



ALMA MATER STUDIORUM
UNIVERSITÀ DI BOLOGNA

CINECA



Accelerating Sparse Linear Solvers in OpenFOAM using RISC-V Vector Extensions

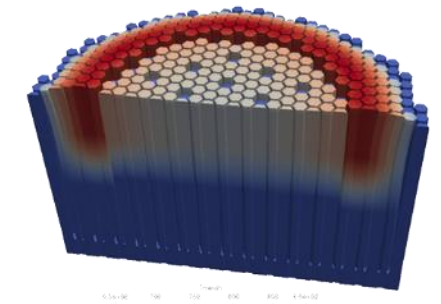
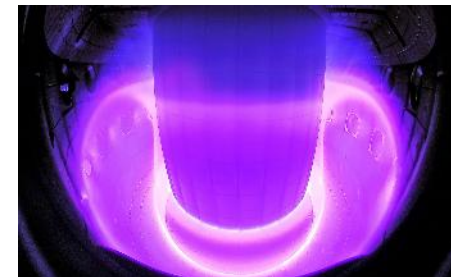
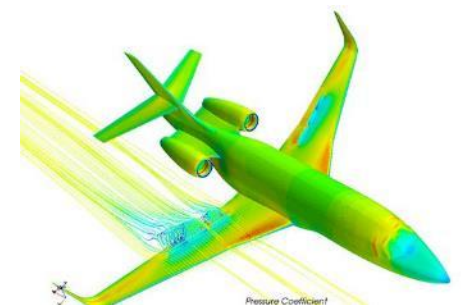
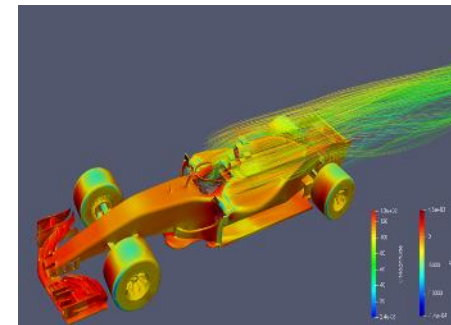
Gabriele Ceccolini¹, Federico Ficarelli², Filippo Barbari², Simone Bnà², Andrea Bartolini¹

¹University of Bologna, Italy, ²CINECA, Italy

gabriele.ceccolini@studio.unibo.it

CFD economy and OpenFOAM framework

Open  FOAM



OpenFOAM® HPC Challenge (OHC-1)]

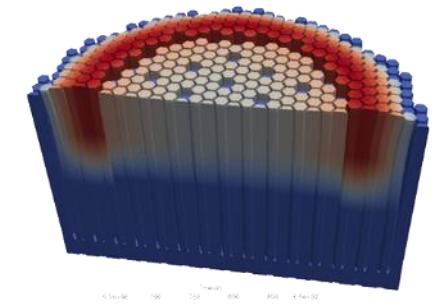
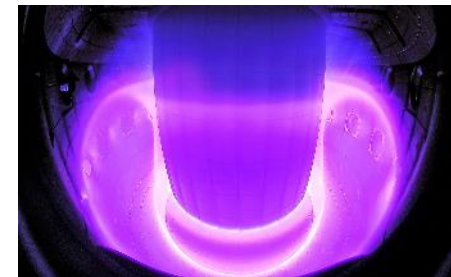
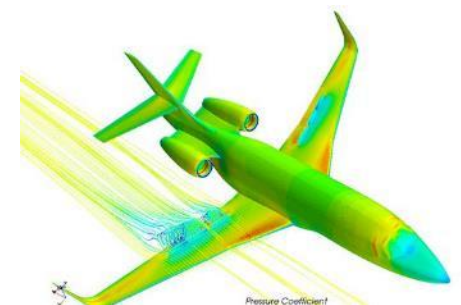
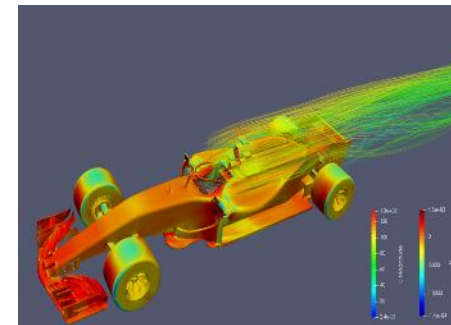
CINECA HPC Annual Report 2024/2025

ENGYS, "2023 CFD Technology Trends Survey Results," ENGYS Blog, 2023

CFD economy and OpenFOAM framework

- Computational Fluid Dynamics (CFD) represents a major fraction of HPC usage (e.g., 8% of total core-hours at CINECA).

Open  FOAM



OpenFOAM® HPC Challenge (OHC-1)]

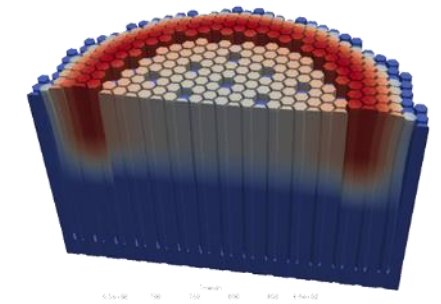
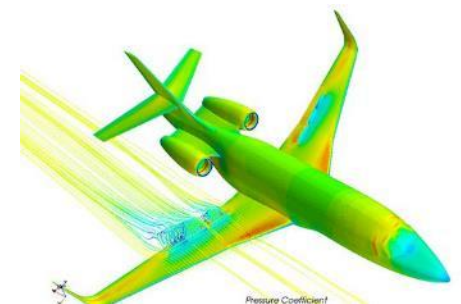
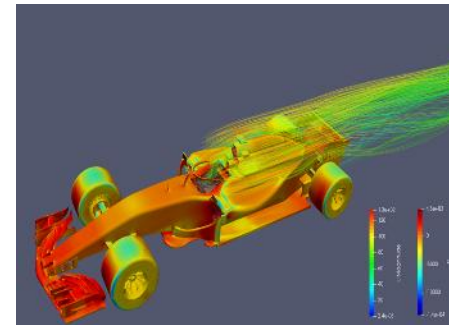
CINECA HPC Annual Report 2024/2025

ENGYS, "2023 CFD Technology Trends Survey Results," ENGYS Blog, 2023

CFD economy and OpenFOAM framework

- Computational Fluid Dynamics (CFD) represents a major fraction of HPC usage (e.g., 8% of total core-hours at CINECA).
- Commercial CFD codes require per-core licenses, becoming extremely expensive for large-scale HPC.

OpenFOAM



OpenFOAM® HPC Challenge (OHC-1)]

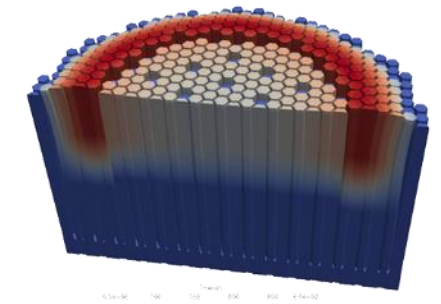
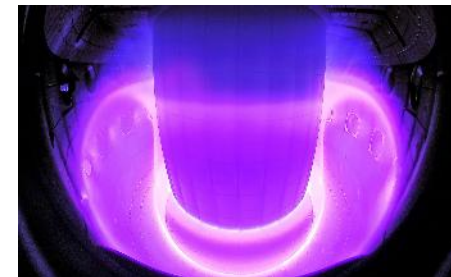
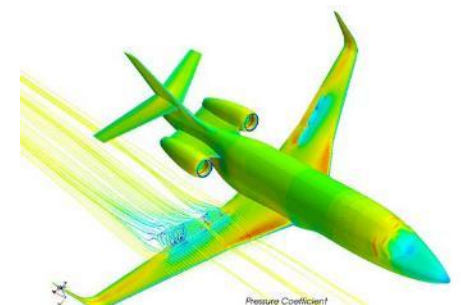
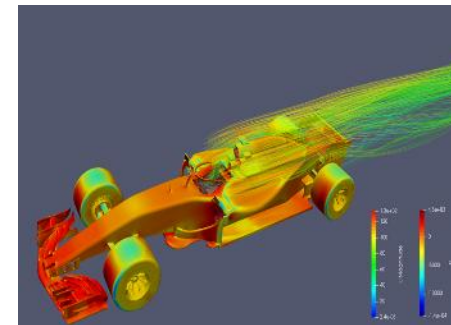
CINECA HPC Annual Report 2024/2025

ENGYS, "2023 CFD Technology Trends Survey Results," ENGYS Blog, 2023

CFD economy and OpenFOAM framework

- Computational Fluid Dynamics (CFD) represents a major fraction of HPC usage (e.g., 8% of total core-hours at CINECA).
- Commercial CFD codes require per-core licenses, becoming extremely expensive for large-scale HPC.
- OpenFOAM is the leading open-source alternative, currently used by 50% of CFD professionals.

OpenFOAM



OpenFOAM® HPC Challenge (OHC-1)]

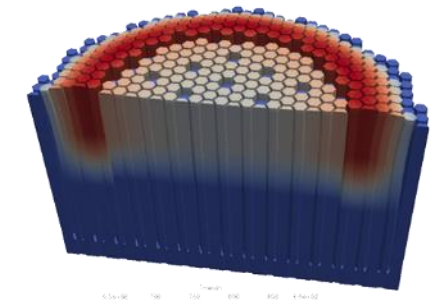
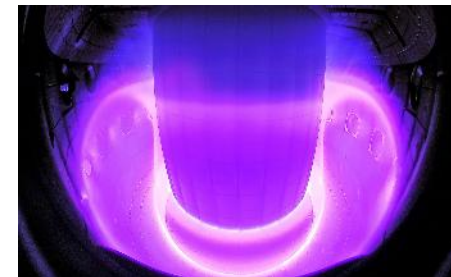
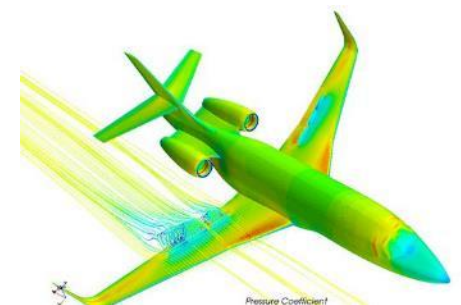
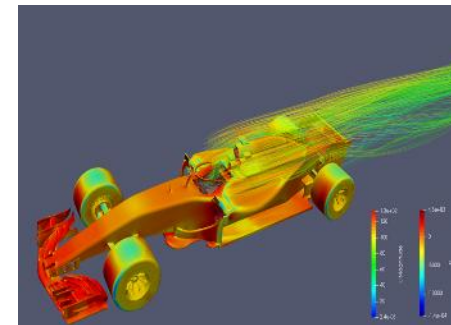
CINECA HPC Annual Report 2024/2025

ENGYS, "2023 CFD Technology Trends Survey Results," ENGYS Blog, 2023

CFD economy and OpenFOAM framework

- Computational Fluid Dynamics (CFD) represents a major fraction of HPC usage (e.g., 8% of total core-hours at CINECA).
- Commercial CFD codes require per-core licenses, becoming extremely expensive for large-scale HPC.
- OpenFOAM is the leading open-source alternative, currently used by 50% of CFD professionals.
- The energy consumed in 2024 by the CFD OpenFOAM jobs in the CPU partition of Leonardo was approximately 182MWh.

OpenFOAM



OpenFOAM® HPC Challenge (OHC-1)]

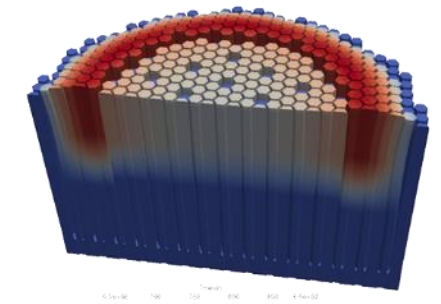
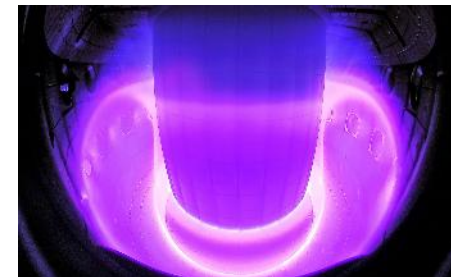
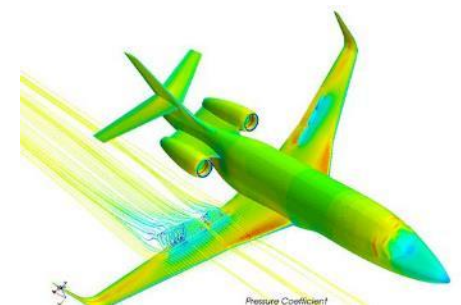
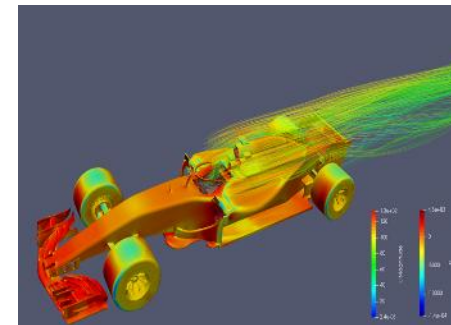
CINECA HPC Annual Report 2024/2025

ENGYS, "2023 CFD Technology Trends Survey Results," ENGYS Blog, 2023

CFD economy and OpenFOAM framework

- Computational Fluid Dynamics (CFD) represents a major fraction of HPC usage (e.g., 8% of total core-hours at CINECA).
- Commercial CFD codes require per-core licenses, becoming extremely expensive for large-scale HPC.
- OpenFOAM is the leading open-source alternative, currently used by 50% of CFD professionals.
- The energy consumed in 2024 by the CFD OpenFOAM jobs in the CPU partition of Leonardo was approximately 182MWh.
- OpenFOAM scales efficiently across thousands of cores (e.g., 32 768 CPU cores at OpenFOAM® HPC Challenge).

OpenFOAM



OpenFOAM® HPC Challenge (OHC-1)]

CINECA HPC Annual Report 2024/2025

ENGYS, "2023 CFD Technology Trends Survey Results," ENGYS Blog, 2023

OpenFOAM current state limitations



OpenFOAM current state limitations

- Based on sparse linear algebra; strongly memory bound.

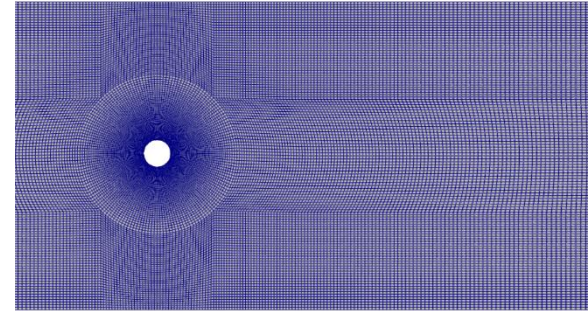
OpenFOAM current state limitations

- Based on sparse linear algebra; strongly memory bound.
- Current Parallelism Status:
 - Limited to domain decomposition and multicore execution using MPI

OpenFOAM current state limitations

\$ blockMesh

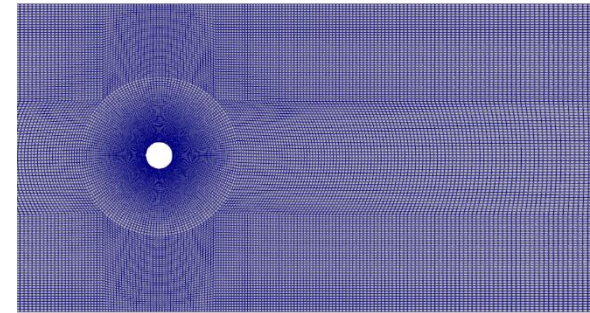
- Based on sparse linear algebra; strongly memory bound.
- Current Parallelism Status:
 - Limited to domain decomposition and multicore execution using MPI



OpenFOAM current state limitations

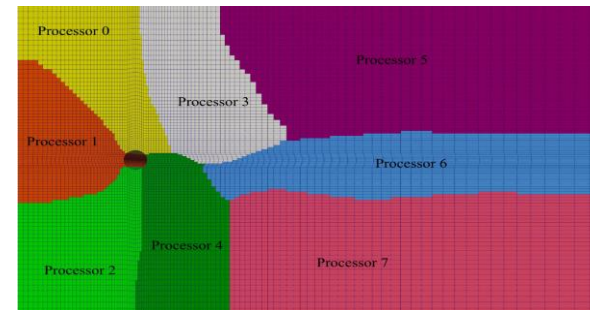
- Based on sparse linear algebra; strongly memory bound.
- Current Parallelism Status:
 - Limited to domain decomposition and multicore execution using MPI

\$ blockMesh



\$ decomposePar

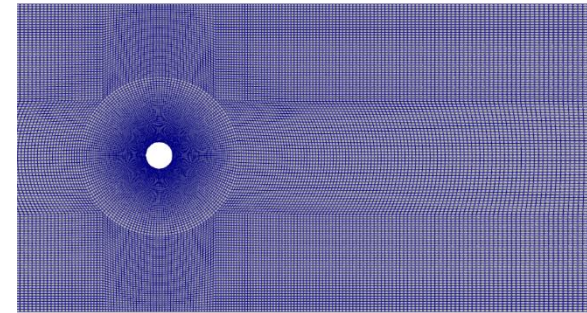
numberOfSubdomains 8;



OpenFOAM current state limitations

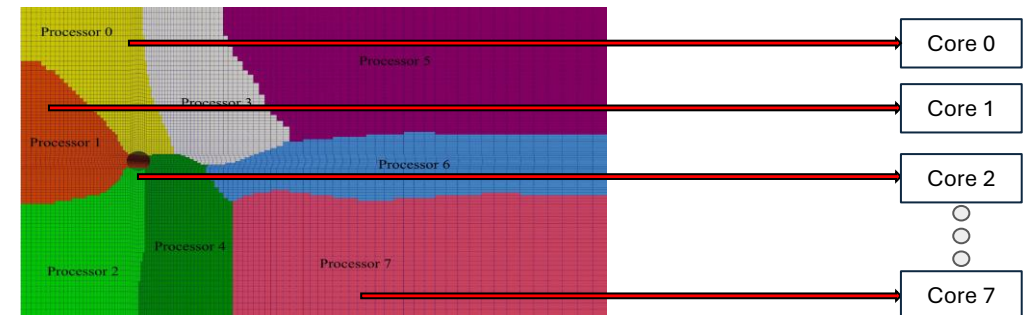
- Based on sparse linear algebra; strongly memory bound.
- Current Parallelism Status:
 - Limited to domain decomposition and multicore execution using MPI

```
$ blockMesh
```

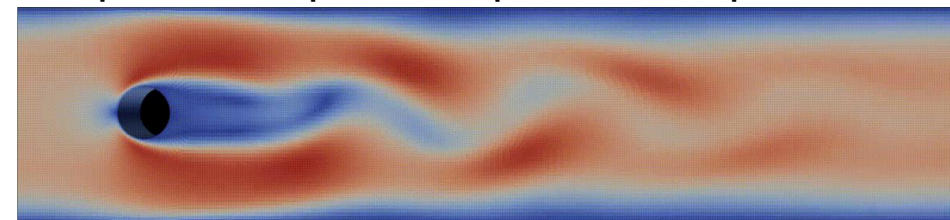


```
$ decomposePar
```

```
numberOfSubdomains 8;
```



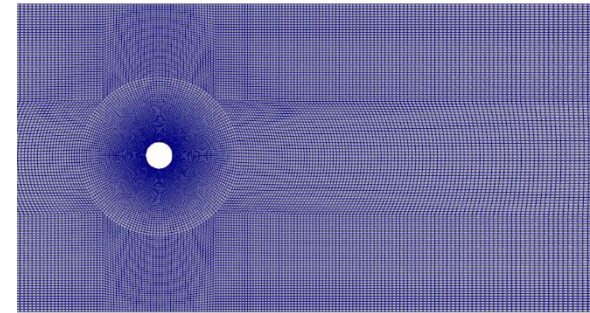
```
$ mpirun -np 8 simpleFoam -parallel
```



OpenFOAM current state limitations

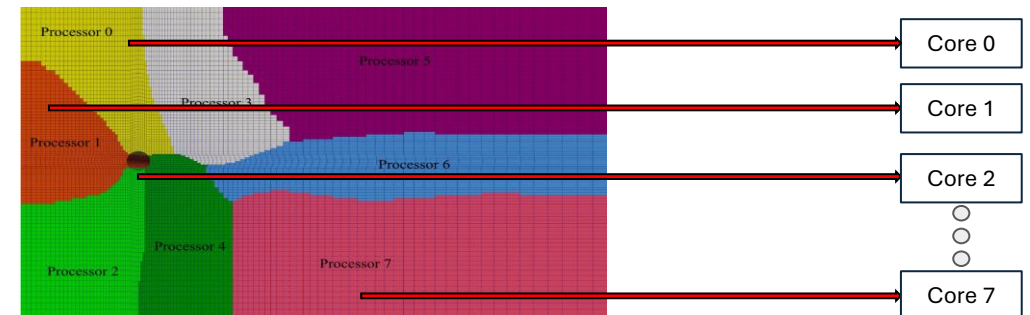
- Based on sparse linear algebra; strongly memory bound.
- Current Parallelism Status:
 - Limited to domain decomposition and multicore execution using MPI
 - Poor vectorization (only compiler auto vectorization limited by internal sparse matrix format).

```
$ blockMesh
```

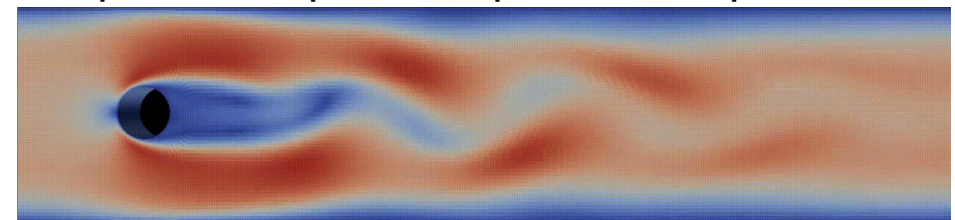


```
$ decomposePar
```

```
numberOfSubdomains 8;
```



