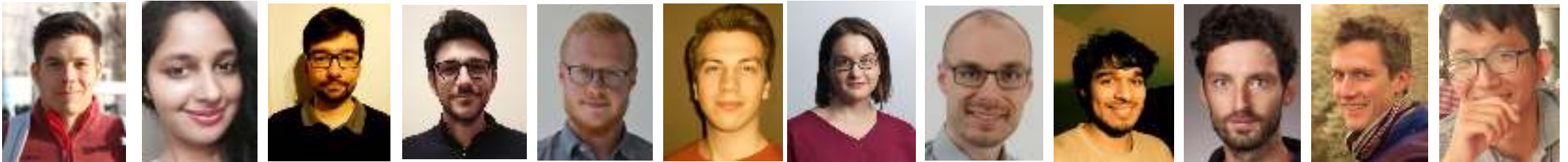


Sustaining Simulation Performance in the US Exascale Computing Project – A tale from the trenches



Approved for public release

Hartwig Anzt, University of Tennessee



Designing an ECP library for sustaining simulation performance

MAGMA SPARSE

MAGMA-sparse as a “child” of MAGMA
explores the development of sparse
linear algebra for NVIDIA GPUs.

Limitations:

- *C code with hand-written build system*
- *Sparse unit testing*
- *Focus on NVIDIA GPUs*
- *Design-specific limitations (flexibility/extensibility)*

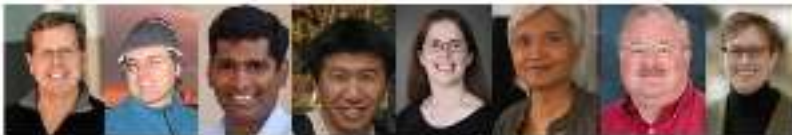
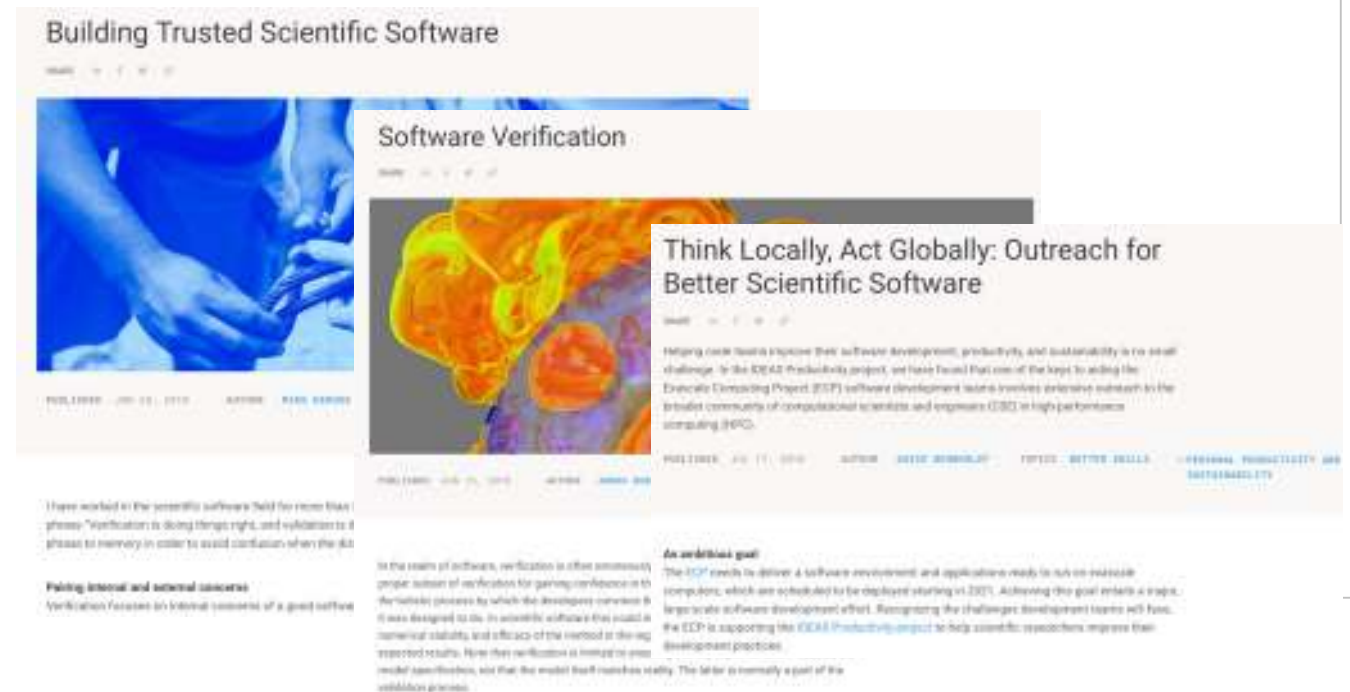
Designing an ECP library for sustaining simulation performance

MAGMA SPARSE

MAGMA-sparse as a “child” of MAGMA explores the development of sparse linear algebra for NVIDIA GPUs.

Design considerations for Ginkgo

- Platform Portability
- Performance
- Rapid integration of new algorithms
- xSDK / E4S Community Policies
- BSSw expertise / experience
- Modern C++
- CI/CD and unit testing
- Open source & permissive licensing



Designing an ECP library for sustaining simulation performance

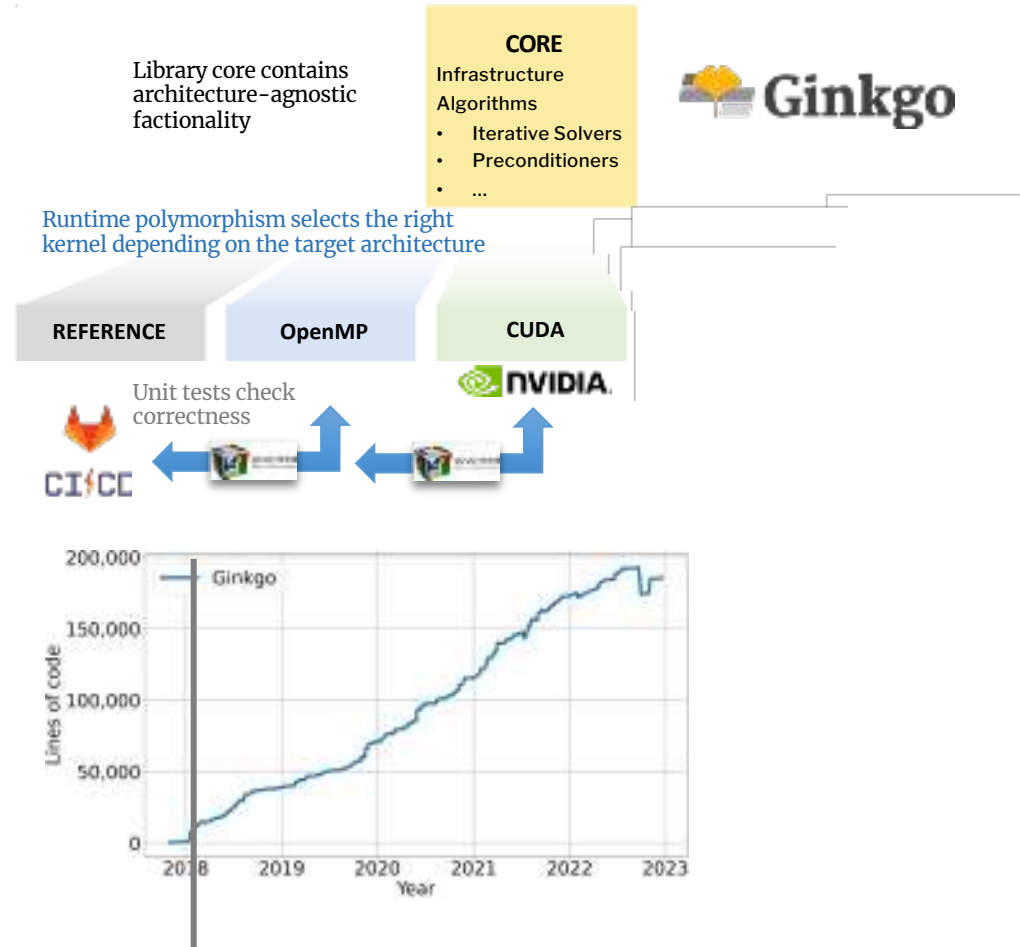
MAGMA SPARSE

MAGMA-sparse as a “child” of MAGMA explores the development of sparse linear algebra for NVIDIA GPUs.

Design considerations for Ginkgo

- Platform Portability
- Performance
- Rapid integration of new algorithms
- xSDK / E4S Community Policies
- BSSw expertise / experience
- Modern C++
- CI/CD and unit testing
- Open source & permissive licensing

Before the first line of code is written, we spend a year on whiteboard discussions.



The LinOp abstraction

Linear Operator Interface

We express everything as Linear Operator

- Automatically use language C++ cross inheritance
- Applications can apply any functionality as a linear operator

Matrix Vector Product

Preconditioner (for matrix A)

Solver (for system $Ax = b$)

$x = A \cdot b$

$x := M^{-1} \cdot b$

$x := S \cdot b$

$M^{-1} := A^{-1}$

$M^{-1} := H(A)$

$S := \mathcal{F}(A)$

All of them can be expressed as

Application of a linear operator $(LinOp)$ $b \mapsto A^{-1} \cdot b$



Library core contains architecture-agnostic factuality

CORE

Infrastructure Algorithms

- Iterative Solvers
- Preconditioners
- ...

Runtime polymorphism selects the right kernel depending on the target architecture

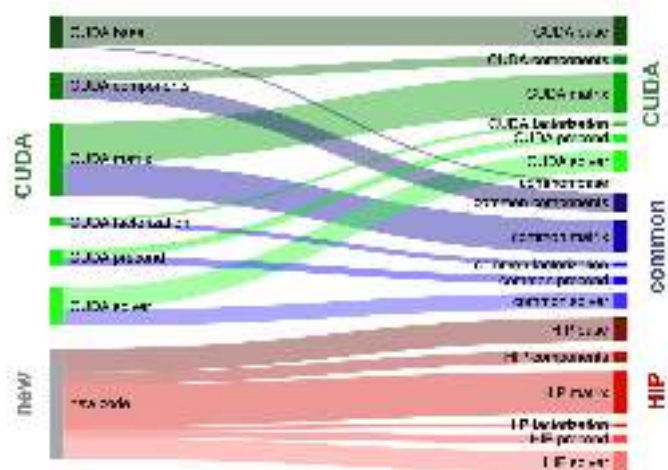
REFERENCE

OpenMP

CUDA

	Functionality	OMP	CUDA
Basic	SpMV	✓	✓
	SpMM	✓	✓
	SpGeMM	✓	✓
Krylov solvers	BiCG	✓	✓
	BiCGSTAB	✓	✓
	CG	✓	✓
	CQS	✓	✓
	GMRES	✓	✓
Preconditioners	IDR	✓	✓
	(Block-)Jacobi	✓	✓
	ILU/IC	✓	✓
	Parallel ILU/IC	✓	✓
	Parallel ILUT/ICT	✓	✓
	Sparse Approximate Inverse	✓	✓

Extending to AMD GPUs



Library core contains
architecture-agnostic
factionality

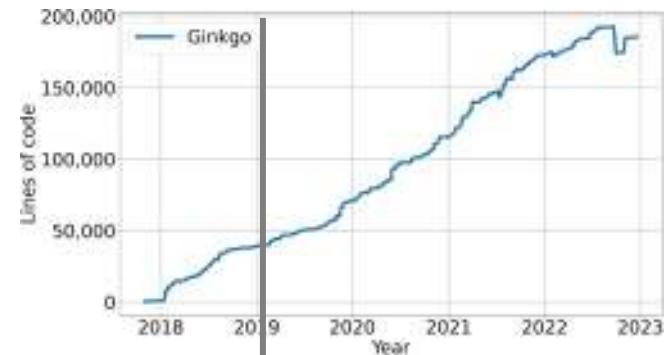
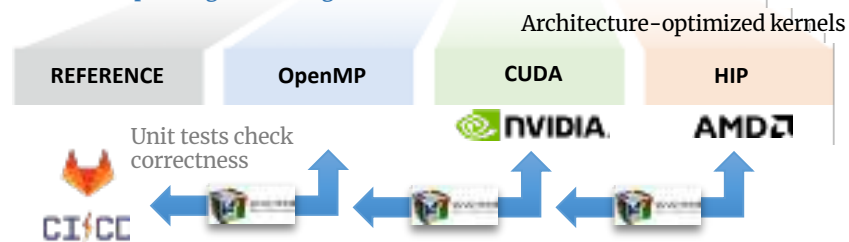
CORE

Infrastructure

Algorithms

- Iterative Solvers
- Preconditioners
- ...

