (FINISHING)
e	end (FINISHED)
c	cancellation (CANCELLED)
d	deleted (DELETED)

When no notification flags are specified, default value of **eb** will be used, i.e., notifications will be sent at the job's beginning and at its end.

No more than 3 e-mail addresses per status change accepted.

rerun

Unique:	yes
Operators:	=
User xRSL:	(rerun=<integer>)
A-REX xRSL:	-"-
Example:	(rerun="2")

Number of reruns (if a system failure occurs).

integer	an integer number
----------------	-------------------

If not specified, the default is 0. Default maximal allowed value is 5. The job may be rerun after failure in any state for which reruning has sense. To initiate rerun user has to use the `arcresume` command.

architecture

Unique:	no
Operators:	= !=
User xRSL:	(architecture=<string>)
A-REX xRSL:	none
Example:	(architecture="i686")

Request a specific architecture.

string	architecture (e.g., as produced by <code>uname -a</code>)
---------------	--

nodeAccess

Unique:	yes
Operators:	=
User xRSL:	(nodeAccess="inbound" "outbound")
A-REX xRSL:	none
Example:	(nodeAccess="inbound")

Request cluster nodes with inbound or outbound IP connectivity. If both are needed, a conjunct request should be specified.

dryRun

Unique:	yes
Operators:	=
User xRSL:	(dryRun="yes" "no")
A-REX xRSL:	- "-
Example:	(dryRun="yes")

If yes, do dry-run: job description is sent to the optimal destination, input files are transferred, but no actual job submission to LRMS is made. Typically used for xRSL and communication validation.

rsl_substitution

Unique:	no
Operators:	=
User xRSL:	(rsl_substitution=(<code><string1></code> <code><string2></code>))
A-REX xRSL:	-"-
Example:	(rsl_substitution=("ATLAS" "/opt/atlas"))

Substitutes `<string2>` with `<string1>` for **internal** RSL use.

<code>string1</code>	new internal RSL variable
<code>string2</code>	any string, e.g., existing combination of variables or a path

Use this attribute to define variables that simplify xRSL editing, e.g. when same path is used in several values, typically in `inputFiles`. Only one pair per substitution is allowed. To request several substitution, concatenate such requests. Bear in mind that substitution must be defined **prior** to actual use of a new variable `string1`.

After the substitution is defined, it should be used in a way similar to shell variables in scripts: enclosed in round brackets, preceded with a dollar sign, **without quotes**:

```
(inputfiles=("myfile" $(ATLAS)/data/somefile))
```

Unlike the `environment` attribute, `rsl_substitution` definition is only used by the client and is valid inside xRSL script. It can not be used to define environment or shell variable at the execution site.

environment

Unique:	no
Operators:	=
User xRSL:	(environment=(<code><VAR></code> <code><string></code>) [<code><VAR></code> <code><string></code>] ...)
A-REX xRSL:	-"-
Example:	(environment=("ATLSRC" "/opt/atlas/src") ("ALISRC" "/opt/alice/src"))

Defines execution shell environment variables.

<code>VAR</code>	new variable name
<code>string</code>	any string, e.g., existing combination of variables or a path

Use this to define variables at an execution site. Unlike the `rsl_substitution` attribute, it can not be used to define variables on the client side.

count

Unique:	yes
Operators:	=
User xRSL:	(count=<integer>)
A-REX xRSL:	- "-"
Example:	(count="4")

Specifies amount of sub-jobs to be submitted for parallel tasks.

countpernode

Unique:	yes
Operators:	=
User xRSL:	(countpernode=<integer>)
A-REX xRSL:	- "-"
Example:	(countpernode="2")

Specifies amount of sub-jobs per node to be submitted for parallel tasks. Note: The count attribute must be specified when this attribute is specified.

exclusiveexecution

Unique:	yes
Operators:	=
User xRSL:	(exclusiveexecution="yes" "no")
A-REX xRSL:	- "-"
Example:	(exclusiveexecution="yes")

Specifies whether the node should be allocated for exclusive use by the job.

jobreport

Unique:	yes
Operators:	=
User xRSL:	(jobreport=<URL>)
A-REX xRSL:	- "-"
Example:	(jobreport="https://grid.uio.no:8001/logger")

Specifies an URL for an accounting service to send reports about job to. The default is set up in the cluster configuration.

URL	URL
-----	-----

It is up to users to make sure the requested accounting service accepts reports from the set of clusters they intend to use.

credentialserver

Unique:	yes
Operators:	=
User xRSL:	(credentialserver=<URL>)
A-REX xRSL:	-"-
Example:	(credentialserver="myproxy://myproxy.nordugrid.org;username=user")

Specifies an URL which Grid Manager may contact to renew/extend delegated proxy of job. Only MyProxy servers are supported.

URL	URL of MyProxy server
-----	-----------------------

It is up to a user to make sure the specified MyProxy server will accept requests from Grid Manager to renew expired credentials. URL may contain options `username` and `credname` to specify user name and credentials name which Grid Manager should pass to MyProxy server. If `username` is not specified DN of user credentials is used instead.

priority

Unique:	yes
Operators:	=
User xRSL:	(priority=<integer>)
A-REX xRSL:	-"-
Example:	(priority="80")

Specifies priority given to this job during staging of input and output files when the new data staging framework is used by A-REX. Values are limited to between 1 (lowest priority) and 100 (highest priority). Default if this attribute is not specified is 50.

9.5.3 Server-side attributes

Note

It is strongly advised to avoid specifying following attributes in the client xRSL.

The following attributes are a part of the internal ARC job representation, and must be provided by ARC client tools and passed to the A-REX. Developers of new ARC client tools and utilities must make sure these attributes are added to the user job description before it is submitted to a A-REX.

sstdin

A-REX xRSL:	(sstdin=<filename>)
Example:	(sstdin="myinput.dat")

Internal attribute for the standard input. Can also be spelled `stdin`. Only needed for GRAM compatibility, not used by ARC as such.

- `filename` - standard input file name

action

A-REX xRSL:	(action="request" "cancel" "clean" "renew" "restart")
Example:	(action="request")

Action to be taken by the gatekeeper: submit the job, cancel job execution, clear the results of the job (also cancels the job), renew the proxy of the job, or restart the job from a previous failed state.

savestate

A-REX xRSL:	(savestate="yes" "no")
Example:	(savestate="yes")

If yes, input RSL is stored in a temporary file at the gatekeeper. Must be always set as yes in the current implementation. Only needed for GRAM compatibility, not used by ARC as such.

lrmstype

A-REX xRSL:	(lrmstype=<string>)
Example:	(lrmstype="pbs")

LRMS type, indicating which submission script is to be invoked.

- string - LRMS type

hostName

A-REX xRSL:	(hostname=<string>)
Example:	(hostName="grid.quark.lu.se")

Name (e.g. as returned by the Linux `hostname` command) of the client machine from which the submission was made.

- string - client host name, as passed by the ARC client

jobid

A-REX xRSL:	(jobid=<string>)
Example:	(jobid="grid.quark.lu.se:2119/jobmanager-ng/157111017133827")

Unique job identification string, needed for cancellation and clean-up.

- string - global job ID

It can also be provided during submission of the job and should be unique to a computing element (cluster).

clientxrs1

A-REX xRSL:	(clientxrs1=<string>)
Example:	(clientxrs1="&(executable=/bin/echo)(arguments=boo)")

Job description xRSL string as submitted by the user, before being pre-processed by the client.

- **string** - original xRSL description submitted by the user

This attribute is added by the User Interface during pre-processing, and is used for job re-submission in order to repeat brokering and matchmaking.

clientsoftware

A-REX xRSL:	(clientsoftware=<string>)
Example:	(clientsoftware="nordugrid-arc-0.5.39")

Version of ARC client used to submit the job.

- **string** - version string

This attribute is added by the User Interface during pre-processing.

delegationid

A-REX xRSL:	(delegationid=<id>)
Example:	(delegationid="a6f3e8fc920de023fbb432ea520")

X.509 delegation identifier for all data staging operations and job itself.

- **id** - delegation identifier obtained during delegation procedure (for applicable job submission interfaces)

This attribute is added by the User Interface during pre-processing.

