



Resource aggregation for task-based Cholesky Factorization on top of heterogeneous machines

Terry Cojean, Abdou Guermouche, Andra Hugo, Raymond Namyst,
Pierre-André Wacrenier

► To cite this version:

Terry Cojean, Abdou Guermouche, Andra Hugo, Raymond Namyst, Pierre-André Wacrenier. Resource aggregation for task-based Cholesky Factorization on top of heterogeneous machines. HeteroPar'2016 workshop of Euro-Par, Aug 2016, Grenoble, France. hal-01181135v3

HAL Id: hal-01181135

<https://inria.hal.science/hal-01181135v3>

Submitted on 23 Aug 2016

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.



Distributed under a Creative Commons Attribution 4.0 International License

Resource aggregation for task-based Cholesky Factorization on top of heterogeneous machines [★]

T. Cojean¹, A. Guermouche¹, A. Hugo², R. Namyst¹, P.A. Wacrenier¹

¹ INRIA, LaBRI, University of Bordeaux, Talence, France
`firstname.lastname@inria.fr`

² University of Uppsala, Sweden
`andra.hugo@it.uu.se`

Abstract. Hybrid computing platforms are now commonplace, featuring a large number of CPU cores and accelerators. This trend makes balancing computations between these heterogeneous resources performance critical. In this paper we propose *aggregating several CPU cores* in order to execute larger parallel tasks and thus improve the load balance between CPUs and accelerators. Additionally, we present our approach to exploit internal parallelism within tasks. This is done by combining two runtime systems: one runtime system to handle the task graph and another one to manage the internal parallelism. We demonstrate the relevance of our approach in the context of the dense Cholesky factorization kernel implemented on top of the StarPU task-based runtime system. We present experimental results showing that our solution outperforms state of the art implementations.

Keywords: Multicore; accelerator; GPU; heterogeneous computing; task DAG; runtime system; dense linear algebra; Cholesky

1 Introduction

Due to recent evolution of High Performance Computing architectures toward massively parallel heterogeneous multicore machines, many research efforts have been devoted to the design of runtime systems able to provide programmers with portable techniques and tools to exploit such hardware complexity. The availability of mature implementations of task based runtime systems (*e.g.* OpenMP or Intel TBB for multicore machines, PaRSEC [6], Charm++ [12], KAAPI [9], StarPU [4] or StarSs [5] for heterogeneous configurations) has allowed programmers to rely on dynamic schedulers and develop powerful implementations of parallel libraries (*e.g.* Intel MKL³, DPLASMA [7]).

However one of the main issues encountered when trying to exploit both CPUs and accelerators is that these devices have very different characteristics

[★] This work is supported by the French National Research Agency (ANR), under the grant ANR-13-MONU-0007

³ <https://software.intel.com/en-us/intel-mkl>

and requirements. Indeed, GPUs typically exhibit better performance when executing kernels applied to large data sets, which we call *coarse grain kernels* (or tasks) in the remainder of the paper. On the contrary, regular CPU cores typically reach their peak performance with fine grain kernels working on a reduced memory footprint.

To work around this granularity problem, task-based applications running on such heterogeneous platforms typically adopt a medium granularity, chosen as a trade-off between coarse-grain and fine-grain kernels. A small granularity would indeed lead to poor performance on the GPU side, whereas large kernel sizes may lead to an under-utilization of CPU cores because (1) the amount of parallelism (*i.e.* task graph width) decreases when kernel size increases and (2) the efficiency of GPU increases while a large memory footprint may penalize CPU cache hit ratio. This trade-off technique is typically used by dense linear algebra hybrid libraries [14,2,7]. The main reason for using a unique task granularity in the application lies in the complexity of the algorithms dealing with heterogeneous task granularities even for very regular applications like dense linear libraries. However some recent approaches relax this constraint and are able to split coarse-grain tasks at run time to generate fine-grain tasks for CPUs [17].

The approach we propose in this paper to tackle the granularity problem is based on resource aggregation: instead of dynamically splitting tasks, we rather aggregate resources to process coarse grain tasks in a parallel manner on the critical resource, the CPU. To deal with Direct Acyclic Graphs (DAGs) of parallel tasks, we have enhanced the StarPU [4,10] runtime system to cope with parallel tasks, the implementation of which relies on another parallel runtime system (*e.g.* OpenMP). This approach allows us to delegate the division of the kernel between resources to a specialized library. We illustrate how state of the art scheduling heuristics are upgraded to deal with parallel tasks. Although our scheme is able to handle arbitrary clusters, we evaluate our solution with fixed-size ones. We show that using our solution for a dense Cholesky factorization kernel outperforms state of the art implementations to reach a peak performance of 4.6 Tflop/s on a platform equipped with 24 CPU cores and 4 GPU devices.

