

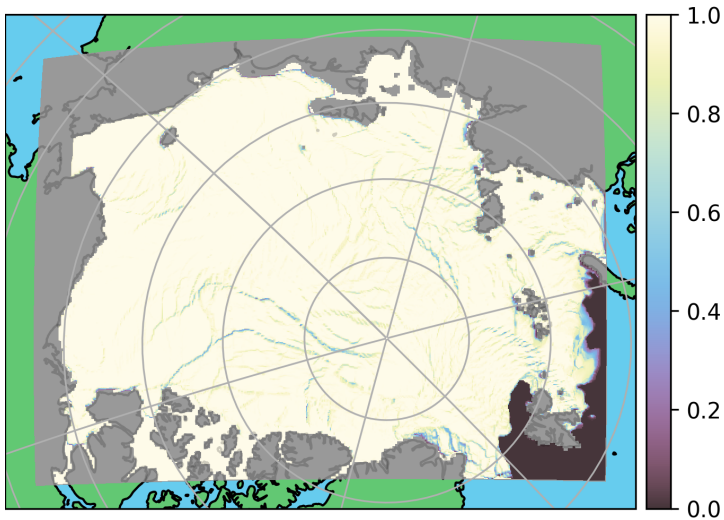
Development of a GPU-accelerated, Finite Element based Dynamical Core for Sea Ice

Robert Jendersie¹ Christian Lessig² Thomas Richter¹

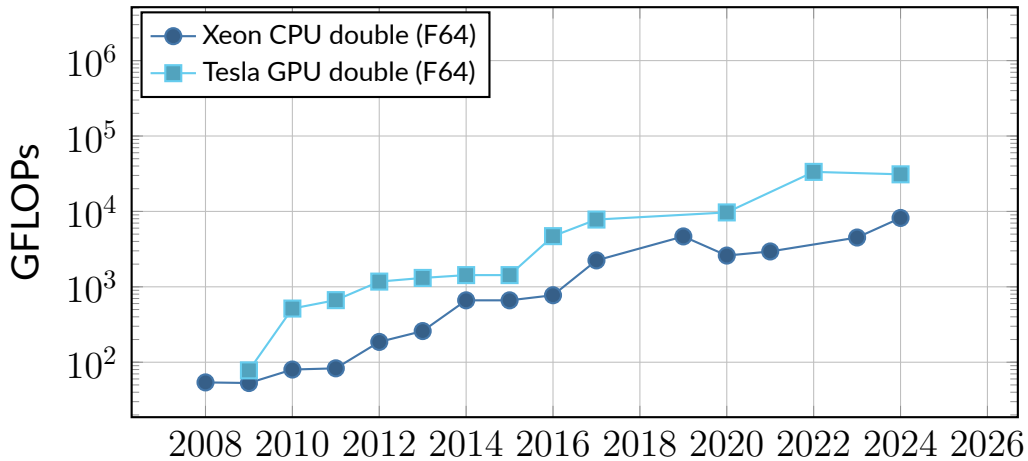
¹Otto-von-Guericke Universität Magdeburg

²ECMWF Bonn

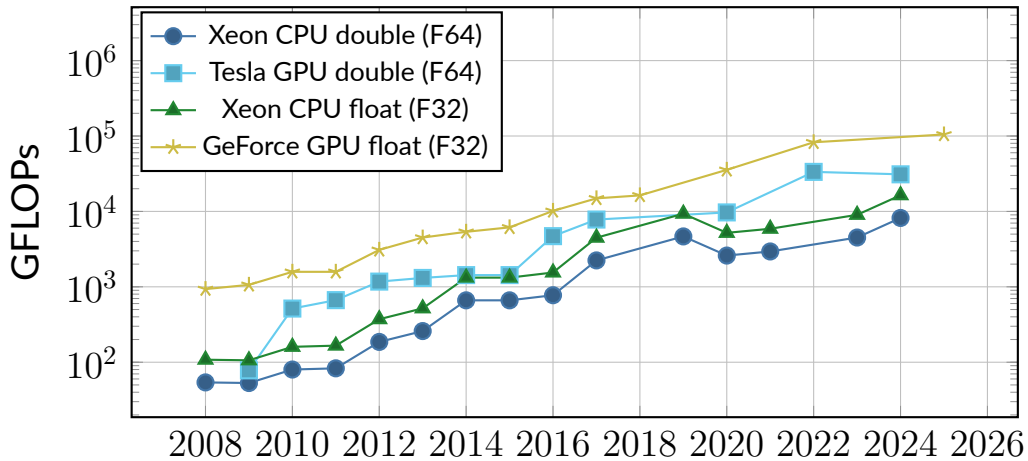
High-Resolution Sea-Ice



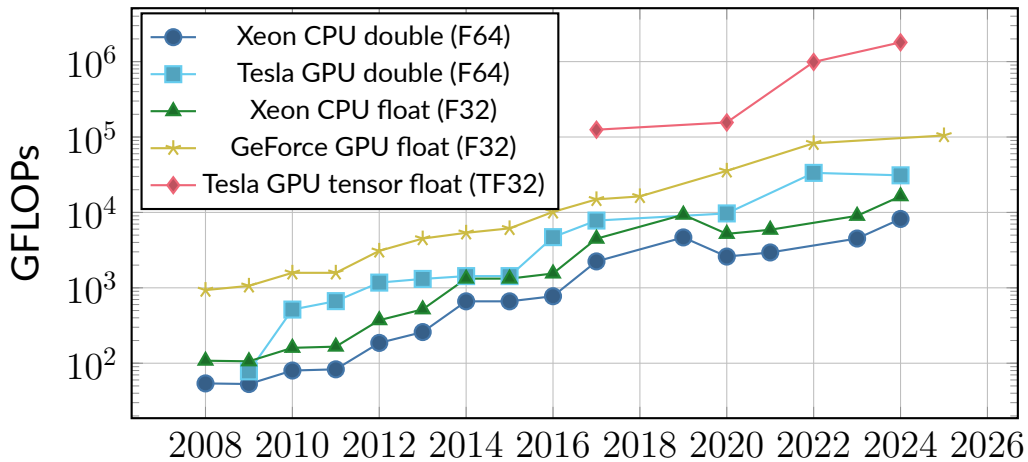
Why GPUs?



Why GPUs?



Why GPUs?



neXtSIM-DG Model

- Momentum and advection

$$\rho_{\text{ice}} H \partial_t \mathbf{v} = \text{div } \boldsymbol{\sigma}(\mathbf{v}, A, H) + F, \quad (1)$$

$$\partial_t H + \text{div}(\mathbf{v} H) = S_H \quad \partial_t A + \text{div}(\mathbf{v} A) = S_A \quad (2)$$

- viscous plastic rheology¹ (with mEVP iteration²)

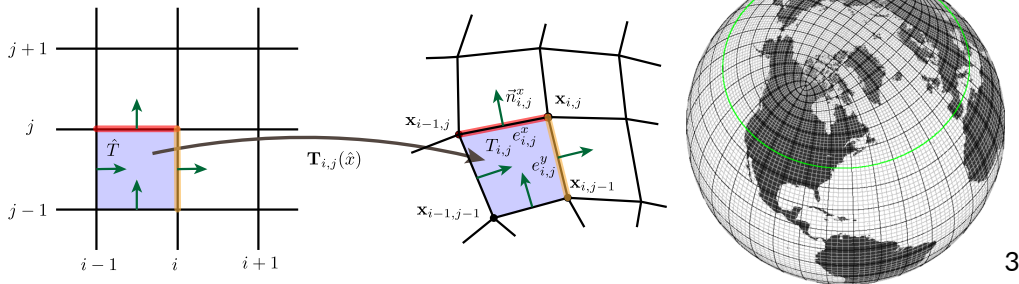
$$\boldsymbol{\sigma}^{(p)} = \frac{\alpha}{1 + \alpha} \boldsymbol{\sigma}^{(p-1)} + \frac{1}{1 + \alpha} \boldsymbol{\sigma}_{vp}(\mathbf{v}^{(p-1)}, A, H) \quad (3)$$

$$\boldsymbol{\sigma}_{vp}(\mathbf{v}, A, H) = \eta(\nabla \mathbf{v} + \nabla \mathbf{v}^T) + \zeta \text{div}(\mathbf{v}) I - \frac{P}{2} I,$$

¹Hibler, "A Dynamic Thermodynamic Sea Ice Model", 1979.

