

25–26 Nov. 2025

ANALYSIS AND OPTIMIZATION OF MEMORY REQUIREMENTS FOR STELLA (AND) CONTINUOUS PERFORMANCE MONITORING FOR GVEC

Joan Vinyals Ylla-Català
Valentin Seitz
Marta Garcia Gasulla

joan.vinyals@bsc.es
valentin.seitz@bsc.es
marta.garcia@bsc.es

3rd Annual Meeting of EUROfusion HPC ACHs



This work has been carried out within the framework of the EUROfusion Consortium, funded by the European Union via the Euratom Research and Training Programme (Grant Agreement No 101052200 — EUROfusion). Views and opinions expressed are however those of the author(s) only and do not necessarily reflect those of the European Union or the European Commission. Neither the European Union nor the European Commission can be held responsible for them.

25–26 Nov. 2025

STELLA

Valentin Seitz

3rd Annual Meeting of EUROfusion HPC ACHs



This work has been carried out within the framework of the EUROfusion Consortium, funded by the European Union via the Euratom Research and Training Programme (Grant Agreement No 101052200 — EUROfusion). Views and opinions expressed are however those of the author(s) only and do not necessarily reflect those of the European Union or the European Commission. Neither the European Union nor the European Commission can be held responsible for them.



Goal: Enabling multiscale simulations

Stella is MPI only.

Scalability in terms of real-space resolution of Stella was majorly hindered by memory usage.

→ Detailed analysis of memory scaling of high resolution input case:

Run different configurations, measure memory allocations per function and build analytical model with Extra-P

- `fields::allocate_arrays`: 0.5GB*p
- `init_g::ginit_noise`: 2GB + 0.09GB*p
- `dist_fn::init_kperp2`: 0.04GB*p

→ linear dependence on p means memory replicated across processes

The models estimated the memory requirements for target run

The screenshot shows the Stella software interface with a table titled 'allocation_size'. The table lists various memory allocation calls and their associated metrics. The columns are: Selection, Model (Default Model), Metric (allocation_size), Set Callpath, Annotations, Value, RSS, Adj. R², and SMAPE. The table is filtered to show 'All' results.

Selection	Model	Metric	Set Callpath	Annotations	Value	RSS	Adj. R ²	SMAPE
	Default Model	allocation_size	malloc		-3.928x10 ⁻¹² + 672...	2.260...	1.000	2.538...
			vpamu_grids::init_mu_grid					
			malloc		-2.440x10 ⁻¹¹ + 279...	1.821...	1.000	3.665...
			gauss_quad::get_laguerre_grids					
			malloc		-3.421x10 ⁻¹³ + 64...	1.086...	1.000	2.498...
			malloc		-8.553x10 ⁻¹⁴ + 16...	6.787...	1.000	2.498...
			extended_zgrid::init_extended_zgrid					
			malloc		-1.072x10 ⁻⁰⁷ + 7.4...	4.020...	1.000	3.111...
			volume_averages::init_volume_averages					
			malloc		-4.816x10 ⁻¹⁰ + 8.4...	1.085...	1.000	3.002...
			dist_fn::init_dist_fn					
			dist_fn::allocate_arrays					
			malloc		2.294x10 ¹⁰	0.000	1	0.000
			dist_fn::init_kperp2					
			malloc		-3.346x10 ⁻⁰⁷ + 4.4...	2.416...	1.000	1.247...
			dist_fn::enforce_single_valued_kperp2					
			malloc		-4.277x10 ⁻¹⁴ + 8.0...	1.697...	1.000	2.498...
			dist_fn::init_vperp2					
			malloc		-3.219x10 ⁻¹² + 288...	1.654...	1.000	9.869...
			gyro_averages::init_bessel					
			malloc		6.451x10 ⁰⁹	0.000	1	0.000
			dist_redistribute::init_redistribute					
			dist_redistribute::init_kxyz_to_vmu_redistribute					
			malloc		1.147x10 ¹⁰	4.661...	1	7.907...
			redistributes::init_redist					
			malloc		1.196x10 ¹⁰	1.242...	1	3.887...
			dissipation::init_dissipation					
			dissipation::read_parameters					
			file_utils::input_unit_exist					
			malloc		1140.000	0.000	1	0.000
			hyper::read_parameters_hyper					
			file_utils::input_unit_exist					
			malloc		588.000	0.000	1	0.000
			sources::init_sources					
			sources::read_parameters					
			file_utils::input_unit_exist					
			malloc		880.000	0.000	1	0.000
			fields::init_fields					
			fields::allocate_arrays					
			malloc		-2.048x10 ⁻⁰⁹ + 3.2...	1.762...	1.000	2.266...
			fields_fluxtube::init_fields_fluxtube					
			malloc		-6.880x10 ⁻⁰⁶ + 5.8...	6.185...	1.000	2.047...
			init_constants					
			malloc		3.225x10 ⁰⁹	1.559...	1	1.706...



