



TSDuck Contributor Guide

Thierry Lélégard

Version 3.46-4811, September 2026

Contents

Preface	5
1. Contributing to TSDuck Development	6
1.1. Transparency of contributions	6
1.2. License and copyright	6
1.3. Contributor workflow	7
1.3.1. Initial setup	7
1.3.2. Contributing code	7
1.3.3. Testing your code	8
1.4. Integrator workflow	8
1.4.1. Method 1: using the GitHub command-line tool	9
1.4.2. Method 2: using Git only	9
2. Testing TSDuck	11
2.1. Testing overview	11
2.2. Organization of the tests	11
2.3. The TSDuck library test suite	11
2.4. The TSDuck tools and plugins test suite	13
2.4.1. Structure of the test suite	13
2.4.2. Adding new tests	13
2.4.3. Testing a development version	14
2.4.4. Running PSI/SI tests only	15
2.4.5. Large files	15
3. Automation	16
3.1. GitHub workflows	16
3.2. Continuous integration	16
3.3. Nightly builds	17
3.4. Cleanup of long-standing issues	17
4. Publishing new releases	18
4.1. Automated release creation	18
4.2. Manual release creation	19
4.2.1. Building the various binaries	19
4.2.2. Creating the GitHub release	19
4.3. Post-release tasks	20
4.3.1. Creating the HomeBrew release	20
4.3.2. Creating the WinGet release	20
4.3.3. Creating the FreeBSD port	20
4.3.4. Updating the version number	21
5. Maintaining TSDuck Code	22
5.1. Modularity and self-registration	22
5.1.1. Principles of independence	22
5.1.2. Source files structure	23
5.1.3. Windows constraints	24
5.1.4. Self registration of modules	25
5.2. Adding PSI/SI tables or descriptors	27

5.2.1. Code base selection	27
5.2.1.1. Classification of descriptors	28
5.2.1.2. Extended descriptors	28
5.2.1.3. Private descriptors (DVB)	28
5.2.1.4. Table-specific descriptors	29
5.2.2. Affiliation to a standard	29
5.2.3. Declaring identifiers	30
5.2.4. XML definition	30
5.2.5. Documentation file	31
5.2.6. C++ class	32
5.2.6.1. Structure registration	32
5.2.6.2. Programming reference	32
5.2.7. Names file	33
5.2.8. Tests	33
5.3. TSP design	35
5.3.1. Plugin Executors	35
5.3.2. Transport packets buffer	35
5.4. Hardware support	36
5.4.1. Standard tuner receiver devices (DVB, ATSC, ISDB)	36
5.4.2. Dektec devices	39
5.4.3. HiDes devices	39
5.5. AstroMeta-based modulators (formerly VATEk)	40
6. Coding guidelines	41
6.1. Rationale for coding guidelines	41
6.2. Classification of coding guidelines	41
6.3. Generic coding guidelines	42
6.3.1. Generic coding rules and recommendations	42
6.3.1.1. Software architecture	42
6.3.1.2. Source code structure	43
6.3.1.3. Revision control system	44
6.3.1.4. Internationalization	45
6.3.1.5. Modularity and compatibility	45
6.3.1.6. Naming conventions	46
6.3.1.7. Coding principles	47
6.3.1.8. Secure coding	48
6.3.1.9. Software evolution	52
6.3.1.10. Compilation errors and warnings	53
6.3.1.11. Makefiles	53
6.3.1.12. Unit testing	55
6.3.1.13. Integration of open source software	56
6.3.2. Generic coding conventions	57
6.3.2.1. Control characters	57
6.3.2.2. Character encoding	58
6.4. C++ coding guidelines	58
6.4.1. C++ coding rules and recommendations	59
6.4.1.1. Language selection	59

6.4.1.2. Modularity	59
6.4.1.3. Global data	61
6.4.1.4. Naming and syntax formatting	63
6.4.1.5. Coding style	65
6.4.1.6. Strict typing	72
6.4.1.7. Assertions	76
6.4.1.8. Secure coding	77
6.4.1.9. C++ classes	86
6.4.1.10. C++ constructors and destructors	92
6.4.1.11. C++ operators	96
6.4.1.12. C++ object management	99
6.4.2. C++ coding conventions	107
6.4.2.1. Source code formatting	107
6.4.2.2. Modularity	107
6.4.2.3. Naming conventions	107
6.4.2.4. Syntax formatting conventions	110
6.4.2.5. Doxygen self-documentation	113
Appendix A: Licenses	116
A.1. TSDuck license	116
A.2. Third-party libraries	116
Appendix B: References	117
B.1. Acronyms and abbreviations	117
Bibliography	117

Preface

TSDuck is a free and open-source toolkit to manipulate MPEG Transport Streams.

This document is the TSDuck Contributor Guide. It provides guidelines on how to contribute to the development of TSDuck. It is also a reference for TSDuck maintainers: it describes the automated test and release workflows and provides coding guidelines.

Structure of this guide:

- The [chapter 1](#) is a guide for advanced users who wish to contribute to TSDuck development, submit new features, improvements, or bug fixes.
- The [chapter 2](#) describes the test-driven development approach which is used in TSDuck.
- The [chapter 3](#) describes the various automation processes which speed up and simplify TSDuck development.
- The [chapter 4](#) explains the process to publish new releases.
- The [chapter 5](#) covers various technical topics in the TSDuck source code.
- The [chapter 6](#) lists the coding guidelines which are used in the development of TSDuck. All maintainers and contributors are required to follow these guidelines.

License

TSDuck is released under the terms of the license which is commonly referred to as "BSD 2-Clause License" or "Simplified BSD License" or "FreeBSD License". This is a liberal license which allows TSDuck to be used in a large number of environments. See the [appendix A](#) for more details.

Documentation format

The TSDuck guides are built using [asciidoc](#), from a set of text files which are maintained alongside the source code, in the same [git](#) repository.

This guide is formatted for HTML. The file [tsduck-contrib.html](#) is monolithic and self-sufficient, without reference to external images. Therefore, this HTML file can be downloaded, saved, and copied, as long as the license and content are not modified.

A PDF version [tsduck-contrib.pdf](#) is also available. However, due to limitations in the PDF generator of [asciidoc](#), the rendering is sometimes not as good as the HTML document.

Documentation set

The TSDuck documentation set is made of:

1. [TSDuck User Guide](#), using TSDuck commands and plugins (also from [tsduck.io](#) and in [PDF](#) format)
2. [TSDuck Builder Guide](#), building and installing TSDuck (also from [tsduck.io](#) and in [PDF](#) format)
3. [TSDuck Developer Guide](#), using TSDuck from C++, Python, Java applications (also from [tsduck.io](#) and in [PDF](#) format)
4. [TSDuck Contributor Guide](#), contributing to TSDuck development (also from [tsduck.io](#) and in [PDF](#) format)
5. [TSDuck Programming Reference](#), Doxygen-generated reference of all TSDuck classes.

Chapter 1. Contributing to TSDuck Development

TSDuck development is managed using Git. The [reference repository](#) is on GitHub.

Code contributions from external developers are welcome and will be reviewed (without guaranteed response time however). Contributions shall be submitted using [pull requests on GitHub](#) exclusively.

This chapter summarizes the main actions to help developers and integrators to work with pull requests. This is the minimal set of actions.

More details can be found on [GitHub documentation](#). Several articles also describe the GitHub standard fork & pull request workflow. We specifically [recommend this one](#).

1.1. Transparency of contributions

All commits in a Git pull request shall have a clear and transparent identification of the author. The author name shall be the true first and last names of the contributor. No pseudo or other forms or anonymity is allowed. Preferably (although not required), the author's e-mail should be a real address where the contributor can be contacted.

This requirement for transparency is not arbitrary. There is a reason for it. The world of Digital TV is roughly divided in industries, service providers, TV operators, and pirates. A technical toolbox such as TSDuck is useful to everyone, equally. But it must be clear to everyone that TSDuck is made by engineers for engineers. TSDuck shall remain a fully transparent project: open source, identified web sites, identified authors and contributors.

In the world of Pay-TV, anonymity equals piracy. This may seem unfair but this is the way it is perceived by the industry. So, to maintain the trust in TSDuck, let's keep anonymity away from it. We hope that all contributors understand this position.

1.2. License and copyright

The TSDuck project is free and open-source. It must and will remain so. TSDuck is released under the terms of the license which is commonly referred to as "BSD 2-Clause License" or "Simplified BSD License" or "FreeBSD License". See the file [LICENSE.txt](#) at the root of the TSDuck source tree or [\[BSD-2C\]](#) for more details.

All contributions shall be released under the terms of the same license. Any code which does not comply with the BSD 2-Clause License will be rejected.

The author of a piece of source code retains copyright on the code he/she wrote.

All source files shall contain the following header template:

```
//-----  
//  
// TSDuck - The MPEG Transport Stream Toolkit  
// Copyright (c) YEAR, FIRST-NAME LAST-NAME  
// BSD-2-Clause license, see LICENSE.txt file or https://tsduck.io/license  
//  
//-----
```

Open-source developers deserve credits for their work. The full real name of the author of a source file shall be mentioned in the copyright line of the header.

- Each source file which is created, entirely or mostly written by a developer shall have the full name of that developer alone in the copyright line.
- If an existing source file is significantly modified, refactored, improved, by a different developer, add the

name of that developer in the copyright line, after other developers names.

- For small bug fixes or local improvements in a source file, don't alter the copyright line.

Additionally, regardless of the size of their contribution, with or without reference in the copyright of a source file, all contributors shall have their full real name in the file `CONTRIBUTORS.txt` at the root of the TSDuck source tree. Please keep the file sorted by last name.



Most software developers, including the author of TSDuck, are not lawyers and sometimes use legal terms such as "license" and "copyright" in an improper way. If this is the case in the above paragraphs, please propose a more appropriate form in the TSDuck support area (see [\[TSDuck-Issues\]](#)).

1.3. Contributor workflow

1.3.1. Initial setup

The first requirement is to create a GitHub account, if you do not already have one.

Initially, create your own fork of the TSDuck repository. Go to the TSDuck [reference repository](#) and click on the "Fork" button.

Clone your GitHub forked repository on your development system. Use one of the two following commands.

```
$ git clone https://USERNAME@github.com/USERNAME/tsduck.git
$ git clone git@github.com:USERNAME/tsduck.git
```

In the first case, you will need to provide a GitHub personal access token each time you push to GitHub. In the second case, you need to first upload your SSH public key to GitHub and then simply push without password.

You may want to track more precisely the master branch of the reference repository. See more details in the [above mentioned article](#).

1.3.2. Contributing code

To facilitate merging, each contribution should be provided in a specific branch. Let's call it `newfeature` here:

```
$ git branch newfeature
$ git checkout newfeature
```

Then, do your coding work.

If the contribution brings new features, be sure to document them in the TSDuck user guide. The user guide is made of text files in Asciidoc format in the directory `doc/user`.

New features and bug fixes should also be documented in the file `CHANGELOG.txt`.

When all modifications are ready, commit and push the work to GitHub:

```
$ git push origin newfeature
```

Finally, create the pull request. Go to your forked repository on GitHub, something like <https://github.com/username/tsduck>, and select the branch `newfeature`. Select the "Pull requests" tab and click on the green button "New pull request". Select the branch for the pull request and click on "Create pull request".

The pull request is transmitted to the project main repository `tsduck/tsduck`. The Continuous Integration (CI)

process is automatically started on your code. During this CI process, the code is compiled and tested on Linux, macOS and Windows using various compilers and variants of the C++ standard. Then, all TSDuck tests are run on your code.

In case of failure of the CI job, you will be notified by mail and the pull request is retained on hold. You may then review the failures, update your code, commit and push on your branch again. The new commits are added to the pull request and the CI job is run again.

Most users work in one environment, one operating system, one compiler. Even if the code works in this environment, it may fail to build or run in another environment, operating system or compiler. This is why the CI process is useful because you can review the impact of your modifications in other environments.

1.3.3. Testing your code

Before pushing your code and creating the pull request, you will test your code. Additionally, when you add new features or support for new signalization, tables or descriptors, it is recommended to update the TSDuck test suite.

See the [chapter 2](#) for more details.

To update the TSDuck test suite according to your new code, fork the [tsduck-test](#) repository, update it, and create pull requests on this repo using the same method as the main [tsduck](#) repository.

Pay attention to the interactions between the [tsduck](#) and [tsduck-test](#) repositories.

The [tsduck-test](#) repository contains tests and reference outputs for those tests. When you update the TSDuck code, the test reference output may need to be updated accordingly. You do that in your fork of the [tsduck-test](#) repository. When you create a pull request on the main [tsduck](#) repository, the CI job checks the origin of the pull request. In your case, this is your [username/tsduck](#) forked repository. The CI job checks if you also have a [username/tsduck-test](#) forked repository. If it exists, it is used to run the test suite. If you do not have a fork of the test repository, the reference [tsduck/tsduck-test](#) repository is used.

Consequently, the recommended workflow depends on the type of code contribution you provide.

- If you provide a simple code update which has no impact on the test suite, then you should fork the [tsduck/tsduck](#) repository only. Your code will be tested against the [tsduck/tsduck-test](#) repository to make sure it does not break the project.
- If your contribution is more substantial and needs an update of the test suite, then you need to fork the [tsduck/tsduck](#) and [tsduck/tsduck-test](#) repositories. Once your code and tests are complete, create the commits and push the two repositories. At the end, create the pull requests on the two repositories. The CI job for the [tsduck](#) repository will then use your [username/tsduck-test](#) repository for the test suite. If all tests pass on all operating systems, your contributions in [tsduck](#) and [tsduck-test](#) will be merged.

One last point: If you maintain your fork of [USERNAME/tsduck-test](#), be sure to keep it synchronized with the reference [tsduck/tsduck-test](#) repository because your [USERNAME/tsduck-test](#) will always be used in your CI jobs. If one day, you submit a small code update which did not need any update in the test suite and your [USERNAME/tsduck-test](#) is not up-to-date, your CI job may fail.

1.4. Integrator workflow

The TSDuck maintainer has to review the pull requests and, if they are satisfactory, merge them into the master branch of the project. Additional review and fix may be necessary before pushing the contribution.

There are two ways to do this. We now recommend the first one, using [gh](#), the GitHub command-line tool.

1.4.1. Method 1: using the GitHub command-line tool

The GitHub command-line tool is named `gh`. It is an encapsulation of the most useful Git operations for specialized tasks on GitHub repositories. It is developed by GitHub and available as a standard package on most distros:

- Ubuntu / Debian: `apt install gh`
- Fedora / Red Hat: `dnf install gh`
- macOS: `brew install gh`
- Windows: `winget install github.cli`

To integrate a pull request number *NNN*, fetch it in a local branch named *NNN*:

```
$ gh pr checkout NNN -b NNN
```

To merge the pull request into the `master` branch:

```
$ git checkout master
$ git merge NNN
```

1.4.2. Method 2: using Git only

On your local development system, configure your TSDuck development git repository to track all pull requests. In the file `.git/config`, add the following line in section `[remote "origin"]`:

```
[remote "origin"]
... existing lines ...
fetch = +refs/pull/*/head:refs/pull/origin/*
```

To integrate a pull request number *NNN*, fetch it in a local branch named *NNN*:

```
$ git fetch origin
$ git checkout -b NNN pull/origin/NNN
```

To merge the pull request into the `master` branch:

```
$ git checkout master
$ git merge NNN
```

Alternatively, if you know that the pull request is correct and you want to directly merge it:

```
$ git fetch origin
$ git merge pull/origin/NNN
```

However, in the context of the TSDuck repository, this method creates problems and we no longer recommend it.

With this configuration in `.git/config`, the command `git fetch` always fetches *all* pull requests from the beginning. In the general case, this is not a problem. However, the TSDuck repository went through a history rewrite in August 2024. The original user's and developer guides were maintained in Microsoft Word files. They are binary files which are badly managed by git. The TSDuck repository accumulated 2 GB of history. Each `git clone` command triggered 2 GB of data transfer. Each local repository had 2 GB of disk space in the `.git` subdirectory. To solve this, the documentation was migrated to AsciiDoc, a text format. The history of the

repository was rewritten from the beginning without those binary files. The entire history was reduced to 40 MB. The clone operations were faster, the disk space was optimized.

However, the old pull requests could not be rewritten. Using `fetch = +refs/pull/*` in `.git/config`, the command `git fetch` downloads the history of all pull requests before the history rewrite. This results in retrieving again the old 2 GB of history. Since we are only interested in working on recent pull requests, after the history rewrite, this method is no longer recommended.

Chapter 2. Testing TSDuck

2.1. Testing overview

TSDuck is highly flexible, allowing an unlimited number of configurations. Testing it is consequently challenging. Moreover, some use cases are difficult to automate or require specific hardware. Testing live TS reception using all possible DVB tuners or Dektec devices requires a lot of material, human and time resources which are far beyond the capabilities of a free open-source project.

From a strict industrial standpoint, TSDuck is consequently not fully tested.

However, on a pragmatic standpoint, test suites have been setup to extract the best of what could be done with limited resources.

Thus, although it is impossible to guarantee that a given release of TSDuck is bug-free, we may be relatively confident that no major regression has been introduced.

2.2. Organization of the tests

The code of TSDuck is divided in two parts, a large C++ library (`tsduck.dll`, `libtsduck.so`, or `.dylib`) and a collection of small command line tools and plugins.

Similarly, the tests for TSDuck are divided in two parts.

- The TSDuck library has its own unitary test suite based on a custom framework named "TSUnit". This test suite is part of the `main tsduck` repository for TSDuck in directory `src/utest`.
- The tools and plugins are less easy to test. They work on large transport stream files which would clutter the TSDuck repository. The repository `tsduck-test` contains those tests and the relevant scripts and data files.

The two test suites are fully automated.

2.3. The TSDuck library test suite

In the main TSDuck repository, the directory `src/utest` contains the source file for one single program named `utest`. This program is divided in many source files (or test suites). Each source file contains many unitary tests. The test infrastructure is based on a custom framework named "TSUnit".



The structure of TSUnit is freely inspired by equivalent frameworks such as CppUnit and JUnit. It has been adapted to TSDuck specificities and focuses on fast and easy development of test suites.

The `utest` executable is built twice, once using the TSDuck shared library and once using the static library. On UNIX systems (Linux, macOS, BSD), both versions of the test suite are built and run using `make test` as illustrated below

```
$ make -j10 test
.....
[CXX] utestZlib.cpp
[LD] /home/tsduck/bin/release-x86_64-vmubuntu/utest
[LD] /home/tsduck/bin/release-x86_64-vmubuntu/utest_static
[UTEST] /home/tsduck/bin/release-x86_64-vmubuntu/utest

OK (666 tests, 32317 assertions)

[UTEST] /home/tsduck/bin/release-x86_64-vmubuntu/utest_static
```

```
OK (663 tests, 32294 assertions)
```



The number of tests and assertions changes with new versions of TSDuck. When new features are added, the corresponding new tests are added.

Note that the statically linked version contains slightly less tests. The missing tests are related to plugin activations and they can run only in a shared library environment.

On Windows, the Visual Studio project builds two executables named `utests-tsduckdll.exe` and `utests-tsducklib.exe`. They can be run manually or from Visual Studio:

```
PS D:\tsduck> .\bin\Release-x64\utests-tsduckdll.exe
```

```
OK (666 tests, 32369 assertions)
```

```
PS D:\tsduck> .\bin\Release-x64\utests-tsducklib.exe
```

```
OK (663 tests, 32346 assertions)
```

```
PS D:\tsduck>
```

Note that the number of assertions is slightly different on Windows because some system-specific tests are different.

By default, the test program runs all tests and reports failures only. But `utest` also accepts a few command options which make it appropriate to debug individual features.

The available command line options are:

- `-d` Debug messages are output on standard error
- `-l` List all tests but do not execute them.
- `-t name` Run only one test or test suite.

First, the option `-l` is used to list all available tests:

```
$ utest -l
AlgorithmTest
  AlgorithmTest::EnumerateCombinations
ArgsTest
  ArgsTest::Accessors
  ArgsTest::AmbiguousOption
.....
XMLTest::Validation
Xoshiro256ssTest
  Xoshiro256ssTest::Random
ZlibTest
  ZlibTest::AllLevels
  ZlibTest::Reference1
  ZlibTest::Reference4
  ZlibTest::Reference9
```



The left-most names are test suite names. They represent one source file in `src/utest`. It is possible to run all tests

in a test suite or one specific test. For instance:

```
$ utest -t ArgsTest
$ utest -t XMLTest::Validation
```

The option `-d` is used to produce debug message (see examples in the test source files).

Thus, `utest` alone can be used as an automated fairly complete non-regression test suite for the TSDuck library or as a debug environment for a given feature under development (in all good "test-driven development" approaches, the code is written at the same time as its unitary test).

2.4. The TSDuck tools and plugins test suite

The Git repository `tsduck-test` contains high-level tests for the TSDuck tools and plugins.

This test suite is fully automated and works on test files only. No specific hardware (DVB tuner or Dektec device), no live stream is tested. This is a known limitation.

The principle is simple. Reference input files are stored in the repository, mainly transport stream files which were once captured from real live sources. Test scripts contain deterministic commands which should always produce the same outputs (data, logs, etc). These reference outputs are also stored in the repository.

The test suite runs using a given version of TSDuck and its outputs are compared with the reference outputs. If differences are detected, there is a potential problem.

Since different versions of TSDuck may produce slightly different outputs, a given version of the test suite formally applies to one version of TSDuck only. Git tags are aligned in both repositories (or should be...) to indicate the target version.

2.4.1. Structure of the test suite

In short, execute the script `run-all-tests.sh` to run the complete test suite.

The repository contains the following subdirectories:

<code>tests</code>	Contains one script per test or set of tests. The name for test <i>NNN</i> is <code>test-NNN.sh</code> . Each test script can be executed individually. All tests are executed using the script <code>run-all-tests.sh</code> .
<code>common</code>	Contains utilities and common script.
<code>input</code>	Contains input data files for the tests.
<code>reference</code>	Contains reference output files for the various tests. There is one subdirectory <code>test-NNN</code> per test which contains all output files for that test.
<code>tmp</code>	Contains output files which are created by the execution of the tests. These files are typically compared against reference output files in <code>reference</code> . These files are temporary by definition. The subdirectory <code>tmp</code> is present on test machines only and is excluded from the Git repository.

2.4.2. Adding new tests

To add a new test:

- Allocate a new test number and document the purpose of the new test in the file `README.md`.
- Add input files in subdirectory `input`. For test *NNN*, all input files should be named `test-NNN.*`. There is generally zero or one input file per test, sometimes more.
- Create the script `test-NNN.sh` in subdirectory `tests`. Use other existing test scripts as templates.

- Run the command `tests/test-NNN.sh --init`. If the test is properly written, this creates the reference output files in the subdirectory `reference/test-NNN`. Manually check the created files, verify that they are correct. Be careful with this step since these files will be used as references.
- Run the same command without the `--init` option. This time, the output files are created in `tmp` and are compared with files in `reference`. Verify that all tests pass. Errors may appear if the test script is not properly written or if the output files contain unique, non-deterministic, time-dependent, system-dependent or file-system-dependent information. Make sure the output files are totally reproduceable in all environments. At worst, add code in the test script to remove any information from the output files which is known to be non-reproduceable.

Sometimes, TSDuck is modified in such a way that an output file is modified on purpose. Usually, this starts with a failed test. When analysing the test failure, it appears that the modification of the output is intentional. In that case, re-run the command `tests/test-NNN.sh --init` to update the reference output files. Do not forget to manually validate them since they will act as the new reference.



The reference output files are stored in the Git repository. Therefore, the best way to have a quick overview of what changed in the output reference files is simply `git diff`.

2.4.3. Testing a development version

By default, the test suite uses the TSDuck command from the system path. Typically, it will use the installed version.

To test a development version, the two Git repositories `tsduck` and `tsduck-test` shall be checked out at the same level, side by side in the same parent directory. First, TSDuck shall be rebuilt in its repository.

Then, when the option `--dev` is specified to a test script or to `run-all-tests.sh`, the test suite automatically uses the TSDuck executables from the development repository.

```
$ ./run-all-tests.sh --dev
---- Testing TSDuck - The MPEG Transport Stream Toolkit - version 3.40-4133 from
/home/tsduck/bin/release-x86_64-vmubuntu
---- test-001: Test test-001.tsanalyze.full.txt passed
---- test-001: Test test-001.tsanalyze.full.log passed
.....
---- test-179: Test test-179.tstabdump.txt passed
---- test-179: Test test-179.tstabdump.log passed
---- ALL 2545 tests passed
$
```

The test suite is made of `bash` shell scripts. On Windows systems, the suite can be run from a "Git Bash" window for instance.

On UNIX systems (Linux, macOS, BSD), in the `tsduck` project, the `make` target `test-suite` runs the test suite on the binaries which were just built. It assumes that the work areas of the git repositories `tsduck` and `tsduck-test` are located side by side in the same parent directory. This target updates the `tsduck-test` repository from its remote origin and runs the test suite.

```
$ make test-suite
Already up to date.
---- Testing TSDuck - The MPEG Transport Stream Toolkit - version 3.40-4133 from
/home/tsduck/bin/release-x86_64-vmubuntu
---- test-001: Test test-001.tsanalyze.full.txt passed
---- test-001: Test test-001.tsanalyze.full.log passed
.....
---- test-179: Test test-179.tstabdump.txt passed
```

```
----- test-179: Test test-179.tstabdump.log passed
----- ALL 2545 tests passed
```



2.4.4. Running PSI/SI tests only

Some tests scripts in the suite are dedicated to some PSI/SI tables or descriptors. The scripts for such tests are usually simple and only invoke the standard script `standard-si-test.sh` using a specific XML input file which contains the PSI/SI tables or descriptors to test.

Running the complete test suite takes time. In some cases, when modifying code which may affect all PSI/SI, it is useful to run all related tests, without running the complete test suite. This can be done using option `--si-only`.

```
$ ./run-all-tests.sh --si-only
```

On UNIX systems (Linux, macOS, BSD), in the `tsduck` project, the `make` target `si-test-suite` runs that subset of the test suite.

```
$ make si-test-suite
```

2.4.5. Large files

The `tsduck-test` repository contains large files, typically transport stream files.

Initially, these files were not stored inside the regular GitHub repository. Instead, they used the [Git Large File Storage](#) (LFS) feature of GitHub. However, using LFS on GitHub happened to be a pain, as experienced by others and explained in [this article](#).

As a consequence, the transport stream files were re-integrated into the Git repository as regular files. But we now limit their size to 20 MB.

Chapter 3. Automation

TSDuck is a free open-source project which is maintained on spare time by very few developers (only one at the time of this writing). As a consequence, the maintenance of the project is driven by the lack of time and resource. To face these constraints, automation is essential.

This is briefly explained in [this presentation](#).

For future maintainers or contributors, this section lists a few automation procedures which are used by the project.

3.1. GitHub workflows

Automation is provided through *GitHub Actions*, the CI/CD infrastructure which is hosted at GitHub. This CI/CD infrastructure is provided for free by GitHub for open-source projects. See [\[GITHUB-CI\]](#) for the complete documentation of GitHub Actions.

A project which is hosted at GitHub can define its own *workflows*. A workflow is defined in a YAML file in the predefined subdirectory `.github/workflows`.

A workflow defines the conditions under which it is run:

- Automatically on events such as "push" or "pull requests".
- Automatically at scheduled times, for instance every day or every week.
- Manually from the GitHub web site.

TSDuck contains several GitHub workflows for automation. The most significant ones are described in the following sections.

Most TSDuck workflows can be manually run in addition to some automatic condition. To manually run a workflow, log in to GitHub, go to the project page, select the "Actions" tab, click on the workflow name in the left pane. Then, click on the "Run workflow" button on the right side. If the workflow defines input parameters, a pane is opened with the possible options. This feature is specifically useful with the "Release" workflow (see [section 4.1](#)).

3.2. Continuous integration

Each time a commit or a pull request is pushed to GitHub, the workflow `continuous-integration.yml` is automatically run to verify that the project can be built on most supported platforms and contains no regression.

- The workflow is run on Linux Ubuntu, macOS, and Windows (64 and 32 bits).
- The compilation is done using one or more revisions of the C++ language.
- On Linux, the builds are repeated using two compilers, `gcc` and `clang`.
- In each configuration of platform, compiler and language level:
 - TSDuck is fully built.
 - The unitary tests are run.
 - The full test-suite is downloaded and run.
 - The sample programs are built using a typical TSDuck installation of the binaries which were just built (Linux and macOS only).
- The documentation is rebuilt: user guide, developer guide, programming reference using doxygen. Missing documentation for new features can be found in the log. This is done on Ubuntu only.

Note that this workflow is quite stringent. It fails on any compilation warning on any operating system or

compiler. It fails on the slightest test failure in any configuration. It fails if any documentation is missing in the code, such as an undocumented parameter in a function.

3.3. Nightly builds

Every night, if there were any modification in the TSDuck repo during the last day, "nightly builds" are automatically produced and published on the [TSDuck web site](#). No manual action is required, everything is automated.

The workflow `nightly-build.yml` is automatically run every night (UTC). Note that the load of the GitHub CI/CD pipeline is increasing every year and the scheduled time is sometimes delayed by several hours. Therefore, there is no guarantee on the publication time.

- On Linux Ubuntu and Windows (64-bit Intel only), the TSDuck binary packages are built for the current state of the master branch in the repository.
- The produced binaries are installed on the CI/CD system.
- The full test-suite is downloaded and run.
- The complete set of documentation (user guide, developer guide, programming reference) is built and packaged in an archive.
- The Linux Ubuntu binaries, the Windows binaries and the documentation package are published as "artefacts" of the workflow. They are publicly downloadable from [GitHub](#).

At the end of the nightly build workflow, when all build jobs are completed, the update job triggers a signal on the [TSDuck web site](#). A PHP script is automatically run on the web site to retrieve, download and publish the latest [nightly binaries](#) and documentation.

3.4. Cleanup of long-standing issues

The [issues area on GitHub](#) is used to report problems, ask questions, and support any discussion about TSDuck. When an issue is obviously completed, because a complete answer was provided or a fix is pushed, the issue is closed. Sometimes, a plausible response or fix is provided but some feedback is expected from the user to confirm this. When a positive feedback is provided, the issue is closed.

However, some users never provide a feedback after their problem is solved. In that case, the issue remains open forever.

To solve this, there is a label named "close pending". When a plausible response, solution or fix is provided, the maintainer of the project sets the "close pending" label on the issue. It remains open. However, if the issue is not updated in the next 150 days, it will be automatically closed.

This is achieved by the workflow `cleanup-issues.yml`. This workflow is scheduled every week on Sunday at 02:00 UTC. It runs the Python script `pkg/github/close-pending.py` which automatically closes all issues with label "close pending" and no update within the last 150 days.

Chapter 4. Publishing new releases

Since TSDuck is maintained with little resource and time, there is no predefined roadmap and no product management cycle. Thanks to the continuous integration process (see [section 3.2](#)), any state of the repository can be used as a release candidate, as long as there is no issue in the CI process and no major development or fix is in progress.

As a rule of thumb, a new release is produced every 3 to 6 months, when enough new features or fixes are available, based on the [CHANGELOG file](#).

Publishing a release is roughly divided in two steps:

- Create and identify the release on GitHub:
 - Create a tag in the Git repository.
 - Build binary installers for Linux (Ubuntu, Debian, Fedora, Red Hat) and Windows on Intel x64 and Arm64 platforms.
 - Create the GitHub "release" with a descriptive text and publish the binaries.
- Post-release tasks (see [section 4.3](#)):
 - Publish the release on HomeBrew (for macOS) and WinGet (for Windows).
 - Update the source code with the next release number.

4.1. Automated release creation

The workflow [release.yml](#) is the preferred way to create a release. This is a versatile workflow which can build and publish binaries and documentation for TSDuck in several ways. Creating a release is only one of its features.

The workflow [release.yml](#) is run manually when the project is ready for a new release.