²Bouillon et al., "The elastic-viscous-plastic method revisited", 2013.

neXtSIM-DG Discretization⁴



3

- cG elements for v , dG elements for A , H , σ and explicit time-stepping

³<https://www.geomar.de/en/ocean-models>

⁴Richter et al., "A dynamical core based on a discontinuous Galerkin method for higher-order finite-element sea ice modeling", 2023.

Implementation

```

void StressUpdateHighOrder(Matrix&  $S^{11}$ , Matrix&  $S^{12}$ , Matrix&  $S^{22}$ ,
    const Matrix&  $E^{11}$ , const Matrix&  $E^{12}$ , const Matrix&  $E^{22}$ ,
    const Matrix&  $H$ , const Matrix&  $A$ ) {
    // update stress for every element
    for (int i = 0; i < N; ++i) {
        Vector h = max{0,  $H_{i,*}$   $\text{PSI}_A$ }
        Vector a = min{1, max{0,  $A_{i,*}$   $\text{PSI}_A$ }}
        Vector  $e^{11}$  =  $E_{i,*}^{11}$   $\text{PSI}_S$ 
        Vector  $e^{12}$  =  $E_{i,*}^{12}$   $\text{PSI}_S$ 
        Vector  $e^{22}$  =  $E_{i,*}^{22}$   $\text{PSI}_S$ 

        Vector  $P_D$  =  $P^* \cdot h \cdot \exp(-20(1-a))$ 
            *  $(\Delta_{\min}^2 + \frac{5}{4}(\mathbf{E}_{i,*}^{11} * \mathbf{E}_{i,*}^{11} + \mathbf{E}_{i,*}^{22} * \mathbf{E}_{i,*}^{22}) + \frac{3}{2}\mathbf{E}_{i,*}^{11} * \mathbf{E}_{i,*}^{22} + \mathbf{E}_{i,*}^{12} * \mathbf{E}_{i,*}^{12})^{-\frac{1}{2}}$ 

         $S_{i,*}^{11}$  =  $(1 - \alpha^{-1})S_{i,*}^{11} + \alpha^{-1}T_i^{-1}(P_D * (\frac{5}{8}e^{11} + \frac{3}{8}e^{22}) - \frac{1}{2}P)$ 
         $S_{i,*}^{12}$  =  $(1 - \alpha^{-1})S_{i,*}^{12} + \alpha^{-1}T_i^{-1}(P_D * \frac{1}{4}e^{12})$ 
         $S_{i,*}^{22}$  =  $(1 - \alpha^{-1})S_{i,*}^{22} + \alpha^{-1}T_i^{-1}(P_D * (\frac{5}{8}e^{22} + \frac{3}{8}e^{11}) - \frac{1}{2}P)$ 
    }
}

```

Implementation

```

void StressUpdateHighOrder(Matrix&  $S^{11}$ , Matrix&  $S^{12}$ , Matrix&  $S^{22}$ ,
    const Matrix&  $E^{11}$ , const Matrix&  $E^{12}$ , const Matrix&  $E^{22}$ , // input fields  $N \times \text{DoF}$ 
    const Matrix&  $H$ , const Matrix&  $A$ ) {
    // update stress for every element
    for (int i = 0; i < N; ++i) {
        Vector h = max{0,  $H_{i,*}$   $\text{PSI}_A$ }
        Vector a = min{1, max{0,  $A_{i,*}$   $\text{PSI}_A$ }}
        Vector  $e^{11}$  =  $E_{i,*}^{11} \text{PSI}_S$ 
        Vector  $e^{12}$  =  $E_{i,*}^{12} \text{PSI}_S$ 
        Vector  $e^{22}$  =  $E_{i,*}^{22} \text{PSI}_S$ 

        Vector  $P_D$  =  $P^* \cdot h \cdot \exp(-20(1-a))$ 
            *  $(\Delta_{\min}^2 + \frac{5}{4}(\mathbf{E}_{i,*}^{11} * \mathbf{E}_{i,*}^{11} + \mathbf{E}_{i,*}^{22} * \mathbf{E}_{i,*}^{22}) + \frac{3}{2}\mathbf{E}_{i,*}^{11} * \mathbf{E}_{i,*}^{22} + \mathbf{E}_{i,*}^{12} * \mathbf{E}_{i,*}^{12})^{-\frac{1}{2}}$ 

         $S_{i,*}^{11}$  =  $(1 - \alpha^{-1})S_{i,*}^{11} + \alpha^{-1}T_i^{-1}(P_D * (\frac{5}{8}e^{11} + \frac{3}{8}e^{22}) - \frac{1}{2}P)$ 
         $S_{i,*}^{12}$  =  $(1 - \alpha^{-1})S_{i,*}^{12} + \alpha^{-1}T_i^{-1}(P_D * \frac{1}{4}e^{12})$ 
         $S_{i,*}^{22}$  =  $(1 - \alpha^{-1})S_{i,*}^{22} + \alpha^{-1}T_i^{-1}(P_D * (\frac{5}{8}e^{22} + \frac{3}{8}e^{11}) - \frac{1}{2}P)$ 
    }
}

```

Implementation

```

void StressUpdateHighOrder(Matrix&  $S^{11}$ , Matrix&  $S^{12}$ , Matrix&  $S^{22}$ ,
    const Matrix&  $E^{11}$ , const Matrix&  $E^{12}$ , const Matrix&  $E^{22}$ ,
    const Matrix&  $H$ , const Matrix&  $A$ ) {
    // update stress for every element
    for (int i=0; i < N; ++i) {
        Vector h = max{0,  $H_{i,*} \text{PSI}_A$ }
        Vector a = min{1, max{0,  $A_{i,*} \text{PSI}_A$ }}
        Vector  $e^{11}$  =  $E_{i,*}^{11} \text{PSI}_S$  // eval fields at gauss points
        Vector  $e^{12}$  =  $E_{i,*}^{12} \text{PSI}_S$ 
        Vector  $e^{22}$  =  $E_{i,*}^{22} \text{PSI}_S$ 

        Vector  $P_D$  =  $P^* \cdot h \cdot \exp(-20(1-a))$ 
            *  $(\Delta_{\min}^2 + \frac{5}{4}(\mathbf{E}_{i,*}^{11} * \mathbf{E}_{i,*}^{11} + \mathbf{E}_{i,*}^{22} * \mathbf{E}_{i,*}^{22}) + \frac{3}{2}\mathbf{E}_{i,*}^{11} * \mathbf{E}_{i,*}^{22} + \mathbf{E}_{i,*}^{12} * \mathbf{E}_{i,*}^{12})^{-\frac{1}{2}}$ 

         $S_{i,*}^{11}$  =  $(1 - \alpha^{-1})S_{i,*}^{11} + \alpha^{-1}T_i^{-1}(P_D * (\frac{5}{8}e^{11} + \frac{3}{8}e^{22}) - \frac{1}{2}P)$ 
         $S_{i,*}^{12}$  =  $(1 - \alpha^{-1})S_{i,*}^{12} + \alpha^{-1}T_i^{-1}(P_D * \frac{1}{4}e^{12})$ 
         $S_{i,*}^{22}$  =  $(1 - \alpha^{-1})S_{i,*}^{22} + \alpha^{-1}T_i^{-1}(P_D * (\frac{5}{8}e^{22} + \frac{3}{8}e^{11}) - \frac{1}{2}P)$ 
    }
}