Goal: Enabling multiscale simulations

Another Problem in Stella: Available Parallelism with MPI as it's often used in velocity space parallel mode

Distributed velocity space:

`nvpa * nmu * nspec =`

$(2 \times 48) \times 12 \times 2 = 2304$ processes max $\rightarrow$ 20 Nodes in MN5 when every process has exactly one velocity space point. (not really ideal?)

Distributed real space:

`naky * nakx * nzed * ntubes * nspec =`

$1115 * 558 * 64 * 1? * 2 = 80e6$ processes $\rightarrow$ More than full MN5 of parallelism



Strategy: Port Parts to OpenMP

The memory models indicate that 2 processes per node would work on current platforms, so we chose to do add OpenMP parallelism.

Focus: timestepping

~60% of runtime for timestep OpenMP parallelized and partially verified before major refactor upstream.

Greatest missing region is and advance_implicit(20%)



Whats next?

Merge OpenMP porting upstream.

With OpenMP mostly in the timestep, we find that the computation of the response terms and the corresponding matrices is a major bottleneck. While this is initially only done in the initialization, it's also required once the time step changes (which can happen quite often).

The next step is to find solution strategies to initialization of response terms (varying matrix sizes).

25–26 Nov. 2025

GVEC

Joan Vinyals Ylla-Català

3rd Annual Meeting of EUROfusion HPC ACHs



This work has been carried out within the framework of the EUROfusion Consortium, funded by the European Union via the Euratom Research and Training Programme (Grant Agreement No 101052200 — EUROfusion). Views and opinions expressed are however those of the author(s) only and do not necessarily reflect those of the European Union or the European Commission. Neither the European Union nor the European Commission can be held responsible for them.



Background

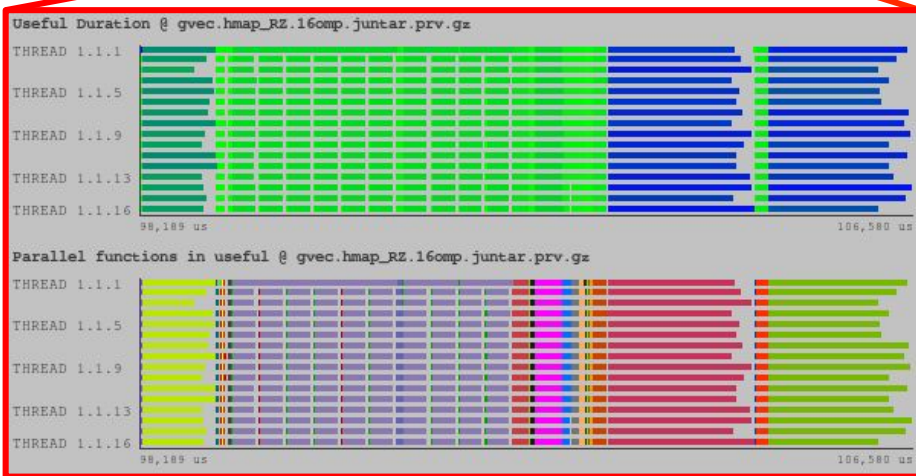
- **Code:** GVEC (Galerkin Variational Equilibrium Code) is an open-source software for the generation of three-dimensional ideal magnetohydrodynamic (MHD) equilibria.
- **Programming:** Fortran, OpenMP, MPI
- **Inputs:** hmap_RZ, hmap_axisNB
- **Analysis:** Scalability of OpenMP threads, with two inputs.
- **Analysis platform:** MareNostrum 5 (GPP partition, 112 CPUs/node, [more info](#))