It has several inputs parameters which describe what to do:

- Build on Intel x64 and/or Arm64 targets. Both are created by default.
- Build on which Linux distros (Ubuntu, Debian, Fedora, Red Hat), using which container version. The container names are selected from <https://hub.docker.com/>, in the category "Operating Systems". The corresponding container is used to build the binary. By default, TSDuck is built on all distros, using the containers which are flagged as [latest](#). Selecting [none](#) as container name inhibits the build on the corresponding distro.
- Build on Windows. Only one runner is used, either an Intel or an Arm one. This single runner is used to build all targets. So, it does not make any difference for the build. The difference is in the tests. The test suite is run on the selected runner, testing the binary installer for this target. If you want to build only one target, it is better to use a runner of the same target to be able to run the test suite.
- Build the documentation set. This is disabled by default. This is typically used to immediately publish an update of the documentation on <https://tsduck.io/> without waiting for the the next "Nightly build". Use [nightly-builds](#) as indicated below.
- Create which type of release. The proposed types are:
 - [none](#): No release is created. This is the default. All binaries are left as artifacts of the workflow run and can be manually downloaded if necessary.
 - [nightly-builds](#): No new release is created. Instead, the binaries and documentation are published on <https://tsduck.io/> in the "Nightly builds" download section.
 - [update-last](#): No new release is created. The build binaries are added to the last published release. This can be useful to add variants of an operating system. For instance, you can create a new release with the Ubuntu packages being built for Ubuntu 24. Then, afterwards, you may re-run the workflow, building for

Ubuntu 25 (and no other distro) and update the release with the new packages.

- **draft**: Create a new GitHub release in "draft" state. This is the recommended procedure. When the workflow is completed and the draft release is published, you may review it and manually remove the "draft" state from the release. At this point, the release will become official.
- **pre-release**: Create a new GitHub pre-release. This option is currently not used with TSDuck. This is typically used in projects with "release candidates".
- **official**: Create an official and final release.

4.2. Manual release creation

Although the workflow `release.yml` is now the preferred way of creating a new release, it is still possible to manually create a release, with the support of several scripts.

Manually creating a TSDuck release is relatively easy. Although the build time for all binaries can be long, everything is automated and run in the background without requiring the maintainer's attention.

4.2.1. Building the various binaries

Binaries must be built for Windows (64 bits only) and a few major Linux distros (Fedora, RedHat, Ubuntu, Debian, and 32-bit Raspian). Additional binaries are now built for some Arm64 Linux distros.

The idea is to have a central system (Linux, macOS or Windows) from which the builds are orchestrated. The script `pkg/build-remote.sh` can be used to build TSDuck binaries on various remote systems and collect the resulting installers.

The remote systems can be physical systems or virtual machines running on the central system. When a virtual machines is not active, it is automatically booted, the binaries are built and collected, and then the virtual machines is shut down. The currently supported hypervisors to managed virtual machines are VirtualBox, VMware and Parallels Desktop (macOS). All systems shall be preconfigured so that the central user is authorized to connect on each of them using `ssh` without password (use public key authentication, not passwordless accounts!)

The maintainer shall typically have a personal script, outside the repository, which calls `build-remote.sh` on all systems.

Typical example of such a script, with one physical remote system (a Raspberry Pi) and 5 virtual machines on the local host:

```
$HOME/tsduck/pkg/build-remote.sh --host raspberry
$HOME/tsduck/pkg/build-remote.sh --host vmfedora --parallels Fedora
$HOME/tsduck/pkg/build-remote.sh --host vmredhat --parallels RedHat
$HOME/tsduck/pkg/build-remote.sh --host vmubuntu --parallels Ubuntu
$HOME/tsduck/pkg/build-remote.sh --host vmdebian --parallels Debian
$HOME/tsduck/pkg/build-remote.sh --host vmwindows --parallels Windows --windows --directory
Documents/tsduck --timeout 20
```

So, building a release is as simple as running that script. After "some time", all binaries and build log files are available in the `pkg/installers` subdirectory.

Be sure to check that all binaries are present. Check the log files if some of them are missing.

4.2.2. Creating the GitHub release

Once the binary installers are collected, simply run this command to create and publish the release on GitHub:

```
$ pkg/github/release.py --create
```

The current state of the repo on GitHub is used as base for the release. A tag is created with the version number. The release is created with a explanatory description. All binaries are published as assets of that release.

4.3. Post-release tasks

4.3.1. Creating the HomeBrew release

The GitHub release contains binaries for Linux and Windows. On macOS, TSDuck is distributed through [HomeBrew](#). The binaries for all macOS versions, Intel and Arm platforms, are built inside the HomeBrew CI/CD pipeline when an updated formula is pushed in a pull request.

Recently, TSDuck has been placed by the Homebrew maintainers in the ["autobump" list of packages](#). The original projects for the 2600+ packages in this list are regularly monitored by a Homebrew workflow. Whenever one of these projects publishes a new version, the update of the formula is automatically done.

This means that a new official TSDuck release on GitHub is automatically detected by HomeBrew and the macOS versions are automatically published.

So, there is nothing specific to do to publish a new TSDuck release on HomeBrew.

In case of problem, if HomeBrew does not detect the new release, it is still possible to trigger its creation using the manual procedure which was used before TSDuck was included in the list of "auto-bumped" projects: create the pull request for the new TSDuck version in HomeBrew using the following command:

```
$ pkg/homebrew/brew-bump-formula-pr.sh -f
```

The option `-f` forces the creation of the pull request. Without it, it works in "dry run" mode. It can be safe to start with a dry run, as a test.

When the HomeBrew CI/CD pipeline has finished to process the pull request, the new version of TSDuck is available for all flavors of macOS.

4.3.2. Creating the WinGet release

WinGet is somewhat similar to HomeBrew, for the Windows world. Because WinGet manages binary packages only, the NSIS installer for TSDuck shall be built first and published on GitHub as a release. Then, the WinGet manifest for TSDuck will point to that release on GitHub.

The PowerShell script `pkg/winget/release.ps1` is used to create and publish a new TSDuck release on winget. It must be run from a Windows system, after the creation of the GitHub release. The Windows binaries on GitHub are declared in WinGet.

4.3.3. Creating the FreeBSD port

On FreeBSD systems, TSDuck is managed as a "port". This is how TSDuck can be installed using the command `pkg install tsduck`.

Each time a new release is created, the maintainer shall request an update of the TSDuck port. See the detailed description in the [pkg/freebsd](#) directory.

4.3.4. Updating the version number

Once a release is published, the minor version number of TSDuck `TS_VERSION_MINOR` must be updated in the source file `src/libtscore/tsVersion.h`.

This is currently not automated and shall be manually updated before the first commit following the publication of a new release.

Chapter 5. Maintaining TSDuck Code

This chapter is meant to help TSDuck maintainers.

It covers various technical topics such as the modular structure of the project, adding new signalization tables and descriptors, the internal design of the `tsp` utility, the support of hardware devices, etc.

This chapter does not pretend to cover all technical topics in TSDuck and remains a work in progress. As usual, remember that "the answer is in the code".

The following figure illustrates the TSDuck software architecture.

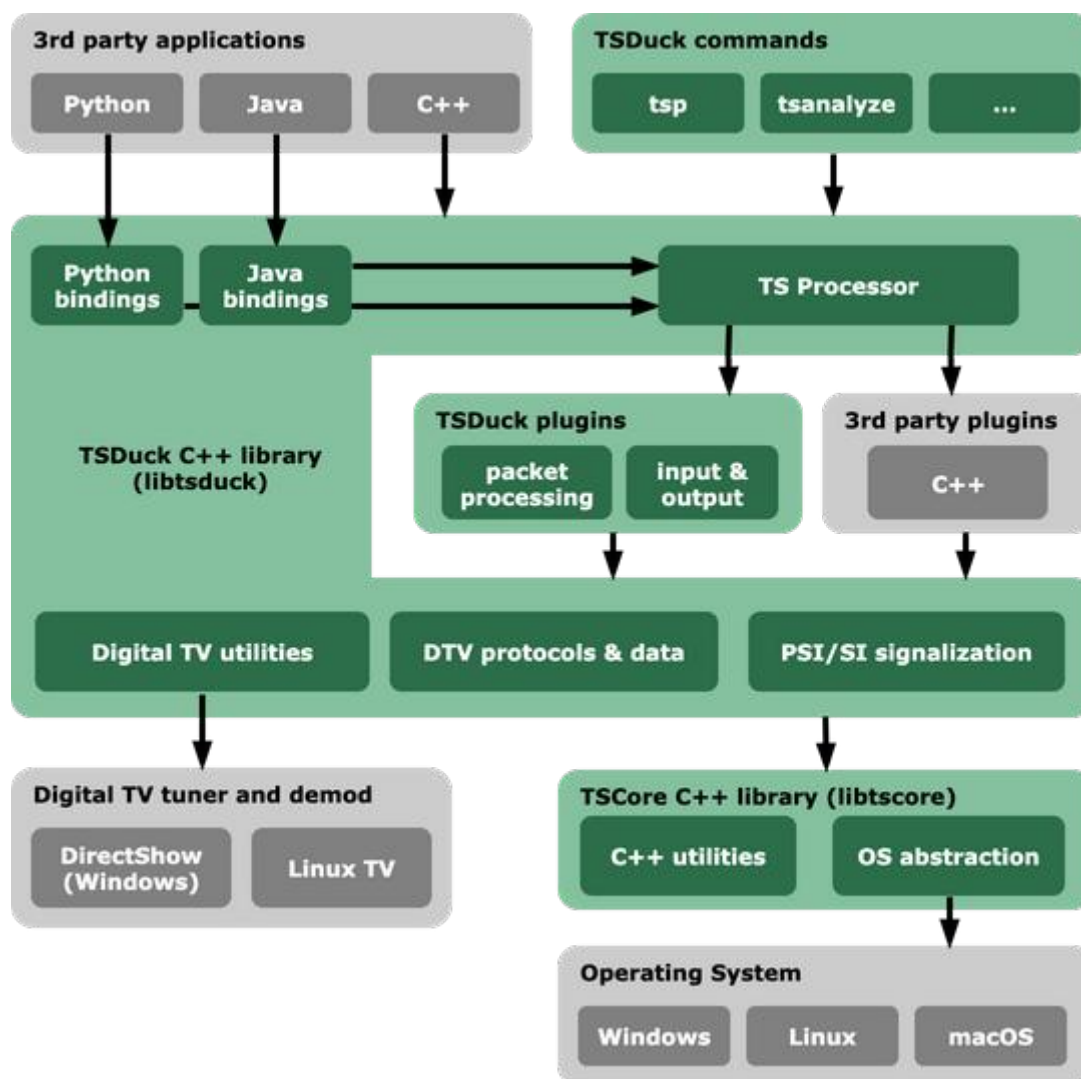


Figure 1. TSDuck software architecture

5.1. Modularity and self-registration

5.1.1. Principles of independence

For the sake of scalability and maintenance, TSDuck is designed in a modular way. A "module" is mostly a C++ class which is made of two files, a `.h` header file and a `.cpp` body file.



The C++20 revision of the language has introduced a new concept named `module` which is different. In this chapter, we use the term "module" in its general meaning.

When a module is not designed to create multiple instances of a class and could be implemented as inter-dependent functions managing a repository of global data, the *singleton* design pattern is used. See [section 6.4.1.3](#) for more details on singletons and shared data.

In rare cases, a module is just a collection of independent support functions, for operating system features for instance.

To avoid name clashing, all C++ declarations (classes, functions, types, constants) are placed in the namespace `ts`. All macros have a name starting with `TS_`.

The TSDuck libraries are roughly split into three hierarchical layers:

- **TSCore:** A collection of C++, operating system, and application utilities. This layer is not related to Digital TV. It is independently packaged in the TSCore library for the benefit of any application.
- **DTV:** A collections of C++ classes to manipulate Digital TV concepts. They are the building blocks for any C++ application in the broadcast or streaming field.
- **Plugins:** The plugin infrastructure which is used by TSDuck applications such as `tsp` and `tsswitch`. The corresponding classes can also be used by external applications which run chains of plugins to process transport streams.

These layers have a strict hierarchical organization. Classes in "TSCore" cannot reference classes in the next layers. Classes in "DTV" can reference classes from "TSCore" but not from "Plugins", etc.

The TSDuck commands are separate executables and the plugins are separate shared libraries. They are independently built and may reference anything from the TSDuck shared libraries.

A third common shared library is built on systems where Dektec devices are supported:

- **TSDehtec:** A collections of C++ classes to manipulate Dektec devices. Dektec provides a library named DTAPI to access their devices. It is available in binary form only and its source code is proprietary. The binary for DTAPI is also included in `libtsdehtec`. Because DTAPI was compiled by Dektec without `dllexport` directives on Windows, it cannot be used outside `libtsdehtec`. The `tsdehtec` command and the two `dehtec` plugins only use the TSDuck C++ classes from `libtsdehtec`, which act as proxies to DTAPI.

5.1.2. Source files structure

The `src` directory is structured as follow:

- `libtscore` contains all source files for the TSCore library. Any new source file inside this directory tree is automatically built as part of the TSCore library.
- `libtsduck` contains all source files for the TSDuck library. Any new source file inside this directory tree is automatically built as part of the TSDuck library.
- `libtsdehtec` contains all source files for the TSDehtec library. This is the code which is specific to Dektec devices. This shared library is not built on systems where Dektec devices are not supported.
- `tsplugins` contains all plugins. Any new source file inside this directory tree is automatically built as one separate shared library. This means that a plugin must be inside one single source file. However, it can rely on classes from the TSDuck library for most features.
- `tstools` contains all commands. Any new source file inside this directory tree is automatically built as one separate executable. This means that a command must be inside one single source file.

The additional subdirectories `utest` and `utils` contain the source code of unitary tests and packaging utilities which are not part of the TSDuck product itself.

There are several thousands of source files inside `libtscore`, `libtsduck`, and `libtsdehtec`. They are organized according to logical and technical constraints.

On a technical perspective, there are several strict rules based on the name of the subdirectories. By default, all

source files at any depth of subdirectories are compiled, included in the TSCore or TSDuck library. All `.h` header files are public by default. They are installed with the TSDuck development environment for third-party applications.

There are some exceptions, based on the name of the subdirectories:

- **private**: All source files in a subdirectory named **private**, anywhere below `src/libtscore` or `src/libtsduck`, is considered as private to the corresponding library. The header files will not be installed with the TSDuck development environment and the corresponding classes are not visible from applications.
- **linux, mac, freebsd, netbsd, openbsd, dragonflybsd, bsd, unix, windows**: A subdirectory with any of these names is specific to the corresponding operating system and is compiled only on that operating system. If there are header files, they are installed on the corresponding operating system. The names **unix** and **bsd** are more generic and are included on all UNIX systems and all BSD systems, respectively.

5.1.3. Windows constraints

All classes, functions, and objects in the TSDuck libraries are gathered into two shared libraries, `libtscore.so` and `libtsduck.so` on Linux and BSD, `libtscore.dylib` and `libtsduck.dylib` on macOS, `tscore.dll` and `tsduck.dll` on Windows,

A C++ class is exported out of a shared library through a set of mangled symbols, one per method or static field. This is straightforward on UNIX systems; this is only a matter of build commands, without impact on the source code.

On Windows, however, this is more complicated and painful (everything is more painful on Windows). By default, a symbol is not exported out of a DLL. A C++ class can be referenced only from inside the DLL but an external application which is linked against that DLL cannot reference any class in the DLL. Linking this application will result in "undefined symbol" errors.

It is possible to declare a class or function as exported from the DLL. However, this cannot be done using the build commands. Specific directives shall be present in the header files, only on Windows, and differently depending whether you include the header to build the library ("export") or to use the library from an application ("import").

All that dirt is hidden in TSDuck using the macros `TSCOREDLL` and `TSUCKDLL` which translate to an export directive when building the corresponding DLL, an import directive when using the DLL, and nothing on non-Windows system.

A similar macro `TSDEKTECDLL` exists for the `libtsdektec` library.

As a consequence, under `src/libtscore` and `src/libtsduck` (with the exception of **private** subdirectories), all public declarations in header files must have a `TSCOREDLL` or `TSUCKDLL` attribute. If this attribute is omitted, nothing changes on non-Windows systems and even on Windows systems when that declaration is not referenced from outside the DLL (it is not needed for references from inside the DLL). You may have the illusion that it works, until you compile an application on Windows which references that declaration and you get an "undefined symbol" error from the linker.

Take care to use the right attribute: `TSCOREDLL` in header files under `src/libtscore` and `TSUCKDLL` in header files under `src/libtsduck`. In case of mismatch, expect weird linking effects.

The `TSCOREDLL` and `TSUCKDLL` attributes must be used in classes, functions and variables which are defined at namespace level. For instance, in the TSDuck library:

```
namespace ts {  
    class TSUCKDLL UString { ... };  
    TSUCKDLL size_t GetProcessVirtualSize();  
}
```

Inline functions must have this attribute too. In the general case, the corresponding code will be inlined in the

application, referencing nothing in the DLL. However, when compiling for debug, without optimization, the code is not inlined and the application code references the symbol in the DLL.

```
namespace ts {
    TSDUCKDLL inline int SmallFunction() { ... }
}
```

Inside a class, the methods and fields do not need a `TSCOREDLL` or `TSDUCKDLL` attribute. They benefit from the attribute of the class. Inner classes, on the other hand, must have a `TSCOREDLL` or `TSDUCKDLL` attribute. See the example below.

```
namespace ts {
    class TSDUCKDLL Outer
    {
    public:
        // Method: no attribute.
        bool serialize() const;

        // Inner class: need an attribute.
        class TSDUCKDLL Inner { ... };
    };
}
```

As a general rule, private inner classes do not need a `TSCOREDLL` or `TSDUCKDLL` attribute. Being private, they cannot be referenced by the applications. However, they can be referenced inside inlined public methods of the same class. See the following example. The private inner class is not directly accessible from an application. However, when the code of the public method `ts::Outer::cool()` is inlined in the application, the binary of that application contains a reference to the symbol `ts::Outer::Inner::Foo()`. Without a `TSDUCKDLL` attribute on the inner class, even though it is private, an "undefined symbol" error will be generated.

Therefore, in doubt, also use a `TSCOREDLL` or `TSDUCKDLL` attribute in private inner classes.

```
namespace ts {
    class TSDUCKDLL Outer
    {
    public:
        void cool() { Inner::Foo(); }
    private:
        // Private inner class needs TSDUCKDLL if referenced by public inlined method.
        class TSDUCKDLL Inner
        {
        public:
            static void Foo();
        };
    };
}
```

Template classes and functions do not need any `TSCOREDLL` or `TSDUCKDLL` attribute. Their instantiations are completely generated inside the application.

5.1.4. Self registration of modules

The TSDuck code base is huge. For instance, there are hundredths of MPEG tables and descriptors with one class per table or descriptor. There are more than one hundred plugins. All tables, descriptors, plugins must be known from some central repository in order to be created when necessary.

New tables, descriptors, or plugins are regularly added, possibly by independent contributors. The challenge is to add new features or new objects without modifying the core code, without explicitly referencing them from some central structure.

To achieve this, TSDuck extensively uses the mechanism of *self-registration* of a module. To add a new descriptor, for instance, it is sufficient to create the `.h` and `.cpp` source files of the corresponding class, nothing else. The presence of the source files is sufficient to declare the descriptor.

This is possible thanks to a set of `TS_REGISTER` macros. As an example, inside the `.cpp` file of a descriptor, the macro `TS_REGISTER_DESCRIPTOR` must be placed at the top level of the source file, with a few parameters describing the descriptor. The same principle applies to extensions, plugins, tables, formatting functions, etc.

The `TS_REGISTER` macros generate code which is executed during the initialization of the application, before `main()` is called. See [section 6.4.1.3](#) for more details on initialization.

Therefore, just because the corresponding object file is present in the library, it will be initialized with the application. And of course, this initialization code will register the feature (extension, plugin, table, descriptor, etc.) into the right repository.

The tables below list all `TS_REGISTER` macros by category. Refer to the [TSDuck Programming Reference](#) for a complete documentation.

General-purpose registrations

The following macros are general-purpose. They register application-specific features into the TSCore library. This is how the TSCore library remains independent from the TSDuck library, without explicit reference to higher-level layers.

<code>TS_REGISTER_FEATURE</code>	Register a specific mandatory or optional feature, for instance a specific library or hardware acceleration. Depending on the parameters which are provided to the macro, the feature will be added to options <code>--version</code> or <code>--support</code> (in command <code>tsversion</code>). The user will be able to check if the feature is supported or not, and display the version.
<code>TS_REGISTER_CHRONO_UNIT</code>	TSDuck uses the C++17 standard template type <code>std::chrono::duration</code> (or <code>cn::duration</code>) for all durations. In addition to standard durations (milliseconds, seconds, hours, etc.), TSDuck defines all sorts of durations, especially for PCR, PTS, DTS and other timing units which are used in the DTV domain. In the base class <code>UString</code> , several methods format durations with the appropriate units. The macro <code>TS_REGISTER_CHRONO_UNIT</code> registers a new duration unit and how to format it.
<code>TS_REGISTER_BITRATE_CALCULATOR</code>	Register a C++ function which computes theoretical bitrates for a given type of modulation.

PSI/SI registrations

The following macros are used by developers who implement new tables or descriptors.

<code>TS_REGISTER_TABLE</code>	Register a fully implemented PSI/SI table. This macro is typically used in the <code>.cpp</code> file of a table.
<code>TS_REGISTER_SECTION</code>	Register a known table with a display functions but no full C++ class. This macro is typically used in the <code>.cpp</code> file of a CAS-specific module or TSDuck extension.
<code>TS_REGISTER_DESCRIPTOR</code>	Register a fully implemented PSI/SI descriptor. This macro is typically used in the <code>.cpp</code> file of a descriptor.
<code>TS_REGISTER_CA_DESCRIPTOR</code>	Register a display function for a <i>CA_descriptor</i> . This macro is typically used in the <code>.cpp</code> file of a CAS-specific module or TSDuck extension.

Plugins registrations

The following macros register a C++ class as a TSDuck plugin. They are typically used in the `.cpp` file of a plugin. If the plugin can be used in several roles (input, output, packet processing), there are as many difference C++ classes as roles and one macro is used per class.

<code>TS_REGISTER_INPUT_PLUGIN</code>	Register an input plugin class in the plugin repository.
<code>TS_REGISTER_OUTPUT_PLUGIN</code>	Register an output plugin class in the plugin repository.
<code>TS_REGISTER_PROCESSOR_PLUGIN</code>	Register a packet processing plugin class in the plugin repository.

Extension registrations

The following macros are typically used in an independant TSDuck extension. They are used to "hook" the extension inside the running TSDuck library. All previous `TS_REGISTER` macros are also used when necessary in extensions but the following ones are more specifically dedicated to extensions.

<code>TS_REGISTER_EXTENSION</code>	Register a TSDuck extension. This macro is typically used in the <code>tslibext_XXX.so</code> shared library of the extension named <code>XXX</code> . This is an optional macro, an extension can work without it. However, it helps identifying which extensions are loaded.
<code>TS_REGISTER_XML_FILE</code>	Register an extension XML model file for the PSI/SI tables and descriptors of that extension. The content is merged with the XML model of all tables and descriptors which are supported by TSDuck.
<code>TS_REGISTER_NAMES_FILE</code>	Register an extension <code>.names</code> file. All definitions are merged with the definitions which are provided by TSDuck (table ids, descriptor ids, etc.)
<code>TS_REGISTER_SECTION_FILTER</code>	Register a section filter, a class implementing <code>TablesLoggerFilterInterface</code> . Commands such as <code>tstables</code> and plugins such as <code>tables</code> are able to filter sections based on some criteria: table id, network id, etc. An extension can add its own filters based on its own criteria on its own private tables. The code of a section filter can define new command line options which are added to commands such as <code>tstables</code> and plugins such as <code>tables</code> . It also analyzes sections to filter them according to the value of the command line options.

5.2. Adding PSI/SI tables or descriptors

Adding support for new PSI/SI tables or descriptors is a welcome contribution to TSDuck. Users from various continents, using different standards, or participating in standardization processes, are in the best position to implement new tables or descriptors.

We recommend to release these contributions in open source, as part of the TSDuck project.

This section summarizes the main steps when implementing new tables or descriptors.

See [chapter 1](#) for more details on the contribution process.

5.2.1. Code base selection

The main recommendation to start with is: do not develop a new table or descriptor from scratch. Use an existing and proven one as code base and adapt to your new structure.

To identify the right existing structure as code base, use some of these criteria:

- Short single-section table vs. long multi-section table.
- Flat vs. structured section, with descriptors or not, including substructures with descriptors (e.g. PMT).

- Use of strings, DVB strings, ATSC strings, ISDB strings.
- Public descriptor, DVB private descriptor, table-specific descriptor.

5.2.1.1. Classification of descriptors

Currently, TSDuck implements nearly 300 different descriptors from various standards. Because the number of possible descriptor *tags* (a.k.a. descriptor ids) is limited to 256 values, there is no room for all possible descriptors. For this reason, the various standard organizations use tricks such as *extended descriptors*, *private descriptors*, or *table-specific descriptors*.

In a MPEG/DVB context, the allocation of descriptor ids is the following:

0x00-0x3E	MPEG-defined descriptors
0x3F	MPEG extension descriptor
0x40-0x7E	DVB-defined descriptors
0x7F	DVB extension descriptor
0x80-0xFE	Private descriptors
0xFF	Reserved, not allocated, typically used as "null" value

5.2.1.2. Extended descriptors

MPEG and DVB separately define the concept of *extended descriptors*. Because of the shortage of descriptor ids, each of the two standards has defined an *extension_descriptor*, typically using their last allocated descriptor id. This descriptor is a generic envelope for specialized descriptors.

The first byte of the descriptor payload is an *extended_descriptor_id* which identifies the actual descriptor type. This allows the definition of up to 256 additional descriptors.

TSDuck does not use any specific type for the generic *extended_descriptor*. Instead, there is a distinct type for each form of *extended_descriptor*. Each of them has its own XML element and C++, just like any other descriptor. The extended descriptor mechanism is only considered as a binary serialization detail, not a different type of descriptor.

If you have to implement a MPEG-defined extended descriptor, you may use the *HEVC_timing_and_HRD_descriptor* as code base.

If you have to implement a DVB-defined extended descriptor, you may use the *supplementary_audio_descriptor* as code base.

5.2.1.3. Private descriptors (DVB)

In the DVB standard, descriptor ids 0x80 to 0xFE are "private". They are reserved for use by private entities, typically TV operators, broadcast equipment vendors, Conditional Access System (CAS) vendors.

To determine which semantics should be associated with a given descriptor id in that range, DVB defines a *private_data_specifier_descriptor* which contains a 32-bit *private_data_specifier* (PDS). DVB allocates a unique PDS to any private organization which requests it. See [ETSI-101-162](#) and [the DVB services](#) for more details on allocated PDS values.

In a descriptor list, all private descriptors which come after a *private_data_specifier_descriptor* are defined by the private organization which is identified in the *private_data_specifier_descriptor*. If several private descriptors from distinct defining entities must be placed in the same descriptor list, several *private_data_specifier_descriptor* are allowed to switch from one entity to another.

In a DVB context, a descriptor with a tag in the range 0x80-0xFE and no preceding *private_data_specifier_descriptor*

is illegal.

That being said, in practice, the experience has exhibited families of bugs:

- Rogue signalization: Some organizations define their own private descriptors, with descriptor ids in the range 0x80-0xFE, without officially allocated PDS value. They use their own arbitrarily defined PDS value, hoping that no one else will use the same PDS value.
- Signalization bugs: Some broadcasters "forget" to insert the right *private_data_specifier_descriptor* before a well-defined private descriptor. (This is the reason why the option `--default-pds` exists in several TSDuck commands and plugins).
- Implementation bugs: Some implementers of receivers (set-top box, TV set) ignore the PDS rule or forget to check the previous PDS. They blindly interpret some private descriptor based on some expected descriptor id in the range 0x80-0xFE. If the same private descriptor id is used in the context of another PDS, the receiver incorrectly interprets the binary descriptor.

All these bugs are real and were regularly found during the development and usage of TSDuck. If you implement a private descriptor, be sure to follow the rules.

You may use the EACEM-defined *eacem_stream_identifier_descriptor* as code base.

Logical channel number descriptors: Most of the commonly used private descriptors are some forms of *logical_channel_number_descriptor*. The Logical Channel Number (LCN) is the usual concept of TV channel number, the oldest and most traditional way of identifying a TV channel. Surprisingly, neither MPEG nor DVB defined it. Therefore, operators or equipment vendors have to define their own way of identifying LCN's. For this reason, there is a wide range of variants of private *logical_channel_number_descriptor* which all contain the same kind of information. If you implement such a descriptor, with a similar implementation, your class should be a subclass of `AbstractLogicalChannelDescriptor`. Use *eacem_logical_channel_number_descriptor* as code base.

5.2.1.4. Table-specific descriptors

So-called *table-specific descriptors* are specific descriptors which exist only in the context of a couple of specific tables. They usually re-use the tag of a standard descriptor, typically in the MPEG-defined range. Of course, it is assumed that the standard descriptor, the tag of which has been hijacked, will never be used in those specific tables to avoid ambiguities.

Let's take an example, the *target_IP_address_descriptor*. This is a DVB-defined descriptor which can be used only inside an INT or a UNT, two DVB-defined tables. The *target_IP_address_descriptor* uses tag 0x09, which is normally used by a MPEG-defined *CA_descriptor*. When TSDuck analyzes a descriptor list and encounters a tag 0x09, it usually starts to analyze a *CA_descriptor*, except when the table is an INT or a UNT, in which case it analyzes a *target_IP_address_descriptor*.

This situation is supported by TSDuck. If you have to implement such a table-specific descriptor, use *target_IP_address_descriptor* as code base.

5.2.2. Affiliation to a standard

Each table or descriptor is defined either by a standard body or an organization, committee or private company. Check if other PSI/SI from this organization is already implemented in TSDuck. This is important because source files for PSI/SI are organized by standard.

Tables are implemented in the directory `src/libtsduck/dtv/tables`, using one subdirectory per standard. The current subdirectories are `atsc`, `dvb`, `isdb`, `mpeg`, `scte`. Currently, only renown standard bodies define tables.

Descriptors are implemented in the directory `src/libtsduck/dtv/descriptors`, using one subdirectory per standard. The current subdirectories are the same as tables, plus various organizations such as `eacem`, `dtg`, `aom`, or `uwa`, plus private companies which define private DVB descriptors such as `eutelsat`, or `sky`.

Try to find the right subdirectory for your new structure. Create another directory if required.

In that subdirectory, you will have to create four or five files (the last is optional). For instance, the MPEG-defined *ISO_639_language_descriptor* is implemented as:

```
src/libtsduck/dtv/descriptors/mpeg:

    tsISO639LanguageDescriptor.xml
    tsISO639LanguageDescriptor.adoc
    tsISO639LanguageDescriptor.h
    tsISO639LanguageDescriptor.cpp
    tsISO639LanguageDescriptor.names
```

More details follow in the next sections.

5.2.3. Declaring identifiers

Your table or descriptor must have a 8-bit identifier. You need to add it in the TSDuck source code.

Table ids are defined in file `src/libtsduck/dtv/signalization/tsTID.h`, in enum list `TID`. Descriptor ids are defined in file `src/libtsduck/dtv/signalization/tsDID.h`, in enum list `DID`. The ids are grouped by standard, be sure to add it at the right place.

In the case of a table, if that table is expected on some predefined PID, also add this PID in file `src/libtsduck/dtv/transport/tsTS.h`, in the enum list `PID`.

In the case of a private DVB descriptor, your descriptor is valid only after a *private_data_specifier_descriptor* which contains the *private_data_specifier* (PDS) of the organization which defines the descriptor. Check if that PDS value is present in file `src/libtsduck/dtv/signalization/tsPDS.h`, in enum list `PDS`. Add it if not present.

For TSDuck to display meaningful identifiers, the source tree contains *names files*, with a `.names` extensions. These files associate a unique value with a name. There are several sections (for PID, TID, DID, for instance). In each section, a value can be present only once and values must be declared in ascending order.

