



NumPex ExaSoft

25th Sep 2025

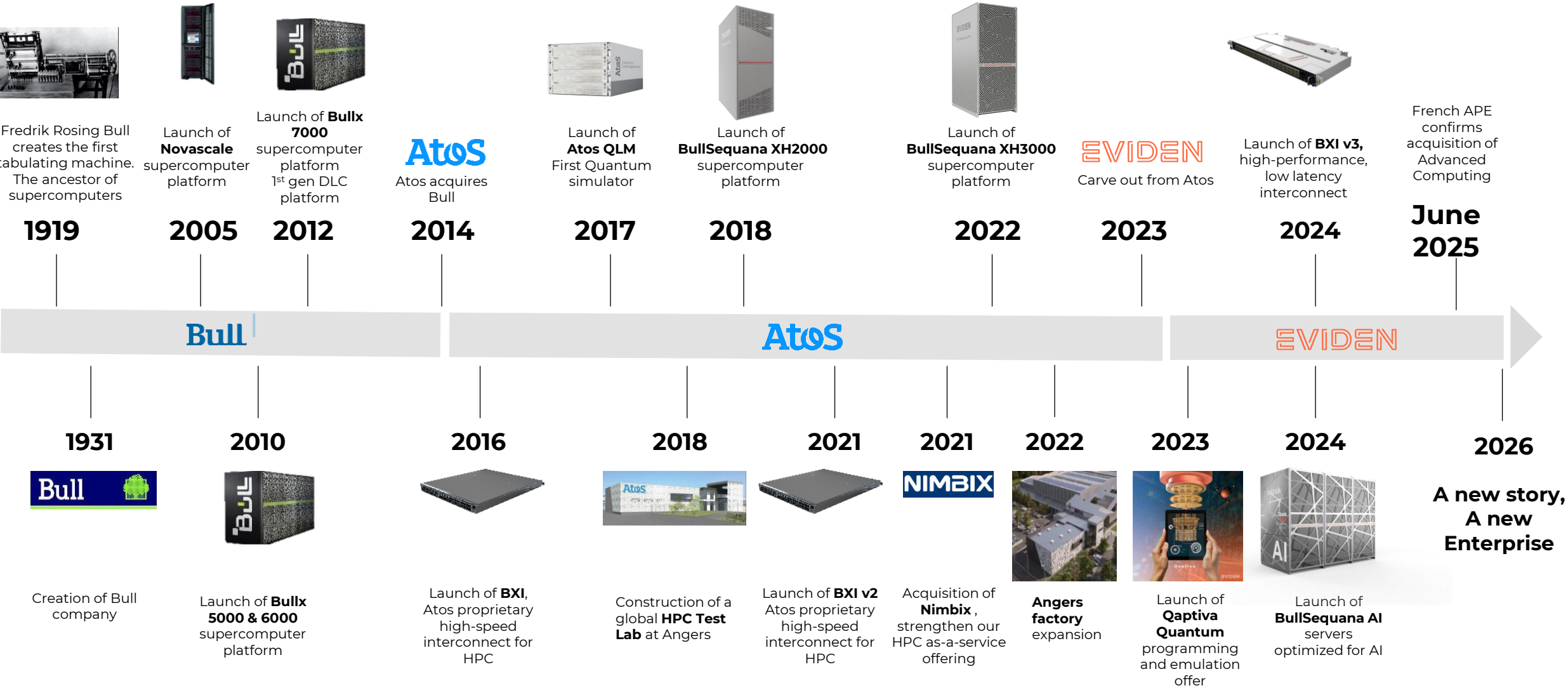
David Goudin

david.goudin@eviden.com

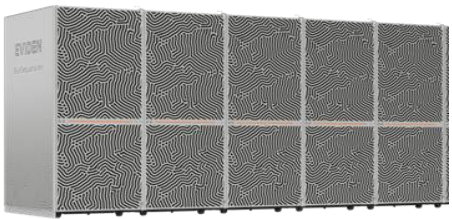


Eviden - Europe's Supercomputing Powerhouse

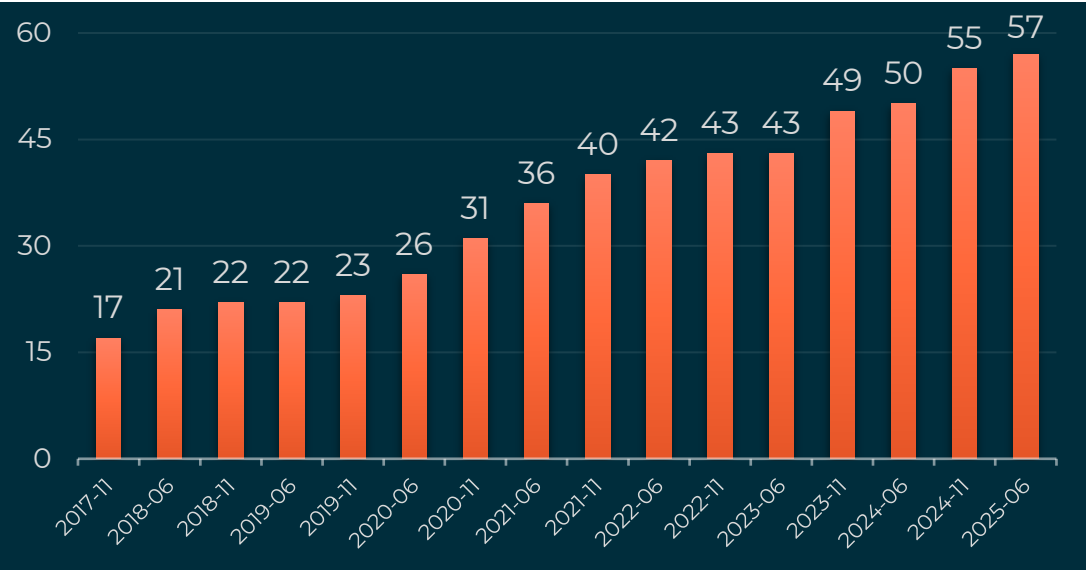
A decade of HPC leadership powering European innovation and sovereignty



Eviden: Tracked record in TOP500



A growing number of Eviden systems in TOP500 since 2017



- 55 as “EVIDEN” manufacturer
- +1 as “ParTec/Eviden” manufacturer (JEDI)
- +1 as “NVIDIA” (DCAI Gefion)

#1

- The number of systems in **Europe** (49)
- The number of systems in **South America** (5)
- The number of systems in **India** (3)

#3

The number of systems in TOP500 (WW)



3 systems in TOP20



#1 and #2 in Green500



Hosting the 1st European Exascale System



50 pre-built and interchangeable modules

- 20 IT containers
- 15 power feed containers
- 10+ logistic & support containers (lobby, workshop, warehouse)
- over 2,300m² to form a complete turnkey data center