Structure & Focus of analysis - Input: hmap_RZ



We observe an iterative pattern

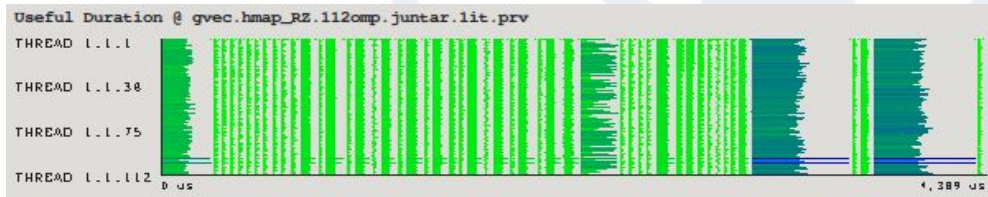
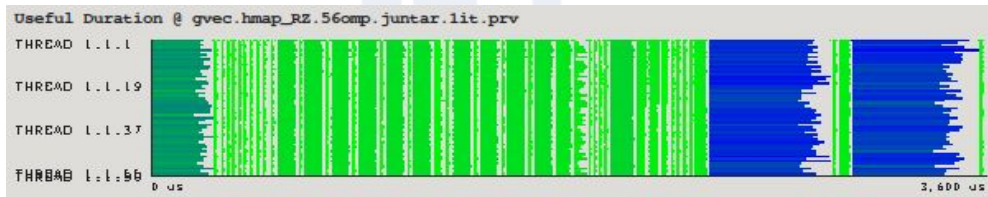
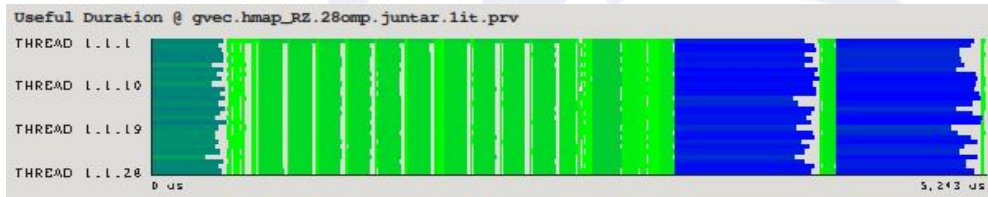
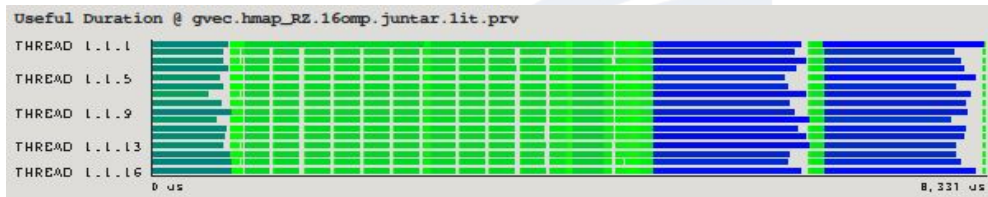
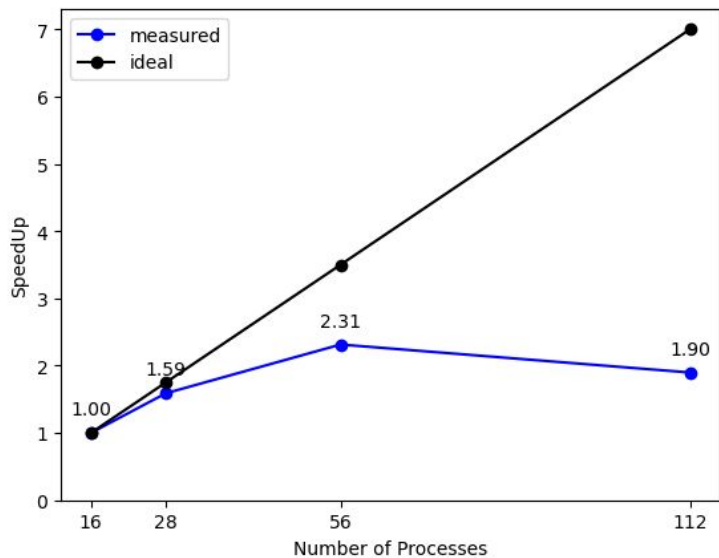


Within this iterative pattern we observe a sub-structure, which correspond to different OpenMP Parallel regions.