Add the table, descriptor, or PDS name in the files `src/libtsduck/dtv/signalization/tsTID.names`, `tsDID.names`, or `tsPDS.names`. Carefully read the comments at the beginning of each section. They explain the encoding of each unique value.

For table ids, the value includes the standard and the optional *CAS_id* (useful for ECM and EMM only).

For descriptor ids, the value includes the standard, the PDS for private descriptors, or the *table_id* for table-specific descriptors.

If you implement a MPEG-defined or DVB-defined extended descriptor, add the corresponding *extended_descriptor_id* in `src/libtsduck/dtv/signalization/tsDID.h`, in enum lists with `XDID_MPEG_` and `XDID_DVB_` symbols.

5.2.4. XML definition

You must define an XML representation for your table or descriptor in a `.xml` file. Use the selected code base as reference.

This XML file is an *XML model file*, as defined in the TSDuck user guide.

A table shall be defined as one XML element inside the following envelope:

```
<?xml version="1.0" encoding="UTF-8"?>
<tsduck>
  <_tables>
    <my_table_name ...>
      <_any in="_metadata"/>
```

```
...
    </my_table_name>
  </_tables>
</tsduck>
```

Note the mandatory `<_any in="_metadata"/>`.

A descriptor shall be defined as one XML element inside the following envelope:

```
<?xml version="1.0" encoding="UTF-8"?>
<tsduck>
  <_descriptors>
    <my_descriptor_name ...>
      ...
    </my_descriptor_name>
  </_descriptors>
</tsduck>
```

For attributes and element names, preferably use the exact same names as defined in the standard for your table or descriptor.

Do not blindly copy the binary structure in the XML description. Define an XML equivalent representation.

For instance, a common pattern for optional fields in binary structures is to define a one-bit *foo_flag* and a subsequent optional *foo* field. The *foo* field is typically present only when *foo_flag* is 1. Do not define *foo_flag* in the XML structure. Just define a *foo* attribute and document it as optional. It is the work of your serialization and deserialization functions to interpret the presence or absence of the XML attribute *foo* as a *foo_flag* value.

The template value of XML attributes is a short informal type declaration. For integer values, always start the description string with `uintN` or `intN`, when *N* is the size in bits of the binary field. This `uintN` or `intN` is used by the automatic XML-to-JSON translation to generate a JSON number instead of a JSON string.

Your `.xml` file will be automatically grabbed by the TSDuck build system and integrated into the final configuration files.

5.2.5. Documentation file

You must describe the documentation for your table or descriptor in `asciidoc` format in a `.adoc` file. This file will be automatically grabbed by the TSDuck build system and integrated into the TSDuck user guide, in HTML and PDF format.

Even if you do not know this format, this is pretty simple. All `.adoc` files for the signalization use the same pattern. The following is the file for the *ISO_639_language_descriptor*.

```
==== ISO_639_language_descriptor

Defined by MPEG in <<ISO-13818-1>>.

[source,xml]
----
<ISO_639_language_descriptor>
  <!-- One per language -->
  <language code="char3, required" audio_type="uint8, required"/>
</ISO_639_language_descriptor>
----
```

The first line starts with `====` and contains the table or descriptor name. This is the section title in the user guide.

The sentence *"Defined by XXX in <<YYY>>"* is mandatory and indicates where to find the reference documentation for the structure. The syntax *<<YYY>>* is a reference in the bibliography. There must be a corresponding ** [[YYY]]* entry in the file `doc/user/20F-app-references.adoc`, for instance:

```
* [[[ISO-13818-1]]] ISO/IEC 13818-1:2018 | ITU-T Recommendation H.222 (2017):
  "Generic coding of moving pictures and associated audio information: Systems" (also known as "MPEG-2
  System Layer").
```

If the reference does not exist yet in the bibliography, add it. Keep the references sorted in alphabetical order.

Add a copy of the XML description of the table or descriptor.

- Remove the enclosing `<tsduck>`, `<_tables>`, `<_descriptors>` structures, just keep your structure.
- Make sure to reindent its top-level tag to the first column.
- In tables, remove the `<_any in="_metadata"/>`. It is meaningless for the user.
- Add any comment or formatting which makes the result more informative to the user.

5.2.6. C++ class

The C++ header (`.h`) and body (`.cpp`) files for the table or descriptor class are mandatory. Start with the selected code base and carefully replace the structure names.

5.2.6.1. Structure registration

In the `.cpp` file, there is a fundamental macro: `TS_REGISTER_TABLE()` for tables and `TS_REGISTER_DESCRIPTOR()` for descriptors. This is a C++ trick which automatically registers your structure in the PSI/SI repository during the initialization of the module. If you omit this macro, your table will not be recognized.

The registration macro may take various forms depending on the type of structure (standard descriptor, table-specific descriptor, extended descriptor, etc.) Be careful to select a code base with the same characteristics in order to copy the same type of registration.

5.2.6.2. Programming reference

All public structures and fields in the C++ header file must be documented using Doxygen tags. See examples in existing structures. This is the way your structure will become documented in the TSDuck Programming Reference (doxygen format).

No public element shall be left undocumented. To verify this, generate the documentation and check any error. Undocumented elements are reported.

- On UNIX systems (Linux, macOS, BSD), run `make doxygen`.
- On Windows systems, run the PowerShell script `doc\doxy\build-doxygen.ps1`.

In the initial descriptor of your C++ class, make sure it is properly identified with the right group and standard. For instance:

```
//!
//! Representation of a Program Association Table (PAT).
//! @see ISO/IEC 13818-1, ITU-T Rec. H.222.0, 2.4.4.3
//! @ingroup libtsduck table
//!
```

or:

```

//!
//! Representation of an ISO_639_language_descriptor
//! @see ISO/IEC 13818-1, ITU-T Rec. H.222.0, 2.6.18.
//! @ingroup libtsduck descriptor
//!

```

The directive `@ingroup` is used by Doxygen to assign the class in the right group. The group name `libtsduck` indicates that the class belongs to the TSDuck library, and not the TSCore library. The group names `table` and `descriptor` are self-explanatory.

The directive `@see` is important in three ways.

1. It is included in the Doxygen documentation.
2. It helps the future maintainers of the code to find the right documentation and directly the section number where to look.
3. It is also used in the automatic generation of the PSI/SI signalization reference section of the [TSDuck Developer Guide](#).

5.2.7. Names file

If necessary, you may provide a `.names` file. This is useful when a field of your structure can get distinct values with distinct meanings. When displaying a structure, it is more convenient for the user to get a meaningful name rather than a value.

A `.names` file is organized in several sections. By convention, use section names which start with the XML name of your structure, followed by a dot.

Example of the file `tsISO639LanguageDescriptor.names`, for the `ISO_639_language_descriptor`:

```

[ISO_639_language_descriptor.audio_type]
Bits = 8
0x00 = undefined
0x01 = clean effects
0x02 = hearing impaired
0x03 = visual impaired commentary

```

In the C++ source file, use the inherited static method `DataName()` to retrieve a meaningful name, with optional formatting of the value before or after the name.

Example of the file `tsISO639LanguageDescriptor.cpp`:

```

void ts::ISO639LanguageDescriptor::DisplayDescriptor(TablesDisplay& disp, PSIBuffer& buf, const
UString& margin, DID did, TID tid, PDS pds)
{
    ...
    disp << " , Type: " << DataName(MY_XML_NAME, u"audio_type", buf.getUInt8(), NamesFlags::FIRST) <<
    std::endl;
}

```

Your `.names` file will be automatically grabbed by the TSDuck build system and integrated into the final configuration files.

5.2.8. Tests

There are lots of traps and pitfalls in the coding of a table or descriptor. It is crucial to test it thoroughly.

First, become familiar with the TSDuck test suite as described in [section 2.4](#).

Once you have cloned your forked versions of the two repositories, `tsduck` and `tsduck-test`, side by side in the same parent directory, you can implement a test for your table or descriptor.

This kind of test is standardized. The idea is to start from an XML file containing several samples of your table or descriptor. Then, invoke the common script `standard-si-test.sh`.

This standard test compiles the XML file in binary, decompiles it to generate XML and JSON, recompiles the output, inject the tables in a transport stream, extract them in text form, etc. All intermediate results are kept as reference.

This kind of test is interesting in two ways. First, during the initial test, after development, it is a good tool to debug the serialization, deserialization, binary and XML. Second, the reference outputs will track any future regression.

For instance, the test 027 is the reference test for SCTE 35 tables and descriptors. All tested structures are in the file `tsduck-test/input/test-027.xml`. The test script `tsduck-test/tests/test-027.sh` is very simple:

```
#!/usr/bin/env bash
source $(dirname $0)/../common/testrc.sh
test_cleanup "$SCRIPT.*"
source "$COMMONDIR"/standard-si-test.sh $SCRIPT.xml
```



In practice, *all* test scripts for that kind of PSI/SI test are identical. Only the input `.xml` file changes.

If your table or descriptor belongs to a set of structures which are already tested in an existing test, you may simply add your tested XML definitions in the existing test and update its reference output.

Otherwise, especially if you plan to implement several structures, you may create a new test. Just use existing tests with `standard-si-test.sh` as a starting point.

Pay attention to the XML structures you want to test. Keep in mind that you test one given structure in all possible ways, regardless of real applications. Your tested structures do not need to carry meaningful values. You test the *syntax* of your table or descriptor, not its *semantics*. You just want to test code, nothing else.

Here are some guidelines:

- If you test a descriptor, you don't care about which table it is in. Use a `<CAT>` for instance, a table which only contains descriptors and nothing else.
- If you test a table which contains descriptors, use any kind of simple descriptors, *ISO_639_language_descriptor* for instance. You do not care if such a descriptor does not make sense in your table.
- If you test a table which contains descriptors, test each descriptor list with zero, one, two descriptors.
- Test optional fields in structures where they are present and in other structures where they are omitted.
- More generally, when your code takes different steps or branches in the presence of different forms of input, test all possible forms of input.
- Test adjacent fields with different values. If two flags are in consecutive bits in the binary structure, test once with a `true/false` combination and once with a `false/true` combination.
- Use integer values which use the full width of a binary field to detect incorrect truncation or size errors. For instance, in a `uint32` field, use value `0xDEADBEEF`, for instance, not 0 or 1.

When the result is satisfactory, submit a pull request for each repository, `tsduck` and `tsduck-test`. See [section 1.3](#) for more details on that.

5.3. TSP design

This section is a brief description of the design and internals of `tsp`. It contains some reference information for `tsp` maintainers.

`tsp` is designed to clearly separate the technical aspects of the buffer management and dynamics of a chain of plugins from the specialized plugin processing (TS input, TS output, packet processing).

5.3.1. Plugin Executors

Each plugin executes in a separate thread. The base class for all plugin threads is `ts::tsp::PluginExecutor`. Derived classes are used for input, output and packet processing plugins.

5.3.2. Transport packets buffer

There is a global buffer for TS packets. Its structure is optimized for best performance.

The input thread writes incoming packets in the buffer. All packet processors update the packets and the output thread picks them at the same place. No packet is copied or moved in memory.

The buffer is an array of `ts::TSPacket`. It is a memory-resident buffer, locked in physical memory to avoid virtual memory paging (see class `ts::ResidentBuffer`).

The buffer is managed in a circular way. It is divided into logical areas, one per plugin thread (including input and output). These logical areas are sliding windows which move when packets are processed.

Inside a `ts::tsp::PluginExecutor` object, the sliding window which is currently assigned to the plugin thread is defined by the index of its first packet (`_pkt_first`) and its size in packets (`_pkt_cnt`).

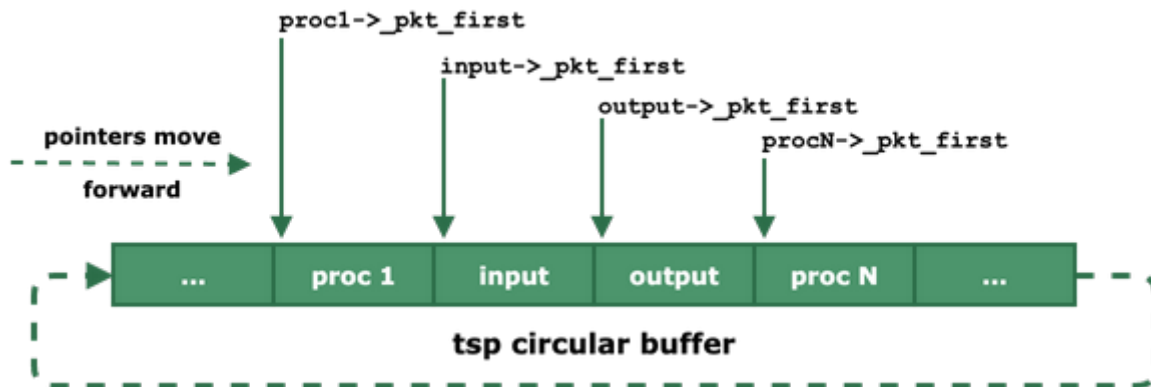


Figure 2. Flat (non-circular) view of the buffer

When a thread terminates the processing of a bunch of packets, it moves up its first index and, consequently, decreases the size of its own area and accordingly increases the size of the area of the next plugin.

The modification of the starting index and size of any area must be performed under the protection of a mutex. There is one global mutex for simplicity. The resulting bottleneck is not so important since updating a few pointers is fast.

When the sliding window of a plugin is empty, the plugin thread sleeps on its `_to_do` condition variable. Consequently, when a thread passes packets to the next plugin (ie. increases the size of the sliding window of the next plugin), it must notify the `_to_do` condition variable of the next thread.

When a packet processor decides to drop a packet, the synchronization byte (first byte of the packet, normally 0x47) is reset to zero. When a packet processor or the output executor encounters a packet starting with a zero byte, it ignores it. Note that this is transparent to the plugin code in the shared library. The check is performed by the `ts::tsp::ProcessorExecutor` and `ts::tsp::OutputExecutor` objects. When a packet is marked as dropped, the

plugin is not invoked.

All `ts::tsp::PluginExecutor` are chained in a ring. The first one is input and the last one is output. The output points back to the input so that the output executor can easily pass free packets to be reused by the input executor.

The `_input_end` flag indicates that there is no more packet to process after those in the plugin's area. This condition is signaled by the previous plugin in the chain. All plugins, except the output plugin, may signal this condition to their successor.

The `_aborted` flag indicates that the current plugin has encountered an error and has ceased to accept packets. This condition is checked by the previous plugin in the chain (which, in turn, will declare itself as aborted). All plugins, except the input plugin may signal this condition. In case of error, all plugins should also declare an `_input_end` to their successor.

5.4. Hardware support

This section describes some specificities of the various types of hardware devices which are supported by TSDuck and their impact on the code.

The TSDuck user guide has a complete chapter on each of these classes of devices. It describes the devices, their naming, usage, supported platforms, required drivers, etc. In the Developer Guide, we focus on the structure of the code, for TSDuck maintainers.

5.4.1. Standard tuner receiver devices (DVB, ATSC, ISDB)

The code for tuners is located in `src/libtsduck/dtv/broadcast`.

The base class for all tuners is `TunerBase`. It exposes the common interface of all tuners. It is normally never directly used by an application.

A tuner is accessed by an application using the subclass `Tuner`.

When opening a `Tuner` instance, depending on the provided name, the tuner is either a physical device or a software emulator. The `Tuner` instance creates either an internal `TunerDevice` object or an internal `TunerEmulator` object. These two classes are also subclasses of `TunerBase`. Once a `Tuner` is opened, all services are redirected to the internal `TunerDevice` or `TunerEmulator` object.

The subclass `TunerEmulator` is a pure software implementation which is available on all systems. The name of such a device in a XML file which describes all the capabilities of the "tuner" and the available frequencies and corresponding modulation parameters. The TSDuck user guide describes the tuner emulator in details.

The `TunerDevice` class is highly dependent on the operating system. The source files are different and placed in system-specific directories. The subdirectories `linux` and `windows` contain actual implementations for these operating systems. The subdirectories `mac` and `bsd` contain stub versions which report errors when used.

Linux

The Linux implementation is relatively straightforward using the Linux TV API. This API evolves on a regular basis. New modulations are added.

It is sometimes useful to review the evolutions of the Linux TV API to detect new modulations and new parameters. In the `broadcast/linux` subdirectory, there is a Bash script named `get-linux-tv-api-history.sh` which extracts from the Linux kernel repository a snapshot of each new version of the Linux TV, with a `diff` file of each version.

Windows

The Windows implementation is more problematic. The tuner devices are managed by DirectShow, a Microsoft framework for multimedia. The specific subsystem of DirectShow for TV receiver devices is BDA (Broadcast Device

Architecture). Most of the time, the hardware vendors provide BDA drivers for their receivers. Windows does not include any predefined BDA driver.

The DirectShow framework is complex. It is based on the concept of "graphs of filters". Each DirectShow filter is a COM object which exposes standard interfaces. Media processing starts with building a graph with filters, and connecting these filters together using the standard interfaces they expose.

Unlike Linux, there is no standard interface to the various BDA drivers. Each vendor defines its own interface. A vendor also provides its own DirectShow filters, typically a "Tuner Filter" and an optional "Receiver / Capture Filter". The proprietary tuner filter communicates with the proprietary BDA driver. An application can only use the standard interface of the DirectShow filters.

However, a filter cannot be used alone. It must be part of a DirectShow graph. This graph must be complete and consistent, with no dangling interface. This is the condition to "start the graph". Starting the graph is required to read incoming packets from the tuner.

To make things worse, there is no single way to build a tuner reception graph. It depends on the proprietary filters. Do they include a receiver / capture filter for instance? Some combinations of filters work, some others don't. Depending on the tuner device, various types of graphs shall be used, or different filters of the same types.

When building a graph, the application "connects" filters together, by pairs. A filter can accept or reject the connection. Therefore, in the absence of deterministic method, TSDuck uses a flexible approach. Several "standard" types of graphs are tested. At each node of the graph, we search all available filters of the required type for that node. Normally, only one combination of filters can be built, with all connections being accepted by all filters. This is the method which is used by TSDuck, testing all possible graphs for the tuner until one is accepted.

This looks complicated (and it is) but the standard use case for a tuner on Windows is simpler. When you purchase a tuner, it comes with all software: BDA driver, DirectShow filters, and TV watching application. Since the application is provided by the tuner vendor, for that tuner, there is only one way to build a graph for that tuner and it is well known by the tuner vendor. With TSDuck, we have a less common use case. The application is the same for all tuners and expects to uniformly read the transport stream from all of them. Simple on Linux, awfully complex on Windows.

Another annoying property of DirectShow is the absence of data output filter. A DirectShow graph is designed to process media. Typically, it reads media from somewhere (a tuner in our case), processes them, filters them, and outputs the result on screen. In short, DirectShow is designed to watch TV... There is no way to do such a simple thing as reading TS packets from the tuner device. Therefore, we were obliged to develop our own DirectShow filter. This filter is inserted in the graph and collects the TS packets using the DirectShow standard interfaces. This custom DirectShow filter is implemented in the class `SinkFilter`, with its source code in the `broadcast/windows` subdirectory.

Because everything is annoying in DirectShow, it uses a "push" model. This means that, once the graph is started, packets are pushed from the tuner all along the graph. The applications which needs packets, on the contrary, have a "pull" model. When packets are needed, it "reads" packets, like reading packets on the network, reading data from a pipe, etc. Therefore, the passing mechanism is asynchronous. When incoming TS packets are pushed up to our `SinkFilter`, the filter passes them through a message queue. The application, on the other hand, consumes packets in pull mode from the message queue. This is performed in the Windows version of the `TunerDevice` class.

Because everything is *really* annoying in DirectShow, there are other constraints. For instance, a tuner reception graph cannot work without a demux filter. In our case, we don't demux anything because TSDuck wants the complete transport stream. Anyway, DirectShow needs it or you won't get any packet. Additionally, even though it would be logical to insert the demux in the middle of the data flow, DirectShow wants it in a dead-end branch of the graph. Don't try to understand... To paraphrase a well-known cryptographer, it seems that DirectShow has been "designed by the E-commerce Department of the Ministry of Silly Walks".

The following diagram represents a typical DirectShow reception graph for a tuner. It helps understanding the code. Given the complexity of this code, it is not recommended to modify it, unless completely and absolutely

necessary.

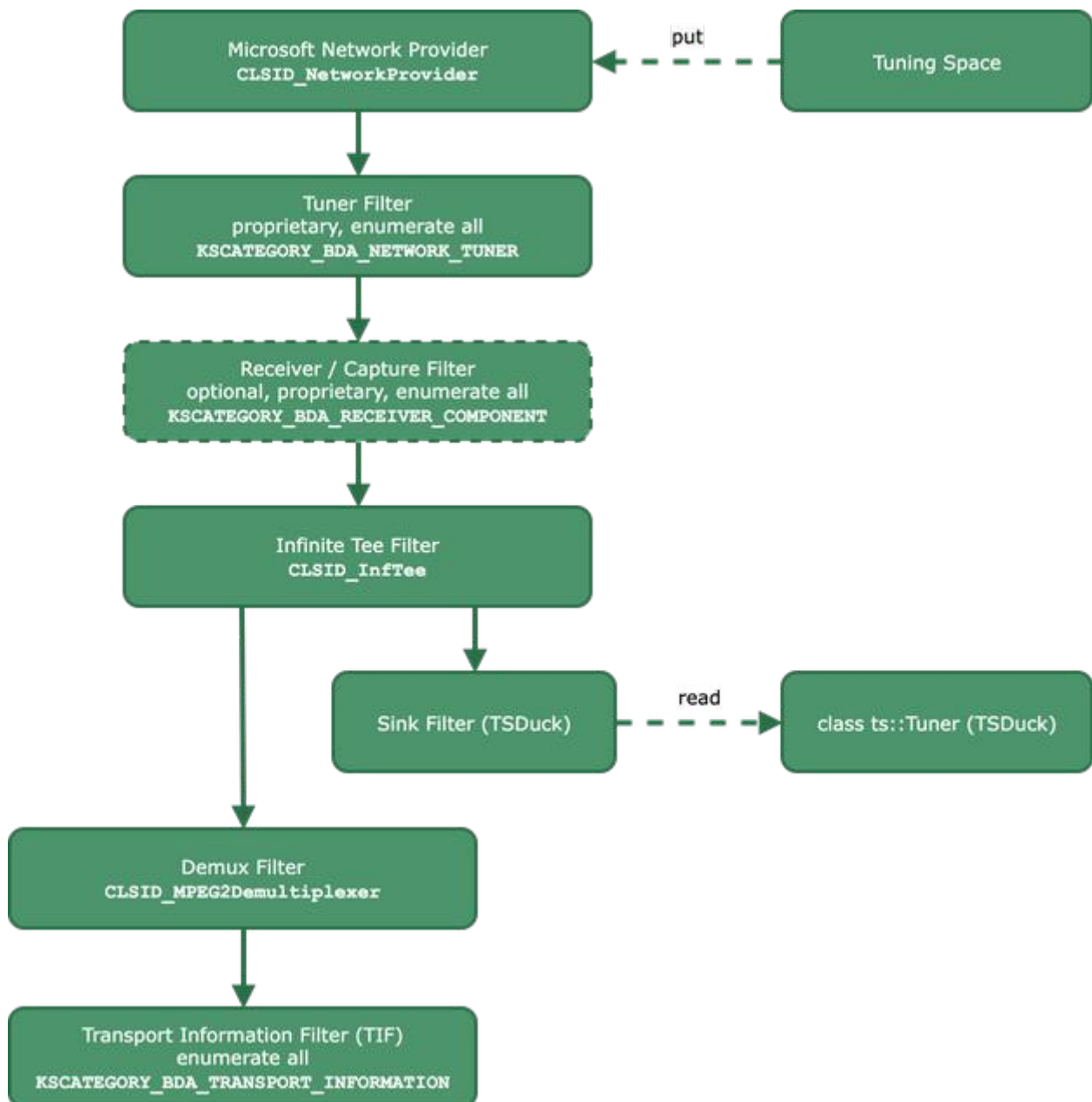


Figure 3. Usage of Windows DirectShow interfaces in TSDuck

As a Grande Finale, let's note that DirectShow received no update, no improvement, no new features, for years. It looks like abandonware, without any new alternative. Recent modulations are not supported on Windows and will probably never be.

In short, if you want to receive broadcast transport streams from a tuner, use Linux...

ISDB-T 204-byte packets

The ISDB-T and ISDB-Tb standards have a specific form of 204-byte packets. In other modulations, the standard 188-byte TS packet is followed by a 16-byte Reed-Solomon outer forward error correction code. These 16 bytes are not interesting and dropped by the demodulator. In ISDB-T and ISDB-Tb, the 16-byte trailer is split in two equal parts. The Reed-Solomon FEC is reduced to 8 bytes, at the end of the trailer. The first 8 bytes of the trailer contain ISDB information such as modulation layer, frame counters, system identification. This information can be analyzed by TSDuck, when available. However, neither the Linux TV nor the Windows DirectShow API's are able to return 204-byte packets. Therefore, it is not possible to analyze ISDB-T information from the `dvb` plugin.

5.4.2. Dektec devices

Using Dektec devices requires kernel-mode device drivers which must be separately installed. Building TSDuck code does not depend on the device drivers. They are only required to use Dektec devices. The TSDuck user guide describes how to install the Dektec drivers.

TSDuck does not call the device drivers directly. Dektec provides a user-mode API for C++ called "DTAPI". The DTAPI is well documented and well designed. Its interface is the same over all operating systems. This library is provided for free by Dektec in binary format only. The source code is proprietary and not publicly available.

On Windows, the DTAPI shall be installed first. This is automated in the PowerShell script `install-dektec.ps1`, which is also automatically invoked by `install-prerequisites.ps1`. The "Dektec Windows SDK" installs the device drivers and the DTAPI system-wide.

On Linux, the DTAPI is not "installed". It is automatically downloaded from <https://dektec.com> during the build process of the shared library `libtsdektec`. The DTAPI is locally installed in the `bin` build area of the TSDuck project. It is therefore deleted as part of `make clean`.

Several versions of the DTAPI binaries are installed, depending on release vs. debug mode and, sometimes, the version of the compilers. The TSDuck build process automatically selects the appropriate version.

On Linux, the DTAPI is provided as one big pre-linked object file (`.o`). On Windows, the DTAPI is provided as one static library (`.lib`). The only supported architectures are Intel x86 and x64. No binary for other architectures are available.

On Linux, the DTAPI references symbols from `glibc` and distros such as Alpine Linux, which uses the `musl` libc instead of `glibc`, are not supported.

On Windows, because of some silly import/export rules, the code in a DLL can be called from outside that DLL only if it was compiled with the right export attributes (see [section 5.1.3](#) for more details). Because Dektec did not compile the DTAPI with these export attributes, the DTAPI can be called only from the same shared library it is in. This is the reason why the DTAPI and all C++ classes which directly reference it are in the same shared library `libtsdektec`. Tools such as `tsdektec` never reference the DTAPI, they reference TSDuck classes. The DTAPI remains hidden inside the shared library `libtsdektec`. For consistency, the same code structure is used on Linux, even though there is no technical reason which prevents the DTAPI from being called outside the shared library `libtsdektec`.

The C++ classes for the two `dektec` plugins, input and output, are in `libtsdektec`. The shared library `tsplugin_dektec` does nothing but registering these C++ classes as plugins. Similarly, the `tsdektec` command just calls a C++ class named `ts::DektecControl` which is inside `libtsdektec`.

The shared library `tsplugin_dektec` and the executable `tsdektec` are the only binaries which reference `libtsdektec`. These three components form the Dektec subsystem of TSDuck. The shared library `libtsdektec` is the only component in TSDuck which contains proprietary code (the DTAPI library). If, for some reason, it is not possible to include proprietary code in some distribution of TSDuck, it is possible to remove `libtsdektec`, `tsplugin_dektec`, and `tsdektec`, and ship them separately, through more permissive channels.

The TSDuck programming environment, which is used by third-party applications for Digital TV, has interfaces for `libtscore` and `libtsduck`, but not for `libtsdektec`. This shared library is specific to `tsplugin_dektec` and `tsdektec`.

5.4.3. HiDes devices

Using HiDes devices requires kernel-mode device drivers which must be separately installed. The drivers are provided by ITE, the vendor of the main chip in HiDes modulators. Building TSDuck code does not depend on the device drivers. They are only required to use HiDes devices. The TSDuck user guide describes how to install the HiDes drivers.

The ITE drivers for Linux and Windows are very different, there is no common structure, no common API or structure. There is also no stable common userland API above the drivers. TSDuck directly interfaces the ITE

drivers, with different implementations on Linux and Windows. TSDuck recreates a stable intermediate layer which provides a system-independent API to access HiDes devices. This code is located in `src/libtsduck/dtv/hides`. The C++ class `HiDesDevice` provides a uniform access to HiDes devices. The subdirectories `linux`, `windows`, `mac`, `bsd` contain the system-specific variants of this class. Note that the macOS and BSD implementations simply return errors and error messages.

The code for the `hides` output plugin and the `tshides` command is located in the standard locations `src/tsplugins` and `src/tstools` respectively. To access the HiDes devices, they call the classes `HiDesDevice` and `HiDesDeviceInfo` which are part of the TSDuck library.

5.5. AstroMeta-based modulators (formerly VATek)

AstroMeta-based USB devices do not need a dedicated device driver. They are accessed through the portable `libusb` library which is available on all operating systems. At application level, AstroMeta devices are accessed through an open-source user-mode library which is provided by AstroMeta.

Most of the TSDuck code for AstroMeta support is located in `src/libtsduck/dtv/vatek`. The `vatek` plugin is implemented in the class `VatekOutputPlugin` and is completely embedded inside the TSDuck library. The main code of the `tsvatek` utility is in `src/tstools` but this is a simple wrapper around the class `VatekControl` which is inside the TSDuck library as well.

All code in `src/libtsduck/dtv/vatek` is a free contribution from AstroMeta (formerly VATek). In case of problem, before trying to modify the code, it is recommended to contact the author, Richie Chang, initially from VATek, then AstroMeta, [a83912a](#) on GitHub.

Chapter 6. Coding guidelines

6.1. Rationale for coding guidelines

In the computing community, there are many ways to develop software, many programming languages, many paradigms (functional, object oriented, etc.) and many conceptions of the art of programming. Each world has its own set of coding guidelines. But each world does have coding guidelines.

In software development, industrial or open-source, the keys to success are the management of the global cost of the software lifetime, the reliability and the security of the software.

- **Costs:** In the lifetime of a software product, the initial development often costs less than the maintenance. Consequently, an important requirement of software development is reducing the maintenance cost. And sometimes this is at the expense of the initial development cost. Writing good and maintainable code may cost a little bit more initially but the final cost will be beneficial. Writing good and maintainable code requires a set of coding guidelines which are followed by the whole development team.
- **Reliability:** The reliability of the software is planned from the design and is achieved through the quality of the code. The quality of the code also requires a common set of coding guidelines.
- **Security:** The security of the code has been quite challenged by hackers in the last twenty years. Most attacks take advantage of flaws in the code, either plain bugs or complex vulnerabilities in code which is otherwise functionally correct. The experience on those vulnerabilities results in a set of secure coding guidelines.

This chapter defines the coding guidelines for the TSDuck project.

Year after year, the code base of TSDuck has proven to become more robust. The maintenance process improved the code quality through regular refactoring operations, preserving or improving the modular structure of the project. Modifications, improvements or implementation of new features were always developed in a very short time and limited to a local set of source files. Applying the present set of guidelines has been of great help in this continuously improving robustness.



A previous document named *TSDuck Coding Guidelines* contained more generic coding rules. This document was derived from a more general one which was written by the author for professional software development. Because it was too generic, it is no longer relevant in the TSDuck documentation set. For reference, a copy is [still available online](#).

6.2. Classification of coding guidelines

There is no unique term for coding guidelines. In the industrial, academic or open-source software, coding guidelines are sometimes named *coding standard* or *coding rules*. But they all refer to the same type of document and serve the same purpose.

