

Energetically Efficient Acceleration EEA-Aware For Scientific Applications of Large-Scale On Heterogeneous Architectures

John A. García. H.^{1,2,4}, Esteban Hernandez B.^{1,3}, Philippe O. A. Navaux², Carlos. J. Barrios H.¹
^{1,2,4}jgarciahenao@acm.org, ^{1,3}ejhernandezb@udistrital.edu.co, ²navaux@inf.ufrgs.br, ¹cbarrios@uis.edu.co

¹ High Performance and Scientific Computing Center, SC3
Universidad Industrial de Santander, UIS - Bucaramanga, Colombia

² Parallel and Distributed Processing Group, GPPD
Universidade Federal do Rio Grande do Sul, UFRGS - Porto Alegre, Brasil

³ Interoperability of Academic Networks Group - GIIRA,
Universidad Distrital Francisco Jose de Caldas, UD - Bogota, Colombia

⁴ Laboratoire d'Informatique, Signaux et Systèmes de Sophia Antipolis - I3S,
Université Nice Sophia Antipolis, UNS - Nice, France

Challenge

Heterogeneous parallel programming has two main problems on large computation systems:

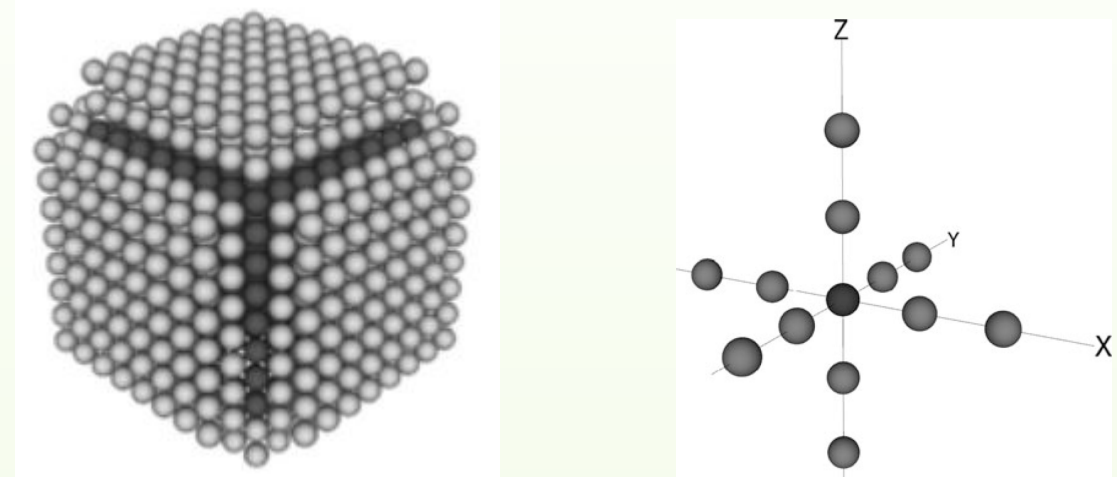
- The increase of power consumption on supercomputers in proportion to the amount of computational resources used to obtain high performance.
- The underuse these resources by scientific applications with improper distribution of tasks.

"Select the optimal computational resources and make a good mapping of task granularity is the fundamental challenge for build the next generation of Exascale Systems"

Analysis of Energy Efficiency and Portability of large Scale Scientific Application:

Case Study: Ondes3D®

Simulation of Seismic Wave Propagation: It's a crucial tool for analyzing geophysical strong movements.



