



The University of Texas at Austin



HPSF
CONFERENCE 2025

May 7th, 2025

Performance Portable Fast Ewald Summation

Kokkos User Group | Chicago, IL

Gabriel Kosmacher, Ziyu Du, Joar Bagge, George Biros



Predictive
Engineering &
Computational Science



- Consider periodic N -body sums of the form

$$u(\mathbf{x}_i) = \sum_{j=1}^{N_s} \sum_{\mathbf{p} \in \mathcal{P}} G(\mathbf{x}_i - \mathbf{y}_j + \mathbf{p}) f(\mathbf{y}_j) \quad i = 1, \dots, N_t.$$

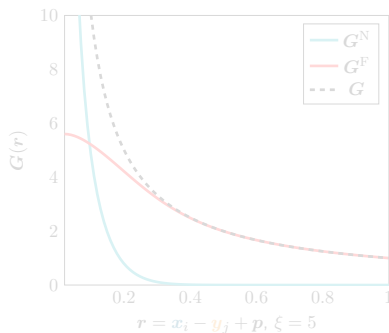
- Ewald's method splits

$$u(\mathbf{x}_i) = u^N(\mathbf{x}_i) + u^F(\mathbf{x}_i)$$

by setting $G(\mathbf{r}) = G^N(\mathbf{r}; \xi) + G^F(\mathbf{r}; \xi)$.

- For the electrostatic potential,

$$G(\mathbf{r}) = \frac{1}{\|\mathbf{r}\|} = \underbrace{\frac{\operatorname{erfc}(\xi \|\mathbf{r}\|)}{\|\mathbf{r}\|}}_{G^N} + \underbrace{\frac{\operatorname{erf}(\xi \|\mathbf{r}\|)}{\|\mathbf{r}\|}}_{G^F}.$$



- Consider periodic N -body sums of the form

$$u(\mathbf{x}_i) = \sum_{j=1}^{N_s} \sum_{p \in \mathcal{P}} G(\mathbf{x}_i - \mathbf{y}_j + \mathbf{p}) f(\mathbf{y}_j) \quad i = 1, \dots, N_t.$$

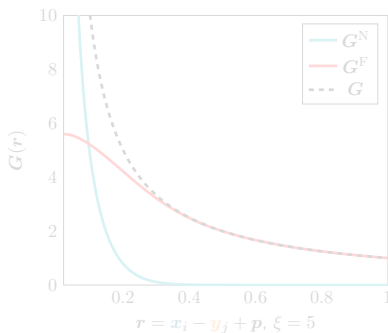
- Ewald's method splits

$$u(\mathbf{x}_i) = u^N(\mathbf{x}_i) + u^F(\mathbf{x}_i)$$

by setting $G(\mathbf{r}) = G^N(\mathbf{r}; \xi) + G^F(\mathbf{r}; \xi)$.

- For the electrostatic potential,

$$G(\mathbf{r}) = \frac{1}{\|\mathbf{r}\|} = \underbrace{\frac{\operatorname{erfc}(\xi \|\mathbf{r}\|)}{\|\mathbf{r}\|}}_{G^N} + \underbrace{\frac{\operatorname{erf}(\xi \|\mathbf{r}\|)}{\|\mathbf{r}\|}}_{G^F}.$$



- Consider periodic N -body sums of the form

$$u(\mathbf{x}_i) = \sum_{j=1}^{N_s} \sum_{\mathbf{p} \in \mathcal{P}} G(\mathbf{x}_i - \mathbf{y}_j + \mathbf{p}) f(\mathbf{y}_j) \quad i = 1, \dots, N_t.$$