```

Implementation

```

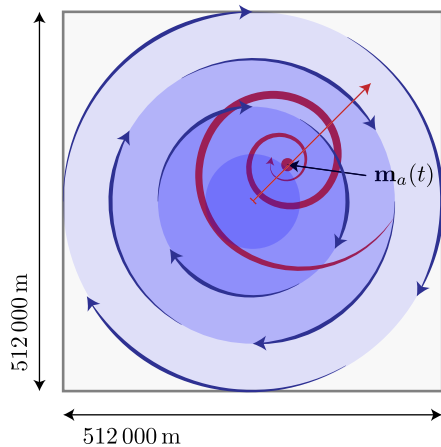
void StressUpdateHighOrder(Matrix&  $S^{11}$ , Matrix&  $S^{12}$ , Matrix&  $S^{22}$ ,
    const Matrix&  $E^{11}$ , const Matrix&  $E^{12}$ , const Matrix&  $E^{22}$ ,
    const Matrix&  $H$ , const Matrix&  $A$ ) {
    // update stress for every element
    for (int i = 0; i < N; ++i) {
        Vector h = max{0,  $H_{i,*} \text{PSI}_A$ }
        Vector a = min{1, max{0,  $A_{i,*} \text{PSI}_A$ }}
        Vector  $e^{11}$  =  $E_{i,*}^{11} \text{PSI}_S$  // eval fields at gauss points
        Vector  $e^{12}$  =  $E_{i,*}^{12} \text{PSI}_S$ 
        Vector  $e^{22}$  =  $E_{i,*}^{22} \text{PSI}_S$ 

        Vector  $P_D$  =  $P^* \cdot h \cdot \exp(-20(1-a))$ 
            *  $(\Delta_{\min}^2 + \frac{5}{4}(\mathbf{E}_{i,*}^{11} * \mathbf{E}_{i,*}^{11} + \mathbf{E}_{i,*}^{22} * \mathbf{E}_{i,*}^{22}) + \frac{3}{2}\mathbf{E}_{i,*}^{11} * \mathbf{E}_{i,*}^{22} + \mathbf{E}_{i,*}^{12} * \mathbf{E}_{i,*}^{12})^{-\frac{1}{2}}$ 

         $S_{i,*}^{11}$  =  $(1 - \alpha^{-1})S_{i,*}^{11} + \alpha^{-1}\mathbf{T}_i^{-1}(P_D * (\frac{5}{8}e^{11} + \frac{3}{8}e^{22}) - \frac{1}{2}P)$ 
         $S_{i,*}^{12}$  =  $(1 - \alpha^{-1})S_{i,*}^{12} + \alpha^{-1}\mathbf{T}_i^{-1}(P_D * \frac{1}{4}e^{12})$  // mapping of parametric mesh
         $S_{i,*}^{22}$  =  $(1 - \alpha^{-1})S_{i,*}^{22} + \alpha^{-1}\mathbf{T}_i^{-1}(P_D * (\frac{5}{8}e^{22} + \frac{3}{8}e^{11}) - \frac{1}{2}P)$ 
    }
}

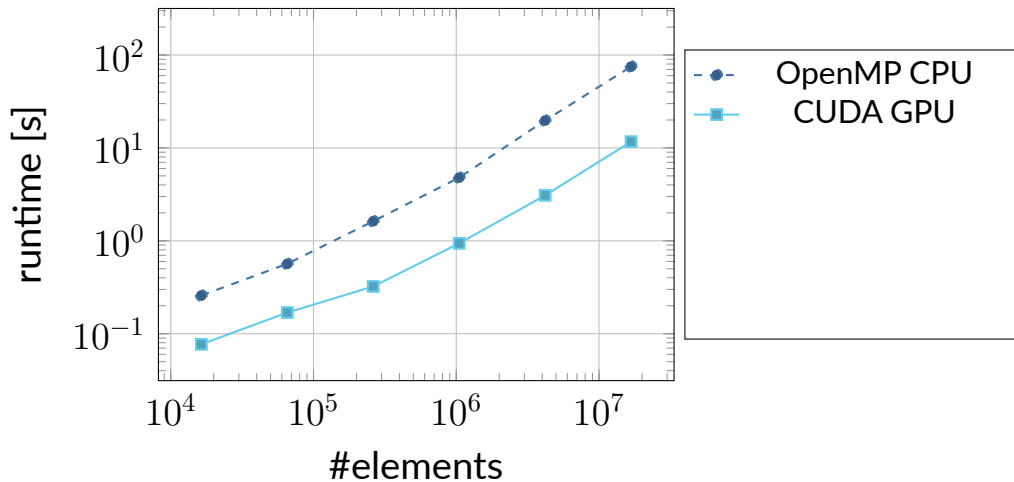
```

Benchmark - Cyclone in a Box⁵

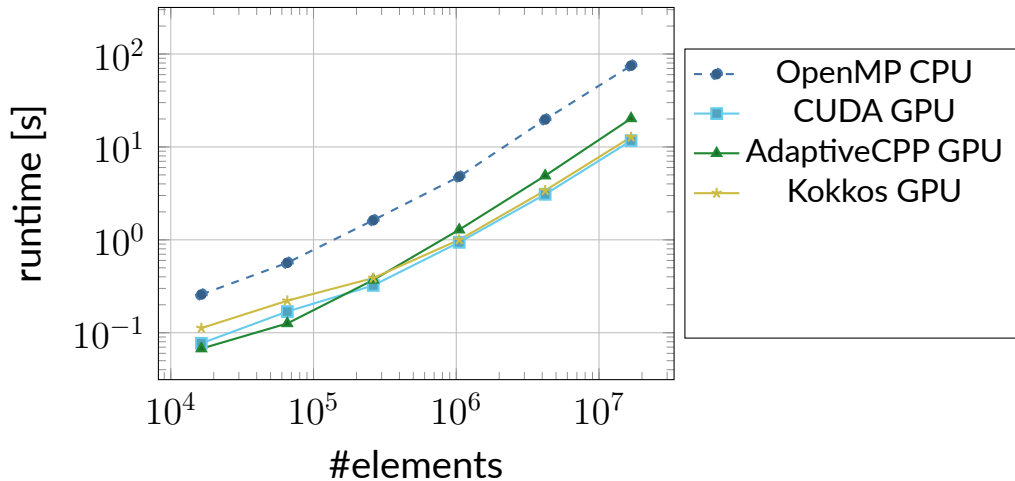


⁵Mehlmann et al., "Simulating Linear Kinematic Features in Viscous-Plastic Sea Ice Models on Quadrilateral and Triangular Grids With Different Variable Staggering", 2021.

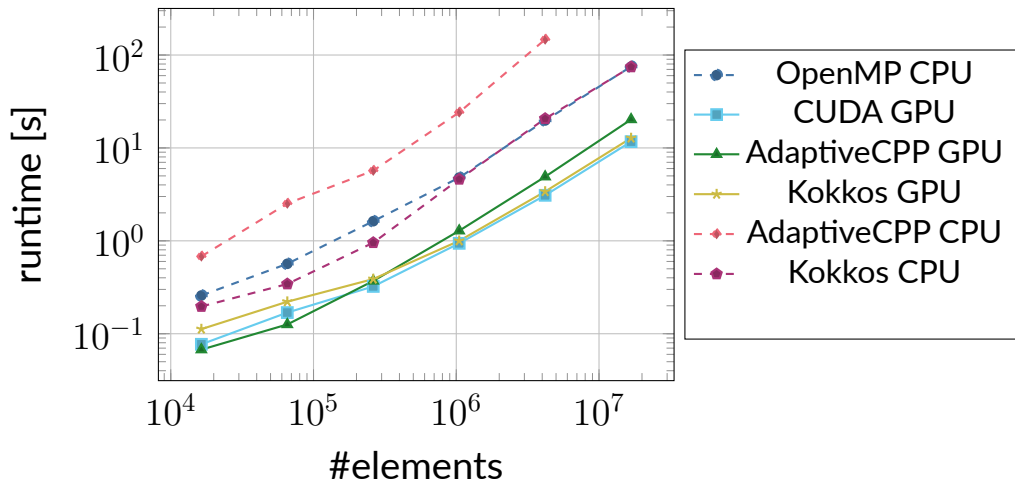
Comparison F64 - A100 / 2× EPYC 7402, 48 cores



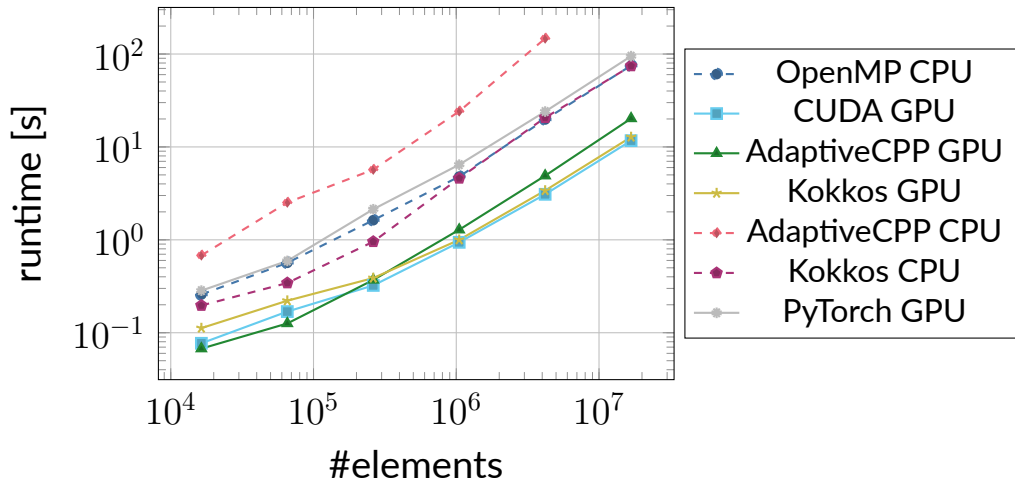
Comparison F64 - A100 / 2× EPYC 7402, 48 cores