We observe that the repetitive pattern doesn't change across iterations, therefore, our FoA will be one iteration.

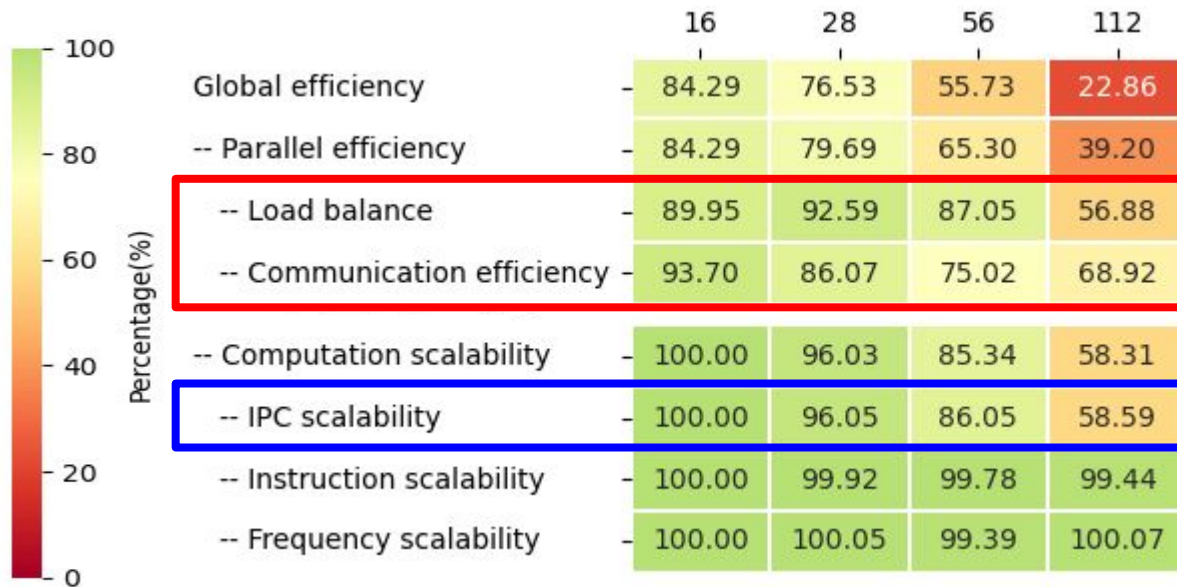


Scalability - Input: hmap_RZ





Efficiency Metrics - Input: hmap_RZ



Further analysis:
showed that the load
imbalance was caused
by **differences in IPC**
across processes.