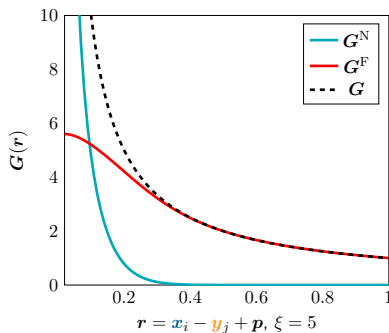
- Ewald's method splits

$$u(\mathbf{x}_i) = u^N(\mathbf{x}_i) + u^F(\mathbf{x}_i)$$

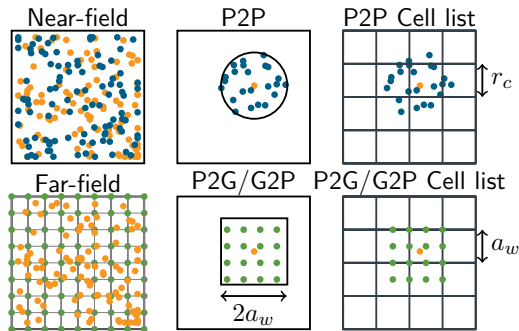
by setting $G(\mathbf{r}) = G^N(\mathbf{r}; \xi) + G^F(\mathbf{r}; \xi)$.

- For the electrostatic potential,

$$G(\mathbf{r}) = \frac{1}{\|\mathbf{r}\|} = \underbrace{\frac{\operatorname{erfc}(\xi \|\mathbf{r}\|)}{\|\mathbf{r}\|}}_{G^N} + \underbrace{\frac{\operatorname{erf}(\xi \|\mathbf{r}\|)}{\|\mathbf{r}\|}}_{G^F}.$$



- *Padded cell lists* $\mathcal{C}^N$ accelerate **P2P**, **P2G**, and **G2P** by assigning each Kokkos team T to a cell β so each thread $t_x \in T$ processes spatially close data.
- $\mathcal{C}^N$ is constructed in two passes through an unsorted list of particles



- For the stokes potential,

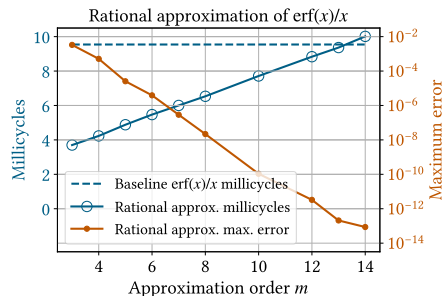
$$\mathbf{G}^N(\mathbf{r}) = \left(\mathbf{I} + \frac{\mathbf{r} \otimes \mathbf{r}}{\|\mathbf{r}\|^2} \right) \left(\operatorname{erfc}(\xi \|\mathbf{r}\|) + \frac{\|\mathbf{r}\| 2\xi e^{-\xi^2 \|\mathbf{r}\|^2}}{\sqrt{\pi}} \right) - \mathbf{I} \frac{4\xi e^{-\xi^2 \|\mathbf{r}\|^2}}{\sqrt{\pi}}.$$

Table: FP64 functions on the NVIDIA H200 GPU, measured with NVIDIA Nsight Compute.

Operation	Millicycles	Flop-equivalents	Actual flops
DFMA	0.121	2	2
DADD	0.127	2.10	1
DMUL	0.127	2.10	1
DIV	2.11	35.0	15 (2)
RSQRT	1.82	30.2	8
EXP	3.07	50.8	30
ERFC	16.1	267	112

- For the stokes potential,

$$\mathbf{G}^N(\mathbf{r}) = \left(\mathbf{I} + \frac{\mathbf{r} \otimes \mathbf{r}}{\|\mathbf{r}\|^2} \right) \left(\operatorname{erfc}(\xi \|\mathbf{r}\|) + \frac{\|\mathbf{r}\| 2\xi e^{-\xi^2 \|\mathbf{r}\|^2}}{\sqrt{\pi}} \right) - \mathbf{I} \frac{4\xi e^{-\xi^2 \|\mathbf{r}\|^2}}{\sqrt{\pi}}.$$



Algorithm Design

The window $w(\mathbf{r}) = \sum_{i=1}^d w_0(r_i)$ is identical in each dimension. How can we use this structure to build an efficient parallel scheme?