Rank	System	Cores	Rmax (PFlop/s)	Rpeak (PFlop/s)	Power (kW)
1	El Capitan - HPE Cray EX255a, AMD 4th Gen EPYC 24C 1.8GHz, AMD Instinct MI300A, Slingshot-11, TOSS, HPE DOE/NNSA/LLNL United States	11,039,616	1,742.00	2,746.38	29,581
2	Frontier - HPE Cray EX235a, AMD Optimized 3rd Generation EPYC 64C 2GHz, AMD Instinct MI250X, Slingshot-11, HPE Cray OS, HPE DOE/SC/Oak Ridge National Laboratory United States	9,066,176	1,353.00	2,055.72	24,607
3	Aurora - HPE Cray EX - Intel Exascale Compute Blade, Xeon CPU Max 9470 52C 2.4GHz, Intel Data Center GPU Max, Slingshot-11, Intel DOE/SC/Argonne National Laboratory United States	9,264,128	1,012.00	1,980.01	38,698
4	JUPITER Booster - BullSequana XH3000, GH Superchip 72C 3GHz, NVIDIA GH200 Superchip, Quad-Rail NVIDIA InfiniBand NDR200, RedHat Enterprise Linux, EVIDEN EuroHPC/FZJ Germany	4,801,344	793.40	930.00	13,088

Angers Factory

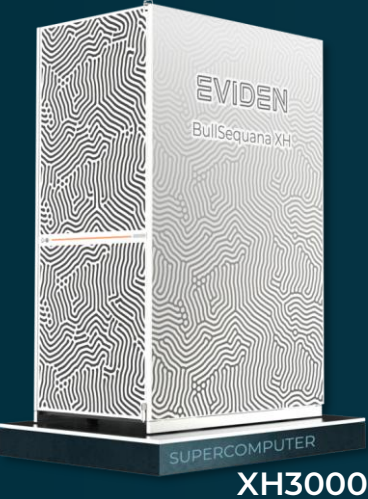


High Level Portfolio



Eviden Own IP Portfolio for HPC/AI/Quantum Computing

Overview



BullSequana SuperComputers



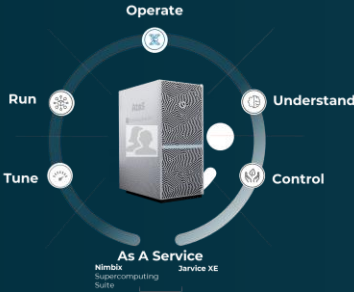
BullSequana SMP



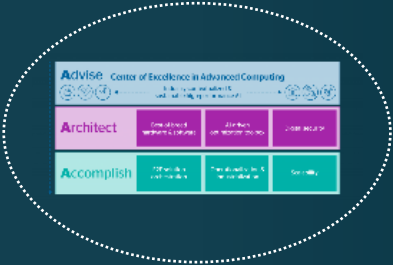
BXI Interconnect



Quantum Computing



HPC Software Suites



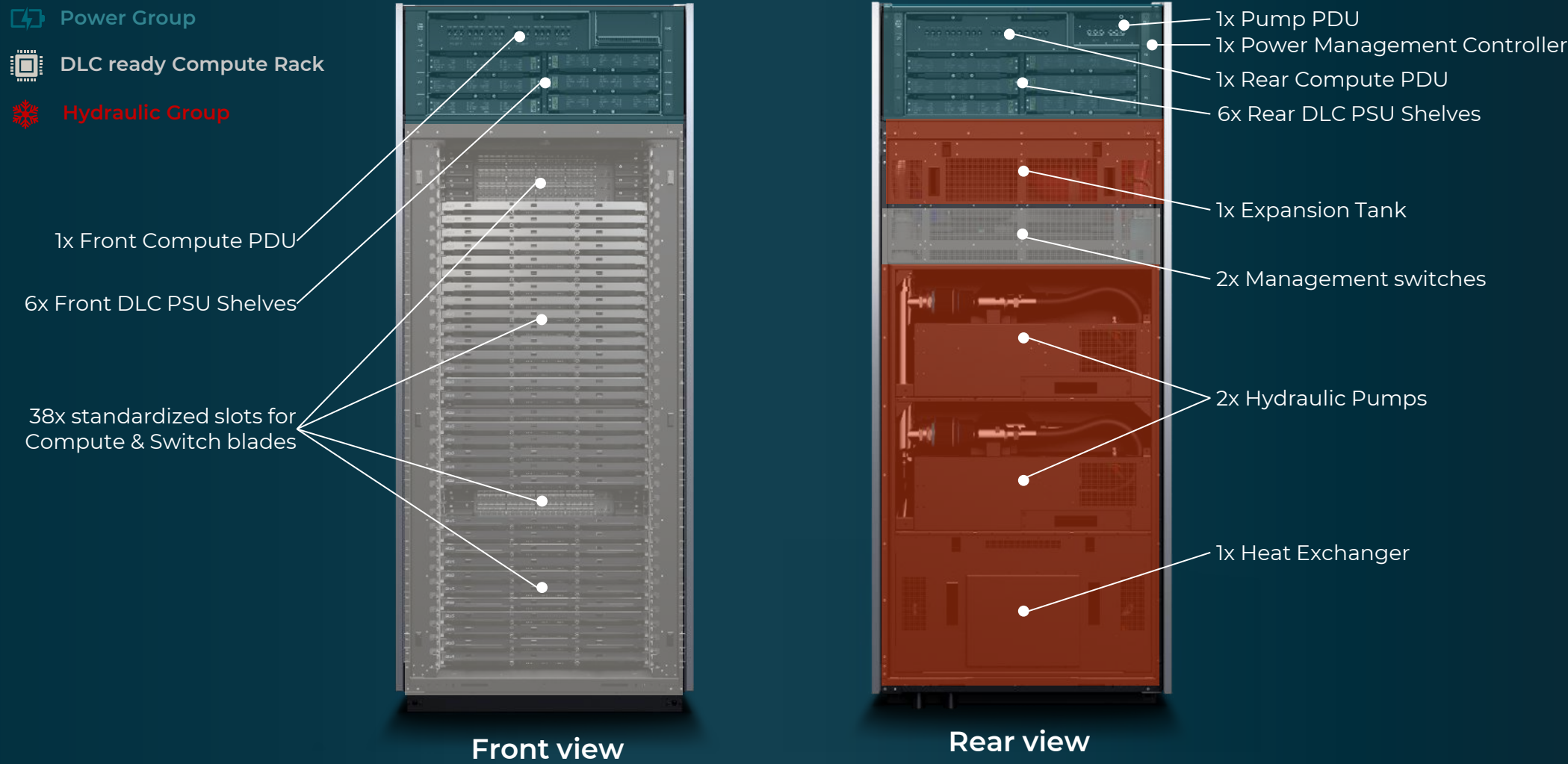
Large Scale AI



As a Service

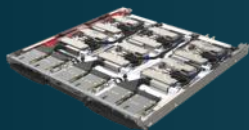
BullSequana XH3000

Infrastructure overview



BullSequana XH3000 Blades = A lot of different options

CPU Compute blades



BullSequana XH3140
(codename *TRIO*)