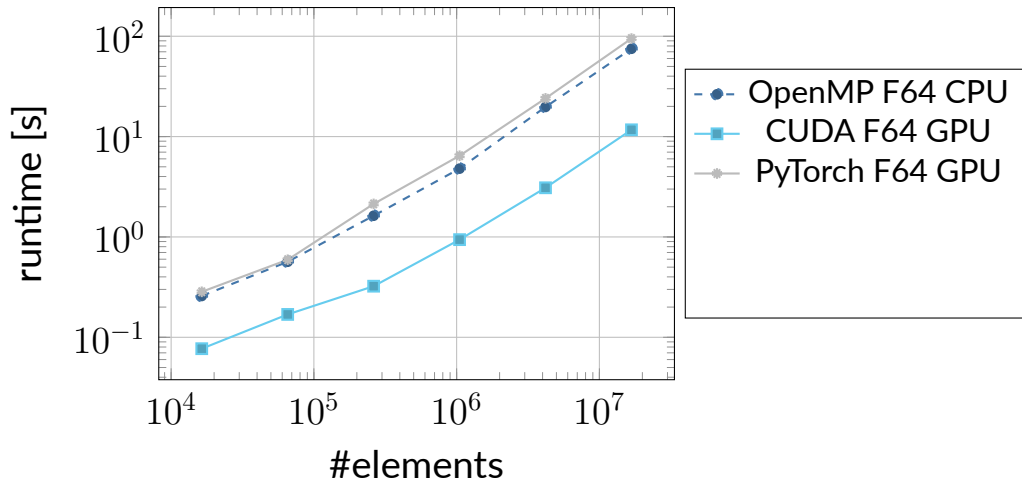
Comparison F64 - A100 / 2× EPYC 7402, 48 cores



Comparison F64 - A100 / 2× EPYC 7402, 48 cores

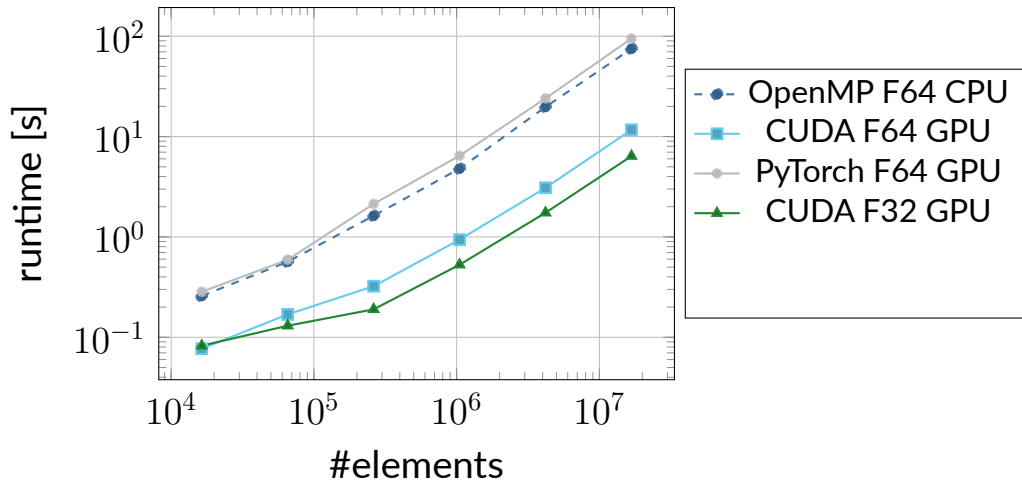


Comparison - Mixed Precision⁶



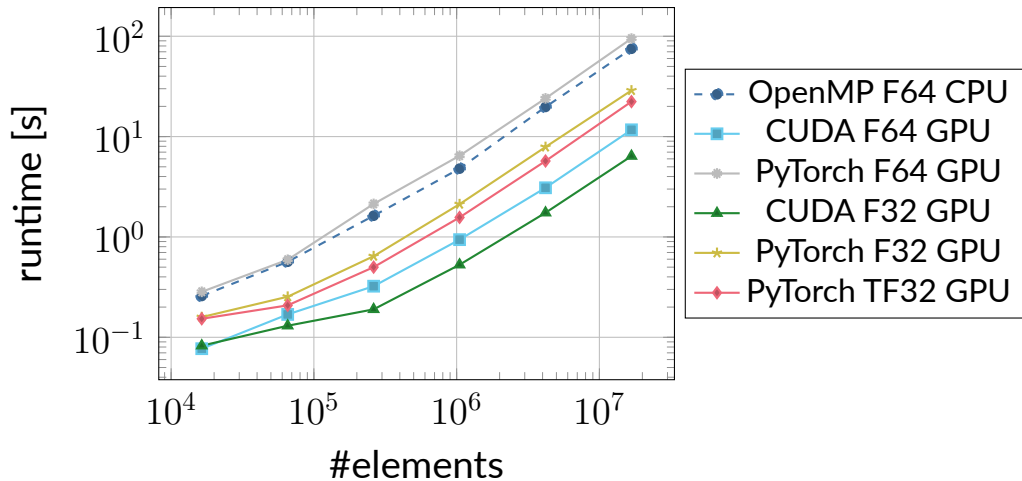
⁶Accuracy investigated in: Jendersie, Lessig, and Richter, "A GPU-parallelization of the neXtSIM-DG dynamical core (v0.3.1)", 2024

Comparison - Mixed Precision⁶



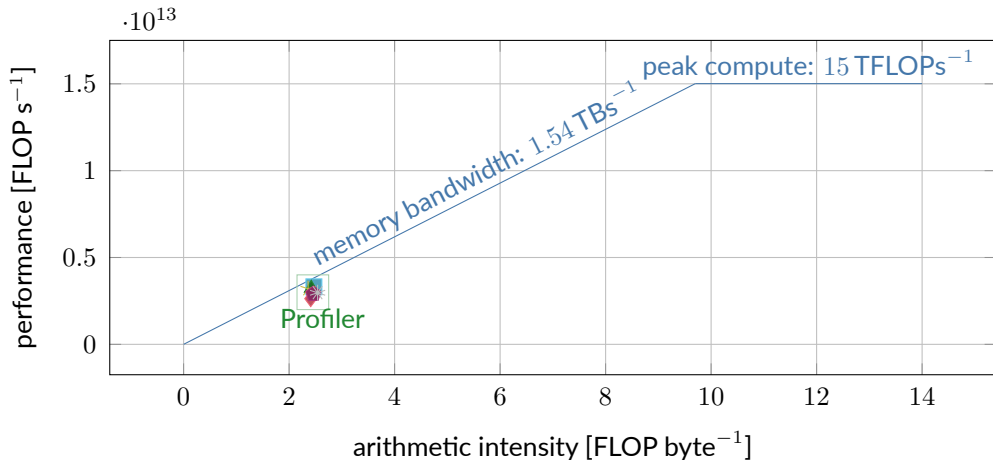
⁶Accuracy investigated in: Jendersie, Lessig, and Richter, "A GPU-parallelization of the neXtSIM-DG dynamical core (v0.3.1)", 2024