P2G-BASE: Too many atomics!

P2G-GRID: Redundant work!

P2G-HYBRID: Best of both worlds?

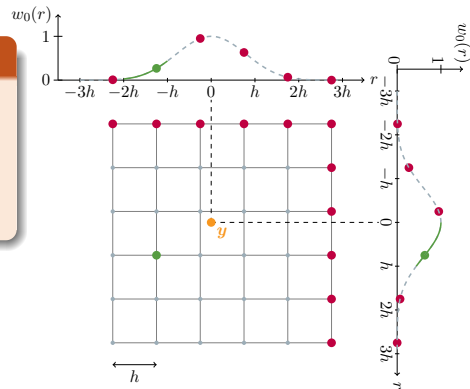
Algorithm Design

The window $w(\mathbf{r}) = \sum_{i=1}^d w_0(r_i)$ is identical in each dimension. How can we use this structure to build an efficient parallel scheme?

P2G-BASE: Too many atomics!

P2G-GRID: Redundant work!

P2G-HYBRID: Best of both worlds?



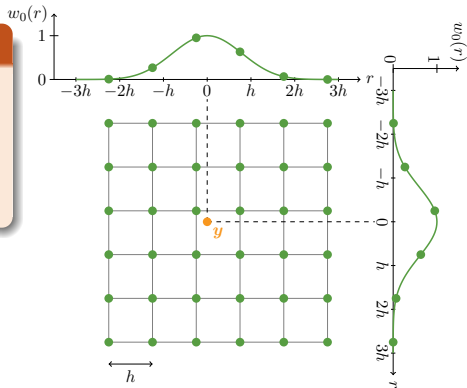
Algorithm Design

The window $w(\mathbf{r}) = \sum_{i=1}^d w_0(r_i)$ is identical in each dimension. How can we use this structure to build an efficient parallel scheme?

P2G-BASE: Too many atomics!

P2G-GRID: Redundant work!

P2G-HYBRID: Best of both worlds?



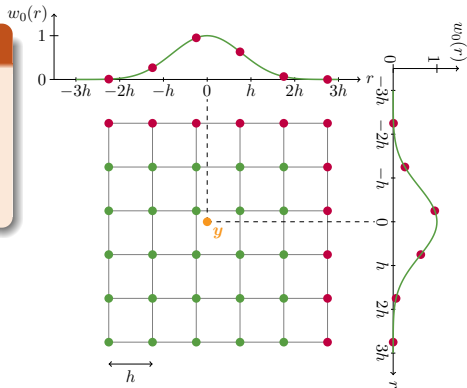
Algorithm Design

The window $w(\mathbf{r}) = \sum_{i=1}^d w_0(r_i)$ is identical in each dimension. How can we use this structure to build an efficient parallel scheme?

P2G-BASE: Too many atomics!

P2G-GRID: Redundant work!

P2G-HYBRID: Best of both worlds?



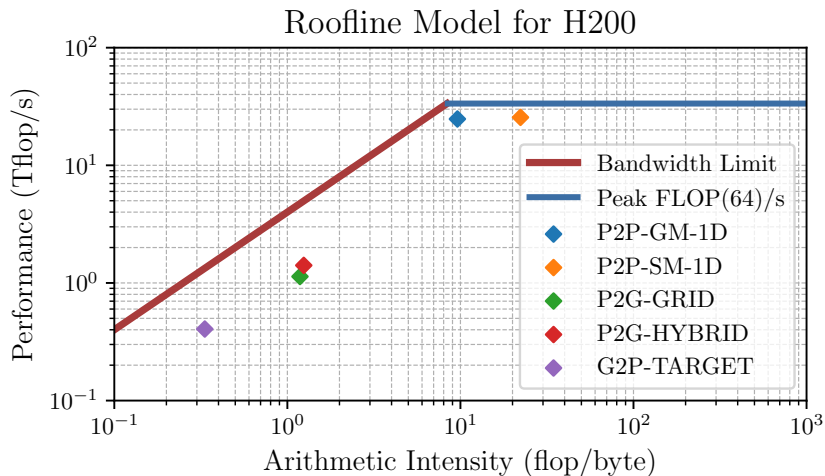


Table: P2P results on an NVIDIA H200 GPU.
Higher $\frac{N_s}{\mu s}$ is better.