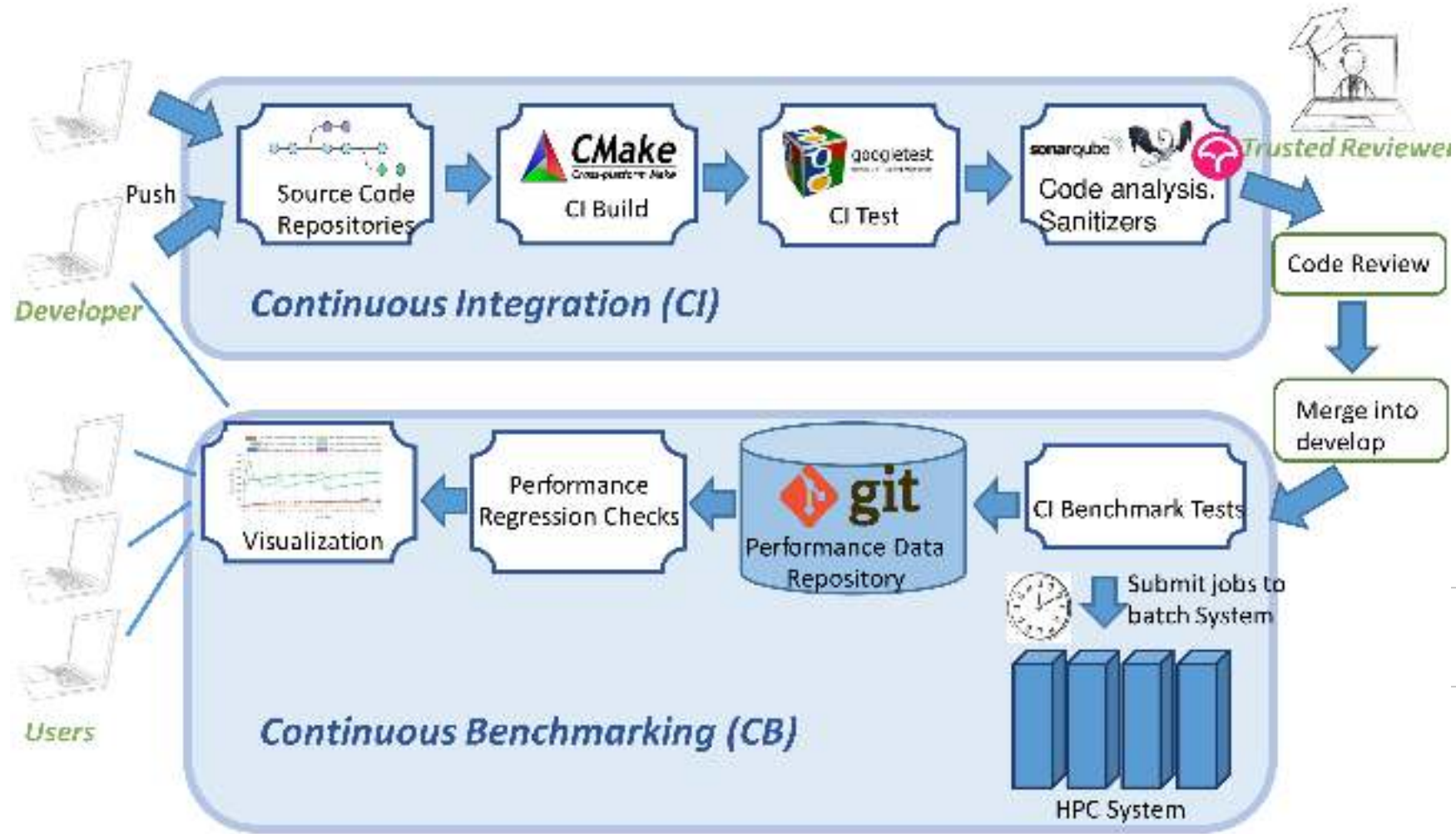
Runtime polymorphism selects the right kernel depending on the target architecture



	Functionality	OMP	CUDA	HIP
Basic	SpMV	✓	✓	✓
	SpMM	✓	✓	✓
	SpGeMM	✓	✓	✓
Krylov solvers	BICG	✓	✓	✓
	BICGSTAB	✓	✓	✓
	CG	✓	✓	✓
	CGS	✓	✓	✓
	GMRES	✓	✓	✓
Preconditioners	IDR	✓	✓	✓
	CBlock-Jacobi	✓	✓	✓
	ILU/IC		✓	✓
	Parallel ILU/IC	✓	✓	✓
	Parallel ILU/ICT	✓	✓	✓
	Sparse Approximate Inverse	✓	✓	✓



Sustainable software development & CI/CD



Sustainable software development & CI/CD

Private AMD system

Build

build/amd/mvapich2/clang/room45/debug/shared

build/amd/mvapich2/gcc/room45/release/shared

build/amd/rocm/clang/room45/debug/shared

build/amd/openmp/clang/room502/release/sh...

build/clang-cuda101/openmp/gcc/cuda/release...

build/cuda92/openmp/gcc/23/release/shared

build/cuda100/openmp/ch2/gcc/all/debug/shared

build/cuda100/openmp/intel/cuda/debug/static

build/cuda104/openmp/gcc/cuda/debug/shared

build/rocm/cpu/release/static

build/rocm/cpu/release/shared

build/rocm/cpu/level_zero_dgpu/release/shared

build/rocm-cuda-somaid/openmp/gcc/clang/release...

build/rocm-cuda/openmp/clang/cuda/release/shared

build/rocm-cuda/openmp/gcc/cuda/debug/static

build/rocm-cuda/openmp/gcc/clang/release/shared

build/rocm-cuda/openmp/clang/cuda/release/static

build/rocm-cuda/openmp/intel/clang/debug/shared

Test

test/cuda100/openmp/intel/cuda/debug/sh...

Code_quality

no circular deps

sonarqube.cov...

Private NVIDIA system

Private Intel systems

"Quick" overview CI jobs

Private or NHR@KIT systems

Build

build/amd/rocm/clang/room502/release/sh...

build/amd/openmp/clang/room40/release/static

build/amd/openmp/gcc/room502/debug/static

build/clang-cuda101/openmp/clang/cuda/debug/...

build/cuda100/openmp/clang/all/release/static

build/cuda100/openmp/intel/cuda/release/shared

build/cuda100/openmp/intel/cuda_wgmp/release...

build/cuda101/openmp/clang/all/release/static

build/cuda101/openmp/gcc/all/debug/shared

build/cuda102/openmp/clang/gpu/release/static

build/cuda102/openmp/clang/all/debug/shared

build/cuda102/openmp/intel/cuda/debug/static

build/cuda106/openmp/ch2/gcc/cuda/debug/shared

build/cuda106/openmp/clang/cuda/release/static

build/rocm/cpu/release/static

build/rocm-cuda-somaid/openmp/clang/clang/sh...

build/rocm-cuda-somaid/openmp/clang/clang/sh...

build/rocm-cuda-somaid/openmp/clang/clang/sh...

Test

test/cuda100/mvapich2/gcc/cuda/debug/shared

test/cuda100/openmp/clang/cuda/release/static

Code_quality

clang-tidy

secret-build

mya

watch-build

warnings

NHR@KIT GPU partition (with SLURM)

NHR@KIT dedicated CI node

Extra jobs, ran before merging



8

Sustainable software development & CI/CD

NVIDIA

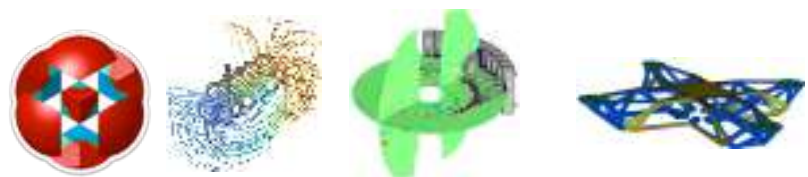
CPU only

Intel

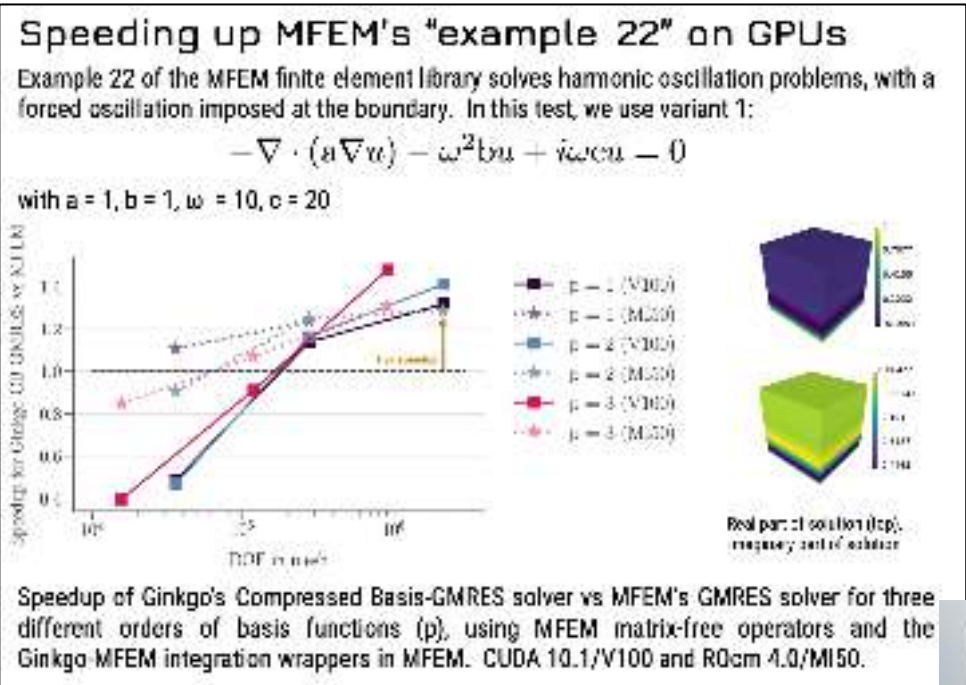
AMD

Build	Test	Deploy
✓ build/cuda90/openmpi/gnu5/llvm39 ↻	✓ test/cuda90/openmpi/gnu5/llvm39 ↻	✓ deploy/cuda90/openmpi/gnu5/llvm39 ↻
✓ build/cuda91/mvapich2/gnu6/llvm40 ↻	✓ test/cuda91/mvapich2/gnu6/llvm40 ↻	✓ deploy/cuda91/mvapich2/gnu6/llvm40 ↻
✓ build/cuda92/mvapich2/gnu7/llvm50/intel2017 ↻	✓ test/cuda92/mvapich2/gnu7/llvm50/intel2017 ↻	✓ deploy/cuda92/mvapich2/gnu7/llvm50/intel2017 ↻
✓ build/cuda100/mvapich2/gnu7/llvm60/intel2018 ↻	✓ test/cuda100/mvapich2/gnu7/llvm60/intel2018 ↻	✓ deploy/cuda100/mvapich2/gnu7/llvm60/intel2... ↻
✓ build/cuda101/openmpi/gnu8/llvm7/intel2019 ↻	✓ test/cuda101/openmpi/gnu8/llvm7/intel2019 ↻	✓ deploy/cuda101/openmpi/gnu8/llvm7/intel2019 ↻
✓ build/cuda101/openmpi/gnu8/llvm11/intel2019 ↻	✓ test/cuda101/openmpi/gnu8/llvm11/intel2019 ↻	✓ deploy/cuda101/openmpi/gnu8/llvm11/intel2019 ↻
✓ build/cuda102/ompi/gnu6/llvm8/intel2019 ↻	✓ test/cuda102/ompi/gnu6/llvm8/intel2019 ↻	✓ deploy/cuda102/ompi/gnu6/llvm8/intel2019 ↻
✓ build/cuda110/mvapich2/gnu9/llvm9/intel2020 ↻	✓ test/cuda110/mvapich2/gnu9/llvm9/intel2020 ↻	✓ deploy/cuda110/mvapich2/gnu9/llvm9/intel2020 ↻
✓ build/cuda114/openmpi/gnu11/llvm12 ↻	✓ test/cuda114/openmpi/gnu11/llvm12 ↻	✓ deploy/cuda114/openmpi/gnu11/llvm12 ↻
✓ build/nocuda/mvapich2/gnu5/llvm39/intel2018 ↻	✓ test/nocuda/mvapich2/gnu5/llvm39/intel2018 ↻	✓ deploy/nocuda/mvapich2/gnu5/llvm39/intel2018 ↻
✓ build/nocuda/openmpi/gnu9/llvm8 ↻	✓ test/nocuda/openmpi/gnu9/llvm8 ↻	✓ deploy/nocuda/openmpi/gnu9/llvm8 ↻
✓ build/oneapi/openmpi ↻	✓ test/oneapi/openmpi ↻	✓ deploy/oneapi/openmpi ↻
✓ build/rocm40/openmpi/gnu5/llvm50 ↻	✓ test/rocm40/openmpi/gnu5/llvm50 ↻	✓ deploy/rocm40/openmpi/gnu5/llvm50 ↻
✓ build/rocm45/mvapich2/gnu8/llvm8 ↻	✓ test/rocm45/mvapich2/gnu8/llvm8 ↻	✓ deploy/rocm45/mvapich2/gnu8/llvm8 ↻
✓ build/rocm502/openmpi/gnu11/llvm11 ↻	✓ test/rocm502/openmpi/gnu11/llvm11 ↻	✓ deploy/rocm502/openmpi/gnu11/llvm11 ↻

Input from the “first customer”



MFEM is a *free, lightweight, scalable* C++ library for finite element methods.