Available

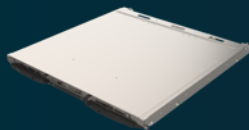
6x  **intel** Emerald Rapids



BullSequana XH3420
(codename *MATO*)

Available

6x  **AMD** Turin



BullSequana XH3610
(codename *MILO*)

In development

6x  **SIPEARL** Rhea 1

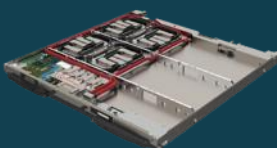
GPU Compute blades



BullSequana XH3145-H
(codename *TIYA*)

Available

1x  **NVIDIA** HGX 4-H100



BullSequana XH3515-HMP
(codename *NEBA MaxP*)

Available

1x  **NVIDIA** HGX CG4



BullSequana XH3515-HMQ
(codename *NEBA MaxQ*)

Available

2x  **NVIDIA** HGX CG4

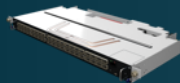


BullSequana XH3406-3MQ
(codename *MANA MaxQ*)

In development

6x  **AMD** MI300-A

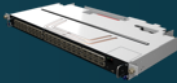
High-speed Interconnect switch blades



BullSequana XH32B2-0100
(codename *W3B1M*)

Available

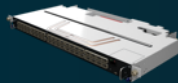
 **EVIDEN** BXI V2



BullSequana XH35IN-0800
(codename *W3NM*)

Available

 **NVIDIA** InfiniBand NDR



BullSequana XH35ET-0200
(codename *W3EM*)

Available

 **NVIDIA** High-speed Ethernet

Yes it is ExaSoft project Meeting so no
more Hardware....

A brief overview of Eviden Teams involved in Applications, Performances, Optimizations

A need of good interactions between Teams

Exploration Center

Explore new technologies and project first performances

CEPP

Applications optimizations
Collaborative Projects

A&P

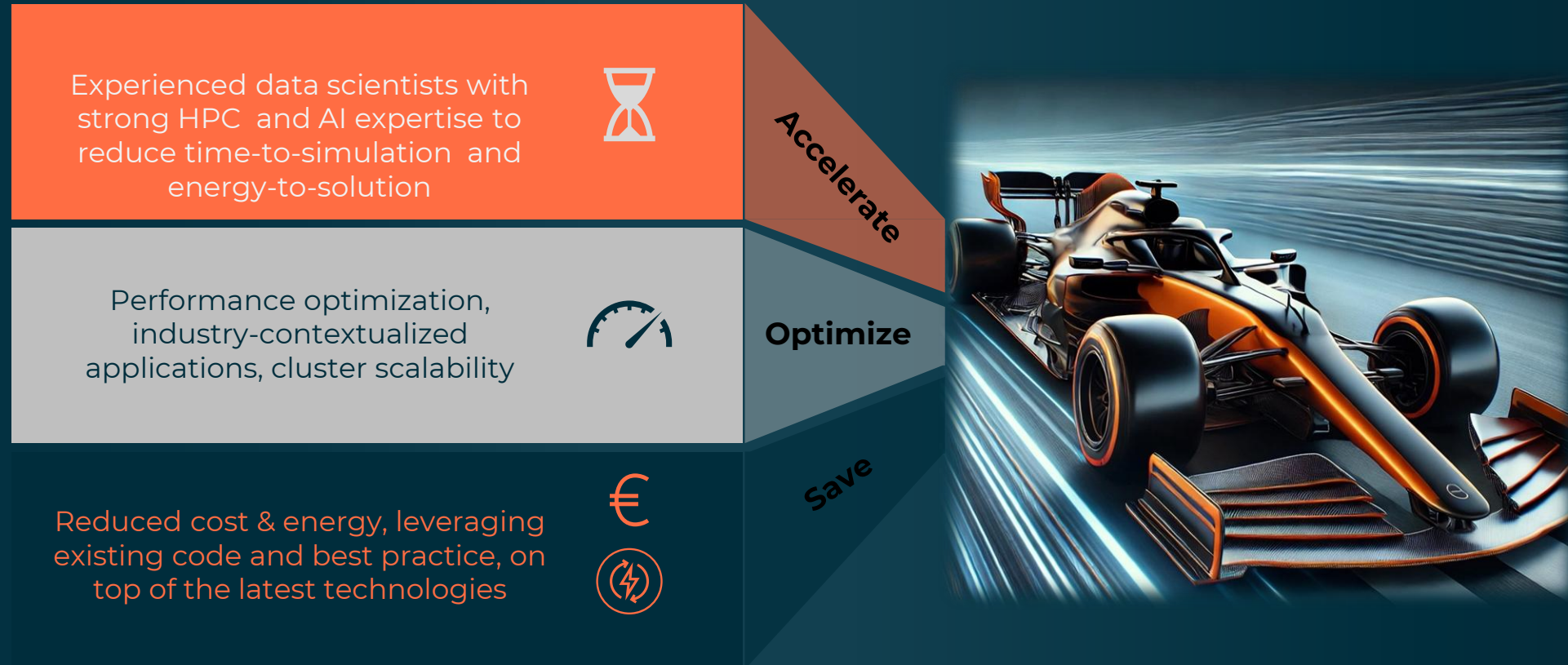
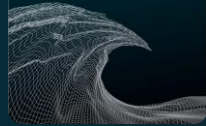
Projection of benchmarks performances for Deals on architectures

R&D SW Teams

Development of SW Portfolio
Collaborative Projects

CEPP – Center for Excellence in Performance Programming

CEPP Acceleration
Accelerate workload,
give value to simulation!



Center for Excellence in Performance Programming (CEPP)

CEPP and R&D Collaborations - HPC AI and Quantum collaboration projects

European technology for exascale



Quantum computing



Cloud



Centre of Excellence for computing applications



Customer Centre of Excellence



International collaboration



Code Migration: ICON OpenACC to OpenMP

Case : DKRZ



Acceleration
Program

CEPP



Methodology

- Translate OpenACC directive to OpenMP
- Comparison between OpenACC and OpenMP

1st Analyze

Analyse the existing ICON OpenACC implementation to assess compatibility and potential challenges in porting to OpenMP.