N_t	method	$s = 256$		
		(b_x, b_y)	$\frac{N_t}{\mu s}$	Flops
10^6	GM-1D	(64, 1)	14.8	71%
	GM-2D	(1, 32)	11.8	56%
	SM-1D	(32, 1)	15.2	72%
	SM-2D	(64, 2)	12.4	59%
$4 \cdot 10^6$	GM-1D	(64, 1)	15.3	73%
	GM-2D	(8, 32)	11.9	57%
	SM-1D	(32, 1)	15.5	74%
	SM-2D	(64, 2)	12.4	59%

Table: P2G results on an NVIDIA H200 GPU.
Higher $\frac{N_s}{\mu s}$ is better.

N_s	method	$P = 8, E = 10^{-9}$		
		$\frac{N_s}{\mu s}$	Spd	Mops
10^6	BASE	10.1	—	22%
	GRID	11.9	1.2×	26%
	HYBRID	63.8	6.3×	30%
$4 \cdot 10^6$	BASE	8.5	—	18%
	GRID	11.1	1.3×	24%
	HYBRID	59.9	7.1×	28%

Table: P2P results on an NVIDIA H200 GPU.
Higher $\frac{N_s}{\mu s}$ is better.

N_t	method	$s = 256$		
		(b_x, b_y)	$\frac{N_t}{\mu s}$	Flops
10^6	GM-1D	(64, 1)	14.8	71%
	GM-2D	(1, 32)	11.8	56%
	SM-1D	(32, 1)	15.2	72%
	SM-2D	(64, 2)	12.4	59%
$4 \cdot 10^6$	GM-1D	(64, 1)	15.3	73%
	GM-2D	(8, 32)	11.9	57%
	SM-1D	(32, 1)	15.5	74%
	SM-2D	(64, 2)	12.4	59%

Table: P2G results on an NVIDIA H200 GPU.
Higher $\frac{N_s}{\mu s}$ is better.

N_s	method	$P = 8, E = 10^{-9}$		
		$\frac{N_s}{\mu s}$	Spd	Mops
10^6	BASE	10.1	—	22%
	GRID	11.9	1.2×	26%
	HYBRID	63.8	6.3×	30%
$4 \cdot 10^6$	BASE	8.5	—	18%
	GRID	11.1	1.3×	24%
	HYBRID	59.9	7.1×	28%

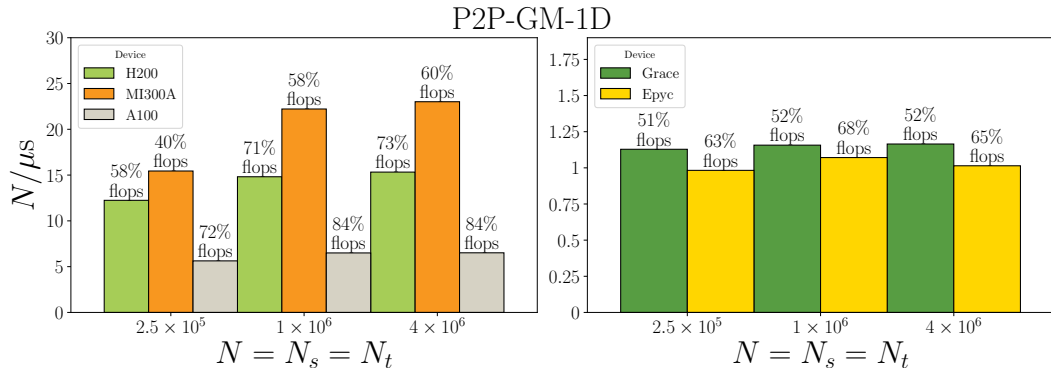


Figure: Particle-to-particle performance portability (higher is better), $s = 256$.

