

Document Title	Design guidelines for using parallel processing technologies on Adaptive Platform
Document Owner	AUTOSAR
Document Responsibility	AUTOSAR
Document Identification No	884

Document Status	published
Part of AUTOSAR Standard	Adaptive Platform
Part of Standard Release	R25-11

Document Change History			
Date	Release	Changed by	Description
2025-11-27	R25-11	AUTOSAR Release Management	• No content changes
2024-11-27	R24-11	AUTOSAR Release Management	• Minor narrative text updates
2023-11-23	R23-11	AUTOSAR Release Management	• Removal of Deterministic Execution
2022-11-24	R22-11	AUTOSAR Release Management	• No content changes
2021-11-25	R21-11	AUTOSAR Release Management	• No content changes (only converted to LaTeX)
2020-11-30	R20-11	AUTOSAR Release Management	• No content changes
2019-11-28	R19-11	AUTOSAR Release Management	• No content changes • Changed Document Status from Final to published
2019-03-29	R19-03	AUTOSAR Release Management	• Minor changes





2018-10-31	R18-10	AUTOSAR Release Management	• Minor changes
2018-03-29	R18-03	AUTOSAR Release Management	• Minor changes
2017-10-27	R17-10	AUTOSAR Release Management	• Initial release

Disclaimer

This work (specification and/or software implementation) and the material contained in it, as released by AUTOSAR, is for the purpose of information only. AUTOSAR and the companies that have contributed to it shall not be liable for any use of the work.

The material contained in this work is protected by copyright and other types of intellectual property rights. The commercial exploitation of the material contained in this work requires a license to such intellectual property rights.

This work may be utilized or reproduced without any modification, in any form or by any means, for informational purposes only. For any other purpose, no part of the work may be utilized or reproduced, in any form or by any means, without permission in writing from the publisher.

The work has been developed for automotive applications only. It has neither been developed, nor tested for non-automotive applications.

The word AUTOSAR and the AUTOSAR logo are registered trademarks.

Table of Contents

1	Introduction	5
1.1	Contents	5
1.2	Prereads	5
1.3	Relationship to other AUTOSAR specifications	5
2	Definition of terms and acronyms	6
2.1	Acronyms and abbreviations	6
3	Related Documentation	7
4	Scope	8
4.1	Definition of parallel processing "technologies"	8
4.2	Audience	8
5	Architectural design	9
5.1	Background	9
5.1.1	Evolving parallel processing technologies	9
5.1.2	Distributed, concurrent, and parallel	9
5.1.3	TLP/DLP/PLP	10
5.2	Service-based parallel processing	11
5.2.1	Layered architectural view	11
5.2.2	Accelerator-model	13
5.2.3	CPU/co-processor-model	13
5.3	Rationale: decoupling of parallel processing specific knowledge from application development	14
5.4	Adaptive Platform methodology consideration	14
6	Non-functional design topics	15
6.1	Performance	15
6.1.1	Interface granularity and communication overhead	15
6.1.2	Data handling and throughput balancing	15
6.2	Safety considerations	16
6.3	Prospects	16
6.3.1	Adaptive Platform standard application services	16
6.3.2	More parallel processing within Adaptive Application	16

1 Introduction

1.1 Contents

This document specifies the guidelines for using parallel processing technologies on the Adaptive Platform, or Parallel Processing Guidelines, in short.

The purpose of this document is to provide design guidelines for using parallel processing technologies on the Adaptive Platform. The focus is on software, especially the application layer including the services. General hardware discussions are also included to build the base for software.

1.2 Prereads

This document is one of the high-level conceptual documents of AUTOSAR.

Useful pre-reads are [\[1\]](#) [\[2\]](#) [\[3\]](#).

1.3 Relationship to other AUTOSAR specifications

Refer to Contents and Prereads.

2 Definition of terms and acronyms

Acronyms and abbreviations which have a local scope and therefore are not contained in the AUTOSAR glossary [1] must be defined in this chapter.