2nd Isolate

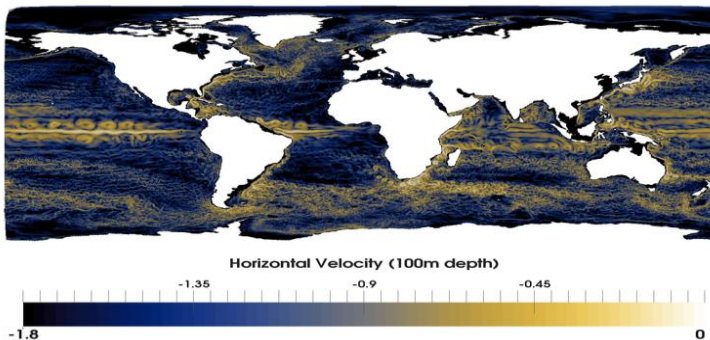
Isolate the nh_solve function for kernel extraction, leveraging work in the EECLiPs project

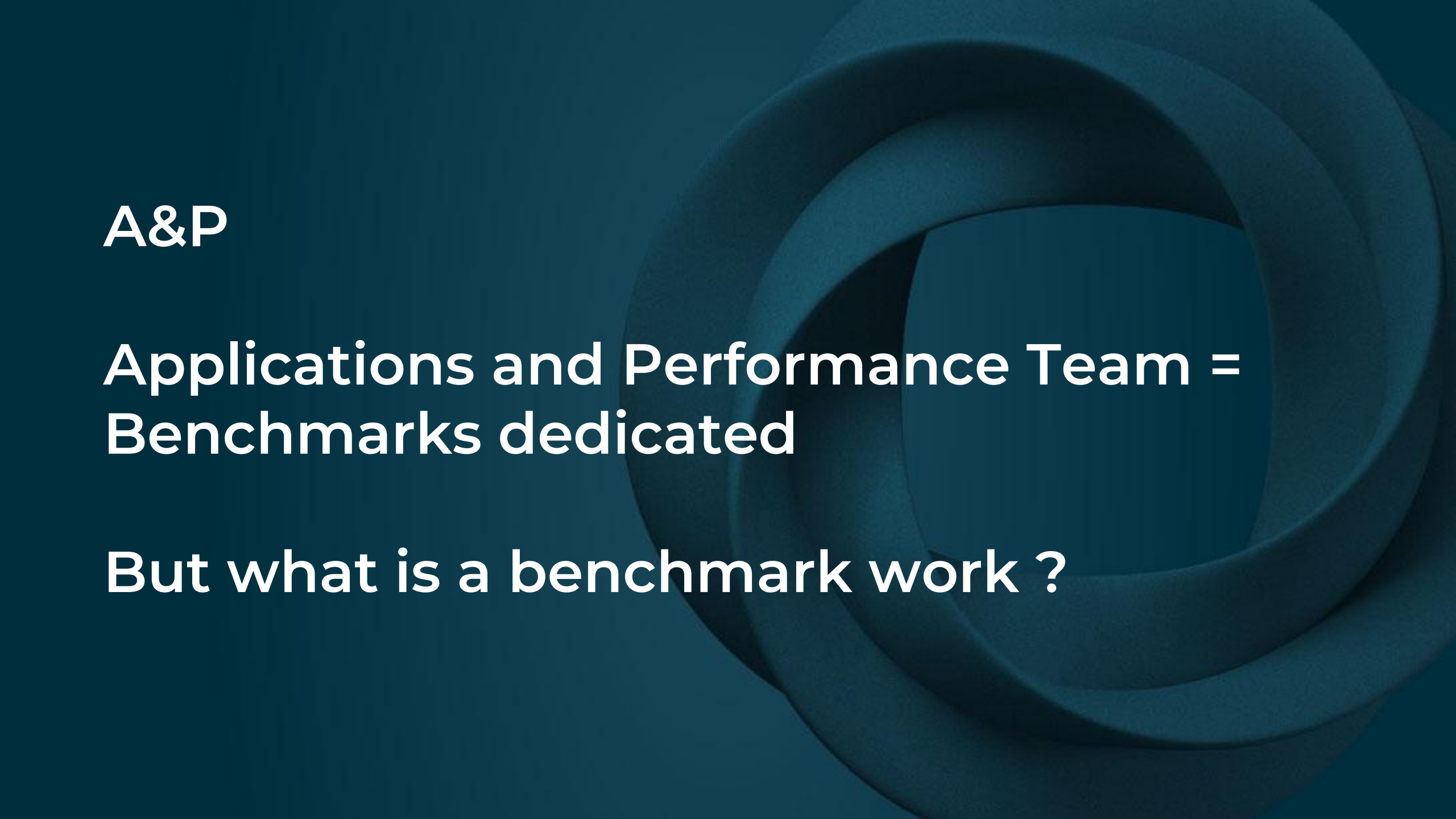
3rd Convert

Convert the extracted nh_solve kernel to utilize OpenMP targeting constructs.

4nd Evaluate

Evaluate the performance, usability, and porting effort of the OpenMP version against the OpenACC original, across a comprehensive range of GPUs.



The background of the slide features a series of concentric, overlapping circles in various shades of dark teal and blue, creating a tunnel-like or ripple effect that draws the eye towards the center.

A&P

Applications and Performance Team =
Benchmarks dedicated

But what is a benchmark work ?

Applications & Performance Group

Typical “applications” ?

- Typical “Exascale” applications:
 - Abinit, BigDFT, CP2K, GROMACS, LAMMPS, NAMD, OpenFOAM, Quantum Espresso, WarpX, YALES2
- 4 to 6 applications and even more (e.g. Alice Recoque Exascale Supercomputer)
- Kernels to test key system metrics:
 - Compute: HPL (single node & full system)
 - « System »: HPCG (single node & partition)
 - Memory: Stream (BabelStream for GPU)
 - Network: OSU MPI benchmark
 - Latency
 - Bandwidth
 - I/O: IOR

Applications & Performance Group

What is benchmark work ?

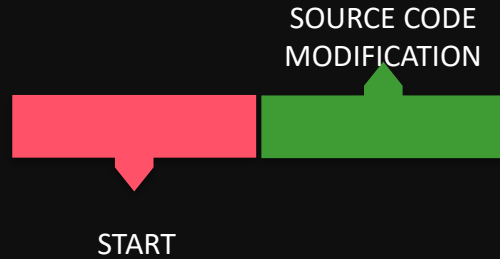


START

- Ensure benchmark conditions:
 - Reference: measure on existing machine (usually provided by clients)
 - Hardware requirements: Memory, #process/core, GPU, interconnect, filesystem
 - Dependencies: Software, licenses
 - Measurements:
 - test case(s),
 - Reduced test case(s)
 - Application kernel(s) (small part of applications representing most of runtime)
 - Goal: validation, reference time, expectations

Applications & Performance Group

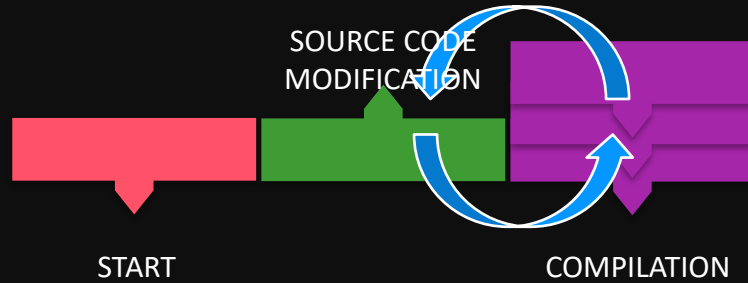
What is benchmark work ?