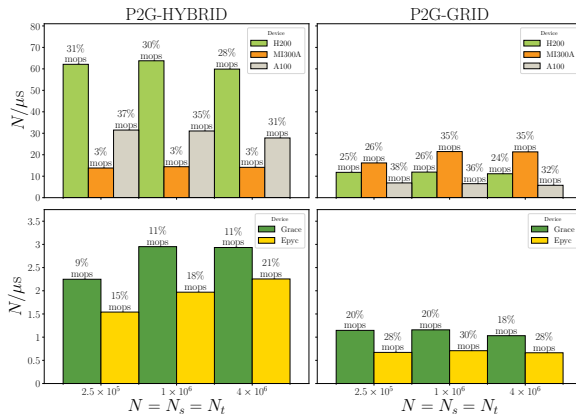


Figure: Particle-to-grid performance portability (higher is better). $E = 10^{-9}$, $s = 224$.

- Introduced algorithms for spectral Stokes Ewald sum with performance analysis.
- Library implemented in PyKokkos, MPI results show preliminary overheads.
- Limitations (i.e., future work):
 - No single-precision tests.
 - CPU kernels and MPI communication not optimized.
 - Nonuniformity effects not studied.
 - Fix shared memory slowdowns on AMD GPUs.
 - No multi-GPU NUMA scaling.

Thank You!

Gabriel Kosmacher
gkosmacher@utexas.edu

- Introduced algorithms for spectral Stokes Ewald sum with performance analysis.
- Library implemented in PyKokkos, MPI results show preliminary overheads.
- Limitations (i.e., future work):
 - No single-precision tests.
 - CPU kernels and MPI communication not optimized.
 - Nonuniformity effects not studied.
 - Fix shared memory slowdowns on AMD GPUs.
 - No multi-GPU NUMA scaling.

Thank You!

Gabriel Kosmacher
gkosmacher@utexas.edu

- **Particle-to-particle (P2P):**

$$\mathbf{u}^N(\mathbf{x}_i) := \sum_{j \in [N_s]} \sum_{\mathbf{p} \in \mathcal{P}} \mathbf{G}^N(\mathbf{x}_i - \mathbf{y}_j + \mathbf{p}) f(\mathbf{y}_j) \quad \mathcal{O}(27sN_t)$$

- **Particle-to-grid (P2G):**

$$\phi_h(\mathbf{g}_\ell) := \sum_{j \in [N_s]} \sum_{\mathbf{p} \in \mathcal{P}} w(\mathbf{g}_\ell - \mathbf{y}_j + \mathbf{p}) f(\mathbf{y}_j) \quad \mathcal{O}(N_s P^3)$$

- **Fourier Grid Convolution (FGC):**

$$\mathbf{v}_h(\mathbf{g}_\ell) := \text{IFFT} \left\{ \hat{\mathbf{G}}^F(\boldsymbol{\kappa}_\ell) [\hat{w}(\boldsymbol{\kappa}_\ell)]^{-2} \text{FFT} \{ \phi_h(\mathbf{g}_\ell) \} \right\} \quad \mathcal{O}(N_g \log N_g)$$