Library core contains architecture-agnostic factuality

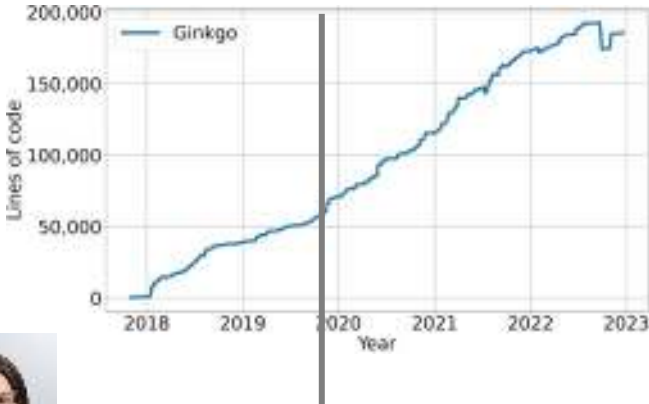
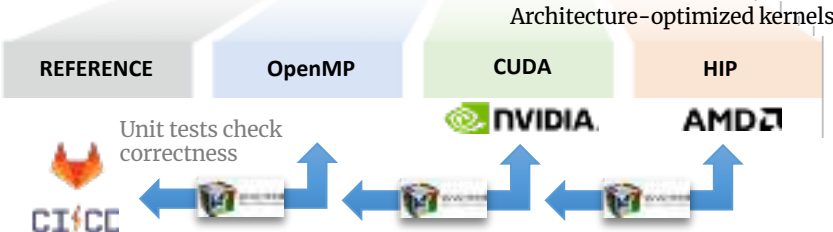
CORE

Infrastructure Algorithms

- Iterative Solvers
- Preconditioners
- ...



Runtime polymorphism selects the right kernel depending on the target architecture



Functionality		OMP	CUDA	HIP
Basic	SpMV	✓	✓	✓
	SpMM	✓	✓	✓
	SpGeMM	✓	✓	✓
Krylov solvers	BICG	✓	✓	✓
	BICGSTAB	✓	✓	✓
	CG	✓	✓	✓
	CGS	✓	✓	✓
	GMRES	✓	✓	✓
Preconditioners	IDR	✓	✓	✓
	(Block-)Jacobi	✓	✓	✓
	ILU/IC	✓	✓	✓
	Parallel ILU/IC	✓	✓	✓
	Parallel ILUT/ICT	✓	✓	✓
	Sparse Approximate Inverse	✓	✓	✓

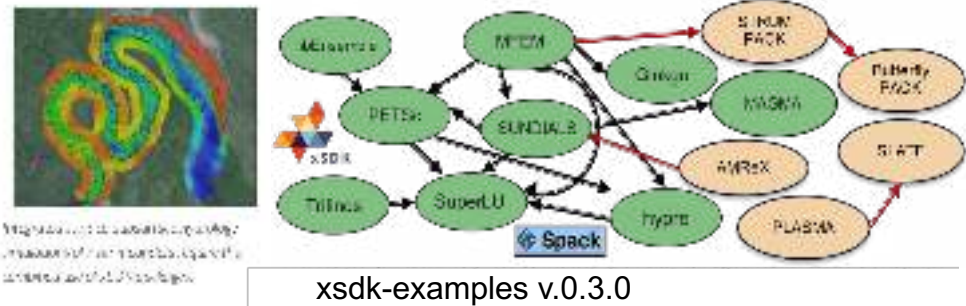


Utilities	On-Device Matrix Assembly	✓	✓	✓
	MC64/RCM reordering	✓	✓	✓
	Wrapping user data	✓	✓	✓



Part of the xSDK effort

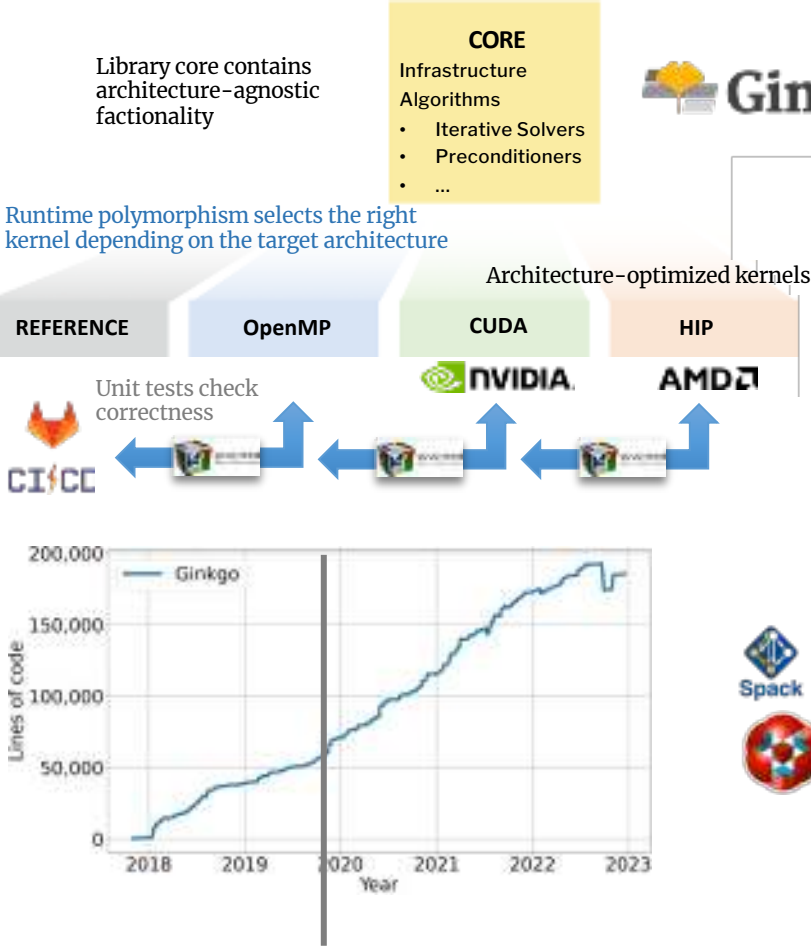
xSDK: Extreme-scale Scientific Software Development Kit



The xSDK provides infrastructure for and interoperability of a **collection of related and complementary software elements**—developed by diverse, independent teams throughout the high-performance computing (HPC) community—that provide the building blocks, tools, models, processes, and related artifacts for rapid and efficient development of high-quality applications.

- November 2022**
- 26 math libraries
 - 2 domain components
 - 16 mandatory xSDK community policies
 - Spack xSDK installer

- xSDK community policies:**
- 16 mandatory policies,
 - 8 recommended policies,
 - 4 Spack variant guidelines
 - Available on Github
- <https://xsdk.info/policies/>

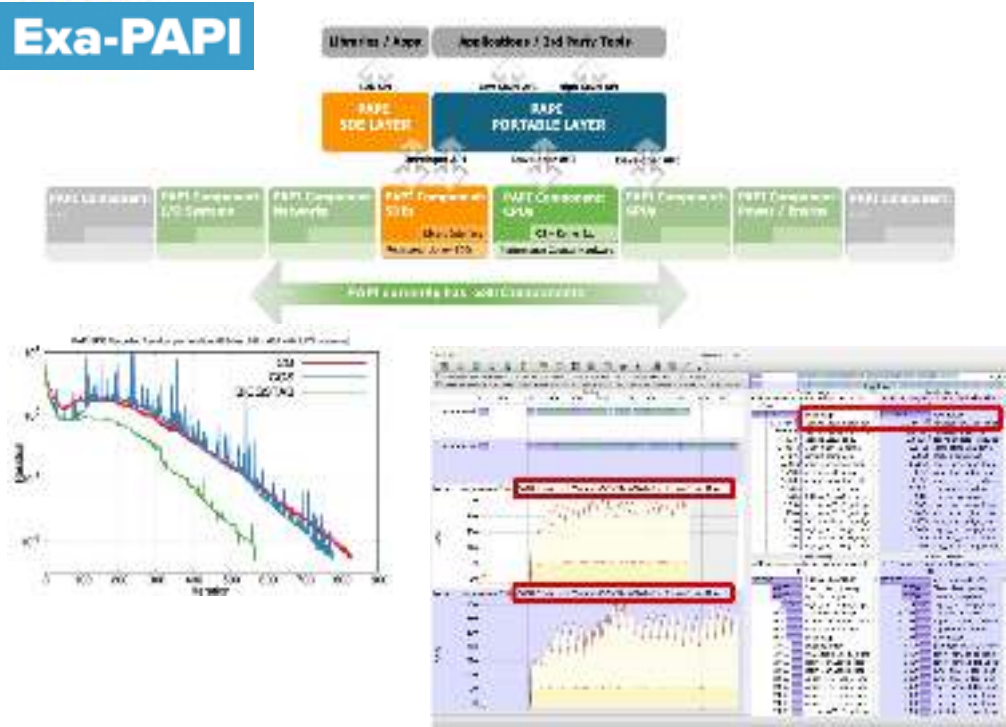


Functionality		OMP	CUDA	HIP
Basic	SpMV	✓	✓	✓
	SpMM	✓	✓	✓
	SpGeMM	✓	✓	✓
Krylov solvers	BICG	✓	✓	✓
	BICGSTAB	✓	✓	✓
	CG	✓	✓	✓
	CQS	✓	✓	✓
	GMRES	✓	✓	✓
Preconditioners	IDR	✓	✓	✓
	(Block-)Jacobi	✓	✓	✓
	ILU/IC	✓	✓	✓
	Parallel ILU/IC	✓	✓	✓
	Parallel ILUT/ICT	✓	✓	✓
Utilities	Sparse Approximate Inverse	✓	✓	✓
	On-Device Matrix Assembly	✓	✓	✓
	MC64/RCM reordering	✓	✓	✓
Wrapping user data		✓	✓	✓



Adding profiling functionality

Exa-PAPI



Library core contains architecture-agnostic functionality

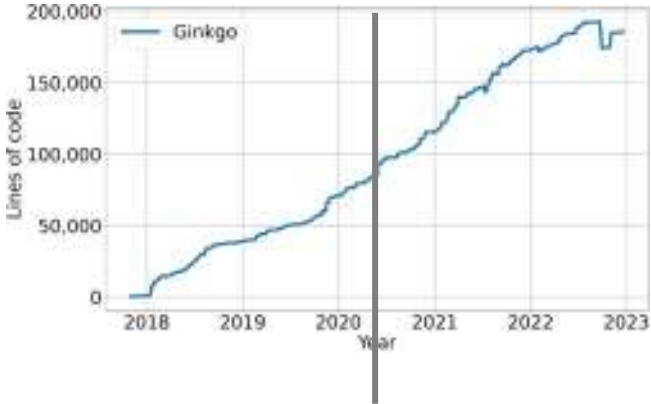
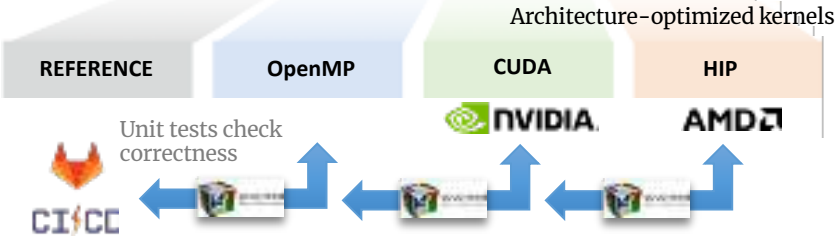
CORE

Infrastructure Algorithms

- Iterative Solvers
- Preconditioners
- ...



Runtime polymorphism selects the right kernel depending on the target architecture



	Functionality	OMP	CUDA	HIP
Basic	SpMV	✓	✓	✓
	SpMM	✓	✓	✓
	SpGeMM	✓	✓	✓
Krylov solvers	BICG	✓	✓	✓
	BICGSTAB	✓	✓	✓
	CG	✓	✓	✓
	CGS	✓	✓	✓
	GMRES	✓	✓	✓
Preconditioners	IDR	✓	✓	✓
	(Block-1)Jacobi	✓	✓	✓
	ILU/IC	✓	✓	✓
	Parallel ILU/IC	✓	✓	✓
	Parallel ILU/IC/CT	✓	✓	✓
	Sparse Approximate Inverse	✓	✓	✓

Utilities	On-Device Matrix Assembly	✓	✓	✓
	MC64/RCM reordering	✓		
	Wrapping user data		✓	
	Logging		✓	
	PAPI counters		✓	



Ginkgo

- ... but also docker image contributions and bug fixes!

The screenshot shows the GitHub repository page for 'fix-cuda-backend-location-fix-2'. The repository is owned by 'fix-cuda-backend-location-fix-2' and has a single commit. The description states that the fix is for the CUDA backend location. The bug fix section mentions that the fix is for the CUDA backend location. The commit history shows a single commit by 'fix-cuda-backend-location-fix-2'.