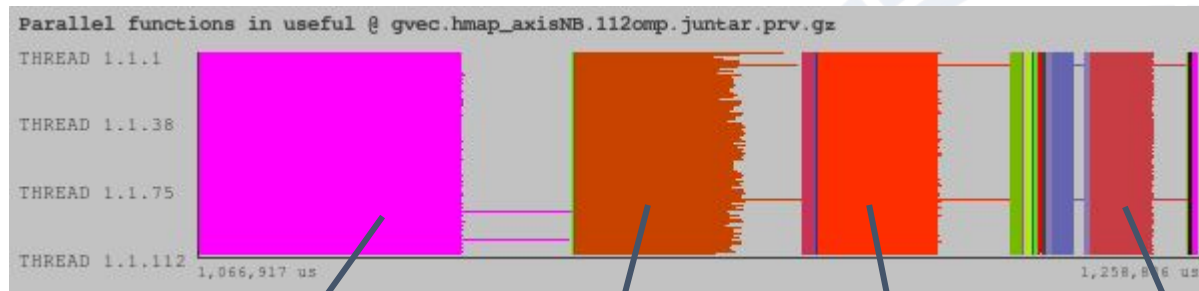


Efficiency metrics - hmap_RZ vs hmap_axisNB

Input	hmap_RZ				hmap_axisNB			
NUM THREADS	16	28	56	112	16	28	56	112
Global efficiency	84.29	76.53	55.73	22.86	0.89	0.88	0.88	0.58
-- Parallel efficiency	84.29	79.69	65.30	39.20	99.11	99.05	97.77	66.13
-- Load balance	89.95	92.59	87.05	56.88	99.73	99.38	99.11	83.66
-- Communication efficiency	93.70	86.07	75.02	68.92	99.37	99.67	98.65	79.04
-- Computation scalability	100.00	96.03	85.34	58.31	0.90	0.89	0.90	0.87
-- IPC scalability	100.00	96.05	86.05	58.59	83.52	83.42	83.37	82.05
-- Instruction scalability	100.00	99.92	99.78	99.44	1.06	1.06	1.06	1.06
-- Frequency scalability	100.00	100.05	99.39	100.07	101.63	101.01	102.02	100.53



Load Balance - Input: hmap_axisNB



Parallel Eff.	70.41%	71.12%	63.05%	64.45%
- Communication Eff.	99.25%	99.59%	99.61%	99.51%
- Load Balance Eff.	70.94%	71.41%	63.30%	64.77%

Further analysis: that all unbalanced regions were caused by variations in IPC across processes.



Background

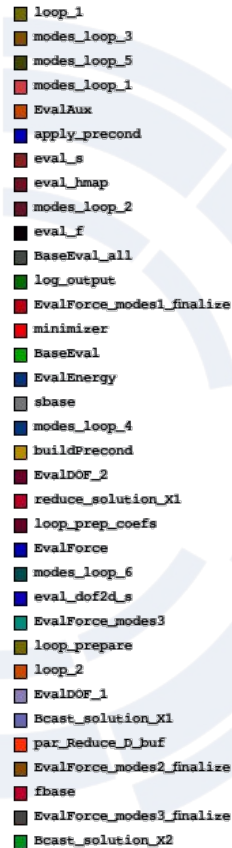
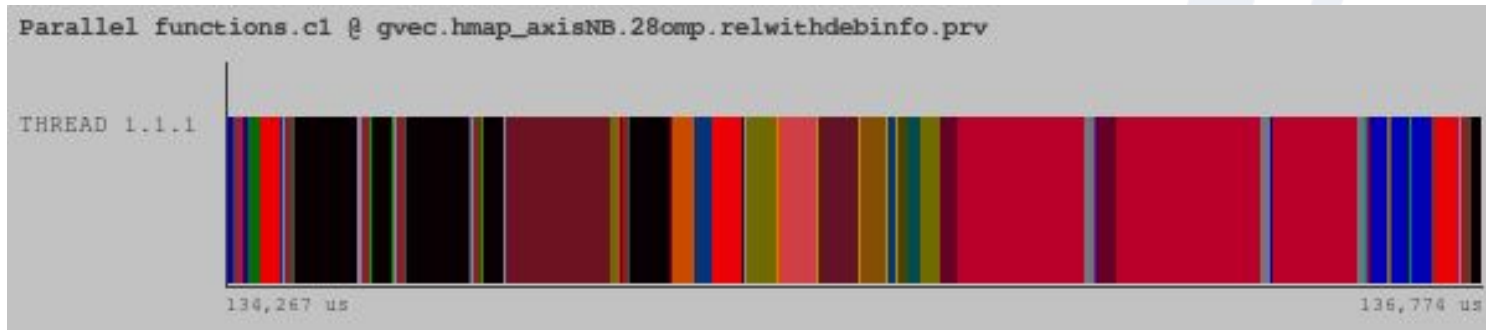
- At this point, based on the analysis comparison of both inputs, the code developer implemented some changes, to improve its performance.





Continuous performance monitoring - Manual Instrumentation

- We used the code annotations they already had implemented in the code to plug in TALP DLB.
- To provide the user with more insight on regions performance.