- Modify source code (if authorized):
 - Bug correction
 - Cache alignment
 - Introduce more parallelism
 - Overlap communication and compute
 - Accelerator support

Applications & Performance Group

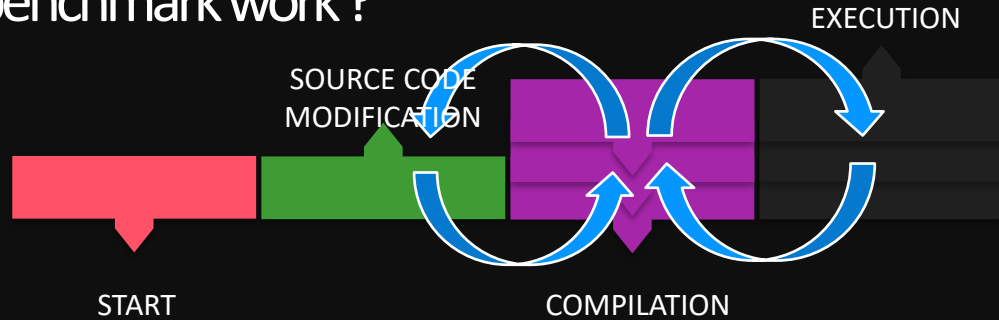
What is benchmark work ?



- For each benchmark, compilation is optimized with
 - Compiler provider : Intel, AMD, NVIDIA, LLVM, GNU, ...
 - Compiler version: client's reference, up-to-date, ...
 - Compiler options: instruction sets, optimization level, ...
 - Optimized dependencies (libraries, ...)
- A compilation could take several seconds or several hours !

Applications & Performance Group

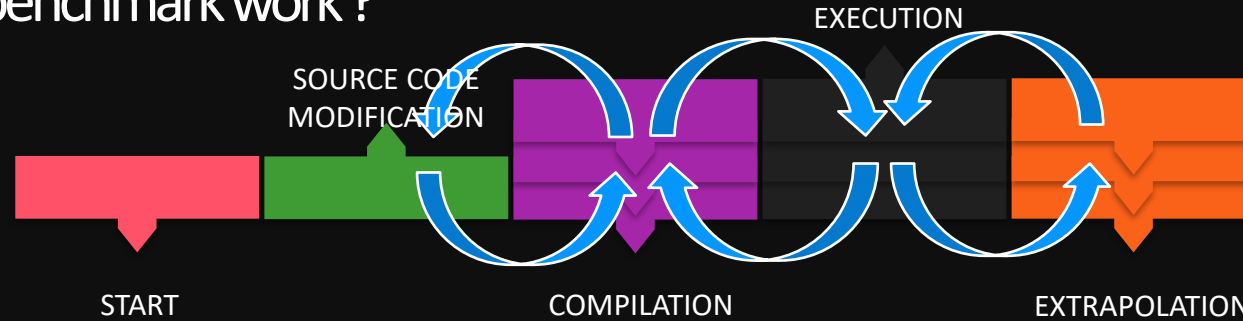
What is benchmark work ?



- For each benchmark, execution should be optimized playing with
 - Kernel environment / System tuning
 - Processes binding
 - Submission command
 - OpenMP/OpenACC/libraries tuning
 - MPI tuning
 - Application tuning (application parameters).
- Each test require a new execution and each bottleneck required a new tuning !
- Perform the tasks on various systems - if possible reduced test case (or kernel) on targeted system's sample

Applications & Performance Group

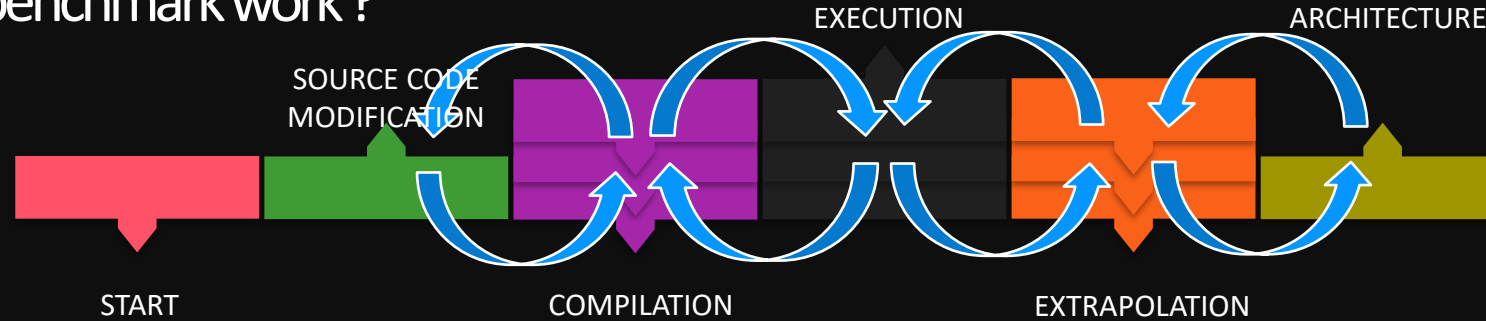
What is benchmark work ?



- For each benchmark, when the best performance is found (sic), benchmark behavior is characterized to extrapolate time on another platform. The characterization required to run several new tests playing with:
 - Frequency => Determines compute characteristics
 - Binding => Determines Memory characteristics
 - Communication (non-blocking/blocking) and profile => Determines communication and imbalance
- This characterization gives input to make the extrapolation:
 - Extrapolate each component (compute/memory/network) from tested system to target
- Do the exercise both on requested metrics AND reduced test case/kernel
- Work with provider(s) (Intel, AMD, and/or NVIDIA, ...) to best match future components characteristics

Applications & Performance Group

What is benchmark work ?



- This step is a target since the START: optimization path varies according to the goal!
 - (Performance, energy, density...)
- Brainstorming between application specialists and system architects to maximize client criteria:
 - Pure Performance
 - Performance/price
 - Performance/watt
 - Performance/m²
 - Energy
 - Price
 - ...