2.1 Acronyms and abbreviations

Abbreviation	Description
ADAS	Autonomous Driving and Assistance System
ARA	AUTOSAR Run-time for Applications
SOA	Service-Oriented Architecture

Term	Description
Machine	see [1] AUTOSAR Glossary

Table 2.1: Glossary-defined Technical Terms

3 Related Documentation

- [1] Glossary
AUTOSAR_FO_TR_Glossary
- [2] Methodology for Adaptive Platform
AUTOSAR_AP_TR_Methodology
- [3] Explanation of Adaptive Platform Design
AUTOSAR_AP_EXP_PlatformDesign
- [4] Explanation of Safe API for hardware accelerators
AUTOSAR_AP_EXP_SafeHardwareAccelerationAPI

4 Scope

4.1 Definition of parallel processing "technologies"

In this document, the meaning of parallel processing technologies is loosely defined. This is so on purpose, with hopes to provide design principles for parallel and related processing (see Distributed, concurrent, and parallel).

The term "parallel processing technologies" in this document, therefore, covers both hardware and software. In term of hardware, multicore, manycore, DFP (Data-Flow Processor), GPU (Graphical Processing Unit), FPGA (Field Programmable Gate Array), or alike; in terms of software, multi-thread programming, pragma based techniques like OpenMP¹, various template programming such as TBB², accelerator programming language like OpenCL³, and even various message passing APIs like MPI⁴ that are not by themselves parallel processing technologies but are tightly related to. The technologies also include various tooling that assists in designing and implementing the parallel processing technologies into an Adaptive Platform based system.

It is not a purpose of this document to list all the existing parallel processing technologies, to explain what they are, nor to guide how to use the technologies themselves. Nevertheless, the document may contain some references to the technologies as minimum as deemed necessary to describe the design guidelines.

4.2 Audience

This specification is for multiple domains of Adaptive Platform related designers and developers, namely the system designer who decides hardware/software partitioning, hardware designer who design and/or select computing hardware resources, a software designer who design overall software system architecture, Adaptive Platform developers, and developers of Adaptive Platform services running on ARA.

Adaptive Application developers, on the other hand, who may not directly use parallel processing technologies and only design sequential, single-threaded application, may find this irrelevant, if his/her software architect follows the architectural design guideline described in this document. However, nowadays it is becoming difficult to write an application without some form of multi-thread programming, and it is likely to be more so in future, so essentially everyone concerned with Adaptive Platform is advised to reference this document.

¹<http://www.openmp.org/>

²<https://www.threadingbuildingblocks.org/>

³<https://jp.khronos.org/opencl>

⁴<http://www.mcs.anl.gov/research/projects/mpi/>

5 Architectural design

5.1 Background

5.1.1 Evolving parallel processing technologies

The parallel processing technologies are still rapidly evolving, both in hardware and software. In hardware, GPGPU (General Purpose GPU) is one of them but never the only one - various manycore processors, data flow processors, FPGA, and some dedicated accelerators are emerging, and it seems there are more to come, including the evolutions of these existing technologies.

The picture looks similar in software. Starting with the threading library offered by POSIX and C++ Standard libraries that the Adaptive Platform supports, and other threading libraries such as TBB, MTAPI¹, compiler directives based threading like OpenMP, accelerator programming language like OpenCL and CUDA® (proprietary), HLS² compiler based FPGA programming, and various parallelization compilers/tools, such as graph or process network-based tools, which generally use threading underneath but technologically not limited to, and there are even model based parallelization tools that can take a Simulink® model as input. Also, there are various message passing APIs, like MPI, that often work along with these technologies. There are higher-level libraries such as OpenCV³, OpenVX⁴ - though these are not by themselves parallel processing libraries, they generally use parallel processing technologies underneath to accelerate the processing. At last, similar in a sense that they are higher-level, there are C++ AMP⁵ and SYCL⁶. To further complicate the matter, OpenMP 4 now supports accelerators. And this is not a complete list.