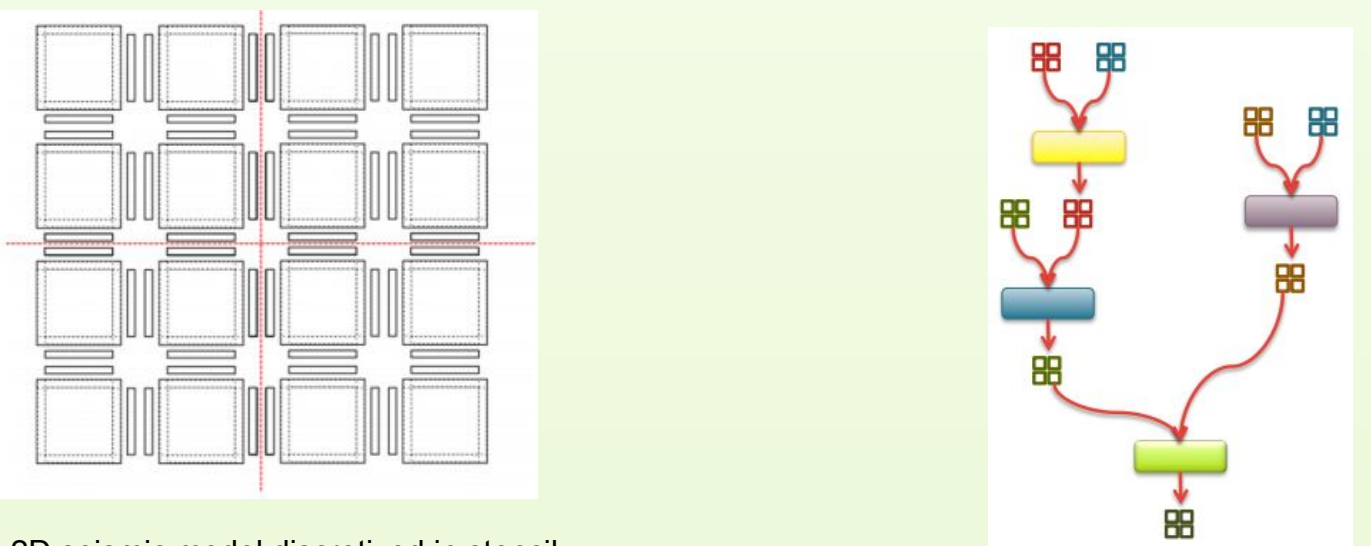
3D Seismic Model with blocks discretization
Spatial Data Dependencies Stencil

Ref John A. García H. et al. Energetically Efficient Acceleration EEA-Aware. Degree work to obtain the title of Master of Science in Systems Engineering and Informatics at UIS 2016.
Ref. Victor Martinez et al. Towards Science Wave Modeling on Heterogeneous Many-Core Architectures Using Task-based Runtime Systems. ISAC-PAD 2015.

Planning workload experimental on Ondes3D® StarPU® task programming library for hybrid architectures.

7.2 Millions Grid (300*300*80) - 625 MB
Block Size 150x150

36 Task processing (20 Speed; 16, Effort)
144 Task of communication



Ref John A. García H. et al. Energetically Efficient Acceleration EEA-Aware. Degree work to obtain the title of Master of Science in Systems Engineering and Informatics at UIS 2016.
Ref. Victor Martinez et al. Task-based programming on Low-power Manycore Architectures: Seismic Simulation on Jetson TK1. CARLA 2015.
Ref. Olivier Aumage. Porting BRIGGS Ondes3D on top of StarPU: Task Parallelism on Heterogeneous platforms. NABA - University of Bordeaux. March 2013.

Heterogeneous Architectures

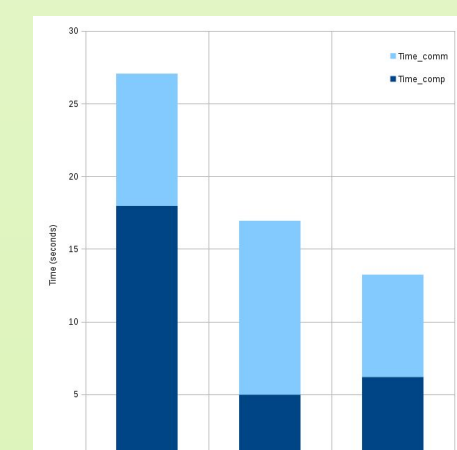
Results of running Ondes3D® with StarPU® on different machines.

Jetson TK1

Desktop: Salmon

Mean Power Consumption on Jetson TK1(15 Watts):
10 Watts by GPU use.
5 Watts by CPU use.

Time to Solution



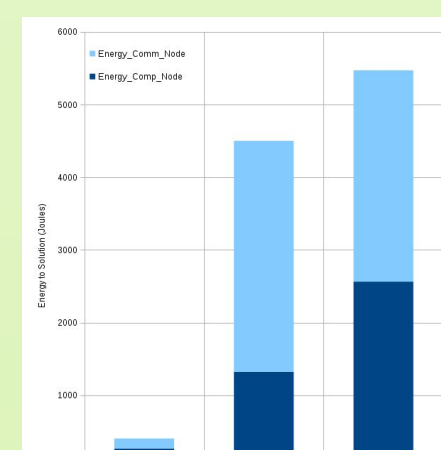
Jetson(27.06 seconds):
17.97 seconds in processing tasks.
9.09 seconds in transfer tasks.

Desktop(16.94 seconds):
4.03 seconds in processing tasks.
11.96 seconds in transfer tasks.

Node (13.23 seconds):
6.20 seconds in processing tasks.
7.03 seconds in transfer tasks.