9.5.4 Examples

User-side xRSL script

```
&
(* test run: if "yes", only submits RSL without actual job start *)
  (dryRun="no")
(* some local variables defined for further convenience *)
  (rsl_substitution=("TOPDIR" "/home/johndoe"))
  (rsl_substitution=("NGTEST" "${TOPDIR}/ngtest))
  (rsl_substitution=("BIGFILE" "/scratch/johndoe/100mb.tmp"))
(* some environment variables, to be used by the job *)
  (environment=("ATLAS" "/opt/atlas") ("CERN" "/cern"))
(* the main executable file to be staged in and submitted to the PBS *)
  (executable="checkall.sh")
(* the arguments for the executable above *)
  (arguments="pal")
(* files to be staged in before the execution *)
  (inputFiles = ("be_kaons" ""))
```

(continues on next page)

(continued from previous page)

```

    ("file1" gsiftp://grid.uio.no$(TOPDIR)/remfile.txt)
    ("bigfile.dat" $(BIGFILE) )    )
(* files to be given executable permissions after staging in *)
    (executables="be_kaons")
(* files to be staged out after the execution *)
    (outputFiles=
        ("file1" "gsiftp://grid.tsl.uu.se/tmp/file1.tmp")
        ("100mb.tmp" "rls://rls.nordugrid.org:39281/test/bigfile")
        ("be_kaons.hbook" gsiftp://ce1.grid.org$(NGTEST)/kaons.hbook)    )
(* user-specified job name *)
    (jobName="NGtest")
(* standard input file *)
    (stdin="myinput.dat")
(* standard output file *)
    (stdout="myoutput.dat")
(* standard error file *)
    (stderr="myerror.dat")
(* A-REX logs directory name *)
    (gmlog="gmlog")
(* flag whether to merge stdout and stderr *)
    (join="no")
(* request e-mail notification on status change *)
    (notify="bqfe john.doe@gmail.com jane.doe@mail.org")
(* maximal CPU time required for the job, minutes for PBS*)
    (CpuTime="60")
(* maximal time for the session directory to exist on the remote node, days *)
    (lifeTime="7")
(* memory required for the job, per count, Mbytes *)
    (Memory="200")
(* wall time to start job processing *)
    (startTime="2002-04-28 17:15:00")
(* disk space required for the job, Mbytes *)
    (Disk="500")
(* required architecture of the execution node *)
    (architecture="i686")
(* required run-time environment *)
    (runTimeEnvironment="APPS/HEP/Atlas-1.1")
(* number of re-runs, in case of a system failure *)
    (rerun="2")

```

Server-side xRSL script

Note that a client tool must do matchmaking and modify correspondingly the job document before submitting it to the matching resource. Specifically, a client tool has to:

- expand all the `rsl_substitution` values
- add double quotes to **all** strings
- insert queue attribute in case such is missing
- make sure every logical or comparison operator is expanded and replaced with a deterministic `=` statement
- streamline blank spaces

Note

Comment lines are removed from the Server-side xRSL! Below they are shown to explain details.

```
&
(* saves RSL in a temporary file if "yes" *)
    ("savestate" = "yes" )
(* job submission to be performed if action is "request" *)
    ("action" = "request" )
(* submission host name *)
    ("hostname" = "cel.grid.org" )
(* client software version *)
    ("clientsoftware" = "nordugrid-arc-0.6.0.3" )
(* walltime value added by the client, in seconds *)
    ("walltime" = "3600" )
(* test run: if "yes", only submits RSL without actual job start *)
    ("dryRun" = "no" )
(* some local variables defined for further convenience *)
    ("rsl_substitution" = ("TOPDIR" "/home/johndoe" ) )
    ("rsl_substitution" = ("NGTEST" "/home/johndoe/ngtest" ) )
    ("rsl_substitution" = ("BIGFILE" "/scratch/johndoe/100mb.tmp" ) )
(* some environment variables, to be used by the job *)
    ("environment" = ("ATLAS" "/opt/atlas" ) ("CERN" "/cern" ) )
(* executable *)
    ("executable" = "checkall.sh" )
(* arguments *)
    ("arguments" = "pal" )
(* files to be staged in before the execution *)
    ("inputfiles" = ("checkall.sh" "279320" )
        ("myinput.dat" "39806" )
        ("be_kaons" "8807" )
        ("file1" "gsiftp://grid.uio.no/home/johndoe/remfile.txt" )
        ("bigfile.dat" "104857600" )
    )
(* files to be given executable permissions after staging in *)
    ("executables" = "checkall.sh" "be_kaons" )
(* files to be staged out after the execution *)
    ("outputfiles" = ("file1" "gsiftp://grid.tsl.uu.se/tmp/file1.tmp" )
        ("100mb.tmp" "rls://rls.nordugrid.org:39281/test/bigfile" )
        ("be_kaons.hbook" "gsiftp://cel.grid.org/home/johndoe/ngtest/kaons.hbook" )
        ("myoutput.dat" "" )
        ("myerror.dat" "" )
    )
(* user-specified job name *)
    ("jobName" = "NGtest" )
(* standard input file *)
    ("stdin" = "myinput.dat" )
(* standard output file *)
    ("stdout" = "myoutput.dat" )
(* standard error file *)
    ("stderr" = "myerror.dat" )
(* flag whether to merge stdout and stderr *)
    ("join" = "no" )
(* request e-mail notification on status change *)
    ("notify" = "bqfe john.doe@gmail.com jane.doe@mail.org" )
(* specific queue to submit the job *)
    ("queue" = "atlas" )
(* CPU time required for the job, converted into seconds *)
```

(continues on next page)

(continued from previous page)

```
( "cputime" = "3600" )
(* maximal time for the session directory to exist on the remote node, seconds *)
( "lifetime" = "604800" )
(* memory required for the job, per count, Mbytes *)
( "memory" = "200" )
(* wall time to start job processing *)
( "startTime" = "20020128171500" )
(* disk space required for the job, Mbytes *)
( "disk" = "500" )
(* required architecture of the execution node *)
( "architecture" = "i686" )
(* required run-time environment *)
( "runtimeenvironment" = "APPS/HEP/Atlas-1.1" )
(* number of re-runs, in case of a system failure *)
( "rerun" = "2" )
(* original client xRSL with expanded string substitutions; shortened here *)
( "clientxrsl" = "&(" "dryrun" = "no" ) ( ... ) ( "rerun" = "2" ) )
```

9.6 ARC Client Config Reference

Some default valued used by *client tools* can be refedined in the client configuration file located in `~/.arc/client.conf`.

All currently supported client configuration options should be defined in the [common] block.

9.6.1 certificatepath

Synopsis: certificatepath = path

Description: Specify the location of client certificate file. Environmental variable X509_USER_CERT redefines this value.

Default: \$HOME/.globus/usercert.pem

Example:

```
certificatepath=/home/user/credentials/cert.pem
```

9.6.2 keypath

Synopsis: keypath = path

Description: Specify the location of client secret key file. Environmental variable X509_USER_KEY redefines this value.

Default: \$HOME/.globus/userkey.pem

Example:

```
keypath=/home/user/credentials/key.pem
```

9.6.3 cacertificatesdirectory

Synopsis: cacertificatesdirectory = path

Description: Specify the location of CA certificates directory Environmental variable X509_CERT_DIR redefines this value.

Default: /etc/grid-security/certificates

Example:

```
cacertificatesdirectory=/home/user/.globus/certificates
```

9.6.4 proxypath

Synopsis: proxypath = path

Description: Specify the location of proxy certificate (both for generation and usage) Environmental variable X509_USER_PROXY redefines this value.

Default: /tmp/x509up_u\${UID}

Example:

```
proxypath=/home/user/credentials/proxy.pem
```

9.6.5 vomsespath

Synopsis: vomsespath = path

Description: Path to file or directory that holds client VOMS configuration to generate proxy certificates

Environmental variables X509_VOMS_FILE and X509_VOMSES redefine this value If missing arcproxy will search for vomses in the following locations:

- ~/.arc/vomses
- ~/.voms/vomses
- /etc/vomses
- /etc/grid-security/vomses

Default: undefined

Example:

```
vomsespath=/home/user/credentials/vomses
```

9.6.6 defaultvoms

Synopsis: defaultvoms = vo[:command]

Description: Default value for --voms (-S) arcproxy option that is used to define VO and optionally FQANs used during proxy certificate generation

This option is **multivalued**.

Default: undefined

Example:

```
defaultvoms=atlas:/atlas/Role=pilot
defaultvoms=nordugrid.org:all
defaultvoms=ops.ndgf.org
```

9.6.7 rejectdiscovery

Synopsis: rejectdiscovery = service

Description: Specify the FQDN or URLs of the services that should be rejected during service discovery process by CLI tools (arcsb, arctest, arcresb)

This option in **multivalued**.

Default: undefined

Example:

```
rejectdiscovery=bad.service.org
rejectdiscovery=bad2.service.org
```

9.6.8 rejectmanagement

Synopsis: rejectmanagement = service

Description: Specify the FQDN or URLs of the CEs that should be skipped during the job management (e.g. arcstat, arckill)

This option in **multivalued**.

Default: undefined

Example:

```
rejectmanagement=bad3.service.org
rejectmanagement=bad4.service.org
```

9.6.9 brokername

Synopsis: brokername = broker

Description: Specify the broker used in resource discovery. The full list of installed brokers can be obtained running arcsb -P

Default: Random

Example:

```
brokername=FastestQueue
```

9.6.10 brokerarguments

Synopsis: brokerarguments = args

Description: Specify broker arguments (if applicable to specified broker)

Default: undefined

Example:

```
brokername=PythonBroker
brokerarguments=ACIXBroker.ACIXBroker:https://cacheindex.ndgf.org:6443/data/index
```

9.6.11 timeout

Synopsis: timeout = seconds

Description: Amount of time to wait for a service to respond before considering it dead.

Default: 20

Example:

```
timeout=60
```

9.6.12 joblist

Synopsis: joblist = path

Description: Path to the jobs database that holds all extra data about submitted jobs to be used during further job management

Default: \$HOME/.arc/jobs.dat

Example:

```
joblist=/home/user/arcjobs.dat
```

9.6.13 joblisttype

Synopsis: joblisttype = type

Description: Type of the backend used for jobs database.

Warning

IT IS STRONGLY advised to keep default SQLITE type of backend.