There are several types of coding guidelines. Some guidelines are generic; they describe general software development practices. Some guidelines are more specific; they apply to details in the writing of source code for specific programming languages.

Some rules are strictly necessary to write good code and are not negotiable in any environment.

Some others are simple coding conventions, for instance the naming conventions for identifiers. Even within one single programming language, many different coding conventions exist and are strongly advocated by their respective supporters. These types of conventions can be initially discussed but, once they are selected, they must be adopted by all developers in the organization or project.

Several sets of conventions can be individually satisfactory but using several sets of conventions in the same software project is not. Good and maintainable software must be easily understandable and consistent coding conventions are a key part of the understanding of the code.

This document adopts the following classification:

- **Rule:** A mandatory regulation which must be applied without exception in all contexts. It cannot be negotiated or modified. A rule is preceded by **[Rule]** in the text.
- **Recommendation:** A regulation which must be applied everywhere it is possible. But there are circumstances where it is difficult to apply. Situations where a recommendation must be applied or, on the contrary, may be omitted shall be documented. A recommendation is preceded by **[Recommendation]** in the text.
- **Convention:** A mandatory regulation which must be applied without exception. Alternative conventions could have been possible in the first place but a single one had to be chosen. A convention is preceded by **[Convention]** in the text.

We use the word *guideline* as a generic term for rule, recommendation and convention.

In this document, we group conventions in separate sections. This separation is made on purpose to highlight the fact that rules and recommendations are not negotiable while conventions may differ between projects.

6.3. Generic coding guidelines

This section describes generic coding practices, independently of any specific programming language or framework.

Some of these guidelines may appear very generic in fact and evaluating the correctness of a code regarding these guidelines may be sometimes subjective or biased. However, this document prefers to be helpful rather than normative. And while generic advices can hardly be considered as normative, they may be quite helpful to the developer.

6.3.1. Generic coding rules and recommendations

6.3.1.1. Software architecture

[Rule] *The software architecture shall be driven by the following keywords: Simplicity, Clarity, Modularity, Independence, Reusability, Maintainability.*

- **Simplicity:** The simpler a system is, the more reliable it is. Apply the well-known KISS principle: Keep It Simple and Stupid.
- **Clarity:** The software architecture shall be immediately understandable. And so must be the source code. Anyone shall be able to understand the source code without effort, especially new contributors.
- **Modularity:** Break the design in smaller modules which are easier to maintain and understand.
- **Independence:** The modules should be potentially used independently of others. Modularity is not the synonym of independence. Some modular systems have so many dependencies between modules that the modularity is only an artificial division of a very big spaghetti-like module. Typical pitfalls which break the independence of modules are the usage of global variables, excessive usage of implicit state, module interdependencies, etc. Drawing a dependency graph of the modules can give a first impression; if there are loops in the graph, the independence principle is broken.
- **Reusability:** Reusability is the key to software cost reduction and speed of maintenance. A module shall be written in a generic way. It should be reusable in environments which are different from the original one. Each time you write a module, try to understand what is specific to your situation and what could be used outside of it. Write the module for the general case with a customized behavior from parameters and express your situations by applying parameters to your generic code. The DRY principle ("Don't Repeat Yourself") is an application of reusability.
- **Maintainability:** This is a corollary to all the preceding points. The maintainability is the key to success in software development, industrial or open-source.

6.3.1.2. Source code structure

This section describes the common rules on the structure of source files. Additional rules may be defined for various programming languages.

[Recommendation] A source file should be limited to a maximum of 500 lines. For complex modules which can be hardly split into smaller pieces, 1000 lines is the maximum.

These limits are raw file size, including blank and comment lines, not only source code (but excluding the standard legal headers).

As explained before, simplicity is a key to maintainability. If your file is larger than these limits, it is probably poor quality and you should consider redesigning it.

This is a recommendation rather than a rule because there are pathological cases where it makes sense to list many cases, or rich generic classes with a lot of features and doxygen comments. We must admit that there are some of them in the TSDuck source code.

[Rule] All source files shall start with a legal header which references the copyright and license information for the project.

All lines in the header must start with the appropriate syntax for comment lines in the language of the source file.

Application to TSDuck: the text of the legal header is the following:

```
TSDuck - The MPEG Transport Stream Toolkit
Copyright (c) 2005-2026, author's first name and last name
BSD-2-Clause license, see LICENSE.txt file or https://tsduck.io/license
```

[Rule] A source file shall contain at least 30% comments. More precisely, after removing the initial standard legal header and all blank lines, at least 30% of the remaining lines in the file shall contain non-empty comments.

As explained before, clarity is a key to maintainability. A new maintainer shall not need to analyze the code to understand it, reading it should be sufficient. The code, its structure and its environment shall be described in comments. Comments must be everywhere, not only in the file header. Write comments both for you and for future maintainers. Explain why you invoke a function; do not assume that the maintainer knows what it is. If you use a specific code structure for a good reason, explain it in comments to avoid a future maintainer to "optimize" it later, breaking your intent.

Comments shall explain the code, not paraphrase it. Do not explain *what* the code does (no need to explain that `i++` increments the variable `i`). Explain *why* it does it that way, at that point of the code.

This 30% comment requirement does not take into account the standard legal headers. The 30% proportion of comments shall apply to the description of the specific code and its environment, after removing all standard and non-significant headers. The standard comments which are processed by automated documentation tools such as Doxygen or Javadoc are included in the 30% comment requirement since they describe the actual code.

In the TSDuck source code, C++ source files (`.cpp` files) contain slightly over 30% comments and blank lines. In the C++ header files (`.h` files), this proportion is raised to 60%, mostly thanks to Doxygen comments.

[Rule] Comments should be written at the same time as the code and not after the code is completed.

When you write the code, you know precisely what you are doing at this specific time. This is the right time to explain what you are doing in comments. When the code is completed, some important thoughts you had when writing the code are already gone.

[Rule] Use source code self-documentation tools such as Doxygen or Javadoc.

These tools use specially formatted comments to produce a set of HTML pages or PDF documents with the documentation of all the code. With these tools, it is no longer necessary to write separate detailed design documents; the documentation is automatically extracted from the code.

This is both a cost reduction and a guarantee that the code is consistent with the documentation.

Document everything correctly. A common pitfall is to build the minimum comment structure so that the documentation tool does not complain. The result is a short and obscure description of a function and a list of parameters without description. This is useless. Generate the documentation and read it or have it read by someone else. If the result is not satisfactory, complete the self-documentation comments.

[Rule] Maintain the code self-documentation.

Another pitfall is to modify the code later without applying the corresponding modifications (if necessary) to the self-documentation comments.

6.3.1.3. Revision control system

[Rule] All source files shall be managed in a revision control system.

The TSDuck project is managed using Git. The reference repository is on GitHub within the organization named [tsduck](#) [TSDuck-Source]. Additional Git repositories are available in the same GitHub organization for unitary tests, third-party device drivers, etc.

[Rule] Each developer shall be accurately identified in the revision control logs. The identification includes the real name of the developer and his e-mail address.

This is necessary for accurate history tracking.

Example: Git tries to guess the real name and e-mail address of the author of each commit. However, the results are not always brilliant. A safe way to set the proper user identification for git is through the following shell commands:

```
$ git config --global user.name "Firstname Lastname"
$ git config --global user.email "username@domain.name"
```

The modifications are permanent. They are saved in the file `$HOME/.gitconfig`.

[Rule] Keep the source code repository clean and well organized. Specifically, ensure that dirty or non-original files are never pushed into the repository: temporary files, binary files, compilation products and more generally everything that can be regenerated from source files in the repository.

Most revision control systems have automatic ways of excluding dirty files, typically specifying file naming templates.

With Git, excluding dirty files is achieved through `.gitignore` files. See the man page `gitignore(5)` for more details.

Executables and objects are compilation products, not sources. They should be ignored in most cases. But UNIX executable files have no defined suffix. It is not possible to set up a global rule to ignore them (like ignoring all `*.exe` in Windows). Thus, in each directory where one or more executables are produced, there must be a local `.gitignore` file which contains the exact names of the executables to ignore.

Binary objects or libraries (`*.o`, `*.a`, etc.) are ignored by the global rule, at the root directory of the project. When a directory contains a third-party binary for which no source is available, this binary must not be ignored. There must be a local `.gitignore` file which contains a "not" (!) directive to avoid ignoring this file.

Example: Content of a local `.gitignore` file illustrating these rules:

```
# Executable files to ignore
my_app
other_app

# Third-party binary library to include in the repository (not to be ignored)
```

```
!libthird.a
```

6.3.1.4. Internationalization

[Rule] All elements in a source file shall be written in English. Comments are in English. Identifiers use English words.

Whether you like it or not, English has replaced Latin as the standard international communication language for decades. As of today, Klingon is not yet eligible as a valid international communication language.

[Rule] Text messages and output shall be written in English.

Generally speaking, if there is a need for internationalization, external localization files shall be used for non-English messages. There are standard methods for this in the various environments and programming languages. This rule is particularly important for GUI applications which are used by the general public.

Since TSDuck is a very technical project, it is used by engineers who have at least some understanding of English. Consequently, there is no internationalization of the messages in TSDuck.

6.3.1.5. Modularity and compatibility

[Rule] Modularity is a contract definition. Define the contract and respect the contract.

A module shall be clearly defined by a public interface. This interface defines a contract between the module and its users. This contract shall be clearly described in the module interface, typically in the self-documentation comments. Whenever the module is modified, respect the original contract.

[Recommendation] Always preserve ascending compatibility.

A contract is a contract and you should not be not allowed to break it. If you make substantial modifications in an existing module, you may add features but you should not remove features or change the contract of existing features.

As a general rule, you must ensure that all existing code that could have used your module still compile and work correctly without modification.

For projects with a long history, this recommendation is not always easy to apply. When adding many features or restructuring poorly designed modules, preserving ascending compatibility may create chaos in the code, making it more difficult to maintain in the future.

So, the right balance shall be found. It also depends on the nature of the project. A standard library which is used by thousands of developers and maintained by a large team shall prioritize ascending compatibility. Libraries with very few maintainers and not too many users may choose to prioritize the maintainability of the code, at the expense of ascending compatibility.

TSDuck belongs the latter category. TSDuck went through several code cleanup cycles, taking advantage of new standard features in C++11, C++17, then C++20. In these operations, old features were removed to simplify the code base.

[Recommendation] Use Object Oriented Design (OOD) wherever it is possible.

TSDuck is developed in C++ and takes advantage of all OOD features of the language.

The Java and Python bindings follow the same principles.

[Rule] Do not expose implementation details.

The interface of a module or class shall expose only public declarations which are part of the contract. Implementation details shall be hidden. Not only undocumented, but logically hidden. Even if you do not document them, someone will discover them in some definition or header file and may use them. Application code trying to use implementation details should not compile.

In C++, implementation details shall be implemented as *private* fields and methods in classes.

[Rule] *Expose opaque data types only.*

This is an application of the preceding rule. If you need to expose the existence of a data type, you should hide its internal structure. Some fields may be strictly private, for internal use of your module. Some fields may be read-only for the user. Some fields may be modifiable by the user but with some control from your implementation.

And in the future, you may need to redesign the structure of the data type but you will also need to preserve the contract, preserve the logical access to these fields.

Thus, do not expose public fields. Use *accessors*, also named *getter* and *setter* methods.

If you think that this is a loss in performance, you are wrong. Most compilers have the ability to inline the code of a function. If you define an inline accessor method, the generated code will be only a data access. You get the same performance as a public exposure of the field but with a better contract and better maintainability of the code.

If you think that coding and documenting accessors is boring, well, you are right. But many IDE's can do most of the work automatically for you.

[Rule] *Global variables are evil. Do not use them.*

Global variables break the contract of a module. They also break the independence principle. And they also break the reusability principle.

[Recommendation] *Static variables are evil. Avoid using them.*

Static variables (in C parlance) are slightly different from global variables since they are not shared between modules. But in the OOD approach, a static variable is shared by all instances of its module. Most of the time, this breaks the reusability principle.

If some kind of persistent unique data is necessary, prefer the singleton design pattern over the static variable.

6.3.1.6. Naming conventions

This section lists a few generic naming conventions which apply to all programming languages. For rules which specifically apply to C++, see [section 6.4.2.3](#).

Usually, naming is addressed by interchangeable *conventions*. However, these conventions shall be managed by higher level common sense rules.

[Rule] *Use meaningful names which are immediately understandable by the reader or maintainer.*

Example:

```
int logicalFileIndex;    // OK, self-explanatory
int logical_file_index;  // OK, just another naming convention
int lfi;                 // Questionable. Is LFI a really widely used acronym for Logical File Index?
                        // Or is it just your own convention?
int i;                   // OK if this is just a loop index in a small function.
                        // WRONG if this is a Logical File Index
```

[Rule] *Do not use names which differ only in letter case, differ by one character only or too similar.*

This is confusing for the reader and error-prone for the maintainer.

Counter-example: The following sets of identifiers can be easily confused.

```
int broadcastIndex; // lower case 'c'
int broadCastIndex; // upper case 'C'
```

```
int counter;  
int counters[10];    // trailing 's'  
  
int index1:          // digit 1  
int indexl;          // lower case 'l'  
int indexI;          // upper case 'i'
```

6.3.1.7. Coding principles

[Rule] RTFM. Read The F... Manual. Read the documentation of all functions, classes or libraries you use.

This is such an obvious rule that it should not be worth mentioning. However, experience proves that it is not always followed.

Each time, *all the time, really*, read the corresponding documentation before invoking a function. Even if this is a function you personally developed some time ago. Understand the error cases, the returned error codes, the side effects, the exact usage of the parameters, etc.

Get prepared for reading documentation. Have a shell window ready for a `man` command, a browser with prepared bookmarks pointing to the online help, Doxygen or Javadoc pages. When coding on a secure network which is not connected to the Internet, make sure that a copy of all online documentations is available on a server within the secure perimeter.

[Rule] If it ain't broken, don't fix it!

Sometimes, a programmer reads some code and finds it clunky. Some programmers cannot resist the temptation to "fix" the ugliness of the code. This is a dangerous practice. In some cases, this modification may introduce a bug and turns a clunky but perfectly functional code into a buggy one. In front of an existing and operational code, in the absence of other intent such as refactoring or generalization, the question shall never be "is it elegant?" but "does it work?" And if it works, do not modify it.

As a consequence of this rule, do not rewrite existing code for the sole purpose of applying coding guidelines. When the coding guidelines evolve, it would be both counter-productive and dangerous to review and modify the existing code base to retrofit the new coding guidelines. The coding guidelines apply only to new code or code which is modified for valid maintenance reasons.

[Rule] Do not uselessly over-optimize without profiling.

Do not make hypotheses on the performance of the code. Do not write more complicated code because you think it will be faster. You do not know what the compiler will generate; only the compiler knows. Let the compiler optimize the code. Recent compilers are very good at optimization.

Even if a local construct is slightly more efficient, you do not know the actual impact on the global execution time. It is dangerous to write more complicated code to gain 0.001% of execution time. On the contrary, because the code you write is more complex, it is more error-prone and more complex to maintain.

Write the code using a clean structure. When the final application code is ready, do some profiling in the real context of the application. This will give you the actual bottlenecks in the code. Then, you will know which parts of the code are worth optimizing.

[Rule] Write endian-neutral code. Make sure your code works identically on big-endian and little-endian processors.

Never assume that a piece of code will work forever on the same type of processor. Even embedded code will eventually be ported on other platforms.

Some CPU architectures are little-endian (Intel), some are big-endian (SPARC, IBM s390) and some can work in either mode (Arm, PowerPC, MIPS). To make provision for future ports, make sure that your code is not dependent on a specific endianness.

The endianness applies only to integer, floating-point data, UTF-16 and UTF-32 characters strings. ANSI and UTF-8 characters strings, cryptographic keys and other types of raw data do not have any endianness; they are just suites of bytes.

The endianness only matters for external data interchange or storage. Write or use *serialization* and *deserialization* libraries to switch back and forth between the native and external representations of data. Each programming language has specific ways of handling this gracefully.

Do not make any assumption about performance. For instance, if you handle external data in little-endian representation and your code targets a little-endian CPU, do not avoid serialization routines simply because you think that this would add a useless overhead. Well-designed serialization libraries do not add any overhead when the native and target representations are identical. In this case, the serialization and deserialization routines are inlined empty functions. But using them in the source code is beneficial. When you port the code to another platform with the opposite endianness, the program will work correctly without modification.

It is difficult, and sometimes impossible, to test that a program is truly endian-neutral, especially when the native and external data representations are identical. When possible, test your code on a little-endian and on a big-endian platform from the beginning.

In the case of MPEG transport stream, most data are serialized in big-endian order. TSDuck provides low-level serialization functions for individual values and high-level classes for serialized data buffer management. Make sure to always use them.

Nowadays, most commonly used platforms are little-endian. The only well supported big-endian platform is IBM s390x (this is one of the three best available Linux distros, with Intel and Arm). It is difficult and expensive to obtain IBM s390x hardware. However, TSDuck is periodically tested on emulated IBM s390x using [qemu](#).

[Recommendation] Use design patterns.

A number of design patterns are available in the literature. These patterns are proven constructs. Do not invent hazardous new designs, use proven design patterns. See [\[GAMMA\]](#) for instance.

6.3.1.8. Secure coding

[Rule] Do not assume anything.

This is the common root of many security flaws.

- Do not assume that the caller has checked a piece of data, check it.
- Do not assume that this object is in such state because you think that this is logical in this context, check it.
- Do not assume that a connection is established, check it.
- Do not assume that a buffer has such minimal size, check it.
- Do not assume that this pointer is non-zero, check it.

Etc. etc. etc.

[Rule] Do not trust any other code.

This is a corollary of the preceding rule: do not assume that a function does what its documentation says it should do. Specifically, if there is a simple way to improve the robustness and security of your code by checking what the other code did or returned, then check.

[Rule] Avoid temporary solutions or quick & dirty fixes.

If you are tempted to write a temporary fix to a problem, do not do that. Take your time and implement the good and final solution right away. Otherwise, because of the daily workload, you will delay the final solution, then you will forget about it and finally your temporary (and most certainly insecure) fix will become permanent.

Counter-example: A well-known product had a random number generation routine which was used to initialize cryptographic operations. The body of this routine was merely something like `return 3;`. It is clear that the

developer intended to implement a proper solution later. But it never came.

[Rule] Inputs are evil. Never trust inputs. Check all inputs.

When you write a module, the inputs you receive come from an untrusted world by definition. You must always check all inputs. Addresses shall be validated (check null pointers, alignments and range constraints). Numerical values shall be checked for allowed ranges. The size of memory areas and buffers shall be checked. C-strings shall be checked for null termination. The format of special strings shall be checked (date and time, request, etc.)

All functions shall check the validity of all inputs, always, without exception.

[Rule] Check all intra-application inputs also, at all levels.

Do not assume that the inputs of a function are valid simply because the caller is part of the same application. The caller code may have bugs. The caller code may be modified later and introduce new bugs. Your function can be reused in another context (remember that we encourage reusability) and the new caller may have bugs.

[Rule] Buffer overflow is your enemy. Check buffer sizes all the time.

If there is any theoretical, practical or syntactical possibility that some data could be larger than the memory area where you are going to copy it, check the data size, the index or whatever.

Use the capabilities of your programming language to define safe buffer objects with well-designed primitives instead of using raw memory buffers.

[Rule] Pay attention to the stack usage in a multi-threaded environment.

In a multi-threaded environment, all threads are not equal. The size of the stack of the main thread is often "unlimited", or at least extremely large, in the multi-mega-byte range. On the contrary, the size of the stack of all other threads is constrained. The default size depends on the platform but it is in the multi-kilo-byte range, much smaller than the main thread.

When you write a function, always remember that it will be potentially used in various environments, single-threaded or multi-threaded, with small stacks or large stacks.

Allocating very large objects on the stack may lead to unpredictable results if they overflow the current stack. If you are lucky enough, the implementation allocates non-accessible *guard* pages of memory before and after the stack of all threads. Thus, if you overflow a stack by a few bytes, the program will crash with an access violation error. However, if you allocate a very-very large object on the stack, the start of the object may step over the guard pages and lie into the stack of another thread. Thus, accessing the highest parts of your local object, you will corrupt local objects in other functions executing in another thread. Finding such a bug is a nightmare.

[Rule] Limit of size of local variables and all data which are allocated on the stack.

This is a corollary of the previous rule. Allocating large objects on stack may overflow the stack of a given thread. Large local objects shall be dynamically allocated on the heap and released before going out of scope.

[Rule] Always explicitly specify the size of the stack of all threads you create.

This is another corollary of the same rule. If you do not explicitly specify the size of the stack of a thread, it will get the default size which is defined by the implementation. And this size is typically different from one platform to another. This means that your code is not deterministic. It may run on one platform and crash on another.

Be very conservative when computing the required size of the stack of a thread. Do not try to be too clever. It is useful to carefully compute the accumulated sizes of all local objects in all nested functions. But what should you do with the result? Clearly, do not use it as the stack size. The stack is cluttered with many stack frames or temporary working data which are allocated by the compiler. This extra size depends on the platform, the ABI, the compiler, the compilation options, etc. In practice, you cannot accurately compute the required stack size. So, use various sources of information. Compute your local data sizes, run tests and measure the maximum stack size and, at the end, double it just in case.

When you define an API, there are cases where the user passes the address of a *handler* or *callback*, some code which is developed by the user but which is invoked by your library. If your library creates internal threads and

some user handler or callback is invoked in the context of these threads, how do you compute the stack size of the threads you create? This cannot be transparent to the user. In the API of your library, you must give a way to specify the required stack usage of the user handler. The user shall be able to pass the handler code address *and* the corresponding required stack usage of his handler. Internally, to compute the stack size of your threads, you must add your own usage and the user's usage.

[Rule] Always check error codes. Take appropriate action and error processing.

Most functions return an error code or some indication that the processing failed somehow. Read the documentation of the function and always handle error cases.

This is specifically important for memory allocation routines.

But there are less trivial cases.

Example: Consider file management.

- We probably always check the error code which is returned by a file opening or creation operation because we anticipate the risk that the file does not exist or we have no permission to create it.
- When reading from a file, we generally always check the returned error code because we can reach the end of file at any time.
- However, how many of us check the error code which is returned by a *write* operation? Nonetheless, writing to successfully opened files can fail at any time, because the volume is full, because the device has been removed, etc. In some cases, failing to write critical data such as transaction information or cryptographic keys without noticing can have disastrous effects.

[Rule] In case of error, fail safely.

When writing complex functions or module, you may invoke hundreds of external functions. Any of them can return an error situation. In the meantime, you may have opened, allocated, locked or somehow initialized many resources. If each error processing is individual, the code is bloated and potentially insecure.

Write a common error processing path which checks which resources should be cleaned and which safely cleans them.

Object-oriented languages can greatly help you in this process. Carefully designed destructors should always automatically clean up resources.

[Rule] Don't write spaghetti-like code, write state machines.

Avoid deeply nested `if ... then ... else` structures. Avoid multiple processing of the same cases in various places of the code. When applicable (and it is more often that you may think), redesign the code as a state machine.

[Rule] Physically clean up secure information from memory.

Passwords, encryption keys and less critical information are stored in memory while being processed. Once the local processing is completed, overwrite the corresponding memory with binary zeroes or random data. If some memory was dynamically allocated on the heap or on the stack, clean it up before freeing it from the heap or returning from the current function.

[Rule] Make sure that secure information is never accidentally written to disk in clear form.

Most high-level operating systems such as UNIX or Windows use virtual memory. When the physical memory of the system is about to fill, the system swaps or paginates, which means that some memory pages belonging to running processes are temporarily written to disk. If those memory pages contain sensitive data such as passwords or cryptographic keys, this becomes a security breach.

You must ensure that this never happens.

With virtual memory systems, make sure that all data buffers which contain sensitive data are locked in physical memory. Thus, the corresponding memory pages will never be written to disk.

Example: On Linux systems, use `mlock()` to lock a memory area in physical memory.

[Rule] Sensitive information shall be exchanged or written to disk in encrypted form only.

Persistent sensitive information usually needs to be stored or exchanged. Storage shall be performed in encrypted form only. Exchanging sensitive information shall be performed using an approved cryptographic protection layer, either individually or using a secure tunnel such as TLS 1.3.



All versions of SSL and versions 1.0 and 1.1 of TLS are considered insecure. TLS 1.2 shall be considered as deprecated.

Encryption keys shall be adequately protected. Passwords and other similar data which need to be compared but not retrieved in plain form shall not be stored. Instead a salted hash is stored.

Use approved encryption algorithms only. Also use approved security protocols and approved key management methods only. Never use your own encryption algorithm or security protocol. In case of doubt, ask an expert.

[Rule] Usage of encryption shall be validated by a cryptologist or security expert.

Cryptography is a smoking gun. Encryption and security protocols are hard to secure. Naïve implementations or careless straightforward usage of encryption primitives are known to be vulnerable to various types of attacks.

All the following mechanisms shall be validated by a cryptologist or a security expert before being used:

- Encryption, decryption.
- Signature generation and verification.
- Encryption keys and other random data generation.
- Key management and protection.
- Security protocols and secure data exchange.
- Usage of hash, nonce and challenge.
- Implementation of cryptographic primitives, choice of cryptographic libraries.

[Rule] Take care of reentrancy in reusable components.

Basic reusable components increase the productivity and the maintainability of the code. However, using the same component from concurrent threads may lead to race conditions and crashes.

Clearly document whether a component is thread-safe or not.

You may internally implement the required locking feature from the beginning to make the component thread-safe. However, this may have several drawbacks. There may be useless performance degradation in applications which do not share the component in multiple threads. Additionally, the concepts of threads and their associated locking features may be quite different between systems and it is not always easy to write portable thread-safe code.

One option is to write a simple and non-thread-safe component and then write a thread-safe variant of the component. Do not re-implement the component in the thread-safe variant. Encapsulate the calls to services of the non-thread-safe component in the services of the thread-safe variant.

Object-oriented languages with template or genericity features such as C++ and Java may help to write a unique portable thread-safe implementation. The component focuses on its core features only and accepts synchronization primitives or objects as template or generic parameters.

[Rule] Use the principle of "least privilege".

This is a fundamental rule of secure software. It means that a code shall not run with any privilege it does not absolutely require for its nominal execution. In case of security breach, when an attacker takes control of your code, this attacker shall not be able to take advantage of extra privilege.

Example: If an application manipulates files on behalf of a specific user, the process shall have no more privilege than this user. More generally, the application shall run with the identity and privileges of the user it works for.

[Rule] Do not run any code using the root account or any kind of privileged user account.

This is a corollary of the previous rule. Only the operating system or low-level system services need system-wide privileges. There is usually no good reason to use the *root* or *administrator* account or whatever privilege to run an application. If you really think that you need to be *root*, you are wrong. If you really insist, see next rule.

[Rule] Drop privileges as soon as possible.

This is another corollary of the same rule. There are rare cases where your application needs to have elevated privileges for a few very specific operations. In that case, use the following scenario:

- Start the application with the strictly minimum set of privileges which are required to perform the privileged operations.
- Perform the privileged operations right at the beginning of the application.
- Drop all extra privileges immediately after.

Example: An HTTP server shall listen on TCP port 80 by default. This port number lies in the system range of ports. On UNIX systems, you need to be root to open this port. If your application implements an HTTP server on this port, it shall be started as root. At the very beginning of the application, opens the port 80. Do not do anything else. Do not work on files or system resources which can lead to security problems if the application is misused. When the TCP port 80 is successfully open, immediately switch to some non-privileged user account like *httpd*. The open file descriptor of the socket on port 80 is still valid and can be used from the unprivileged user account.

[Recommendation] Lock your application in an isolated enclave without access to the rest of the system.

This is yet another corollary of the above rule. When your application is successfully attacked and control is taken by an attacker, the "pwned" application shall not be able to access the rest of the operation system.

Define the required environment for your application; define the corresponding boundaries and lock the environment within these boundaries inside a subsystem without access to the rest of the system. The available confinement mechanisms depend on the operating system.

Example:

- *chroot()* on Linux and some other UNIX systems.
- Linux containers (LXC, Docker).
- Virtual machines under control of a hypervisor (VMware, VirtualBox, XEN, HyperV, etc.)

6.3.1.9. Software evolution

[Rule] Avoid copy/paste source code. Use refactoring instead.

Most of the value in software engineering comes from the software reusability and its corollary, development and maintenance cost reduction. Software reusability means generic and reusable code. But writing clean generic code right from the beginning is not natural. The process of producing generic code is typically extending and generalizing existing code into generic software.

This is named *refactoring*, one of the most important concepts in modern software engineering.

Creating new code from existing code using copy / paste is duplicating, not refactoring. Duplicating code means duplicating bugs, duplicating maintenance costs and more generally increasing the entropy of the system.

Remember: Copy/paste propagates bugs but does unfortunately not propagate fixes.

[Recommendation] Avoid creating too many branches in the source code repository.

The presence of multiple persistent branches in the source code repository is a wide-scale variant of the copy / paste syndrome. It multiplies the maintenance costs and ruins the source code consistency. Fixing old bugs becomes a nightmare because of the multiple and divergent branches where the bug is present. Creating a branch may save some time on the very short term but the extra cost on the mid and long term is overwhelming.

Exception: Creating short-lived branches for exploration, integration or tests may be accepted under the firm condition that the branch will be rapidly merged into the main trunk and deleted.

6.3.1.10. Compilation errors and warnings

[Rule] Use the most paranoid "warning mode" of the compiler and fix all warnings without exception.

Compilation warnings are never gratuitous. A compiler warning always draws attention to a potential bug. Even if the code works as expected on the tested platform (i.e. the set of compiler & target execution system) the incriminated section of code may fail on another platform. Ignoring tons of uncorrected compilation warnings is a very bad practice. When porting to a new platform, you may have to fix hundredths of bugs that would have been avoided if the warnings were considered in time.

Even warnings which appear to be really harmless after analysis are dangerous because they create a culture of ignoring warnings. A listing of hundredths of "harmless" warnings hides new warnings which may indicate actual bugs.

The most paranoid warning mode depends on the compiler. Different compilers have different options. Some compilers find warnings that some other compilers don't. As a general rule, fix all warnings from all compilers. Your code must not generate any warning on any platform.

[Rule] Turn compilation warnings into errors.

Many compilers have an option which turns warnings into errors. Use it in all code without exception. Thus, the presence of one single warning prevents the code generation.

[Rule] Understand an error or a warning before fixing it.

This seems obvious but this is too often not applied in practice.

Fixing a compilation error or warning is not simply making it disappear. It is tempting to simply find a modification in the source code which gets the compilation message away. But is this the right fix? Does this remove the potential bug that was behind? Did you even understand the potential bug?

Example: You may get an error about a pointer type mismatch or a warning about a comparison between signed and unsigned integers. Of course, casting one object to the appropriate type removes the message.

And sometimes, this is the appropriate fix. But sometimes, this is not.

You need to deeply understand the situation to decide whether you should use a type cast or modify the code structure. Do you understand why `if (a < b)` produces a warning when `a` is a signed integer and `b` an unsigned one (or the opposite)? Do you understand that the underlying risks include an index mismatch leading to a buffer overflow, one of the major sources of security breaches?

If you don't, do not fix the code yourself, get help first.

6.3.1.11. Makefiles

This section describes the rules to use GNU Make to build native software on UNIX systems (Linux, macOS, BSD) only.

[Recommendation] The command `make` should be usable at any point in the source tree. A makefile shall be present in all directories of the source tree. The name of the makefile shall be `Makefile`.

For most projects, the source tree is complex. Sub-projects, sub-systems, test suites are implemented as sub-directory trees. Some developers work on specific sub-trees while others work the entire project tree (integrators

for instance).

Consequently, running `make` should be possible anywhere in the source tree.

This recommendation applies to "logical entities" only. A library may contain too many files to use one single directory. The source files are spread all over a tree of subdirectories. However, if the compilation of the library is homogeneous, it is acceptable to have a makefile at the root of the library source tree only.