Node (4570.4 Joules):
2892.1 Joules by using CPU y GPU.
2778.3 Joules by component in an idle state.
203.4W (N_run) + 210W (N_idle) = 413 Watts

Energy Consumed on Execution



Jetson(405.9 Joules):
2778.3 Joules by using CPU.
135.3 Joules by using GPU.
10W (GPU) + 5W (CPU) = 15 Watts

Desktop(4500.5 Joules):
2298.3 Joules by using GPU.
2202.2 Joules by using CPU.
135.6W (GPU) + 130W (CPU) = 265.6 Watts

Node (4570.4 Joules):
2892.1 Joules by using CPU y GPU.
2778.3 Joules by component in an idle state.
203.4W (N_run) + 210W (N_idle) = 413 Watts

Which is the task distribution to be mapping on heterogeneous architectures CPU-GPU,
that determine an efficient computational acceleration and low power consumption
during the time of scientific application execution?

Ref John A. García H. et al. Energetically Efficient Acceleration EEA-Aware. Degree work to obtain the title of Master of Science in Systems Engineering and Informatics at UIS 2016.

More Information:

www.sc3.uis.edu.co/
www.inf.ufrgs.br/gppd

Project Repository:

Stable:
<http://forge.sc3.uis.edu.co/redmine/projects/eea-uis>
Testing:
<https://github.com/jagh/enerGPU>