5.1.2 Distributed, concurrent, and parallel

In most cases, AUTOSAR systems are distributed systems. A distributed system is concurrent, meaning multiple tasks are running at the same time. Each subsystem in the distributed system has some sort of processing elements, typically CPUs (but not necessarily) - therefore at the whole this is a multi-processor system, capable of both concurrent and parallel computing. Note that if the Adaptive Platform runs on a single-core processor *Machine* without any other computing elements, parallel processing is not possible, although concurrency still is, as OS provides the threading mechanism to switch processing (threads) triggered by some event.

¹The Multicore Association <http://www.multicore-association.org/workgroup/mtapi.php>

²High Level Synthesis

³<http://opencv.org/>

⁴<https://www.khronos.org/openvx/>

⁵<http://download.microsoft.com/download/4/0/E/40EA02D8-23A7-4BD2-AD3A-0BFFFB640F28/CppAMPLanguageAndProgrammingModel.pdf>

⁶<https://www.khronos.org/sycl>

Parallel processing may occur at different computing layers, from bit-level, instruction-level, thread-level, (and/or) task-level. The definition of "task-level" differs among computing models and methodologies. In AUTOSAR Adaptive Platform software point of view, however, parallel processing strictly applies either on thread-level, process level, or on [Machine](#) (platform) level. It is also noteworthy to recognize that a process, in the context of the Adaptive Platform that is based on POSIX multi-process OS, is just a container for threads and not an execution entity like a thread by itself. The container provides an enclosure for a certain unique set of resources, which include some logical memory accessible by the process. This is also the same for the machines. It is always the thread-concurrency and parallelism (if more than two processors are available) at the software level that is directly executed on top of the AUTOSAR Adaptive Platform OS.

There may be other processing elements that are either incapable of directly executing Adaptive Platform but offer some useful computing. GPU, FPGA, DFP (Data Flow Processor), and manycore processor, are representative examples, although some of them can execute some executive or OS, and even Adaptive Platform itself, fully or partially. If they are incapable of running Adaptive Platform at least partially, then the parallel processing capability can only be accessed by some kind of specific interfaces from a thread running on Adaptive Platform, regardless of the mechanism behind the interface. They are still programmable in one way or another - but just not in the way ARA and C++ bindings that the Adaptive Platform defines.

The important architectural design consideration here is that parallel processing at large is a system level topic. Distributed, concurrent, and parallel processing are highly interrelated. One example is that a well-designed multi-threaded program may run concurrently on a single-core processor, or in parallel on a multi-core processor, or even distributed over two machines provided it uses some processor/[Machine](#) transparent thread communication.

5.1.3 TLP/DLP/PLP

In general, there are three types of parallel processing: Task Level Parallelism (TLP), Data Level Parallelism (DLP), and Pipeline Level Parallelism (PLP). The TLP refers performing multiple tasks at the same time as they are (mostly) independent and (mostly) do not depend on each other. DLP refers to performing the same calculation with multiple, (mostly) non-interdependent sets of (large) data. The same calculation is multiplexed with the different set of data. PLP refers to executing multiple inter-dependent tasks in a pipeline fashion. Each task is assigned to a pipeline stage according to the data dependency of the task input/output.

The three types of parallelism exist in multiple layers of the system. There can be system level parallelism, like two Adaptive Platform machines may have TLP or even PLP. Another example may be that multiple-camera-based 360-degree real-time object recognition may be realized by multiple Adaptive Platform machines performing DLP against a large data set of (virtual) vehicle surrounding video image. The point here is that it is critical for a vehicle system designer to understand the system overall data

flow and processing loads and allocate Adaptive Platform machines accordingly. This can be termed as an Adaptive Platform [Machine](#) level parallelism.

The next physical level below is OS-thread level parallelism. The three types of parallelism described above can be implemented using OS threads.

Yet another physical level below is the instruction level parallelism. This is generally in the field of processor and compiler technologies. For example, a VLIW⁷ processor architecture has multiple execution units that allow concurrent execution of multiple instruction streams, in either TLP or DLP fashion. A SIMD (Single Instruction Multiple Data) co-processor instruction extension enables DLP at the instruction level. A GPGPU, in general, is a form of instruction level DLP in "[5.2.2 Accelerator-model](#)". The SIMD extension, on the other hand, is DLP in "[5.2.3 CPU/co-processor-model](#)". A manycore processor, including most of DFP (Data Flow Processor), offer in general MIMD (Multiple Instruction Multiple Data) instruction level parallelism. Since it is not the same single instruction as in the case of SIMD, MIMD can be used to implement all three forms of parallelism, namely TLP, DLP, and PLP.