Allowed values: SQLITE,, XML

Default: SQLITE

Example:

```
joblisttype=XML
```

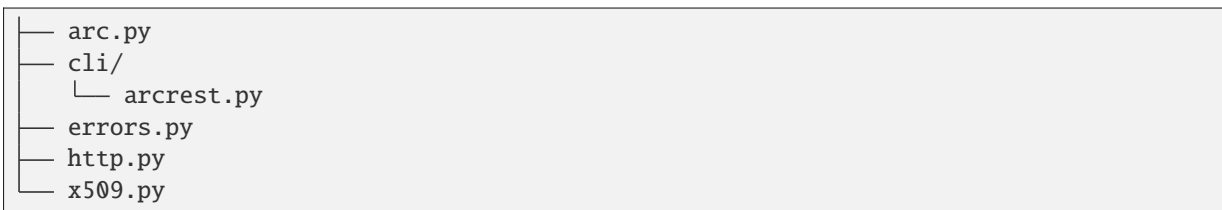
9.7 ARC SDK Documentation

9.8 PyARCREST

PyARCREST is a python client for interacting with the ARC CE REST interface. It offers a high-level interface for managing jobs, delegations, file transfers, and certificate operations. With pyARCREST you can perform various operations such as creating, monitoring, and manipulating jobs, as well as handling errors and certificate operations.

The core functionalities of the pyARCREST include: - Creating and managing jobs within the ARC CE REST interface. - Handling job status and information retrieval. - Facilitating file transfers and delegation operations. - Offering a command-line interface for executing REST API service operations. - Providing error handling for job operations and file handling. - Supporting customizable HTTP client classes for sending requests to specified URLs. - Enabling X.509 certificate operations such as creating, signing, and parsing certificates.

9.8.1 STRUCTURE



9.8.2 MODULES

Module arc.py

Description: This module provides a high-level interface for interacting with the ARC CE REST interface, supporting automatic handling of multiple API versions.

It includes methods for

- creating jobs
- managing job status and info
- file transfers,
- delegation operations

Relations with other modules

This module interacts with other modules by making requests to the *http* module for HTTP operations, handling errors through the *errors* module, and performing certificate operations through the *x509* module for signing CSR and handling PEM files. The *http* module handles the actual HTTP requests and responses, while the *errors* module deals with error handling. The *x509* module is used for certificate-related operations.

Module cli/arcrest.py

This module is a command-line interface for executing various operations on a REST API service. It allows users to interact with resources such as jobs and delegations by providing commands like getting info, listing, submitting, killing, renewing, and deleting, among others.

Relations with other modules

This module interacts with the ARCRest module by calling various methods such as getJobsList, getJobsInfo, getJobsStatus, killJobs, cleanJobs, getAPIVersions, restartJobs, getCEInfo, createDelegation, getDelegation, refreshDelegation, and deleteDelegation based on the provided arguments in the main function.

Module errors.py

This module defines various error classes related to job operations and file handling within an ARC system, including parsing errors, matchmaking issues, and missing result files.

Module http.py

This module defines a customizable HTTP client class that handles sending requests with different methods, headers, and parameters to a specified URL. It also supports HTTPS connections and error handling for connection issues.

Module x509.py

This module provides functions to handle X.509 certificates, including checking if a certificate is in the old format, creating and signing proxy certificates, parsing PEM files, generating keys, and converting between different formats.

9.8.3 CLASSES

class Result [arc.py]:

Stores a value and an error state.

class ARCRest [arc.py]:

Handles REST API interactions for ARC middleware jobs.

class ARCRest_1_0 [arc.py]:

Handles job creation, submission, and delegation list retrieval for API.

class ARCRest_1_1 [arc.py]:

Performs actions related to creating and submitting jobs for processing.

class Transfer [arc.py]:

Manages transfer job information and cancellation event handling.

class ARCJob [arc.py]:

Maps and updates ARC job information from external sources.

class TransferQueue [arc.py]:

Synchronize access to a shared queue using barriers and locks.

class TransferQueueEmpty [arc.py]:

Custom exception raised when transfer queue is empty.

class X509Error [errors.py]:

Base error class for X.

class HTTPClientError [errors.py]:

Custom error for handling HTTP client errors.

class ARCError [errors.py]:

Base error class for arc module exceptions.

class ARCHTTPErrors [errors.py]:

Represents an HTTP error for job operations.

class DescriptionParseError [errors.py]:

Handles errors related to description parsing.

class DescriptionUnparseError [errors.py]:

Handles errors occurring during description unparsing.

class InputFileError [errors.py]:

Handles errors related to input file in job descriptions.

class NoValueInARCResult [errors.py]:

Custom error representing a missing value in ARC result.

class MatchmakingError [errors.py]:

Handles errors related to matchmaking process within the application.

class InputUploadError [errors.py]:

Handles errors that occur during input file upload process.

class MissingResultFile [errors.py]:

Represents an error for missing job result file.

class MissingOutputFile [errors.py]:

Represents a missing job output file for error handling.

class MissingDiagnoseFile [errors.py]:

Represents a missing job diagnose file for error handling.

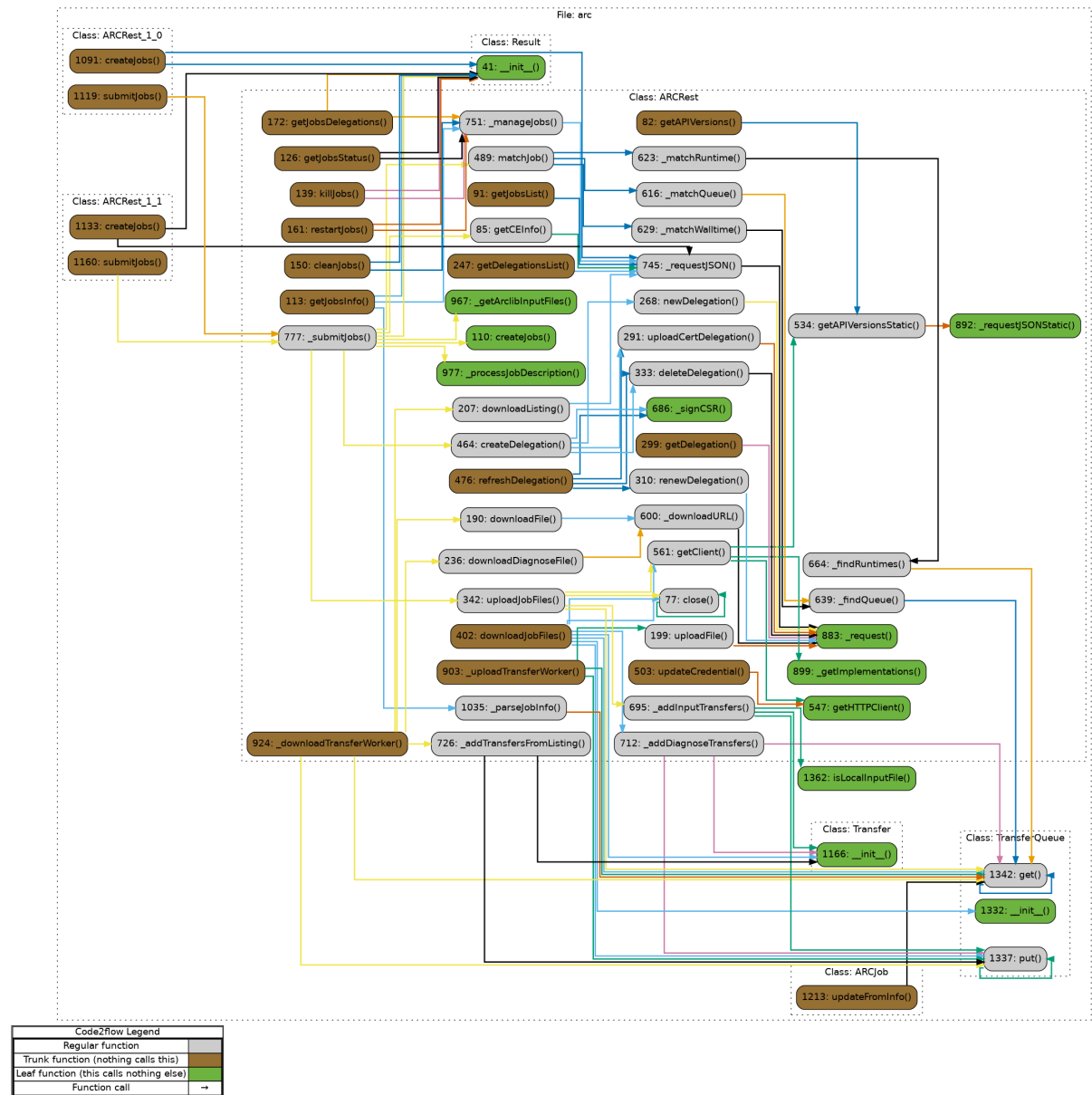
class HTTPClient [http.py]:

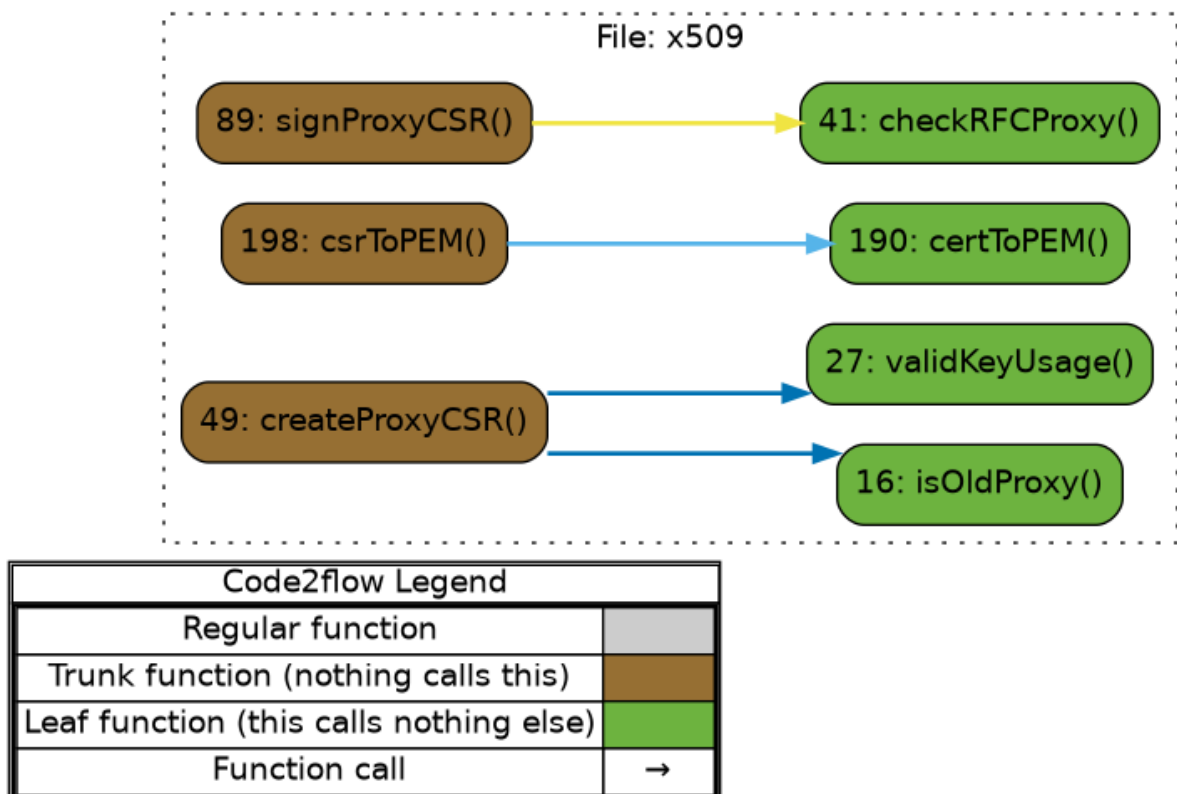
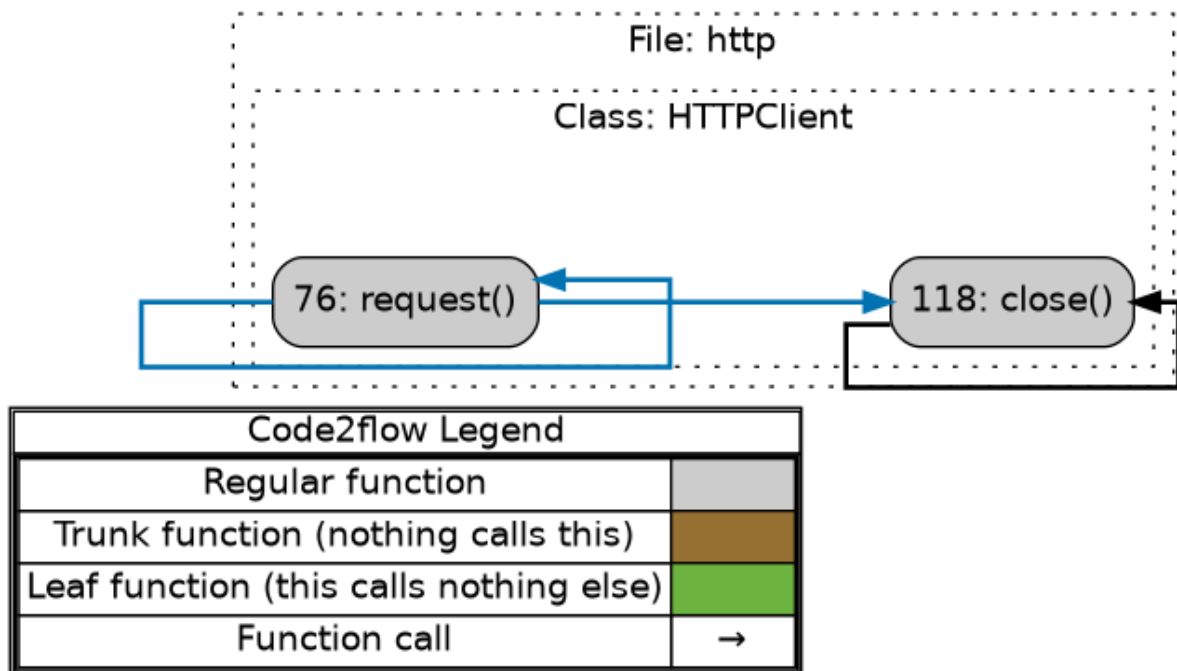
Connect and communicate with HTTP and HTTPS servers with optional proxy.

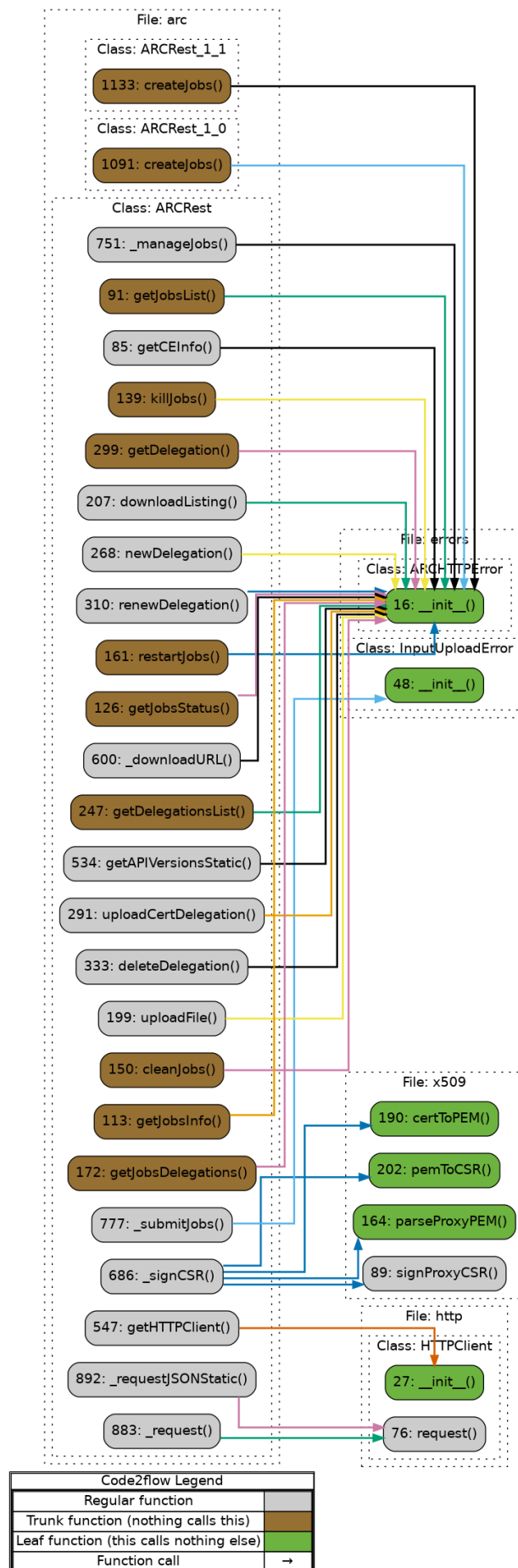
9.8.4 EXECUTION FLOWS

MODULES

MODULE RELATIONSHIPS







ARC MISCELLANEOUS PAGES

These pages contain information that does not fit in the formal ARC documentation. That can be for instance more dynamic contents, like overview of testing campaigns related to releases, or how-to pages.

10.1 About the Nordugrid ARC Releases

This page collects general information concerning the releases of the Nordugrid ARC software.

An ARC release is defined as the ARC source package AND the corresponding binary packages on the *supported platforms*.

10.1.1 Release categories

The source package is identified by its major, minor and patch version number. The major and minor releases are scheduled releases, the patch releases are unplanned fast-track releases to fix specific unexpected problems with the previous minor release.

An example of a scheduled release is NorduGrid ARC 7.1.0 where 7 refers to the major release series, and 1 to the minor. Whenever there is a specific patch release, addressing only one isolated issue with the previous minor release, the last digit is used, for instance 7.1.1.

The properties of a major are:

- may break backward compatibility
- introduces new components, features
- obsoletes components
- has longer term planning
- 3-6 months preparation
- can include alpha, beta, rcx test releases
- bumps the major number in the version number
- the contents mainly follow a high level ARC development roadmap

The properties of a minor release are:

- includes bugfixes
- can include new features and/or enhancements
- can include new components
- does not include anything that is backward incompatible
- the release bumps the minor number in the version number

The Nordugrid ARC releases can come in three types: scheduled, express and binary update of a release.

Scheduled releases:

- Minor releases are scheduled regularly monthly or bi-monthly
- Major releases are scheduled irregularly as needed, with frequency ranging from 1 to three years.
- All developments that at the time of the release are already merged into master are included in the scheduled release
 - this means that backward incompatible changes are kept in separate development branches, and only merged before a major release is imminent.

Express releases:

- is an unplanned urgent release to fix a security issue or a critical bug
- maximum two weeks preparation
- only development needed to fix the issue is included into the release in order to not delay the release or introduce new bugs
- a separate branch is created from the point on master of the last release, and the release tag is created on this branch
- bumps the third digit in the ARC release number by one: 7.1.0 -> 7.1.1.