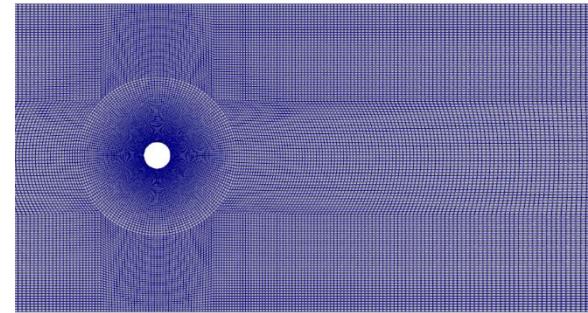
```
$ mpirun -np 8 simpleFoam -parallel
```



OpenFOAM current state limitations

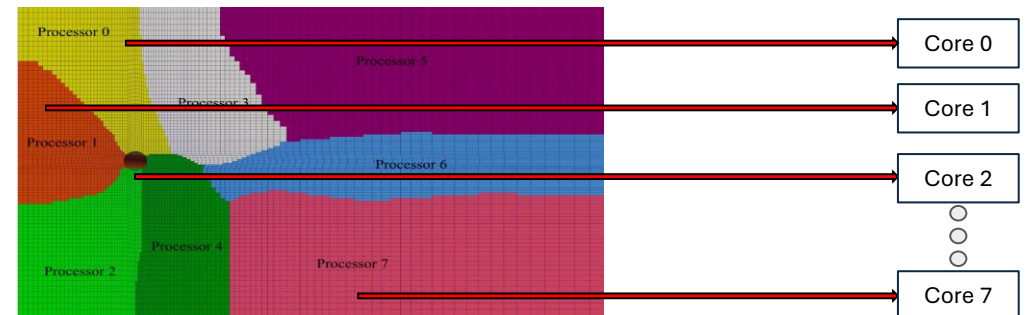
- Based on sparse linear algebra; strongly memory bound.
- Current Parallelism Status:
 - Limited to domain decomposition and multicore execution using MPI
 - Poor vectorization (only compiler auto vectorization limited by internal sparse matrix format).
 - GPU support is limited (e.g. SPUMA) and not integrated with the upstream version

```
$ blockMesh
```

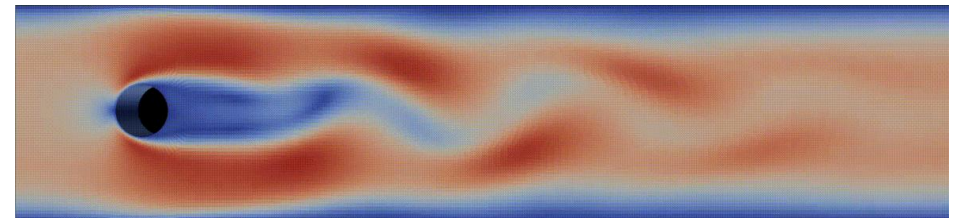


```
$ decomposePar
```

```
numberOfSubdomains 8;
```



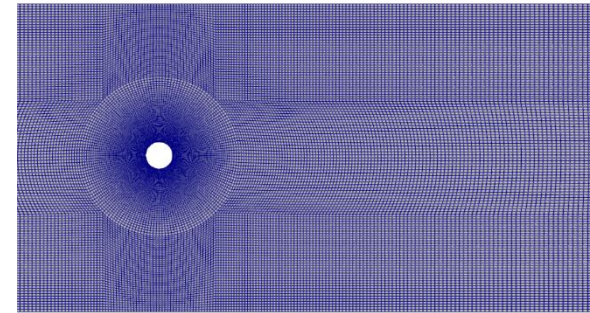
```
$ mpirun -np 8 simpleFoam -parallel
```



OpenFOAM current state limitations

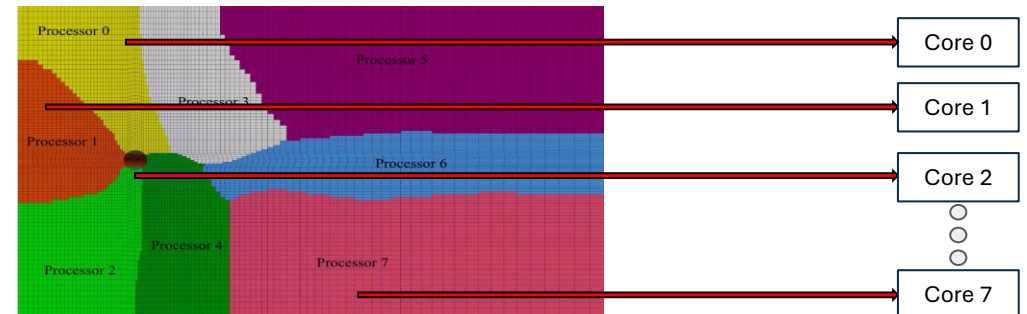
- Based on sparse linear algebra; strongly memory bound.
- Current Parallelism Status:
 - Limited to domain decomposition and multicore execution using MPI
 - Poor vectorization (only compiler auto vectorization limited by internal sparse matrix format).
 - GPU support is limited (e.g. SPUMA) and not integrated with the upstream version
 - **Virtually, no use of data-level parallelism!**

```
$ blockMesh
```

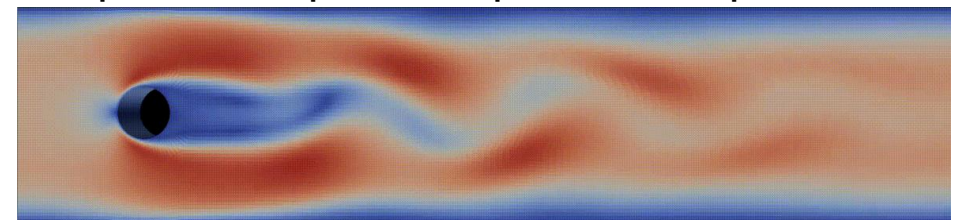


```
$ decomposePar
```

```
numberOfSubdomains 8;
```



```
$ mpirun -np 8 simpleFoam -parallel
```



Linear solvers and SpMV



Linear solvers and SpMV

- CFD mathematics is based on the discretization of fluid dynamics PDEs (e.g. Navier Stokes)

$$\frac{\partial \rho}{\partial t} + \nabla \cdot (\rho \mathbf{u}) = 0$$

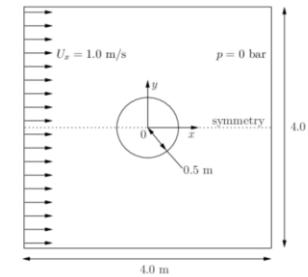


Figure 2.15: Geometry of flow round a cylinder

Linear solvers and SpMV

- CFD mathematics is based on the discretization of fluid dynamics PDEs (e.g. Navier Stokes)
- The spatial domain is discretized into a mesh of control volumes to derive one linear equation per cell.

$$\frac{\partial \rho}{\partial t} + \nabla \cdot (\rho \mathbf{u}) = 0$$

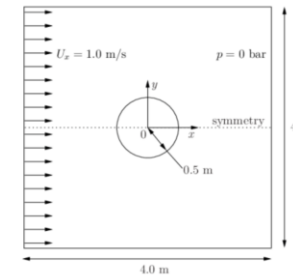


Figure 2.15: Geometry of flow round a cylinder

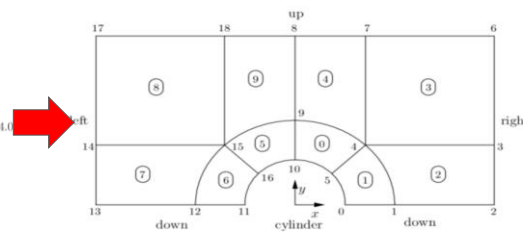


Figure 2.16: Blocks in cylinder geometry

Linear solvers and SpMV

- CFD mathematics is based on the discretization of fluid dynamics PDEs (e.g. Navier Stokes)
- The spatial domain is discretized into a mesh of control volumes to derive one linear equation per cell.
- This approach collapses the entire problem into **solving a large sparse linear system**.

$$\frac{\partial \rho}{\partial t} + \nabla \cdot (\rho \mathbf{u}) = 0$$

$$Ax = b$$

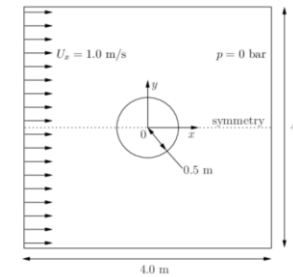


Figure 2.15: Geometry of flow round a cylinder

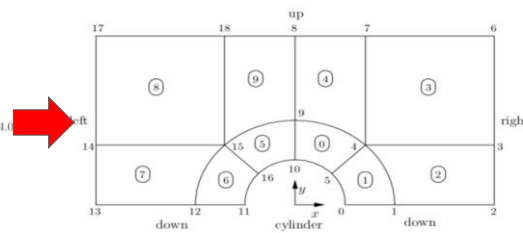
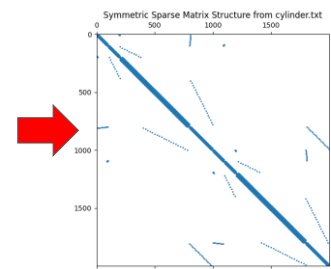


Figure 2.16: Blocks in cylinder geometry



Linear solvers and SpMV

- CFD mathematics is based on the discretization of fluid dynamics PDEs (e.g. Navier Stokes)
- The spatial domain is discretized into a mesh of control volumes to derive one linear equation per cell.
- This approach collapses the entire problem into **solving a large sparse linear system**.
- We use linear iterative solvers computationally based on **sparse matrix vector product (SpMV)**

$$\frac{\partial \rho}{\partial t} + \nabla \cdot (\rho \mathbf{u}) = 0$$

$$Ax = b$$

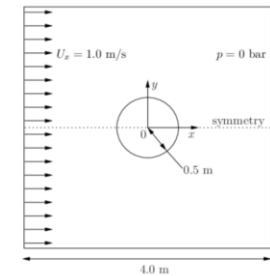


Figure 2.15: Geometry of flow round a cylinder

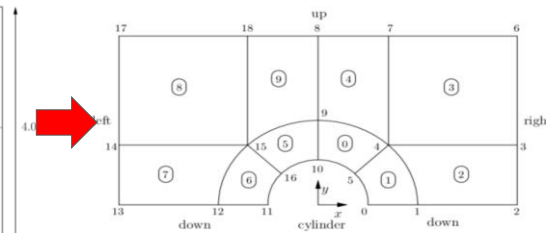
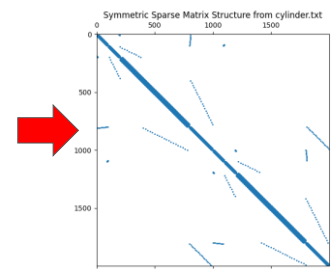


Figure 2.16: Blocks in cylinder geometry



```

r0 := b - Ax0
if r0 is sufficiently small, then return x0 as the result
p0 := r0
k := 0
repeat
    αk :=  $\frac{\mathbf{r}_k^T \mathbf{r}_k}{\mathbf{p}_k^T \mathbf{A} \mathbf{p}_k}$ 
    xk+1 := xk + αk pk
    rk+1 := rk - αk A pk
    if rk+1 is sufficiently small, then exit loop
    βk :=  $\frac{\mathbf{r}_{k+1}^T \mathbf{r}_{k+1}}{\mathbf{r}_k^T \mathbf{r}_k}$ 
    pk+1 := rk+1 + βk pk
    k := k + 1
end repeat
return xk+1 as the result
    
```

Linear solvers and SpMV

- CFD mathematics is based on the discretization of fluid dynamics PDEs (e.g. Navier Stokes)
- The spatial domain is discretized into a mesh of control volumes to derive one linear equation per cell.
- This approach collapses the entire problem into **solving a large sparse linear system**.
- We use linear iterative solvers computationally based on **sparse matrix vector product (SpMV)**
- SpMV is one of the most computationally intensive kernel of a CFD code: good candidate for optimization**

$$\frac{\partial \rho}{\partial t} + \nabla \cdot (\rho \mathbf{u}) = 0 \longrightarrow Ax = b$$

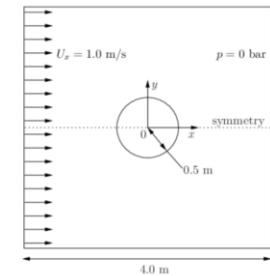


Figure 2.15: Geometry of flow round a cylinder

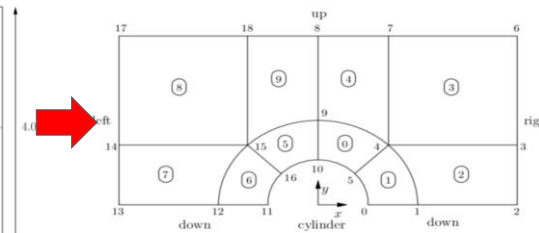
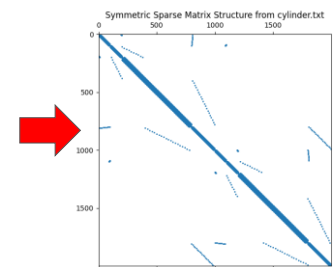
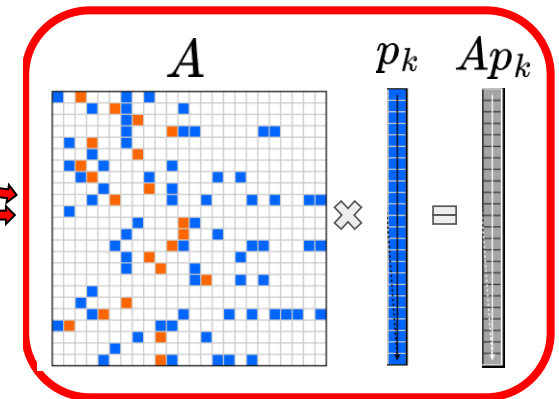


Figure 2.16: Blocks in cylinder geometry



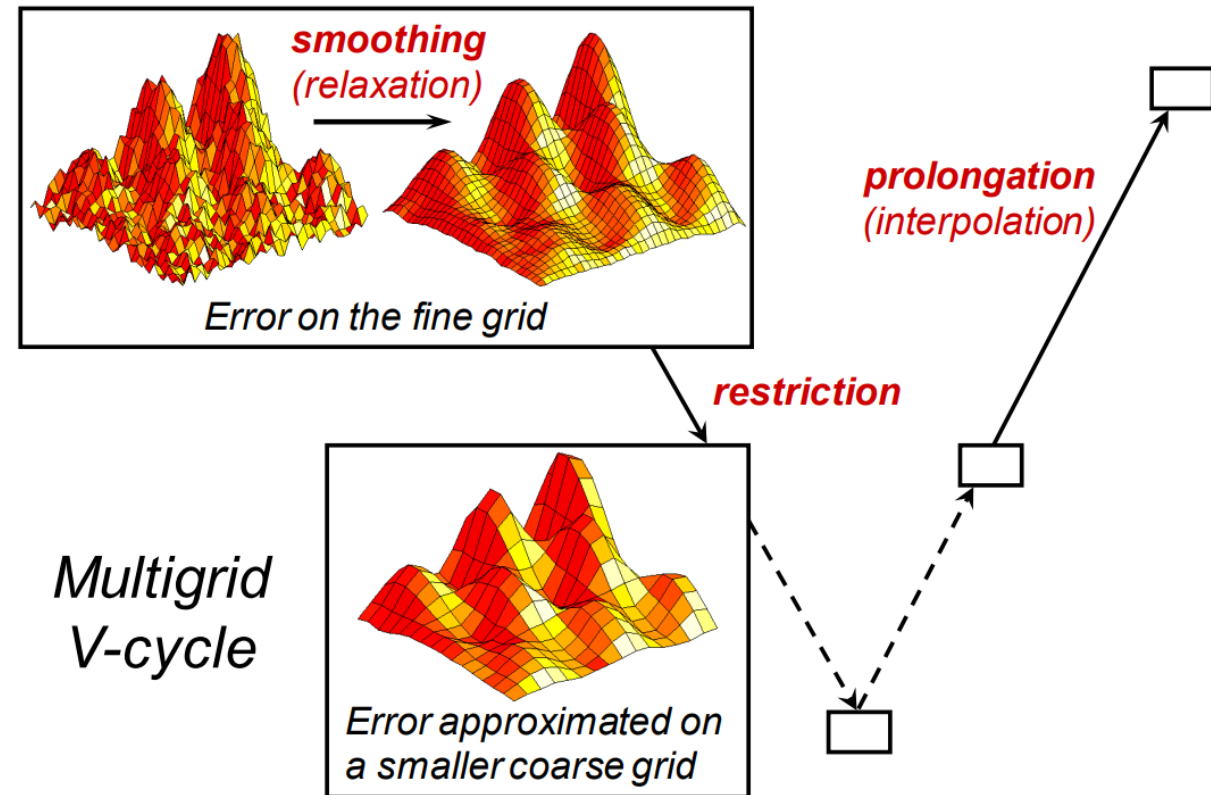
```

r_0 := b - Ax_0
if r_0 is sufficiently small, then return x_0 as the result
p_0 := r_0
k := 0
repeat
    alpha_k := (r_k^T r_k) / (p_k^T (Ap_k))
    x_{k+1} := x_k + alpha_k p_k
    r_{k+1} := r_k - alpha_k (Ap_k)
    if r_{k+1} is sufficiently small, then exit loop
    beta_k := (r_{k+1}^T r_{k+1}) / (r_k^T r_k)
    p_{k+1} := r_{k+1} + beta_k p_k
    k := k + 1
end repeat
return x_{k+1} as the result
    
```



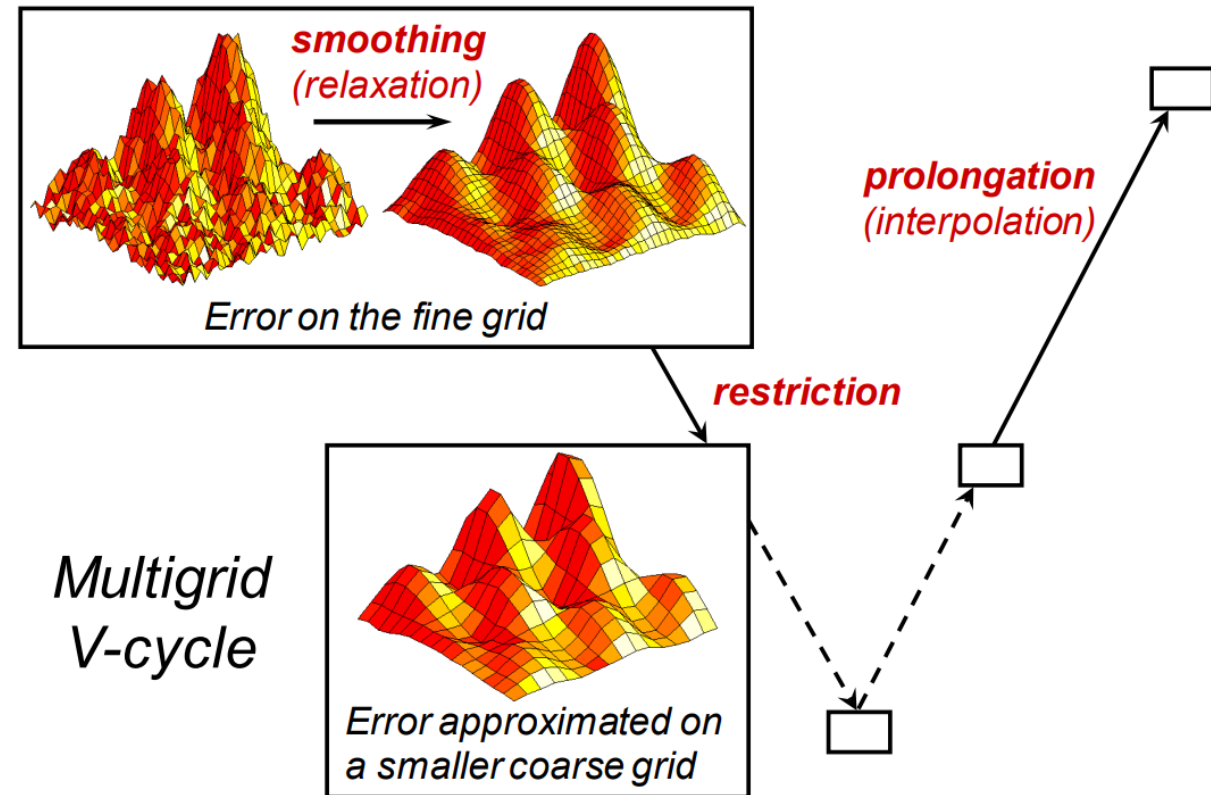
Multigrid methods and Smoothers

- Standard linear solvers effectively dampen efficiently high-frequency errors but stall on low-frequency errors across the mesh.