Also, regardless of the physical levels, namely Adaptive Platform [Machine](#) level, OS-thread level, or Instruction level, the TLP, DLP, and PLP are not always used independently. For example, for the multi-stage processing of large data, the combination of DLP and PLP are popular.

5.2 Service-based parallel processing

With the background provided in [5.1](#), the key concept of this guideline is to utilize the SOA of Adaptive Platform. That is, to push the use of parallel processing technologies underneath non-platform services, leaving the [Adaptive Application](#) free from the various parallel processing technologies used. The service 'implementation', on the other hand, will be specific to the choice of parallel processing technologies used. We call this model of parallel processing as "[5.2 Service-based parallel processing](#)".

This model allows the maximum reuse of [Adaptive Application](#) that requires high-performance computing in realizing its functionality. The heavy lifting part will be separated into non-platform services, and the implementation of services are free to utilize the full capability of the parallel processing technologies of choice, provided they conform to the safety/security requirements of the project.

5.2.1 Layered architectural view

Figure [5.1](#) illustrates the overall architecture of the service-based parallel processing. The example is based on some ADAS domain application, but it is not the intention to limit the domain in any way.

⁷Very Long Instruction Word

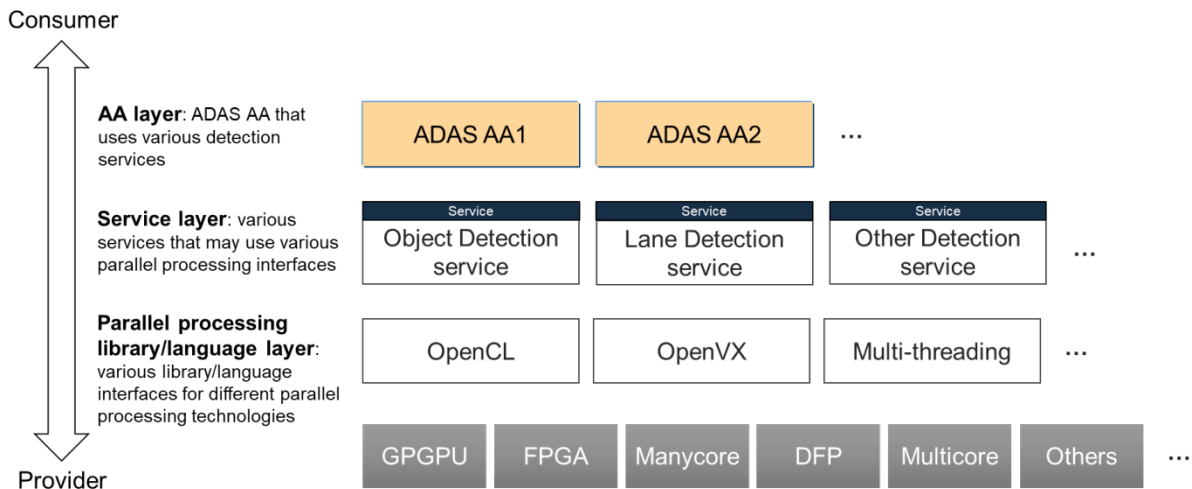


Figure 5.1: Parallel processing consumer-provider layered view example

The overall picture shows the producers and consumers of various parallel processing services, in a top-down layered fashion.

The Adaptive Application layer is the Adaptive Application that uses various services. The Adaptive Application does (or should) not know which of the services it is using uses parallel processing underneath.