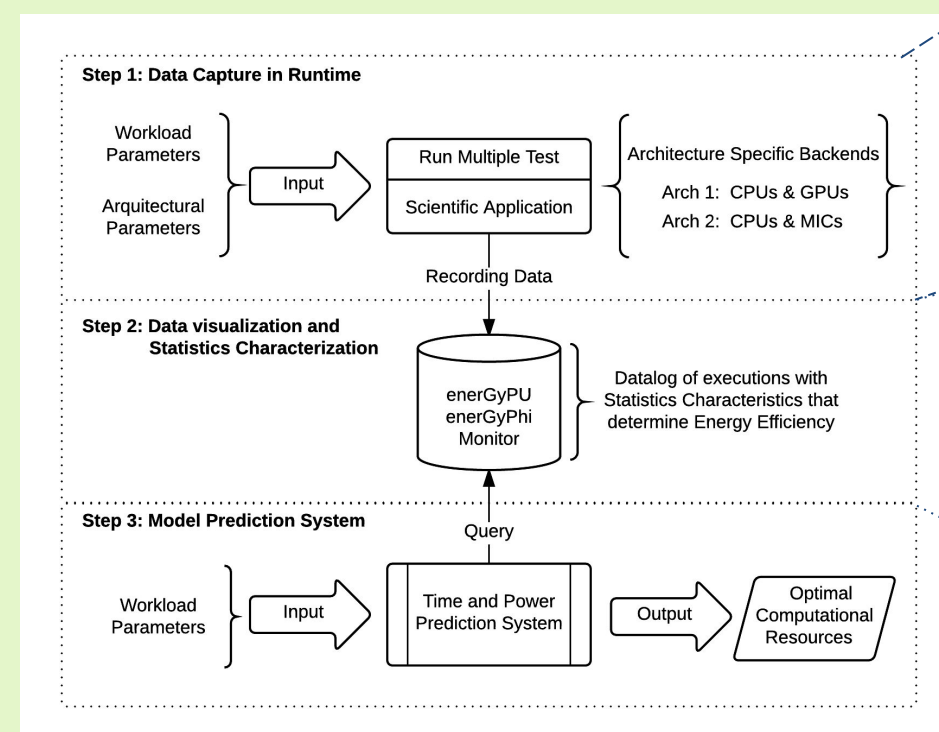


References

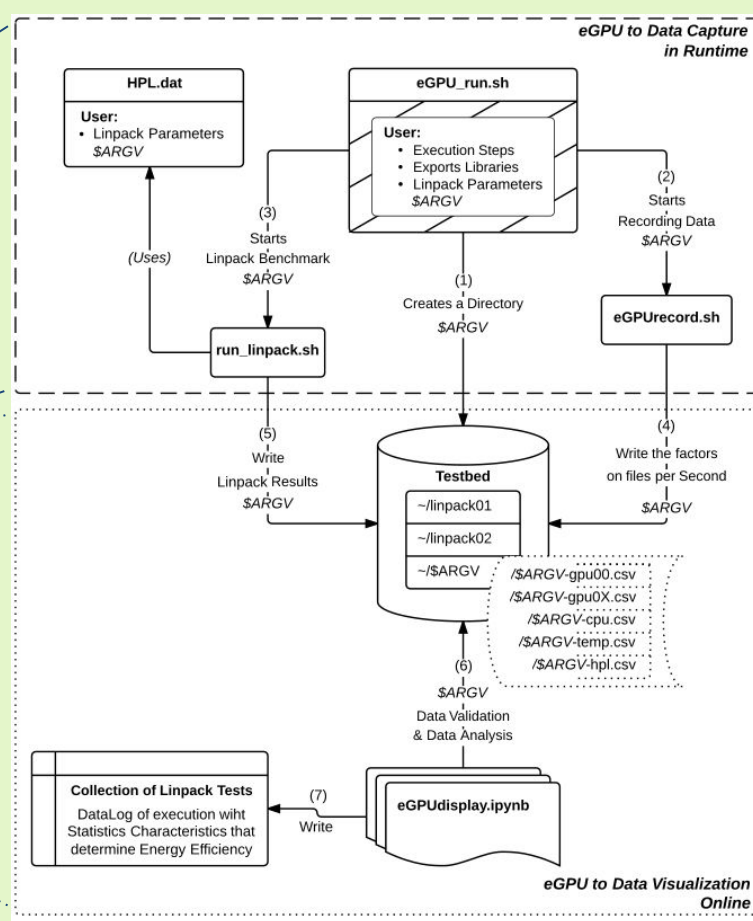
- John A. García H. et al. Energetically Efficient Acceleration EEA-Aware. Degree work to obtain the title of Master of Science in Systems Engineering and Informatics at UIS 2016.
- John A. García H. et al. enerGPU and enerGPhi Monitor for Power Consumption and Performance Evaluation on Nvidia Tesla GPU and Intel Xeon Phi. CCGrid 2016 - IEEE/ACM.
- John A. García H. et al. eGPU for Monitoring Performance and Power Consumption on Multi-GPUs. XIII Workshop de Processamento Paralelo e Distribuído, WSPPD 2015 - IEEE.
- Raj Jain. The art of computer systems performance analysis - techniques for experimental design, measurement, simulation, and modeling. Wiley professional computing. Wiley, 1991.
- K. Bergman et al. Exascale Computing Study: Technology Challenges in Achieving Exascale Systems. Peter Kogge, Editor & Study Lead. Technical report, 2008.

EEA-Aware Structure

This research proposes an integrated energy-aware scheme called Efficiently Energetic Acceleration (EEA) for scientific applications of large-scale on heterogeneous architectures. The EEA energy-aware scheme has a workflow of three steps as shown in EEA-Aware Scheme. In the first step, the data is captured in runtime and executed the enerGPU or enerGPhi monitor tool in parallel with the scientific application using different combinations of computational resources, applying the power capping technique for nodes with multi-GPUs. In the second step, the data visualization and statistics characterization used a separate level of enerGPU and enerGPhi monitor tool for analysis the key factors by results of each experiment in terms of energy efficiency and estimated the power levels. A deep description of monitor structure and utilization is present by García John et al. in [2], [3]. Finally using the data collected by monitors cost functions are build and model prediction system running for obtaining the optimal computational resources in a static time to mapping parallel task granularity of scientific applications on heterogeneous architectures.



EEA-Aware Scheme



enerGPU Monitor Structure

Step 1: Experimental Procedures and Results

At the first level the global parameters have chosen the workload and computational architecture according to the combination of resources that will be used. The experimental procedures were executed with a set of tests of HPL code variants using different workload and architectural parameters on Cluster nodes GUANE. The experimental procedures was chosen following the fractional factorial design principle proposed by Raj Jain [4].

The computational resources used: 1 GUANE Node on 'A' settings.

Configuraciones	A	B	C
Node type	SL390s	SL390s	SL390s
Number of nodes	8	8	5
Processor Intel	Xeon E5645	Xeon E5645	Xeon E5640
Processor Model	Sandy Bridge	Sandy Bridge	Sandy Bridge
Processor by node (#)	2	2	2
Core/Architectural	6	6	4
Thread/Core (#)	2	2	2
GPUS Nvidia	Tesla M2075	Tesla M2050	M2050
GPUS Model	Fermi	Fermi	Fermi
GPUs by node (#)	8	8	8
CUDA core (#)	448	448	448
Memory DDR4 (GB)	104	104	104
SAS disk (GB)	200	200	200
Gigabit Ethernet (Gbps)	10	10	10
InfiniBand (IB)	1	1	1