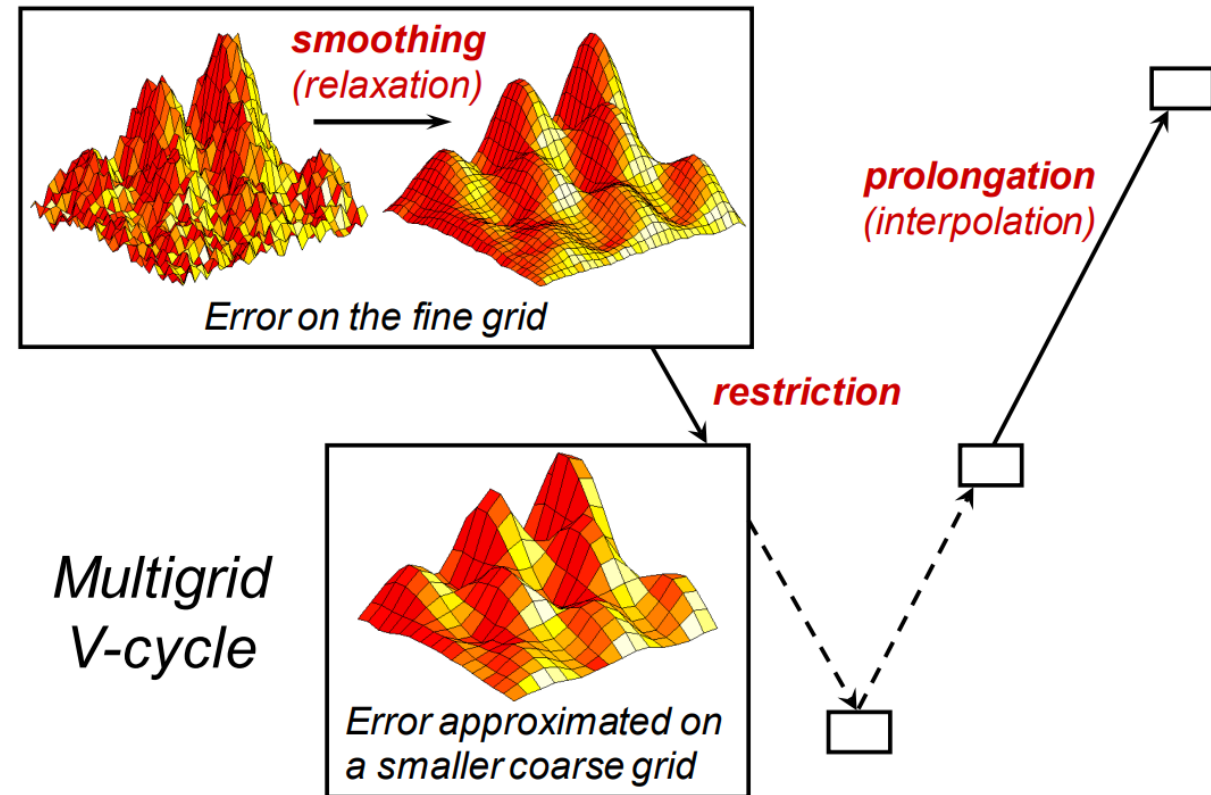
Multigrid methods and Smoothers

- Standard linear solvers effectively dampen efficiently high-frequency errors but stall on low-frequency errors across the mesh.
- To solve this, modern CFD solvers use the **Multigrid Method**, projecting the problem onto progressively coarser meshes.



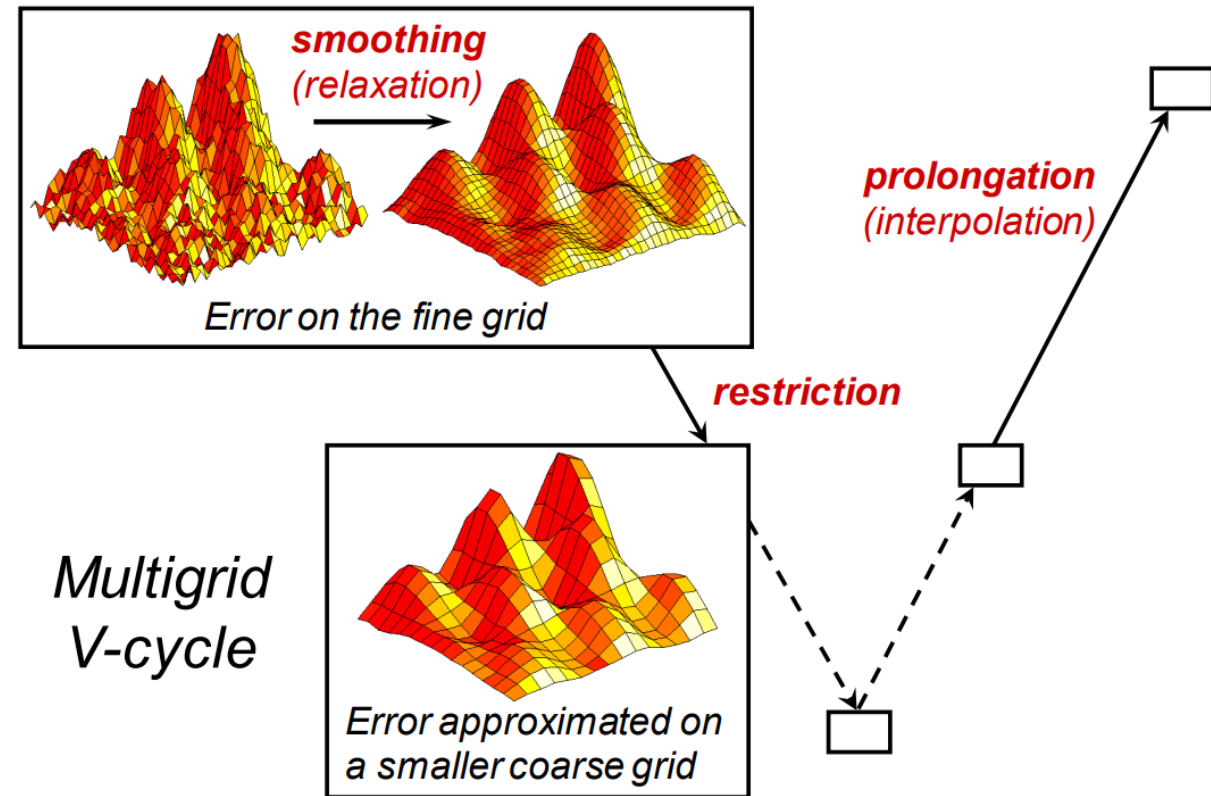
Multigrid methods and Smoothers

- Standard linear solvers effectively dampen efficiently high-frequency errors but stall on low-frequency errors across the mesh.
- To solve this, modern CFD solvers use the **Multigrid Method**, projecting the problem onto progressively coarser meshes.
- At each grid level, a linear solver acts as a "**smoother**," rapidly damp out high-frequency errors specific to that scale.



Multigrid methods and Smoothers

- Standard linear solvers effectively dampen efficiently high-frequency errors but stall on low-frequency errors across the mesh.
- To solve this, modern CFD solvers use the **Multigrid Method**, projecting the problem onto progressively coarser meshes.
- At each grid level, a linear solver acts as a "**smoother**," rapidly damp out high-frequency errors specific to that scale.
- As the grid coarsens, large low-frequency errors appear as high-frequency ones, allowing the smoother to finally damp out them.

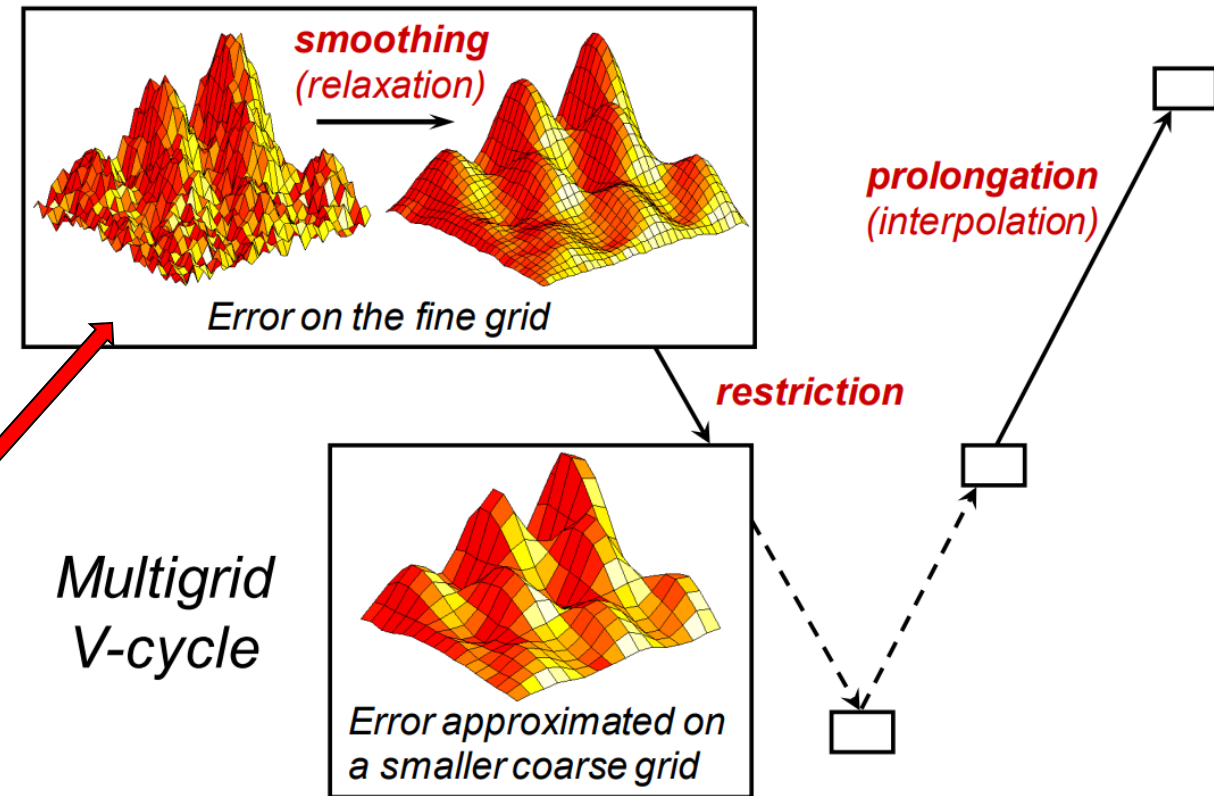


*Multigrid
V-cycle*

*Error approximated on
a smaller coarse grid*

Multigrid methods and Smoothers

- Standard linear solvers effectively dampen efficiently high-frequency errors but stall on low-frequency errors across the mesh.
- To solve this, modern CFD solvers use the **Multigrid Method**, projecting the problem onto progressively coarser meshes.
- At each grid level, a linear solver acts as a "**smoother**," rapidly damp out high-frequency errors specific to that scale.
- As the grid coarsens, large low-frequency errors appear as high-frequency ones, allowing the smoother to finally damp out them.
- **In our test case, this smoothing operation accounts for ~50% of the total compute time, making it the most intensive phase and our primary optimization target.**

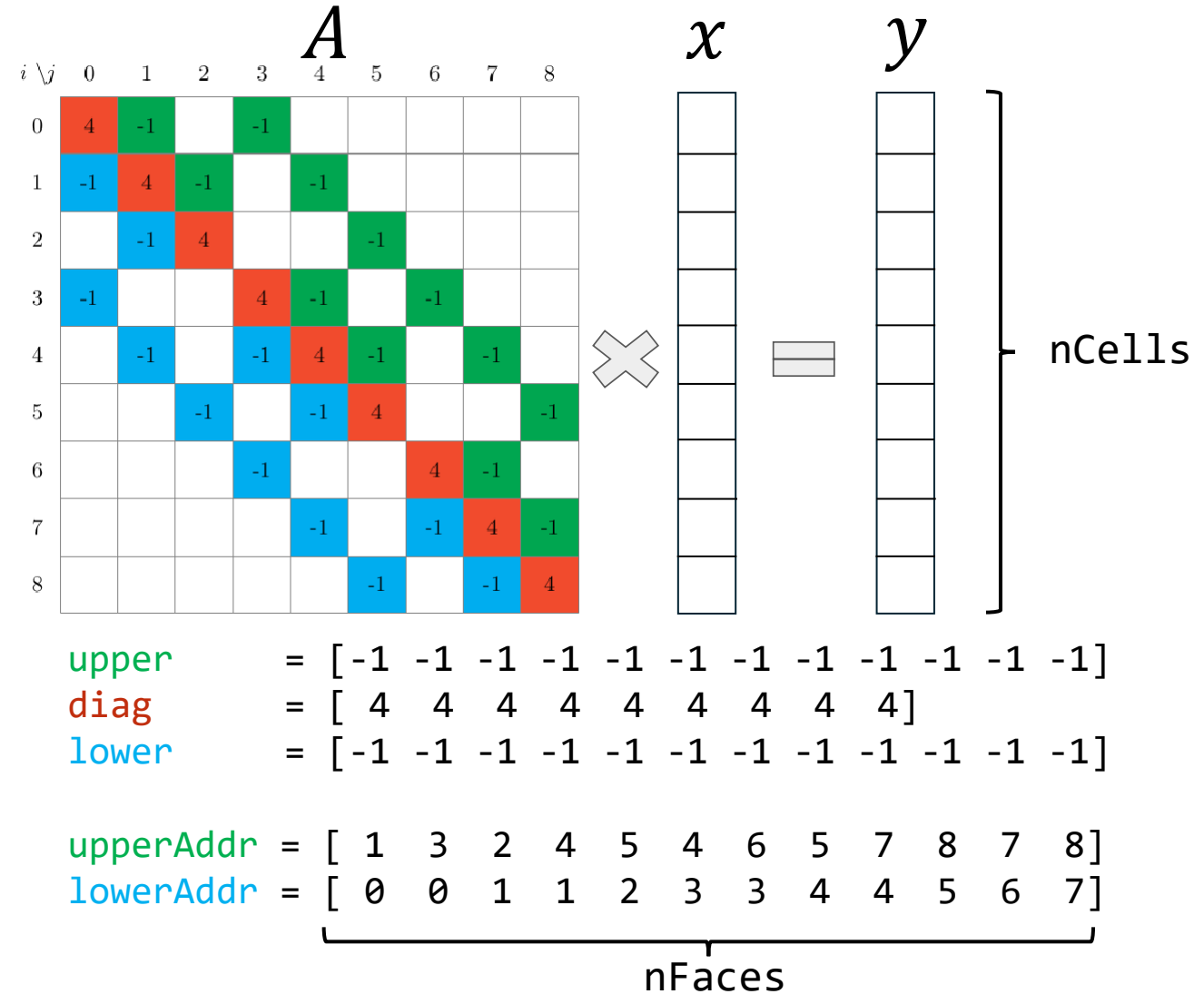


OpenFOAM native sparse matrix format



OpenFOAM native sparse matrix format

- LDU-COO is the internal sparse matrix format in OpenFOAM

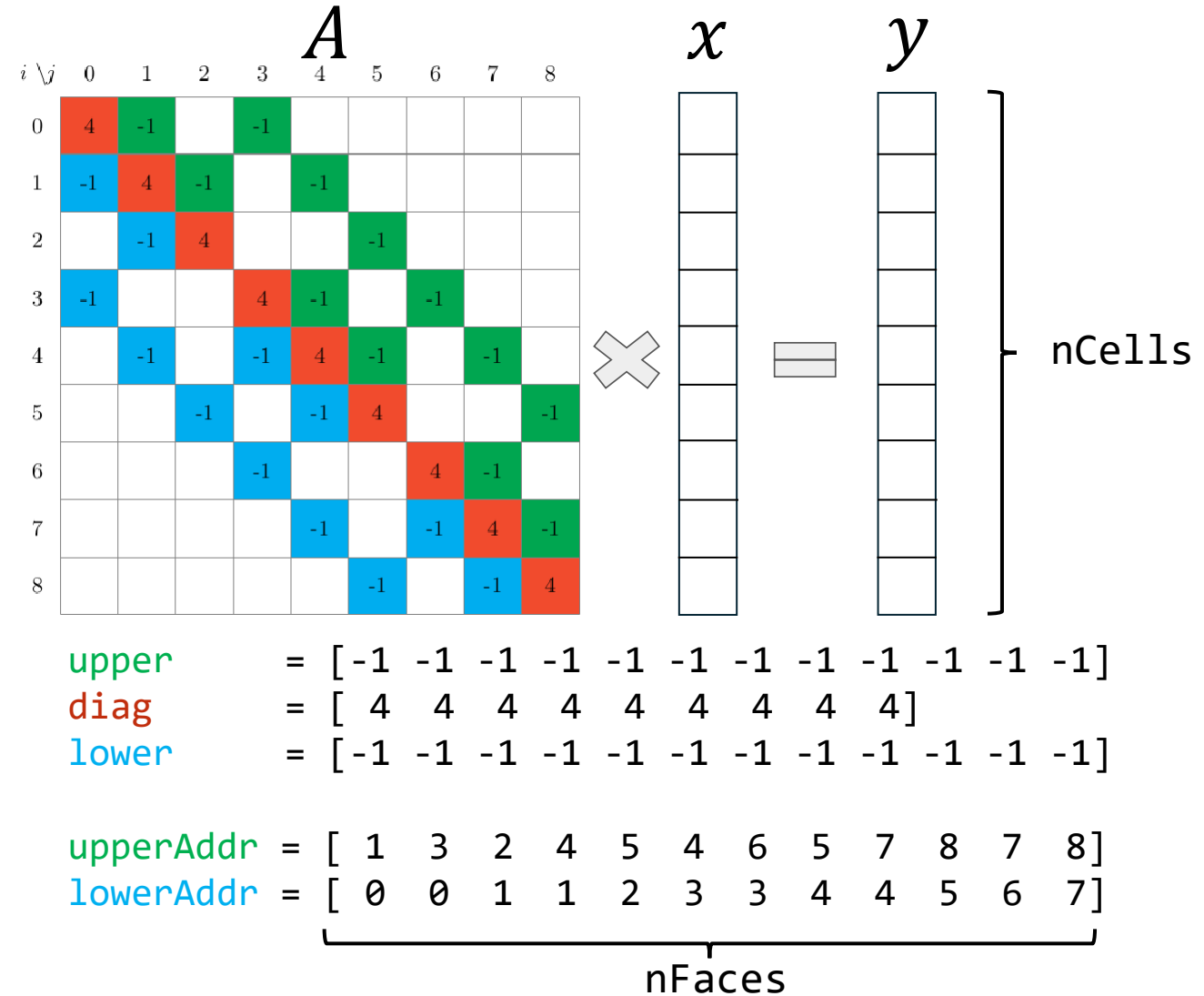


OpenFOAM native sparse matrix format

- LDU-COO is the internal sparse matrix format in OpenFOAM
- Hostile to vectorization
 - Indirect access on x and y

```
// 1. Diagonal contribute
for (size_t i = 0; i < nCells; i++) {
    y[i] = diag[i] * x[i];
}

// 2. Off-Diagonal contribute
for (size_t f = 0; f < nFaces; f++) {
    y[lowerAddr[f]] += upper[f] * x[upperAddr[f]];
    y[upperAddr[f]] += lower[f] * x[lowerAddr[f]];
}
```

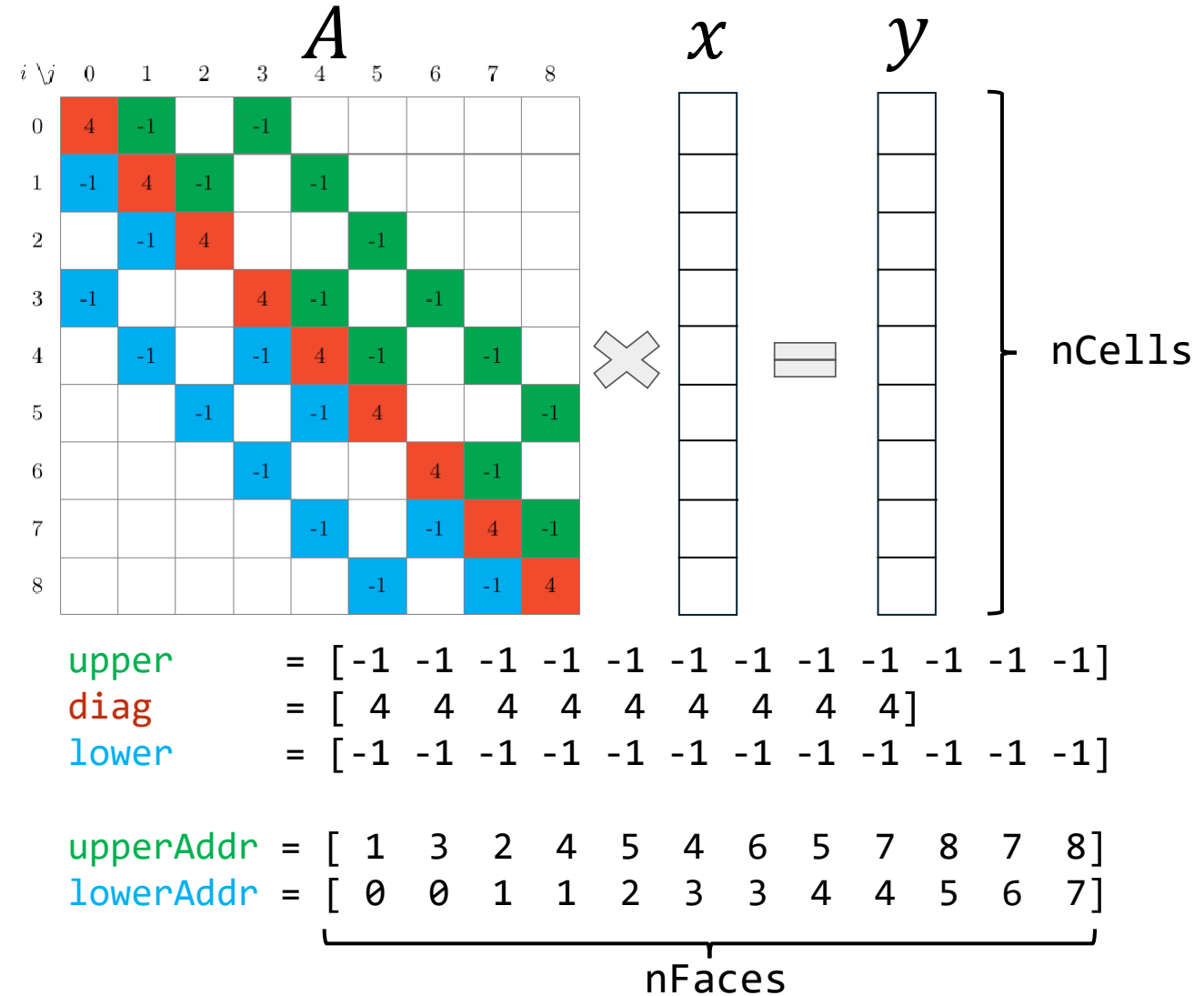


OpenFOAM native sparse matrix format

- LDU-COO is the internal sparse matrix format in OpenFOAM
- Hostile to vectorization
 - Indirect access on x and y
 - Race conditions on y update