[Recommendation] As a general rule, when run in a specific directory, the command `make` shall recursively operate on all subdirectories.

In the source tree, any directory and its entire tree of subdirectories shall be considered as a consistent subsystem to build. See the preceding rule.

Consequently, running `make` somewhere in the source tree shall build the corresponding subsystem, meaning the current directory and all subdirectories, recursively.

The setup from a common file shall provide the required tools to simplify this process (see next rule).

When a directory is simply an intermediate level where there is nothing to build locally, simply include the common file and nothing more. The first target in the common file shall recurse into the subdirectories.

[Recommendation] All makefiles shall include a common file. This file contains the common set of rules, variables, options and targets which are used by all projects.

The actual content of this common makefile is outside the scope of this section. Use a relative path to include the common file since the source code repository can be checked out at different locations on different systems.

Example:

```
include ../../Makefile.inc
```

In the TSDuck source code, there is a file named `Makefile.inc` at the top level of the source tree. It is included by all makefiles, at all levels.

[Rule] Do not try to explicitly enumerate the header file dependencies of C++ files. Setup an automatic dependency resolution system.

It is difficult to manually maintain such dependencies. Include directives are added and removed from nested header files and there is no reliable way to manually maintain this on the long run. This can lead to subtle bugs when a module is not recompiled simply because its dependency on a recently updated header file was missing.

The idea is to automatically generate a text file `module.dep` which contains a makefile dependency rule for each `module.cpp`. The `module.dep` file lists all header files which are directly or indirectly included by `module.cpp`. Generating such a file is a feature of the C preprocessor (at least with GCC and Clang).

All `.dep` are automatically included in the makefile. The makefile automatically regenerates obsolete `.dep` before including them. This is possible using the GCC or Clang compilers and GNU Make.

In TSDuck, this system is implemented in the top-level `Makefile.inc`.

Other building systems such as `cmake`, `qmake`, MSBuild (Microsoft Visual Studio) implement their own automatic dependency resolution system. Use the one that suits you or use the above method with plain makefiles, but do not try to maintain the dependencies manually.

[Rule] Do not try to explicitly enumerate all source files. Setup an automatic file discovery system.

This rule is similar to the previous one. In a large system, with many source files (there are 2300 source files in TSDuck), it is easy to forget to add new or moved files.

Some source files are not explicitly referenced in the application. In TSDuck, this is the case of all descriptors, tables, filters or plugins. There are hundredths of them and they register themselves in a central repository upon initialization. This means that omitting to build such a source file will not prevent the application from being built. However, subtle bugs will emerge from the absence of a module.

This is why this rule is not a matter of laziness, avoiding to list source files. This is a functional and security requirement.

Tools such as GNU make have features to explore the source files. There are heavily used in the TSDuck makefiles. Simply adding a source file somewhere in the directory tree under `src/libtsduck` automatically includes it in the build process of the TSDuck library. Similarly, adding a source file in `src/tsplugins` or `src/tstools` automatically builds a new plugin or tool.



Automatic discovery of source files is sometimes a disputed topic. Some developers advocate the explicit listing of all source files in `Makefile` or `CMakeLists.txt`. This is usually the result of a lack of experience in very large and dynamic projects.

6.3.1.12. Unit testing

[Rule] Use unitary and non-regression tests automation.

Use a unitary testing framework (Cunit, CppUnit, Junit, etc.) For each module, build the corresponding unit tests and integrate them in the framework. Cover most common legal, illegal and limit cases in unit tests.

With TSDuck, the test suite is split in two parts. The first part addresses low-level unitary tests, typically for C++ classes, using a dedicated testing framework named `TUnit`. The second part includes non-regression tests for TSDuck commands and plugins. See [chapter 2](#) for more details.

Use a continuous integration framework (Jenkins, for instance) to run all unitary tests nightly and report failures.

With TSDuck, continuous integration is managed using GitHub Actions. See [chapter 3](#) for more details.

[Rule] Use a Test-Driven Development (TDD) approach.

Using TDD, each module is written in parallel with its unitary tests. Modern methods such as Agile or eXtreme Programming (XP) advocate that the module and its tests shall be developed in parallel by distinct developers.

The typical scenario is the following:

- Create the module in the main development area. The module is initially empty. Make sure it is compiled and inserted in the main product (library, executable, whatever).
- Create the corresponding test in the unitary tests area. Make sure that the test suite is correctly built and executed, although the test suite is initially empty.
- Develop one feature of the module.
- Develop the corresponding unitary tests.
- Run the tests and debug all potential problems.
- Write another feature and its tests.
- And so on.

This approach has several advantages:

- The developer is forced to define a clear API of the module from the beginning. Otherwise, the unitary tests are difficult to define.
- In case of failure, the unitary tests provide a simple environment to debug.
- The module is immediately inserted in the automated non-regression tests. During further developments of the module, if you later break something that was previously tested, the new bug will be automatically caught.
- The integration time is reduced since each module is individually tested at development time.

[Rule] Add new tests for each fixed bug.

Each time a bug is found beyond unit testing, in integration phase or even production, there is a bug in the code. But there is also a bug in the unit test suite because it failed to catch the bug in the code.

The proper fix scenario is:

- Fix the unit test suite first: Add a unit test which exhibits the bug.
- Run the unit test suite and verify that it fails (the bug is caught).
- Fix the application code.
- Re-run the unit test suite and verify that it passes.

This method is also well suited to fight the "resurrecting bug syndrome". We have all encountered situations where a bug was fixed in the past but reappears later. This is usually due to some human error with the configuration management system. By enriching the unit test suite with test cases for all known bugs, we are able to detect the resurrection of outdated buggy versions.

6.3.1.13. Integration of open source software

[Rule] Always check the license of Free and Open Source Software (FOSS) before using it. In this context, "using" means integrating and linking with FOSS modules as well as copy / paste of FOSS source code.

TSDuck is open-source software. So, it may seem weird to worry about integrating other open source software. However, all open source software are not equal in terms of usage and integration because they use distinct, and sometimes incompatible, licenses.

Today, a vast amount of generic or reusable software is available for free ("*free as a beer*") in the form of Free ("*free as in freedom*") and Open Source Software.

Several topics must be studied when using FOSS:

- The license of *your* software (FOSS or not).
- The license of the FOSS you plan to use.
- Your usage of that FOSS.
- The planned deployment of your software.

Many different types of licenses exist for FOSS. The Wikipedia article "[comparison of free and open source software licenses](#)" lists more than 40 different licenses. Some of them are obscure or written by legal illiterate programmers and are difficult to understand. Actual legal cases which are based on these licenses are disputed or even contradictory. Therefore, do not underestimate the difficulty.

Generally speaking, there are two main categories of FOSS licenses: *permissive* licenses (BSD License for instance, as used by TSDuck) and *invasive* or *copyleft* licenses (GPL for instance). While the former may be safely used in proprietary software with some care, the latter cannot.

The way the target FOSS is used is also important. A license may apply differently when the FOSS is copied / pasted in source form, statically linked or used as a side component. The LGPL variant of the GPL is typically designed for dynamically linked libraries. But some libraries are LGPL while others are GPL. And some cases are borderline. Dynamic reference (`dlopen()`) of GPL'ed shared libraries on Linux is a disputed subject for instance.

Because the TSDuck source code is released under the BSD-2-Clause license, it is usually permitted to copy and paste a few lines of code from another project which is released under another permissive license. This license must be compatible with the BSD-2-Clause license. Don't forget to check and take any requested action, such as mentioning the original author, software, copyright or license. However, it is not permitted to copy source code which is released under a restrictive license such as GPL or LGPL. These licenses are not compatible with permissive licences such as the BSD licences.

Finally, the type of deployment of your software matters. For proprietary software which is distributed outside the

company where it is developed (commercially or even for free), the license of the FOSS applies. But for internal proprietary tools which are never distributed, you may usually do what you want. This rule does not apply to community projects such as TSDuck but it is worth remembering.

As a rule of thumb, FOSS with a permissive license such as TSDuck may freely use other FOSS with another permissive license. However, using FOSS with a restrictive license must be done with care:

- LGPL libraries can be accessed through dynamic linking only. Static linking or dynamic loading (`dlopen()`) are prohibited.
- GPL applications can be started through `exec()` system calls or equivalent.
- GPL libraries cannot be used.

As an example, `libreadline` is one of the few libraries with a GPL license, not LGPL. As a consequence, a BSD-licensed software such as TSDuck cannot use it. Fortunately, there is an equivalent library named `libedit` with almost the same features as `libreadline` and a BSD license. Therefore, TSDuck uses `libedit` instead of `libreadline`.



In practice, `libedit` has been developed only because of the unfriendly license of `libreadline`. Similarly, many new software now avoid the GPL license.

In all cases, this subject is too complicated for us, the developers. In case of doubt, get a legal advice first.

6.3.2. Generic coding conventions

This section is present to fulfil the required separation of immutable *rules* and *recommendations* from potentially replaceable *conventions*, as explained in [section 6.2](#).

6.3.2.1. Control characters

[Convention] Use exclusively line-feeds (LF, '\n', 0x0A) as line delimiters.

This is the usual UNIX vs. Windows or LF vs. CR-LF line format. All compilers on Windows understand files with LF as line delimiters. On the contrary, some compilers or interpreters in the UNIX world have problems with CR-LF line delimiters. The `bash` shell, for instance, considers the CR as part of the line.

Most editors and IDE's have an option to force the usage of LF as line delimiters, even on Windows systems.

Automatic clean-up scripts should be used on a regular basis on the code base to convert CR-LF sequences into LF.

Exception: When working with Git, it is possible to automatically switch back and forth between LF and CR-LF during check-out and commit on all or selected types of text files. This feature is managed through the `autocrlf` configuration option and `.gitattributes` files in the repository tree. If you use that feature, make sure that text files are committed with LF lines. If you want some Windows-specific text files (such as `.sln` or `.vcxproj` files) to be committed with CR-LF, make sure to properly reference them in a global `.gitattributes` file, as used in TSDuck.

[Convention] Do not use the tabulation character, use spaces.

In source code, the indentation is essential. It is tempting to use tab characters to indent. However, the size of a tab character may vary between tools. Most system tools use 8 characters per tab. But many IDE's use 4 characters (the usual single indentation width). When editing a source file, depending on the typing and automatic indentation features, the result is often a mixture of tabs and spaces in the indentations. When such a file appears correctly in an IDE with 4 spaces per tab, for instance, the same file appears totally messed up in an editor with 8 spaces per tab.

Most editors and IDE's have an option to force the usage of multiple spaces instead of a tab.

Automatic clean-up scripts should be used on a regular basis on the code base to convert tab characters into

spaces (but a uniform tab width must be applied and may not match the original environment).

Exception: The presence of an actual tabulation character is required in some specific file formats such as makefiles. But most editors and IDE's can handle this exception.

6.3.2.2. Character encoding

There is always some misunderstanding between the characters in a text and their binary representation in a file. In the most general case, a character is universally represented by a 32-bit UNICODE value. But files are rarely represented in UTF-32 format.



In practice, UTF-32 is untransformed UNICODE. Some environments, such as Windows, erroneously name "UNICODE" a 16-bit representation which is actually UTF-16 with surrogate characters. UTF-16 means Unicode Transformation Format - 16 bits.

When a file is transferred between systems or accessed from a shared disk by multiple heterogeneous systems, the physical content of the file remains the same but the various operating systems or production tools often interpret this physical content in different ways, leading to unpredictable results.

The following rules attempt to mitigate this problem.

[Convention] Restrict the character set of source files to the 7-bit ASCII set.

There are few reasons to use characters outside the 7-bit ASCII set in software source files. The syntax of all modern programming languages uses ASCII only. All identifiers and comments shall be written in English which is ASCII only. Any people name in comment (author, developer, etc.) shall be used without accent or local "decoration".



Avoid "national pride" of non-English characters. Contributors shall accept to drop them from their name, if any. As an example, the main author of TSDuck is named "Lelégard", with an accent on the second "e", but the name is simply stored as "Lelegard" (no accent) in all source files. All contributors are kindly requested to accept this rule.

[Convention] Use UTF-8 encoding when internationalization is required.

It is sometimes required to use international texts. Always use specific internationalization files for non-English languages. The actual format of these files depends on the toolset. Unless specified otherwise by a toolset, always use the UTF-8 encoding for these files. When present, the option "UTF-8 without BOM" should be used.



BOM: Byte Order Mark, a feature of UTF-16 and UTF-32 which is normally useless in UTF-8. A specific leading binary sequence has been defined as "UTF-8 BOM". It can be used to assert that the file format is UTF-8. However, several tools are unable to process this sequence correctly. Therefore, it is recommended to avoid it.

When the syntax of a file format allows the explicit specification of the character encoding, use it to specify UTF-8.

Example: All XML files shall be saved in UTF-8 format and start with:

```
<?xml version="1.0" encoding="UTF-8"?>
```

Exception: Some tools may impose or generate files using their specific encoding. Always use the character encoding which is the safest one for such tools.

6.4. C++ coding guidelines

This section defines the coding guidelines which are specific to the C++ language.

6.4.1. C++ coding rules and recommendations

6.4.1.1. Language selection

[Recommendation] Unless it is impossible for a very good reason, do not use C, use C++.

C is an old and unsafe language. C++ is much safer than C. It is possible to improve the safety of C using all advanced features of ANSI C99 and the most paranoid warning mode of the C compiler, but C will never be as safe as C++.

Good reasons to use C instead of C++ include:

- Maintenance of a legacy application written in C.
- Embedded system development on a platform without C++ compiler.

All other reasons are usually bad reasons.

If you need at least two essential reasons to use C++ instead of C, one is named *constructors* and the other one is named *destructors*. These features are the essential bases of safe coding and they are not available in C.

For the record, the very first version of the ancestor of TSDuck, back in 2005, so-called "TSDuck V1", was developed in C. When the code base grew, this quickly appeared to be a fatal mistake. All the code was scrapped and rewritten in C++, starting "TSDuck V2".

[Rule] Performance is never a good reason for not using C++.

There is a common misconception that C++ code is slower than C code. This is wrong. The confusion comes from the fact that C++ has much more features than C and seems more complicated. But more features do not mean more generated code.

Most C++ features are source-level structure and safety. If you are compiler aware, you realize that the generated code overhead is insignificant. It is quite possible to write good object-oriented C++ code with performance in mind.

As an example, a good object design involves lots of very small methods. But if you declare them as inline, you get the performance of C with the features of C++.

6.4.1.2. Modularity

[Rule] For each file `module.cpp` which is not a main program, there must be exactly one corresponding `module.h` which contains the declaration of the public interface of the module. No internal or private definitions shall be present in the header. No part of the public interface of the module shall be declared in another header file.

This is the natural C++ implementation of a module. A module must have exactly one public interface and one implementation. There is one file for each.

Exception: In some cases, there is no need for an implementation and the `.cpp` file is not present, everything is present in the header file. This can be the case where the module contains declarations only, or template methods or classes, or when all functions and methods are short enough to be declared inline.

[Rule] Each module shall contain only one C++ class. If the class has inner classes, the inner classes are declared and defined in the same module at the top-level class.

This simplifies maintenance.

It is not possible to declare an inner class in a different header file from its enclosing class. To keep the one-to-one association between `.h` and `.cpp` files, the definitions of the inner class methods are implemented into the same `.cpp` file as the enclosing class.



Inner classes (also known as nested classes) shall not be confused with subclasses.

[Rule] If a data type is used only inside a module and is consequently private to the module, it shall be declared within the `module.cpp` file and not in any header file.

If the type is private, it should not be useable by any other module. If the type were declared in a header file, even a specific one, different from the `module.h` public interface, it could be included by another module and used outside its normal scope.

In the case of C++ "private" declarations, the syntax of the language requires to declare them in the class declaration, meaning in the header file of the class. However, being private, these declarations cannot be used outside their outer class definition, which preserve the data type hiding rule.



In the early times of the C language, in the "Kernighan & Ritchie" era, it was common to declare all data types in one header file, outside of the source files. This practice is clearly outdated and should be banned. This separation of data types and code is purely based on the syntax, not on the semantics and is a violation of the principle of modularity.

[Rule] A module implementation file `module.cpp` shall start with an `#include` directive of its own interface file `module.h`.

Thus, the successful compilation of `module.cpp` validates two important points:

- The header file is self-sufficient.
- The header file is consistent with its implementation.

[Rule] A header file must allow multiple inclusions without generating errors.

You cannot manage how your header will be used. The compilation of a user application which includes your header file shall never generate an error simply because of your header.

Use case: A user code explicitly includes your header `module.h` as well as some other header `other.h`. If this `other.h` happens to also include your `module.h`, there must be no error.

In C++, the `#pragma once` is defined for this purpose. The old C tradition of conditional compilation using `#ifndef MODULE_H` is outdated and should be avoided.

Example: C++ header file `module.h`:

```
#pragma once
// ... file content here ...
```

[Rule] The structure of a header file shall be self-sufficient. The users of the header file must not have to specify other header files in order to use it. There must be no ordering requirement of the `#include` directives for the user.

This means that a header file must include all other header files which are needed by the declarations it contains.

[Rule] The header file for a C module shall be safely compiled when included from a C++ module. A C module shall be safely linkable with C++ modules.

In a C header file, do not use C++ reserved names.

In a C header file, all language constructs which are specific to C or otherwise incompatible with C++ shall be conditionally compiled using the predefined macro `__cplusplus` (with two leading underscores). This macro is defined by the compiler when the compilation occurs in a C++ environment.

In a C header file, all C declarations which generate a reference to a linker symbol (function or data) shall be enclosed in `extern "C" {...}` when compiled in C++.

Example: The following structures are inserted at the beginning and end of the declarations in the C header file:

```
#if defined (__cplusplus)
extern "C" {
#endif

/* all C declarations here ... */

#if defined (__cplusplus)
} /* extern "C" */
#endif
```

[Rule] Use the anonymous namespace to define elements (data, classes, functions, etc.) which are internal to a module.

This avoids name clashes and namespace pollution.

This is the C++ equivalent of **static** declarations in C. Note that **static** has a different meaning in C++ and should not be used for that purpose.

Example:

```
namespace {
    void SomeInternalMethod()
    {
        ...
    }
}
```

6.4.1.3. Global data

[Rule] The order of initialization of C++ modules is undefined. Never use non-constexpr global variables, at compilation unit level, at namespace level, or as static class members. Use the singleton design pattern for variable objects or local static data for constant objects.

All global data are constructed at the initialization of the application, before **main()** is called. Data with the **constexpr** attributes are guaranteed to be built at compilation time, their content is always valid. Most other global data must be dynamically built, their constructor must be executed.

In each compilation unit (i.e. object file), the compiler generates one initialization routine which calls all constructors of all global data which are defined in this unit. When the application or shared library is built, the linker is somehow responsible for creating one big initialization routine which will call the initialization routines of all compilation units, one by one.

The core problem is that there is no way to explicitly control the order of execution of these initialization routines. This is a limitation of the C++ language. Other languages such as Ada have specific pragmas to control this order, not C++.

Executing a constructor may involve complex processing, calling external functions or referencing external data from other modules. When a constructor is called during the initialization phase and references an object from another module, and this other object was not yet constructed because the initialization routine of its compilation unit will be called later, weird things happen (usually a crash).

The most disturbing aspect is the unreliability of the process. The application may work when built with compiler A and fail when built with compiler B. Or the application may have worked from the beginning and then suddenly crash after adding a module or an object. In all cases, the core problem is indeterminism of initialization order. In other words, when it works, it's only by chance.

This is why you should *never* use non-`constexpr` global variables, at compilation unit level, at namespace level, or as static class members.

The solution to somewhat control the order of initialization of global objects is to force the construction of an object when it is used for the first time. This includes cases where the compilation unit into which this object is defined is not yet initialized. Yes, it is possible to force the construction of an object before its compilation unit.

There are several techniques for this. Starting with C++11, the semantic of *local static data* (static objects which are defined inside a function) has been carefully specified.

Block variables with static or thread storage duration are initialized the first time control passes through their declaration (unless their initialization is zero- or constant-initialization, which can be performed before the block is first entered). On all further calls, the declaration is skipped. If multiple threads attempt to initialize the same static local variable concurrently, the initialization occurs exactly once (similar behavior can be obtained for arbitrary functions with `std::call_once`). Usual implementations of this feature use variants of the double-checked locking pattern, which reduces runtime overhead for already-initialized local statics to a single non-atomic boolean comparison.

— cppreference.com

As an example, consider the following code:

```
const std::vector<int>& ConfigTable()
{
    static const std::vector<int> table {1, 2, 3, 4, 5};
    return table;
}
```

The static data `table` is constructed exactly once, when `ConfigTable()` is called for the first time. All subsequent calls return the same reference to the same object.

Moreover, if the first and second call to `ConfigTable()` are simultaneously performed from two distinct threads, the rules of the language impose that some hidden form of synchronization is applied. Only one construction is safely performed. The second call waits for the completion of the construction to proceed.

Therefore, the coding rule can be illustrated by the following example and counter-example:

Good code:

Bad code:

```
const std::vector<int>& ConfigTable()
{
    static const std::vector<int> table {1, 2, 3};
    return table;
}

void something()
{
    for (auto i : ConfigTable()) {
        ...
    }
}
```

```
const std::vector<int> ConfigTable {1, 2, 3};

void something()
{
    for (auto i : ConfigTable) {
        ...
    }
}
```

The *singleton* design pattern is a variant of this technique. A class has no public constructor so that no application can build its own instance of the class. Instead, the class contains a static method which returns a reference to the sole instance of the class.

```
class Singleton
{
private:
    // Constructor cannot be used outside the class.
    Singleton();
public:
    // Return a reference to the single instance of the class.
    static Singleton& Instance();
};

Singleton& Singleton::Instance()
{
    static Singleton the_one;
    return the_one;
}
```

In the TSDuck library, this pattern is implemented using the two macros `TS_SINGLETON()` (in a header file) and `TS_DEFINE_SINGLETON()` (in the corresponding compilation unit).

[Rule] Never use static variables in inline functions.

If you do that, there will be one static variable per module where the function is used, breaking the semantic of the static attribute.

6.4.1.4. Naming and syntax formatting

Most of these topics are covered by coding conventions in [section 6.4.2](#). However, a few topics are clearly *rules* rather than *conventions* because they impact the reliability and good understanding of the code.

[Rule] Always use a namespace when declaring entities. Never define anything in the default namespace.

This avoids the name clashes during the link. In TSDuck code, everything is defined in the namespace named `ts` or in one of its inner namespaces.

[Rule] The directive `using namespace` is strictly forbidden. There is no exception.

This avoids ambiguities. This also avoids some portability issues, especially with Visual C++ on Windows platforms.

All C++ standard entities must be referenced with the `std::` prefix.

Remember that there is no need to explicitly specify the namespace to reference an entity which is declared in the current namespace or an outer one. Therefore, inside the source code of TSDuck, it is not necessary to repeat the prefix `ts::` everywhere. Only third-party applications need to use it.

[Rule] Explicitly use the `::` prefix for predefined entities which are defined in the default namespace.

This indicates an explicit reference to the default namespace and avoids ambiguities with entities which are declared in the current namespace or an outer one.

Example:

```
open("file");    // WRONG: may conflict with an open() method of this class
::open("file");  // OK, there is no ambiguity
```

[Rule] Do not place code after a comment on the same line.

This is confusing. The trailing code can easily remain unnoticed.

Example:

```
// comment line is OK
a = 1; // comment after code is OK
b = x /* WRONG: code after comment is confusing */ + 3;
```

See also next rule.

[Rule] The comments always start with a `//` and extend up to the end of line. The usage of the C comment syntax `/*...*/` is discouraged.

A common problem with the C syntax `/*...*/` is the potential absence of closing `*/`. In some cases, the subsequent lines of codes are silently "swallowed" by the comment, up to the end of the next comment. Such bugs are very difficult to track.

By using the `//` syntax, all comments automatically terminate at the end of line.

Counter-example:

```
/* innocent comment */
a = 1;
/* WRONG: unterminated comment
b = 2;
/* the preceding comment actually ends here -->*/
c = 3;
```

In this example, the instruction `b = 2;` is excluded from the compiled code. Good luck to find the bug!

[Rule] Using the `{ }` braces in conditional and loop statements is mandatory, even if there is only one or no instruction in the block.

Omitting the braces is dangerous for the maintainability of the code. If the indentation is incorrect or if the unique statement in the block is complicated or spans multiple lines, omitting the braces makes the code hard to understand.

Good code:

Bad code:

Even worse:

```
if (x > a) {
    x = 0;
}

for (i = 0; i < max; i++) {
    print(a[i]);
}

while (readFile()) {
}
```

```
if (x > a)
    x = 0;

for (i = 0; i < max; i++)
    print(a[i]);

while (readFile());
```

```
if (x > a) x = 0;
```

Notes and anecdotes:

- Yes, Python is bad. Know your history. Using indentation for structuring code was abandoned in the 70's after the Fortran era.
- Failing to apply this rule was the root cause of the famous security vulnerability nick-named "gotofail" on macOS and iOS.
- The Linux kernel coding rules specifically prohibit the usage of braces when there is only one instruction. Do these guys ever heard of "gotofail"?

6.4.1.5. Coding style

[Rule] Avoid numerical constants, use `static constexpr` declarations for these values.

This is the C++ way. Preprocessing macros shall not be used for constants in C++. Use preprocessing macros only for conditional compilation or for numerical values which are evaluated in the context of conditional compilation.

Exception: Usual constants such as 0 and 1 which are used in obvious contexts (reset, increment) shall be used without names since they are self-explanatory.

[Rule] The value of all preprocessing macros shall be enclosed into parentheses (in the absence of other syntactic constraints).

This avoids compilation ambiguities.

Example:

```
#define TS_GOOD (TS_FOO + 3)
#define TS_BAD TS_FOO + 3    // WRONG: what about "a = 5 * TS_BAD" ?
```

[Rule] In all preprocessing macros, the parameters shall be enclosed into parentheses when referenced in the definition.

This avoids compilation ambiguities.

Example:

```
#define TS_GOOD(x) (2 * (x))
#define TS_BAD(x) (2 * x)    // WRONG: what about "TS_BAD(a + b)" ?
```

[Rule] In all preprocessing macros, each parameter shall be referenced exactly once in the definition.

Macro preprocessing is only text substitution. When the actual parameter of a macro has side effects, it is evaluated as many times as referenced in the macro definition while the caller expects exactly one evaluation.

Example:

```
#define TS_GOOD(x) (f((x)))      // OK:   TS_GOOD(a++) => a incremented
#define TS_BAD1(x) (g((x), (x))) // WRONG: TS_BAD1(a++) => a incremented twice
#define TS_BAD2(x) (h())        // WRONG: TS_BAD2(a++) => a not modified
```

Exception: There are macros which are reserved to specific usages where the actual parameters are not general-purpose expressions but must be lvalues or special syntactic structures.

Alternative: When it is required to reference a macro parameter several times, you cannot use a macro. You should implement a real function and declare it `inline`. See next rule.

[Rule] Preprocessing macros should be avoided for code structures. Use inline function definitions.

This is the C++ way. This avoids all macro substitution pitfalls which are described in the preceding rules.

There is no difference in terms of code size or performance. The generated code for inline functions is expanded inline, just like a preprocessing macro. If the inline function is not used, no code is generated.

[Rule] The preprocessing directives `#ifdef` and `#ifndef` should not be used. Use `#if` followed by an expression instead.

This is more consistent with directives containing multiple conditions or `#elif` directives. For complex conditions, the code is more compact and more readable.

Good code:

```
#if defined(A)
    #define X 1
#elif defined(B) && defined(C)
    #define X 2
#endif
```

Bad code:

```
#ifdef A
    #define X 1
#else // not A
    #ifdef B
        #ifdef C
            #define X 2
        #endif // C
    #endif // B
#endif // A
```

[Rule] Use the preprocessing directive `#error` wherever there is no valid default alternative.

The `#error` will draw the attention of the developer when the code is compiled in a new environment which was not previously addressed. The developer who does the porting can precisely locate where there is some specific work to do.

Example 1: A simple list of mutually exclusive alternatives, ensuring that one branch is taken.

```
#if defined(TS_WINDOWS)
    using SystemError = ...;
#elif defined(TS_UNIX)
    using SystemError = ...;
#else
    #error "unknow O/S, please update SystemError in this header file"
#endif
```

[Rule] When declaring multiple variables with the same base type, create multiple individual declarations, one per line.

The following two code excerpts are perfectly valid according to the C and C++ standards and have identical effects. However, in the bad code example, it is not clear that the line declares objects which are not `int` (the last two variables are a pointer and an array). It is also not clear that `j` is pre-initialized to zero while `i` is not.

Good code:

```
int i;
int j = 0;
int* p;
int a[10];
```

Bad code:

```
int i, j = 0, *p, a[10];
```

[Rule] *The order of evaluation of the operators in C and C++ is complex. Do not hesitate to add extra parentheses in order to make the code more readable.*

Do you know what `++p->a` means? Does it increment the pointer first and then dereference it? Or does it dereference the pointer first and then increment the pointed field? Actually, the second alternative is right. But how many maintainers will know that?

By the way, did you even know that the self-increment operator `++` has distinct precedencies when it is used as a prefix or a suffix?

Example: Adding theoretically useless but helpful parentheses:

```
++p->a;    // WRONG: nobody really knows what this mean
++(p->a);  // GOOD: same as "++p->a" but more readable
(++p)->a;  // not the same as "++p->a" so parentheses are required anyway
```

[Rule] *A `switch` statement must not contain any implicit "fall through" path, i.e. all cases must end with a `break`.*

The `switch/case` implicit fall-through path is a rare and confusing construct. Most of the time, no `break` at the end of a `case` is a forgotten one, i.e. a bug. By allowing implicit fall-through paths, we cannot clearly identify if a missing `break` is a bug or a feature. So, to be safe, we forbid it.

Exception: It is allowed to have no `break` when there is no instruction at all after the `case`. This situation is not a real fall-through, this is a `case` with multiple equivalent values.

Example:

```
switch (x) {
  case 1:
    print("something");
    // WRONG: implicit fall-through path is confusing, bug or feature ?
  case 2:
    print("else");
    break;
  case 3: // OK, not a real fall-through but a multiple case entry
  case 4:
    print("ok");
    break;
  default:
    print("error");
    break;
}
```

In TSDuck, when allowed as a compiler option, the implicit fallthrough paths are detected and rejected using the appropriate warning configuration.

[Rule] *In a `switch` statement, if a "fall through" path is really necessary, it must be explicitly declared using a `[[fallthrough]]` attribute.*

These constructs shall remain rare and reserved to cases where any other structure would be even more confusing.

Example:

```
switch (x) {
  case 1:
    print("something ");
    [[fallthrough]];
  case 2:
    print("else");
    break;
  default:
    print("error");
    break;
}
```

[Rule] A `switch` statement must end with a `default` entry.

This ensures that unexpected values are always processed, typically with an error processing.

This is especially useful when the `switch` is applied on all values of an `enum` type. You may think that the `default` alternative is useless since all values are handled. However, what will happen when you update the `enum` type and add a value? If several `switch` statements are used in various modules, you will probably forget to update the list of cases in one of them. Without `default` alternative, the new value will be silently ignored and the bug will be either hard to find or (worse) remain undetected. In the presence of a `default` alternative which triggers error code, the run-time error will immediately detect the bug.

In TSDuck, when allowed as a compiler option, the absence of `default` alternative is detected and rejected using the appropriate warning configuration.

[Rule] If a local variable is required in a case or default entry in a switch, declare a local block using `{ }` for the entry.

Do not declare a local variable in a `switch case` without enclosing `{ }` block. The scope of the local variable extends to the subsequent entries in the `switch`, which is usually not what you indented.

Note that the C and C++ languages have distinct interpretations here.