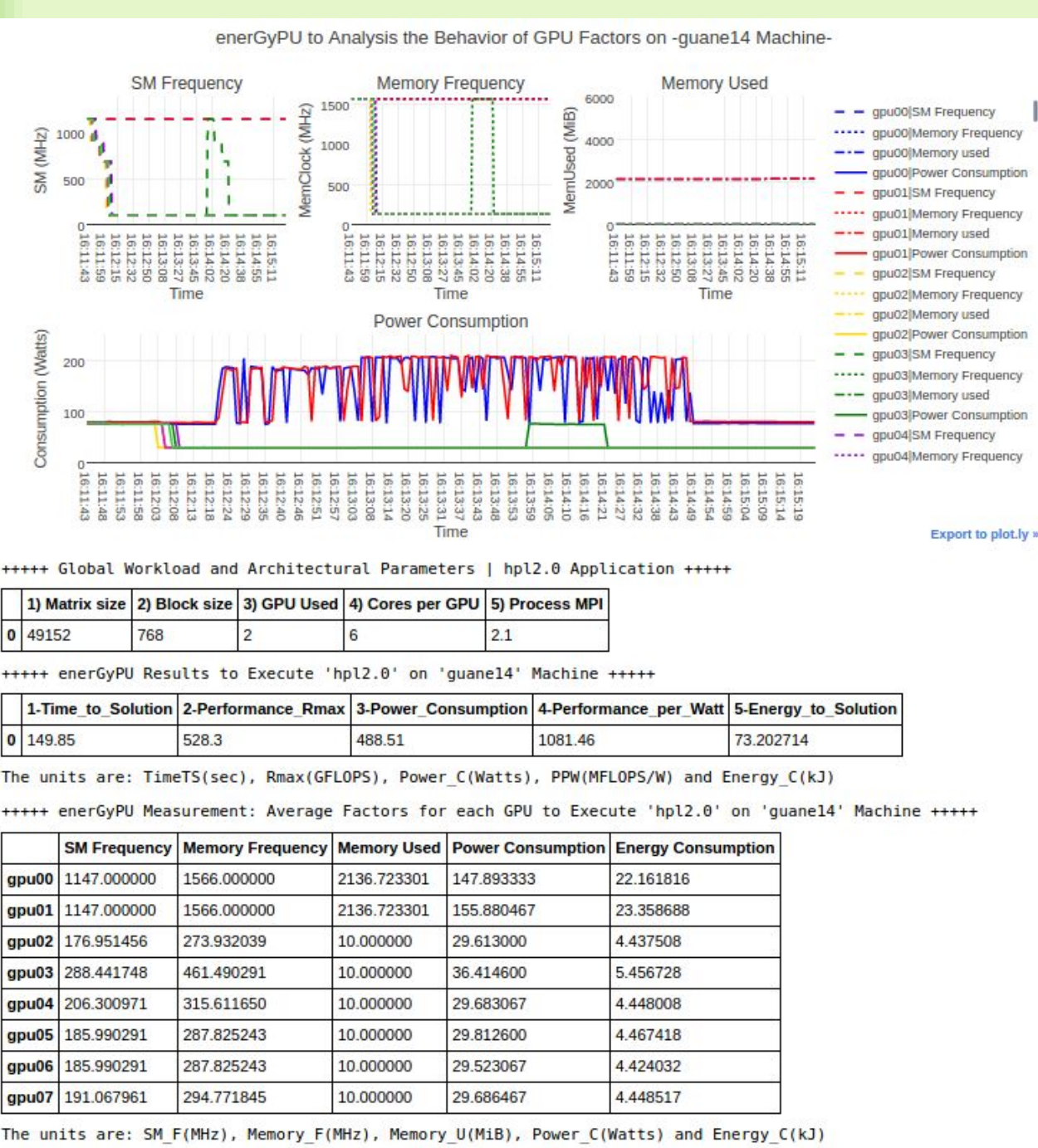
The Linpack used: HPL2.0 version configured for Tesla GPUs.

Ref. Massimiliano Patric. Accelerating Linpack with CUDA on heterogeneous clusters. ACM, 2009.