2 Related Work

A number of research efforts have recently been focusing on redesigning HPC applications to use dynamic runtime systems. The dense linear algebra community has massively adopted this modular approach over the past few years [14,2,7] and delivered production-quality software relying on it. For example, the MAGMA library [2], provides Linear Algebra algorithms over heterogeneous hardware and can optionally use the StarPU runtime system to perform dynamic scheduling between CPUs and GPUs, illustrating the trend toward delegating scheduling to the underlying runtime system. Moreover, such libraries often exhibit state-of-the-art performance, resulting from heavy tuning and strong optimization efforts. However, these approaches require that accelerators process a large share of the total workload to ensure a fair load balancing between resources. Additionally,

all these approaches rely on an uniform tile size, consequently, all tasks have the same granularity independently from where they are executed leading to a loss of efficiency of both the CPUs and the accelerators.

Recent attempts have been made to resolve the granularity issue between regular CPUs and accelerators in the specific context of dense linear algebra. Most of these efforts rely on heterogeneous tile sizes [15] which may involve extra memory copies when split data need to be coalesced again [11]. However the decision to split a task is mainly made statically at submission time. More recently, a more dynamic approach has been proposed in [17] where coarse grain tasks are hierarchically split at runtime when they are executed on CPUs. Although this paper succeeds at tackling the granularity problem, the proposed solution is specific to linear algebra kernels. In the context of this paper, we tackle the granularity problem with the opposite point of view and a more general approach: rather than splitting coarse grained tasks, we aggregate computing units which cooperate to process the task in parallel. By doing so, our runtime system does not only support sequential tasks but also parallel ones.

However, calling simultaneously several parallel procedures is a difficult matter because usually they are not aware of the resource utilization of one another and they may thus oversubscribe threads to the processing units. This issue has been first tackled within the Lithe framework [13] a resource sharing management interface that defines how threads are transferred between parallel libraries within an application. This contribution suffered from the fact that it does not allow to dynamically change the number of resources associated with a parallel kernel. Our contribution in this study is a generalization of a previous work [10], where we introduced the so-called scheduling contexts which aim at structuring the parallelism for complex applications. Actually, our runtime system is able to cope with several flavors of inner parallelism (OpenMP, Pthreads, StarPU) simultaneously. In this paper, we showcase the use of OpenMP to manage internal task parallelism.

3 Background

We integrate our solution to the StarPU runtime system as it provides a flexible platform to deal with heterogeneous architectures. StarPU [4] is a C library that provides programmers with a portable interface for scheduling dynamic graphs of tasks onto a heterogeneous set of processing units called workers in StarPU (*i.e.* CPUs and GPUs). The two basic principles of StarPU are firstly that tasks can have several implementations, for some or each of the various heterogeneous processing units available in the machine, and secondly that necessary data transfers to these processing units are handled transparently by the runtime system. StarPU tasks are defined as multi-version kernels, gathering the different implementations available for CPUs and GPUs, associated to a set of input/output data. To avoid unnecessary data transfers, StarPU allows multiple copies of the same registered data to reside at the same time in different memory locations as long as it is not modified. Asynchronous data prefetching

is also used to hide memory latencies allowing to overlap memory transfers with computations when possible.

StarPU is a platform for developing, tuning and experimenting with various task scheduling policies in a portable way. Several built-in schedulers are available, ranging from greedy and work-stealing based policies to more elaborated schedulers implementing variants of the Minimum Completion Time (MCT) policy [16]. This latter family of schedulers is based on auto-tuned history-based performance models that provide estimations of the expected lengths of tasks and data transfers. The performance model of StarPU also supports the use of regressions to cope with dynamic granularities.