Binary update of a release:

- is sometimes needed due to a change in the external dependencies and ARC releases have to be rebuilt
- this “semi release” does not affect the ARC source code, no re-tagging takes place, only new binaries are built
- can happen for both major, minor or emergency releases

10.1.2 Testing

Testing needs to be done on several Linux distributions as ARC depends on a lot of software that has different versions on different distributions with different file structure.

The level of importance of the different Linux distributions is evaluated as follows:

1. Rocky/AlmaLinux/CentOS Stream version 9.x architecture x86_64 is The Most Important version.
2. Latest Ubuntu LTS, Latest Ubuntu, any architecture
3. Fedora and other Ubuntu
4. The rest.

Some OSes are more used server-side (RHEL-flavour) and some client-side (Fedora).

The testing of Nordugrid ARC is performed in levels, where level 1 and 2 are always performed, and level 3 is performed on the best-effort or as-needed -basis.

Level 1

All supported platforms are built on a nightly basis and [publically available](#). Follow [these instructions](#).

In addition ARC is built and deployed for the most important platforms on the [GitLab CI platform](#). The deployment includes a basic functional test. Both supported platforms and tests are subject to ongoing development. Currently the build and deploy is performed on each commit to master.

Level 2

The release candidate is deployed on volunteer sites (minimally 1) and on testing infrastructure available at the time. In particular the functionality affected by the development in the current release should be tested.

The test-site tracks his contribution to the test on [in the GitLab Wiki](#).

Level 3

If necessary and possible, larger testing campaigns are organized using the release candidate in question (involving users and sysadmins outside the testing team). This campaign will follow a plan organized centrally by the Nordugrid ARC team. The sites report back on any issues seem, and may be asked to perform tests, and give a simple report back on the test results.

The tests and their results are recorded on a dedicated pages in the [Testing area of the ARC documentation](#)

10.1.3 Release notes

Release notes are published [here in our docs](#) and <http://www.nordugrid.org/arc/releases/> and are distributed in the following channels:

- News section at <http://www.nordugrid.org>
- RSS feed: <http://www.nordugrid.org/news.xml>
- RSS feed: <http://www.nordugrid.org/arc/release-news.xml>
- LinkedIn Nordugrid page: <http://www.linkedin.com/company/nordugrid>

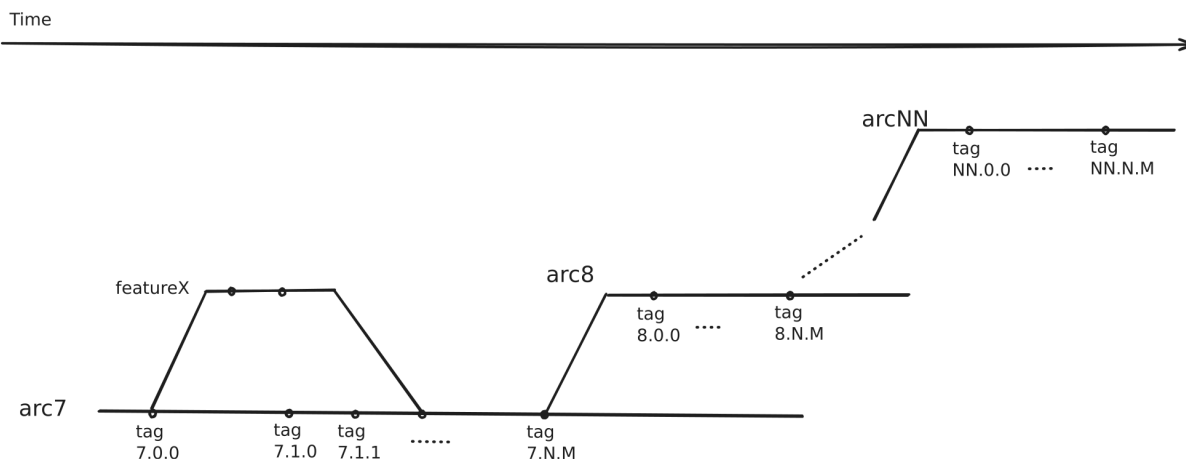
The following Web pages are updated upon a release:

- ARC Web page: <http://www.nordugrid.org/arc/>
- ARC releases: <http://www.nordugrid.org/arc/releases/> (updated automatically when release notes are uploaded)
- Nordugrid repositories, in case new one is created: [repository document within this documentation](#) and <http://download.nordugrid.org/repos-6.html>
- Wikipedia article: http://en.wikipedia.org/wiki/Advanced_Resource_Connector

10.1.4 ARC Git branching strategy

As of ARC 7 we do no longer have any master branch. The arc7 branch is the new default branch. The arc6 branch is the previous master branch containing the ARC 6 source code.

Our branching strategy:



Default branch is whichever is the current major release: arc7 or arc8 etc. This is set in GitLab and No main (or master) or next branch.
 We create a new arcX branch once development starts for the next major release (as before)
 We never merge arcX for next major release into current major release (as before)
 (no merge next into master, or in new naming convention no merging arc8 into arc7).
 Tagging for a release happens on the current default branch (as before)
 We create feature branches as before from central repo or from forks (as before).

10.2 ARC 7 Testing Area

There are 1 slurm test-sites set up for users to test e.g. token submission or ARCREST interface to:

- <https://source.coderefinery.org/nordugrid/arctestsite-slurm-el9-arc7>

Instructions on how to interact with this ARC-CE site can be found in the repos README file.

In addition there is 1 condor site which can receive jobs, but does not have a repo page: - arctestcluster-condor-almalinux9-arc7-ce1.cern-test.uiocloud.no

Contact us in one of our support channels if you want to do any testing on any of these sites.

10.3 Getting in touch

- ARC support chat: We have moved to Matrix: https://matrix.to/#/#arc_support:matrix.org
- Nordugrid ARC email list
 - Discussion forum (general support): nordugrid-discuss@nordugrid.org
 - (Join/leave by sending a mail to nordugrid-discuss-join@nordugrid.org/nordugrid-discuss-leave@nordugrid.org)

10.4 Release testing

- ARC-6.2.0
- ARC-6.3.0

10.5 Release management

10.5.1 People involved in preparing for a release

Table 10.1: People

Name	Experties/responsibility
Anders Waananen	Packaging, autotools, admin access to Bugzilla and nordugrid.org (package downloads http://download.nordugrid.org/repos/6/)
Balazs Konya	Overall coordination
Mattias Ellers	packaging and builds, EPEL and Debian uploads (Globus)
Maiken Pedersen	Release manager
Oxana Smirnova	Bugzilla, Web site, documentation

10.5.2 Release workflow

- Decide on release date and timeline on Nordugrid technical coordination weekly meeting
 - The aim is a monthly scheduled minor release
 - The timeline contains code-freeze date. From this date on no new merge requests will as a rule be added to master.
 - Together with the developers a set of minimal tests are defined for the current release in planning.
 - Clarify if any documentation changes are needed.
 - The release manager enters the test-descriptions on the GitLab wiki test-page: <https://source.coderefinery.org/nordugrid/arc/wikis/home>
- Release manager sends out a heads-up email announcing the release in preparation and its timeline right after the decision of the release is made
 - to: nordugrid-discuss
 - Inform that documentation updates should be ready by code-freeze date
- The day after code-freeze, the agreed test-sites should install the nightly build from the nightly build repository: <http://download.nordugrid.org/builds/index.php?pkgname=nordugrid-arc&type=master> and report on the tests and any issues encountered
 - If there are any issues the release date might be postponed.
 - * If a postponement is necessary a new timeline is drawn
 - * TO-DO: Should a new email go out to nordugrid-discuss or will that just be noise?
 - If no issues are found the release is ready for the next step
- After the definite code-freeze the translations are run
 - The release manager gives a green light to run translations
 - Translations are run. Responsible: Oxana Smirnova

- The release manager creates a tag on master
 - Informs the build-manager about the commit-hash either on skype or by email to nordugrid-core mail list
- The builds are created and repositories populated. Responsible: Anders Waananen
- The list of supported platforms is checked for updates. Responsible: Anders Waanaen
- Once the source tarballs are available in the nordugrid repo, the build in the Fedora build system and for Debian is started. Responsible: Mattias Ellert.
 - The builds should not be made available before the release is announced. TO-DO: Is this possible?
- Prepare the release announcement (see more below). Responsible: Maiken Pedersen
- Notify release manager that the builds are ready to be pushed to the repo. Responsible: Anders Waananen
- Notify responsible in order to publish the release announcement on the web. Responsible: Maiken Pedersen
 - Publish the release announcement. Responsible: Oxana Smirnova
- Notify responsible in order to push the packages to the nordugrid repo. Responsible: Maiken Pedersen
 - Push packages to nordugrid repo. Responsible: Anders Waananen
- Notify responsible in order to push builds to Fedora and Debian repo. Responsible: Maiken Pedersen
 - Push packages to Fedora and Debian repos. Responsible: Mattias Ellers
- Ensure there are consistent binary packages in agreed repositories and linked from the release page. Responsible: Maiken Pedersen
- Ensure there is corresponding release version in Bugzilla. Responsible: Anders Wannanen

10.5.3 The release announcement should contain

- Most important changes/highlights
- List of fixed bugs
- List of new features (if relevant)
- List of backward incompatible changes (if relevant - only for major releases)
- Link to documentation in addition to installation links
- Link to the GitLab ARC repo
- Getting-in-touch information

10.6 Changelogs/list of bugs

10.6.1 Main changes in ARC 7 compared to ARC 6

ARC 7 represents the next major release of the Advanced Resource Connector Middleware building upon the technology preview components that were already included in the previous major ARC 6 version released in 2019 and has been developed further during the past years. ARC 7 consolidates the technology developments that took place on the ARC 6 line of the middleware. ARC 7 continues streamlining the ARC CE interfaces and the codebase by implementing intensive functionality and code cleanup. In ARC 7 the REST interface is promoted as the primary job management and information retrieval interface of the ARC CE and token-based authentication and authorization is treated on equal level with X509 proxies. ARC 7, despite being a major release, remains as much backward compatible as possible with ARC 6 CEs. An ARC6 server configuration, modulo the deprecated components is expected to work with ARC 7 release.