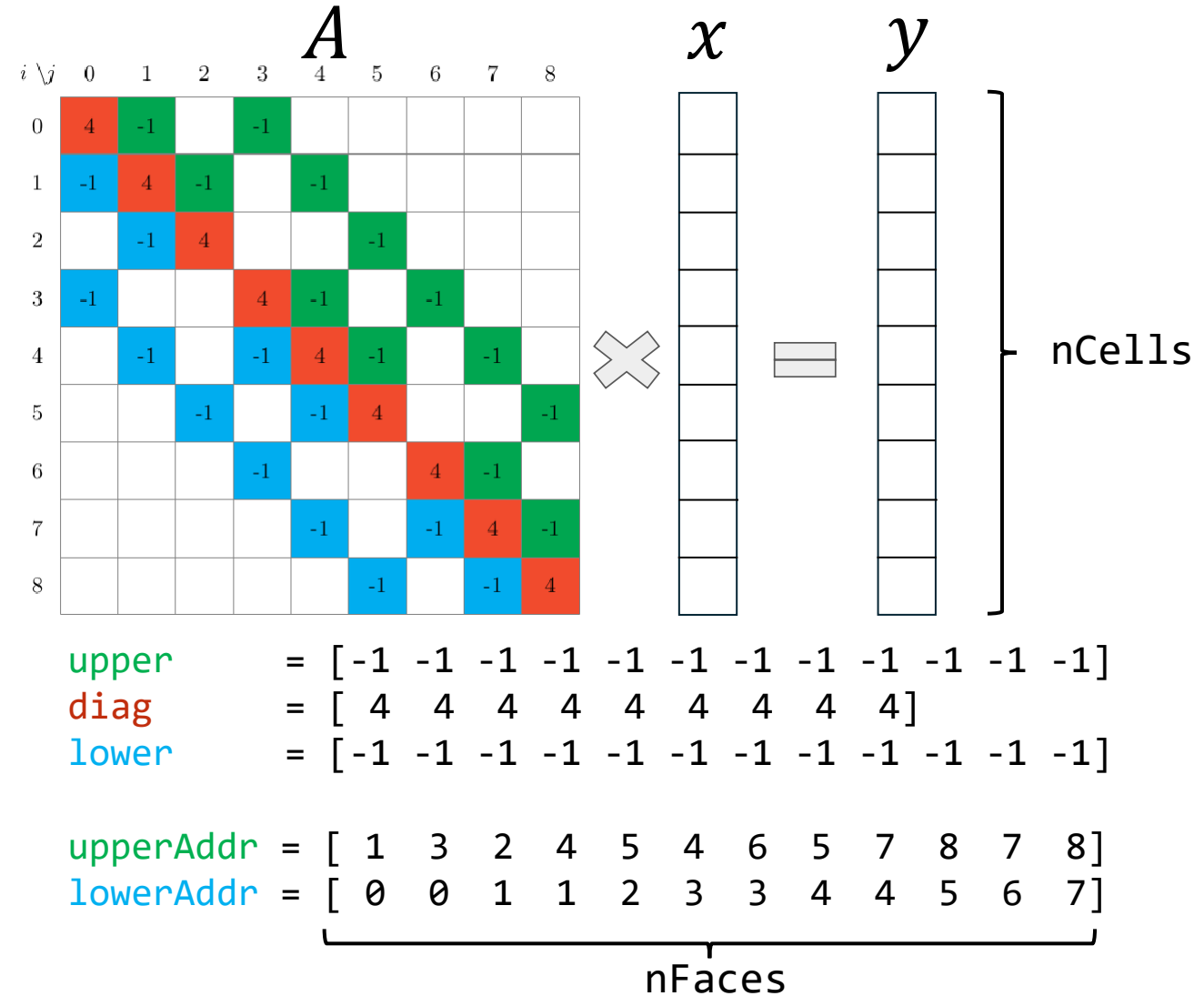
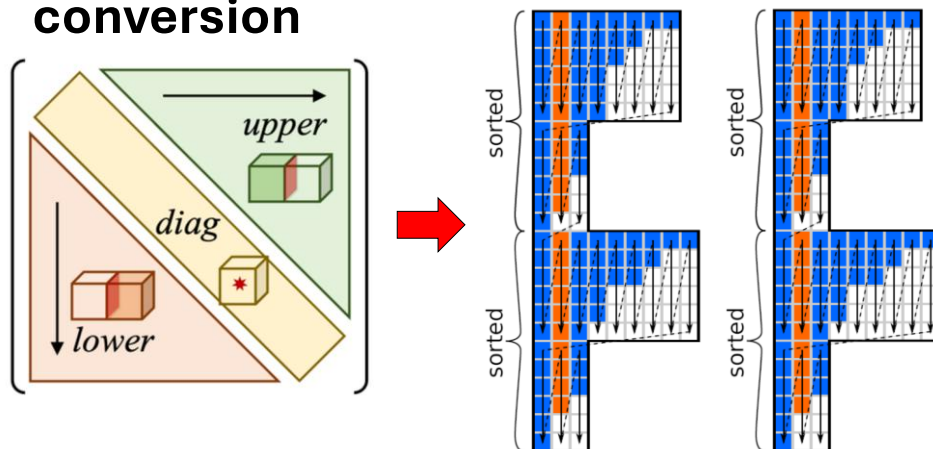
```
// 1. Diagonal contribute
for (size_t i = 0; i < nCells; i++) {
    y[i] = diag[i] * x[i];
}

// 2. Off-Diagonal contribute
for (size_t f = 0; f < nFaces; f++) {
    y[lowerAddr[f]] += upper[f] * x[upperAddr[f]];
    y[upperAddr[f]] += lower[f] * x[lowerAddr[f]];
}
```



OpenFOAM native sparse matrix format

- LDU-COO is the internal sparse matrix format in OpenFOAM
- Hostile to vectorization
 - Indirect access on x and y
 - Race conditions on y update
- **Solution -> runtime matrix format conversion**



Application profiling



ALMA MATER STUDIORUM
UNIVERSITÀ DI BOLOGNA

CINECA

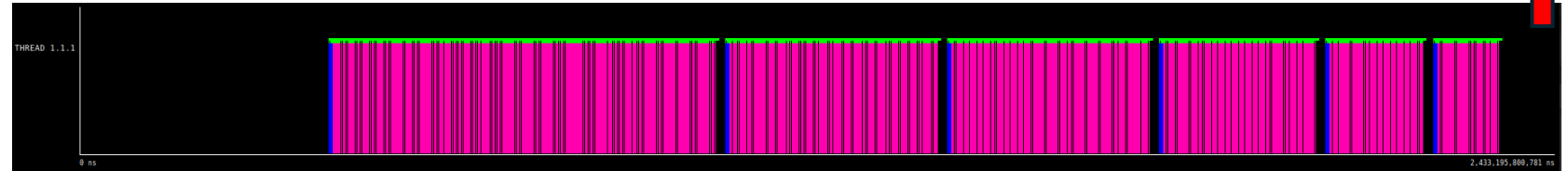


RISC-V Summit Europe 2026 11.06.2026

V. Pillet et al. (1995) PARAVR: A tool to visualize and analyze parallel code <http://tools.bsc.es/paraver>

Application profiling

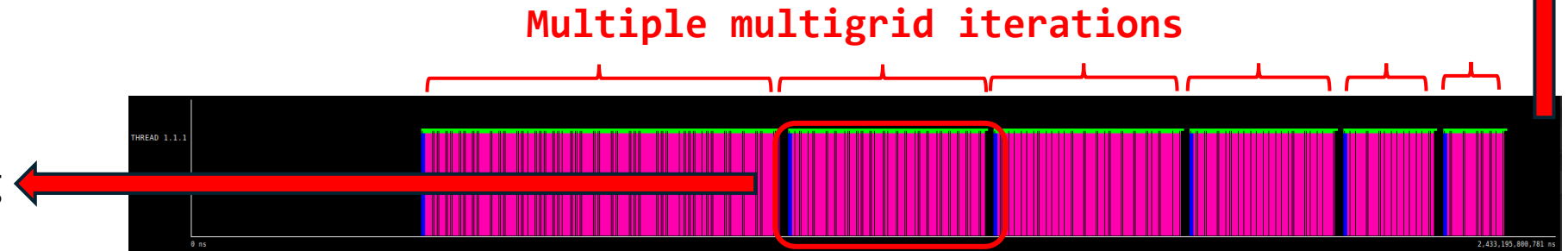
Full CFD simulation
timeline profiled with
Paraver



Application profiling

Full CFD simulation
timeline profiled with
Paraver

Multiple multigrid method
iterations are called during
a CFD simulation

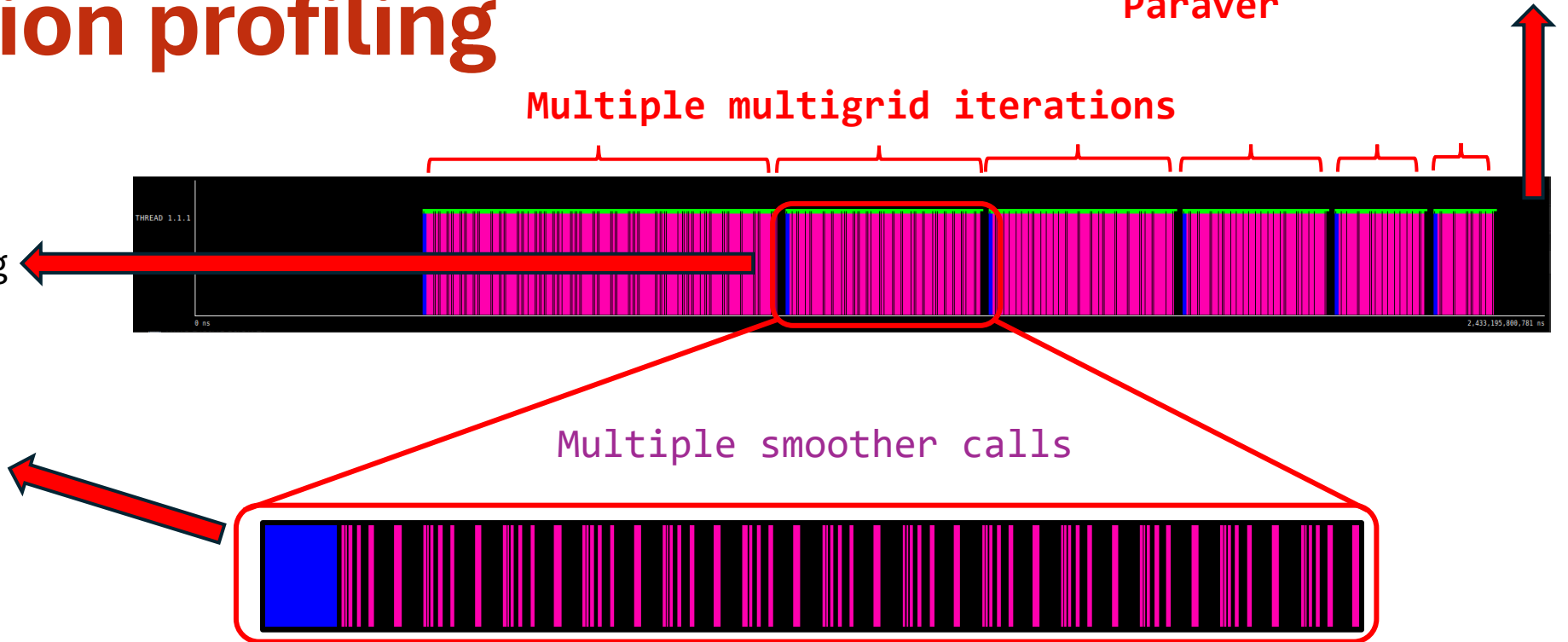


Application profiling

Full CFD simulation
timeline profiled with
Paraver

Multiple multigrid method
iterations are called during
a CFD simulation

During a single multigrid
iteration, the smoother is
called tens to hundreds
of times



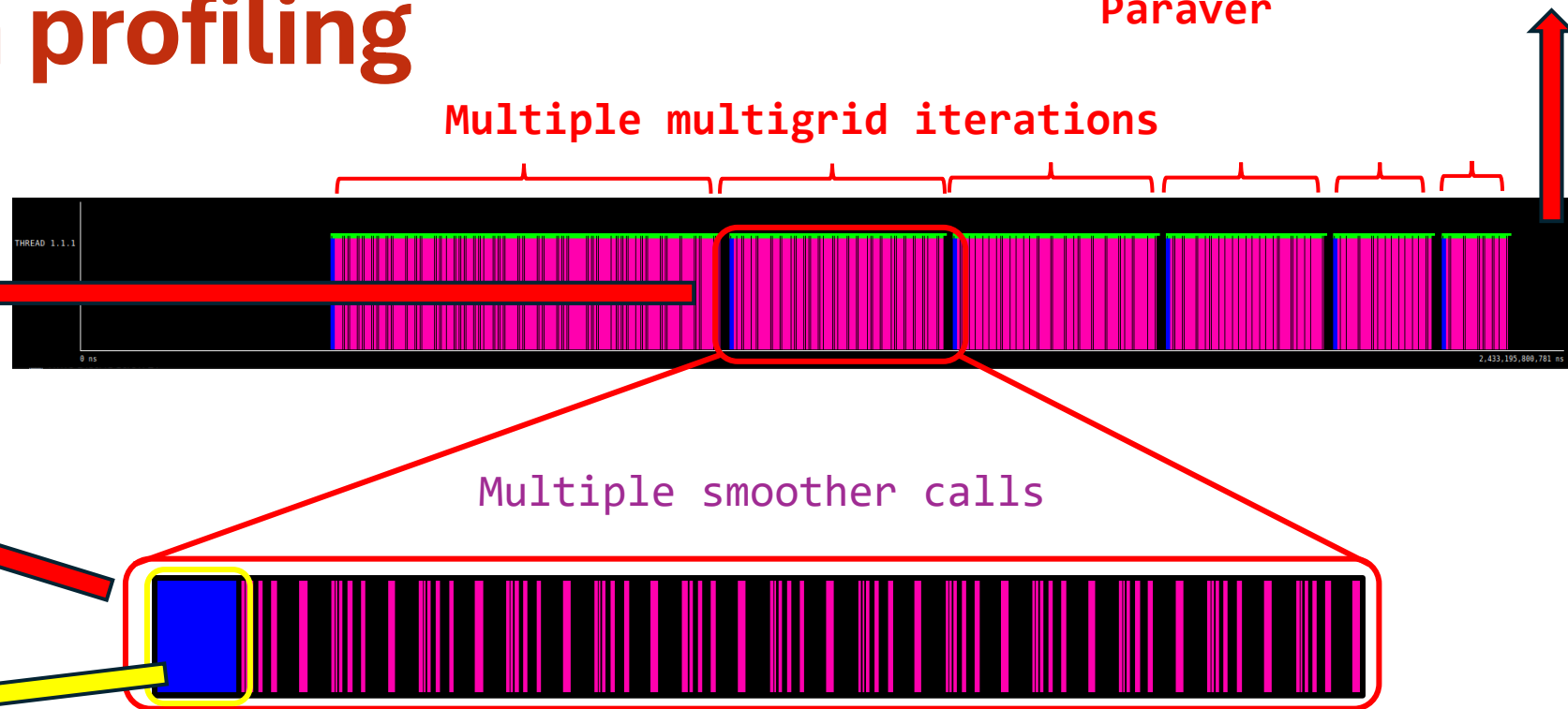
Application profiling

Full CFD simulation
timeline profiled with
Paraver

Multiple multigrid method
iterations are called during
a CFD simulation

During a single multigrid
iteration, the smoother is
called tens to hundreds
of times

Format conversion
overhead



Application profiling

Full CFD simulation
timeline profiled with
Paraver

Multiple multigrid method
iterations are called during
a CFD simulation

During a single multigrid
iteration, the smoother is
called tens to hundreds
of times

Format conversion
overhead

Multiple multigrid iterations

Multiple smoother calls

$$(D^{-1})_{ii} = \frac{1}{A_{ii}}$$

Smoother

$$\psi^{(k+1)} = \psi^{(k)} + \omega D^{-1}(b - A\psi^{(k)})$$

SpMV

SpMV

$A\psi^{(k)}$

V. Pillet et al. (1995) PARAVAR: A tool to visualize and
analyze parallel code <http://tools.bsc.es/paraver>

Application profiling

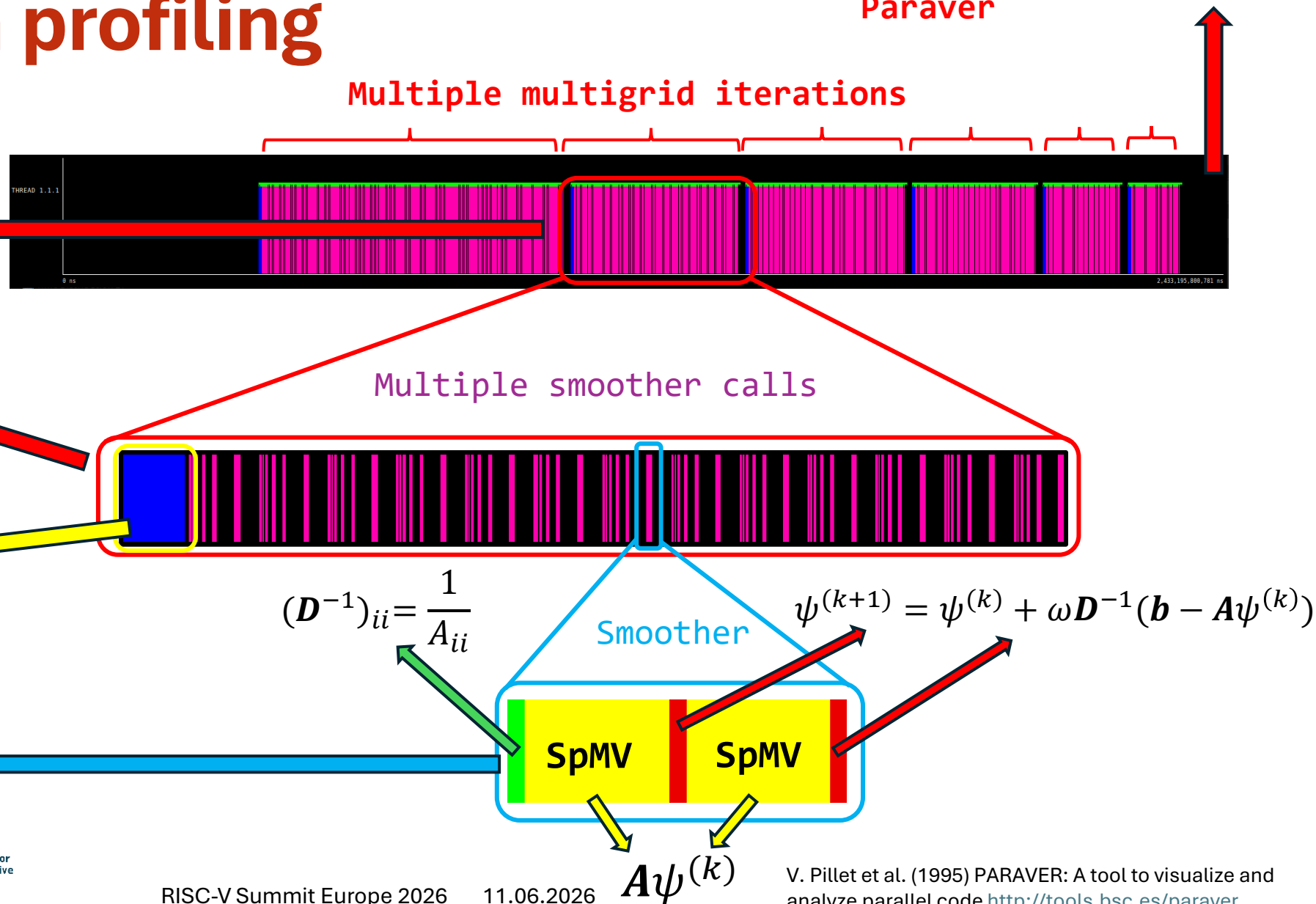
Full CFD simulation
timeline profiled with
Paraver