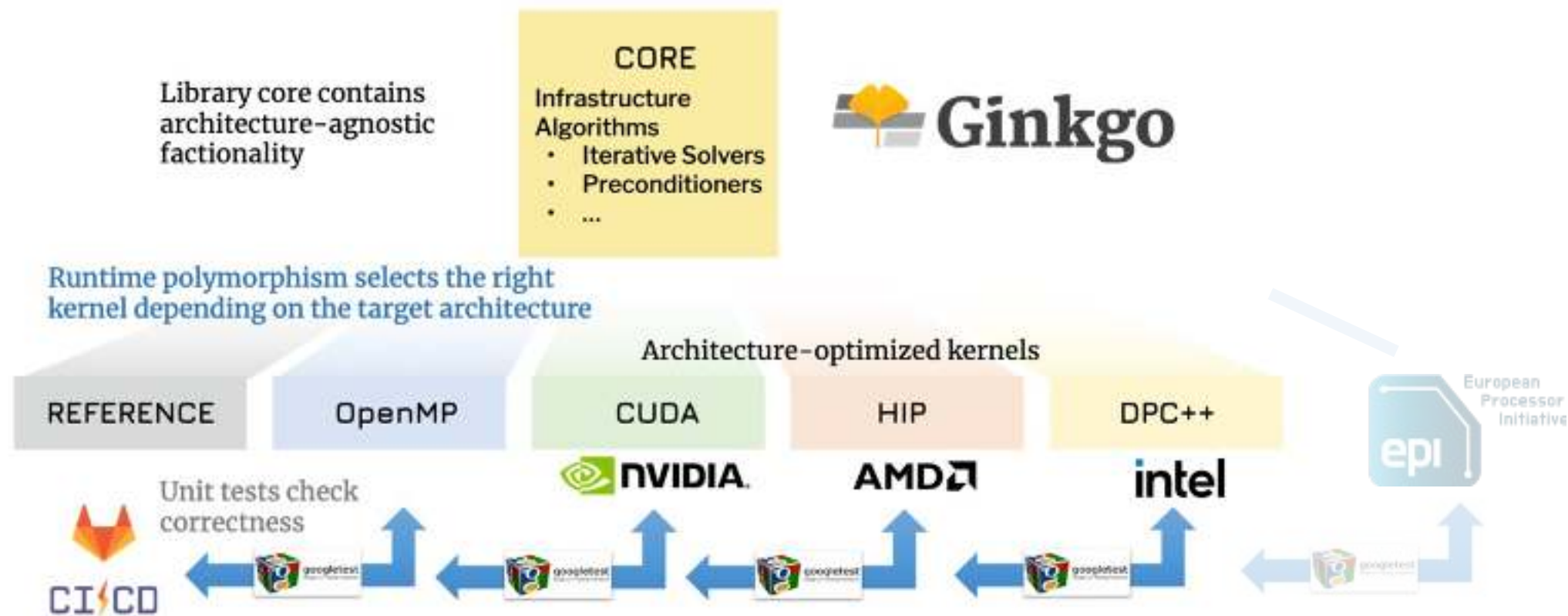
- The options are also required at linking time? Unused in files without kernels?
- Any example of other projects integrating AoT in a CMake setup?

Arcticus at ALCF
Intel Xe-HP GPU (Arctic Sound 2x)
7,680 x2 Cores @900 MHz
FP32 peak 13,820 x2 GFlop/s

Transitional system to Aurora
Exascale supercomputer at ANL



Portability as central design principle



Focus efforts as lightweight tool in ECP to address challenges



Focus efforts

- Mixed precision
- batched



- Address recent hardware trends (tensor cores, etc.)
- Address hardware requirements

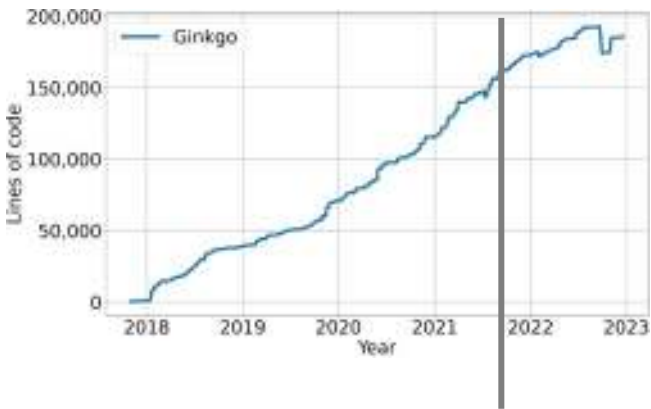
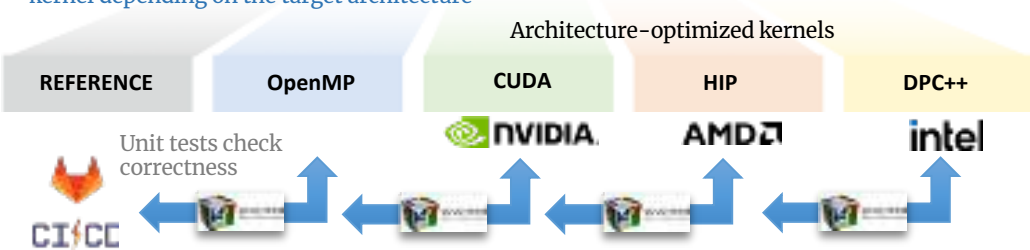
Library core contains architecture-agnostic functionality

CORE
Infrastructure Algorithms

- Iterative Solvers
- Preconditioners
- ...



Runtime polymorphism selects the right kernel depending on the target architecture



oneAPI Industry Collaboration with bi-weekly meetings

Functionality		OMP	CUDA	HIP	DPC++
Basic	SpMV	✓	✓	✓	✓
	SpMM	✓	✓	✓	✓
	SpGeMM	✓	✓	✓	✓
Krylov solvers	BICG	✓	✓	✓	✓
	BICGSTAB	✓	✓	✓	✓
	CG	✓	✓	✓	✓
	CGS	✓	✓	✓	✓
	GMRES	✓	✓	✓	✓
Preconditioners	IDR	✓	✓	✓	✓
	(Block-)Jacobi	✓	✓	✓	✓
	ILU/IC	✓	✓	✓	✓
	Parallel ILU/IC	✓	✓	✓	✓
	Parallel ILUT/ICT	✓	✓	✓	✓
	Sparse Approximate Inverse	✓	✓	✓	✓

Utilities	On-Device Matrix Assembly	✓	✓	✓	✓
	MC64/RCM reordering	✓			
	Wrapping user data		✓		
	Logging		✓		
	PAPI counters		✓		



Mixed precision focus effort



Focus efforts

- Mixed precision
- batched



- Address recent hardware trends (tensor cores, etc.)
- Address hardware requirements

Advances in Mixed Precision Algorithms: 2021 Edition

by the ECP Multiprecision Effort Team (Lead: Hartwig Anzt)

Ahmad Abdelfattah, Hartwig Anzt, Alan Ayala, Erik G. Boman, Erin Carson, Sebastien Cayrols, Terry Cojean, Jack Dongarra, Rob Falgout, Mark Gates, Thomas Grützmacher, Nicholas J. Higham, Scott B. Kruger, Sherry Li, Neil Lindquist, Yang Liu, Jennifer Lee, Piotr Luszczek, Pratik Nayak, Daniel Osei-Kuffuor, Sri Pranesh, Sivasankaran Rajamanickam, Tobias Ribizel, Barry Smith, Kasia Swirydowicz, Stephen Thomas, Stanimire Tomov, Yaohung M. Tsai, Ichi Yamazaki, Urike Meier Yang

Available access Research article First published online March 18, 2021

A survey of numerical linear algebra methods utilizing mixed-precision arithmetic

Ahmad Abdelfattah, Hartwig Anzt, and Urike Meier Yang

Volume 35, Issue 4 | <https://doi.org/10.1177/10646420211003311>

Library core contains architecture-agnostic factuality

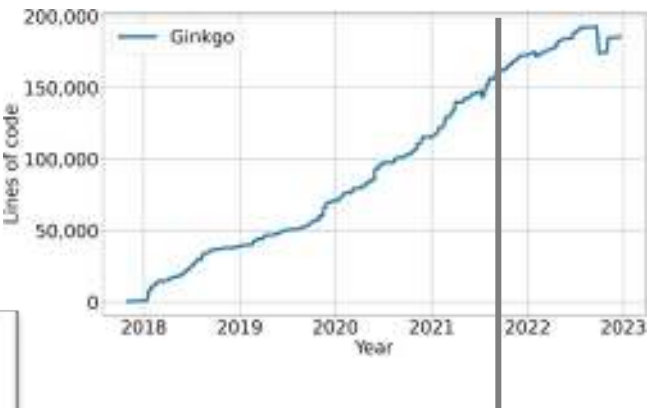
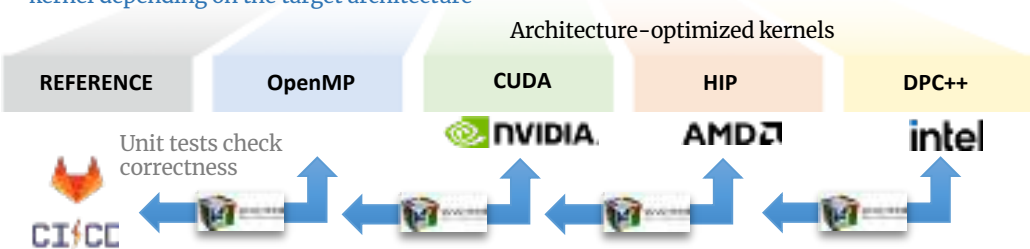
CORE

Infrastructure Algorithms

- Iterative Solvers
- Preconditioners
- ...



Runtime polymorphism selects the right kernel depending on the target architecture



oneAPI Industry Collaboration with bi-weekly meetings

Functionality	OMP	CUDA	HIP	DPC++
Basic				
SpMV	✓	✓	✓	✓
SpMM	✓	✓	✓	✓
SpGeMM	✓	✓	✓	✓
Krylov solvers				
BICG	✓	✓	✓	✓
BICGSTAB	✓	✓	✓	✓
CG	✓	✓	✓	✓
CQS	✓	✓	✓	✓
GMRES	✓	✓	✓	✓
IDR	✓	✓	✓	✓
(Block-)Jacobi	✓	✓	✓	✓
Preconditioners				
ILU/IC	✓	✓	✓	✓
Parallel ILU/IC	✓	✓	✓	✓
Parallel ILUT/ICT	✓	✓	✓	✓
Sparse Approximate Inverse	✓	✓	✓	✓

Utilities	On-Device Matrix Assembly	✓	✓	✓	✓
	MC64/RCM reordering	✓			
	Wrapping user data		✓		
	Logging		✓		
	PAPI counters		✓		



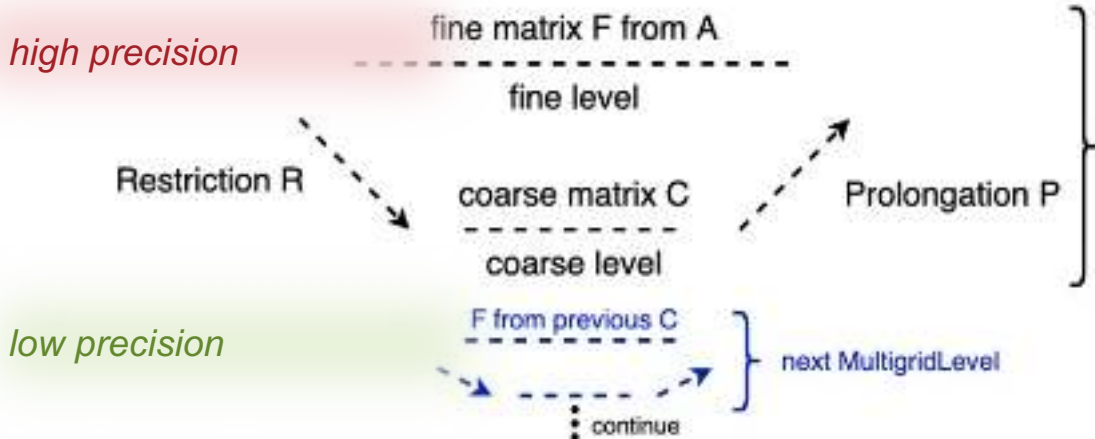
Mixed precision AMG on GPUs

- **Preconditioning iterative solvers**

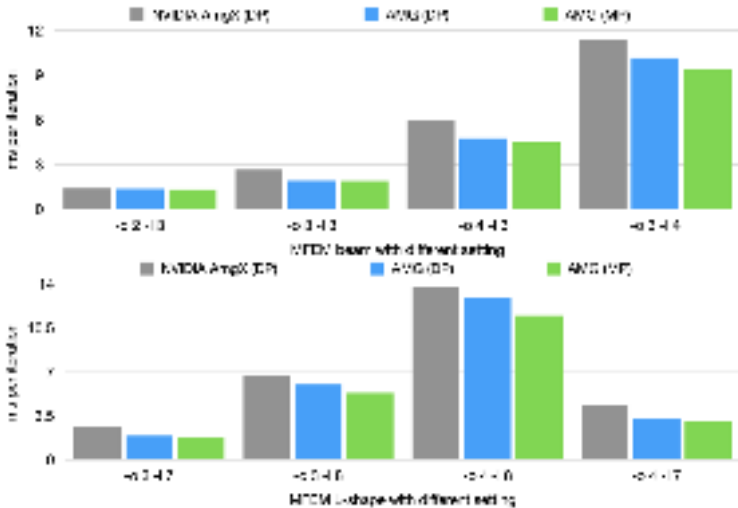
- Idea: Approximate inverse of system matrix to make the system “easier to solve”: $P^{-1} \approx A^{-1}$
and solve $Ax = b \iff P^{-1}Ax = P^{-1}b \iff \tilde{A}x = \tilde{b}$

- **Mixed Precision Multigrid Preconditioner**

```
1 multigrid::build()
2   .with_max_levels(10u) // equal to NVIDIA/AMGX 11 max levels
3   .with_min_coarse_row(64u)
4   .with_pre_smoother(sm, sm_f)
5   .with_mg_level(pgm, pgm_f)
6   .with_level_selector(
7     [](const size_type level, const LinOp*) -> size_type {
8       // Only the first level is generated by MultigridLevel(double).
9       // The subsequent levels are generated by MultigridLevel(float)
10      return level >= 1 ? 1 : 0;
11    })
12   .with_coarest_solver(coarest_solver_f)
```



Mike Tsai



Mixed precision AMG on GPUs



- Focus efforts
- Mixed precision
 - batched



Library core contains architecture-agnostic functionality

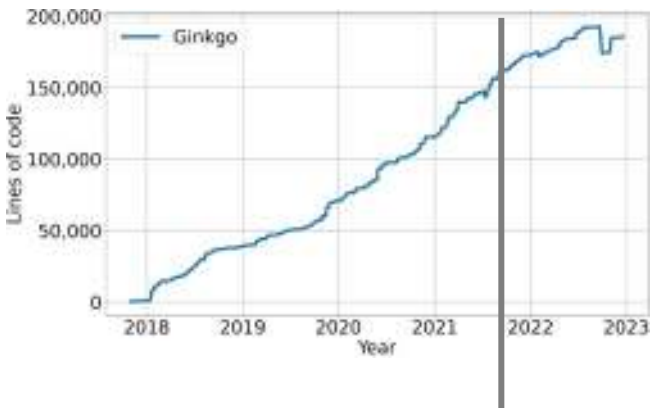
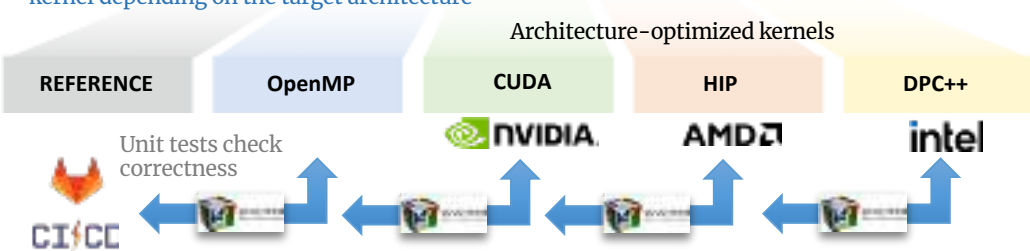
CORE

Infrastructure Algorithms

- Iterative Solvers
- Preconditioners
- ...



Runtime polymorphism selects the right kernel depending on the target architecture



oneAPI Industry Collaboration with bi-weekly meetings

Functionality		OMP	CUDA	HIP	DPC++
Basic	SpMV	✓	✓	✓	✓
	SpMM	✓	✓	✓	✓
	SpGeMM	✓	✓	✓	✓
	BICG	✓	✓	✓	✓
Krylov solvers	BICGSTAB	✓	✓	✓	✓
	CG	✓	✓	✓	✓
	CGS	✓	✓	✓	✓
	GMRES	✓	✓	✓	✓
Preconditioners	IDR	✓	✓	✓	✓
	(Block-1) Jacobi	✓	✓	✓	✓
	ILU/IC	✓	✓	✓	✓
	Parallel ILU/IC	✓	✓	✓	✓
	Parallel ILUT/ICT	✓	✓	✓	✓
	Sparse Approximate Inverse	✓	✓	✓	✓

AMG	AMG preconditioner	✓	✓	✓	✓
	AMS solver	✓	✓	✓	✓
	Parallel Graph Match	✓	✓	✓	✓

Utilities	On-Device Matrix Assembly	✓	✓	✓	✓
	MC64/RCM reordering	✓			
	Wrapping user data		✓		
	Logging		✓		
	PAPI counters		✓		

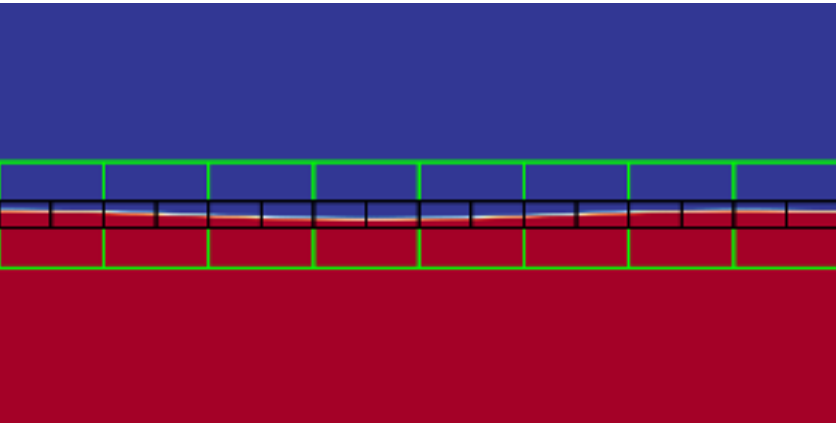


Batched focus effort – Combustion Simulations

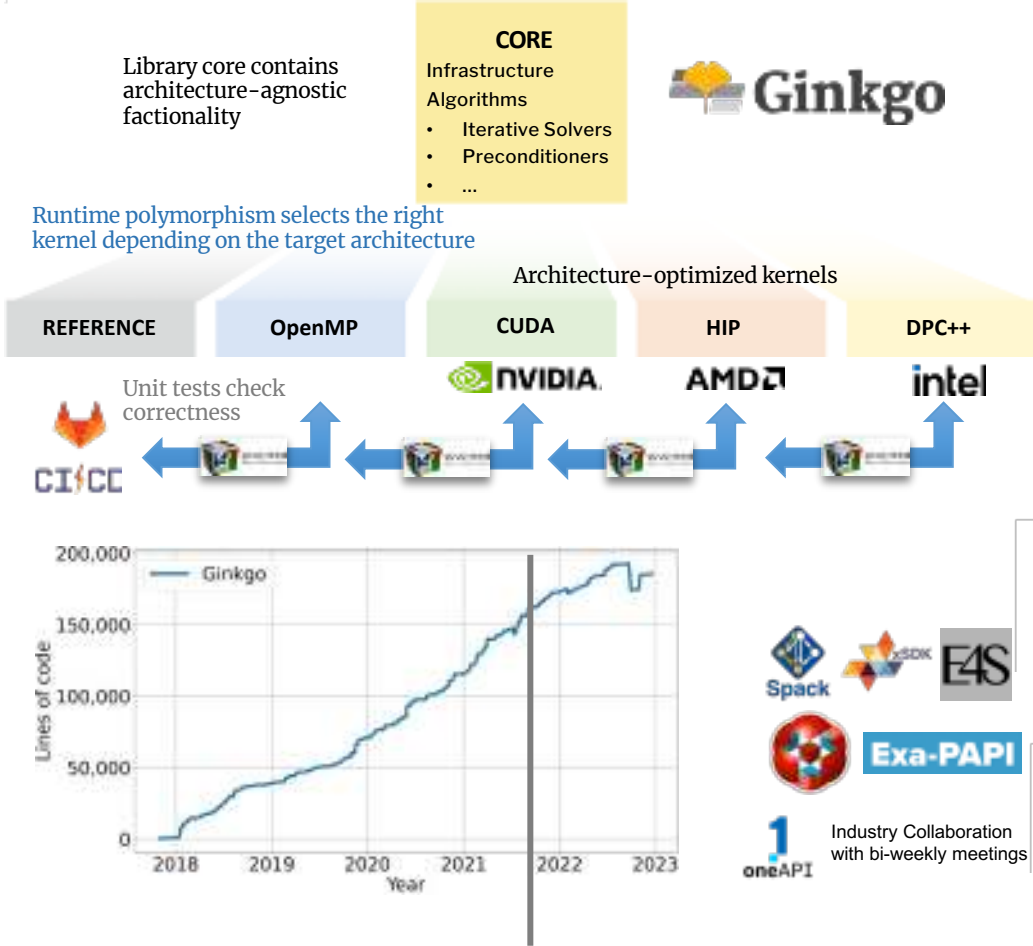
Batched iterative solvers for SUNDIALS / PeleLM

PeleLM is a parallel, adaptive mesh refinement (AMR) code that solves the reacting Navier-Stokes equations in the low Mach number regime. The core libraries for managing the subcycling AMR grids and communication are found in the [AMReX source code](https://amrex-combustion.github.io/PeleLM/overview.html).

<https://amrex-combustion.github.io/PeleLM/overview.html>



Problem	Size	Non-zeroes (A)	Non-zeroes (L+U)
amrzone_1u	54	2,332 (90%)	2,754 (94%)
amr19	22	438 (90%)	442 (91%)
gr12	33	978 (90%)	1,018 (93%)
gr30	54	2,560 (88%)	2,860 (98%)
amrzone	144	6,135 (90%)	20,307 (98%)
amrzone	10	91 (91%)	91 (91%)



Functionality	OMP	CUDA	HIP	DPC++
Basic				
SpMV	✓	✓	✓	✓
SpMM	✓	✓	✓	✓
SpGeMM	✓	✓	✓	✓
Krylov solvers				
BICG	✓	✓	✓	✓
BICGSTAB	✓	✓	✓	✓
CG	✓	✓	✓	✓
CQS	✓	✓	✓	✓
GMRES	✓	✓	✓	✓
IDR	✓	✓	✓	✓
(Block-)Jacobi	✓	✓	✓	✓
ILU/IC	✓	✓	✓	✓
Parallel ILU/IC	✓	✓	✓	✓
Parallel ILU/ICT	✓	✓	✓	✓
Sparse Approximate Inverse	✓	✓	✓	✓

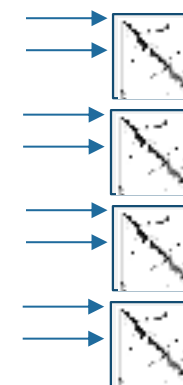
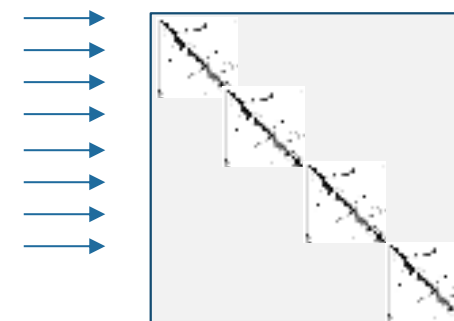
AMG preconditioner	✓	✓	✓	✓
AMG solver	✓	✓	✓	✓
Parallel Graph Match	✓	✓	✓	✓

On-Device Matrix Assembly	✓	✓	✓	✓
MC64/RCM reordering	✓			
Wrapping user data		✓		
Logging		✓		
PAPI counters		✓		



Batched focus effort – Combustion Simulations

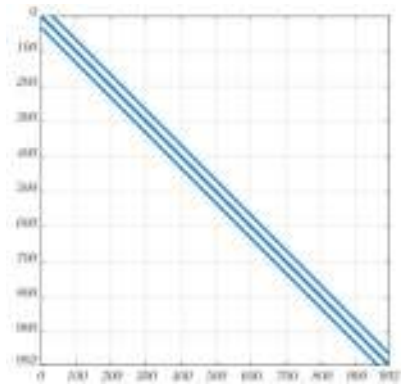
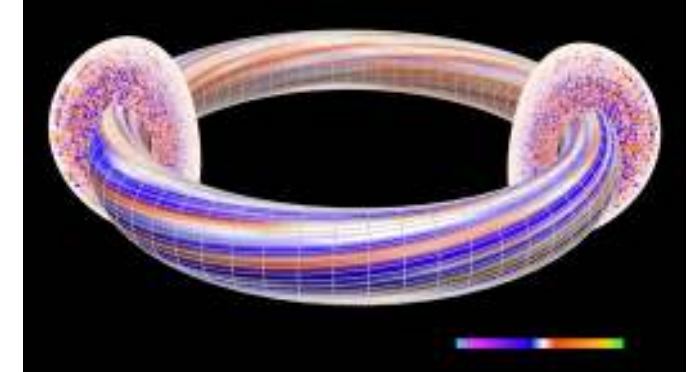
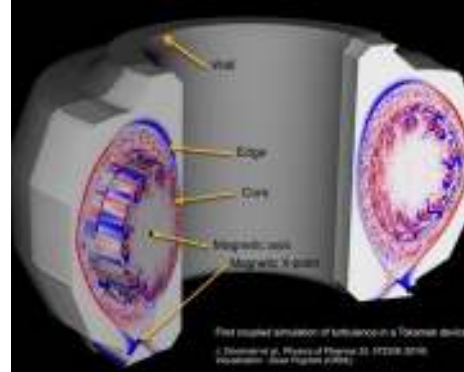
- Many sparse problems of medium size have to be solved concurrently.
 - $\sim 50 - 2,000$ unknowns, $< 50\%$ dense;
 - All sparse systems may share the same sparsity pattern;
 - An approximate solution may be acceptable (e.g., inside a non-linear solver);
- One solution is to arrange the individual systems on the main diagonal of one large system.
 - Convergence determined by the “hardest” problem;
 - No reuse of sparsity pattern information;
 - Global synchronization points;
- Better approach: design batched iterative solve functionality that solves all problems concurrently.
 - Problem-dependent convergence accounted for;
 - No global synchronization;
 - Reuse of sparsity pattern information;



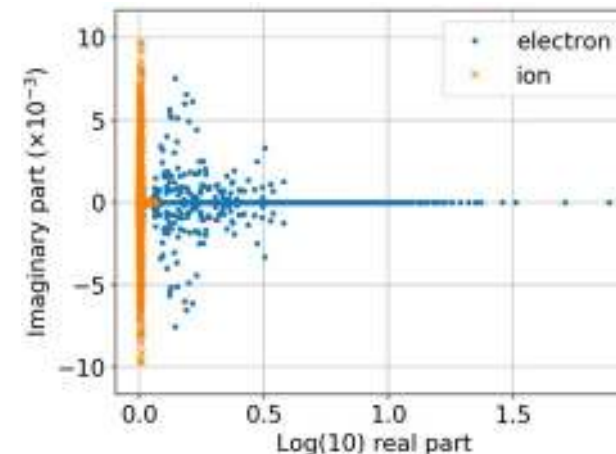
Batched focus effort – Fusion Plasma Simulations

XGC is a gyrokinetic particle-in-cell code, which specializes in the simulation of the edge region of magnetically confined thermonuclear fusion plasma. The simulation domain can include the magnetic separatrix, magnetic axis and the biased material wall. *XGC* can run in total-delta-f, and conventional delta-f mode. The ion species are always gyrokinetic except for ETG simulation. Electrons can be adiabatic, massless fluid, driftkinetic, or gyrokinetic.