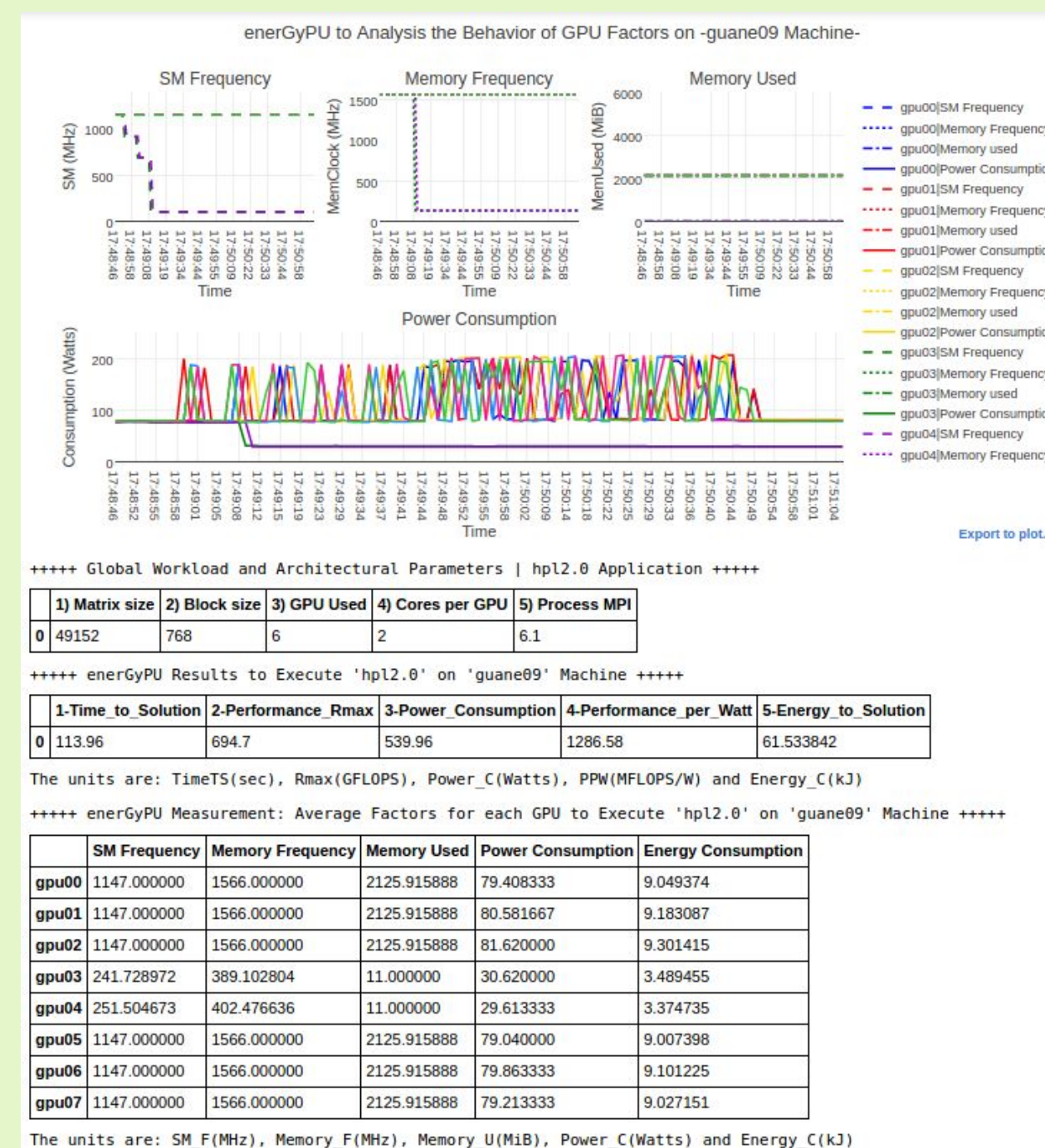
Experiments to solve HPL2.0	ES	PC	ET	ES	PC	ET	ES	PC	ET
Global Workload Parameters	49152			49140			73728		
Mx	768	2048	2048	768	2048	2048	1024	1536	1536
Bx									
TaskSize (MBytes)	384	768	768	384	960	960	576	864	864
TaskSystem	48	24	24	48	20	20	72	48	48
Architectural Parameters for Nvidia GPU									
GPUs/node	6	8	1	4	8	1	4	8	1
Cores/node	2	1	12	3	1	12	3	1	12
MPUs/node	6	8	1	4	8	1	4	8	1
enerGPU Results									
enerGPUTime (sec)	120	211	1	507	309	915	331	493	1350
enerGPUEnergy (sec)	6	6	5	6	6	6	6	6	6
HPL2Time (sec)	113.96	205.23	531.31	222.47	361.28	909.34	327.9	486.76	1341.09
Rmax (GFLOPS)	694.7	131.10	149	695	425.6	170	814.8	548.9	198.8
Power (Watts)	520.97	428	327.16	425.86	466.25	327.30	425.43	700.63	327.16
PPW (MJ/OPS/W)	1286.79	603.93	443.92	1297.44	625.90	494.95	1873.23	755.43	569.37
Efficiency (KJ)	61.52	131.10	179.13	96.7	243.13	312.38	142.77	341.04	469.25

Representative results of enerGPU to analyze the best Energetic to Evolution (ES), the worst Power Consumption (PC) and the worst Execution Time (ET) for each matrix problem.