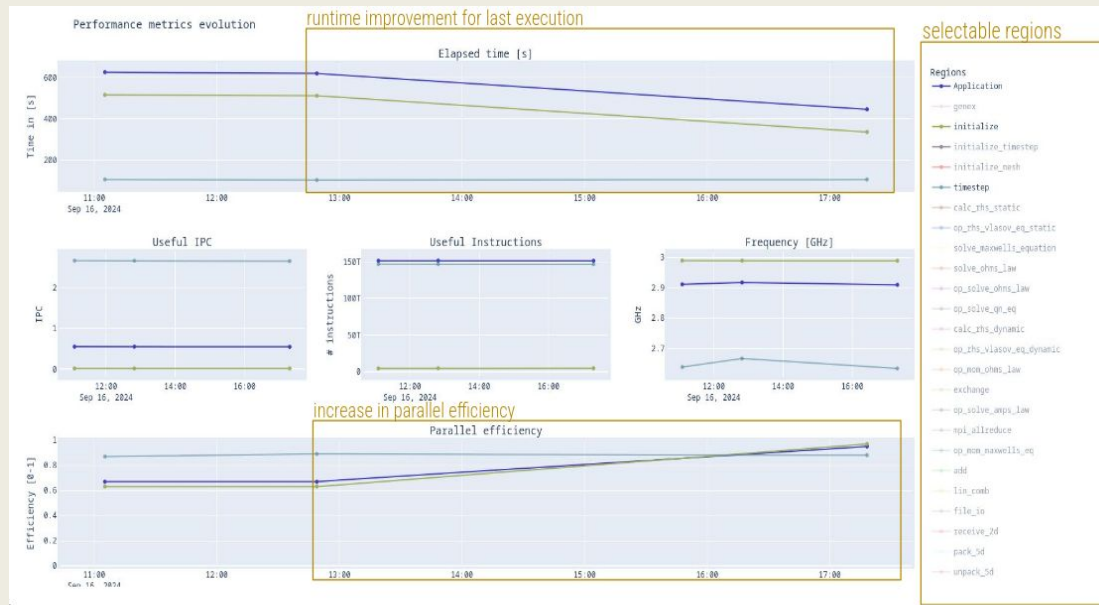


Continuous performance monitoring - TALP Pages

Metrics	8xOpenMP	36xOpenMP	72xOpenMP
Global Efficiency	0.95	0.65	0.43
- Parallel efficiency	0.95	0.68	0.48
- OpenMP Parallel efficiency	0.95	0.68	0.48
-- OpenMP Scheduling efficiency	1	1	1
-- OpenMP Load balance	0.95	0.68	0.48
-- OpenMP Serialization efficiency	1	1	0.99
- Computation Scalability	1	0.96	0.89
- Instructions scaling	1	1	1
- IPC scaling	1	1.03	1.03
- Frequency scaling	1	0.93	0.87
Useful IPC	1.98	2.04	2.03
Frequency [GHz]	2.25	2.1	1.95
Elapsed time [s]	44.02	14.2	10.85

- With TALP Pages we get performance scaling efficiency tables for every commit like the one above.

Example from another code:



More examples:

https://tu-dresden.de/zih/das-department/ressourcen/dateien/toolworkshop/2024_09_19-ValentinSeitz.pdf?lang=en

Integrate TALP Pages to the GitLab CI/CD infrastructure of GVEC, to provide region specific efficiency and execution time evolutions across commits.

25–26 Nov. 2024

ANALYSIS AND OPTIMIZATION OF MEMORY REQUIREMENTS FOR STELLA (AND) CONTINUOUS PERFORMANCE MONITORING FOR GVEC

Joan Vinyals Ylla-Català
Valentin Seitz
Marta Garcia Gasulla

joan.vinyals@bsc.es
valentin.seitz@bsc.es
marta.garcia@bsc.es

3rd Annual Meeting of EUROfusion HPC ACHs



This work has been carried out within the framework of the EUROfusion Consortium, funded by the European Union via the Euratom Research and Training Programme (Grant Agreement No 101052200 — EUROfusion). Views and opinions expressed are however those of the author(s) only and do not necessarily reflect those of the European Union or the European Commission. Neither the European Union nor the European Commission can be held responsible for them.