Source: https://xgc.pppl.gov/html/general_info.html

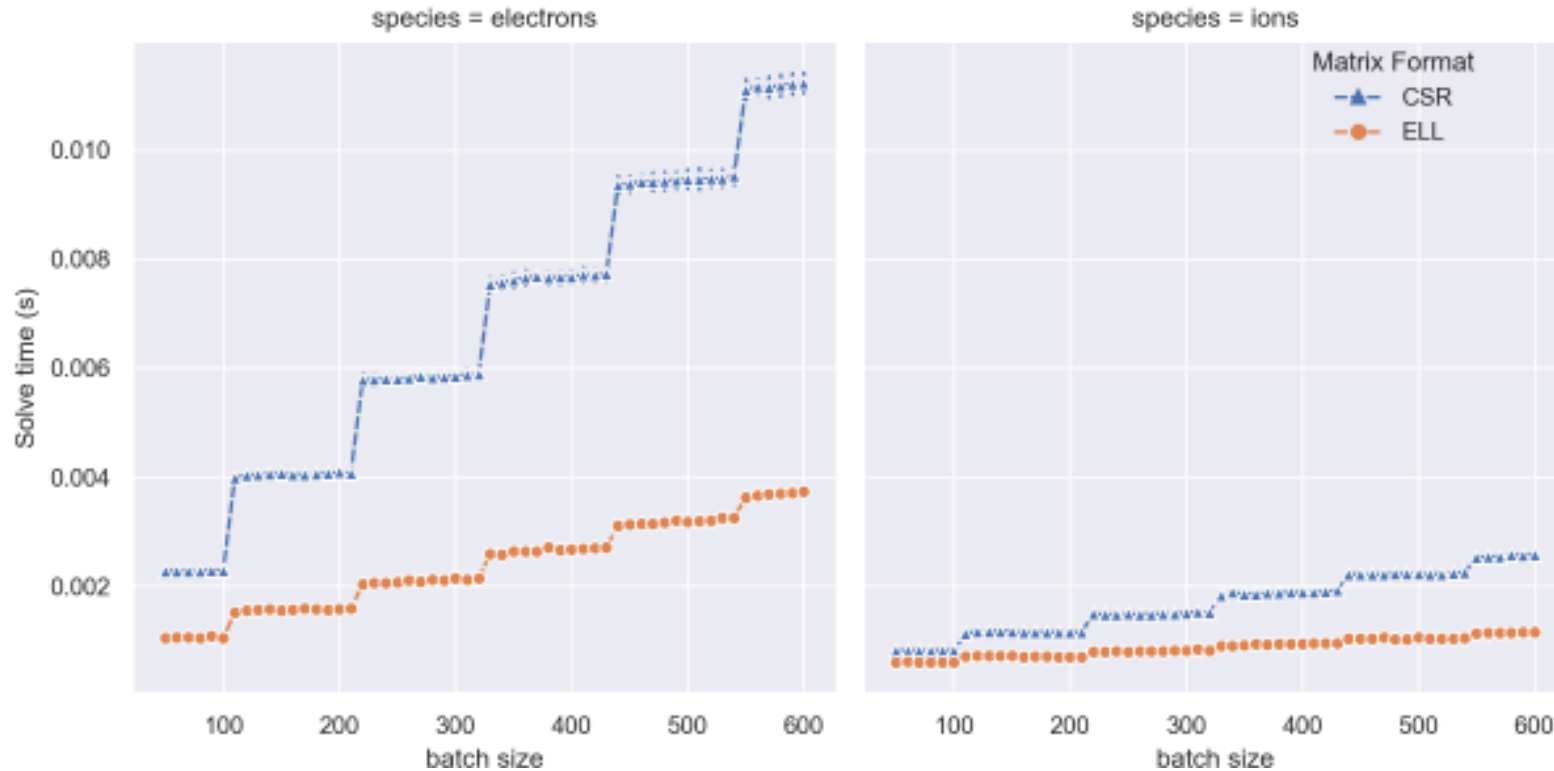


- Two species
- Ions easy to solve
- Electrons hard to solve
- Banded matrix structure
- Non-symmetric, need BiCGSTAB
- $n \sim 1,000$
- $n_z \sim 9,000$



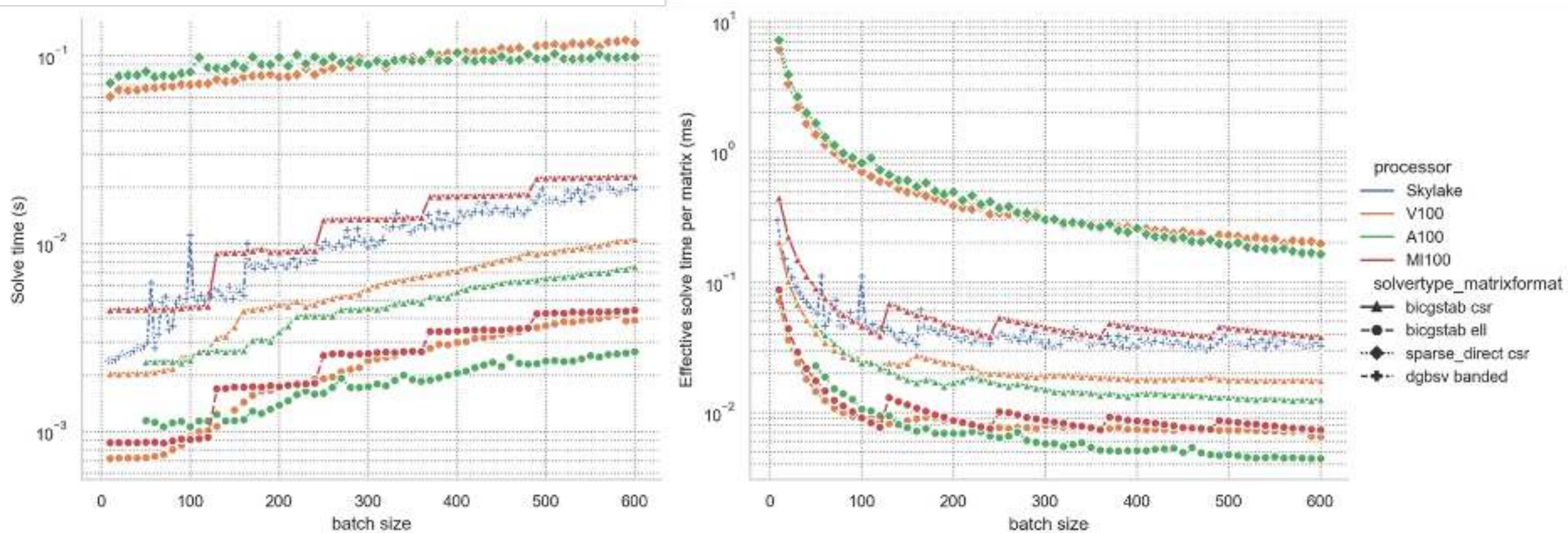
Batched focus effort – Fusion Plasma Simulations

NVIDIA A100 GPU



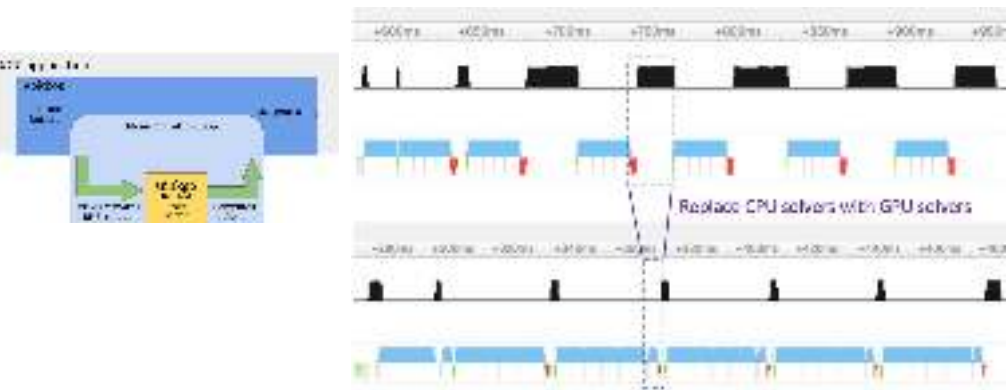
- Ions easy to solve
- Electrons hard to solve
- ELL format more suitable
- “step pattern” for runtime filling up multiprocessors

Batched focus effort – Fusion Plasma Simulations



Aditya Kashi, Pratik Nayak, Dhruva Kulkarni, Aaron Scheinberg, Paul Lin, and Hartwig Anzt. **Batched sparse iterative solvers on gpu for the collision operator for fusion plasma simulations.** In *2022 IEEE International Parallel and Distributed Processing Symposium (IPDPS)*, pages 157–167. IEEE, 2022.

Batched focus effort – Fusion Plasma Simulations



XGC collision operator solve LAPACK vs. Ginkgo: XGC pe459_d3d_EM_healload test case

- XGC pe459_d3d_EM_healload (Aaron's test case; used for Summit, Perlmutter and Crusher scaling studies)
- Preliminary study on 32 nodes of Perlmutter (128 A100s)
 - 2 polaroid planes (216k nodes per plane); 22.4M ptxGPU, 89.6M ptxnode (ptx_num=700k)
 - Ran 20 time steps; collisions calculated every other time step

Time step	Time	GPU
MAIN_LOOP	17.5s	1.0s
COLLISION_OPERATOR	17.5s	1.0s
COLLISION_OPERATOR_GPU	17.5s	1.0s
COLLISION_OPERATOR_GPU_HEALLOAD	17.5s	1.0s
COLLISION_OPERATOR_GPU_HEALLOAD_GPU	17.5s	1.0s
COLLISION_OPERATOR_GPU_HEALLOAD_GPU_HEALLOAD	17.5s	1.0s
COLLISION_OPERATOR_GPU_HEALLOAD_GPU_HEALLOAD_GPU	17.5s	1.0s

With CPU LAPACK debug

- COLLISION_OPERATOR_STEP_SOLVE reduced from 17.5s to 1.0s (94.3% reduction)
- COLLISION_OPERATOR_STEP_SOLVE reduced from 17.5s to 1.0s (94.3% reduction)
- COLLISION_OPERATOR_STEP_SOLVE reduced from 17.5s to 1.0s (94.3% reduction)
- COLLISION_OPERATOR_STEP_SOLVE reduced from 17.5s to 1.0s (94.3% reduction)

Replacing CPU LAPACK debug by GPU Ginkgo

- COLLISION_OPERATOR_STEP_SOLVE reduced from 17.5s to 1.0s (94.3% reduction)
- COLLISION_OPERATOR_STEP_SOLVE reduced from 17.5s to 1.0s (94.3% reduction)
- COLLISION_OPERATOR_STEP_SOLVE reduced from 17.5s to 1.0s (94.3% reduction)
- COLLISION_OPERATOR_STEP_SOLVE reduced from 17.5s to 1.0s (94.3% reduction)

XGC collision operator solve performed on GPU

© Doug Kothe

CORE
Infrastructure Algorithms
• Iterative Solvers
• Preconditioners
• ...