Experiment: hpl2.0-guane14-2015073016-49152768622



Experiment: hpl2.0-guane09-2015073117-49152768626



Step 2: Architectural Characterization

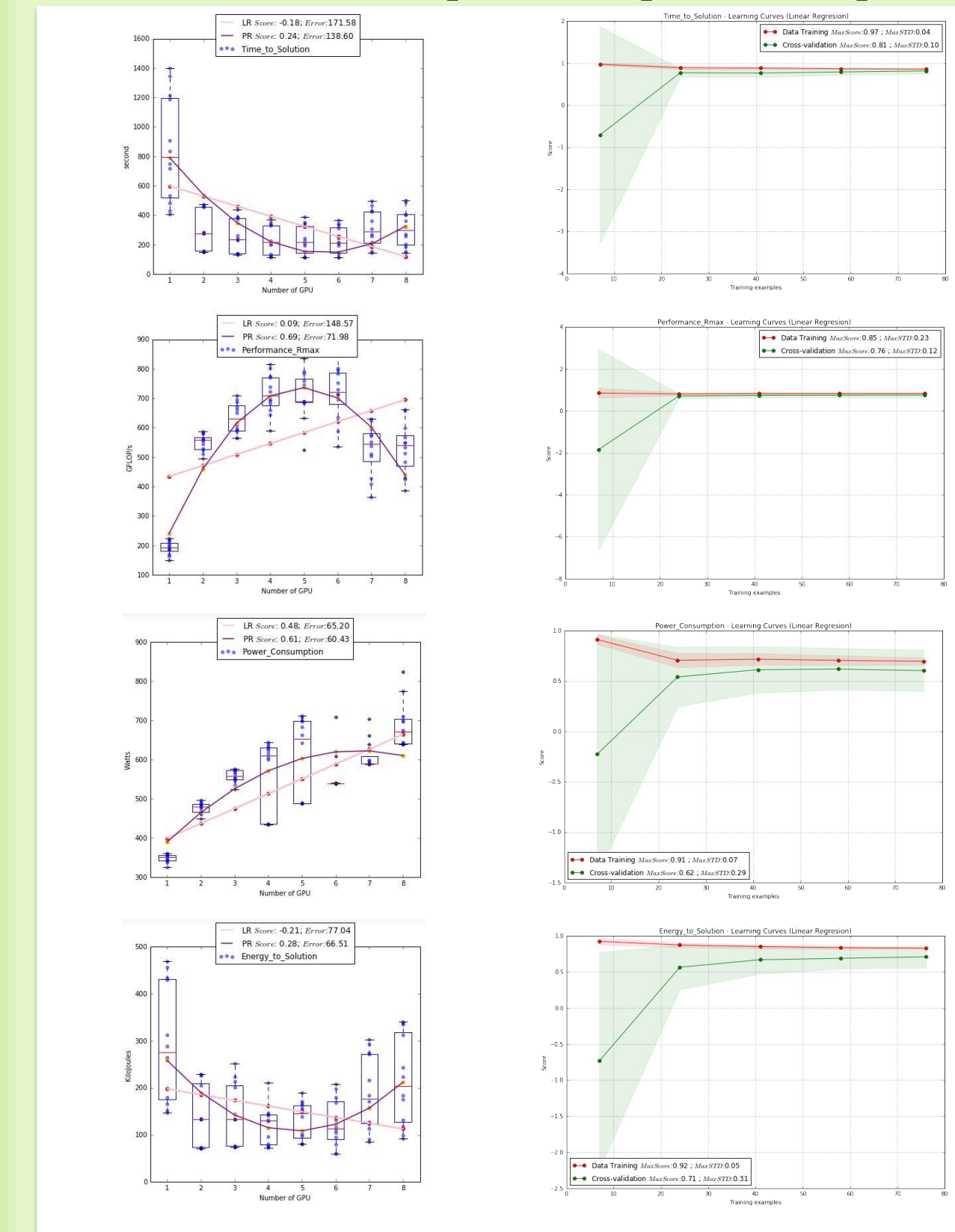
The second level uses the enerGPU monitor in the post-processing for data visualization and statistical characterization of each architecture. In which the Figure below shows that the accuracy using the key factors is much more accurate than just using the number of GPUs. In addition, this monitor displays information via sequence data, statistics, histograms and tables showing results in terms of energy efficiency, for each experiment.