The Service layer, in the context of this guideline, consists of the services which use parallel processing technologies. There are non-platform services that use `ara::com`. They provide C++ interface library generated from the service definition, which is used by the Adaptive Application. Note that these services may very well use other services internally. One example is that one may design a pre-processing or low-level sensing service, and a meta-data provider service that uses the output of the former service. Another example may be that one may design various detection services and a predictor service that use the detection services to predict the object in a future horizon. Also, note that there may well be some common higher-level library or engine used by the Service layer. Such a library may use some parallel processing library underneath. One example may be the relationship between OpenCV and OpenCL. OpenCV provides the vision processing framework/library, which underneath (can) use OpenCL to use programmable accelerators. The OpenCV library may be used by multiple services. This is similar for the relationship between OpenVX and OpenCL - however, unlike OpenCV, OpenVX is designed so that the OpenVX interface implementations can directly access the specific accelerators, without OpenCL in between. Therefore, it is drawn to be in the Parallel processing library/language layer in the figure.

The Service layer uses Parallel processing library/language layer, which can vary depending on the choice of parallel processing technologies used in the service implementation. The programming interface for this layer varies as discussed in "5.1.1 Evolving parallel processing technologies", and it is just not semantically possible to have a single unified interface to generalize all or even most of the different interface/languages, without severely impacting the performance benefit, which contradicts the purpose of employing the parallel processing in the first place.

5.2.2 Accelerator-model

The Parallel processing library/languages layer interacts with the parallel processing hardware in different ways. There are two general models. One is accelerator-model, where the parallel processing library/language calls underneath some form of device drivers that directly controls the parallel processing hardware. The device driver, depending upon the design of OS used, maybe another process or some form of kernel module that executes in the context of OS kernel. The examples of this accelerator model include OpenCL/CUDA, OpenVX, etc. Figure 5.2 shows the combined process and physical architectural views for OpenCL based parallel processing.