Multiple multigrid method
iterations are called during
a CFD simulation

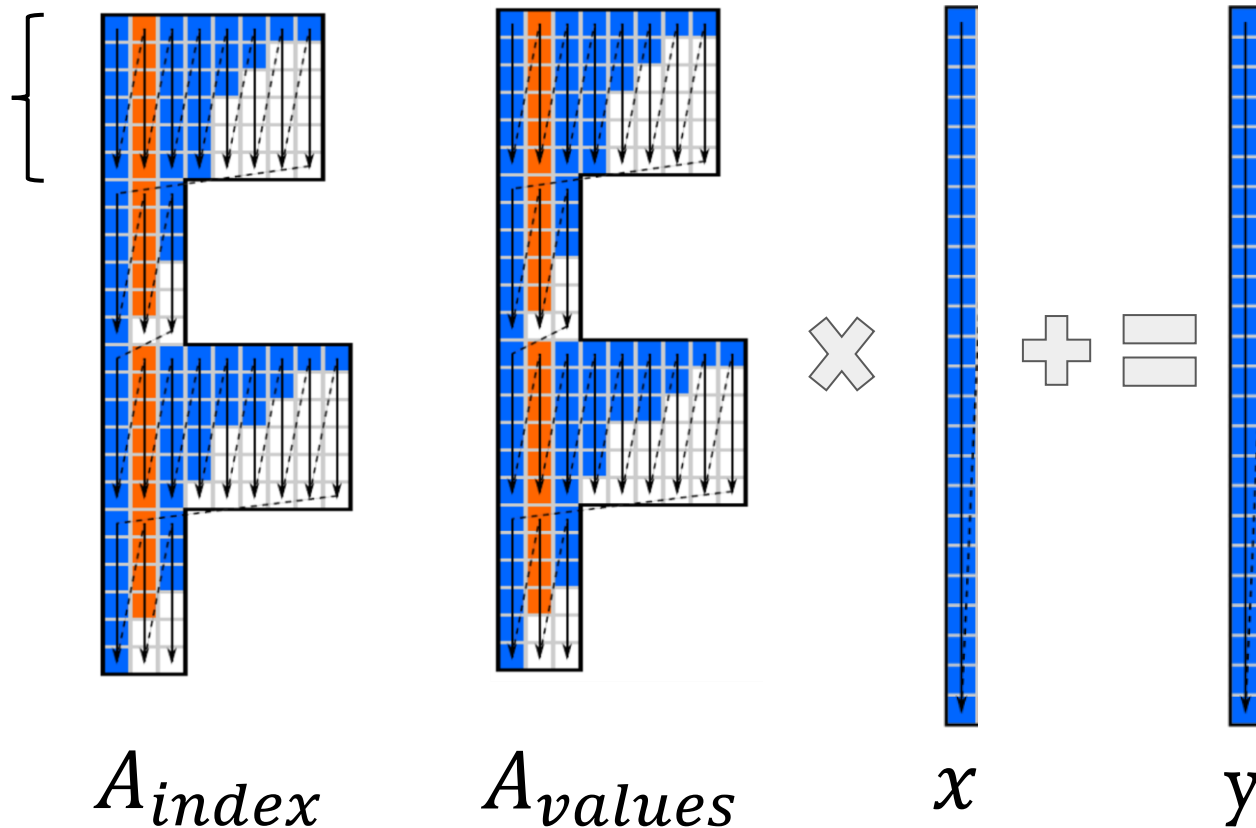
During a single multigrid
iteration, the smoother is
called tens to hundreds
of times

Format conversion
overhead

The smoother represents
around 50% of the total
application runtime, and the
computation is dominated by
SpMV



SpMV SELL-C- σ vectorization

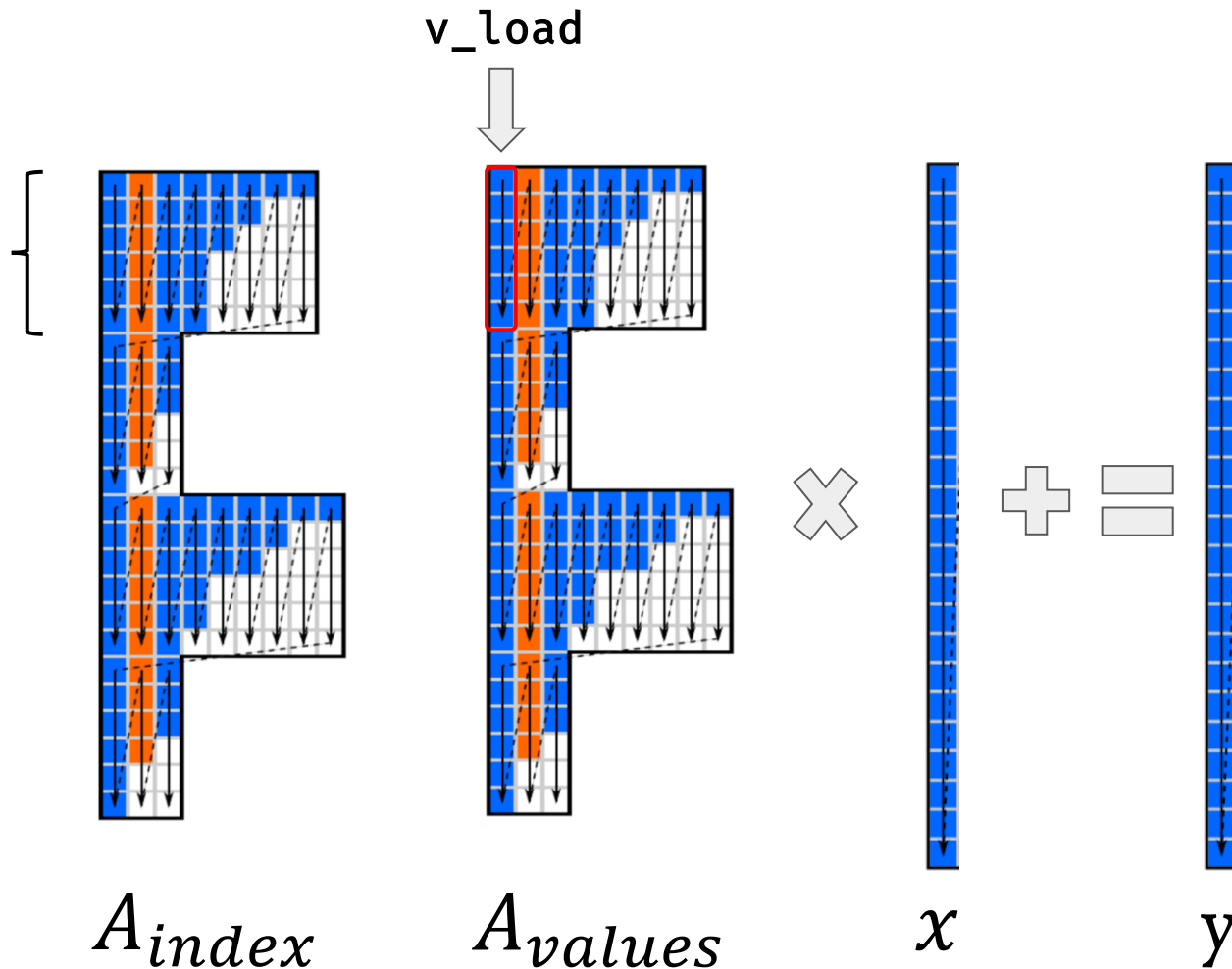


- A_{index} , A_{values} are stored by slice and within the slice by column
- Contiguous and predictable vector loads A_{values} and A_{index}
- No race conditions on y updates
- Especially in long vectors architectures (e.g. EPAC with $vl_{en}=256$) we get huge benefits by hiding memory latency and better use of bandwidth during vector loads and stores [*]

M. Kreutzer et al., "A Unified Sparse Matrix Data Format for Efficient General Sparse Matrix-Vector Multiplication on Modern Processors with Wide SIMD Units," *SIAM J. Sci. Comput.*, 2014.

[*] Pablo Vizcaino et al. (2023) Short reasons for long vectors in HPC CPUs: a study based on RISC-V <https://doi.org/10.1145/3624062.3624231>

SpMV SELL-C- σ vectorization

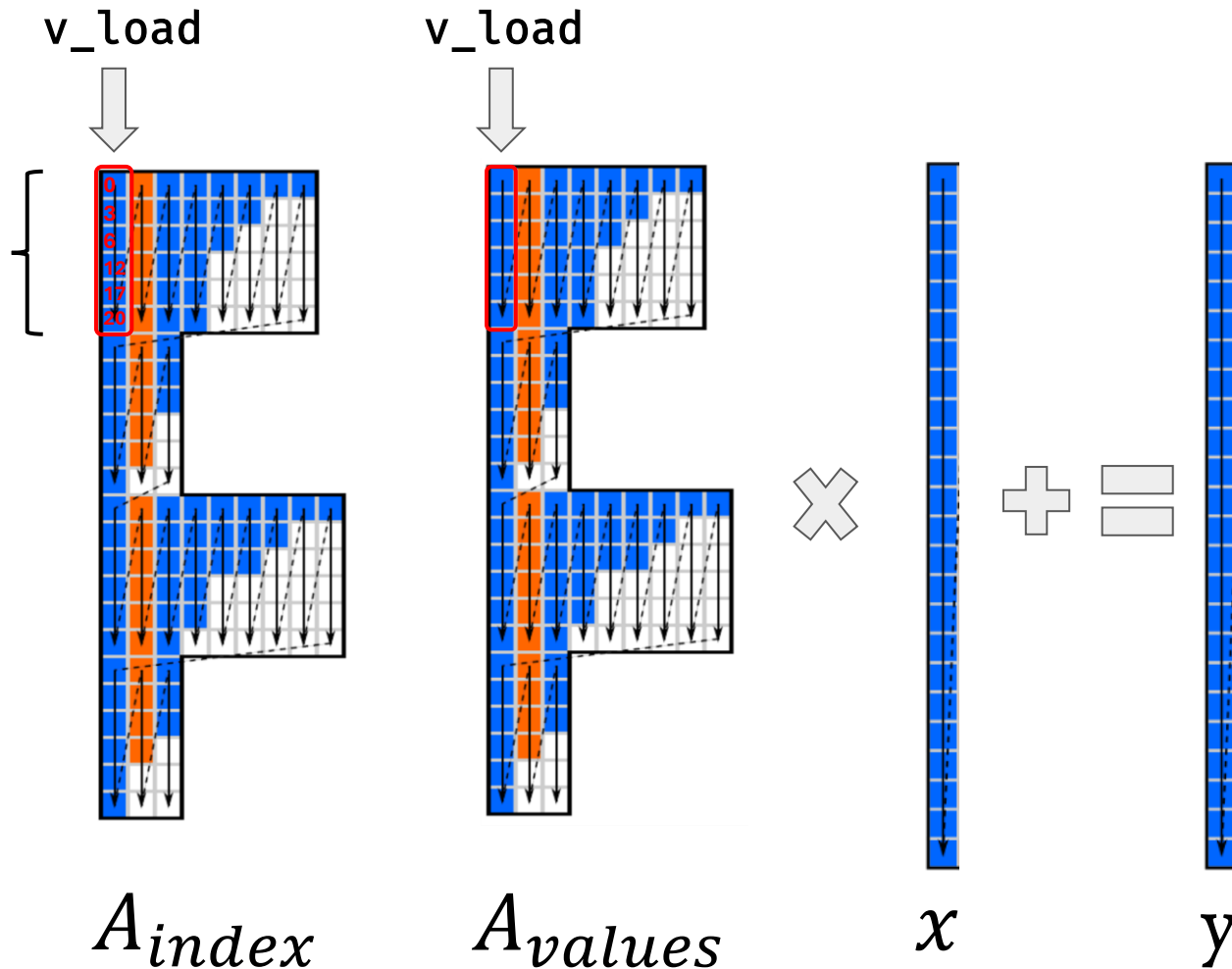


- A_{index} , A_{values} are stored by slice and within the slice by column
- Contiguous and predictable vector loads A_{values} and A_{index}
- No race conditions on y updates
- Especially in long vectors architectures (e.g. EPAC with $vlen=256$) we get huge benefits by hiding memory latency and better use of bandwidth during vector loads and stores [*]

M. Kreutzer et al., "A Unified Sparse Matrix Data Format for Efficient General Sparse Matrix-Vector Multiplication on Modern Processors with Wide SIMD Units," *SIAM J. Sci. Comput.*, 2014.

[*] Pablo Vizcaino et al. (2023) Short reasons for long vectors in HPC CPUs: a study based on RISC-V <https://doi.org/10.1145/3624062.3624231>

SpMV SELL-C- σ vectorization

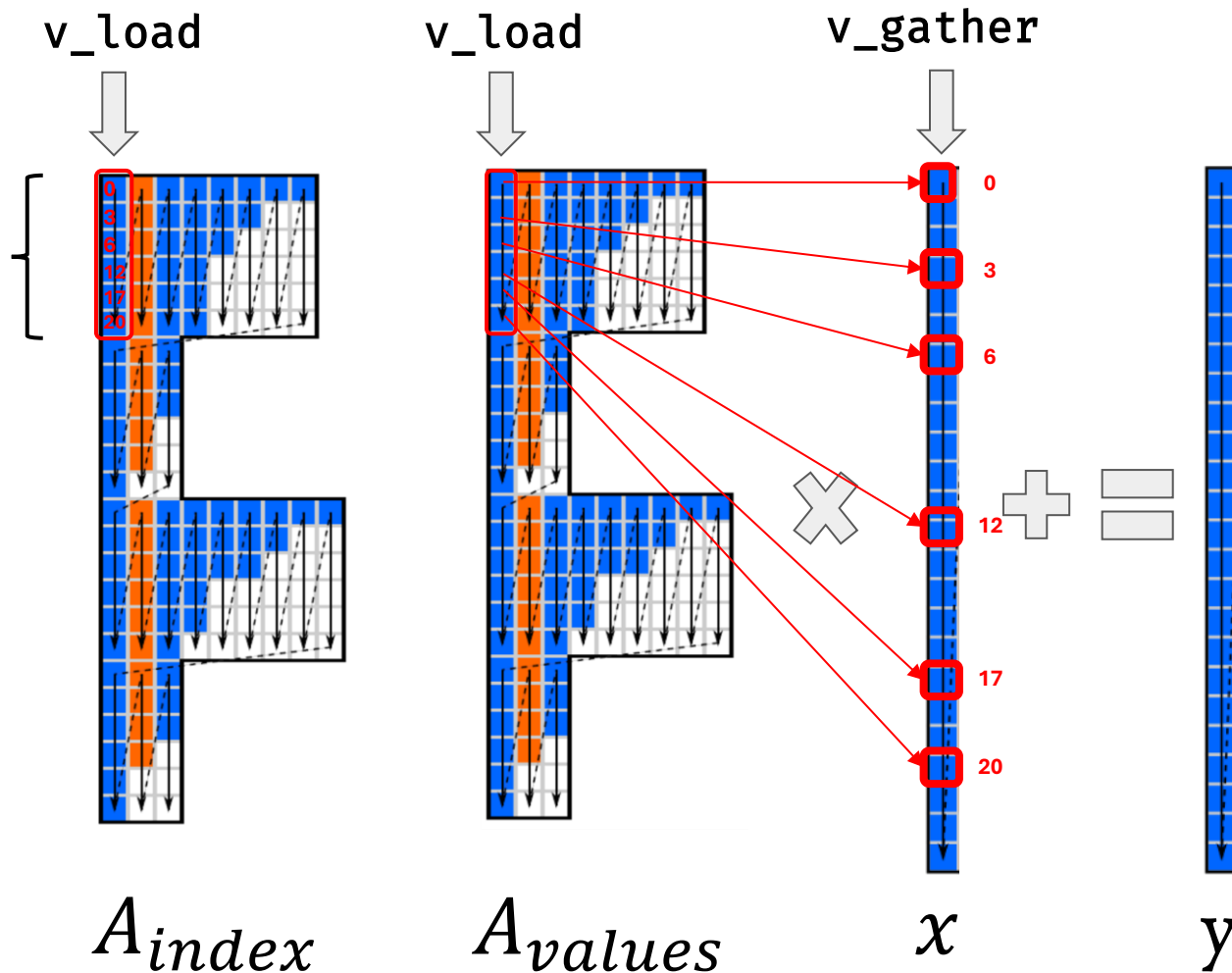


- A_{index} , A_{values} are stored by slice and within the slice by column
- Contiguous and predictable vector loads A_{values} and A_{index}
- No race conditions on y updates
- Especially in long vectors architectures (e.g. EPAC with $vlen=256$) we get huge benefits by hiding memory latency and better use of bandwidth during vector loads and stores [*]

M. Kreutzer et al., "A Unified Sparse Matrix Data Format for Efficient General Sparse Matrix-Vector Multiplication on Modern Processors with Wide SIMD Units," *SIAM J. Sci. Comput.*, 2014.

[*] Pablo Vizcaino et al. (2023) Short reasons for long vectors in HPC CPUs: a study based on RISC-V <https://doi.org/10.1145/3624062.3624231>

SpMV SELL-C- σ vectorization

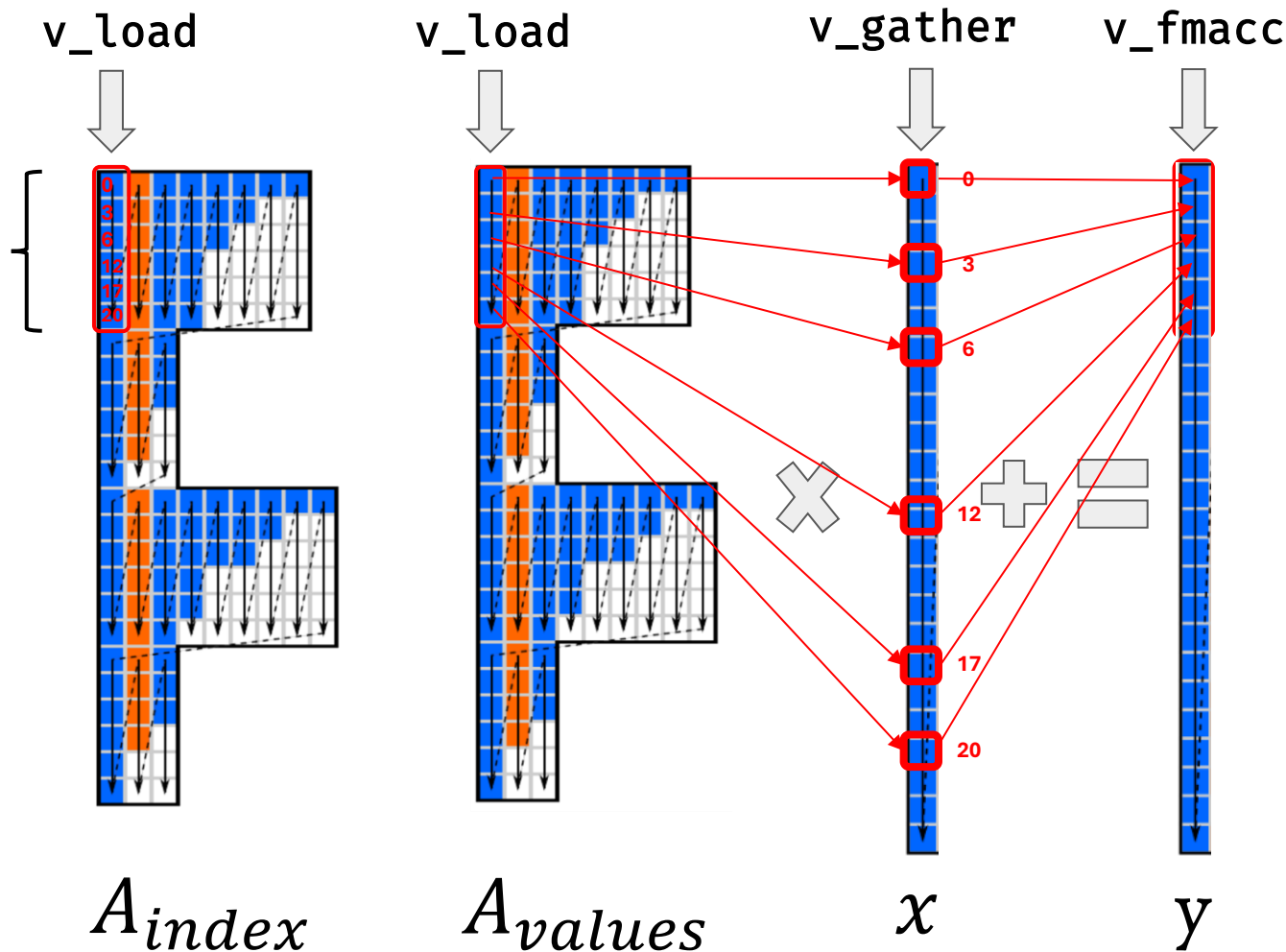


- A_{index} , A_{values} are stored by slice and within the slice by column
- Contiguous and predictable vector loads A_{values} and A_{index}
- No race conditions on y updates
- Especially in long vectors architectures (e.g. EPAC with $vlen=256$) we get huge benefits by hiding memory latency and better use of bandwidth during vector loads and stores [*]

M. Kreutzer et al., "A Unified Sparse Matrix Data Format for Efficient General Sparse Matrix-Vector Multiplication on Modern Processors with Wide SIMD Units," *SIAM J. Sci. Comput.*, 2014.