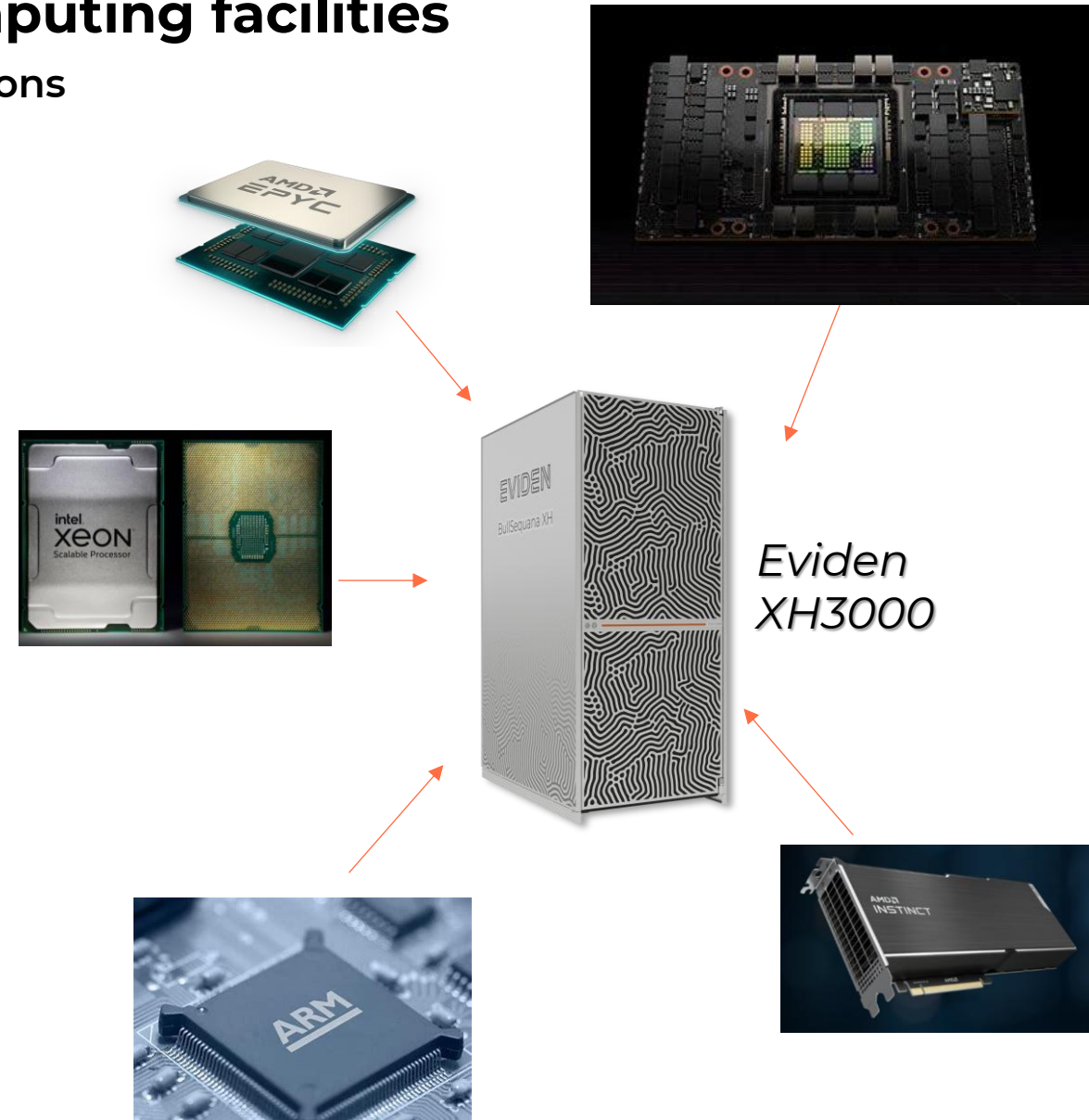
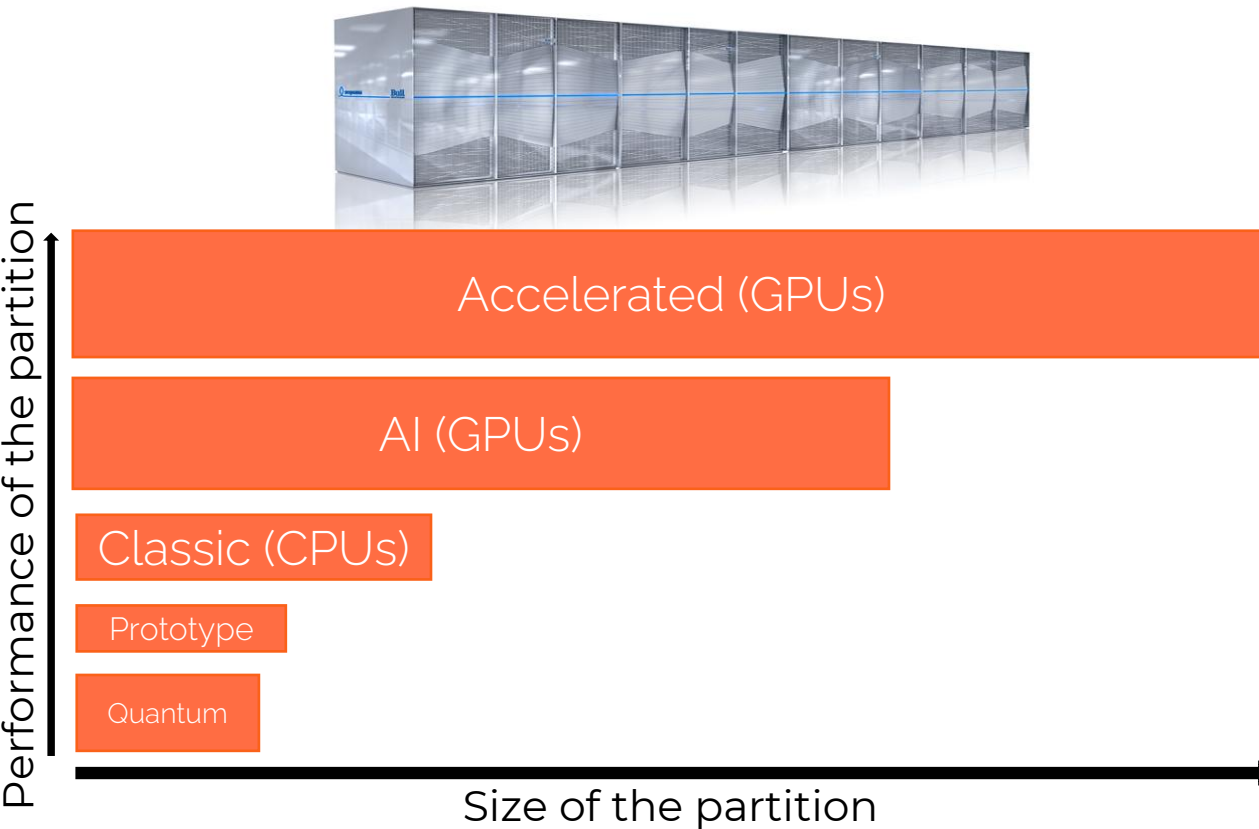
The background is a dark teal color. On the right side, there is a large, abstract graphic consisting of several concentric circles. In the center of these circles is a square. The circles and the square have a 3D effect, with shadows and highlights that make them appear to be floating or layered. The text is positioned on the left side of the image, overlapping the circles.