4 A runtime solution to deal with nested parallelism

We introduce a set of mechanisms which aim at managing nested parallelism (i.e. task inner parallelism) within the StarPU runtime system. We consider the general case where a parallel task may be implemented on top of any runtime system. We present in Figure 1a the standard architecture of a task-based runtime system where the task-graph is provided to the runtime and the ready tasks (in purple) are dynamically scheduled on queues associated with the underlying computing resources. We propose a more flexible scheme where tasks may feature *internal parallelism* implemented using any other runtime system. This idea is represented in Figure 1b where multiple CPU devices are grouped to form *virtual resources* which will be referred to as *clusters*: in this example, each cluster contains 3 CPU cores. We will refer to the main runtime system as the *external runtime system* while the runtime system used to implement parallel tasks will be denoted as the *inner runtime system*. The main challenges regarding this architecture are: 1) how to constrain the inner runtime system’s execution to the selected set of resources, 2) how to extend the existing scheduling strategies to this new type of computing resources, and 3) how to define the number of *clusters* and their corresponding resources. In this paper, we focus on the first two problems since the latter is strongly related to online moldable/malleable task scheduling problems which are out of the scope of this paper.

Firstly, we need to aggregate cores into a cluster. This is done thanks to a simple programming interface which allows to group cores in a compact way with respect to memory hierarchy. In practice, we rely on the `hwloc` framework [8], which provides the hardware topology, to build clusters containing every computing resource under a given level of the memory hierarchy (*e.g.* Socket, NUMA node, L2 cache, ...). Secondly, forcing a parallel task to run on the set of resources corresponding to a cluster depends on whether or not the inner runtime system has its own pool of threads. On the one hand, if the inner runtime system offers a multithreaded interface, that is to say the execution of the parallel task requires a call that has to be done by each thread, the inner runtime system can directly use the StarPU workers assigned to the cluster. We show in Figure 2a how we manage internal SPMD runtime systems. In this case, the parallel task is inserted in the local queue of each StarPU worker. On the other hand, if the

inner runtime system features its own pool of threads (*e.g.* as most OpenMP implementations), StarPU workers corresponding to the cluster need to be paused until the end of the parallel task. This is done to avoid oversubscribing threads over the underlying resources. We describe in Figure 2b how the interaction is managed. We allow only one StarPU worker to keep running. This latter called the *master worker* of the cluster, is in charge of popping the tasks assigned to the cluster by the scheduler. When tasks have to be executed, the master worker takes the role of a regular application thread with respect to the inner runtime system. In Figure 2b, the black threads represent the StarPU workers and the pink ones the inner runtime system (*e.g.* OpenMP) threads. The master worker joins the team of inner threads while the other StarPU threads are paused.

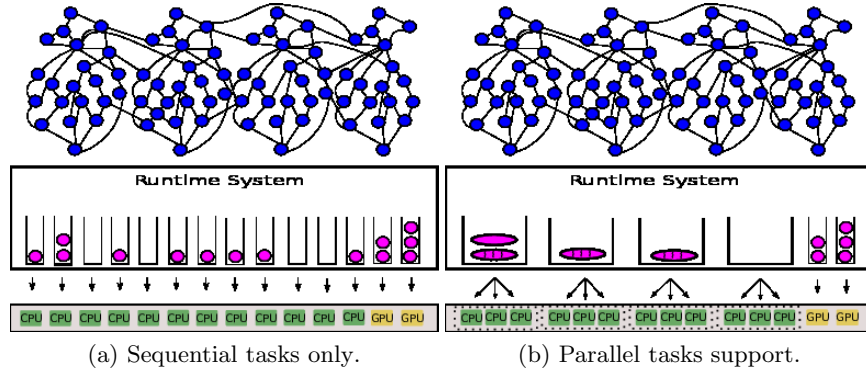


Fig. 1: Managing internal parallelism within StarPU.

Depending on the inner scheduling engine, the set of computing resources assigned to a cluster may be dynamically adjusted during the execution of a parallel task. This obviously requires the inner scheduler (resp. runtime system) to be able to support such an operation. For instance, parallel kernels implemented on top of runtime systems like OpenMP will not allow removing a computing resource during the execution of the parallel task. In this case we refer to the corresponding parallel task as a *moldable* one and we consider resizing the corresponding cluster only at the end of the task or before starting a new one.