Note

The primary interface of ARC 7 CE is now the REST 1.1, all the other older versions of WS interfaces and the legacy gridftp job submission interface are obsoleted. The LDAP-based CE information interface is still available, though marked as legacy. ARC 6 clients can only submit and manage jobs on ARC 7 CE via the REST interface while ARC 7 clients are able to submit and manage jobs on ARC 6 CEs using both the REST interface and the deprecated gridftp jobsubmission interface. For the later, the legacy globus plugin should be installed.

Full list of changes in ARC 7**Server-side configuration changes ARC CE configuration**

The server-side configuration has undergone smaller clean up, some attributes got regrouped and some configuration blocks got removed due to obsoleted functionality. New attributes were also introduced. No renaming of configuration parameters took place. All the config changes are listed below:

- obsoleted and removed configuration blocks: [authtokens], [lrms/ssh], [arex/ws/publicinfo], [arex/ws/argus], [arex/jura/archiving], [gridftpd], [gridftpd/jobs], [gridftpd/filedir], [infosys/glue1], [infosys/glue1/site-bdii], [acix-scanner], [acix-index], [userlist:name], [nordugridmap]
- new configuration block: [infosys/accesscontrol]
- new configuration parameters: authtokens and authokensgen in [authgroup], tokenscopes and authokenmap in [arex], tlsserverorder in [arex/ws]
- new or modified parameter values: modified allowed values for lrms type in [lrms], new supported environment variables for the plugin parameter in [authgroup], benchmark in [lrms], copyurl in [arex/data-staging]
- obsoleted and removed configuration parameters: interactive_contactstring in [infosys/cluster], userlist in [authgroup], use_remote_acix in [arex/data-staging]

Internal restructuring and improvements

- Major redesign of how *A-REX organizes and stores its internal information about jobs* in specific files located under the A-REX control directory.
- Added support for conf.d type ARC configuration: if an arc.conf.d directory exists next to arc.conf file, all files in this directory ending with .conf will be read in alphabetical order and their content merged with the arc.conf file.
- Added support for using the system CAs instead of traditional IGTF CAs, see the *x509_cert_policy* option in the [common] block.

REST 1.1 is the primary ARC CE interface

- REST is used exclusively for both job management and information retrieval. All other previous CE interfaces are deleted or labelled OBSOLETE in ARC 7.
- The REST interface got extended to support bulk job submission of identical job. The optional instances parameter can be used to request the creation of <number> of identical jobs defined by the same identical job description document in the request body. The arcsub cli also got new options -instances-max and -instances-min to specify those numbers.
- JWT token delegation is now supported in REST 1.1. The Delegation protocol depends on the delegation type (x509 or jwt) specified with an optional type parameter. If not explicitly specified the default delegation type is x509 for backward compatibility with REST API 1.0.

Token-based authentication and authorisation

- Token based authentication is now enabled by default together with x509 proxy based authentication. In ARC 7 tokens are treated on equal basis with x509 proxies. As a consequence, the previously available [authtokens] block in arc.conf got obsoleted and removed.
- Extended support for using token information (scopes) for authorization and identity mapping. Added support for managing VO information via tokens.
- *arcctl* comes with a built-in test token issuer that is very useful to testing ARC CE installation with tokens.

Information interfaces and information indices in ARC 7

- REST is the primary interface to obtain CE and job information. EMIES and all the other previous WS-based information publishing is removed. Server-side LDAP is kept but labelled OBSOLETE.
- *ARCHERY* is the only remaining supported information index service in ARC 7
- For backward compatibility with monitoring systems the LDAP-based server side information publishing is kept in ARC 7 nevertheless this information system technology is now labelled as OBSOLETE and target for removal in an upcoming release.
- Job states are only published via the REST interface and those are conforming to either the REST states or the nordugrid schema values.
- The publicinfo xml cluster document is by default available without the need of authentication (as the information is regarded as public). At the same time a new block [infosys/accesscontrol] enables configuration of access to information via the usage of the denyaccess and allowaccess parameters for more control.

Accounting subsystem changes

- The accounting subsystem in ARC7 is updated to improve operational efficiency.
- ARC 7 comes with a redesigned database schema for better handling of production use-cases, cleanup, backup, rotation, improved troubleshooting and performance optimization features
- On upgrade ARC7 will immediately start writing data to the new accounting database without blocking startup and job submission.
- CE admins can migrate records from the old database to the new one at their convenience with *arcctl* , minimizing disruption.
- Republishing with *arcctl* works with both old and new databases.
- Project information is now configurable for APEL accounting
- More consistent handling of benchmark information across various ARC configuration blocks.
- Token information is now part of the accounting records.

ARC CE LRMS support

- With ARC 7 it is still required that the computing cluster behind the ARC CE comes with a batch or LRMS system deployed.
- The fully supported thus recommended LRMS options for ARC 7 are the SLURM and Condor batch systems. Fork as a single node simplistic fall-back solution is also available and fully supported.
- LRMS integration with further LRMS systems such as Boinc, SGE, PBS/PBSPro/OpenPBS, LSF and Loadleveler (LL) are provided as part of ARC 7 however only on a best effort community basis.
- In ARC 7 the Batch system modules are distributed in two separate packages, the “supported” and the “contrib”
- The support for Loadleveler and LSF are labelled as OBSOLETE and removal of these components are expected in an upcoming ARC release.
- The experimental Python-based LRMS implementation, including the remote-SSH LRMS backend have been removed from ARC.

Software package changes

- New package containing the python-based REST client (python3-arcrest)
- New package for the community supported LRMS systems (arc-arex-lrms-contrib)

ARC client-side components

- no new development took place in the arc cli commands apart from the removal of the non-supported interfaces and some minor synchronization of cli behaviour and cli documentation.
- *ARC Control Tool (arccctl)*, the recommended server-side system administration tool to work with various ARC and third-party components, has been under continuous upgrade to meet all the introduced ARC7 internal and external changes. arccctl offers a convenient tool for ARC site admins by wrapping typical complex operations and unifying ARC CE administration under a single utility.
- *pyarcrest*, a python tool to utilize the ARC REST interface has been added to the ARC 7 release bundle. Get it from the python3-arcrest package.

Compatibility with ARC 6 components

- ARC6 clients can submit jobs to ARC7 CEs via the ARC 7 REST interface. For this one must install the arcrest plugin package for ARC 6
- ARC7 clients will work only with REST-based ARC6 CEs.
- Data staging between ARC7 and ARC6 servers should work without any problem.
- An arc.conf file from an ARC 6 CE should work with an ARC 7 CE modulo the deprecated configuration values and blocks. Validation of the ARC configuration is enforced by A-REX startup scripts. If you have obsoleted configuration blocks or parameters in the arc.conf then A-REX will not start, and you will be pointed to the error.

REMOVED and OBSOLETE components

The following components and functionalities have been removed and are not supported any in ARC7 any longer:

- Server-side gridftp-based ARC solutions both for job management and data storage
- Both server and client-side support for EMIES job management interface
- ARC Cache Index (ACIX), that was a distributed system that maintained a catalog of locations of cached files stored in various A-REX caches.
- userlists generation and mapping support by nordugridmap utility
- archiving of the JURA accounting service (got removed already in the 6.8.0 release)
- GLUE1 and site-BDII support within the LDAP-based components of the ARC information system
- Remote batch system backend support via SSH
- Python-based LRMS backend framework (it never progressed beyond technology preview status)
- Support for ARGUS via AREX/WS component
- Support for python2
- LDAP information system support in the ARC7 clients
- EGIS resource discovery in the ARC 7 clients
- The following components and functionalities have been kept but labelled as DEPRECATED and target for near-future removal:
 - OBSOLETE: LDAP-based server-side information system with GLUE2 and nordugrid schema rendering
 - OBSOLETE: gridftp protocol for data staging (still available in ARC 7 and provided as a separate plugin package)
 - OBSOLETE: Support for Loadleveler and LSF batch systems

10.7 Using ARC packages from nightly builds

Recent development version of ARC is available as a [nightly builds packages](#) for many Linux distribution.

To use these nightlies smoothly for installation and updates, the repository should be added to your system. Run the commands below as root.