Comparison - Mixed Precision⁶

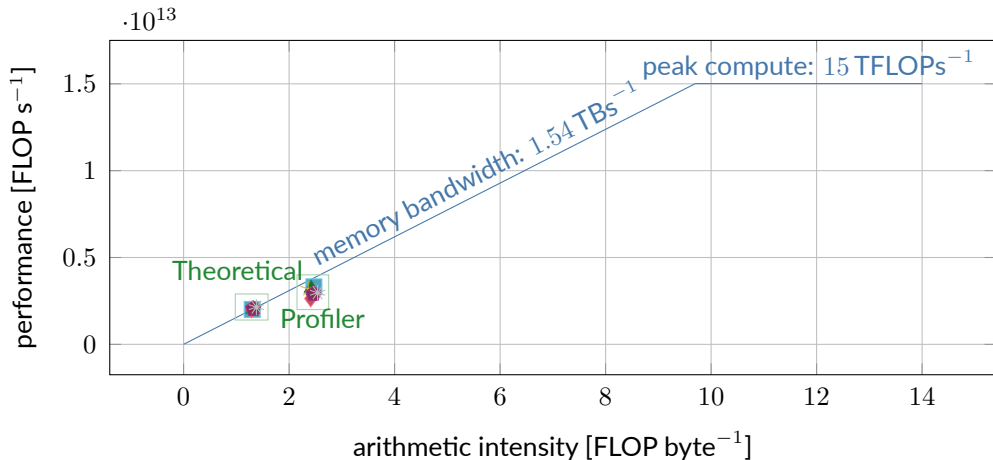


⁶Accuracy investigated in: Jendersie, Lessig, and Richter, "A GPU-parallelization of the neXtSIM-DG dynamical core (v0.3.1)", 2024

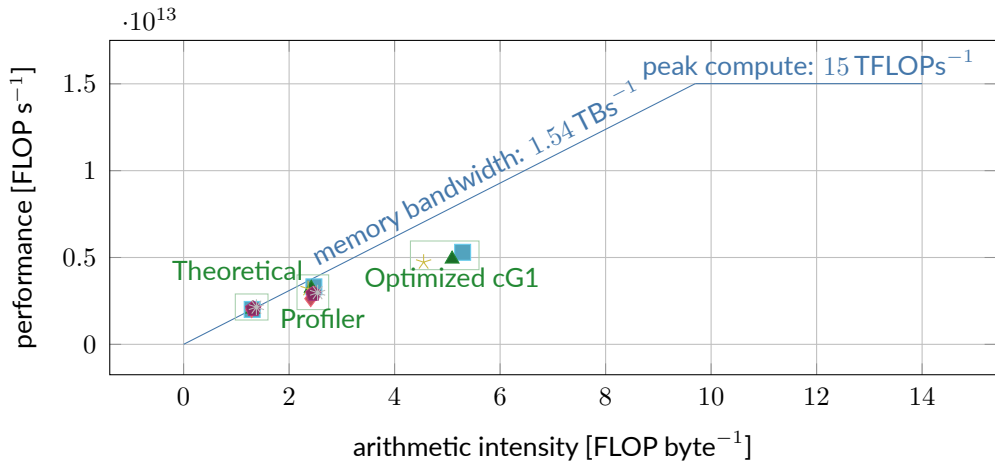
Roofline Analysis - A100



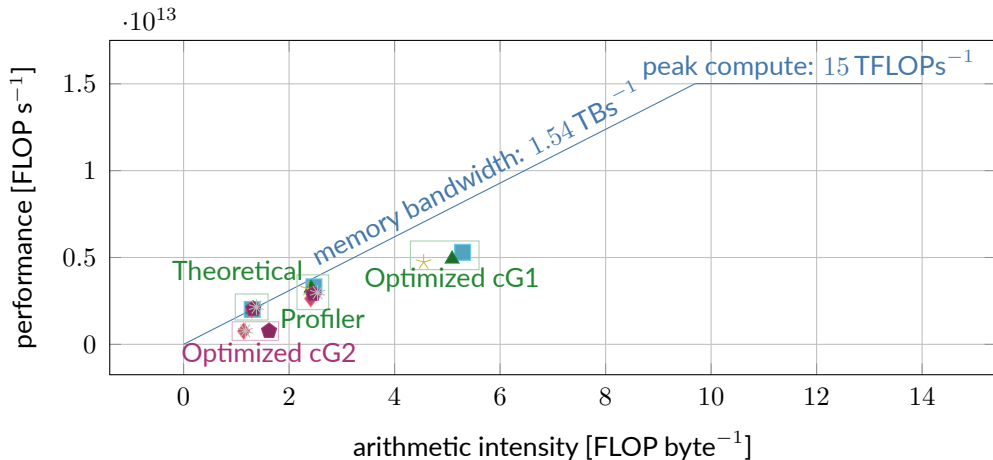
Roofline Analysis - A100



Roofline Analysis - A100



Roofline Analysis - A100



Choice for Full Implementation

Performance

$\text{CUDA} \leq \text{Kokkos} < \text{AdaptiveCPP(Sycl)} < \text{PyTorch}$

Choice for Full Implementation

Performance

$\text{CUDA} \leq \mathbf{\text{Kokkos}} < \text{AdaptiveCPP(Sycl)} < \text{PyTorch}$

Ease of Development

$\text{PyTorch} < \text{Sycl} < \mathbf{\text{Kokkos}} < \text{CUDA}$

Choice for Full Implementation

Performance

$\text{CUDA} \leq \mathbf{\text{Kokkos}} < \text{AdaptiveCPP(Sycl)} < \text{PyTorch}$

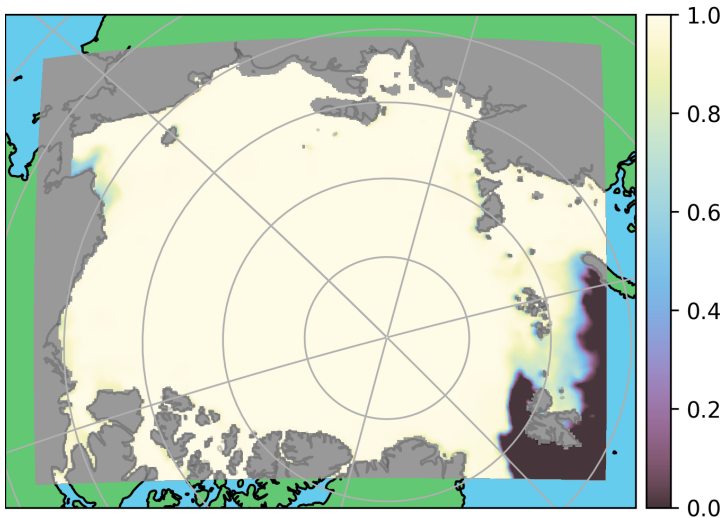
Ease of Development

$\text{PyTorch} < \text{Sycl} < \mathbf{\text{Kokkos}} < \text{CUDA}$

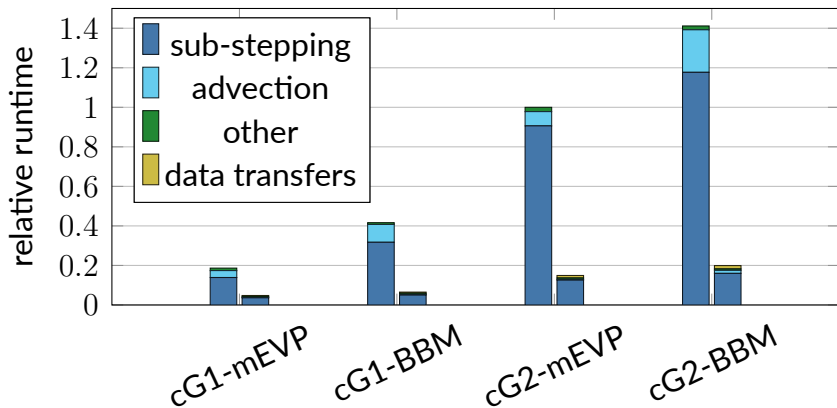
Deployment effort

$\text{CUDA} \leq \mathbf{\text{Kokkos}} < \text{PyTorch} < \text{AdaptiveCPP(Sycl)}$

Test case - Arctic



Performance - Arctic 3.125 km



2 x EPYC 7773X, 122 cores (left) - H100 PCIe GPU (right)

Conclusion

- Kokkos can be decent alternative to CUDA
 - Higher order CG/DG discretization works well on GPU
 - Effective use of tensor cores is difficult
-



Robert Jendersie, Christian Lessig, and Thomas Richter. "A GPU-parallelization of the neXtSIM-DG dynamical core (v0.3.1)". In: (Sept. 2024). DOI: 10.5194/egusphere-2024-2539

This project is supported by Schmidt Sciences, LLC.

Conclusion

- Kokkos can be decent alternative to CUDA
 - Higher order CG/DG discretization works well on GPU
 - Effective use of tensor cores is difficult
-



Robert Jendersie, Christian Lessig, and Thomas Richter. “A GPU-parallelization of the neXtSIM-DG dynamical core (v0.3.1)”. In: (Sept. 2024). DOI: 10.5194/egusphere-2024-2539

This project is supported by Schmidt Sciences, LLC.

Conclusion

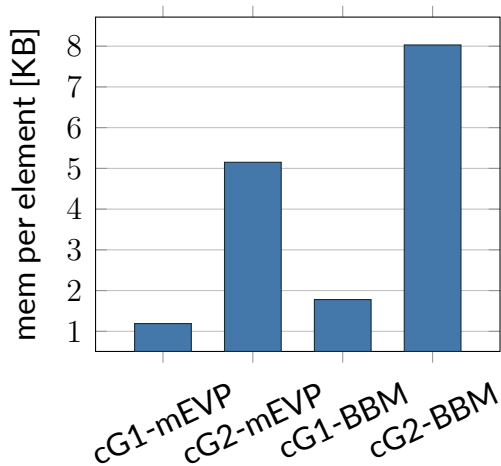
- Kokkos can be decent alternative to CUDA
 - Higher order CG/DG discretization works well on GPU
 - Effective use of tensor cores is difficult
-



Robert Jendersie, Christian Lessig, and Thomas Richter. "A GPU-parallelization of the neXtSIM-DG dynamical core (v0.3.1)". In: (Sept. 2024). DOI: [10.5194/egusphere-2024-2539](https://doi.org/10.5194/egusphere-2024-2539)

This project is supported by Schmidt Sciences, LLC.

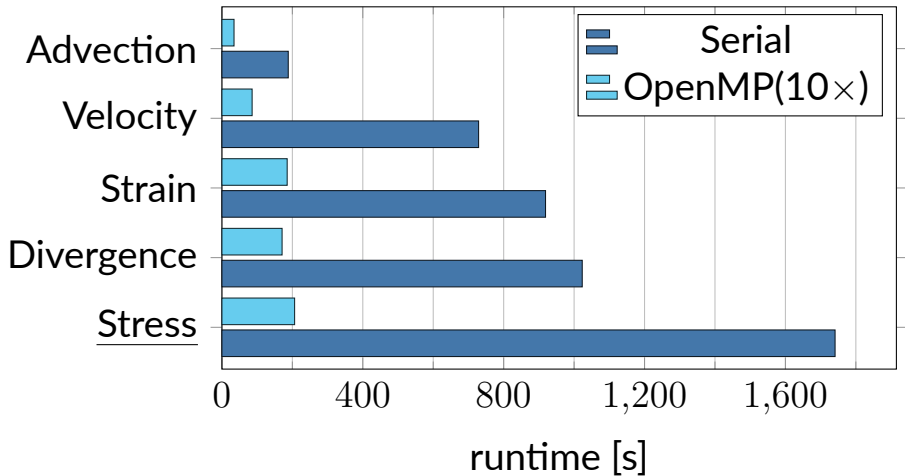
Device Memory Usage



max #elem with 40GB VRAM:

- cG1-mEVP 3.3×10^7
- cG2-BBM 4.9×10^6

Parallelization



Frameworks

- Directive based: OpenMP, OpenACC
- Low level interface: CUDA
- Heterogeneous compute: Sycl, Kokkos
- Machine learning framework: PyTorch

Implementation OpenMP

```
void StressUpdateHighOrder(Matrix&  $S^{11}$ , Matrix&  $S^{12}$ , Matrix&  $S^{22}$ ,  
    const Matrix&  $E^{11}$ , const Matrix&  $E^{12}$ , const Matrix&  $E^{22}$ ,  
    const Matrix&  $H$ , const Matrix&  $A$ ) {  
    // update stress in parallel  
    #pragma omp parallel for  
    for (int  $i=0$ ;  $i < N$ ; ++ $i$ ) {  
        computeStress( $i$ ,  $S^{11}$ ,  $S^{12}$ ,  $S^{22}$ ,  $E^{11}$ ,  $E^{12}$ ,  $E^{22}$ ,  $H$ ,  $A$ ,  $\alpha$ );  
    }  
}
```

Implementation OpenMP

```
void StressUpdateHighOrder(Matrix&  $S^{11}$ , Matrix&  $S^{12}$ , Matrix&  $S^{22}$ ,  
    const Matrix&  $E^{11}$ , const Matrix&  $E^{12}$ , const Matrix&  $E^{22}$ ,  
    const Matrix&  $H$ , const Matrix&  $A$ ) {  
    // update stress in parallel  
    #pragma omp target parallel for  
    for (int  $i=0$ ;  $i < N$ ; ++ $i$ ) {  
        computeStress( $i$ ,  $S^{11}$ ,  $S^{12}$ ,  $S^{22}$ ,  $E^{11}$ ,  $E^{12}$ ,  $E^{22}$ ,  $H$ ,  $A$ ,  $\alpha$ );  
    }  
}
```

Implementation CUDA

```
void StressUpdateHighOrder(Matrix&  $S^{11}$ , Matrix&  $S^{12}$ , Matrix&  $S^{22}$ ,  
    const Matrix&  $E^{11}$ , const Matrix&  $E^{12}$ , const Matrix&  $E^{22}$ ,  
    const Matrix&  $H$ , const Matrix&  $A$ ) {  
    // allocate memory  
    if (init) {  
        double*  $\hat{S}^{11}$  = cudaAlloc( $N \cdot n_s \cdot \text{sizeof}(\text{double})$ )  
        ...  
    }  
    // upload data  
    cudaMemcpy( $S^{11}$ ,  $\hat{S}^{11}$ ,  $N \cdot n_s \cdot \text{sizeof}(\text{double})$ )  
    ...  
    // invoke kernel  
    computeStress<<< $N_{\text{grid}}, N_{\text{block}}$ >>>( $\hat{S}^{11}$ ,  $\hat{S}^{12}$ ,  $\hat{S}^{22}$ ,  $\hat{E}^{11}$ ,  $\hat{E}^{12}$ ,  $\hat{E}^{22}$ ,  $\hat{H}$ ,  $\hat{A}$ ,  $\alpha$ );  
    // fetch data  
    cudaMemcpy( $\hat{S}^{11}$ ,  $S^{11}$ ,  $N \cdot n_s \cdot \text{sizeof}(\text{double})$ )  
    ...  
}
```

Implementation CUDA

```
void StressUpdateHighOrder(Matrix&  $S^{11}$ , Matrix&  $S^{12}$ , Matrix&  $S^{22}$ ,  
    const Matrix&  $E^{11}$ , const Matrix&  $E^{12}$ , const Matrix&  $E^{22}$ ,  
    const Matrix&  $H$ , const Matrix&  $A$ ) {  
    // allocate memory  
    if (init) {  
        double*  $\hat{S}^{11}$  = cudaAlloc( $N \cdot n_s \cdot \text{sizeof}(\text{double})$ )  
        ...  
    }  
    // upload data  
    cudaMemcpy( $S^{11}$ ,  $\hat{S}^{11}$ ,  $N \cdot n_s \cdot \text{sizeof}(\text{double})$ )  
    ...  
    // invoke kernel  
    computeStress<<< $N_{\text{grid}}$ ,  $N_{\text{block}}$ >>>( $\hat{S}^{11}$ ,  $\hat{S}^{12}$ ,  $\hat{S}^{22}$ ,  $\hat{E}^{11}$ ,  $\hat{E}^{12}$ ,  $\hat{E}^{22}$ ,  $\hat{H}$ ,  $\hat{A}$ ,  $\alpha$ );  
    // fetch data  
    cudaMemcpy( $\hat{S}^{11}$ ,  $S^{11}$ ,  $N \cdot n_s \cdot \text{sizeof}(\text{double})$ )  
    ...  
}
```

Implementation CUDA

```
void StressUpdateHighOrder(Matrix&  $S^{11}$ , Matrix&  $S^{12}$ , Matrix&  $S^{22}$ ,  
    const Matrix&  $E^{11}$ , const Matrix&  $E^{12}$ , const Matrix&  $E^{22}$ ,  
    const Matrix&  $H$ , const Matrix&  $A$ ) {  
    // allocate memory  
    if (init) {  
        double*  $\hat{S}^{11}$  = cudaAlloc( $N \cdot n_s \cdot \text{sizeof}(\text{double})$ )  
        ...  
    }  
    // upload data  
    cudaMemcpy( $S^{11}$ ,  $\hat{S}^{11}$ ,  $N \cdot n_s \cdot \text{sizeof}(\text{double})$ )  
    ...  
    // invoke kernel  
    computeStress<<< $N_{\text{grid}}$ ,  $N_{\text{block}}$ >>>( $\hat{S}^{11}$ ,  $\hat{S}^{12}$ ,  $\hat{S}^{22}$ ,  $\hat{E}^{11}$ ,  $\hat{E}^{12}$ ,  $\hat{E}^{22}$ ,  $\hat{H}$ ,  $\hat{A}$ ,  $\alpha$ );  
    // fetch data  
    cudaMemcpy( $\hat{S}^{11}$ ,  $S^{11}$ ,  $N \cdot n_s \cdot \text{sizeof}(\text{double})$ )  
    ...  
}
```

Implementation Kokkos

```
void StressUpdateHighOrder(Matrix&  $S^{11}$ , Matrix&  $S^{12}$ , Matrix&  $S^{22}$ ,  
    const Matrix&  $E^{11}$ , const Matrix&  $E^{12}$ , const Matrix&  $E^{22}$ ,  
    const Matrix&  $H$ , const Matrix&  $A$ ) {  
    // allocate memory  
    if (init) {  
        auto  $\hat{S}^{11}$  = View( $N, N_s$ )  
        ...  
    }  
    // upload data  
    kokkosCopy( $S^{11}$ ,  $\hat{S}^{11}$ )  
    ...  
    // compute in parallel  
    parallel_for( $N$ , KOKKOS_CLASS_LAMBDA(int i){  
        stressUpdate(i,  $\hat{S}^{11}$ ,  $\hat{S}^{12}$ ,  $\hat{S}^{22}$ ,  $\hat{E}^{11}$ ,  $\hat{E}^{12}$ ,  $\hat{E}^{22}$ ,  $\hat{H}$ ,  $\hat{A}$ ,  $\alpha$ )  
    })  
    // fetch data  
    kokkosCopy( $\hat{S}^{11}$ ,  $S^{11}$ )  
    ...  
}
```

Implementation Kokkos

```
void StressUpdateHighOrder(Matrix&  $S^{11}$ , Matrix&  $S^{12}$ , Matrix&  $S^{22}$ ,  
    const Matrix&  $E^{11}$ , const Matrix&  $E^{12}$ , const Matrix&  $E^{22}$ ,  
    const Matrix&  $H$ , const Matrix&  $A$ ) {  
    // allocate memory  
    if (init) {  
        auto  $\hat{S}^{11}$  = View( $N, N_s$ )  
        ...  
    }  
    // upload data  
    kokkosCopy( $S^{11}$ ,  $\hat{S}^{11}$ )  
    ...  
    // compute in parallel  
    parallel_for( $N$ , KOKKOS_CLASS_LAMBDA(int i){  
        stressUpdate(i,  $\hat{S}^{11}$ ,  $\hat{S}^{12}$ ,  $\hat{S}^{22}$ ,  $\hat{E}^{11}$ ,  $\hat{E}^{12}$ ,  $\hat{E}^{22}$ ,  $\hat{H}$ ,  $\hat{A}$ ,  $\alpha$ )  
    })  
    // fetch data  
    kokkosCopy( $\hat{S}^{11}$ ,  $S^{11}$ )  
    ...  
}
```

Implementation Kokkos

```
void StressUpdateHighOrder(Matrix&  $S^{11}$ , Matrix&  $S^{12}$ , Matrix&  $S^{22}$ ,  
    const Matrix&  $E^{11}$ , const Matrix&  $E^{12}$ , const Matrix&  $E^{22}$ ,  
    const Matrix&  $H$ , const Matrix&  $A$ ) {  
    // allocate memory  
    if (init) {  
        auto  $\hat{S}^{11}$  = View( $N, N_s$ )  
        ...  
    }  
    // upload data  
    kokkosCopy( $S^{11}$ ,  $\hat{S}^{11}$ )  
    ...  
    // compute in parallel  
    parallel_for( $N$ , KOKKOS_CLASS_LAMBDA(int i){  
        stressUpdate(i,  $\hat{S}^{11}$ ,  $\hat{S}^{12}$ ,  $\hat{S}^{22}$ ,  $\hat{E}^{11}$ ,  $\hat{E}^{12}$ ,  $\hat{E}^{22}$ ,  $\hat{H}$ ,  $\hat{A}$ ,  $\alpha$ )  
    })  
    // fetch data  
    kokkosCopy( $\hat{S}^{11}$ ,  $S^{11}$ )  
    ...  
}
```

Implementation Sycl

```
void StressUpdateHighOrder(Matrix&  $S^{11}$ , Matrix&  $S^{12}$ , Matrix&  $S^{22}$ ,  
    const Matrix&  $E^{11}$ , const Matrix&  $E^{12}$ , const Matrix&  $E^{22}$ ,  
    const Matrix&  $H$ , const Matrix&  $A$ ) {  
    // define buffers  
    if (init) {  
        auto  $\bar{S}^{11}$  = buffer( $S^{11}$ ,  $N$ ,  $N_s$ )  
        ...  
    }  
    queue.submit([&]() {  
        // declare required data  
        auto  $\hat{S}^{11}$  = accessor( $\bar{S}^{11}$ , read_write)  
        ...  
        // compute in parallel  
        parallel_for( $N$ , [=](int i) {  
            stressUpdate(i,  $\hat{S}^{11}$ ,  $\hat{S}^{12}$ ,  $\hat{S}^{22}$ ,  $\hat{E}^{11}$ ,  $\hat{E}^{12}$ ,  $\hat{E}^{22}$ ,  $\hat{H}$ ,  $\hat{A}$ ,  $\alpha$ )  
        })  
    }  
}
```

Implementation Sycl

```
void StressUpdateHighOrder(Matrix&  $S^{11}$ , Matrix&  $S^{12}$ , Matrix&  $S^{22}$ ,  
    const Matrix&  $E^{11}$ , const Matrix&  $E^{12}$ , const Matrix&  $E^{22}$ ,  
    const Matrix&  $H$ , const Matrix&  $A$ ) {  
    // define buffers  
    if (init) {  
        auto  $\bar{S}^{11}$  = buffer( $S^{11}$ ,  $N$ ,  $N_s$ )  
        ...  
    }  
    queue.submit([&]() {  
        // declare required data  
        auto  $\hat{S}^{11}$  = accessor( $\bar{S}^{11}$ , read_write)  
        ...  
        // compute in parallel  
        parallel_for( $N$ , [=](int i) {  
            stressUpdate(i,  $\hat{S}^{11}$ ,  $\hat{S}^{12}$ ,  $\hat{S}^{22}$ ,  $\hat{E}^{11}$ ,  $\hat{E}^{12}$ ,  $\hat{E}^{22}$ ,  $\hat{H}$ ,  $\hat{A}$ ,  $\alpha$ )  
        })  
    }  
}
```

Implementation Sycl

```
void StressUpdateHighOrder(Matrix&  $S^{11}$ , Matrix&  $S^{12}$ , Matrix&  $S^{22}$ ,  
    const Matrix&  $E^{11}$ , const Matrix&  $E^{12}$ , const Matrix&  $E^{22}$ ,  
    const Matrix&  $H$ , const Matrix&  $A$ ) {  
    // define buffers  
    if (init) {  
        auto  $\bar{S}^{11}$  = buffer( $S^{11}$ ,  $N$ ,  $N_s$ )  
        ...  
    }  
    queue.submit([&]() {  
        // declare required data  
        auto  $\hat{S}^{11}$  = accessor( $\bar{S}^{11}$ , read_write)  
        ...  
        // compute in parallel  
        parallel_for( $N$ , [=](int i) {  
            stressUpdate(i,  $\hat{S}^{11}$ ,  $\hat{S}^{12}$ ,  $\hat{S}^{22}$ ,  $\hat{E}^{11}$ ,  $\hat{E}^{12}$ ,  $\hat{E}^{22}$ ,  $\hat{H}$ ,  $\hat{A}$ ,  $\alpha$ )  
        })  
    }  
}
```

Implementation Sycl

```
void StressUpdateHighOrder(Matrix&  $S^{11}$ , Matrix&  $S^{12}$ , Matrix&  $S^{22}$ ,  
    const Matrix&  $E^{11}$ , const Matrix&  $E^{12}$ , const Matrix&  $E^{22}$ ,  
    const Matrix&  $H$ , const Matrix&  $A$ ) {  
    // define buffers  
    if (init) {  
        auto  $\bar{S}^{11}$  = buffer( $S^{11}$ ,  $N$ ,  $N_s$ )  
        ...  
    }  
    queue.submit([&]() {  
        // declare required data  
        auto  $\hat{S}^{11}$  = accessor( $\bar{S}^{11}$ , read_write)  
        ...  
        // compute in parallel  
        parallel_for( $N$ , [=](int i) {  
            stressUpdate(i,  $\hat{S}^{11}$ ,  $\hat{S}^{12}$ ,  $\hat{S}^{22}$ ,  $\hat{E}^{11}$ ,  $\hat{E}^{12}$ ,  $\hat{E}^{22}$ ,  $\hat{H}$ ,  $\hat{A}$ ,  $\alpha$ )  
        })  
    }  
}
```