From a practical point of view, we integrate in a *callback* function the specific code required to force the inner runtime to run on the selected set of resources. This prologue is transparently triggered before starting executing any sequence of parallel tasks. We call this callback the *prologue callback*. This approach can be used for most inner runtime systems as the programmer can provide the implementation of the prologue callback and thus use the necessary functions in order to provide the resource allocation required for the corresponding cluster. Such a runtime should however respect certain properties: be able to be executed on a restricted set of resources and allow the privatization of its global and static variables. From the user point of view, provided that he has parallel implemen-

tation of his kernels, using clusters in his application is straightforward: he needs to implement the callback and create clusters. In the experimental section, we use this approach to force the MKL library, which relies on OpenMP, to run on the set of resources corresponding to the clusters.

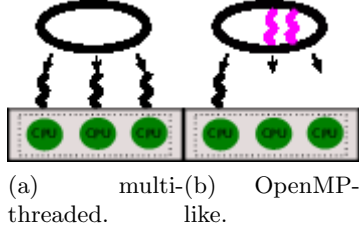


Fig.2: Management of the pool of threads within a cluster.

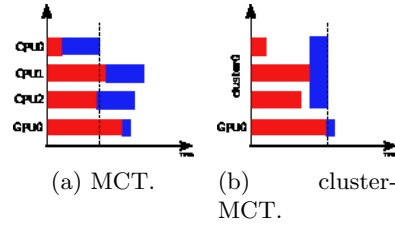


Fig.3: Adaptation of the MCT scheduling strategy.

4.1 Adapting MCT and performance models for parallel tasks

As presented in Section 3, MCT is a scheduling policy implemented in StarPU. The task's estimated length and transfer time used for MCT decisions is computed using performance prediction models. These models are based on performance history tables dynamically built during the application execution. It is then possible for the runtime system to predict for each task the worker which completes it at the earliest. Therefore, even without the programmer's involvement, the runtime can provide a relatively accurate performance estimation of the expected requirements of the tasks allowing the scheduler to take appropriate decisions when assigning tasks to computing resources.

As an illustration, we provide in Figure 3a an example showing the behavior of the MCT strategy. In this example, the blue task represents the one the scheduler is trying to assign. This task has different length on CPU and GPU devices. The choice is then made to schedule it on the CPU0 which completes it first. We have adapted the MCT strategy and the underlying performance models to be able to select a pool of CPUs when looking for a computing resource to execute a task. We have thus introduced a new type of resource: the cluster of CPUs. The associated performance model is parametrized not only by the size and type of the task together with the candidate resource but also by the number of CPUs forming the cluster. Thus, tasks can be assigned to a cluster either explicitly by the user or by the policy depending on where it would finish first. This is illustrated in Figure 3b, where the three CPUs composing our platform are grouped in a cluster. We can see that the expected length of the parallel task on the cluster is used to choose the resource having the minimum completion time for the task. Note that in this scenario, we chose to illustrate a cluster with an OpenMP-like internal runtime system.

This approach permits to deal with a heterogeneous architecture made of different types of processing units as well as clusters grouping different sets of

processing units. Therefore, our approach is able to deal with multiple clusters sizes simultaneously with clusters of one CPU core and take appropriate decisions. Actually, it is helpful to think of the clusters as mini-accelerators. In this work, we let the user define sets of such clusters (mini-accelerators) and schedule tasks dynamically on top of them.

5 Experimental Results

For our evaluation, we use the Cholesky factorization of **Chameleon** [1], a dense linear algebra library for heterogeneous platforms based on the StarPU runtime system. Similarly to most task-based linear algebra libraries, **Chameleon** relies on optimized kernels from a BLAS library. Our adaptation of **Chameleon** does not change the high level task-based algorithms and subsequent DAG. We simply extend the prologue of each task to allow the use of an OpenMP implementation of MKL inside the clusters and manage the creation of clusters. We call **pt-Chameleon** this adapted version of **Chameleon** that handles parallel tasks. The machine we use is heterogeneous and composed of two 12-cores Intel Xeon CPU E5-2680 v3 (@2.5 GHz equipped with 30 MB of cache each) and enhanced with four NVidia K40m GPUs. In StarPU one core is dedicated to each GPU, consequently we report on all figures performance with 20 cores for the **Chameleon** and **pt-Chameleon** versions. We used a configuration for **pt-Chameleon** composed of 2 clusters aggregating 10 cores each (noted 2×10), so that the 10 cores of a CPU belong to a single cluster. In comparison, **Chameleon** uses 20 sequential CPU cores on this platform. Finally, we show on all figures the average performance and observed variation over 5 runs on square matrices.