[*] Pablo Vizcaino et al. (2023) Short reasons for long vectors in HPC CPUs: a study based on RISC-V <https://doi.org/10.1145/3624062.3624231>

SpMV SELL-C- σ vectorization

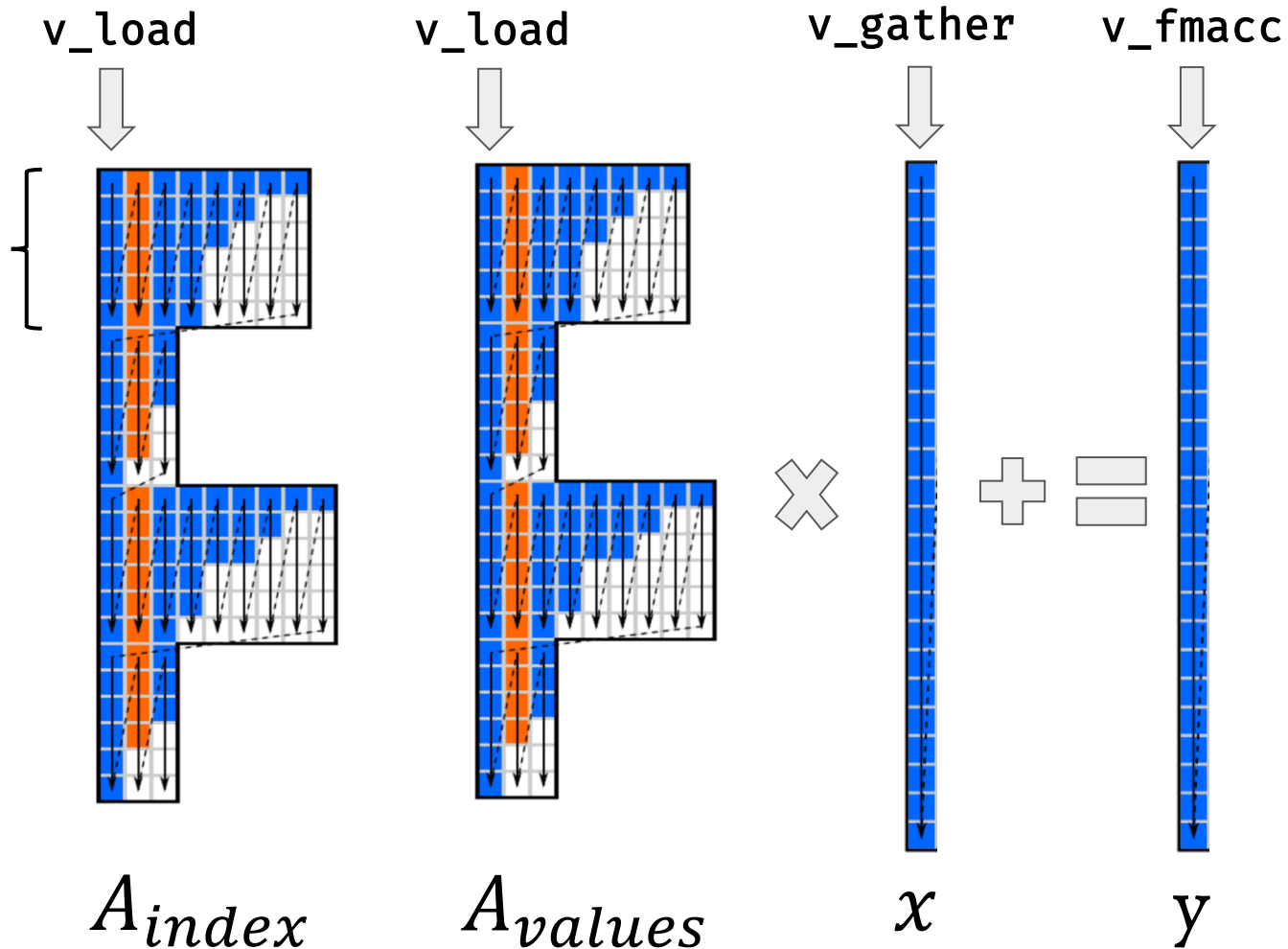


- A_{index} , A_{values} are stored by slice and within the slice by column
- Contiguous and predictable vector loads A_{values} and A_{index}
- No race conditions on y updates
- Especially in long vectors architectures (e.g. EPAC with $vlen=256$) we get huge benefits by hiding memory latency and better use of bandwidth during vector loads and stores [*]

M. Kreutzer et al., "A Unified Sparse Matrix Data Format for Efficient General Sparse Matrix-Vector Multiplication on Modern Processors with Wide SIMD Units," *SIAM J. Sci. Comput.*, 2014.

[*] Pablo Vizcaino et al. (2023) Short reasons for long vectors in HPC CPUs: a study based on RISC-V <https://doi.org/10.1145/3624062.3624231>

SpMV SELL-C- σ vectorization

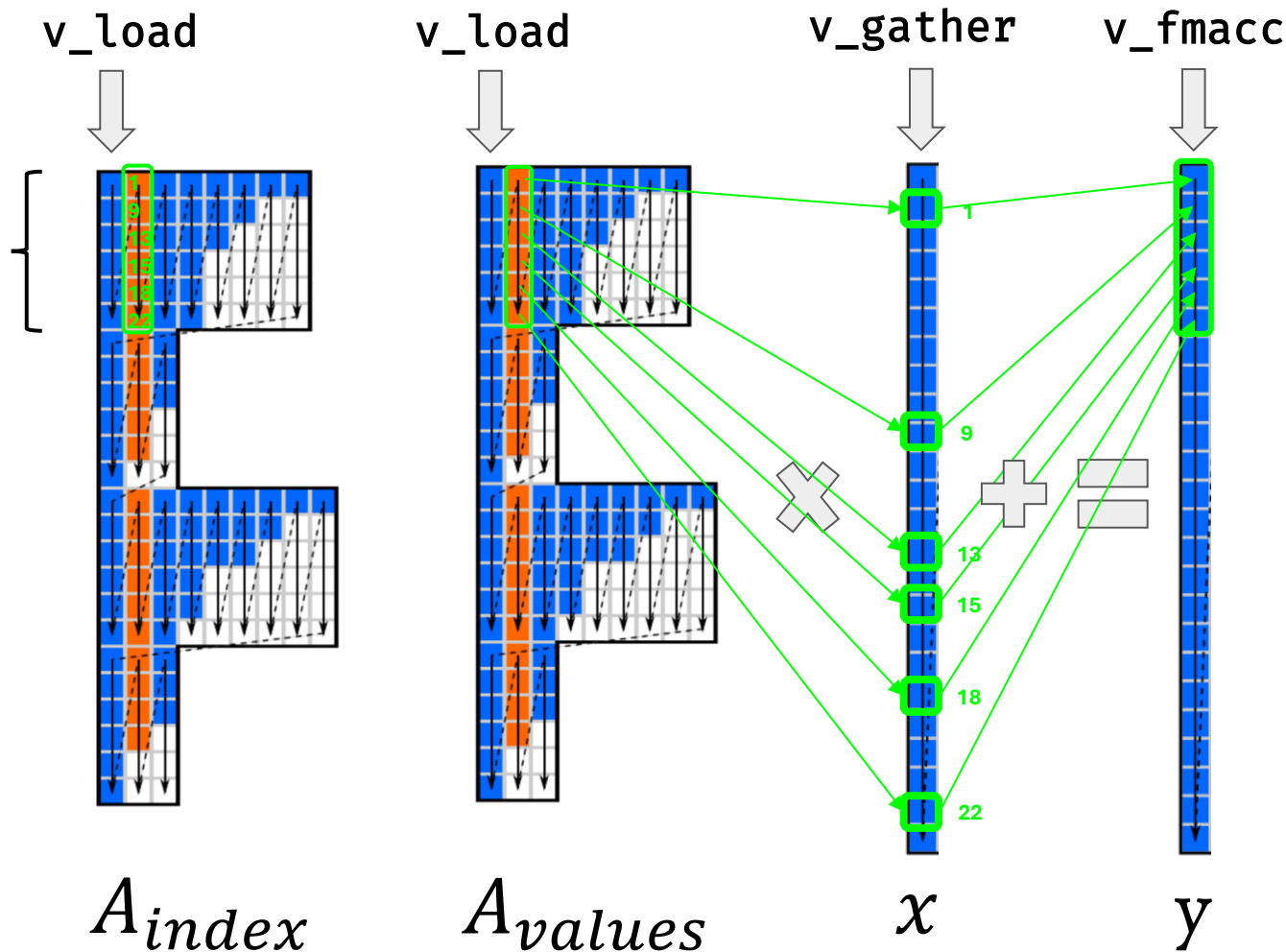


- A_{index} , A_{values} are stored by slice and within the slice by column
- Contiguous and predicable vector loads A_{values} and A_{index}
- No race conditions on y updates
- Especially in long vectors architectures (e.g. EPAC with $vlen=256$) we get huge benefits by hiding memory latency and better use of bandwidth during vector loads and stores [*]

M. Kreutzer et al., "A Unified Sparse Matrix Data Format for Efficient General Sparse Matrix-Vector Multiplication on Modern Processors with Wide SIMD Units," *SIAM J. Sci. Comput.*, 2014.

[*] Pablo Vizcaino et al. (2023) Short reasons for long vectors in HPC CPUs: a study based on RISC-V <https://doi.org/10.1145/3624062.3624231>

SpMV SELL-C- σ vectorization

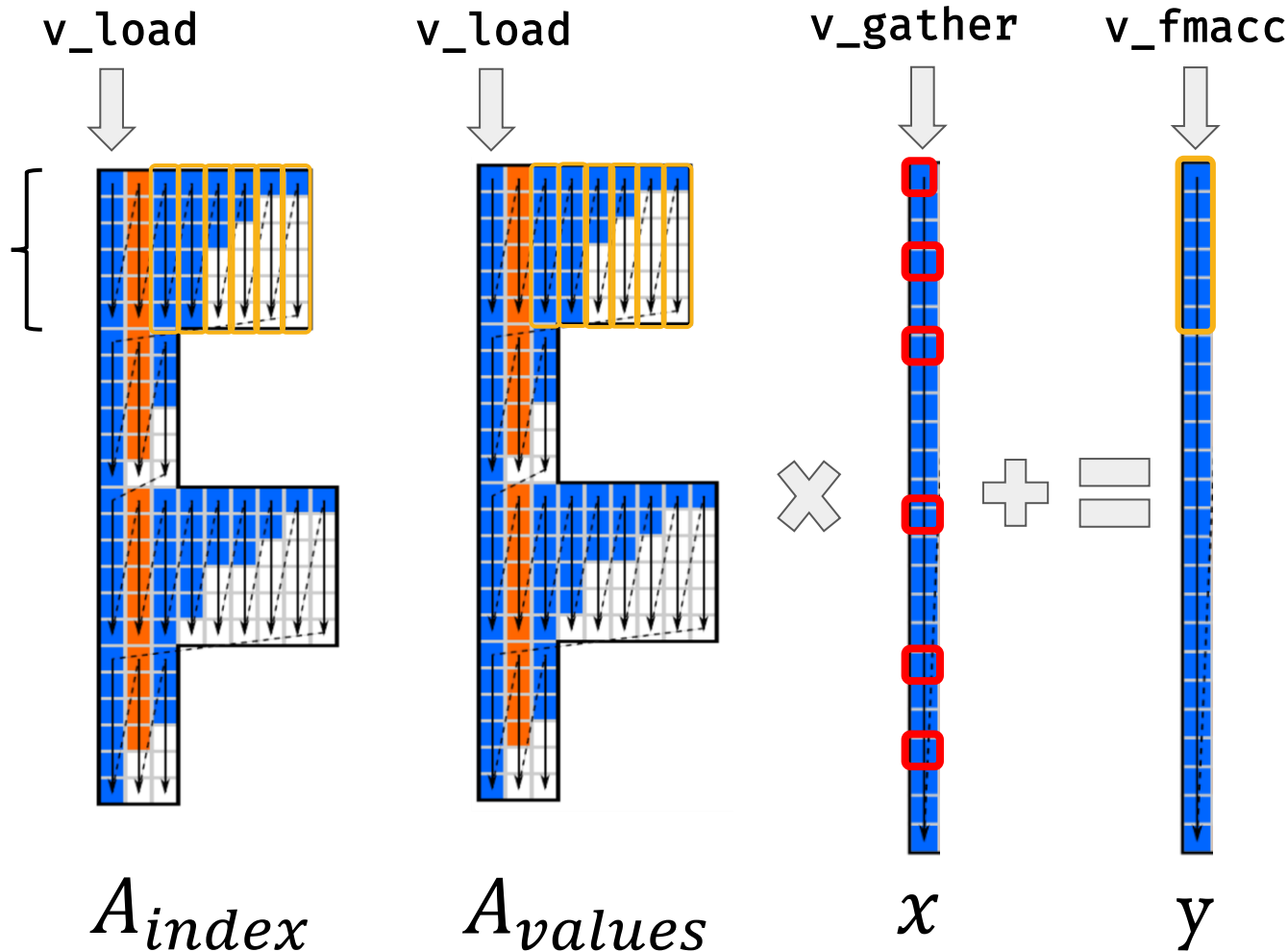


- A_{index} , A_{values} are stored by slice and within the slice by column
- Contiguous and predictable vector loads A_{values} and A_{index}
- No race conditions on y updates
- Especially in long vectors architectures (e.g. EPAC with $vlen=256$) we get huge benefits by hiding memory latency and better use of bandwidth during vector loads and stores [*]

M. Kreutzer et al., "A Unified Sparse Matrix Data Format for Efficient General Sparse Matrix-Vector Multiplication on Modern Processors with Wide SIMD Units," *SIAM J. Sci. Comput.*, 2014.

[*] Pablo Vizcaino et al. (2023) Short reasons for long vectors in HPC CPUs: a study based on RISC-V <https://doi.org/10.1145/3624062.3624231>

SpMV SELL-C- σ vectorization

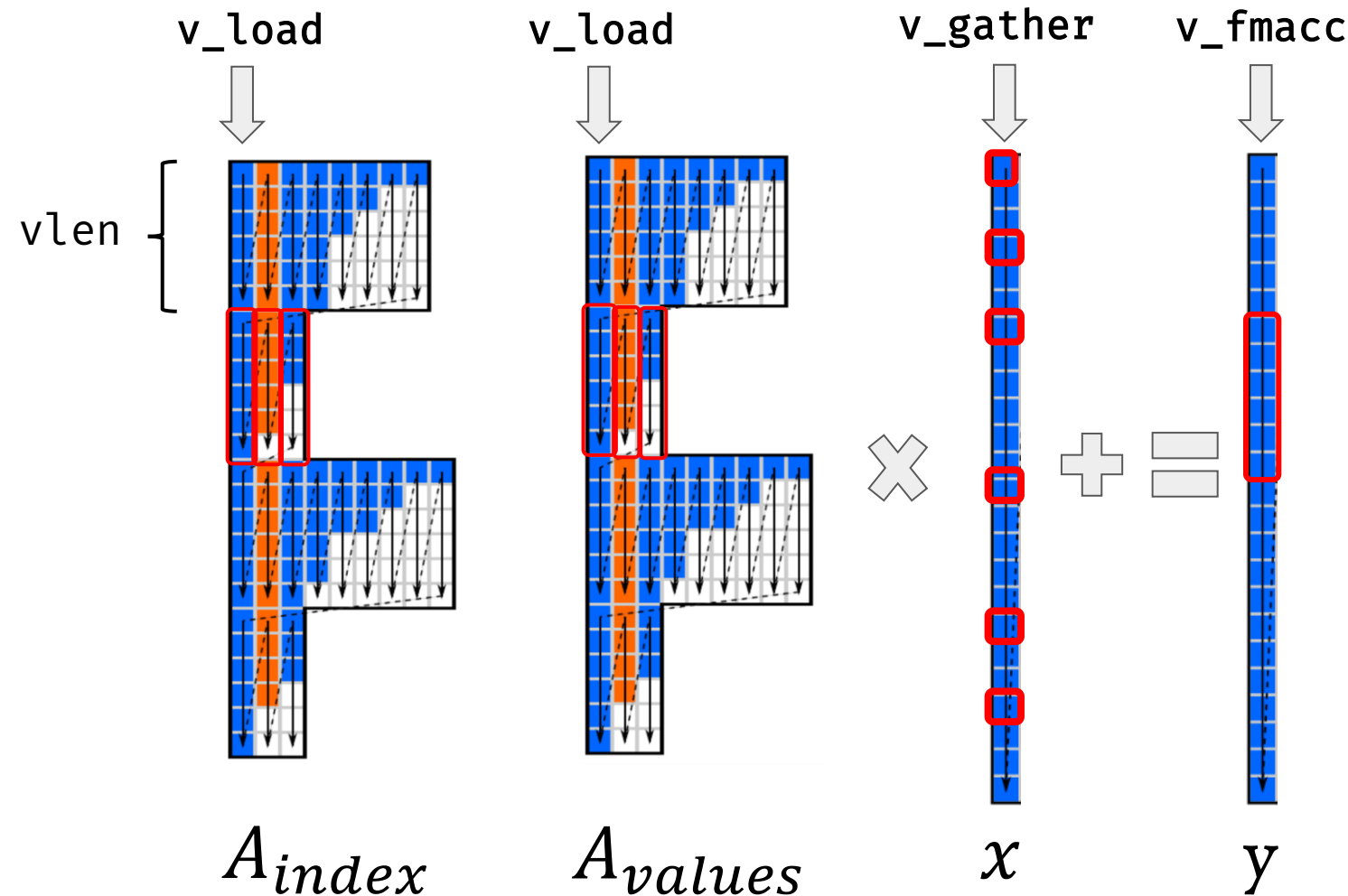


- A_{index} , A_{values} are stored by slice and within the slice by column
- Contiguous and predictable vector loads A_{values} and A_{index}
- No race conditions on y updates
- Especially in long vectors architectures (e.g. EPAC with $vlen=256$) we get huge benefits by hiding memory latency and better use of bandwidth during vector loads and stores [*]

M. Kreutzer et al., "A Unified Sparse Matrix Data Format for Efficient General Sparse Matrix-Vector Multiplication on Modern Processors with Wide SIMD Units," *SIAM J. Sci. Comput.*, 2014.

[*] Pablo Vizcaino et al. (2023) Short reasons for long vectors in HPC CPUs: a study based on RISC-V <https://doi.org/10.1145/3624062.3624231>

SpMV SELL-C- σ vectorization

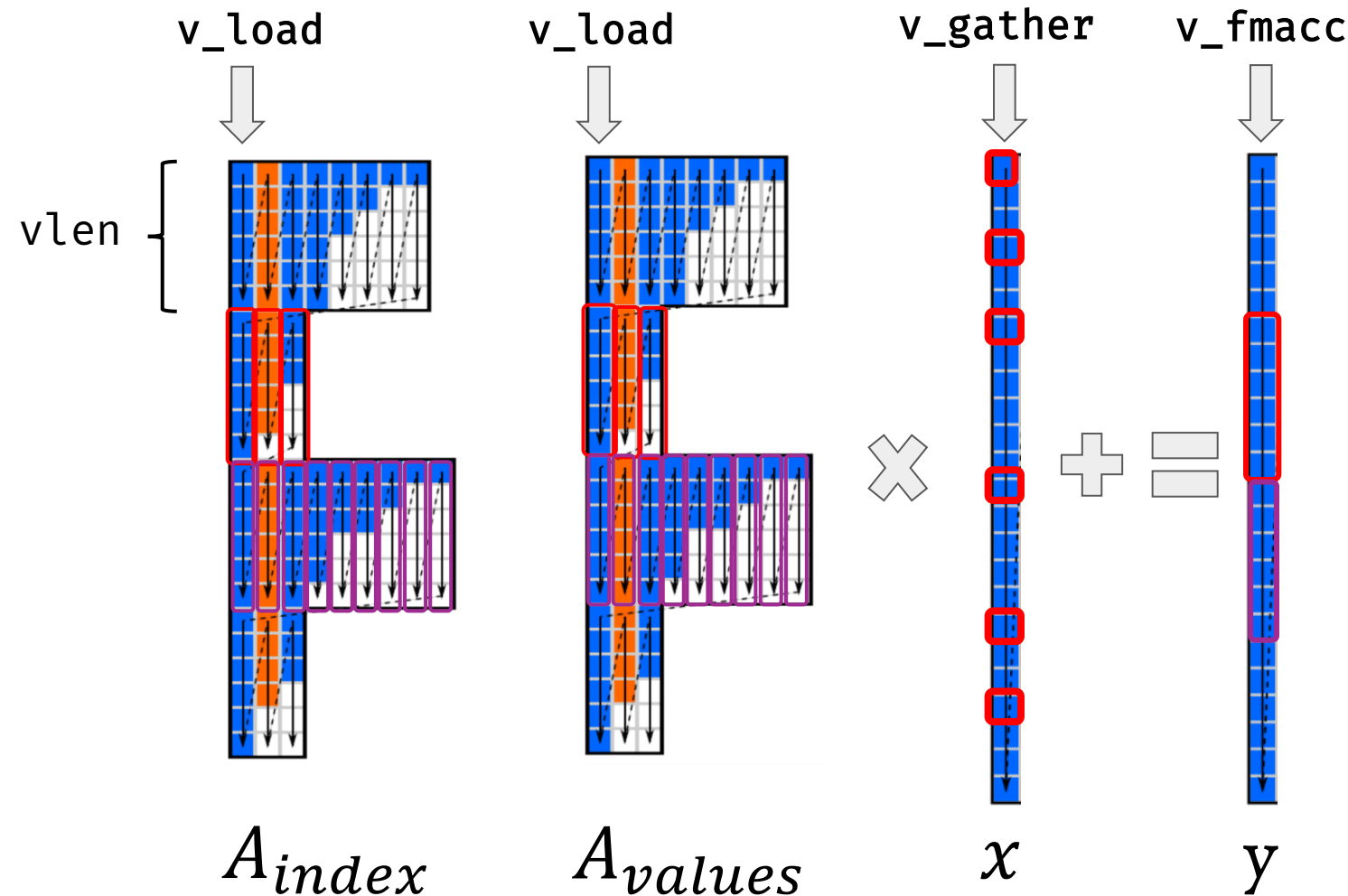


- A_{index} , A_{values} are stored by slice and within the slice by column
- Contiguous and predictable vector loads A_{values} and A_{index}
- No race conditions on y updates
- Especially in long vectors architectures (e.g. EPAC with $vlen=256$) we get huge benefits by hiding memory latency and better use of bandwidth during vector loads and stores [*]

M. Kreutzer et al., "A Unified Sparse Matrix Data Format for Efficient General Sparse Matrix-Vector Multiplication on Modern Processors with Wide SIMD Units," *SIAM J. Sci. Comput.*, 2014.