Good code:

```
switch (x) {
  case 1: { // <== note the "{"
    int i; // local to this entry
    ...
    break;
  } // <== note the "}"
  case 2:
    ...
    break;
  default:
    ...
    break;
}
```

Bad code:

```
switch (x) {
  case 1:
    int i; // C: compilation error
    ...
    break;
  case 2:
    // C++: i is still visible here
    ...
    break;
  default:
    ...
    break;
}
```

[Rule] In conditional expressions, explicitly compare against zero for non-boolean values such as integers or pointers.

The implicit interpretation of an integer or pointer value as a boolean is confusing. Sometimes, the "common interpretation" is even the opposite of the reality.

Example: The common sense is that a function which returns a boolean value would return true on success and false on error because the word "true" reveals a positive feeling while "false" reveals a negative one. Many UNIX system calls do not return a boolean value. They return an integer which is zero on success and a non-zero error code on error.

See how this can be confusing.

Confusing code:

```
if (close(fd)) {
    // You think it's a success ?
    // In fact, it's an error
}
```

Cleaner code:

```
if (close(fd) != 0) {
    // error processing
}
```

[Rule] Never use boolean operators (!, &&, ||) on non-boolean values such as integers or pointers.

Same reasons as the preceding rule.

[Rule] Do not use assignments in conditional expressions, even for boolean variables.

This is confusing. The lazy reader will misinterpret `if (a = b)` as `if (a == b)`. So, assign first, then use the resulting variable in the conditional expression:

Confusing code:

```
if (big = isLargeFile(f)) {
}
```

Cleaner code:

```
big = isLargeFile(f);
if (big) {
}
```

In TSDuck, when allowed as a compiler option, an assignment in a boolean expression is detected and rejected using the appropriate warning configuration.

Anecdote: A long time ago, someone attempted to introduce a backdoor in the Linux kernel, pushing code containing something like `if (proc->uid = 0) {...}`. The careless reader could think that the kernel code tested if the caller was root. In practice, the code forced the caller to become root, which was a severe security breach. If you need only one reason for this coding rule, use this anecdote.

[Rule] Reduce the scope of variables to the smallest possible block.

This is more readable and easier to maintain.

Example: Loop indexes should be declared inside the `for` statement when possible:

Good code:

```
for (int i = 0; i < max; i++) {
    const int j = 3 * i;
    ...
}
```

Bad code:

```
int i, j;
...
for (i = 0; i < max; i++) {
    j = 3 * i;
    ...
}
```

[Rule] Delay the declaration of variables until they are used for the first time.

This avoids lengthy initial declarations which are not immediately useful. This also gives you a chance to declare the variable as `const` for instance.

Good code:

```
int a;
...
a = ...;
...
const int max = 2 * a + 1;
for (int i = 0; i < max; i++) {
    ...
}
```

Bad code:

```
int a, i, max;
...
a = ...;
...
max = 2 * a + 1;
for (i = 0; i < max; i++) {
    ...
}
```

[Rule] To set the initial value of an object, always prefer the initialization syntax over an assignment.

This is faster to execute. In the case of the initialization by assignment, the default constructor is first invoked, then the assignment operator is invoked.

Example: Objects `b`, `c` and `d` have an identical initial value. But the initialization of `b` is faster.

```
ts::Foo a;
ts::Foo b(a);    // Invoke the copy constructor
ts::Foo c = a;   // Invoke the default constructor, then the assignment operator
ts::Foo d;       // Invoke the default constructor
d = a;           // Invoke the assignment operator
```

In the case of objects `c` and `d`, the default constructor does something which is immediately undone by the assignment operator.

Sometimes, the compiler may optimize this but do not count on it since a deep code analysis of the constructor and assignment is required to allow this optimization. Without this deep source code analysis, the compiler does not know if the semantics and side effects of the default constructor plus assignment are equivalent to the copy constructor.

[Rule] The use of `goto` is prohibited.

The `goto` statement is the most discussed and infamous construct in computer software history.

Possible exception: The only acceptable exception is a common error path within one function, at the end of the function, after the return of the normal (non-error) path.

However, this should be reserved to special cases, when working on low-level system calls, in the presence of multiple potential error cases, when using another more complex structure would be even more confusing. In modern C++ code, using structured objects with a proper destructor should prevent this.

Example: Valid and (possibly) acceptable usage:

```
class Foo
{
public:
    Foo(); // default constructor
private:
    FILE* f1 = nullptr;
    FILE* f2 = nullptr;
};
```

```

Foo::Foo()
{
    // Open each resource
    f1 = fopen("file1", "r");
    if (f1 == nullptr) {
        goto error;
    }
    f2 = fopen("file2", "r");
    if (f2 == nullptr) {
        goto error;
    }

    // Do other initial processing
    // ...
    return;

error:
    if (f1 != nullptr) {
        fclose(f1);
        f1 = nullptr;
    }
    if (f2 != nullptr) {
        fclose(f2);
        f2 = nullptr;
    }
}

```

[Rule] A single function shall not exceed 50 lines of actual code or approximately 100 raw lines including comments.

Very long functions are hard to understand and maintain. They should be redesigned to extract local treatments in other well documented local functions, even if these local functions are called only once.

Exception: A very large `switch` construct with a lot of short entries may be acceptable. In fact, breaking it apart into several functions may be less maintainable.

[Rule] Do not pass parameters of non-elementary type by value.

Passing a parameter by value is legal but it requires copying the parameter value, typically on stack. While this is harmless for elementary types (integers, floats, enums and pointers), it can be costly for structured types.

In C++, this can even be devastating with polymorphic types (classes with at least one virtual function) when the type of the formal parameter is a superclass of the actual parameter. In this case, we get the *slicing problem*: the pointer to the *vtable* and the superclass parts are copied while the derived parts are not. The result is an inconsistent object with virtual methods possibly using fields that do not exist. In C++, the common way is to pass such parameters *by reference* (a C++ specific mechanism, not to be confused with passing *by address* or *by pointer* in C).

Example:

```

class Foo {...};

void f(int x);           // OK: elementary type by value ("in" parameter)
void f(int* x);          // OK: elementary type by address ("out" parameter)
void f(Foo x);           // WRONG: structured type by value, don't do that
void f(const Foo* x);     // OK: structured type by address ("in" parameter)
void f(Foo* x);           // OK: structured type by address ("out" parameter)
void f(const Foo& x);     // OK: structured type by reference (C++)

```

[Recommendation] In **for** loops, prefer the modern C++11 syntax instead of explicit iterators.

The code is more readable, more compact, and less error-prone.

Example:

```
std::list<Foo> flist;

// Using iterators
for (std::list<Foo>::iterator iter = flist.begin(); iter != flist.end(); ++iter) {
    process(*iter);
}

// Modern C++11 syntax
for (auto& f : flist) {
    process(f);
}
```

Exception: There are cases where the iterator is required for intermediate operations of specific processing at the end of each iteration. This is the case when it is necessary to insert or remove elements in the list in an iteration.

[Rule] In modern C++11 **for** loops, use a reference in the iteration parameter.

This saves the useless copy of a temporary object.

Example:

```
std::list<Foo> flist;

// INCORRECT code
for (auto f : flist) {
    // In each iteration, f is a new temporary object which is constructed
    // from the value of the element in the list
    process(f);
}

// correct code
for (auto& f : flist) {
    // f is a reference to the element in the list
    process(f);
}
```

6.4.1.6. Strict typing

To detect as many potential errors as possible directly at compilation time, use strict types and attributes on all declarations and definitions.

[Rule] Use the **const** attribute wherever an entity is used as read-only.

This applies to parameter declarations in function profiles and in object declarations.

Example 1: A value is computed once and reused later without modification.

```
const size_t max = count * sizeof(something) + offsetFoo;
size_t i;
for (i = 0; i < max; i++) {
```

Example 2: If a reference or pointer value is used to access a read-only area (at least through this pointer or reference), use the **const** attribute. This is especially important for function declarations since this establishes a

more precise contract with the caller.

```
void LogMessage(const char* msg, const Time& time);
```

[Rule] Use the `volatile` attribute wherever an entity is potentially asynchronously updated by another thread or some hardware mechanism.

This guarantees that the generated code will not optimize the access to the variable by caching it in a register for instance.

[Rule] Do not use `volatile` on a non-elementary types (i.e. not integer and not boolean).

The compiler can hardly guarantee an atomic access on such types and there is a risk to get randomly corrupted values.

Redesign the code so that accesses to non-elementary types are explicitly synchronized.

[Rule] Be careful when using the `const` attribute on pointers. Do you want a variable pointer to constant data, a constant pointer to variable data, or a constant pointer to constant data? A good practice is to use type names for complex pointer types.

The order of the `const` attribute, the type name and the `*` sign makes the difference. This is often the source of obscure compilation error messages.

Example:

```
const char* a;           // a variable pointer to "const char"
char* const b = ...;     // a constant pointer to "char"
const char* const c = ...; // a constant pointer to "const char"

a = "abcd";             // OK
*a = 'x';               // ERROR
b = "abcd";             // ERROR
*b = 'x';               // OK
c = "abcd";             // ERROR
*c = 'x';               // ERROR
```

It is more readable to use type declarations:

```
using CharPointer = char*;
using ConstCharPointer = const char*;

ConstCharPointer a = nullptr;
const CharPointer b = ...;
const ConstCharPointer c = ...;
```

[Rule] Never use built-in integer types such as `int` or `long`, unless explicitly required by the context.

The C and C++ standards do not define the implementation of the `int`, `short`, `long` types and their variants (`unsigned`, `long long`, etc.).

Specifically, the types `int` and `long` are known to have different sizes on different platforms. Using them reduces the portability of the code by introducing subtle side effects on large values.

Of course, in the presence of a legacy library which uses `int` or some of its variants in an API, you are obliged to use the same type for the data which are used with this API.

Use the following standard types when the size of integer values matter. They are defined in the standard headers `stdint.h` and `inttypes.h`.

```
int8_t  uint8_t
int16_t uint16_t
int32_t uint32_t
int64_t uint64_t
```

[Rule] Use the predefined type `size_t` for values which hold the size in bytes of a C/C++ object.

This is the standard definition in the C/C++ language. The C and C++ standards define `size_t` as the return type for the `sizeof` operator. The modern C/C++ standard libraries use this type in the profile of their functions for size values.

The actual implementation of `size_t` differs from one platform to another. It has typically the same size as a pointer. But there is no unique correspondence with the built-in integer types.

For instance, some 64-bit platforms define `int` as 32-bit and `long` as 64-bit while other 64-bit platforms define both `int` and `long` as 64-bit. But in all cases, `size_t` is 64-bit on these platforms. This is especially true when porting between Windows / Visual Studio and Linux / GCC.

So, there are clearly at least three distinct sets of integer semantics:

- Built-in types (`int` et al.): not portable, no specific semantic, unpredictable, do not use.
- Size of objects (`size_t`).
- Values represented by a given number of bits (`uint8_t` et al.)

Variant: The standard type `ssize_t` declares a signed integer type of the same size of `size_t`.

[Rule] Never use the built-in type `char` for any other purpose than 7-bit ASCII characters.

Do not use `char` to represent array of bytes, use either `int8_t` or `uint8_t`. Do not use `char` to represent characters outside the 7-bit ASCII range. In the context of internationalization, the binary representations vary. This can be a dedicated 8-bit character set (Latin-1, Latin-9, etc.) or some binary representation of UNICODE (UTF-8, UTF-16, UTF-32, etc.)

In all cases, the management of internationalized character strings shall be encapsulated into some dedicated library.

In TSDuck source code, the types `ts::UChar` and `ts::UString` implement Java-like characters and strings (UTF-16 with surrogate pairs).

[Rule] Never use the built-in type `unsigned char`, nowhere, never.

The type `unsigned char` represents nothing. It has no semantics at all.

Either a data is an ASCII character and it is a `char` or it is a byte and it is an `int8_t` or `uint8_t`.

[Rule] Always use literal constants which are type-compatible with the context.

Example:

```
int  i = 17;
long l = 17L;    // 'L' suffix means a literal of type long
char c = '\0';   // do not use integer literal such as: char c = 0;
void* p = nullptr; // C++ notation for a null pointer, not always binary zero
```

It is important to understand why using the `L` suffix is essential in integer literals which are used in the context of an expression the final type of which is `long`. Consider the following example.

```
const int i = 1;
const long l1 = i + 0xFFFFFFFFL;
```

```
const long l2 = i + 0xFFFFFFFF; // same literal without 'L' suffix
```

You may think that `l1` and `l2` have the same value. This may be the case. But they don't on some platforms where `l1` gets `4294967296` (the expected value) while `l2` gets zero.

Why?

On some platforms, `int` is a 32-bit value while `long` is a 64-bit value. In the case of `l2`, the intermediate context into which the literal is evaluated is `int` since `i` is an `int`. So, the literal `0xFFFFFFFF` is *first* evaluated into the context of an `int` and its value is `-1`. The addition is performed on `int` values, giving zero. *Finally*, the result is promoted to long, staying zero.

On the contrary, in the case of `l1`, the literal is explicitly of type `long` because of the trailing `L` suffix. Thus, the intermediate context into which the addition is performed is `long`. This time, `i` is *first* promoted to `long` and *then* the addition is performed on `long` values, giving the expected result.

Portability issue: The `L` suffix is the standard way to represent `long` literals. But we discourage the usage of predefined integer types (see the corresponding rule above) for portability reasons. For integer types which are explicitly smaller than 64 bits (`int8_t`, `int16_t`, `int32_t`, etc.), using `int` literals is usually safe. For explicit 64-bit integer types (`int64_t`, `uint64_t`), it is recommended to use the suffix `LL` in literals. The corresponding value is of type `long long`. The size of this type is not guaranteed either but it is at least 64 bits on all known implementations.

[Rule] Use the keyword `nullptr` for null pointers, not the literal `0`, not the macro `NULL`.

This is the only non-ambiguous way of specifying a null pointer.

The old macro `NULL` was specified for the C language, not C++. Early specifications of the C++ language commanded to use the literal `0` for null pointers. This has led to many ambiguous or incorrect code since this literal has two distinct semantics: integer zero of type `int` and null pointer to any type.

In TSDuck, when allowed as a compiler option, using zero as null pointer literal is detected and rejected using the appropriate warning configuration (e.g. option `-Wzero-as-null-pointer-constant` with GCC and Clang).

[Rule] Never change the order of `enum` values in a type declaration.

The `enum` types are strictly defined in the C and C++ standards. Specifically, in the absence of explicit numerical value, `enum` values are defined as implicitly consecutive numerical values. It is safe for C/C++ code to compare `enum` values using `<` or `>` operators.

If you change the order of the declarations in an `enum` type for simple aesthetic reasons (sort in alphabetical order for instance), you change the semantics of the type and you potentially break the user's code. This is why re-ordering the values of an `enum` type should be avoided.

If you add a new value in an `enum` type, think about why you add it:

- If this is simply a new value, add it at the end of the type, regardless of aesthetic reasons.
- You may insert it somewhere else only if there is a very good reason to insert it between two existing values, for instance because this is new logic state between two previously consecutive states in a strictly ordered state machine.

[Rule] Always use a type definition (using) for each meaningful type. Do not use basic types which somehow reflect the implementation of the type.

Example:

```
using FooIndex = uint16_t;
```

This facilitates the maintenance if the implementation changes someday.

In the above example, if a 16-bit integer becomes too short some day for a `Foo` index, just change the type definition. Otherwise, you would have to change all references to `uint16_t` as `uint32_t` in all variables which have the semantics of a `Foo` index and only those `uint16_t`. Needless to say that the probability to introduce a bug at this stage is very high.

[Rule] In type definitions, use the C++ syntax `using name = definition`. Do not use the old C `typedef` syntax.

The `using` syntax is more clear than `typedef`. The type name is clearly isolated before the `=` sign.

With `typedef`, in some complex cases such as function pointers, the type name is buried into the type definition and not easy to spot.

6.4.1.7. Assertions

Assertions are statements which are inserted in the middle of executable code. They assert a condition and abort the execution if this condition is false. Assertions are included in the generated code under some specific compilation options and removed by the compiler for the final production code.

Assertions are very useful to check the consistency of the code.

[Rule] Assertion must be used to check the internal consistency of code. Assertions must never be used to check input values or other external conditions.

Assertions are debug tools exclusively. Removing them shall not change the semantics of the code.

An assertion shall be used only to check the internal consistency of the code. This means that an assertion shall be used to assert something that you know is always true, regardless of the external environment, if your code is right.

In other words, an assertion contains a condition which is always true if your code is bug-free, even in the presence of incorrect inputs or incorrect environment.

For instance, if your code is such that an index `i` shall always be less than some maximum value at the end of a loop, you may write `assert(i < max)`.

Similarly, if you are about to write a data structure `ds` into a memory area `ma` which is defined as an array of `uint8_t`, you may write `assert(sizeof(ds) <= sizeof(ma))` and then `memcpy(&ma, &ds, sizeof(ds))`.

But, under no circumstances, you should use some external condition such as an input parameter of the current function in an assertion. The semantic of an assertion is that it shall have no effect on correct code. The compiler may include or remove the assertion from the generated code and the execution of correct code should be exactly the same in both cases.

Checking an external or input parameter in an assertion is illegal since an invalid external condition is not a symptom of incorrect code (at least your code). If you use an input parameter in an assertion in correct code, this code will either run or fail in the presence of invalid input, depending on the compilation option.

The correct way of checking inputs is to use plain checking code, not assertions. In the presence of invalid inputs, write the appropriate error management path for the context (return an error code, log an error, perform a default action or whatever is required by the context).

Example:

```
int8_t buf1[(3 * TS_SIZE1) + 4];
int8_t buf2[TS_SIZE1 + TS_SIZE2];

// correct, sizes are compile-time constants
assert(sizeof(buf1) >= sizeof(buf2));
memcpy(&buf1, &buf2, sizeof(buf2));

int8_t* buf3 = (int8_t*)(malloc(someComputedSize));
```

```
// incorrect, depends on system run-time resources
assert(buf3 != nullptr);

// correct, valid run-time check regardless of compilation options
if (buf3 == nullptr) {
    // do some error processing
}

// incorrect, depends on a run-time value
assert(sizeof(buf1) >= someComputedSize);

// correct, valid run-time check regardless of compilation options
if (sizeof(buf1) < someComputedSize) {
    // do some error processing
}
memcpy(&buf1, buf3, someComputedSize);
```

[Rule] Assert everything that could be asserted.

Use assertions as often as necessary. When writing code, always try to locate all the implicit assumptions you make (data size, `enum` values ordering, etc.). Locate all the conditions that must be true and which would corrupt something if incorrect (final index values, counters, etc.) Assert all of this.

When writing code, you easily assume conditions. You know that they are true by design. But you may be wrong (nobody is perfect). And some future maintenance may break the design. So, even if this seems futile, write assertions on critical conditions.

And remember that assertions are automatically removed by the compiler in production code. So, do not hesitate to use assertions, they are performance-free on production code.

[Rule] The debug code which is conditionally compiled (typically using some preprocessor symbol such as `DEBUG`) shall not modify the functional behavior of the code.

The reasons are the same as for the previous rule.

Such debug can typically only log debug information.

Be aware, however, that the presence of debug code may have a significant impact on timing.

6.4.1.8. Secure coding

[Rule] Always initialize data with predictable values. Use zero for integers and `nullptr` for pointers if there is no other meaningful values in the context.

Uninitialized variables are a common source of bugs which may be difficult to track.

Getting different behaviors on different platforms or versions depending on different unmanaged initial values is even worse.

Example:

```
uint32_t x = a + b;
uint32_t y = 0;
uint32_t* p = nullptr;
```

[Rule] Do not use `memset()` to initialize structures. Use appropriate initial value for each type.

All binary zero is often not appropriate as initial value for a data structure.

Good code:

```
struct Foo {
    uint8_t* p;
    size_t   size;
};

Foo obj;
obj.p = nullptr; // not always binary zero
obj.size = 0;
```

Bad code:

```
struct Foo {
    uint8_t* p;
    size_t   size;
};

Foo obj;
memset(&obj, 0x00, sizeof(obj));
```

[Rule] Make no assumption about the size or memory layout of a data structure.

Each platform has its constraints on data alignment. The compiler takes them into account when representing a data structure. The same data structure may have different sizes or memory layout on different platforms.

Example:

```
struct Foo {
    uint8_t a;
    uint32_t b;
} Foo;
```

The size of this structure may be 5 or 8 bytes, depending on the platform. The field **b** may be at offset 1 or 4 from the beginning of the structure.

If your code depends on the fact that **b** is assumed to be at a specific offset, you should assert this. The following example asserts that **b** is at offset 1 after **a**:

```
assert((((char*)&((Foo*)nullptr)->b - (char*)&((Foo*)nullptr)->a) == 1));
```

This is quite a strange expression since the null pointer is explicitly dereferenced. But it is used only to evaluate an address and not a content. Thus, the pointer is never really dereferenced.

For the sake of clarity, the above code can be rewritten as:

```
assert(offsetof(Foo, b) == 1);
```

[Rule] Make no assumptions about the evaluation order of operands or function parameters.

The C and C++ languages do not specify the order of evaluation in expressions and function calls. If any operand or parameter has side effects, the final result is unpredictable and may vary from one platform to another.

Counter-example: Bad code which rely on the order of evaluation:

```
int a = 2;
int b = a * a++; // WRONG: result can be either 4 or 6
f(print(a), print(b)); // WRONG: a and b may be printed in any order
```

[Rule] Do not assume that a null pointer (`nullptr` in C++) is represented by binary zeroes.

The C and C++ standards state that a zero value in a pointer context is interpreted as a *null pointer*. But, on some platforms where zero is a valid address, the compiler may use a non-zero bit pattern to represent the concept of a null pointer.

This rarity makes the problem even more dangerous. If your code relies on the fact that a null pointer is represented by a zero bit pattern, it will work on most platforms. Later, when recompiling the code on a new platform where null pointers are represented differently, the code will no longer work. And it will take ages to find the problem.

[Rule] In a function, never return a pointer or a reference to a local variable or a parameter of this function.

The local variables go out of scope upon return. The returned pointer or reference points to a portion of the stack which is no longer used. Worse, the referenced memory area will be reused by the local variables of another function call.

Consequently, returning pointer or a reference to a local variable or a parameter may lead to subtle random memory corruptions in the later functions. This kind of problem is a nightmare to debug.

[Rule] As a generalization of the preceding rule, always consider the lifetime of an object before passing it using a pointer or a reference to some other entity.

The preceding rule addresses this problem for statically allocated objects in the context of a mono-thread application. But similar problems exist with dynamically allocated objects and multi-threaded applications.

Object lifetime: First, what is the lifetime of an object? This depends on the nature of the object. If the object is directly declared in a function or code block, its birth is its declaration point and its death occurs when the execution flow exits from the function or block. If the object is dynamically allocated (`malloc()` or equivalent in C, `new` in C++), the allocation is its birth and the corresponding deallocation (`free()` or equivalent in C, `delete` in C++) is its death.

Dynamic allocation: When a function dynamically allocates an object and returns this pointer or passes it to multiple modules, the function loses all tracks of the usage of the pointer by other entities. When shall this object be deallocated is a complex question which must be addressed at the application design level.



C++ may mitigate this problem using implementations of the smart pointer design pattern, typically `std::shared_ptr` in modern C++.

Multi-threaded applications: If a function passes the address of one of its local variables to another thread, there is a risk that the function returns while the other thread still uses the referenced data. To protect from this, the function may synchronize its return point with the termination of the other thread (*join* operation). Otherwise, this problem must be addressed at the application design level.

[Rule] When you declare a function accepting some form of array or raw data buffer, always add a parameter to specify the size of the data (input) or maximum size of the buffer (output).

Never assume that the format of the data will reliably define the end of the data.

Example:

```
// WRONG: how much data should I send?
void SendData1(const void* data);

// OK
void SendData2(const void* data, size_t dataSize);
```

Of course, in the second case, the correct declaration does not mean that the implementation is correct. But it is possible to write a correct implementation for this function.

On the other hand, the first function can never be safe and reliable.

[Rule] Never call any function with the following characteristics:

- A parameter is an address where the function is supposed to write to.

- *The function may write more than one element at that address.*
- *No parameter gives the maximum number of elements to write to.*

Many functions write data to a user-supplied buffer without letting the user specify the size of this buffer. These functions are simply badly defined and badly designed. They must be thrown away, simply.

[Rule] *As a corollary of the preceding rule, never use any unsafe predefined string functions such as `sprintf()`, `strcpy()`, `strcat()`, all other `strXXX()` functions of the standard `libc` and many other standard functions.*

Many standard functions are inherited from the old days of the UNIX operating system in the 70's. They are obviously flawed in their definition. Just think twice about using standard functions, especially old ones.

Note that safe alternative versions exist for most of these functions: `snprintf()`, `strncpy()`, etc. Although safer than their ancestors, these functions still require some care. The pain with C-strings is the final null termination character. When a safe standard C-string function (one with an `n` in the middle of the name) copies a string into a buffer and the string is potentially larger than the buffer, the final string is truncated to the size of the buffer. In this case, truncation means no final null character. This also means that the result is *not a valid C-string*. If you use it as a C-string, the results are unpredictable (collecting memory garbage or crash).

Thus, when using the safe `n` C-string functions, you must always either check the return value to detect overflow (if supported by the function) or systematically overwrite the last byte of the buffer with a null character.

Example: The function, here `strncpy()`, does not indicate if a truncation occurred:

```
char s[10];
strncpy(s, longText, sizeof(s));
s[sizeof(s) - 1] = '\0'; // force valid C-string if truncated
```

Example: The function, here `snprintf()`, returns a value which can be used to detect the truncation:

```
if (snprintf(s, sizeof(s), "%s", longText) >= int(sizeof(s))) {
    // result was truncated,
    // either process error or force valid C-string as above
}
```

[Rule] *Zero out pointers after `free (C)` or `delete (C++)`.*

By zeroing the pointer, any accidental usage will result into an access violation (immediate crash) instead of a *reuse-after-free* or *double-deallocation* (subtle memory corruption which may run undetected for sometimes and hard to debug).

Example:

```
Foo* p = new Foo(...);
...
if (...) {
    delete p;
    p = nullptr;
}
*p = ...; // immediate crash if preceding branch was taken
```

[Rule] *Never cast a pointer into an integer.*

First, converting between integers and pointers is bad design. Second, you do not know the size of a pointer and, on some platforms, casting a pointer to an integer results in a loss of information.

If you need to declare a field which stores a generic pointer, use `void*`.

If you need to perform low-level address arithmetic on bytes (instead of C elements), cast the pointers into `uint8_t*`.

[Rule] *Be careful when mixing the `sizeof` operator with pointer or array index arithmetic. Remember that the pointer arithmetic uses the element size as a base while the `sizeof` operator returns a value in bytes.*

Do not use `sizeof` to compute the last index in an array; this works only for arrays of `char` and 8-bit integers.

Counter-example: Incorrectly setting the last element of an array, leading to a crash (at best) or some random memory corruption (at worst).

```
uint32_t a[100];

// WRONG: not the last element but far, far away in memory
a[sizeof(a) - 1] = 0;
```

[Rule] *To determine the number of elements in an array, do not use the explicit size from the declaration of this array. Use `sizeof` in an appropriate formula.*

Otherwise, the maintenance becomes difficult if the declaration of the array is modified.

Example:

```
Foo array[SOME_COUNT];

// WRONG: what if the declaration of array is modified ?
const size_t numberOfElements = SOME_COUNT;

// This method is independent from the declaration of the array
const size_t numberOfElements = sizeof(array) / sizeof(array[0]);
```

[Rule] *Do not declare arrays as local variables.*

This is a corollary of the rules on buffer overflow and stack overflow.

An array can be very large and may overflow the stack of the current thread.

Moreover, using an array is subject to buffer overflow, even if you make your best efforts to avoid that. On a security standpoint, a buffer overflow is much more dangerous on the stack than on the heap. The stack is full of "code addresses", typically return addresses or exception handlers. A buffer overflow on the stack is the typical vulnerability which can be exploited for malware code injection. Even if there are some "code addresses" in the heap, they are less common. This is why a buffer overflow is less dangerous on the heap than on the stack.

This problem is more frequent in C than C++. With C++, we do not use arrays; we use vectors or other container classes. In instance of such a class is typically implemented as a short data structure containing pointers. The actual data are contained in dynamically allocated areas which are transparently managed by the container class.

[Rule] *Never mix signed and unsigned integers.*

Mixing signed and unsigned integers, typically dealing with index or offset values, is extremely dangerous and should be banned. The weird effect appears when offset values are decremented towards zero. After zero, the signed value goes negative while the unsigned one wraps up to $2^n - 1$. The two values are subsequently desynchronized.

Some coding guidelines advocate the usage of unsigned integers to manipulate sizes "because they never become negative". This is not always a good idea. An unsigned value never becomes negative for sure, but it becomes "extremely large" instead, which is not better or even worse in some case. It is true that the C/C++ standards use the type `size_t` for sizes and `size_t` is unsigned. But standards may have bad ideas too.

Counter-example: The following code is a very common way to explore a buffer backward.

```
i = size;
while (--i >= 0) {
    // do something on buffer[i]
}
```

The code is valid when the iterator `i` is signed. However, if `i` is unsigned, of type `size_t` for instance, this code loops forever (or most certainly crashes if `i` is used to access memory).

Do you understand why Java, a language "designed with security in mind", has no unsigned type?

[Rule] Anticipate and prevent integer arithmetical overflow and underflow.

In C/C++, the integer types always implement modular arithmetic. All operations are performed modulo 2^n . Arithmetical operations never overflow, they wrap up or down. In some cases, this behaviour is fine. But in most cases it is not. In fact, integer arithmetic and modular arithmetic are two different domains with distinct usages but C/C++ implements only one of the two. Other programming languages have a different approach. The Ada language, for instance, defines two distinct families of integer types: integer types and modular types. All operations on integer types remain in the range of the type and throw an exception in case of underflow or overflow. Modular types, on the other hand, implement the same semantics as C/C++. Based on the application requirements, the developer chooses the appropriate type. But this is not possible in C/C++ or Java. The developer has to implement all checks manually.

The type of check depends on the operation (addition, subtraction, multiplication) and the signedness of the integer type. Note that an integer division never overflows or underflows but the divider shall be checked to prevent "divide by zero" errors. To check overflows on addition and subtraction, we need to know the maximum and minimum values of the integer type. To check overflows on multiplication with signed types, we need to know the proper function to compute an absolute value.

In C++, the standard header `<limits>` declares the template class `std::numeric_limits`. This class provides a generic mechanism to determine various properties of numeric types, either integer or floating point. The header `<cstdlib>` declares overloaded versions of `std::abs()` for all predefined signed integer types.

The characteristics of an application-defined integer type `Foo` are:

- Minimum value: `std::numeric_limits<Foo>::min()`
- Maximum value: `std::numeric_limits<Foo>::max()`
- True when the type is signed: `std::numeric_limits<Foo>::is_signed`
- Absolute value: `std::abs(value)`

For some reason, in the template class `std::numeric_limits`, `min()` and `max()` have the syntax of a function while `is_signed` has the syntax of a constant. But the functions are inlined and return a constant. So, there is no performance penalty.

Detecting underflow on subtraction with unsigned integers

The following code illustrates a very dangerous situation of underflow.

```
size_t start = ...;           // some starting index inside an array
size_t current = ...;        // current exploration index
size_t size = current - start; // size the explored area
```

When `current` is less than `start`, the operation `current - start` underflows since `size_t` is an unsigned type. The variable `size` receives a very large positive value, close to 2^{32} or 2^{64} , depending on the platform. When `size` is later used as an actual data size, the most likely result is a buffer overflow, the best known security vulnerability, leading to all sorts of malware infections.

So, the code should be rewritten as follow:

```
size = current < start ? 0 : current - start;
```

The test `current < start` anticipates and detects the possible underflow in the subsequent subtraction. In case of detected overflow, the code uses an appropriate alternative (here a zero value).