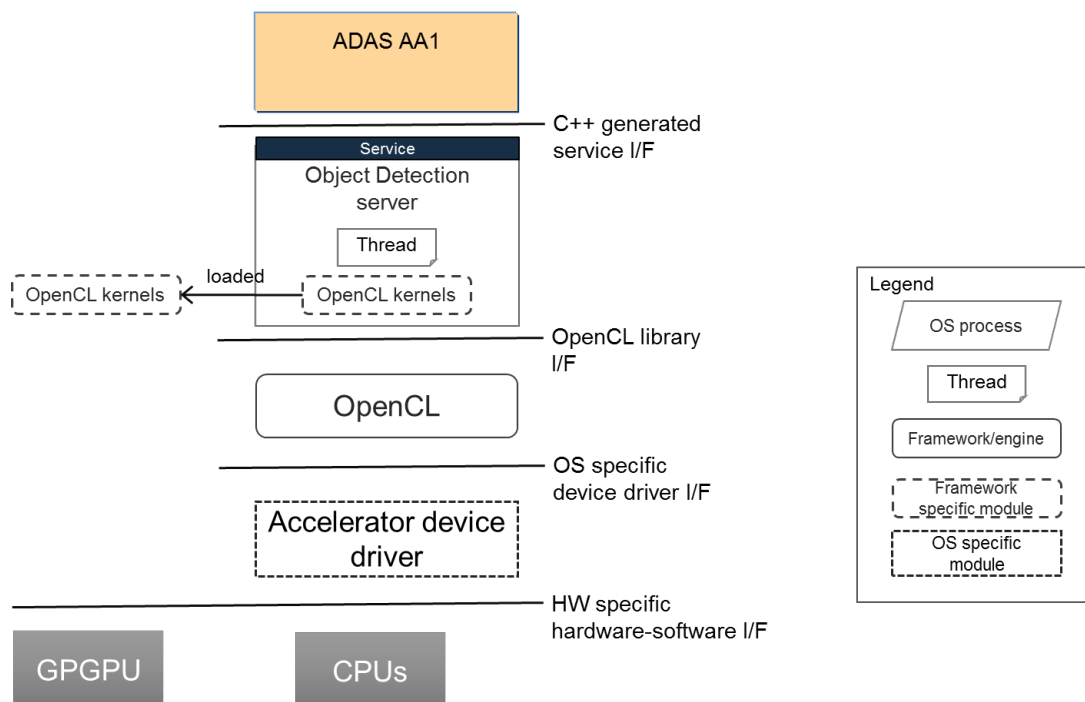


Figure 5.2: Accelerator-model example

5.2.3 CPU/co-processor-model

The other model is CPU/co-processor-model, where the parallel processing is executed directly by the CPUs with or without co-processor support. The most popular example is a threading model, which uses multiple POSIX threads to parallelize the processing. This can be fully handwritten, directive-based like OpenMP, or use some other vendor specific semi/full parallelization compiler technologies. Furthermore, there may be support for utilizing the specialized co-processor instructions, also may be manual or semi-automatic. Figure 5.3 shows the combined process and physical architectural views for threading-based (like POSIX threads) parallel processing.

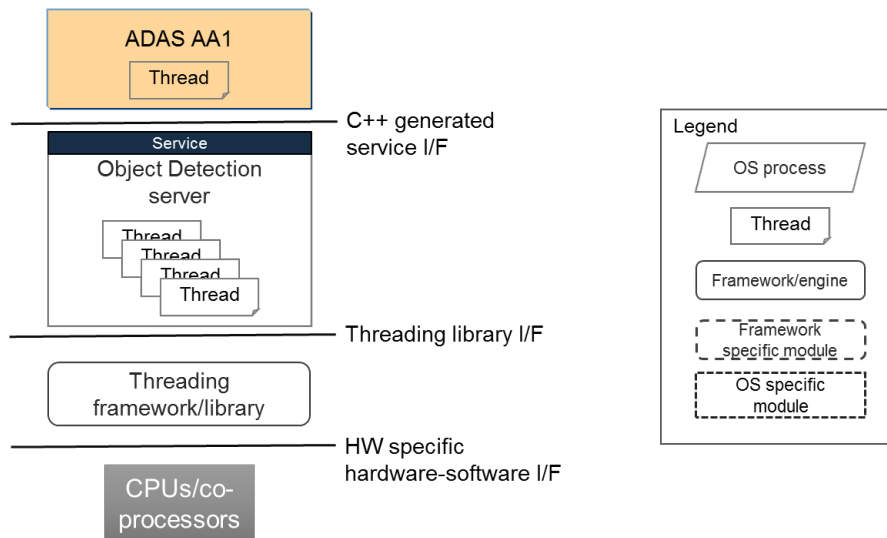


Figure 5.3: CPU/co-processor-model example

5.3 Rationale: decoupling of parallel processing specific knowledge from application development

Understanding the specifics of non-general computing hardware requires specific skills. As previously mentioned, the parallel processing technologies are still actively being developed and evolving, it is hard at the best to understand all these. Some standardization effort, such as OpenCL, aims to ease this problem by setting up a hardware independent API set. However, in order to cover various types of hardware and also to fully exploit the hardware features for best performance, the OpenCL, in general, is very low-level API, essentially requiring the similar level of detailed hardware level knowledge.

Our proposed model of Service-based parallel processing decouples the required knowledge of parallel processing hardware from Adaptive Application developers via Adaptive Platform service interfaces. This frees the Adaptive Application developers from acquiring the specific hardware knowledge every time a new, more efficient, or more suitable hardware is introduced, and also allows the system designer to do so if that the introduction of such hardware yields better system design. At the same time, this decoupling also frees hardware designers to come up with new, innovative parallel processing technologies, as long as they can provide the Adaptive Platform services required by the users.

5.4 Adaptive Platform methodology consideration

Service-based parallel processing approach will not introduce any new Adaptive Platform methodology. It uses already defined service interface description to define the services.

6 Non-functional design topics

6.1 Performance

One of the primary purposes of using parallel processing is to achieve higher performance. Since "Service-based parallel processing" utilizes the SOA of Adaptive Platform, the general performance-related design techniques also apply.

6.1.1 Interface granularity and communication overhead

The granularity of interface is the size of operation unit per API. If the granularity is small, the service has many API. The finer the granularity is, the more flexible the service is in general, because it will allow the different application to optimize its usage.