**Link between Application /
Benchmark and ExaSoft project ?**

Eviden provides Large Modern Supercomputing facilities

HPC/AI data centers with multiple, heterogeneous partitions

- More and more complex and different architectures
- Different type and size of partitions, and objectives
- Different type of calculations (FP64 or not)



Example of Synthetic Benchmarks



- Problem
 - Very often, the most efficient solvers (from Partners: Intel, AMD, Nvidia)) as HPL, HPL-MxP and HPCG are not open source
 - Except for AMD GPUs (rochPL, rochPCG)
 - Evolution of Hardware (like for GPUS, FP64 vs other data formats)
 - Need of Simulation tools associated (At scale, Performance, Energy) with kernels models/performances for future architectures provided by our partners

Existing task based runtime

Unexhaustive

- Shared memory parallelism
 - OpenMP tasks (+ depend clause)
 - Cilk (MIT, Intel)
 - cudaSTF (Nvidia)
 - OmpSs (Barcelona Supercomputing Center)
 -
- Both shared and distributed memory parallelism
 - StarPU (Inria)
 - PaRSEC (University of Tennessee)
 - Charm++ (University of Illinois)
 - High Performance ParalleX (STE||AR group)
 - Legion (University of Stanford)
 -

CUDASTF

OpenMP

OpenCilk

OmpSs-2
Programming Model

Legion

PARSEC

StarPU

Charm++

HIPX
STE||AR GROUP

Collaborations, Solver Part



StarPU

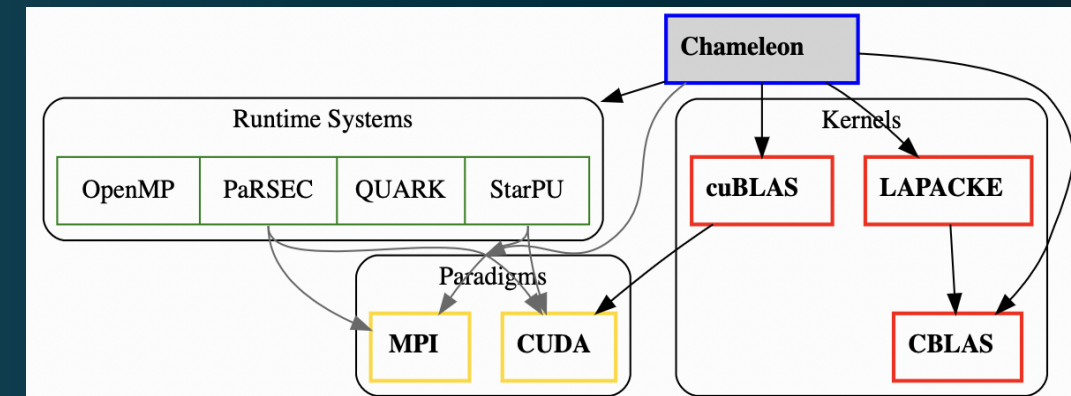
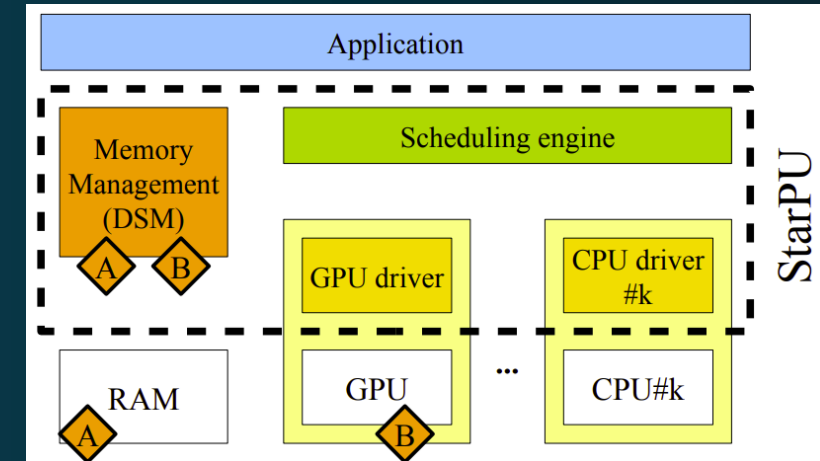
- Task based runtime system, STF model
- Both CPU and GPU (CUDA, OpenCL, HIP, SYCL) drivers
- MPI extension

Chameleon

- Dense linear algebra library on top of runtime systems
- Implements distributed blocked LU algorithm without partial pivoting (among others)

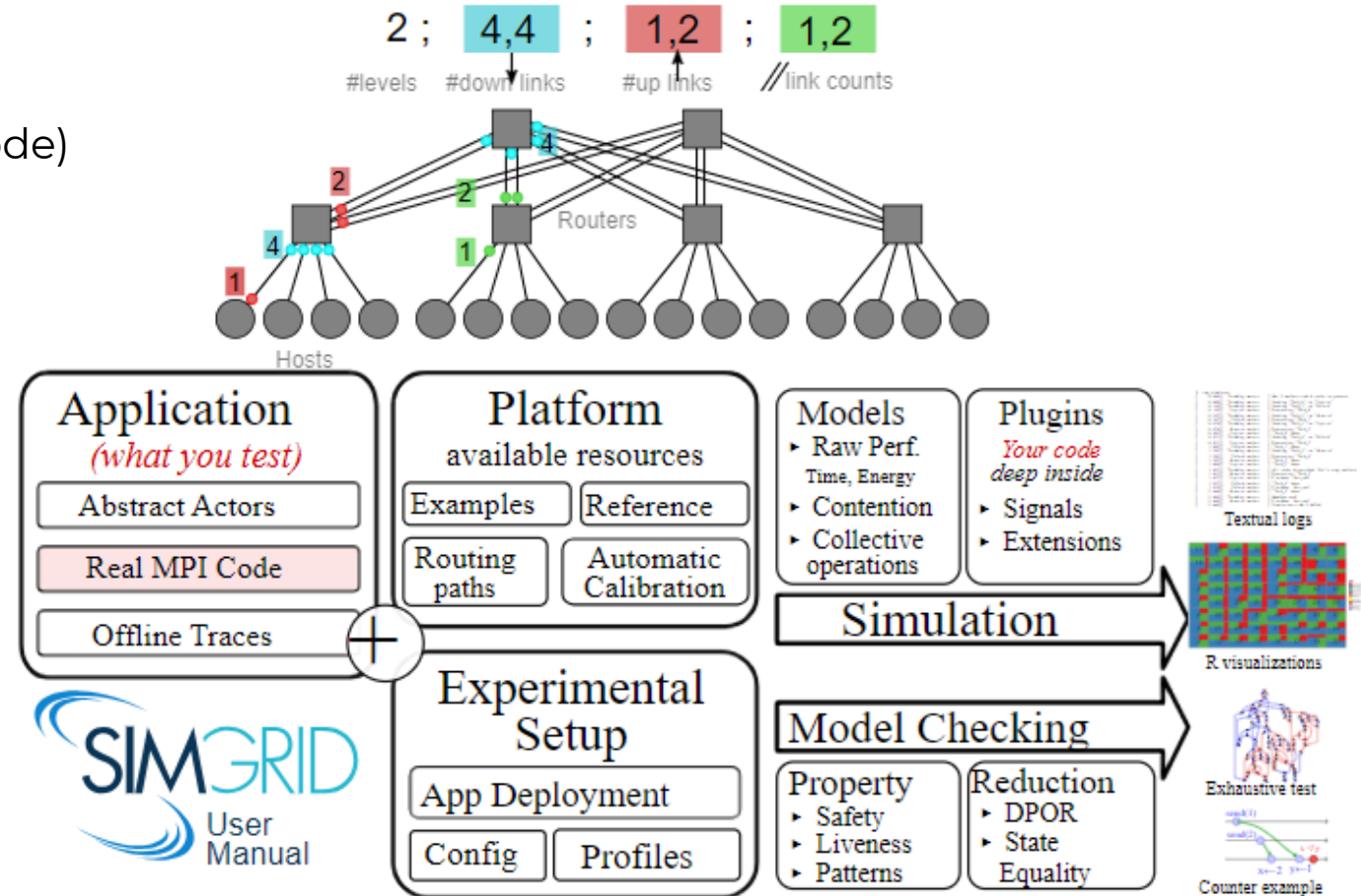
NewMadeleine

- Optimizing Communication Library for High-Performance networks



Collaborations, Simulation Tool Part

- Performance **prediction**
 - On existing and future architectures (one node)
 - At scale (network performances simulation)
- **Sensitivity** analysis
- Performance **tuning**
- Energy efficiency and Energy Performance **projection**