10.7.1 Enabling nightlies repo for RHEL-based distributions

Setup a CRON job (e.g. `/etc/cron.daily/update-arc-nightly-time.sh`) to fetch the latest nightlies date:

```
cat << EOF > /etc/cron.daily/update-arc-nightly-time.sh
#!/bin/bash
curl -s http://builds.nordugrid.org/nightlies/nordugrid-arc/arc7/ | sed -n 's/^.*<a.*>
↪\(.*\)\</a>.*$/\1/p' | sort | tail -1 > /etc/yum/vars/arcnightly
EOF
```

Run the CRON script once manually to have the `arcnightly` variable initialized. Also, make sure the script has executable permissions.

```
chmod +x /etc/cron.daily/update-arc-nightly-time.sh
/bin/bash /etc/cron.daily/update-arc-nightly-time.sh
```

To add the nightlies repository to your RHEL-based system, create a `/etc/yum.repos.d/nordugrid-nightly.repo`:

```
cat << EOF > /etc/yum.repos.d/nordugrid-nightly.repo
[nordugrid-nightly]
name=Nordugrid ARC ARC7 Nightly Builds - \${basearch}
baseurl=http://builds.nordugrid.org/nightlies/nordugrid-arc/arc7/\${arcnightly}/rocky/9/
↪ \${basearch}
enabled=1
gpgcheck=0
EOF
```

Note

Note that you should modify OS release version to match your case. Allowed values are *centos-stream*, *rocky*, *fedora*, *debian*, *ubuntu*.

Check if it works running yum (or dnf), e.g.:

```
dnf makecache
```

Enabling dependent repos for RHEL-based distributions

The Nordugrid repositories for RedHat Enterprise Linux/CentOS depends on the [EPEL](#) Repositories which must also be part of the YUM configuration, so as root user do:

For RHEL8 flavour:

```
dnf install -y epel-release
dnf config-manager --set-enabled powertools
```

For RHEL9 flavour:

```
dnf install -y epel-release
dnf config-manager --set-enabled crb
```

Once the Nordugrid repositories are configured and the dependency above installed, install the alpha/beta/release-candidate packages with:

```
dnf install --enablerepo nordugrid-nightly <list-of-packages>
```

For instance:

```
dnf install --enablerepo nordugrid-nightly nordugrid-arc-arex
```

Please refer to the *ARC Computing Element Installation and Configuration Guide* for package selection and configuration.

10.7.2 Enabling nightlies repo for Debian-based distributions

Nightlies for Debian/Ubuntu are available as standalone packages without repository index files generated.

The suggested approach is to:

- download packages locally
- create the necessary repository package index
- use the repository on the local filesystem

To accomplish this, install the necessary tools for making repository files:

```
apt-get -y install dpkg-dev
```

Regularly fetch latest nightly packages and create the repository index with the following daily CRON job:

```
cat << EOF > /etc/cron.daily/update-arc-nightly-time.sh

#!/bin/bash

arcrelease='ubuntu/24.04/amd64'
latestdate=$( wget -q -O - http://builds.nordugrid.org/nightlies/nordugrid-arc/arc7/ |
sed -n 's/^.*<a.*>\(.*\)\/</a>.*$/1/p' | sort | tail -1 )
repopdir=/srv/nordugrid-nightlies

rm -rf $repopdir; mkdir -p $repopdir; cd $repopdir

wget -q -r -nH --cut-dirs=8 --no-parent -l1 -A "*.deb, *.ddeb" http://builds.
nordugrid.org/nightlies/nordugrid-arc/arc7/$latestdate/$arcrelease/

dpkg-scanpackages . /dev/null 2>/dev/null | gzip -9c > Packages.gz

EOF
```

Create the local repository file `/etc/apt/sources.list.d/nordugrid-nightlies.list`:

```
cat << EOF > /etc/apt/sources.list.d/nordugrid-nightlies.list

deb [trusted=yes] file:/srv/nordugrid-nightlies ./

EOF
```

Check it works running `apt`, e.g:

```
apt-get update
```

The Nordugrid ARC repositories for Ubuntu packaging utility APT can be configured through: `/etc/apt/sources.list` or when supported through a repo specific file: `/etc/apt/sources.list.d/nordugrid.list`.

The repository configuration can be set up automatically by means of installing `nordugrid-release` package or creating the configuration file manually.

Debian version names:

Debian release	Code name
13	trixie
12	bookworm
11	bullseye

Install with dpkg

Install manually

The examples below give you the link for most recent Debian releases. Packages are shown below for amd64 architecture. Replace amd64 for i386 if required for your architecture.

Debian: 13 12 11

Install the source file with dpkg, example shown for Debian 12:

```
[root ~]# wget -q https://download.nordugrid.org/packages/nordugrid-release/releases/
↳ 7/debian/12/amd64/nordugrid-release_7~bpo12+1_all.deb
[root ~]# dpkg -i nordugrid-release_7~bpo12+1_all.deb
```

For a different version of Debian, change the version names appropriately.

For manual APT repository setup for Debian, the APT sources file should contain the following (here shown for Debian 12 Bookworm).

The configurations for the various APT based distributions can be found in the following sections. To enable a specific repository, remove the “#” from the beginning of the line, before the “deb” as shown for the Base Channel.

```
# Base channel - must be enabled
deb http://download.nordugrid.org/repos/7/debian/ bookworm main
deb-src http://download.nordugrid.org/repos/7/debian/ bookworm main

# Updates to the base release - should be enabled
deb http://download.nordugrid.org/repos/7/debian/ bookworm-updates main
deb-src http://download.nordugrid.org/repos/7/debian/ bookworm-updates main

# Scheduled package updates - optional
#deb http://download.nordugrid.org/repos/7/debian/ bookworm-experimental main
#deb-src http://download.nordugrid.org/repos/7/debian/ bookworm-experimental main
```

For a different release version, change the version name accordingly.

```
[root~]# apt-get update
```

```
[root ~]# dnf config-manager --set-enabled powertools && dnf makecache
```

```
[root ~]# dnf config-manager --set-enabled crb && dnf makecache
```

The NorduGrid ARC repositories for RedHat Enterprise Linux / CentOS packaging utility dnf can be configured through: /etc/yum.repos.d/nordugrid.repo

The repository configuration can be set up automatically with dnf by installing the nordugrid-release package or creating the configuration file manually.

Install with dnf

Install manually

The easiest way to configure DNF to use the NorduGrid repository for Red Hat Enterprise Linux, CentOS and similar distributions is to install the nordugrid-release package which can be found in the NorduGrid package repository for the appropriate RHEL/EPEL release.

Links to the release packages:

Rocky Linux/AlmaLinux

CentOS Stream

Fedora

Rocky Linux: 10 9 8

CentOS Stream: 10 9

Fedora: 41 40

Install with dnf (Fedora, CentOS Stream, Rocky Linux, CentOS Linux 8+9) by copying the appropriate link from above

```
[root ~]# dnf install <rhel-repo link>
```

This creates the appropriate repo files in `/etc/yum.repos.d/`.

For manual YUM repository setup, create a file `/etc/yum.repos.d/nordugrid.repo` with the following contents (here using Rocky as example, if you are on Fedora, replace `rocky` with `fedora`)

If you are installing an alpha, beta or release candidate, please set the `nordugrid-testing` to `enabled=1`.

```
[nordugrid]
name=NorduGrid - $basearch - base
baseurl=http://download.nordugrid.org/repos/7/rocky/$releasever/$basearch/
↳base
enabled=1
gpgcheck=1
gpgkey=http://download.nordugrid.org/RPM-GPG-KEY-nordugrid-7

[nordugrid-updates]
name=NorduGrid - $basearch - updates
baseurl=http://download.nordugrid.org/repos/7/rocky/$releasever/$basearch/
↳updates
enabled=1
gpgcheck=1
gpgkey=http://download.nordugrid.org/RPM-GPG-KEY-nordugrid-7

[nordugrid-testing]
name=NorduGrid - $basearch - testing
baseurl=http://download.nordugrid.org/repos/7/rocky/$releasever/$basearch/
↳testing
enabled=0
gpgcheck=1
gpgkey=http://download.nordugrid.org/RPM-GPG-KEY-nordugrid-7
```

For Debian without sudo:

```
[root ~]# wget -q http://download.nordugrid.org/DEB-GPG-KEY-nordugrid-7.asc -O- | apt-
↳key add -
```

For rpm based distributions like Red Hat Enterprise Linux and Fedora:

```
[root ~]# rpm --import http://download.nordugrid.org/RPM-GPG-KEY-nordugrid-7
```

For Ubuntu distributions with sudo:

```
[user ~]$ wget -q http://download.nordugrid.org/DEB-GPG-KEY-nordugrid-7.asc -O- |
↳sudo apt-key add -
```

The Nordugrid ARC repositories for Ubuntu packaging utility APT can be configured through: `/etc/apt/sources.list` or when supported through a repo specific file: `/etc/apt/sources.list.d/nordugrid.list`.

The repository configuration can be set up automatically by means of installing `nordugrid-release` package or creating the configuration file manually.

Ubuntu version names:

Ubuntu release	Code name
25.04	plucky
24.04	noble
22.04	jammy

Install with dpkg

Install manually

The examples below give you the link for most recent Ubuntu releases. Packages are shown below for amd64 architecture. Replace amd64 for i386 if required for your architecture.

Ubuntu: 25.04 24.04 22.04

Install the source file with dpkg, example shown for Ubuntu 22.04:

```
[root ~]# wget -q https://download.nordugrid.org/packages/nordugrid-release/releases/
↪7/ubuntu/24.04/amd64/nordugrid-release_7~noble1_all.deb
[root ~]# dpkg -i nordugrid-release_7~noble1_all.deb
```

For a different version of Ubuntu, change the version names appropriately.

For manual APT repository setup for Ubuntu, the APT sources file should contain the following (here shown for Ubuntu 22.04 Jammy).