Implementation PyTorch

```
void StressUpdateHighOrder(Matrix&  $S^{11}$ , Matrix&  $S^{12}$ , Matrix&  $S^{22}$ ,  
    const Matrix&  $E^{11}$ , const Matrix&  $E^{12}$ , const Matrix&  $E^{22}$ ,  
    const Matrix&  $H$ , const Matrix&  $A$ ) {  
    // allocate memory  
    if (init) {  
        auto  $\hat{S}^{11}$  = Tensor( $N$ ,  $N_s$ , "cuda")  
        ...  
    }  
    // upload data  
     $\hat{S}^{11}$ .copy_( $S^{11}$ )  
    ...  
    // compute on device  
    Tensor  $h$  = max{0,  $\hat{H}^{PSI}$ }  
    Tensor  $a$  = min{1, max{0,  $\hat{A}^{PSI}$ }}  
    ...  
    // fetch data  
     $S^{11}$ .copy_( $\hat{S}^{11}$ )  
    ...  
}
```

Implementation PyTorch

```
void StressUpdateHighOrder(Matrix&  $S^{11}$ , Matrix&  $S^{12}$ , Matrix&  $S^{22}$ ,  
    const Matrix&  $E^{11}$ , const Matrix&  $E^{12}$ , const Matrix&  $E^{22}$ ,  
    const Matrix&  $H$ , const Matrix&  $A$ ) {  
    // allocate memory  
    if (init) {  
        auto  $\hat{S}^{11}$  = Tensor(N, N_s, "cuda")  
        ...  
    }  
    // upload data  
     $\hat{S}^{11}$ .copy_( $S^{11}$ )  
    ...  
    // compute on device  
    Tensor h = max{0,  $\hat{H}^{PSI}$ }  
    Tensor a = min{1, max{0,  $\hat{A}^{PSI}$ }}  
    ...  
    // fetch data  
     $S^{11}$ .copy_( $\hat{S}^{11}$ )  
    ...  
}
```

Implementation PyTorch

```
void StressUpdateHighOrder(Matrix&  $S^{11}$ , Matrix&  $S^{12}$ , Matrix&  $S^{22}$ ,  
    const Matrix&  $E^{11}$ , const Matrix&  $E^{12}$ , const Matrix&  $E^{22}$ ,  
    const Matrix&  $H$ , const Matrix&  $A$ ) {  
    // allocate memory  
    if (init) {  
        auto  $\hat{S}^{11}$  = Tensor( $N$ ,  $N_s$ , "cuda")  
        ...  
    }  
    // upload data  
     $\hat{S}^{11}$ .copy_( $S^{11}$ )  
    ...  
    // compute on device  
    Tensor  $h$  = max{0,  $\hat{H}^{PSI}$ }  
    Tensor  $a$  = min{1, max{0,  $\hat{A}^{PSI}$ }}  
    ...  
    // fetch data  
     $S^{11}$ .copy_( $\hat{S}^{11}$ )  
    ...  
}
```

Development and Deployment

Ease of Development

PyTorch < Sycl < Kokkos < CUDA

Deployment effort

CUDA $\leq$ Kokkos < PyTorch < Sycl(AdaptiveCPP)

Development and Deployment

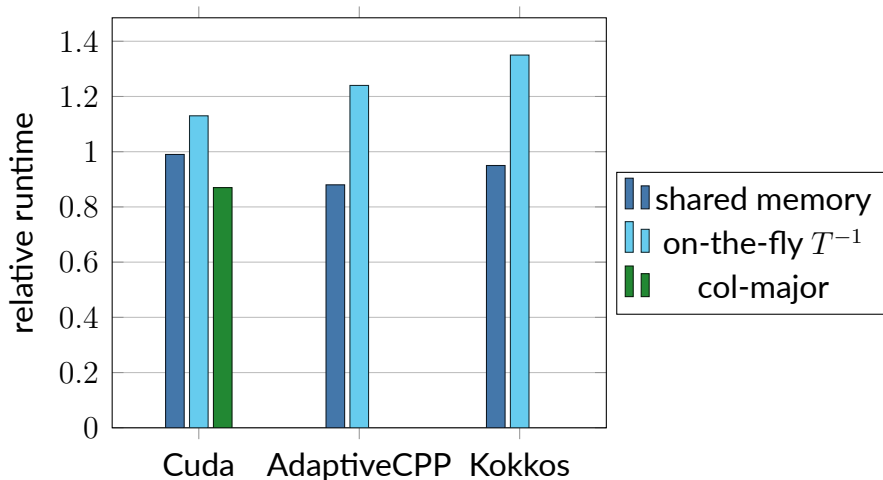
Ease of Development

PyTorch < Sycl < Kokkos < CUDA

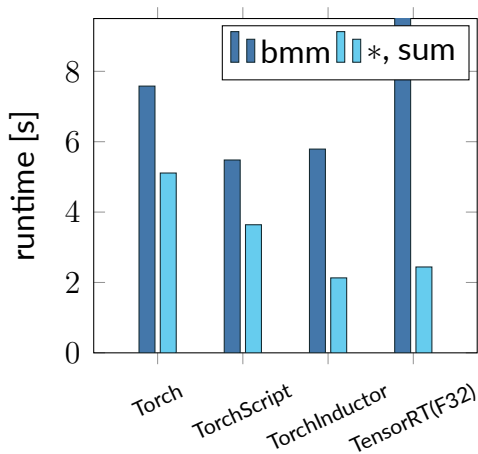
Deployment effort

CUDA $\leq$ Kokkos < PyTorch < Sycl(AdaptiveCPP)

Optimization

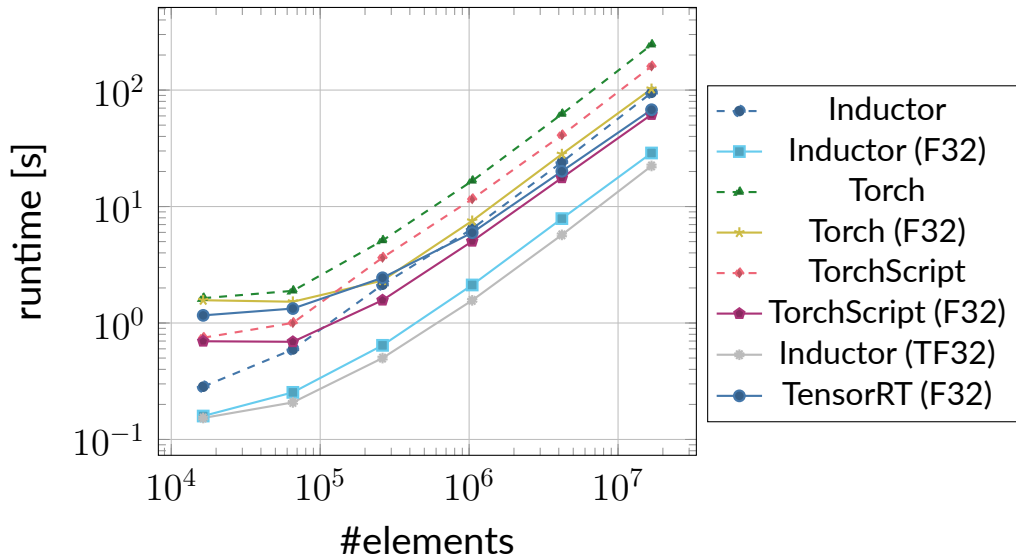


Optimization - PyTorch



```
# m1 ∈ ℝ262144×4×3 , m2 ∈ ℝ262144×3  
# bmm  
def bmm(m1, m2):  
    return torch.bmm(m1, m2.unsqueeze_(2))  
    .squeeze_(2)  
  
# *, sum  
def mulsum(m1, m2):  
    return torch.sum(m1 * m2.unsqueeze_(1),  
                     dim=2)
```

Performance - PyTorch



Performance - Higher Order