	DPOTRF		DTRSM		DSYRK		DGEMM	
	960	1920	960	1920	960	1920	960	1920
1 core (Gflop/s)	27.78	31.11	34.42	34.96	31.52	32.93	36.46	37.27
GPU / 1 core	1.72	5.95	8.72	18.59	26.96	31.73	28.80	30.86
10 cores / 1 core	5.55	7.48	6.75	8.48	6.90	8.63	7.77	8.56

Table 1: Acceleration factor of Cholesky factorization kernels on a GPU and 10 cores compared to one core with tile size 960 and 1920.

We report in Table 1 the acceleration factors of using 10 cores or one GPU compared to the single core performance for each kernel of the Cholesky factorization. We conduct our evaluation using MKL for the CPUs and CuBLAS (resp. MAGMA) for the GPUs. This table highlights a sublinear scalability of using 10 cores compared to using 1 core. For example on our best kernel DGEMM we accelerate the execution by a factor of 7.77 when using 10 cores and this increases to 8.56 with a tile size of 1920. Despite this, we can see that relying on sequential kernels worsens the performance gap between the CPUs and GPUs while relying on clusters makes the set of computing resources more homogeneous. We can obtain an acceleration factor of GPU against CPUs by dividing the second line by the third one. For example, the performance gap for the DGEMM kernel with

a tile size of 960 is $\simeq 29$ when using 1 core compared to a GPU whereas it is $28.80/7.77 \simeq 3.7$ when using 10 cores compared to a GPU. As a consequence, if 28 independent DGEMM of size 960 are submitted on computer of 10 cores and a GPU, the **Chameleon** scheduler assigns all the tasks to the GPU whereas **pt-Chameleon** assigns 6 tasks to the 10 core cluster and 22 tasks to GPUs. Another important aspect which can compensate some loss in efficiency is the **pt-Chameleon** ability to accelerate the critical path. Indeed, a cluster of 10 cores can execute the DPOTRF kernel on a tile size of 960 three times faster than on a GPU. The performance is also almost the same for the DTRSM task.

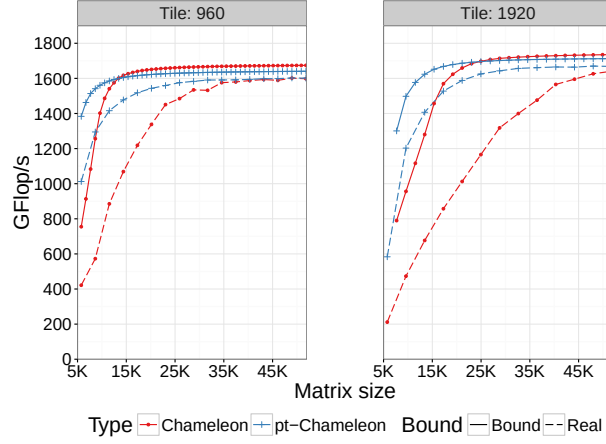


Fig. 4: Comparison of the **pt-Chameleon** and **Chameleon** Cholesky factorization with computed bounds. 20 CPUs and 1 GPU are used.

We show in Figure 4 the performance of the Cholesky factorization for both **Chameleon** and **pt-Chameleon** with multiple tile sizes and their computed make-span theoretical lower bounds. These bounds are computed thanks to the iterative bound technique introduced in [3] which iteratively adds new critical paths until all are taken into account. As these bounds do not take communications with GPU devices into account, they are clearly unreachable in practice. These bounds show that **pt-Chameleon** can theoretically obtain better performance than **Chameleon** on small to medium sized matrices. Indeed, the CPUs are underutilized in the sequential tasks case due to a lack of parallelism whereas using clusters lowers the amount of tasks required to feed the CPU cores. The 5K matrix order point shows a difference of performance of 600 Gflop/s, this is close to the obtainable performance on these CPUs. For both tile sizes on large matrices (*e.g.* 40K), the **Chameleon** bound is over the **pt-Chameleon** one. This is due to the better efficiency of the sequential kernels since the parallel kernels do not exhibit perfect scalability, allowing the CPUs to achieve better performance per core in the sequential case. We observe that for a coarser kernel grain of 1920, the maximum achievable performance is higher, mainly thanks to a better kernel efficiency on GPUs with this size. For DGEMM kernel we can gain close

to 100 Gflop/s (or 10%). We can also note that the gap between **Chameleon** and **pt-Chameleon** bound decreases slightly as we increase the tile size to 1920 thanks to a relatively better gain in efficiency per core compared to the sequential one. Additionally, the real executions are underneath the theoretical bounds. This is due to the fact that transfer time is not taken into account in the bounds. Moreover, the online MCT scheduler can exaggeratedly favor GPUs because of their huge performance bonus in the **Chameleon** case as was shown in [3]. Finally, this figure highlights a constantly superior performance of **pt-Chameleon** over **Chameleon** which achieves up to 65% better performance on a matrix size of 11K for the 960 tile size case and up to 100% better performance on matrices lower than 10K. On those matrix sizes, real **pt-Chameleon** execution is able to go over the theoretical bound of **Chameleon** which demonstrates the superiority of our approach.

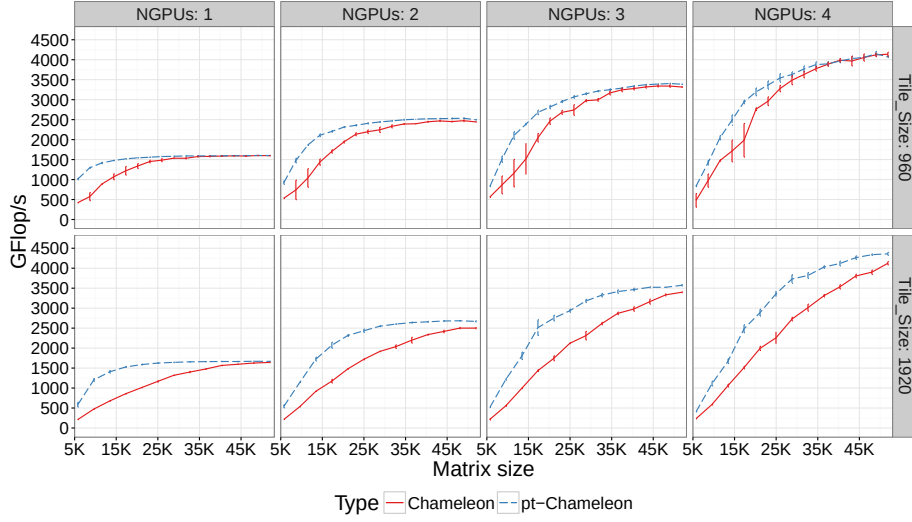


Fig.5: Performance of the Cholesky factorization with **pt-Chameleon** and **Chameleon** with varying number of GPUs and task granularity.

We report in Figure 5 the performance behavior of our implementation of the Cholesky factorization using the **pt-Chameleon** framework, compared to the existing **Chameleon** library. When looking at medium sized matrices we observe that **pt-Chameleon** is able to achieve significantly higher performance than **Chameleon** across all test cases. On those matrices, the **Chameleon** library has some performance variability. This is mainly due to bad scheduling decisions regarding tasks on the critical path in **Chameleon**. Indeed, if an important task is wrongly scheduled on a CPU such as a DPOTRF, we may lack parallelism for a significant period of time. Whereas in the **pt-Chameleon** case using parallel tasks even accelerates the critical path due to a better kernel performance, which makes the approach less sensitive to bad scheduling decisions, lowering **pt-Chameleon**'s variance. Both **Chameleon** and **pt-Chameleon** showcase a good

scalability when increasing the number of GPUs. For example the peak for 1 GPU with a tile size of 960 is at 1.6 Tflop/s and for 2 GPUs it goes up to 2.6 Tflop/s. This improvement is as expected since 1 Tflop/s is the performance of a GPU on this platform with the DGEMM kernel. **Chameleon** scales slightly less than **pt-Chameleon** with a coarse task grain size of 1920. The gap between the two versions increases when increasing the number of GPUs. As shown previously, the scheduler can schedule too many tasks on the GPUs leading to a CPU under-utilization with such a high factor of heterogeneity.

Another factor is the cache behavior of both implementations. Indeed, each processor benefits of 30MB cache and by using one cluster per processor instead of 10 independent processing units we lower by 10 the working set size. Since a tile of 960 weights 7MB whereas a tile of 1920 weights 28MB we are even able to fit entirely a 1920 tile in the LLC. This highlights another constraint: the memory contention bottleneck. We had to explicitly use the *numactl* tool to allocate pages in a round robin fashion on all 4 NUMA nodes, otherwise the behavior of **Chameleon** was very irregular. In fact, even with the interleave setting, we observed that some compute intensive kernels such as DGEMM could become more memory bound for the **Chameleon** case with a matrix size of 43K. To investigate this issue we conducted an experiment using Intel VTune where we allocated the complete matrix on one NUMA node thanks to the *numactl* tool. We saw that for **Chameleon** 59% of the DGEMM kernels were bounded by memory, whereas for **pt-Chameleon** only 13% were bounded by memory. We also observed over two times less cache misses on our **pt-Chameleon** version.

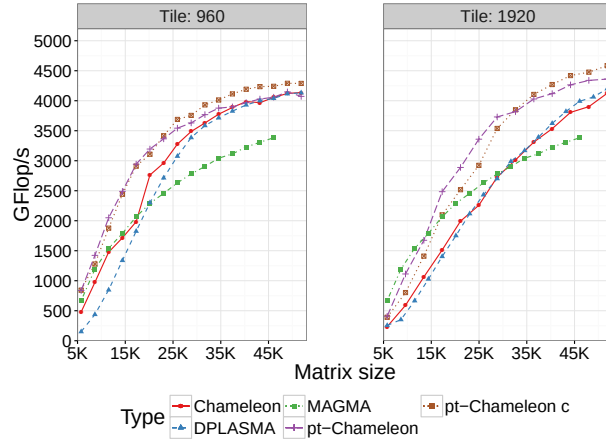


Fig. 6: Comparison of the constrained **pt-Chameleon** with baseline **Chameleon**, **MAGMA** (default parameters and multithreaded MKL) and hierarchical **DPLASMA** (internal blocking of 192 (left) and 320 (right)).

Finally, in Figure 6 we compare **pt-Chameleon** to multiple dense linear algebra reference libraries: **MAGMA**, **Chameleon** and **DPLASMA** using the hierarchical granularity scheme presented in [17]. We make use of a constrained version ($2 \times 10c$) where the DPOTRF and DTRSM tasks are restricted to CPU workers. On

this figure, the **MAGMA** and **DPLASMA** versions use the 24 CPU cores. This strategy is comparable to what is done in [17] where only DGEMM kernels are executed on GPU devices. We observe that using the regular MCT scheduler for small matrices leads to better performance since in the constrained version the amount of work done by CPUs is too large. However, when we increase the matrix size, the constrained version starts to be efficient and leads to a 5% increase in performance on average, achieving a peak of 4.6 Tflop/s on our test platform. We see that the absolute peak is obtained by **pt-Chameleon** and outperforms all the other implementations.

6 Conclusion

One of the biggest challenge raised by the development of high performance task-based applications on top of heterogeneous hardware lies in coping with the increasing performance gap between accelerators and individual cores. One way to address this issue is to use multiple tasks' granularities, but it requires in-depth modifications to the data layout used by existing implementations.

We propose a less intrusive and more generic approach that consists in reducing the performance gap between processing units by forming clusters of CPUs on top of which we exploit tasks' inner parallelism. Performance of these clusters of CPUs can better compete with the one of powerful accelerators such as GPUs. Our implementation extends the StarPU runtime system so that the scheduler only sees virtual computing resources on which it can schedule parallel tasks (*e.g.* BLAS kernels). The implementation of tasks inside such clusters can virtually rely on any thread-based runtime system, and runs under the supervision of the main StarPU scheduler. We demonstrate the relevance of our approach using task-based implementations of the dense linear algebra Cholesky factorization. Our implementation is able to outperform the **MAGMA**, **DPLASMA** and **Chameleon** state-of-the-art dense linear algebra libraries while using the same task granularity on accelerators and clusters.

In the near future, we intend to further extend this work by investigating how to automatically determine the optimal size of clusters. Preliminary experiments show that using clusters of different sizes sometimes leads to significant performance gains. Thus, we plan to design heuristics that could dynamically adapt the number and the size of clusters on the fly, based on statistical information regarding ready tasks.

Acknowledgment We are grateful to Mathieu Faverge for his help for the comparison of **DPLASMA** and **pt-Chameleon**. Experiments presented in this paper were carried out using the PLAFRIM experimental testbed.

References

1. Agullo, E., Augonnet, C., Dongarra, J., Ltaief, H., Namyst, R., Thibault, S., Tomov, S.: A hybridization methodology for high-performance linear algebra software for gpus. *GPU Computing Gems, Jade Edition* 2, 473–484 (2011)

2. Agullo, E., Demmel, J., Dongarra, J., Hadri, B., Kurzak, J., Langou, J., Ltaief, H., Luszczek, P., Tomov, S.: Numerical linear algebra on emerging architectures: The PLASMA and MAGMA projects. *Journal of Physics: Conference Series* 180(1) (2009)
3. Agullo, E., Beaumont, O., Eyraud-Dubois, L., Kumar, S.: Are Static Schedules so Bad ? A Case Study on Cholesky Factorization. In: *Proceedings of IPDPS'16* (2016)
4. Augonnet, C., Thibault, S., Namyst, R., Wacrenier, P.A.: StarPU: A Unified Platform for Task Scheduling on Heterogeneous Multicore Architectures. *Concurr. Comput. : Pract. Exper.* 23, 187–198 (Feb 2011)
5. Ayguadé, E., Badia, R., Igual, F., Labarta, J., Mayo, R., Quintana-Ortí, E.: An Extension of the StarSs Programming Model for Platforms with Multiple GPUs. In: *Euro-Par 2009*. pp. 851–862 (2009)
6. Bosilca, G., Bouteiller, A., Danalis, A., Herault, T., Lemarinier, P., Dongarra, J.: Dague: A generic distributed dag engine for high performance computing. *Parallel Computing* 38(Issues 1), 37 – 51 (2012)
7. Bosilca, G., Bouteiller, A., Danalis, A., Herault, T., Luszczek, P., Dongarra, J.: Dense linear algebra on distributed heterogeneous hardware with a symbolic dag approach. *Scalable Computing and Communications: Theory and Practice* (2013)
8. Broquedis, F., Clet-Ortega, J., Moreaud, S., Furmento, N., Goglin, B., Mercier, G., Thibault, S., Namyst, R.: hwloc: A generic framework for managing hardware affinities in HPC applications. In: *Proceedings of the 2010 18th Euromicro Conference on Parallel, Distributed and Network-based Processing*. pp. 180–186. PDP '10, IEEE Computer Society, Washington, DC, USA (2010), <http://dx.doi.org/10.1109/PDP.2010.67>
9. Hermann, E., Raffin, B., Faure, F., Gautier, T., Allard, J.: Multi-gpu and multi-cpu parallelization for interactive physics simulations. In: *Euro-Par 2010 - Parallel Processing*, vol. 6272, pp. 235–246 (2010)
10. Hugo, A., Guermouche, A., Wacrenier, P., Namyst, R.: Composing multiple starpu applications over heterogeneous machines: A supervised approach. *IJHPCA* 28(3), 285–300 (2014)
11. Kim, K., Eijkhout, V., van de Geijn, R.A.: Dense matrix computation on a heterogeneous architecture: A block synchronous approach. Tech. Rep. TR-12-04, Texas Advanced Computing Center, The University of Texas at Austin (2012)
12. Kunzman, D.M., Kalé, L.V.: Programming heterogeneous clusters with accelerators using object-based programming. *Scientific Programming* 19(1), 47–62 (2011)
13. Pan, H., Hindman, B., Asanović, K.: Composing parallel software efficiently with lithe. *SIGPLAN Not.* 45, 376–387 (June 2010), <http://doi.acm.org/10.1145/1809028.1806639>
14. Quintana-Ortí, G., Quintana-Ortí, E.S., van de Geijn, R.A., Zee, F.G.V., Chan, E.: Programming matrix algorithms-by-blocks for thread-level parallelism. *ACM Trans. Math. Softw.* 36(3) (2009)
15. Song, F., Tomov, S., Dongarra, J.: Enabling and scaling matrix computations on heterogeneous multi-core and multi-gpu systems. In: *Proceedings of ICS'12*. pp. 365–376 (2012)
16. Topcuoglu, H., Hariri, S., Wu, M.Y.: Performance-effective and low-complexity task scheduling for heterogeneous computing. *Parallel and Distributed Systems, IEEE Transactions on* 13(3), 260–274 (Mar 2002)
17. Wu, W., Bouteiller, A., Bosilca, G., Faverge, M., Dongarra, J.: Hierarchical dag scheduling for hybrid distributed systems. In: *29th IEEE International Parallel & Distributed Processing Symposium (IPDPS)*. Hyderabad, India (May 2015)