The configurations for the various APT based distributions can be found in the following sections. To enable a specific repository, remove the “#” from the beginning of the line, before the “deb” as shown for the Base Channel.

```
# Base channel - must be enabled
deb http://download.nordugrid.org/repos/7/ubuntu/ jammy main
deb-src http://download.nordugrid.org/repos/7/ubuntu/ jammy main

# Updates to the base release - should be enabled
deb http://download.nordugrid.org/repos/7/ubuntu/ jammy-updates main
deb-src http://download.nordugrid.org/repos/7/ubuntu/ jammy-updates main

# Scheduled package updates - optional
#deb http://download.nordugrid.org/repos/7/ubuntu/ jammy-experimental main
#deb-src http://download.nordugrid.org/repos/7/ubuntu/ jammy-experimental main
```

For a different release version, change the version name accordingly.

10.8 NorduGrid testing repository information for ARC 7

The NorduGrid ARC packages are available through YUM and APT repositories for several systems. We have release-based repositories that you can follow. This will keep your install to a particular release of NorduGrid ARC with only minor and bug-fixing updates. This page explains how you install NorduGrid ARC packages from the *testing* channel - which is used for alpha, beta, and release candidate releases.

10.8.1 Repository security

The Nordugrid RPM packages and DEB repositories are signed, and in order for the repository tools APT and YUM to verify them you must install the Nordugrid GPG key:

For rpm based distributions like Red Hat Enterprise Linux and Fedora:

```
[root ~]# rpm --import http://download.nordugrid.org/RPM-GPG-KEY-nordugrid-7
```

For Ubuntu distributions with sudo:

```
[user ~]$ wget -q http://download.nordugrid.org/DEB-GPG-KEY-nordugrid-7.asc -O- | sudo apt-key add -
```

For Debian without sudo:

```
[root ~]# wget -q http://download.nordugrid.org/DEB-GPG-KEY-nordugrid-7.asc -O- | apt-key add -
```

10.8.2 Repository configuration - Red Hat Enterprise Linux

The Nordugrid ARC repositories for RedHat Enterprise Linux / CentOS packaging utility dnf can be configured through:

```
/etc/yum/nordugrid-testing.repo
```

The repository configuration must be set up manually when using the testing repo, as per instructions below.

Manual YUM repository setup - Nordugrid repository

For manual YUM repository setup, create a file `/etc/yum.repos.d/nordugrid-testing.repo` with the following contents:

Note

Here we are using Rocky as an example, if you are using other flavours, please replaces rocky with

- Fedora: fedora
- CentosStream: centos-stream
- AlmaLinux: rocky

[nordugrid-testing]

```
name=Nordugrid - $basearch - Testing
baseurl=http://download.nordugrid.org/repos/7/rocky/9/$basearch/testing
enabled=1
gpgcheck=1
gpgkey=file:///etc/pki/rpm-gpg/RPM-GPG-KEY-nordugrid-7
```

[nordugrid-testing-debuginfo]

```
name=Nordugrid - $basearch - Testing - Debug
baseurl=http://download.nordugrid.org/repos/7/rocky/9/$basearch/testing/debug
enabled=0
gpgcheck=1
gpgkey=file:///etc/pki/rpm-gpg/RPM-GPG-KEY-nordugrid-7
```

(continues on next page)

(continued from previous page)

```
[nordugrid-testing-source]
name=NorduGrid - Testing - Source
baseurl=http://download.nordugrid.org/repos/7/rocky/9/source/testing
enabled=0
gpgcheck=1
gpgkey=file:///etc/pki/rpm-gpg/RPM-GPG-KEY-nordugrid-7
```

Check if it works running dnf, e.g.:

```
[root ~]# dnf makecache
```

Install required packages

The Nordugrid repositories for RedHat Enterprise Linux/CentOS depends on the [EPEL](#) Repositories which must also be part of the YUM configuration:

For RHEL8 flavour:

```
dnf install -y epel-release
dnf config-manager --set-enabled powertools
```

For RHEL9 flavour:

```
dnf install -y epel-release
dnf config-manager --set-enabled crb
```

Once the Nordugrid repositories are configured and the dependency above installed, install the alpha/beta/release-candidate packages with:

```
[root~]# dnf install --enablerepo nordugrid-testing <list-of-packages>
```

For instance:

```
dnf install --enablerepo nordugrid-testing nordugrid-arc-arex --refresh
```

Please refer to the *ARC Computing Element Installation and Configuration Guide* for package selection and configuration.

10.8.3 Repository configuration - Debian and Ubuntu

The Nordugrid ARC repositories for Debian and Ubuntu packaging utility APT can be configured through:

```
/etc/apt/sources.list
```

or when supported through a repo specific file:

```
/etc/apt/sources.list.d/nordugrid.list
```

The configurations for the various APT based distributions can be found in the following sections. To enable a specific repository, remove the “#” from the beginning of the line, before the “deb” as shown for the Base Channel.

For the testing repo, you must manually set up the configuration file as per instructions below.

Manual APT repository setup - NorduGrid repository

For manual APT repository setup for Debian, the APT sources file should contain the following (here shown for Debian 12 bookworm):

```
# Scheduled package updates - optional
#deb http://download.nordugrid.org/repos/7/debian/ bookworm-experimental main
#deb-src http://download.nordugrid.org/repos/7/debian/ bookworm-experimental main
```

For manual APT repository setup for Ubuntu, the APT sources file should contain the following (here shown for Ubuntu 22.04 jammy):

```
# Scheduled package updates - optional
#deb http://download.nordugrid.org/repos/7/ubuntu/ jammy-experimental main
#deb-src http://download.nordugrid.org/repos/7/ubuntu/ jammy-experimental main
```

For a different release version, change the version name accordingly.

For Debian:

- 12: bookworm
- 11: bullseye
- 10: buster

For Ubuntu:

- 24.04: noble
- 23.10: mantic
- 22.04: jammy
- 20.04: focal

Install required packages

Run the following command to update the packages database:

```
[root~]# apt-get update
```

Install the packages with (showing example for nordugrid-arc-arex):

```
[root~]# apt-get <list-of-packages>
```

Please refer to the *ARC Computing Element Installation and Configuration Guide* for package selection and configuration.

10.9 Work-in-progress Docs

Hidden area that holds unfinished documents to be build and available in the doc tree, but not yet linked to the right place.

10.9.1 INTERNAL interface - ARC 6

NB! WIP [TODO] Fix references, now in tex style ref to arc-ce sys admin guide

The INTERNAL submission interface is aimed for restrictive HPC sites. When ARC runs in the internal mode, it should do so alongside a local instance of the ARC Control Tower (aCT). aCT pulls jobs from the central job server (like PaNDA for ATLAS), and feeds ARC with new jobs internally. Since there is no external access, there is no need for a web-service, gridftp server or ldap, as the purpose of these components are to facilitate external access. There is neither a need for a host-certificate on such a machine. The only service running on ARC is A-REX. Therefore a stripped-down version of aCT and ARC-CE can be used which is beneficial for installation, configuration and maintenance. Furthermore, in the INTERNAL mode ARC should be installed as a normal user. No user mapping (See Section [myref{sub:access_control}](#)) is needed in this case, as the default behaviour in ARC 6 is to map the user submitting the job to the A-REX user.

Implementation overview

The INTERNAL submission plugin which is part of the ARC client, interacts with the parent plugin classes using the same API as the other plugins such as ARC REST plugin. However, the INTERNAL plugin interacts directly with the A-REX memory and methods, and therefore is integrated as part of the A-REX service which belongs to the ARC-CE code-base. Therefore, both the ARC client and the ARC CE must be installed on the same machine for the INTERNAL submission plugin interface to function.

All interaction between the client and A-REX happens directly via files in the controldir or via A-REX memory.

Actions

Retrieving Service Information

As a site running in the INTERNAL mode is not accessible from the outside, any service retrieval information can only be done from within the site. You may inspect the service information as usual by issuing the `arcinfo` command. The information is extracted by direct access to the `info.xml` file in the controldir. The `sstat` method of the `INTERNALClient` reads the `info.xml` and outputs information in xml-format to the client. An example of the output of `arcinfo` for localhost

```
[root ~]# arcinfo -c localhost
Computing service:
  Information endpoint: file://localhost
  Submission endpoint: file://localhost (status: ok, interface: org.nordugrid.
↪internal)
```

When `arcinfo` is called, the INTERNAL submission interface extracts the site information by directly accessing the `info.xml` file in the controldir. The INTERNAL plugin reads the `info.xml` and outputs information in xml-format to the client, which in turn displays it to the user.

Job submission

When a job is submitted via the INTERNAL submission interface the plugin creates an A-REX job object, which in turn takes care of creating all necessary files (like for instance the ARC job description) and folders (sessiondir) for the job, in addition to creating a job ID. The INTERNAL plugin then places any input files local to the client in the newly created sessiondir. Remaining remote input files are downloaded by the DTR (See Section [myref{sub:datastaging}](#)). Once these files are present in the controldir A-REX adds the job to its joblist, and takes over the handling of the job from there.

Accessing Information About Job

Job information evoked by calling `arcstat` is extracted from a combination of information stored in A-REX memory (job state) and the `job.ID.local` file in the `controldir` (session, stagein and stageout directories).

Controlling Execution Of Job

Killing, cleaning and resubmitting jobs is initiated by direct call the existing `ARexJob` methods: `Kill()`, `Clean()`, `Resume()`. These methods all place files in the `controldir` that the grid-manager acts upon, such as `job.jobid.clean` mark or `job.jobid.cancel` mark.

Delegation Interface [TODO]