[*] Pablo Vizcaino et al. (2023) Short reasons for long vectors in HPC CPUs: a study based on RISC-V <https://doi.org/10.1145/3624062.3624231>

SpMV SELL-C- σ vectorization

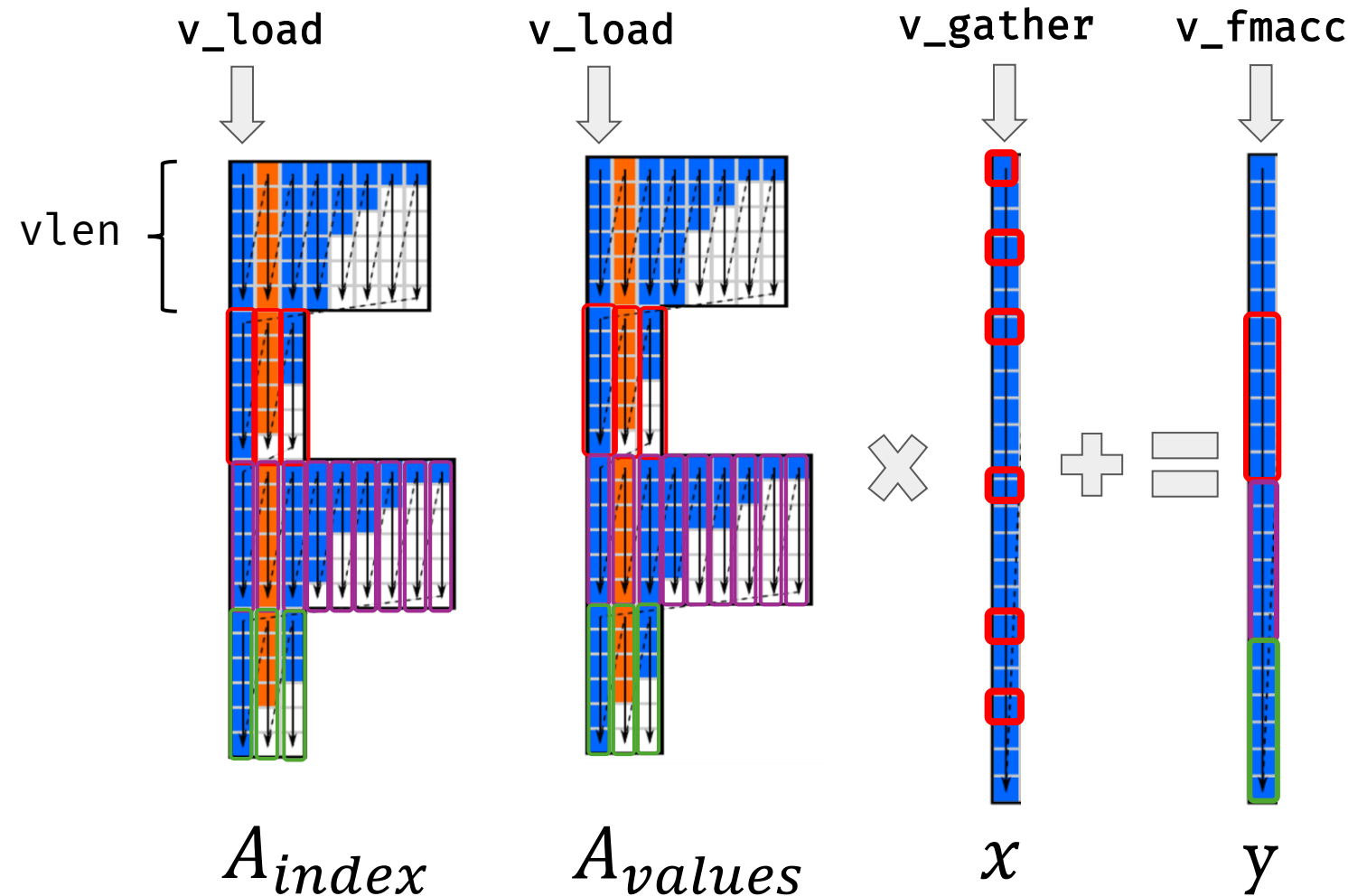


- A_{index} , A_{values} are stored by slice and within the slice by column
- Contiguous and predictable vector loads A_{values} and A_{index}
- No race conditions on y updates
- Especially in long vectors architectures (e.g. EPAC with $vlen=256$) we get huge benefits by hiding memory latency and better use of bandwidth during vector loads and stores [*]

M. Kreutzer et al., "A Unified Sparse Matrix Data Format for Efficient General Sparse Matrix-Vector Multiplication on Modern Processors with Wide SIMD Units," *SIAM J. Sci. Comput.*, 2014.

[*] Pablo Vizcaino et al. (2023) Short reasons for long vectors in HPC CPUs: a study based on RISC-V <https://doi.org/10.1145/3624062.3624231>

SpMV SELL-C- σ vectorization



- A_{index} , A_{values} are stored by slice and within the slice by column
- Contiguous and predictable vector loads A_{values} and A_{index}
- No race conditions on y updates
- Especially in long vectors architectures (e.g. EPAC with $vlen=256$) we get huge benefits by hiding memory latency and better use of bandwidth during vector loads and stores [*]

M. Kreutzer et al., "A Unified Sparse Matrix Data Format for Efficient General Sparse Matrix-Vector Multiplication on Modern Processors with Wide SIMD Units," *SIAM J. Sci. Comput.*, 2014.

[*] Pablo Vizcaino et al. (2023) Short reasons for long vectors in HPC CPUs: a study based on RISC-V <https://doi.org/10.1145/3624062.3624231>

RISC-V Vector target architectures

Sophon - SG2044

(XuanTie C920v2 core) RVV 1.0

- 64-core, general purpose server
- 12 stage out-of-order multiple issues superscalar pipeline design
- **128-bit vector registers (2 DP elements)**
- 64 kB L1, 2 MB L2, 64 MB L3 data cache
- Target: Commercial general-purpose servers
- **Short-vector architecture**



RISC-V Vector target architectures

Sophon - SG2044

(XuanTie C920v2 core) RVV 1.0

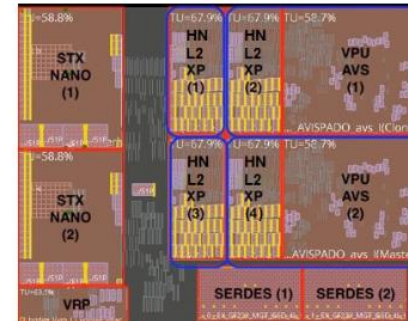
- 64-core, general purpose server
- 12 stage out-of-order multiple issues superscalar pipeline design
- **128-bit vector registers (2 DP elements)**
- 64 kB L1, 2 MB L2, 64 MB L3 data cache
- Target: Commercial general-purpose servers
- **Short-vector architecture**



EPI – EPAC 1.5

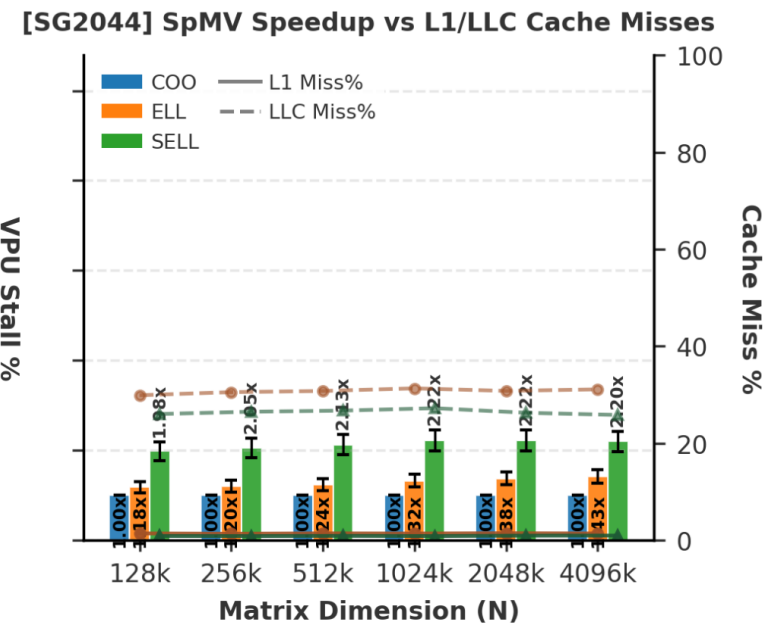
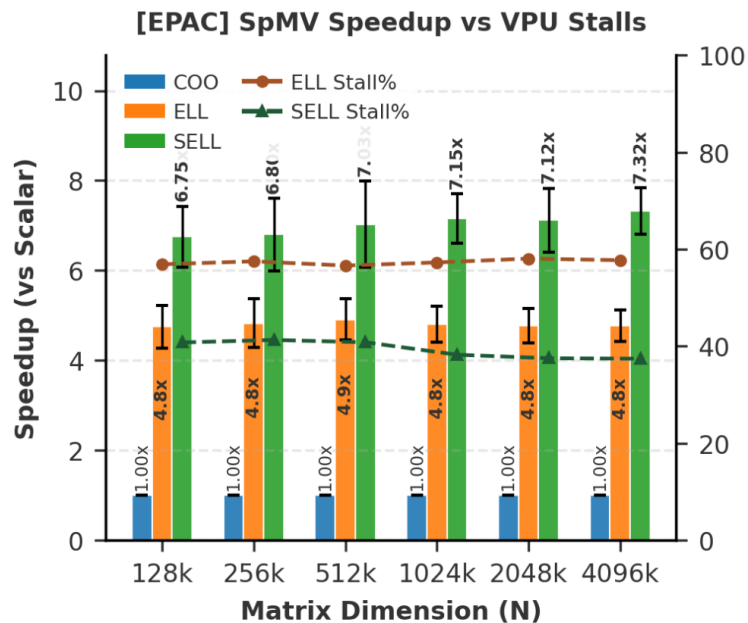
(VEC tile) RVV 0.7.1

- Prototypal vector accelerator
- Scalar in-order RISC-V Avispado core + Vector processing unit (VPU)
- **16384-bit vector registers (256 DP elements)**
- 32 kB L1 (scalar core), 256 kB L2 data cache
- Target: Traditional HPC workloads, FP64
- **Long-vector architecture**



SpMV vectorization speedup

EPI – EPAC 1.5
(Long-vector)
 $vlmax=256$



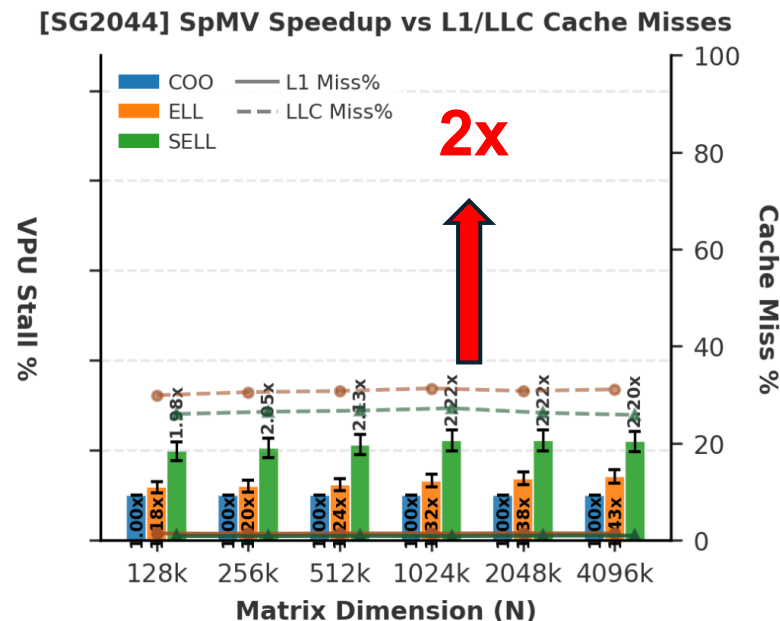
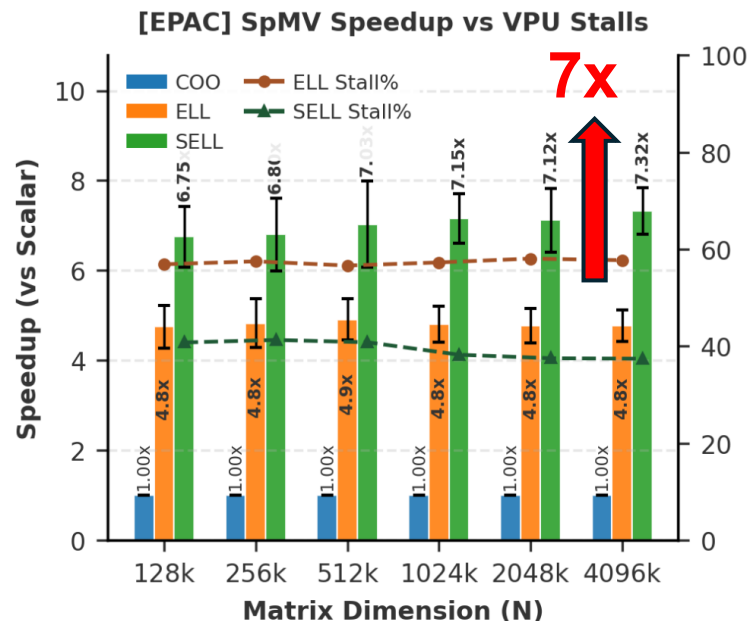
Sophon SG2044
(Short-vector)
 $vlmax=2$



SpMV vectorization speedup

- 7x speedup on EPAC (long vector) and 2x on SG2044 (short vector)

EPI – EPAC 1.5
(Long-vector)
 $vlmax=256$



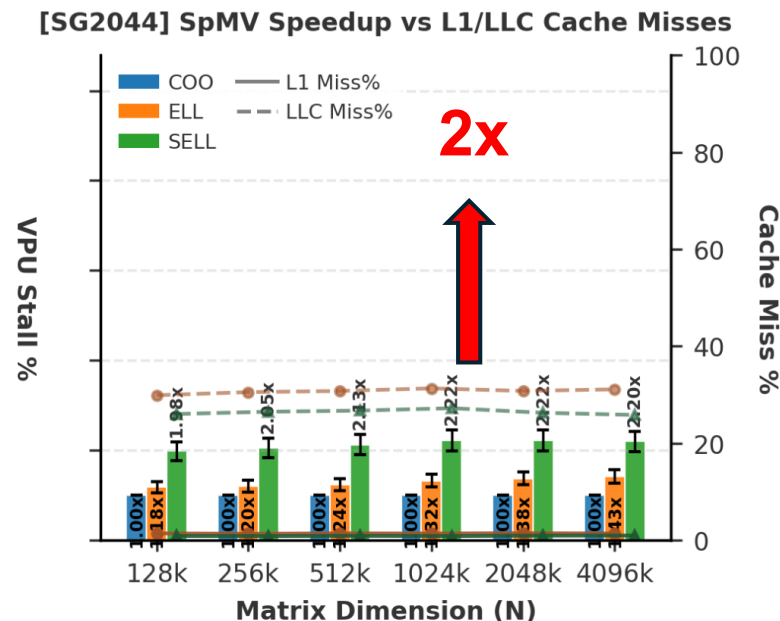
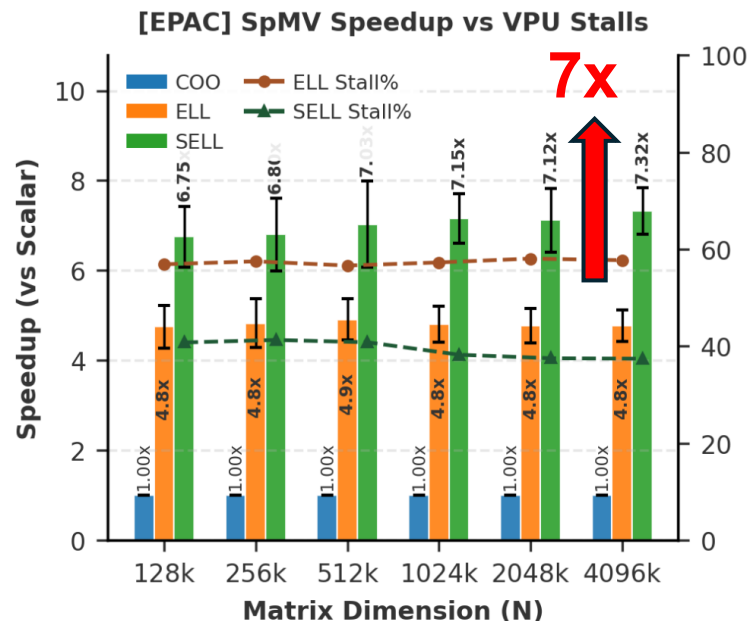
Sophon SG2044
(Short-vector)
 $vlmax=2$



SpMV vectorization speedup

- 7x speedup on EPAC (long vector) and 2x on SG2044 (short vector)
- The long-vector approach proves to be very efficient for SpMV in SELL format.

EPI – EPAC 1.5
(Long-vector)
 $vlmax=256$



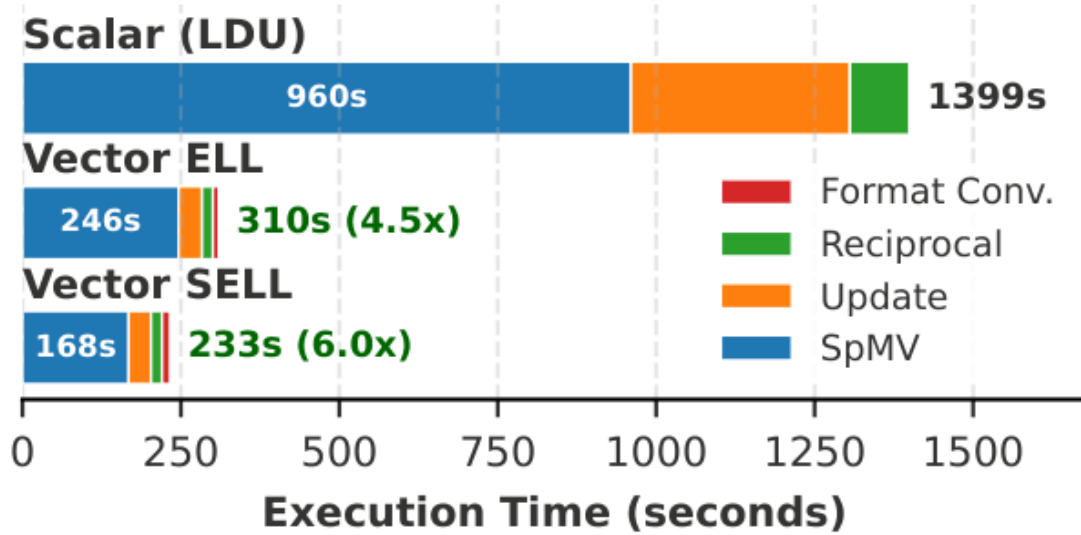
Sophon SG2044
(Short-vector)
 $vlmax=2$



Smoother speedup

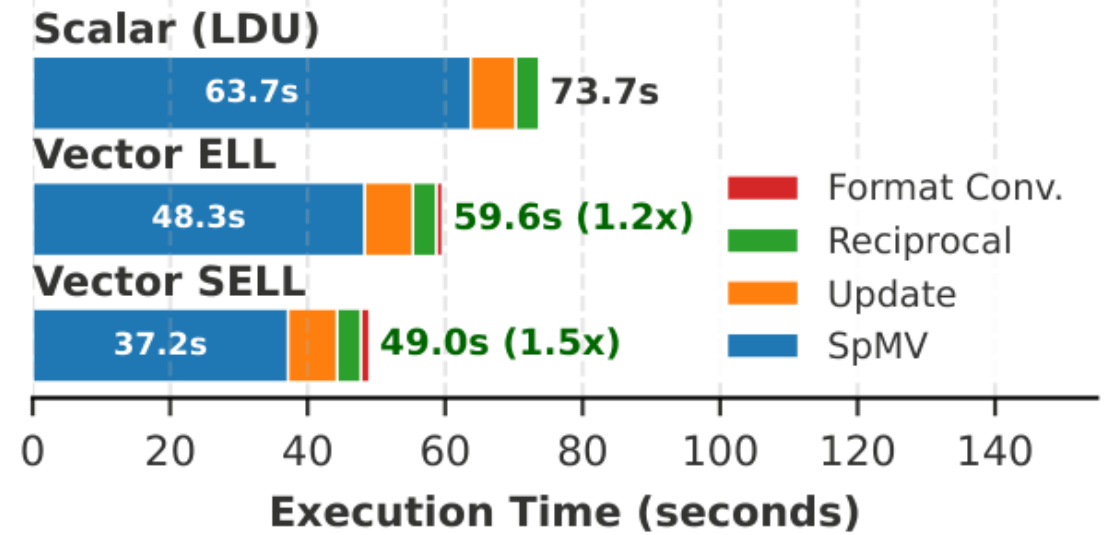
EPI – EPAC 1.5 (Long-vector)

vlmax=256



Sophon SG2044 (Short-vector)