Detecting overflow on addition with unsigned integers

Detecting overflows is a bit more complicated than underflows because the wrapping value, where the overflow occurs, depends on the size of the integer types. The following code illustrates how to detect a potential overflow:

```
size_t start = ...; // some starting index inside an array
size_t size = ...;  // size of an area inside the array
size_t next;        // will receive the index after the area

if (start > std::numeric_limits<size_t>::max() - size) {
    // process arithmetic overflow
}
else {
    next = start + size;
    if (next >= size_of_array) {
        // process buffer overflow
    }
    else {
        // now you can safely process the array at index "next".
    }
}
```

This is the minimum required code to safely move forward into a buffer using index values. Note the different cases for buffer overflow and arithmetic overflow on index computation.

Checking the potential overflow must be done using operations which never overflow or underflow. The operation `std::numeric_limits<size_t>::max() - size` is safe because the integer type is unsigned and `std::numeric_limits<size_t>::max()` is its maximum value.

On the other hand, the naive test `if (start + size > std::numeric_limits<size_t>::max())` is both wrong and useless. It is wrong because the addition may overflow and its value is not reliable. And it is useless because the test is always false since no `size_t` value can be greater than `std::numeric_limits<size_t>::max()`.

Detecting underflow or overflow on addition with signed integers

With signed integers, the addition generates an error in the following cases.

- Overflow: Both operands are positive and the result is negative.
- Underflow: Both operands are negative and the result is positive.

Summary of underflow or overflow detection by type and operation

The following table summarizes the right ways to detect error conditions on the various arithmetical operations on signed and unsigned types. The generic names `MIN`, `MAX` and `ABS()` are used to designate the minimum value, maximum value and function returning the absolute value for the given type. In C, you have to select the right constants and functions. In C++, use the generic mechanisms of the language. The operands are named `a` and `b`; the result is named `c`.

Addition (unsigned):

```
if (a > MAX - b) {
    // ERROR
}
else {
    c = a + b;
}
```

Addition (signed):

```
c = a + b;
if ((a > 0 && b > 0 && c <= 0) || (a < 0 && b < 0 && c >= 0))
{
    // ERROR
}
```

Subtraction (unsigned):

```
if (a < b) {
    // ERROR
}
else {
    c = a - b;
}
```

Subtraction (signed):

```
c = a - b;
if ((a > 0 && b < 0 && c <= 0) || (a < 0 && b > 0 && c >= 0))
{
    // ERROR
}
```

Multiplication
(unsigned):

```
if (b != 0 && a > (MAX / b)) {
    // ERROR
}
else {
    c = a * b;
}
```

Multiplication (signed):

```
if (b != 0 && ABS(a) > ABS(MAX / b)) {
    // ERROR
}
else {
    c = a * b;
}
```

Division (signed or
unsigned):

```
if (b == 0) {
    // ERROR
}
else {
    c = a / b;
}
```



Due to rounding effects, the check on multiplication is a bit too aggressive. It detects all error cases but some marginal valid cases, close to the limit, may be incorrectly detected as errors. However, implementing accurate overflow detection on multiplication in the general case is very costly in terms of performance. The proposed solution is a tradeoff.

In TSDuck, the header file `tsIntegerUtils.h` declares template functions such as `ts::bound_check()`, `ts::bounded_add()` or `ts::bounded_sub()`.

[Rule] Anticipate and mitigate floating point rounding

Floating point types are less prone to underflow and overflow than integers because they are always signed and their minimum and maximum values are very large. If you really need to check for overflow and underflow, use the same technique as for signed integers.

The main problem with floating point types is rounding and the accumulation of imprecisions after an arbitrary large sequence of arithmetical operations.

Never use the equality (`==`) or inequality (`!=`) operators on floating point values because these operators check the (in)equality of binary representations without taking into account rounding and imprecisions.

There are various complex ways to safely compare floating point values. But an accurate implementation should take into account the complete sequence of arithmetical operations which led to the values to compare. For a given sequence of operations, there is a given accumulated imprecision and this imprecision must be taken into account when comparing values. In practice, this is too complicated - and way too slow - to implement.

So, we must adopt a pragmatic approach when comparing floating point values.

Example 1: If we need a relative precision, we should reduce the comparison toward `1.0` and compare it with the *epsilon* value of the floating type. However, since *epsilon* is defined in the C++ standard as the *"the difference between 1 and the least value greater than 1 that is representable in the given floating point type"*, the accumulated impression may have exceeded *epsilon*. So, in practice, it is recommended to compare to some small multiple of *epsilon*.

The following C++ function implements a generic numerical equality operator. It takes an additional parameter named `impr` (for imprecision) which is the small multiple of *epsilon*. The default value is 4 and is adequate in most cases. However, if you compare values which were created by long sequences of arithmetical operations, it is probably better to use a higher value.

```
template <typename T>
bool probablyEqual(T a, T b, int impr = 4)
{
    if (std::numeric_limits<T>::is_integer) {
        // T is an integer type, equality is valid.
        return a == b;
    }
    else {
        // T is a floating point type, equality is fuzzy.
        const T c = std::max(std::abs(a), std::abs(b));
        return c == 0.0 || std::abs(a - b) / c <= impr * std::numeric_limits<T>::epsilon();
    }
}
```



This example is just for clarity. In practice, we would implement one version for integer types and one version for floating point types, using SFINAE.

Usage:

```
float a, b;
...
if (probablyEqual(a, b)) {
    // Consider that a "pragmatically" equals b.
}
```

In TSDuck, the header file `tsFloatUtils.h` declares the template function `ts::equal_float()`.

Example 2: In some contexts, a relative precision is useless and an absolute precision is sufficient in practice. If you handle financial values for instance, the smallest amount of currency is usually a cent.

This means that everything smaller than `0.01` is irrelevant. So, regardless of the values and how they were computed, it is much simpler to write comparisons this way:

```
float a, b;
...
if (std::abs(a - b) < 0.005) {
    // Consider that a "pragmatically" equals b for financial data.
}
```

6.4.1.9. C++ classes

[Rule] The `struct` data types shall be avoided. Use classes instead.

The `struct` data types are allowed in C++ for C compatibility only. They are strictly equivalent to classes except that their fields without explicit scope are public by default (private by default in classes). So, there is no good reason to use `struct` in C++.

[Rule] Always declare a method of a class as `const` when it does not modify the object instance.

The contract is better defined. It asserts that the method will not modify the object. It also allows the method on `const` objects.

Example:

```
class C
{
public:
    void reset();          // this method will potentially modify the object
    void print() const;    // this method will not modify the object
};

void f(const C& c)
{
    c.print();             // OK
    c.reset();             // ERROR: cannot modify a const object
}
```

[Rule] Avoid public data members in class declarations.

Public data members give no control on the way the data members are used. Further maintenance is more difficult if you need to modify the internal structure of the class since you need to maintain a flawed contract.

Instead of publishing the data members, keep them private and export public getters and setters.

Good code:

Bad code:

```
class Foo
{
public:
    int getCount() const {return _count;}
    void setCount(int c) {_count = c;}
private:
    int _count = 0;
};
```

```
class Foo
{
public:
    int count = 0;
};
```

These two implementations have identical performances, thanks to the inlining of the setter and getter.

But the good code is much more maintainable. You may modify the internal structure of the class, you may completely remove the `_count` field, you may add bound checking in the setter, etc. But you will always be able to maintain the same contract (the two public methods `getCount` and `setCount`).

[Rule] When you provide a method with a `const std::string&` parameter, consider the advantages of providing an overload with `const char*` parameter. The same rule applies any other string type.

The type `std::string` provides a constructor from `const char*`. So, you may invoke a method with a `std::string` parameter using a null terminated C-string or a string literal. But, you should inspect how you use the `std::string` parameter within the method. If you simply use its null terminated C-string representation (method `std::string::c_str()`), then it is probably faster to overload the method.

The same rule applies to any string type, for instance `ts::UString`.

Slower code:

```
void f(const std::string& s)
{
    something(s.c_str());
    ...
}

std::string aString;
f(aString);
f("abcd"); // converted to std::string but
           // only used as C-string internally
```

Faster code:

```
void f(const char* s)
{
    something(s);
    ...
}

inline void f(const std::string& s)
{
    f(s.c_str());
}

std::string aString;
f(aString);
f("abcd"); // no conversion
```

[Recommendation] Try to define interface classes when possible. An interface class describes the API for a service which can be implemented by concrete classes.

The characteristics of an interface class are:

- All methods are public.
- All methods are pure virtual.
- There is no field.
- There is no static item.
- There is no constructor.
- There is an empty virtual destructor.

By convention, the name of an interface class ends with the suffix **Interface**.

The concept of *interface* class is not defined by the C++ standard. The C++ standard defines the concept of *abstract* class but it is less restrictive; an abstract class is a class with at least one pure virtual function.

The interface class is the C++ equivalent of the concept of *interface* in Java. The C++ interface class is interesting because it proposes a reasonable approach to multiple inheritance. Because of its characteristics, an interface class is typically limited to a header file.

Example:

```
class MutexInterface
{
public:
    virtual bool acquire() = 0;    // pure virtual method
    virtual bool release() = 0;   // pure virtual method
    virtual ~MutexInterface() {}  // empty virtual destructor
};
```

[Rule] Do not use actual multiple inheritance. A class should inherit from at most one non-interface class. A class may inherit from multiple interface classes.

Multiple inheritance is a powerful but dangerous concept. Using multiple inheritance from more than one concrete class often leads to structures which are complex and difficult to understand and maintain. After several levels of multiple inheritances, the result can be inconsistent sometimes.

The coding rules of some organizations even completely ban all forms of multiple inheritance.

Our coding rules propose a reasonable trade-off which is derived from Java: one "real" superclass and as many interface classes as necessary (the preceding recommendation defines the non-standard concept of interface class).

[Recommendation] At the beginning of the declaration of each subclass, define a type definition with the name **SuperClass for the superclass.**

This convention gives an easy and reliable way to explicitly invoke a method of the superclass. This is the C++ equivalent of the **super** keyword in Java.

If the class hierarchy is refactored and a new intermediate class is inserted, simply modify the type definition for **SuperClass**, there is no need to browse the entire implementation to track explicit invocations of the superclass.

Note that C++ allows multiple inheritance and it is not possible to define a single superclass in the general case. But the preceding rule limits the number of non-interface superclasses to one. This single non-interface superclass is considered as the "actual" superclass which is declared with the **SuperClass** type definition.

Exception: If a class has no non-interface superclass, the **SuperClass** type definition is omitted.

Example:

```
class A
{
public:
    void f();
};

class B: public A
{
public:
    // Standard name for superclass, modify it if you change the inheritance
    using SuperClass = A;
```

```
// Override method
void f()
{
    // Explicit (and maintainable) invocation of superclass
    SuperClass::f();

    // Specific additional processing in subclass
    ....
}
};
```

[Rule] *In the documentation of a method of a class, clearly indicate if the method needs to be explicitly invoked by subclasses which override the method. If necessary, also indicate if the superclass method shall be invoked at the beginning or at the end of the subclass method.*

This is entirely dependent on the implementation of the superclass. Only the developer of the superclass can decide what is necessary.

Example:

```
class File
{
public:
    // Invoke at beginning of overridden method in subclasses
    void open(const std::string& fileName);

    // Invoke at end of overridden method in subclasses
    void close();
};
```

When we specialize the class **File** to add buffering capabilities for instance, the methods **open()** and **close()** must be typically overridden to add the setup and flush of the buffer, respectively. But, while the buffer management is handled in the subclass, the file management shall remain in the superclass.

```
class BufferedFile: public File
{
public:
    using SuperClass = File;

    void open(const std::string& fileName)
    {
        SuperClass::open(fileName);
        // ... setup the buffer
    }

    void close()
    {
        // ... flush the buffer
        SuperClass::close();
    }
};
```

Exception: This rule does not apply to constructors and destructors. The compiler automatically invokes the superclass constructor at the beginning of the subclass constructor. It also automatically invokes the superclass destructor at the end of the subclass destructor.

[Rule] *Never overload methods which differ only in the type of integer or pointer parameters.*

In some contexts, the default integer type conversions give unexpected results. In case of maintenance of the code, you may even accidentally break the contract of the class.

Example:

```
class C
{
public:
    void f(int);
    void f(long); // BAD PRACTICE
};

int a = 1;
long b = 2;

C x;
x.f(a);
x.f(b);
```

In this example, `x.f(a)` invokes `C::f(int)` and `x.f(b)` invokes `C::f(long)`.

But let's assume that the initial version of the class `C` only had one method `C::f(int)`. The user code `x.f(b)` invoked `C::f(int)` since this was the only possible method at that time and the C++ compiler automatically downgraded `b` from a `long` to an `int` before the call.

Now, in the maintenance phase, assume that you add the `C::f(long)` method. When recompiling the existing user code, `x.f(b)` no longer invokes `C::f(int)`. Now it invokes `C::f(long)` since this is a better match. And the two methods may perform very different actions. So, by adding the overload `C::f(long)`, you break the contract and break the ascending compatibility of your class and the compiler does not even tell you it is different now.

Note 1: This rule is also valid for pointers and integers. Using the literal `0` can be indifferently interpreted by the compiler as an integer literal or as the null pointer.

Note 2: The terms *override* and *overload* shall not be confused. A method of a class *overrides* another when it redefines a method with the same name and same profile in a superclass. A method *overloads* another when it redefines a method with the same name and a different profile in the same context.

Example:

```
class A
{
public:
    void f(int i);
};

class B: public A
{
public:
    void f(int i);           // override A::f(int)
    void g(int i);           // overload B::g(const std::string&)
    void g(const std::string& s); // overload B::g(int)
};
```

[Recommendation] When an overloaded method is overridden in a subclass, override all overloaded variants, unless there is a good reason to hide some of them or to add a new one.

Hum? This sounds a bit mysterious...

Counter-example: Explanations are always better with an example.

```

class A
{
public:
    void f(int i);
    void f(const std::string& s);
};

class B: public A
{
public:
    void f(int i);
};

A a;
B b;
a.f(1);           // OK, call A::f(int)
a.f("foo");       // OK, call A::f(const std::string&)
b.f(1);           // OK, call B::f(int)
b.f("foo");       // ERROR, does not compile, B::f(const std::string&) is hidden

```

In class **A**, the method **f()** is overloaded. There is one version taking an integer as parameter and one version using a character string. In class **B**, **f()** is overridden but only one version is declared, the integer one. As a consequence, the string one is hidden.

In C++, when you override a method in a subclass, all overloaded versions (methods with the same name in the superclass) are hidden. Only the explicit declarations in the subclass can be used. This is why **B::f(const std::string&)** is hidden in the above example. You cannot use it.

As a general rule, a subclass specializes a contract. It adds or replaces services but does not remove any. This rule is broken in the example.

Exception: In some cases, it can be legitimate to remove (hide) or add overloaded versions of a method. But this is tightly linked to the semantics of the object.

[Rule] Never change the default values for the parameters of a function after the function is published.

Changing the default values changes the contract.

Consider the following function:

```
void paint(Color c = BLACK);
```

All users which invoke **paint()** without argument are confident that the color will be black. If you later change the default value to **BLUE** and recompile the code, the users of the contract will obtain a different result.

[Rule] When overriding a virtual method with default parameters in a derived class, use the same default values for the parameters.

Counter-example: Let's review the consequences of breaking this rule. Assume that all shapes should be painted in black by default, except the rectangles which should be painted in blue by default.

See the following WRONG implementation:

```

class Shape
{
public:
    // "I see a red door and I want it painted black"
    virtual void paint(Color c = BLACK);

```

```
};

class Rectangle: public Shape
{
public:
    // WRONG: overridden method with other defaults
    virtual void paint(Color c = BLUE);
};
```

See why this implementation gives incorrect results:

```
Rectangle r1, r2;
Rectangle* p1 = &r1;
Shape* p2 = &r2;

p1->paint(); // Rectangle r1 is painted in BLUE
p2->paint(); // Rectangle r2 is painted in BLACK !
```

The default values for the parameters are always evaluated in the context of the caller, not the callee. When an object is accessed through a pointer to a superclass, the context of the caller is the superclass and the superclass' default parameters are applied.

This is why we impose to keep the same default values for parameters.

To solve the specific problem of the example where actual defaults should be different between the superclass and the subclass, a safe alternative would be to use overloading instead of default parameters as illustrated below.

```
class Shape
{
public:
    virtual void paint(Color c);
    virtual void paint() {paint(BLACK);}
};

class Rectangle: public Shape
{
public:
    virtual void paint(Color c);
    virtual void paint() {paint(BLUE);}
};
```

In this case, the parameter **BLUE** is evaluated in the context of the callee, which is **Rectangle::paint()**.

6.4.1.10. C++ constructors and destructors

[Rule] For a class type, always declare both default constructor and copy constructor, or none of them. If any of them should be disabled, explicitly delete them both.

The default and copy constructors are implicit if not defined by the developer. But the default implementation may not be adequate to your class. So, you shall either provide an explicit version or explicitly disable it.

If you don't and if you do not want a default or copy constructor because they are meaningless in your context, the compiler will implicitly generate wrong one for you. If any default object declaration is used, the buggy implicit constructor will be used.

Example: Disabling default and copy constructor:

```
class C
{
public:
    // Only my explicit constructors are valid (definition required)
    C(const std::string&);
    C(int, int);
private:
    // Default and copy constructors are explicitly deleted
    C() = delete;
    C(const C&) = delete;
};
```

[Rule] Always free all resources which are private to the class in a destructor.

Well-designed classes safely prevent resource leaks: memory, open files, locked resources, etc.

[Rule] Always provide a virtual destructor.

If your destructor is not virtual, it will not be invoked when an object is destructed using a delete of a pointer to one of its super-classes.

Counter-example: Class **C1** is a super-class. Subclass **C2** has a non-virtual destructor while class **VC2** has a virtual destructor. This example illustrates how instances of class **C2** can be incorrectly destructed, leading to potential resource leaks.

```
class C1 { ... };
class C2: public C1 { public: ~C2(); };
class VC2: public C1 ( public: virtual ~VC2(); );

// later in the code:
const C1* x = new C2();
const C1* y = new VC2();
delete x; // destructor ~C2 is NOT executed !
delete y; // destructor ~VC2 is executed
```

[Rule] All destructors shall be idempotent.

In rare obscure cases, a destructor may be invoked twice on the same object. This may especially happen if the destructor is explicitly invoked once and will be likely implicitly invoked later when the object goes out of scope.

Thus, the implementation of a destructor shall be coded so that it is safe when invoked more than once. This is the meaning of idempotent.

Counter-example: Non-idempotent destructor:

```
class C
{
private:
    int* _p;
public:
    C(): _p(new int[8])
    {
    }
    virtual ~C()
    {
        delete[] _p; // double deallocation if invoked twice
    }
};
```

Example: The following alternative destructor is idempotent:

```
virtual ~C()
{
    if (_p != nullptr) {
        delete[] _p;
        _p = nullptr;
    }
}
```

[Rule] In a constructor or destructor, never invoke any virtual method of the class.

During the execution of the constructor and destructor, the *vtable* of an object instance points to the virtual methods of the class which defines the current constructor or destructor. On the other hand, during the rest of the lifetime of the object, the *vtable* points to the virtual methods of the actual class of the object. This can be very misleading and hard to debug.

Counter-example: Consider the following class **C1** which incorrectly invokes a virtual method in the constructor and destructor.

```
class C1
{
public:
    virtual void vf(const std::string& s);

    C1() // constructor
    {
        vf("in constructor");
    }

    ~C1() // destructor
    {
        vf("in destructor");
    }

    void f() // some standard method
    {
        vf("in method f()");
    }
};
```

Now consider a subclass **C2** which overwrites the virtual method.

```
class C2: public C1
{
public:
    virtual void vf(const std::string& s);
};
```

Given the nature of virtual methods, for a given instance of **C1** or any of its subclasses, it may be expected that the same virtual method **vf()** will be used in all contexts. But this is not true. Let's assume that **vf()** simply displays a message:

```
void C1::vf(const std::string& s)
{
    std::cout << "C1::f: " << s << std::endl;
```

```

}

void C2::vf(const std::string& s)
{
    std::cout << "C2::f: " << s << std::endl;
}

```

Consider the following applications code:

```

int main()
{
    C2 c;
    c.f();
}

```

The virtual method `vf()` of the instance `c` is invoked three times. But this is not the same `vf()` in all cases. The application displays this:

```

C1::f: in constructor
C2::f: in method f()
C1::f: in destructor

```

Thus, invoking virtual methods in constructors and destructors is error-prone. This may even fail if the methods are pure virtual in the superclass.

Debugging hint: If an application fails with a message similar to "pure virtual method was called", look for invocations of virtual methods in constructors and destructors.

[Rule] All constructors shall properly initialize all its super-classes and member fields.

This is an application of a previous rule: all data shall be initialized.

[Rule] In the definition of a constructor, the field initializers must be in the same order as their declarations in the class declaration.

The order is significant since the field initializers can make a reference to each other. In the recommended paranoid warning mode, some compilers even reject the constructor definition when the field initializers are in a different order from the class declaration.

Example:

```

class C
{
public:
    C(int x);
private:
    int a;
    int b;
    int c;
};

C::C(int x):
    a(0),
    b(x),
    c(x + 1)
{
    ...
}

```

}

[Recommendation] *If a constructor has only one argument or if the second and subsequent arguments have default values, this constructor shall be prefixed by the `explicit` keyword.*

Invoking a constructor with one argument is implicitly used by the compiler to perform an automatic type conversion. In some cases, this is the right thing to do. But in most cases, this isn't. As a general rule, you must analyze the semantic of a constructor with one argument. Does it make sense to convert back and forth between the class of the constructor and the class of the first argument? If the answer is yes, then leave the constructor alone. If the answer is no, use `explicit`. In this context, the keyword `explicit` means that type conversions shall be explicit and the compiler will never use this constructor for implicit type conversions.

Example: In the following code, we declare two classes, `City` and `House`, which implement two different concepts. It does not make sense to implicitly convert a `House` into a `City` or vice-versa. But it is legitimate that an instance of `House` has a property of class `City` because a house is located into a city. So, we provide a `House` constructor with an argument of type `City`, specifying the location of the constructed house.

```
class City { ... };

class House
{
public:
    explicit House(const City& location, const char* name = "");
};
```

Since the second argument of the constructor has a default value, it can be invoked with only one argument of type `City`. To avoid the accidental implicit conversion of a `City` into a `House`, we use the keyword `explicit`.

Let's examine the effect of the absence of `explicit` keyword using the following incorrect code.

```
void selectHouse(const House& h);

City paris;
selectHouse(paris); // incorrect usage, we want to select a house, not a city
```

The function `selectHouse()` is probably used to select a specific existing house. So, calling `selectHouse()` with an argument of type `City` is incorrect and should be rejected by the compiler. Indeed, there is a compilation error here, thanks to the `explicit` keyword. Without this keyword, however, the compiler would implicitly use the constructor to convert the object of type `City` into `House` and `selectHouse()` is called with a new `House` object with all default values in the city of paris. This does not make sense.

Exception: When the class of the constructor and the class of its first argument have equivalent semantics with different representations, it makes sense to allow conversions between the two types. The standard type `std::string` has a constructor with one argument of type `const char*` (i.e. a C-string). The two types are semantically identical, they are character strings. So, converting between the two is both legitimate and useful. When a function uses a parameter of type `std::string`, it is useful to be able to call it with a string literal such as `"foo"` (which has type `const char*`). The compiler implicitly creates a new object of type `std::string` from the string literal. This would not be possible if the `std::string` constructor had the keyword `explicit`.

6.4.1.11. C++ operators

[Rule] *Always provide an assignment operator with the same scope as the copy constructor.*

The two operations, assignment operator and copy constructor, are closely related. They should be both enabled and implemented or both disabled.

Both enabled:

```
class C
{
public:
    // implementation required
    C(const C&);
    C& operator=(const C&);
};
```

Both disabled:

```
class C
{
private:
    // explicitly deleted
    C(const C&) = delete;
    C& operator=(const C&) = delete;
};
```

[Rule] When provided, the copy constructor and the assignment operator shall have consistent implementations.

The following two sequences shall have identical results.

Copy constructor:

```
C a;
C b(a);
```

Assignment operator:

```
C a;
C b;
b = a;
```

But beware that, although consistent, the two implementations are usually not identical. The copy constructor works on an uninitialized object with no pre-existent state. The assignment operator, on the contrary, works on an initialized object which may need some clean-up before copy.

[Rule] Prevent self-assignment in an assignment operator.

A self-assignment `a = a` is a valid but void operation which must be filtered specifically.

Example:

```
class C
{
public:
    C& operator=(const C&);
    ...
};

C& C::operator=(const C& other)
{
    if (this != &other) {
        // implement actual assignment here
        ...
    }
    return *this;
}
```

Consider a class which encapsulates complex dynamic data structures. Assume that you decide that the semantics of the assignment operator is a deep copy. In the assignment operator, you must first free the previous content of the object and then duplicate the new content from the assigned value. If the test `if (this != &other)` is not present, the effect of the self-assignment would be a *reuse-after-free* bug.

[Rule] In the assignment operator of a derived class, make sure that the superclass fields are properly assigned using an explicit invocation to the superclass assignment operator.

Failing to do this may leave the superclass fields in an inconsistent state.

Example:

```
class D: public C
{
public:
    using SuperClass = C;
    D& operator=(const D&);
    ...
};

D& D::operator=(const D& other)
{
    if (this != &other) {
        // assignment of superclass fields through explicit invocation of superclass
        SuperClass::operator=(other);
        // implement actual assignment of subclass fields here
        ...
    }
    return *this;
}
```

[Rule] All assignment operators (=, +=, -=, etc.) shall return *this.

This is required to be consistent with the semantic of these operators.

[Rule] The comparison operators, when implemented in a class, must be consistent with each other.

If you implement the operator ==, be sure to also implement != and ensure that for any instances **a** and **b** of the class:

```
(a != b) == !(a == b)
```

Similarly, if you implement any of <, <=, >, or >=, you must implement them all and ensure that for any instances **a** and **b** of the class:

```
(a < b) == !(a == b) && !(a > b)
(a <= b) == !(a > b)
(a > b) == !(a == b) && !(a < b)
(a >= b) == !(a < b)
```

In practice, it is only necessary to provide a deep implementation for operators == and <. All other operators can be implemented using references to the first two.

[Rule] When you do not use the final result of an expression using the increment (++) or decrement (--) unary operators, prefer the prefixed notation (++i, --i) to the postfix notation (i++, i--).

The postfix notation needs to create a temporary object holding the previous value of the object and use this temporary object as the result of the expression.

For integer or pointer types, this is usually optimized away by the compiler. But for more complex object classes which redefine these operators, the useless construction and destruction of the temporary object cannot be avoided. And this can cost some time for nothing.

This is especially the case for the iterators in the standard template library. The following example illustrates two functionally identical ways of walking through a list of int but the second way is uselessly slower.

```
std::list<int> x;

// Use ++i on iterator: good
for (std::list<int>::iterator i = x.begin(); i != x.end(); ++i) {
    ....
}

// Use i++ on iterator: identical but uselessly slow
for (std::list<int>::iterator i = x.begin(); i != x.end(); i++) {
    ....
}
```

6.4.1.12. C++ object management

[Rule] Never use `malloc()` and `free()` in C++, use the `new` and `delete` operators.

The C memory allocation routines return raw memory, not objects. Casting the result of `malloc()` to a pointer to an object instance is incorrect. This memory area is not a properly initialized object.

On the contrary, the C++ allocation operators manipulate objects. They invoke the proper constructors and destructors.

[Rule] Objects which were allocated using `new[ ]` must be deallocated using `delete[ ]`.

This is a really annoying feature of C++. The operators for allocating single objects and arrays are different and the corresponding `delete` operator must be used. But when a pointer is declared as `Foo*`, there is no implicit indication how the pointed object was allocated. The developer must take care of this.

[Recommendation] Never use dynamically allocated arrays (operator `new[ ]`).

The preceding rule demonstrates that using the operator `new[ ]` may create confusion when it comes to deallocation. Moreover, a previous recommendation explicitly recommends avoiding arrays of C++ objects for other reasons, whether they are statically or dynamically allocated.

Use standard containers from the C++ Standard Template Library (STL) instead of dynamically allocated arrays. The standard container `std::vector`, for instance, is a much better alternative.

If you think that you need an array because you will pass it to a C routine which requires the address of an array, a `std::vector` is still better than a dynamically allocated array. The C++ Standard specifies that the underlying representation of the current content of a `std::vector` must match the representation of an array. Thus, the following sample code is both legal and safe.

```
// A C routine which expects an "array"
void legacyFunction(int* buffer, size_t intCount);

std::vector<int> v;
v.resize(100); // or fill the vector the way you like
legacyFunction(&v[0], v.size());
```

[Rule] Check the size of an instance of `std::vector` before using the index operator `[ ]`.

To reference an element in a vector, the C++ standard specifies that the method `std::vector::at()` performs bound checking while `std::vector::operator[]` does not. The difference is typically for performance reason. It is up to the developer to choose between safety and speed of code for each access.

Example:

```
int x = 0;
std::vector<int> v(4); // declare a vector with 4 elements
```

```
x = v.at(6);           // throw exception std::out_of_range
x = v[6];             // read an invalid value in uninitialized memory
```

Some implementations of the C++ STL add an assertion on the index value in `std::vector::operator[]`. This means that, in case of invalid index value, the application is aborted when compiled in debug mode (assertions on). In production mode, the assertions are off and the invalid index value is unnoticed.

This is normally fine. You should not use invalid index values anyway and the assertion helps the developer in debug mode. But this is not always the case. There are rare cases where using an invalid index is correct and the application should not fail in debug mode. See the following counter-example.

Counter-example:

```
void legacyFunction(int* buffer, size_t intCount);

std::vector<int> v;
legacyFunction(&v[0], v.size());
```

If the vector `v` is empty, `0` is an invalid index value and simply getting the address of `v[0]` may fail in debug mode. However, this is valid code since `v.size()` is zero and any address value, even an incorrect one, is fine as first parameter for `legacyFunction()` because the function will not access the memory area anyway.



This is the case for the Microsoft Visual C++ library.

Be sure to identify that kind of usage and rewrite the code as follow:

```
legacyFunction(v.empty() ? nullptr : &v[0], v.size());
```

[Rule] In function declarations, when a parameter is not an elementary type, not a pointer type or is a template type, always use a reference. If the contract is to not modify the object, use a `const` qualifier.

Not using a reference forces a copy of the object. This can be costly and this may even not compile in the case of an actual template parameter type without copy constructor.

Good code:

```
void f(const Foo& x);

template <typename T>
void g(const T& x);
```

Bad code:

```
void f(Foo x); // force a copy of a Foo object

template <typename T>
void g(T x);   // call will not compile if
               // for T has no copy
               constructor
```

[Rule] When catching an exception, always use a reference to avoid a copy of the object.

This is an application of the preceding rule to the exception handling.

Example:

```
try {
    ....
}
```

```
catch (const std::exception& e) { // reference to const exception object
    ....
}
```

[Rule] Multiple exception handlers in a `try / catch` structure must be ordered properly.

In a `try / catch` structure with multiple exception handlers, an exception is checked against all `catch` clauses until a matching one is found. Only the first matching one is used. So, when exception classes are derived from each other, the most specific `catch` clauses must come first. Otherwise, they are useless.

Example:

```
class A: public std::exception {...};
class B: public A {...};
```

Good code:

```
try {
    ...
}
catch (const B& e) {
    // catch all B exceptions
}
catch (const A& e) {
    // catch all A exceptions, including
    // subclasses
    // of A, except B which are handled above
}
```

Bad code:

```
try {
    ...
}
catch (const A& e) {
    // catch all A exceptions,
    // including all subclasses of A
}
catch (const B& e) {
    // never used since B is a subclass of A
    // and was already caught in previous
    // handler
}
```

[Rule] Use the *guard design pattern to fail safely*. Define guard classes for resources. Use destructors to fail safely.

It is sometimes necessary to reserve, lock or allocate a resource temporarily. This means that something shall be done at the end of a sequence of statements (release, unlock, deallocate). A problem occurs when an exception is thrown or a return is inadvertently executed in the middle of the sequence. In this case, the resource does not get cleaned up.