Library core contains architecture-agnostic functionality

Runtime polymorphism selects the right kernel depending on the target architecture

Architecture-optimized kernels

REFERENCE OpenMP CUDA HIP DPC++

Unit tests check correctness

CI/CD

Lines of code

Year

2018 2019 2020 2021 2022 2023

oneAPI

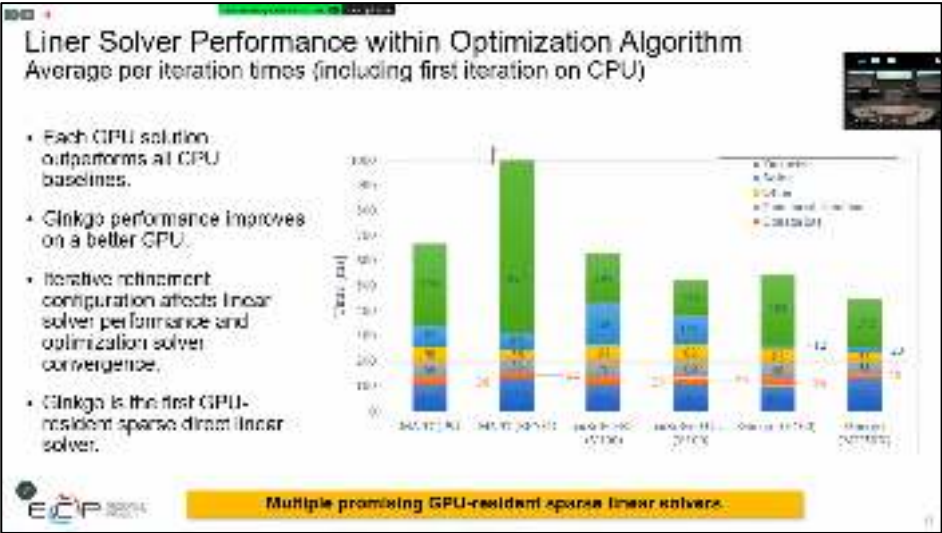
Industry Collaboration with bi-weekly meetings

XGC

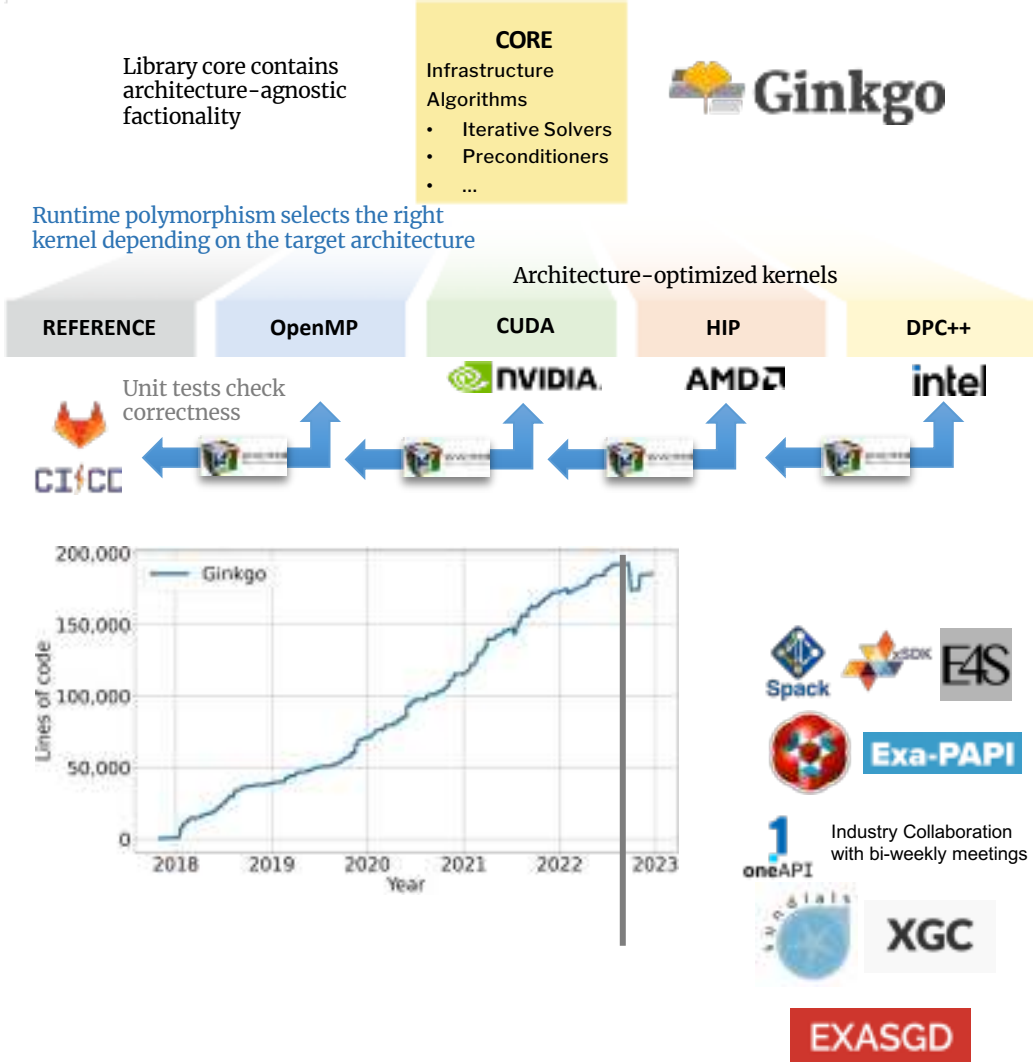
Functionality	OMP	CUDA	HIP	DPC++
Basic				
SpMV	✓	✓	✓	✓
SpMM	✓	✓	✓	✓
SpGEMM	✓	✓	✓	✓
Krylov solvers				
BICG	✓	✓	✓	✓
BICGSTAB	✓	✓	✓	✓
CG	✓	✓	✓	✓
CGS	✓	✓	✓	✓
GMRES	✓	✓	✓	✓
IDR	✓	✓	✓	✓
(Block-)Jacobi	✓	✓	✓	✓
ILU/IC	✓	✓	✓	✓
Parallel ILU/IC	✓	✓	✓	✓
Parallel ILUT/ICT	✓	✓	✓	✓
Sparse Approximate Inverse	✓	✓	✓	✓
Batched				
Batched BICGSTAB	✓	✓	✓	✓
Batched CG	✓	✓	✓	✓
Batched GMRES	✓	✓	✓	✓
Batched ILU	✓	✓	✓	✓
Batched ISAI	✓	✓	✓	✓
Batched Jacobi	✓	✓	✓	✓
AMG				
AMG preconditioner	✓	✓	✓	✓
AMG solver	✓	✓	✓	✓
Parallel Graph Match	✓	✓	✓	✓
Utilities				
On-Device Matrix Assembly	✓	✓	✓	✓
MC64/RCM reordering	✓	✓	✓	✓
Wrapping user data	✓	✓	✓	✓
Logging	✓	✓	✓	✓
PAPI counters	✓	✓	✓	✓



Sparse direct solvers for power grid simulations



EXASGD © Slaven Peles



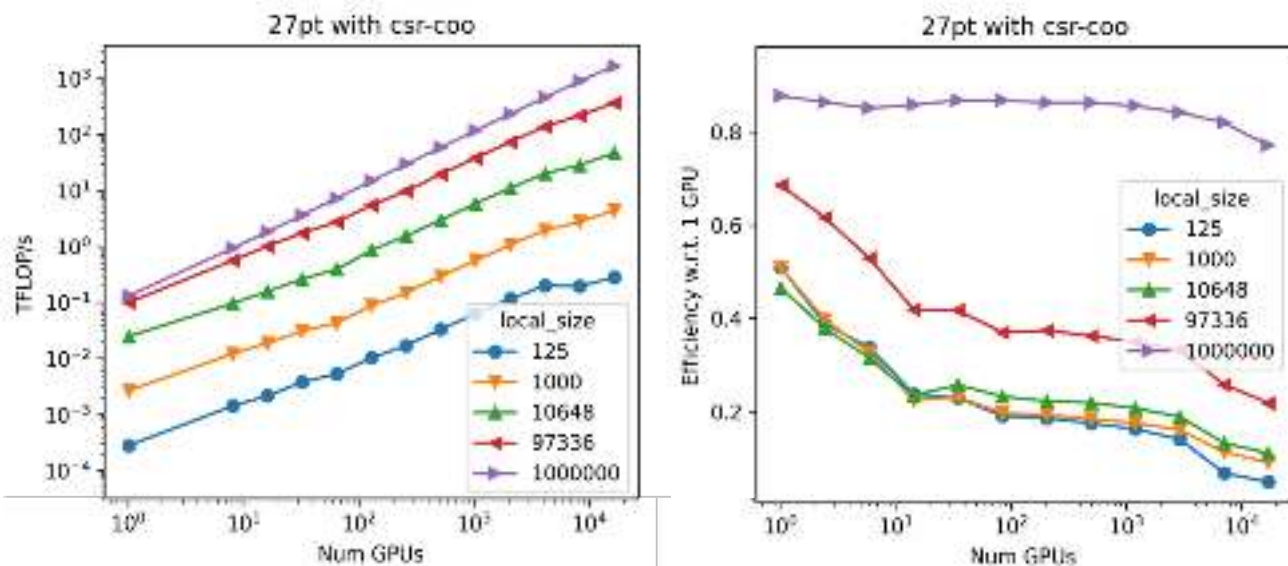
Functionality	OMP	CUDA	HIP	DPC++
Basic				
SpMV	✓	✓	✓	✓
SpMM	✓	✓	✓	✓
SpGeMM	✓	✓	✓	✓
Krylov solvers				
BICG	✓	✓	✓	✓
BICGSTAB	✓	✓	✓	✓
CG	✓	✓	✓	✓
CGS	✓	✓	✓	✓
GMRES	✓	✓	✓	✓
IDR	✓	✓	✓	✓
(Block-)Jacobi	✓	✓	✓	✓
Preconditioners				
ILU/IC	✓	✓	✓	✓
Parallel ILU/IC	✓	✓	✓	✓
Parallel ILUT/ICT	✓	✓	✓	✓
Sparse Approximate Inverse	✓	✓	✓	✓
Batched				
Batched BICGSTAB	✓	✓	✓	✓
Batched CG	✓	✓	✓	✓
Batched GMRES	✓	✓	✓	✓
Batched ILU	✓	✓	✓	✓
Batched ISAI	✓	✓	✓	✓
Batched Jacobi	✓	✓	✓	✓
AMG				
AMG preconditioner	✓	✓	✓	✓
AMG solver	✓	✓	✓	✓
Parallel Graph Match	✓	✓	✓	✓
Sparse direct				
Symbolic Cholesky	✓	✓	✓	✓
Numeric Cholesky	✓	✓	✓	✓
Symbolic LU	✓	✓	✓	✓
Numeric LU	✓	✓	✓	✓
Sparse TRSV	✓	✓	✓	✓
On-Device Matrix Assembly	✓	✓	✓	✓
MC64/RCM reordering	✓	✓	✓	✓
Wrapping user data	✓	✓	✓	✓
Logging	✓	✓	✓	✓
PAPI counters	✓	✓	✓	✓



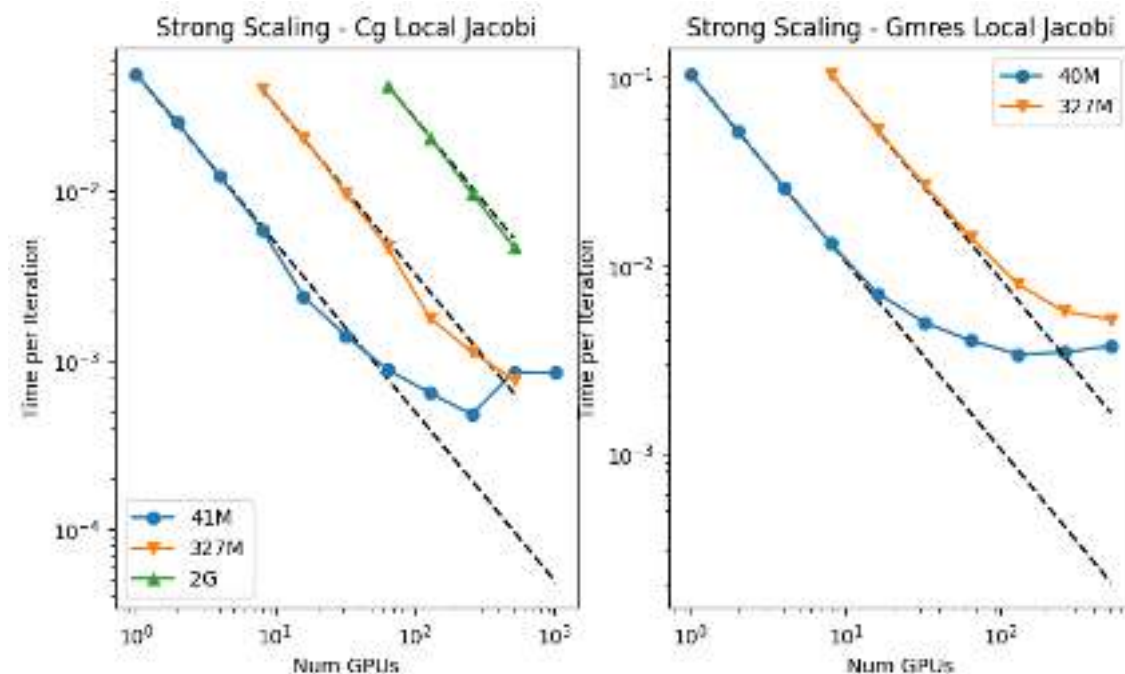
Distributed runs on Frontier (Cray + AMD MI250 GPUs)