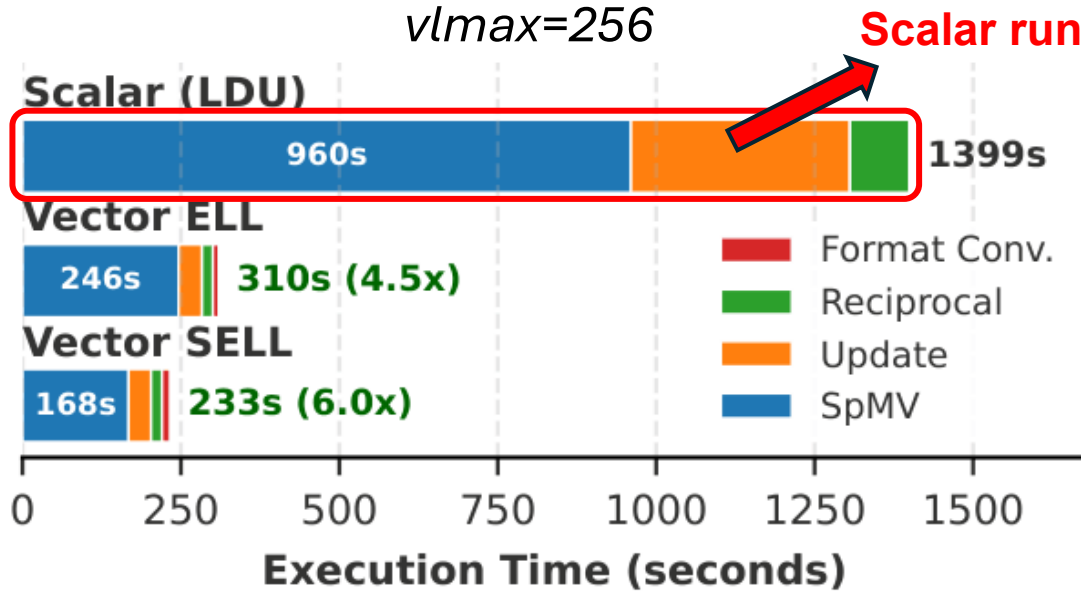
vlmax=2



Smoother speedup

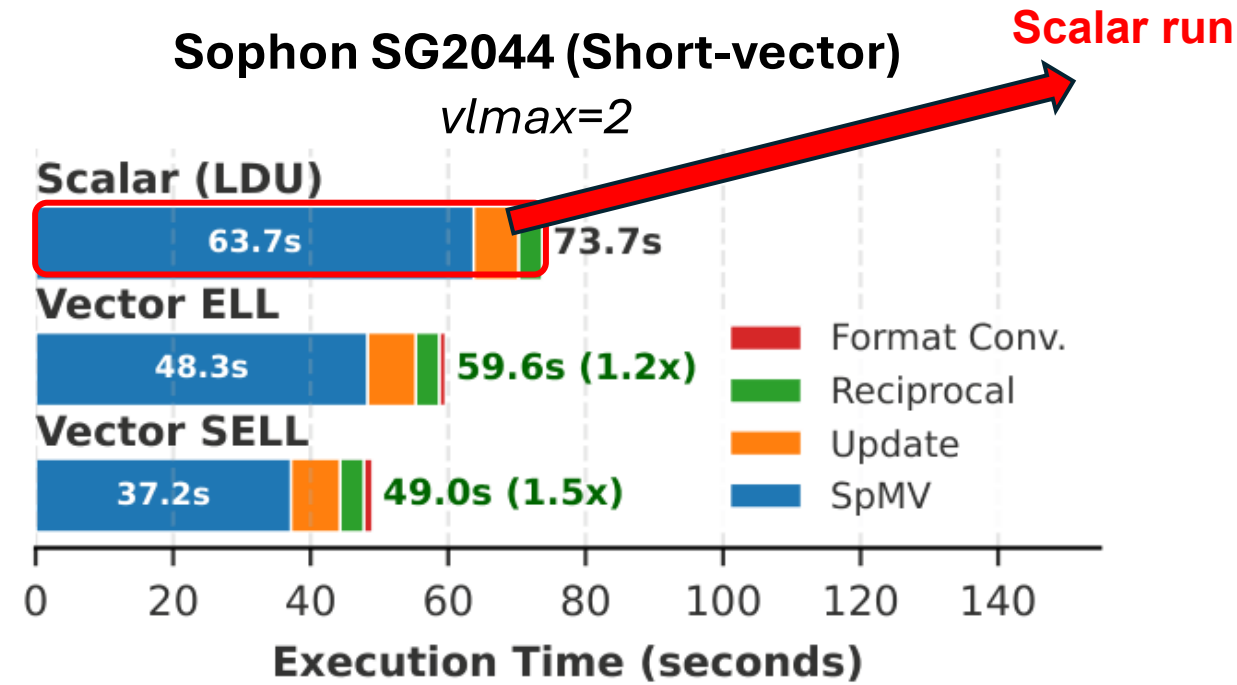
EPI – EPAC 1.5 (Long-vector)

$vlmax=256$



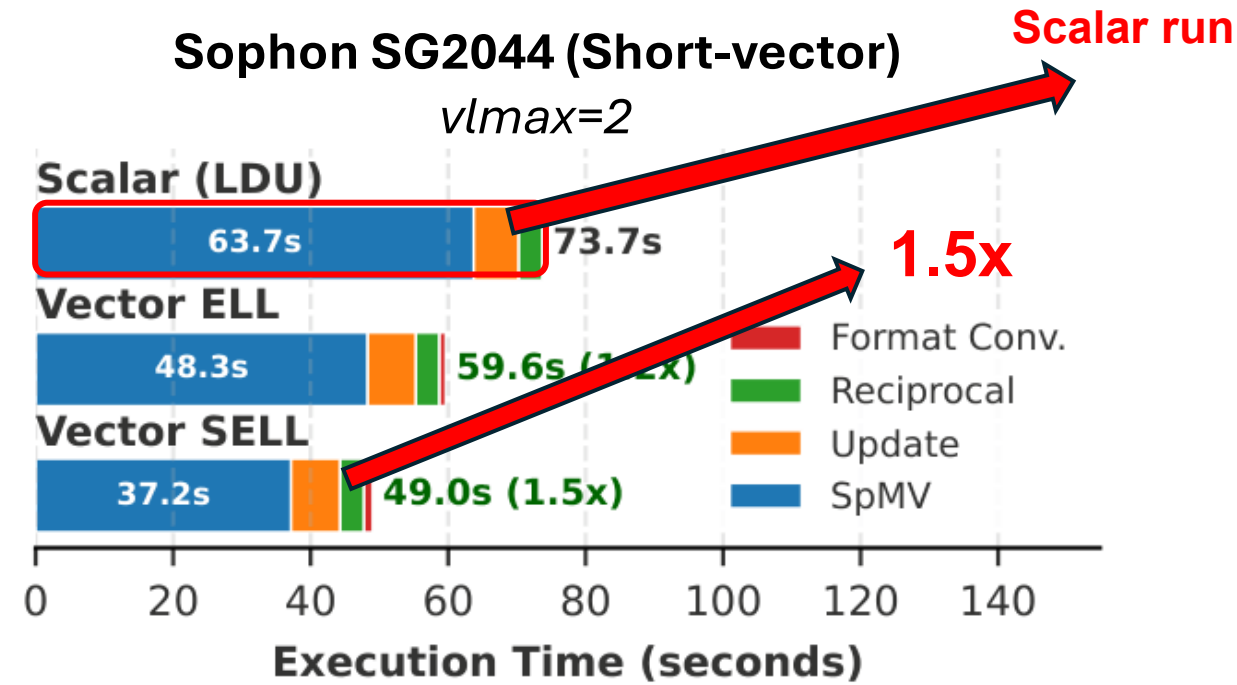
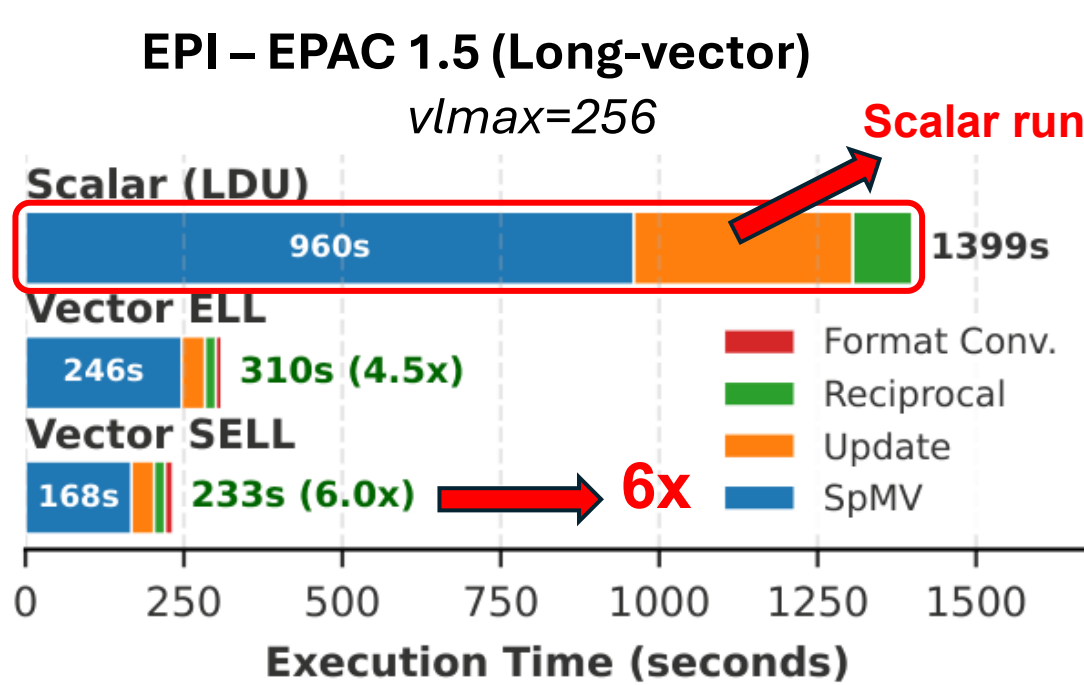
Sophon SG2044 (Short-vector)

$vlmax=2$



Smoother speedup

- 6x speedup on EPAC (long vector) and 1.5x on SG2044 (short vector)
- Marix format conversion is negligible (only 1-2% of total time)

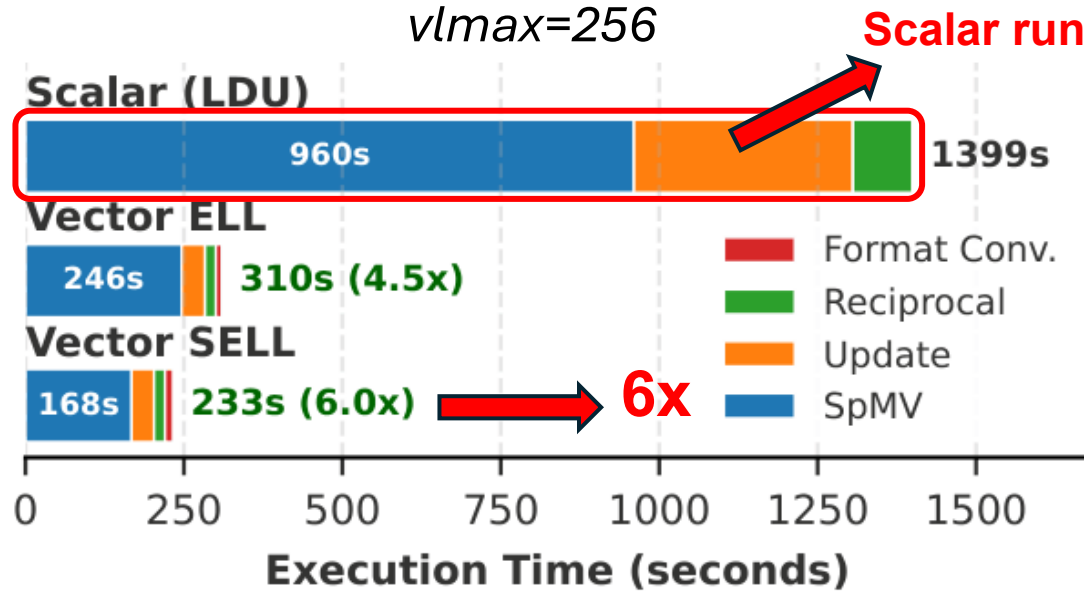


Smoother speedup

- 6x speedup on EPAC (long vector) and 1.5x on SG2044 (short vector)
- Marix format conversion is negligible (only 1-2% of total time)
- The long-vector approach proves to be more efficient for both sparse kernels (SpMV) and dense kernels.

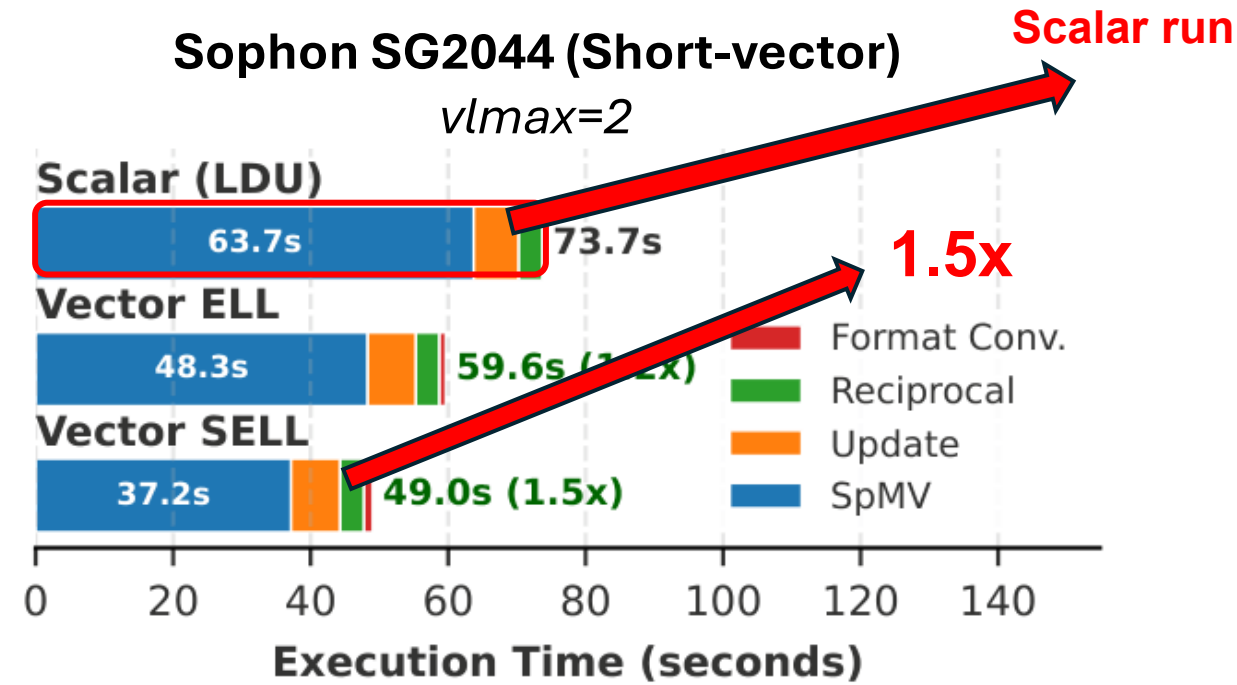
EPI – EPAC 1.5 (Long-vector)

vlmax=256



Sophon SG2044 (Short-vector)

vlmax=2



Conclusions and future works



Conclusions and future works

Conclusions

- Non-invasive OpenFOAM optimization on vector architectures has been proved using run-time format conversion.

Conclusions and future works

Conclusions

- Non-invasive OpenFOAM optimization on vector architectures has been proved using run-time format conversion.
- The effectiveness of long vector architectures on memory-bound kernels and CFD solvers has been proven.



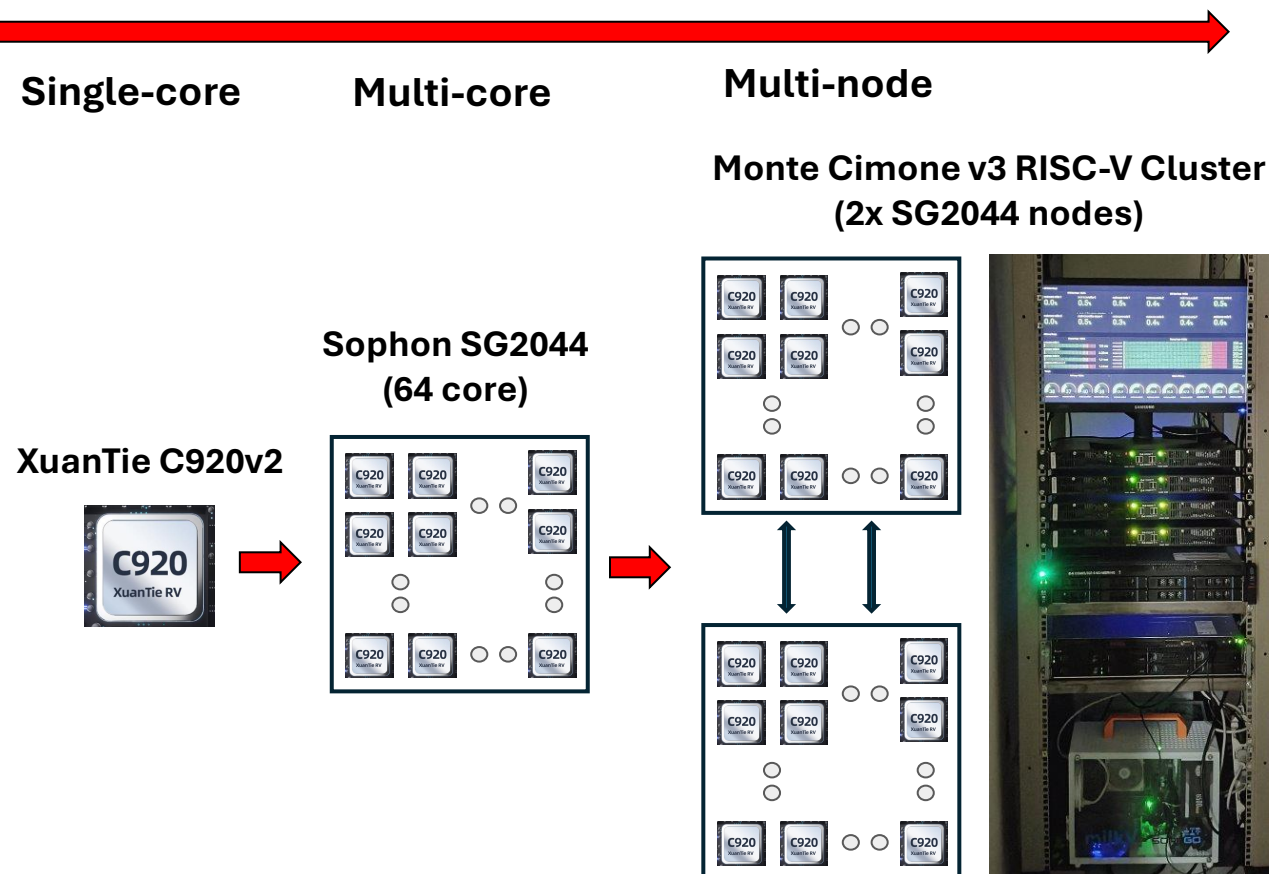
Conclusions and future works

Conclusions

- Non-invasive OpenFOAM optimization on vector architectures has been proved using run-time format conversion.
- The effectiveness of long vector architectures on memory-bound kernels and CFD solvers has been proven.

Next steps

Strong scaling study





Conclusions and future works

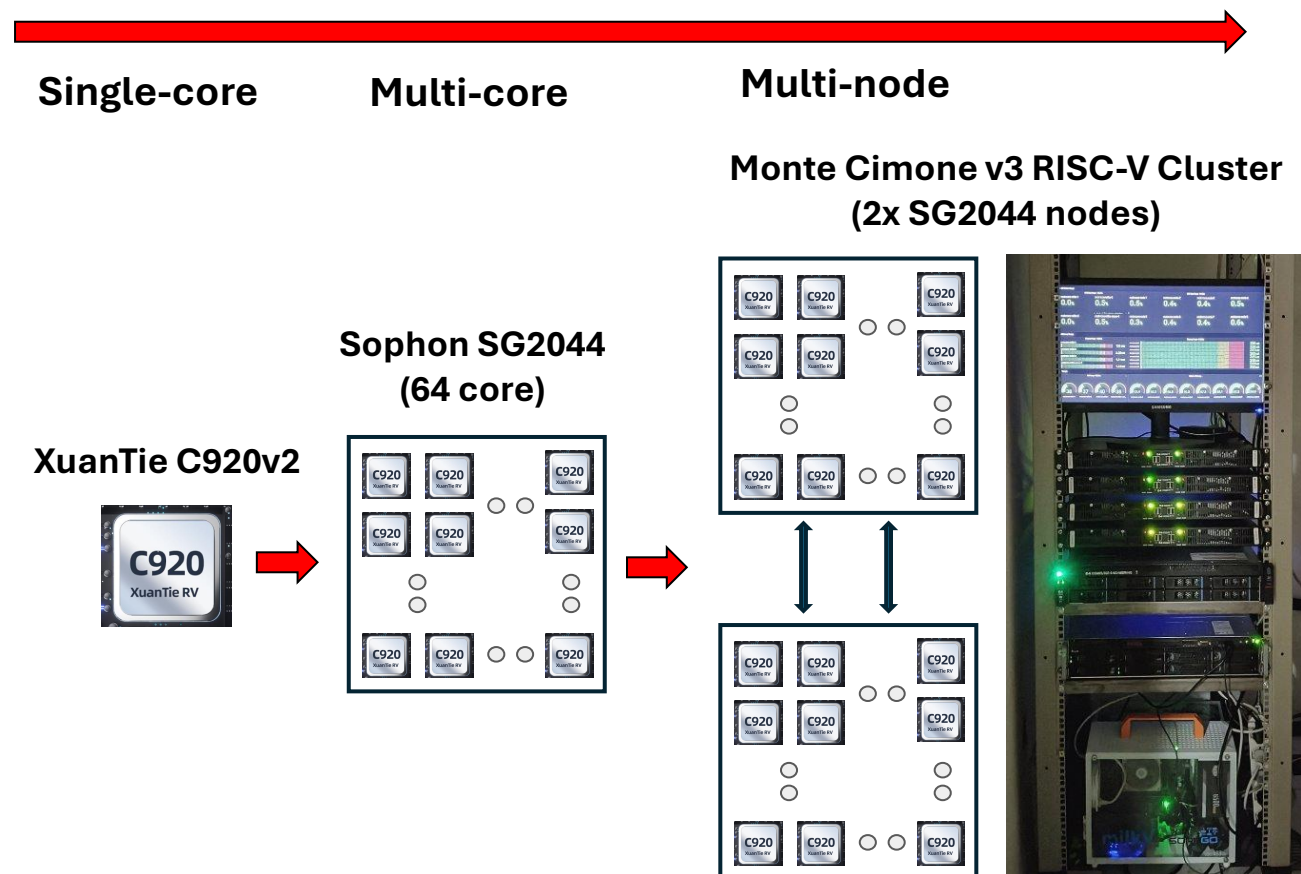
Conclusions

- Non-invasive OpenFOAM optimization on vector architectures has been proved using run-time format conversion.
- The effectiveness of long vector architectures on memory-bound kernels and CFD solvers has been proven.

Next steps

- Complete the vectorization of other components of the multigrid solver

Strong scaling study





Conclusions and future works

Conclusions

- Non-invasive OpenFOAM optimization on vector architectures has been proved using run-time format conversion.
- The effectiveness of long vector architectures on memory-bound kernels and CFD solvers has been proven.

Next steps

- Complete the vectorization of other components of the multigrid solver
- Strong scaling scalability study using Monte Cimone HPC cluster

Strong scaling study