- **Grid-to-particle (G2P):**

$$\mathbf{u}_h^F(\mathbf{x}_i) := h^3 \sum_{\ell \in [N_g]} \sum_{\mathbf{p} \in \mathcal{P}} w(\mathbf{x}_i - \mathbf{g}_\ell + \mathbf{p}) \mathbf{v}_h(\mathbf{g}_\ell) \quad \mathcal{O}(N_t P^3)$$

```
0: parfor  $x_i \in \beta$  do                                // Team gets  $\beta$ , TeamThreadRange over  $x_i$ 
1:   parfor  $c \leftarrow 1, \dots, \lceil s/|\alpha_c| \rceil$  do      // VectorThreadRange over  $y_j$ 
2:      $\forall y_j \in \alpha_c, \text{LOADSCRATCH}(y_j, f(y_j))$           // GM/SM
3:     ____TEAMBARRIER( )
4:     for  $y_j \in \alpha_c$  do
5:        $u_c \leftarrow u_c + G^N(x_i - y_j + p)f(y_j)$           //  $G^N$  is flop intensive!
6:     end for
7:     ____TEAMBARRIER( )
8:   end parfor
9:    $u_h^N(x_i) \leftarrow \text{VECTORREDUCE}(u_c)$                 // Kokkos provided reduction
10: end parfor
```

```

0: parfor  $y_j \in \beta$  do                                     // Team gets  $\beta$ , TeamThreadRange over  $y_j$ 
1:   Compute  $w_1[l], w_2[l], w_3[l]$  for  $l \leftarrow 0, \dots, P-1$  and store in shared memory.
2: end parfor
3: ____TEAMBARRIER( )
4: parfor  $t_x \leftarrow 0, \dots, (P/2)^3 - 1$  do           // TeamThreadRange over  $(P/2)^3$  threads
5:   for  $\alpha \in \text{Colleagues}(\beta)$  do
6:      $g_\ell \leftarrow \text{INDEX2GRID}(t_x, \alpha)$            // Assign  $g_\ell$  to a thread
7:     for  $y_j \in \beta$  do
8:       LOADSCRATCH( $w_1[(g_\ell - a_j)_1], w_2[(g_\ell - a_j)_2], w_3[(g_\ell - a_j)_3], f(y_j)$ )
9:        $\phi_{\text{loc}} \leftarrow \phi_{\text{loc}} + w_1 w_2 w_3 f(y_j)$ 
10:    end for
11:    ATOMICADD( $\phi(g_\ell), \phi_{\text{loc}}$ )                         //  $\mathcal{O}(27N_s)$  atomics
12:  end for
13: end parfor

```

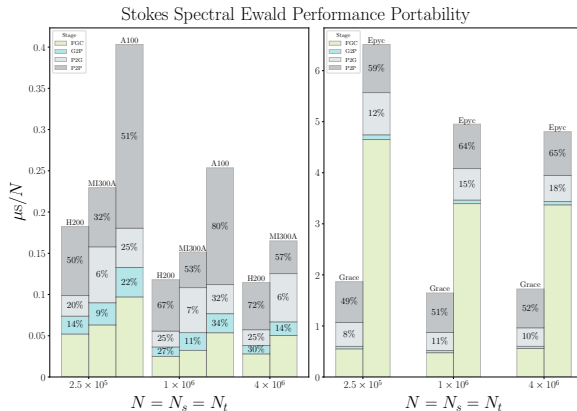


Figure: Spectral Ewald run on multiple machines (lower is better). $E = 10^{-9}$, $s = 224$.

Table: Runtime in *milliseconds* for multi-GPU weak scaling test with N_{GPU} NVIDIA H200 GPUs, MPI communication, and NVIDIA's cuFFTMp library. $E = 10^{-9}$, $s = 224$, $N = 4 \times 10^6 N_{\text{GPU}}$.

Step	$N_{\text{GPU}} = 1$	2	4	8	16
P2P	217.7	223.8	225.4	226.9	226.3
P2G	72.2	75.7	75.7	75.7	75.5
FFT	46.1	581.8	737.4	772.2	686.1
IFFT	1.5	24.0	28.6	30.7	26.1
G2P	39.9	40.6	40.6	40.4	40.7
MPI-SORT	5.6	22.7	23.2	36.4	61.7
MPI-GHOST-SOURCE	0.4	6.8	7.0	6.6	6.4
MPI-GHOST-GRID	0.2	3.2	5.2	4.7	4.6