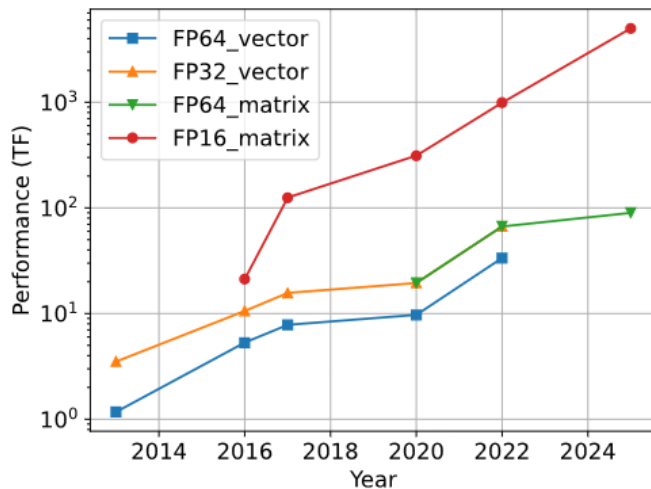
Energy performance simulation & optimization

Goals:

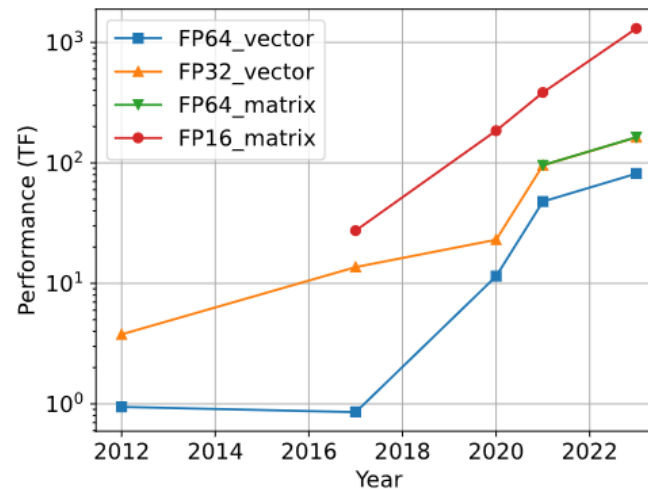
1. Energy-aware **simulation**: understand and describe power usage
 - Simulate the impact of power capping on performance
 - Predict energy consumption at scale
2. Energy-aware performance **optimization**: leverage perf/power trade-offs
 - Improve dynamic scheduling
 - take into account the performance drop of GPU when CPU is used
 - Optimize for efficiency
 - choose dynamically the architecture with the better GFLOPS/W
 - Ensure a global consumption envelope is not exceeded
 - delay tasks that are not on the critical path

Emulation and Mixed-Precision Context, Evolution of data formats

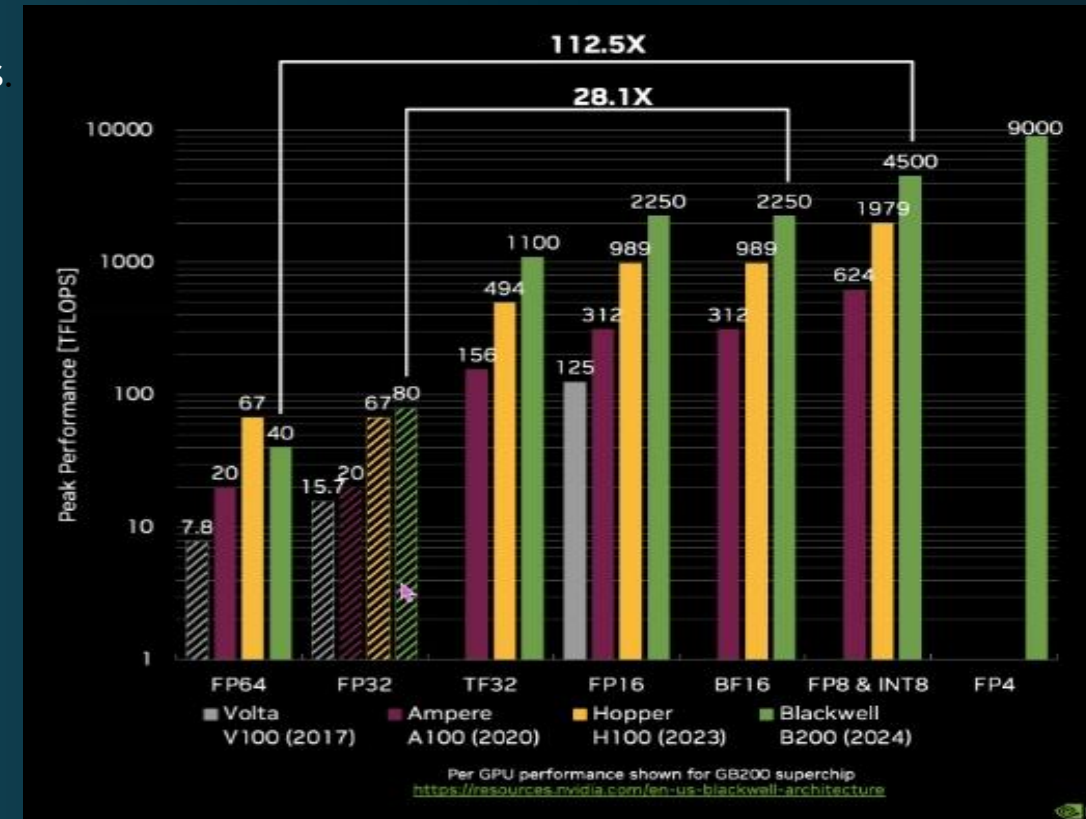
- “Convergence” of AI and HPC hardware (GPUs, CPU Matrix Cores, ...)
- Speedups reachable, better energy efficiency, if:
 - lowering precision of algorithms;
 - or emulating FP64 GEMM with low precision computations.



(a) NVIDIA



(b) AMD



Source: <https://arxiv.org/pdf/2412.19322> (ORNL, 2025).

Possible Links everywhere and every topic of ExaSoft Project for Future Architectures

- Optimization / Compilation
- Numerical libraries
- Runtimes
- Portability
- Profiling Tools
- Energy Optimization
- Extension to other Applications than Synthetic Benchmarks,
- Exascale and Post Exascale Projects