In SOA, increasing the granularity will increase the communication between the client and server in general. Caching mechanisms may circumvent that problem. However, in a real-time system such as Adaptive Platform, caching, depending on the architecture and implementation increases non-determinism, thus not a convenient choice.

In the Adaptive Platform, there are two possible approaches to minimize the communication overhead of services. One is to make the service interface as coarse as possible, especially for the interface that has a high frequency in its usage. For example, instead of providing an interface only for processing one datum, providing another interface for processing a batch of data at a time is recommended. The other approach is to optimize at the service interface library. This means that the service interface may cache some server-side data for client-local processing, and/or simple interfaces that set up or read fields in an object stored locally in the client process heap. The two techniques can be mixed.

6.1.2 Data handling and throughput balancing

The overhead of moving very large data is costly. This is especially true if a lot of copying of data occurs. Often, parallel processing is used to perform processing a large amount of data, and this is often performed against a stream of data, constituting a data-flow processing. It is therefore essential to design the whole chain of data flow, from a data-generating device, a device driver, a primary server to process the raw data, a secondary server to work on the primary server output, then finally an `Adaptive Application` that uses the result of the secondary server. One typical design to achieve the highest throughput is to have all these components forming PLP, each component forming a stage of a pipeline. For the servers that perform heavy computation, DLP and/or PLP is employed. The data that flows between the components have to be propagated in an efficient manner, avoiding copying of the data where possible.

6.2 Safety considerations

Safety is a system design topic. One of the typical issues is that the parallel processing hardware technology does not satisfy ASIL-D but only up to ASIL-B. The software can be developed based on ASIL-D practices but as the hardware is only capable of ASIL-B, as a whole it cannot achieve ASIL-D. The required ASIL for a "subsystem" depends on the system functionality it provides - e.g. parallel processing subsystem is used for ASIL-B system functionality, which computation result is safety-checked by ASIL-D subsystem. Or, one can go duplicate ASIL-B subsystem to achieve ASIL-D (though it may be expensive). The design guidelines for system design to achieve overall functional safety requirements are out of the scope of this document.

6.3 Prospects

This guideline, especially the parallel processing hidden under the service model, should be capable of surviving a long time to come, due to the intrinsic decoupling. There are two areas with foreseeable advancement in future; (1) Adaptive Platform standard application services and (2) more parallel processing directly within Adaptive Application.

6.3.1 Adaptive Platform standard application services

It should be reasonable for one to expect such application services that use parallel processing technologies to be standardized by the Adaptive Platform. This indeed will not occur immediately, nor all services at once - however, even incremental introduction of such services should help both the providers of parallel processing technologies and also users of such. Higher level API standardization, that uses parallel processing technologies underneath, are already emerging in some areas, such as OpenVX.

6.3.2 More parallel processing within Adaptive Application

Following the service-based parallel processing design, Adaptive Application will use multiple services in parallel. As the number of services grows and if the Adaptive Application remains single-threaded, then the Adaptive Application itself can be a bottleneck in the whole processing chain. This will call for more parallel processing within Adaptive Application eventually. The Adaptive Platform already provides threading APIs of currently supported C++ standard and POSIX APIs, however, this may not be sufficient.

AUTOSAR Adaptive Platform adopts the C++ standard, along with the CPP Coding Guidelines to use the language with safety and security in mind. The C++ standard is incrementally introducing parallel processing. The most of both open source and commercial compilers support the standard and widely used in the industry. Therefore, it is foreseeable and perhaps promising to introduce more parallel processing technologies

as the C++ standard progresses. One potentially promising standard is SYCL, as it is purely based on standard C++ with template libraries to implement parallel processing, part of it being introduced in C++17. The single source approach with the standard language and also capable of mixing with normal C++ multi-threading may help to consolidate the situation in future. Moreover, SYCL SC (Safety-Critical)¹ is being actively developed for safety critical areas.

AUTOSAR Adaptive contains Functional Cluster responsible for parallel tasks execution with hardware acceleration - Safe Hardware Acceleration (ara::shwa). For more details please read [4].

¹<https://www.khronos.org/syclsc>