Weak scaling: problem size increases with parallel resources

Weak scaling up to 16k GCDs (8k GPUs)

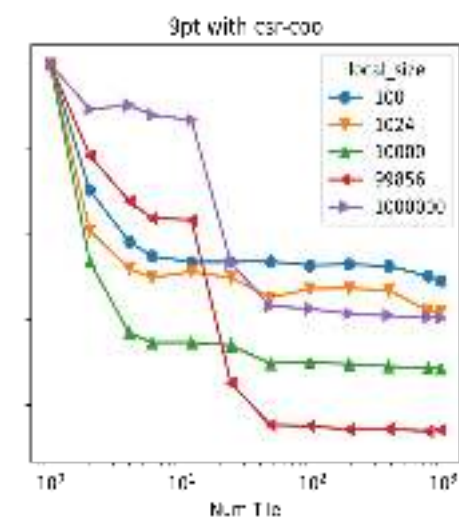
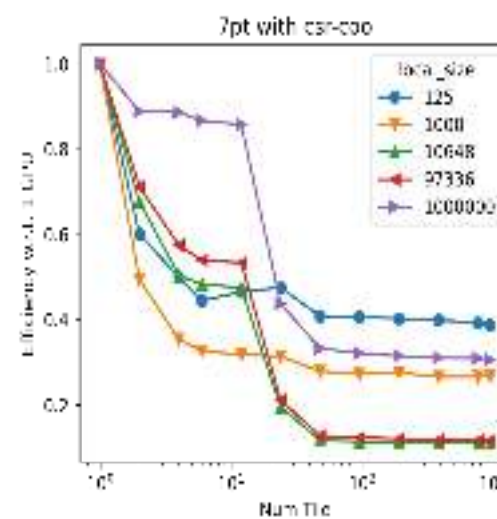
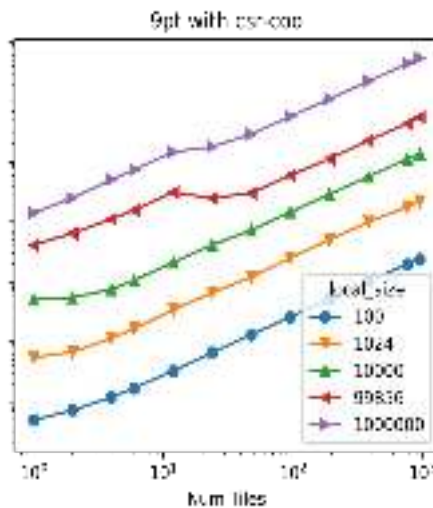
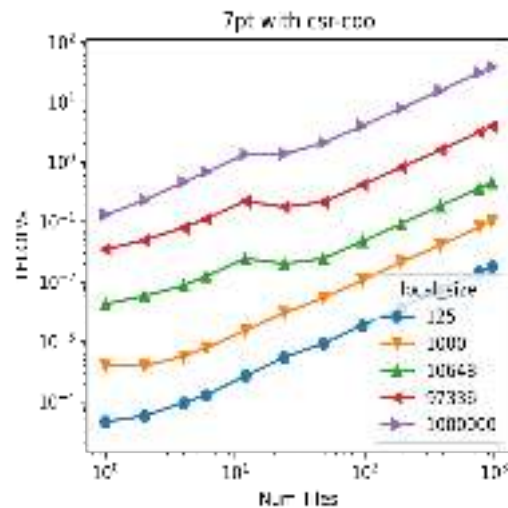
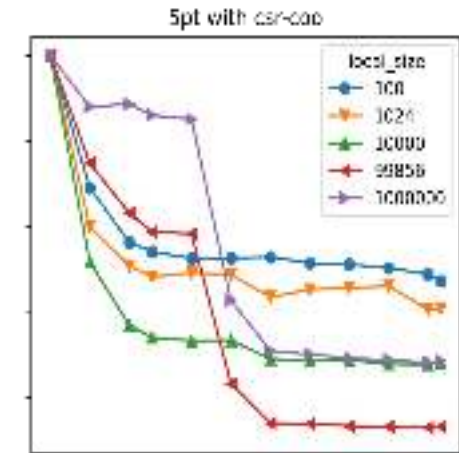
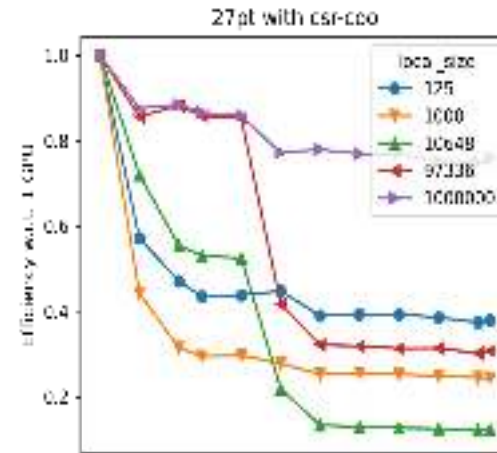
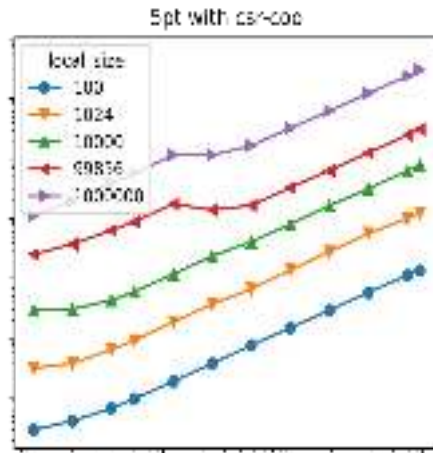
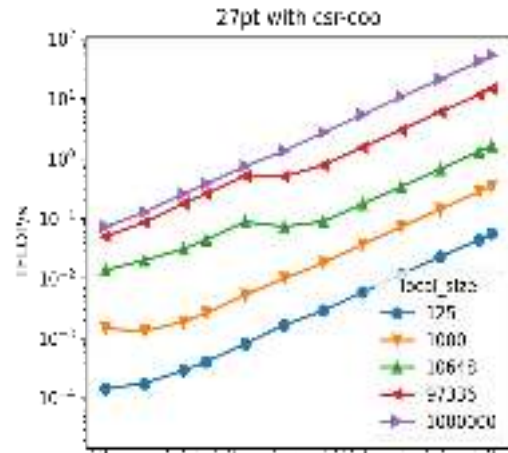


Strong scaling: problem size constant



Distributed runs on Sunspot (Intel PVCA GPUs)

Weak scaling: problem size increases with parallel resources



“Now” – Near completion of ECP

- Sustainable software design ready for the addition of new backends.

- EuroHPC Project MICROCARD uses Ginkgo



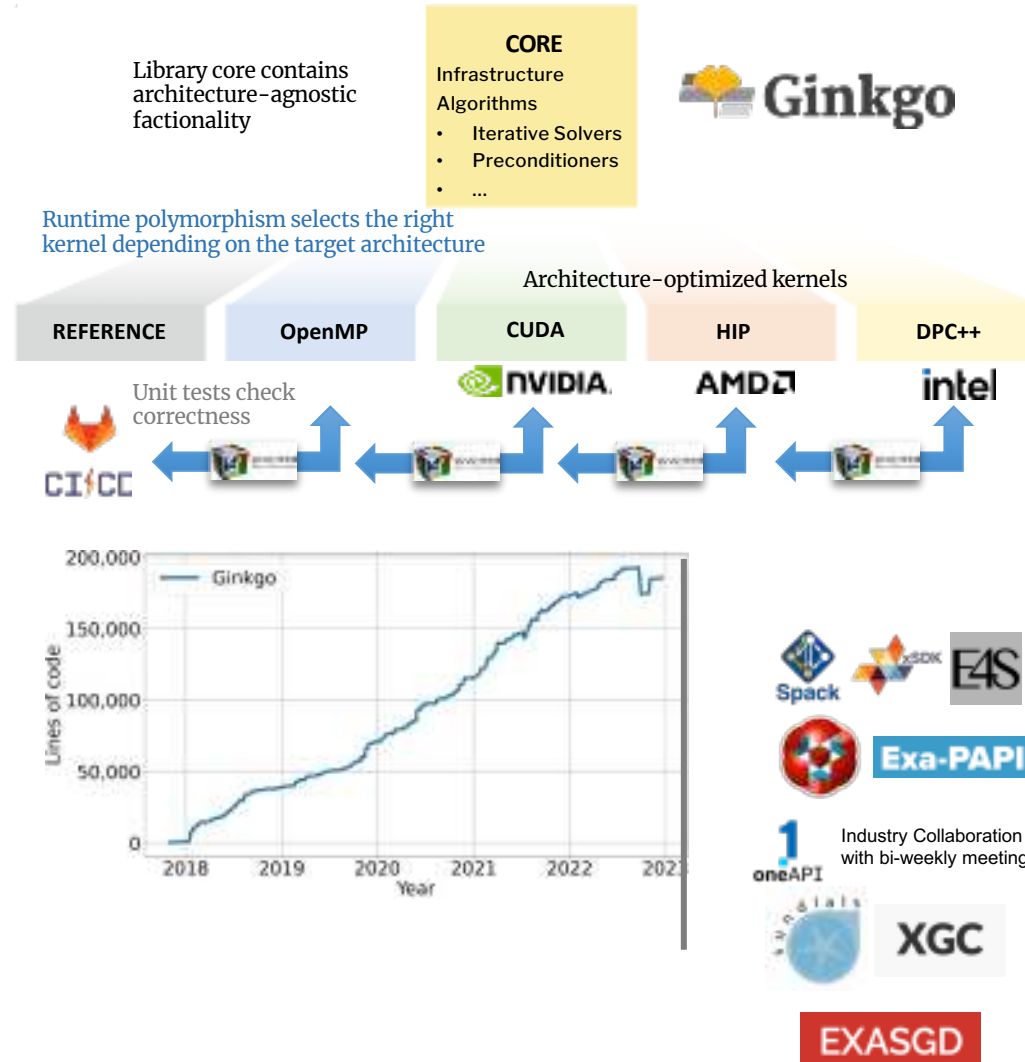
- BMBF PDExa project uses Ginkgo



- BMBF ExaSIM project uses Ginkgo



<https://exasim-project.com>



Functionality	OMP	CUDA	HIP	DPC++
Basic				
SpMV	✓	✓	✓	✓
SpMM	✓	✓	✓	✓
SpGeMM	✓	✓	✓	✓
Krylov solvers				
BICG	✓	✓	✓	✓
BICGSTAB	✓	✓	✓	✓
CG	✓	✓	✓	✓
CGS	✓	✓	✓	✓
GMRES	✓	✓	✓	✓
IDR	✓	✓	✓	✓
(Block-)Jacobi	✓	✓	✓	✓
Preconditioners				
ILU/IC	✓	✓	✓	✓
Parallel ILU/IC	✓	✓	✓	✓
Parallel ILUT/ICT	✓	✓	✓	✓
Sparse Approximate Inverse	✓	✓	✓	✓
Batched BICGSTAB	✓	✓	✓	✓
Batched CG	✓	✓	✓	✓
Batched GMRES	✓	✓	✓	✓
Batched ILU	✓	✓	✓	✓
Batched ISAI	✓	✓	✓	✓
Batched Jacobi	✓	✓	✓	✓
Batched				
AMG preconditioner	✓	✓	✓	✓
AMG solver	✓	✓	✓	✓
Parallel Graph Match	✓	✓	✓	✓
AMG				
Symbolic Cholesky	✓	✓	✓	✓
Numeric Cholesky	✓	✓	✓	✓
Symbolic LU	✓	✓	✓	✓
Numeric LU	✓	✓	✓	✓
Sparse direct				
Sparse TRSV	✓	✓	✓	✓
On-Device Matrix Assembly	✓	✓	✓	✓
MC64/RCM reordering	✓	✓	✓	✓
Wrapping user data	✓	✓	✓	✓
Logging	✓	✓	✓	✓
Utilities				
PAPI counters	✓	✓	✓	✓



Lessons learnt from the Ginkgo development process

- **ECP earmarking roughly half the budget to Software & App development is a game changer.**
 - **Central component for the success of ECP.**
 - This concept needs to – and does become - the blueprint for other nations and projects.
- **Workforce recruitment and workforce retention are the key to success in software development.**
 - Money does not write software. RSEs do. **We need to create attractive career plans.**
 - We need to make research software development attractive to students. **Academic recognition.**
- **Anticipating the future in hardware development accelerates the porting process.**
 - **Blueprints** and **early access systems** both useful.
 - **Interaction with industry** is mutually beneficial.
- **Management, tools, and strategic initiatives, interaction and collegial behavior are important.**
 - Jira/Notion/[...] milestones and deliverables give projects and collaborative interactions a structure and timeline.
 - **Strategic focus groups, conferences, and meetings** bring experts together and **create collaboration.**
 - **Listen to the application needs. Value input and acknowledge collaborators.**