The *guard* design pattern is a technique to avoid this. For a given type of resource which needs to be locally reserved, define an associated guard class which has basically only two operations: a constructor which reserves the resource and a destructor which releases the resource (whatever *reserve* and *release* means for the given resource). For each sequence of code which reserves the resource, simply create a block of `{ }` where a local guard object is defined.

Example: This is how the guard pattern is used.

Unsafe code:

Equivalent safe code using a guard class:

```
Resource myResource;
...
myResource.reserve();
...
myResource.release();
```

```
Resource myResource;
...
{
    ResourceGuard(myResource); // implicit
    reserve()
    ...
    // some exception or return here
    ...
} // implicit release() even on exception or
return
```

Example: This is how the guard pattern can be implemented.

```
class Mutex
{
public:
    void acquire() {...}
    void release() {...}
};

class MutexGuard
{
public:
    MutexGuard(Mutex& mutex): _mutex(mutex) {_mutex.acquire();}
    ~MutexGuard() {_mutex.release();}
private:
    Mutex& _mutex;
};
```

And this is how this implementation is used.

```
Mutex globalMutex;
...
{
    MutexGuard guard (globalMutex); // implicit globalMutex.acquire()
    ...
    // some exception or return here
    ...
} // implicit globalMutex.release() even on exception or return
```

Application: Use the standard guard type `std::lock_guard` on `std::mutex` and `std::recursive_mutex`.

[Rule] Use the safe pointer design pattern to prevent memory leaks and complex memory tracking.

When comparing C++ to Java, the main C++ nightmare is the memory management: when should I free memory?

A *safe pointer* (or *smart pointer*, or *shared pointer*) is a design pattern which replaces the usage of plain pointers by a template class which tracks all accesses to a resource. When no more active reference to the object exists, the object is automatically reclaimed.

Great care shall be taken when designing the safe pointer class (especially the multi-threading aspect). But once available, C++ memory management becomes almost as easy as Java.

In early versions of TSDuck, before using C++11, a dedicated template class was designed. Now, TSDuck uses the standard template type `std::shared_ptr` for all complex pointer managements.

Based on `valgrind` results on Linux platforms, TSDuck has proven to be memory-safe when using safe pointers

and zero explicit deallocation on any arbitrary large number of dynamically allocated objects with a limited life-time.

Warning: There are pathological cases where the safe pointer design pattern is ineffective. For instance, when two objects reference each other but are no longer referenced anywhere else, they are lost. They are not automatically freed by the safe pointers because references exist. They should be both freed at the same time. But this situation is typically the symptom of a poorly designed data model. So, the safe pointer design pattern removes the burden of the technical aspects of the memory management but not the requirement for a correct design.

[Recommendation] Take care of reentrancy in reusable components. Use an abstract mutex interface and template reusable components.

If a class is not thread-safe, it is possible to make it thread-safe on the user's request without much effort and without performance penalty when thread-safety is not required.

Solution 1: Compile-time selection of the thread-safety model.

Define two mutex classes with identical services.

```
class NullMutex
{
public:
    // acquire() and release() are
    // empty and inlined as no code
    void acquire() {}
    void release() {}
};
```

```
class RealMutex
{
public:
    // Implement acquire() and release()
    // in the .cpp file
    void acquire();
    void release();
};
```

The second class implements the actual locking features on one or more environments. Now consider that the reusable component to implement is a safe pointer (see preceding rule). Define the safe pointer class as follow:

```
template <typename T, class MUTEX> class SafePointer {...};
```

In the implementation of the SafePointer class, define an object of the generic type **MUTEX** to implement the locking. Whenever you need a thread-safe or non-thread-safe pointer to objects of class Foo, use one of the following:

```
SafePointer <Foo, RealMutex> threadSafePointer;
SafePointer <Foo, NullMutex> nonThreadSafePointer;
```

The first pointer class uses thread-safe code while the second one uses fast code (the inline empty locking services are optimized away by the compiler). There is exactly zero overhead for the non-thread-safe version but the two versions share the same source code.

Solution 2: Run-time selection of the thread-safety model.

Define a general abstract mutex interface which declares the basic synchronization services as pure

```
virtual functions:
class MutexInterface
{
public:
    virtual void acquire() = 0;
    virtual void release() = 0;
    virtual ~MutexInterface() {}
```

```
};
```

Define two subclasses `NullMutex` and `RealMutex`. The first one implements inline empty services which do nothing. The second one implements the actual locking features on one or more environments.

A reusable component can, based on some run-time condition, allocate either a `NullMutex` instance or a `RealMutex` instance. In the non-thread-safe case, the `acquire()` and `release()` operations are simple indirect calls to an empty routine which returns immediately.

Comparison: Performances of the two versions are almost identical in the thread-safe case. The solution 1 offers the best performance in the non-thread-safe case. But it is slightly more constraining: the selection is done at compile-time, the mutex guard class must be a template one and the code footprint is larger when both thread-safe and non-thread-safe usages of the same reusable component are present in the same application.

In TSDuck, we use standard mutex types such as `std::mutex` or `std::recursive_mutex`. To implement the solution 1, TSDuck also defines the class `ts::null_mutex` which declare the standard methods `lock()`, `unlock()`, and `try_lock()` as empty inline methods. An instance of `ts::null_mutex` can be used wherever a `std::mutex` or `std::recursive_mutex` could be used, but doing nothing at zero cost.

[Rule] Do not use C-style casts. Use C++ cast operators `reinterpret_cast`, `dynamic_cast`, `static_cast` and `const_cast`.

The C++ cast operators provide a better control over which type of casting you want. More appropriate checks are performed by the compiler or even at run-time (`dynamic_cast`). The code is also more readable, your intention is more clearly explained to the maintainer.

The most general cast operator is `reinterpret_cast` and is equivalent to a C-style cast. There is no check at all. It is dangerous and most of the time inappropriate. Usually, we need to cast between related types and `static_cast` or `dynamic_cast` is a better option.

Beware of `const_cast` when it is used to remove the "constness" of an object. By doing so, you break the API contract if the `const` object is provided by your caller. Is the `const_cast` required because you invoke another API which is badly defined (an argument is not declared as `const` but not modified anyway)? Or is the `const_cast` required because you will actually modify the object? While the former situation is legitimate, the latter is not.

[Recommendation] Use `dynamic_cast` to check a subclass type.

If a general method working on a superclass needs some specific processing when the object belongs to a given subclass, this is usually the symptom of a bad design. A virtual method should be defined in the super-class.

But sometimes it is not possible to modify the superclass and it is necessary to use specific code. In this case, `dynamic_cast` is the only reliable and safe way to test the subclass of the object.

Example:

```
class C {...};
class C1 : public C {...};
class C2 : public C {...};

void f(const C& c)
{
    // generic processing on super-class view "c" of the object
    c.someMethod();

    // is this object an instance of C2?
    const C2* p2 = dynamic_cast<const C2*>(&c);
    if (p2 != 0) {
        // yes, the object is an instance of C2
        // specific processing on C2 sub-class view "*p2" of the object
        p2->someMethodOfC2();
    }
}
```

```
}
}
```

Note 1: Be aware that `dynamic_cast` is allowed only if the superclass is polymorphic, i.e. if it has at least one virtual method. A virtual destructor is sufficient.

Note 2: Using `dynamic_cast` requires that the C++ compiler supports Run-Time Type Information (RTTI). All modern C++ compilers on desktop and server platforms support RTTI. But, on some constrained embedded platforms, the C++ compiler may not support RTTI and `dynamic_cast` is rejected by the compiler. In that case, you have to determine the actual subclass by some other way and then use `static_cast` after verifying the actual subclass type.

[Rule] Do not use `exit()` or `std::exit()`, except in case of emergency termination.

In C++, `exit()` only guarantees the proper destruction of global objects. The destructors of pending local objects are not invoked. If those destructors have external effects (flushing a cache, closing a file, etc.), the state of some external resources may not be consistent.

To perform an early exit of the program, throw an exception. Use a dedicated application-defined exception and ensure that no function catches this exception (beware of generic `catch (...)` structures). In multi-threaded applications, this is a little bit more complicated; you need to synchronize the termination of all threads.

In the main program, to return an error code to the operating system, do not call `exit()`, perform a return instruction with the appropriate exit code.

[Rule] Do not use `va_start`, `va_arg` and `va_end`, especially on class types.

These macros were designed for C and not C++. They do not properly invoke constructors.

In C++, variable argument lists are handled in a much safer way using `std::initializer_list` and variadic templates.

In TSDuck, the class `ts::ArgMix` and its subclasses are designed to transparently support type-safe heterogeneous variable argument lists. Sample usages can be found in methods `ts::UString::format()` and `ts::UString::scan()`.

[Rule] Never declare a function with an argument being an array of objects when subclasses exist for this class.

This practice has very dangerous side effects on subclasses. This introduces bugs which are very hard to track.

Example:

```
class A
{
public:
    int a;
    A(): a(1) {}
};

class B: public A
{
public:
    int b;
    B(): b(2) {}
};

void f(A array[], size_t elemCount)
{
    // update second element of array
    if (elemCount >= 2) {
        array[1].a = 47;
    }
}
```

```

    }
}

B x[2];
f(x, 2); // do you think that x[1].a is updated ?

std::cout << x[0].a << ", " << x[0].b << ", "
          << x[1].a << ", " << x[1].b << std::endl;

```

The last instruction prints "1, 47, 1, 2". The function `f()` has corrupted the subclass fields of `x[0]`.

The same buggy effect is obtained when the function is defined using a pointer instead of an array:

```
void f(A* array, size_t elemCount)
```

Alternatives: If you really need that function, you must overload it for all possible subclasses of A. And if new subclasses of A are created in the future, you must remember to overload the function for all new subclasses. Since this is a maintenance challenge, this kind of function should really be avoided anyway.

[Recommendation] Avoid using arrays of C++ objects. Use standard containers from the STL instead.

The preceding rule is a good illustration of the danger of arrays of objects in C++.

There is almost no case where using an array is better than using the standard container `std::vector`, neither in performance nor in functionalities.

[Rule] Never assign objects through dereferencing a pointer to a base class.

The assignment operator cannot be virtual (at least in its strict definition). By using the assignment through dereferencing a pointer to a base class, the actual assignment operator is the one from the superclass. The assignment of the object will be partial (*slicing effect*) and the object can be left in an inconsistent state.

Example:

```

class Fruit
{
public:
    Fruit& operator=(const Fruit&);
    ...
};

class Apple: public Fruit
{
public:
    Apple& operator=(const Apple&);
    ...
};

class Tomato: public Fruit
{
public:
    Tomato& operator=(const Tomato&);
    ...
};

Apple apple;
Tomato tomato;

Fruit* p1 = &apple;
Fruit* p2 = &tomato;

```

```
*p1 = *p2; // Legal but WRONG !
```

The above assignment is legal and compiles. However, since the assignment operator cannot be virtual, it invokes the assignment operator of the superclass, i.e. the method `Fruit::operator=(const Fruit&)`. In practice, the common `Fruit` fields of an `Apple` object will be assigned with the common `Fruit` fields of a `Tomato` object. The specialized `Apple` fields of the object are left unmodified and, thus, possibly inconsistent with the new values of the common `Fruit` fields.

Honestly, this is weird to assign a tomato into an apple, even if both are fruits!

We reach here the limitation of the virtual vs. non-virtual method model. This model works fine as long as the virtual / non-virtual method works on one object only (this object). But when the method is designed to work on two objects globally, as it is the case for the assignment, there is no good solution.

6.4.2. C++ coding conventions

This section is present to fulfil the required separation of immutable *rules* and *recommendations* from potentially replaceable *conventions*, as explained in [section 6.2](#).

6.4.2.1. Source code formatting

[Convention] *The file name extensions by file type are `.h` for C++ header files and `.cpp` for C++ source files.*

For C++ files, several conventions exist: `.h`, `.hpp`, `.hxx`, `.H` for headers, `.cpp`, `.cxx`, `.C` for source files. The selected convention has a large adoption and is good enough.

[Convention] *Indentation: use 4 space characters without tabulation.*

Using less than 4 spaces is not clear enough. Using 8 spaces moves too fast to the right.

6.4.2.2. Modularity

[Convention] *When a module is specialized in the management of one data type (object-oriented design), the base name of the module files is the data type name, using the same lower/upper case letters.*

Example: The class `ts::UString` is declared in file `tsUString.h` and defined in file `tsUString.cpp`.

[Convention] *The file name of a module containing a C++ class is built from the concatenation of all nested namespaces and the class name.*

This way, the class declaration can be easily found.

Example: The class named `ts::xml::Element` is declared in file `tsxmlElement.h` and defined in file `tsxmlElement.cpp`.

[Convention] *Non-inline template methods or methods of template classes are grouped at the end of the header file, in a clearly identified section.*

The portability of the C++ templates is a challenge. Most compilers require the definition of the template methods to be available during the compilation of the modules which use them. So, we need to have them in the header file.

6.4.2.3. Naming conventions

There are many naming conventions in the C and C++ communities. This is sometimes both the most sensitive part for the developers and the less important for the quality of the code. Discussing which naming convention is

the best one is a waste of time. There is no best one. But many are good enough. The quality of the code only depends on the strict and consistent application of one single good enough naming convention.

This section describes the naming conventions for the C++ language in the TSDuck project. Simply use them.

[Convention] All entities which are defined within the project shall be declared within the namespace named `ts`.

This is a trade-off between readability and usability. Using a more significant but longer namespace such as `tsduck` would appear as painful to developers who could be tempted to disobey the "no using namespace directive" rule.

[Rule] Multiple nested namespaces may be defined within the namespace `ts`.

Typically, there is one inner namespace per subproject or logical group of classes.

Developers are encouraged to use short but meaningful namespaces (4 or 5 characters maximum). Using acronyms is acceptable.

Typical nested namespaces are `ts::xml`, `ts::json`, `ts::tlv`.

[Convention] Namespaces are composed of lowercase letters or digits only.

No capital letter, no underscore.

As mentioned in the preceding rules, all entities are in the namespace `ts` or in one of its inner namespaces.

[Convention] Preprocessing macro names (`#define`) are composed of uppercase letters and digits only. Words are separated by underscores. All macro names start with the prefix `TS_`.

Example:

```
#define TS_FOO_VERSION 47
#define TS_INCR(x) (...)
```

[Convention] Use a verb as the central component of the name for functions that do something. Use an imperative form for the name of functions returning a boolean value. Such functions typically start with `is` or `has`, depending on what they return.

Example:

```
namespace ts {
    class Foo
    {
    public:
        Foo(...);
        ~Foo();
        void load(...);
        void save(...) const;
        bool isEmpty() const;
        bool hasChildren() const;
    private:
        // ... internal fields
    };
}
```

[Convention] Class names, type names, public static functions, non-member (global) functions use a mixed case convention, without underscore, starting with an uppercase letter.

See examples below.

[Convention] Non-static public member fields and functions use a mixed case convention, without underscore, starting with a lowercase letter.

When reading code, it is useful to differentiate static and non-static members. Static members are class-wide; their name starts with an uppercase letter. Non-static members applies to one instance; their name starts with an lowercase letter.

See examples below.

[Convention] Public member fields and local variables (in functions) use either a mixed case convention, without underscore, starting with a lowercase letter, or an all-lowercase convention with underscores.

The dual convention is unfortunate in TSDuck. This is the result of history. However, as mentioned in [section 6.3.1.7](#), "do not rewrite existing code for the sole purpose of applying coding guidelines".

[Convention] Public constants use uppercase letters with underscores to separate words.

See examples below.

[Convention] Private entities of any sort use the same convention as their public counterpart but start with an underscore.

When reading code, it is important to understand the scope of an action, public or private. Modifying a public field may break the interface contract of the class while modifying a private field only requires a local analysis. Private fields are identified by their leading underscore.

Example:

```
namespace ts {
    void GlobalFunction();           // global function, not in a class
    typedef Foo* FooPtr;             // type name
    class ClassName                  // class name
    {
    public:
        static const size_t MAX_SIZE = ...; // constant
        int somePublicField;           // public member field
        void someMemberFunction();     // public member function
        static void SomeStaticFunction(); // public static function
    private:
        int _somePrivateField;         // private member field
        int _some_private_field;       // also acceptable
        void _somePrivateMemberFunction(); // private member function
        static void _SomePrivateStaticFunc(); // private static function
    };
}

void ts::ClassName::someMemberFunction()
{
    uint32_t messageIndex;           // local variable
    uint32_t message_index;         // also acceptable
}
```

[Convention] In all mixed case conventions, acronyms are all-uppercase, or all-lowercase at the beginning of the name.

Example:

```
typedef uint16_t TCPOrUDPPort;
TCPOrUDPPort tcpDefaultPort = 80;
```

```
static TCPOrUDPPort _udpDefaultPort = 80;
```

[Convention] Values in enum types are composed of uppercase letters and digits only. Words are separated by underscores. In pure enum types (ie. not enum classes), a common prefix which is derived from the type name is prepended to all values.

Example:

```
namespace ts {  
    enum Counter {  
        COUNTER_ONE,  
        COUNTER_TWO  
    };  
}
```

[Convention] Conditional compilation of debug code is allowed using the symbol `DEBUG`. When `DEBUG` is undefined, no debug code shall be compiled.

This symbol shall not be defined in any header file. It shall be defined by the compilation environment (makefile, IDE compilation profile, etc.)

As an exception, the macro `DEBUG` does not start with the required `TS_` prefix because `DEBUG` is a common symbol which is recognized by many libraries to trigger debug-specific code.

6.4.2.4. Syntax formatting conventions

Just like naming conventions, the syntax formatting conventions are subject to discussion. And, similarly, there is no best one. We must simply adopt one and use it consistently in all code. Most IDE's can be configured to automatically enforce these settings when typing code. This is possible with Emacs, Eclipse, Microsoft Visual Studio or Qt Creator.

[Convention] Add a space character in the following locations: before and after binary or ternary operators, before a group of one or more `(` in an expression, after a group of one or more `)` or `]` in an expression, after a `,`.

Exception: no space before `,` or `;`.

Example:

```
a = (b + c) * ((e / 4) % 3);  
x = y > size ? y : size;  
p = f(a, b, c[i] + 1);
```

[Convention] Do not use any space character before `(` in a function declaration, definition or call. Similarly, do not use any space character before `[` in an array declaration or reference.

Example:

```
void foo(char* name, size_t size);  
char buffer[200];  
buffer[0] = 'a';  
foo(buffer, sizeof(buffer));
```

[Convention] Comment lines must be aligned on the code they refer to.

Good code:

```
void f(int x)
{
    // about the test
    if (x > 1) {
        // about the zero
        x = 0;
    }
}
```

Bad code:

```
void f(int x)
{
    // about the test
    if (x > 1) {
        // about the zero
        x = 0;
    }
}
```

[Convention] Multi-lines comments shall be formatted using `//` on each line. All lines which belong to the same logical comment shall be aligned.

A previous rule explains why single-line C++-style comments are safer than multi-line C-style comments.

Example:

```
// This is a comment
// on several lines.
```

Exception: For self-documentation (Doxygen for instance), use the required conventions for the corresponding documentation extraction tool.

[Convention] The function definition has its initial `{` on a new line. The parameters are listed on one line, if this makes the line not too long.

Example:

```
void FunctionCode(const char* name, size_t size)
{
    ...
}
```

[Convention] In a function declaration, definition or call, when the parameter list is too long to fit on one line, the parameters are aligned on the opening parenthesis and there must be exactly one parameter per line.

Example:

```
MyFunction("abcd", size, 47); // short line

MyFunction(tsCallingSomethingVeryVeryLongAndUnreadable(a + 45 * x, "title"),
           size,
           47); // long statement: one parameter per line
```

In the first example, using one parameter per line would be counter-productive. The code would be bloated and less readable. So, if everything fits on one line, use one line.

In the second example above, putting the two parameters `size`, `47` on the same line could make the careless reader think that there are only two parameters to the function.

So, either put all parameters on one line or one parameter per line, but not a mixture of the two.

[Convention] The conditional and loop statements have their initial `{` on the same line. The closing `}` is

on its own line. The `else if` construction is on one line.

Example:

```
if (x == a) {
    ...
}
else if (y < b) {
    ...
}
else {
    ...
}
```

[Convention] The `switch` statement has its initial `{` on the same line. The various `case` entries and the optional default entry are on individual lines. They are indented from the `switch`.

Example:

```
switch (state) {
    case TS_START:
        print("started");
        break;
    case TS_CONTINUE:
    case TS_END:
        print("ok");
        break;
    default:
        print("error");
        break;
}
```

[Convention] The namespace declaration has its initial `{` on the same line. The closing `}` is on its own line. The content of the namespace is indented.

Example:

```
namespace ts {
    namespace proj {
        void foo();
    }
}
```

[Convention] The `class` declaration has its initial `{` on a new line. The `public`, `protected` and `private` keywords are aligned on the `{`.

Example:

```
class C
{
public:
    void init();
private:
    void _reset();
};
```

Remember that the closing `}` of a class declaration must be followed by a `;`. This is a typical C++ oddity and a common source of compilation syntax errors.

[Convention] *In template class declarations and definitions, the template part is on a separate line with the same alignment as the declaration or definition.*

Example:

```
template <typename T, class MUX>
class SafePointer
{
    ...
};

template <typename T, class MUX>
ts::SafePointer<T,MUX>& operator=(T* p)
{
    ...
}
```

[Convention] *In the definition of a constructor, the field initializers are indented and separated one per line.*

Example:

```
class C
{
public:
    C(int x);
private:
    int _a;
    int _b;
    int _c;
};

C::C(int x):
    _a(0),
    _b(x),
    _c(x + 1) // be careful, no comma on last field initializer
{
    ...
}
```

[Convention] *If a sequence of stream output operators `<<` is too long to fit on one line, the `<<` operators are aligned on multiple lines.*

Example:

```
std::cout << "title: " << title
          << ", message: " << message
          << std::endl;
```

6.4.2.5. Doxygen self-documentation

[Convention] *All Doxygen commands in embedded comments in source files shall be introduced with the `@` character (e.g. `@file`, `@param`, `@return`, etc.)*

There are two syntaxes for Doxygen commands, starting with a `@` or with a `\`. For consistency, all developers shall use one single common convention. The syntax using `@` is preferred since it is compatible with Javadoc.

[Convention] In C++ source files, the Doxygen comment blocks shall be formatted using `///` on each line. All lines which belong to the same logical documentation block shall be aligned.

Example:

```
///  
/// ... Doxygen documentation commands  
///
```

There are two main syntaxes for Doxygen comment blocks in C++, one using the C-style comment `/**` and one using the C++-style comment `///`. For consistency, all developers shall use one single common convention.

A previous rule explains why single-line C++-style comments are safer than multi-line C-style comments.

[Rule] When an entity is in a Doxygen-documented source file, all its components shall be documented without exception.

The brief description of the entity (`@brief`) is mandatory. The detailed description (`@details`) is optional if the entity is so simple that the brief description is sufficient.

Functions and methods shall document all their parameters and their returned value (if any).

Enumeration types shall document all their values individually in addition to the documentation of the type.

Hint 1: The mandatory documentation of all fields, functions, methods, and their parameters can be enforced using the following directive in the `Doxyfile`:

```
WARN_NO_PARAMDOC = YES
```

Hint 2: The `@brief` and `@details` commands do not need to be explicitly present. The first item is implicitly the brief description and the detailed description is implicitly everything that follows a blank line after the brief description. The implicit notation is more readable in the source code for developers and produces the same documentation.

Example: The following two Doxygen blocks are equivalent.

```
///  
/// @brief This is the brief description of the next item.  
/// @details This is the long and detailed description.  
/// The Doxygen commands can be omitted if a blank line  
/// separates the brief and detailed description.  
///  
  
///  
/// This is the brief description of the next item.  
/// This is the long and detailed description.  
/// The Doxygen commands can be omitted if a blank line  
/// separates the brief and detailed description.  
///
```

This behavior is allowed when the `Doxyfile` contains the following directive:

```
JAVADOC_AUTOBRIEF = YES
```

Hint 3: Individual parameters or values can be documented on one line after the element itself if the documentation contains only a brief description.

Example: The following two notations are equivalent. Note that the first one is more compact and probably more readable.

```
//!
//! Flags for the @c Hexa family of functions
//!
enum HexaFlags {
    HEX_HEXA = 0x0001, //!< Dump hexa values
    HEX_ASCII = 0x0002, //!< Dump ascii values
    ....
};

//!
//! Flags for the @c Hexa family of functions
//!
enum HexaFlags {
    //!<
    //!< Dump hexa values
    //!<
    HEX_HEXA = 0x0001,
    //!<
    //!< Dump ascii values
    //!<
    HEX_ASCII = 0x0002,
    ....
};
```

Appendix A: Licenses

A.1. TSDuck license

TSDuck is released under the terms of the license which is commonly referred to as "BSD 2-Clause License" or "Simplified BSD License" or "FreeBSD License". See <http://opensource.org/licenses/BSD-2-Clause>.

Copyright © 2005-2026, Thierry Lelégard

All rights reserved.

Redistribution and use in source and binary forms, with or without modification, are permitted provided that the following conditions are met:

1. Redistributions of source code must retain the above copyright notice, this list of conditions and the following disclaimer.
2. Redistributions in binary form must reproduce the above copyright notice, this list of conditions and the following disclaimer in the documentation and/or other materials provided with the distribution.

THIS SOFTWARE IS PROVIDED BY THE COPYRIGHT HOLDERS AND CONTRIBUTORS "AS IS" AND ANY EXPRESS OR IMPLIED WARRANTIES, INCLUDING, BUT NOT LIMITED TO, THE IMPLIED WARRANTIES OF MERCHANTABILITY AND FITNESS FOR A PARTICULAR PURPOSE ARE DISCLAIMED. IN NO EVENT SHALL THE COPYRIGHT OWNER OR CONTRIBUTORS BE LIABLE FOR ANY DIRECT, INDIRECT, INCIDENTAL, SPECIAL, EXEMPLARY, OR CONSEQUENTIAL DAMAGES (INCLUDING, BUT NOT LIMITED TO, PROCUREMENT OF SUBSTITUTE GOODS OR SERVICES; LOSS OF USE, DATA, OR PROFITS; OR BUSINESS INTERRUPTION) HOWEVER CAUSED AND ON ANY THEORY OF LIABILITY, WHETHER IN CONTRACT, STRICT LIABILITY, OR TORT (INCLUDING NEGLIGENCE OR OTHERWISE) ARISING IN ANY WAY OUT OF THE USE OF THIS SOFTWARE, EVEN IF ADVISED OF THE POSSIBILITY OF SUCH DAMAGE.

A.2. Third-party libraries

TSDuck includes a few third-party libraries, either in source form, binary form or both. For more details about the licenses of these third-party libraries, see the file named `OTHERS.txt` in the TSDuck source code repository.

DTAPI: On Linux and Windows, the TSDuck binary distributions contain the DTAPI library in static form. This library is part of the [Dektec SDK](#). This software is available in binary format only and is distributed under the BSD 2-Clause License. *"Copyright (c) 2017 by Dektec Digital Video B.V."*

LIBSRT: On Windows, the TSDuck binary distribution contains the SRT library in static form. There are two versions of the SRT library, one from [Haivision](#), one from [Robotweax](#). TSDuck can be built with any of them.

Haivision LIBSRT: The [Haivision SRT library](#) is open-source and distributed under the Mozilla Public License v2.0. *"Copyright (c) 2018 Haivision Systems Inc."*

Robotweax LIBSRT: The [Robotweax SRT library](#) is open-source and distributed under the MIT License. *"Copyright (c) 2026 Robotweax GmbH and contributors"*

LIBRIST: On Windows, the TSDuck binary distribution contains the [RIST library](#) in static form. This is an open-source library which is distributed under the BSD 2-Clause License. *"Copyright © 2019-2020 SipRadius LLC. All right reserved."*

LibVatek: On Windows and Linux, the TSDuck binary distribution contains the [LibVatek library](#) in static form. This is an open-source library which is distributed under the BSD 2-Clause License. *"Copyright © 2022, Richie Chang."*

Small Deflate (sdefl): On Windows, the TSDuck binary distribution contains the Small Deflate library in static form. This is an open-source library which is distributed under two licenses, MIT License and Public Domain License. *"Copyright (c) 2020-2023 Micha Mettke."*

Appendix B: References

B.1. Acronyms and abbreviations

ABI	Application Binary Interface
API	Application Programming Interface
BOM	Byte Order Mark (a feature of UTF-16 and UTF-32)
CI/CD	Continuous Integration/Continuous Deployment
DRY	Don't Repeat Yourself (a good practice)
FOSS	Free and Open Source Software
GCC	GNU Compiler Collection
GPL	GNU General Public License
IDE	Integrated Development Environment (e.g. Eclipse, MS Visual Studio)
LGPL	Lesser GPL
MSC	Microsoft C/C++ Compiler
MSVC	Microsoft Visual C/C++
NIH	Not Invented Here (a bad practice)
NTTP	Non-Type Template Parameter (C++)
OOD	Object-Oriented Design
OOP	Object-Oriented Programming
RAII	Resource Acquisition Is Initialization (C++)
RTFM	The Most Important Acronym For Developers (yes, it is)
RTTI	Run-Time Type Information (C++)
SDK	Software Development Kit
SFINAE	Substitution Failure Is Not An Error (C++)
STL	Standard Template Library (C++)
TDD	Test-Driven Development
UTF	Unicode Transformation Format (UTF-8, UTF-16, UTF-32, etc.)
XP	eXtreme Programming

Bibliography

- [BSD-2C] BSD 2-Clause License, <http://opensource.org/licenses/BSD-2-Clause>
- [CERT-C] "SEI CERT C Coding Standard, Rules for Developing Safe, Reliable and Secure Systems", Software Engineering Institute, Carnegie Mellon University, 2016 Edition.
- [CERT-CPP] "SEI CERT C Coding Standard, Rules for Developing Safe, Reliable and Secure Systems in C", Aaron Ballman, Software Engineering Institute, Carnegie Mellon University, 2016 Edition.
- [CPPREF] Online reference of the C++ language and standard library, <http://en.cppreference.com/>
- [GAMMA] "Design Patterns, Elements of Reusable Object-Oriented Software", Erich Gamma, Richard Helm,

Ralph Johnson, John Vlissides, Addison-Wesley, 1995.

- [GITHUB-CI] GitHub Actions documentation, <https://docs.github.com/en/actions>
- [Haivision-SRT] Haivision original SRT library, <https://github.com/Haivision/srt/>
- [ISO-14882] ISO/IEC 14882, "Programming Languages - C", 1998 (the C98 standard).
- [ISO-14882-11] ISO/IEC 14882:2011, "Programming Languages - C", 2011 (the C11 standard).
- [ISO-14882-14] ISO/IEC 14882:2014, "Programming Languages - C", 2014 (the C14 standard).
- [ISO-14882-17] ISO/IEC 14882:2017, "Programming Languages - C", 2017 (the C17 standard).
- [JOSUTTIS] "The C++ Standard Library, A Tutorial and Reference", Nicolai M. Josuttis, Addison-Wesley, 1999.
- [MEYERS-EFF] "Effective C++, Third Edition, 55 Specific Ways to Improve Your Programs and Designs", Scott Meyers, Addison-Wesley, 2005.
- [MEYERS-MORE] "More Effective C++, 35 New Ways to Improve Your Programs and Designs", Scott Meyers, Addison-Wesley, 2008.
- [MEYERS-STL] "Effective STL, 50 Specific Ways to Improve Your Use of the Standard Template Library", Scott Meyers, Addison-Wesley, 2001.
- [Robotweax-SRT] Robotweax alternative SRT library, <https://github.com/Robotweax/srt/>
- [STROUSTRUP] "The C++ Programming Language, Special Edition", Bjarne Stroustrup, Addison-Wesley, 2000.
- [TSDuck] TSDuck Web site, <https://tsduck.io/>
- [TSDuck-Issues] TSDuck issues tracker and discussion forum, <https://github.com/tsduck/tsduck/issues>
- [TSDuck-Source] TSDuck source code repository, <https://github.com/tsduck/tsduck/>