Minimization model with Task and Architectural Parameters

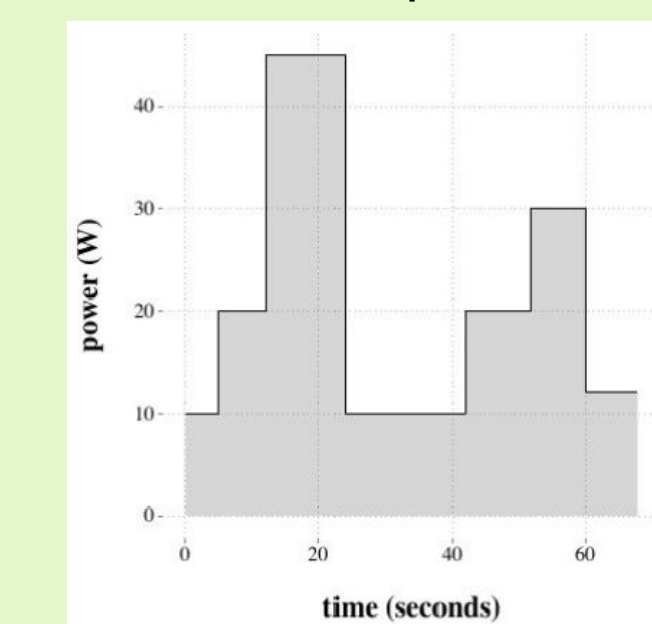
$GPU_{PowerRuntime}(GPU_worker[GPU_Factors], GPU_idle[GPU_Factors])$

Accuracy of Linear and Polynomial Models Among Learning Curves Multivariable Model

$plot1dModel(data[data.columns[Number_GPU]],data[data.columns[enerGPU_results]])$
 $plotMulModel(data[data.columns[Workload_Par]],data[enerGPU_results],enerGPU_results)$



Discretization Power Consumption per Unit Time:



Ref. Gustavo Rostollini et al. GreenHPC: A Novel Framework to Measure Energy Consumption on HPC applications. 2015

Equations to Calculate the Energy Consumption:

$$Energy = \int Power(t) dt \quad (3.1)$$

$$Energy = \sum_{i=1}^N Power(i) * \Delta t(i) \quad (3.2)$$

$$E_{solution} = \sum_{i=1}^N Power_{Node}(i) * \max(\Delta t_{CPU}(i), \Delta t_{GPU}(i)) \quad (3.3)$$

$$Power_{Node}(i) = \sum_{j=1}^{NC} P_{CPU}(i) + \sum_{j=1}^{NG} P_{GPU}(i) + \sum_{j=1}^{NM} P_{RAM}(i) \quad (3.4)$$

$$\Delta t_{CPU}(N) = \max(time_{comm}, time_{comp}) \quad (3.5)$$

Ref. John A. García H. et al. Energetically Efficient Acceleration EEA-Aware. Degree work to obtain the title of Master of Science in Systems Engineering and Informatics at UIS 2016.
Ref. Jack Dongarra. Algorithmic and Software Challenges for Numerical Libraries at Exascale. 2013.

Step 3: EEA Prediction System

The third level uses the projected multivariable regression model results to see metrics such as time, performance, power consumption, power consumption and performance per watt, which are used to execute the HPL for each combination of computational resources and to calculate the best combination of computational resources, as shown in Figure.

Efficiently Energetic Acceleration | EEA Prediction System

In [25]: eeaPredictionSystem(49152)						
**** EEA Prediction System of Computational Resources for Matrix 49152 to Execute HPL-2.0 on GUANE ****						
	Best Prediction	Time_to_Solution	Performance_Rmax	Power_Consumption	Energy_to_Solution	Performance_per_Watt
block_size	768.000000	1024.000000	768.000000	2048.000000	1024.000000	768.000000
Number_task	64.000000	48.000000	64.000000	24.000000	48.000000	64.000000
Task_size	288.000000	384.000000	288.000000	768.000000	384.000000	288.000000
GPU	3.000000	5.000000	5.000000	1.000000	3.000000	3.000000
Cores	4.000000	2.000000	2.000000	12.000000	4.000000	4.000000
Time_to_Solution	83.730042	13.185518	56.792579	751.901756	50.122982	83.730042
Performance_Rmax	662.169353	640.145840	674.374812	118.395452	627.940482	662.169353
Power_Consumption	546.218050	609.774292	607.899581	356.528464	548.092760	546.218050
Energy_to_Solution	55.832045	41.345681	65.305453	241.895380	31.870272	55.832045
Performance_per_Watt	1.218755	1.080187	1.151112	0.406359	1.147830	1.218755

EEA-Aware Scheme for Scalability and Portability of Large Scale Applications on Heterogeneous Architectures:

