

Task-parallel Analysis of Molecular Dynamics Trajectories

Ioannis Paraskevatos
Rutgers University
Piscataway, New Jersey, USA

Andre Luckow
Ludwig-Maximilians-University
Munich, Germany

Mahzad Khoshlessan
Arizona State University
Tempe, Arizona, USA

George Chantzalexou
Rutgers University
Piscataway, New Jersey, USA

Thomas E. Cheatham
University of Utah
Salt Lake City, Utah, USA

Oliver Beckstein
Arizona State University
Tempe, Arizona, USA

Geoffrey C. Fox
Indiana University
Bloomington, Indiana, USA

Shantenu Jha
Rutgers University and Brookhaven
National Laboratory.

ABSTRACT

Different parallel frameworks for implementing data analysis applications have been proposed by the HPC and Big Data communities. In this paper, we investigate three task-parallel frameworks: Spark, Dask and RADICAL-Pilot with respect to their ability to support data analytics on HPC resources and compare them with MPI. We investigate the data analysis requirements of Molecular Dynamics (MD) simulations which are significant consumers of supercomputing cycles, producing immense amounts of data. A typical large-scale MD simulation of a physical system of $O(100k)$ atoms over μsecs can produce from $O(10)$ GB to $O(1000)$ GBs of data. We propose and evaluate different approaches for parallelization of a representative set of MD trajectory analysis algorithms, in particular the computation of path similarity and leaflet identification. We evaluate Spark, Dask and RADICAL-Pilot with respect to their abstractions and runtime engine capabilities to support these algorithms. We provide a conceptual basis for comparing and understanding different frameworks that enable users to select the optimal system for each application. We also provide a quantitative performance analysis of the different algorithms across the three frameworks.

KEYWORDS

Data analytics, MD Simulations Analysis, Parallel MD analysis, task-parallel

ACM Reference Format:

Ioannis Paraskevatos, Andre Luckow, Mahzad Khoshlessan, George Chantzalexou, Thomas E. Cheatham, Oliver Beckstein, Geoffrey C. Fox, and Shantenu Jha. 2018. Task-parallel Analysis of Molecular Dynamics Trajectories. In *Proceedings of 47th International Conference on Parallel Processing, August 13-16, 2018, University of Oregon, Eugene, Oregon, USA (ICPP)*. ACM, New York, NY, USA, 10 pages. <https://doi.org/>

Permission to make digital or hard copies of part or all of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. Copyrights for third-party components of this work must be honored. For all other uses, contact the owner/author(s).

ICPP, 2018

© 2018 Copyright held by the owner/author(s).

ACM ISBN .

<https://doi.org/>

1 INTRODUCTION

Different frameworks for implementing parallel data analysis have been developed by the HPC (MPI, OpenMP) and Big Data communities (Spark, Dask) [1]. MPI is the most widely used programming abstraction on HPC resources. It assumes a SPMD execution model where each process executes the same program and it is highly optimized for high-performance computing and communication, which along with synchronization need explicit implementation. Big Data frameworks utilize higher-level MapReduce [2] programming models avoiding explicit implementations of communication. In addition, the MapReduce [2] abstraction makes it easy to exploit data-parallelism as required by many analytics applications. Several recent publications applied HPC techniques to advance traditional Big Data applications and Big Data frameworks [1]. The application of Big Data frameworks to HPC analytics applications has received less attention.

Task-parallel applications are parallelized by partitioning a workload into a set of self-contained units of compute, which require minimal communication. Depending on the application, these tasks can be independent, i. e., with no inter-task communication, or coupled with varying degrees of data and compute dependencies. Big Data applications exploit task parallelism for data-parallel parts (e. g., map operations), but also require some coupling, for computation of aggregates (the reduce operation). The MapReduce [2] abstraction popularized this important execution pattern. Typically, the reduce operation includes shuffling of intermediate data from a set of nodes to node(s) where the reduce operation executes. There is a recognized need to optimize in particular communication-intensive parts of Big Data frameworks using established HPC techniques for interprocess communication, e. g. shuffle operations [3] and other forms of communication [4, 5].

Spark [6] and Dask [7] are two well-known Big Data frameworks. Both provide MapReduce abstractions and are optimized for parallel processing of large data volumes, interactive analytics and machine learning. Their runtime engines can automatically partition data, generate parallel tasks, and execute them on a cluster of nodes. In addition, Spark offers in-memory capabilities allowing caching data that are read multiple times, making it suited for interactive analytics and iterative machine learning algorithms. Dask also provides a MapReduce API (Dask Bags). Furthermore, Dask's

API is more versatile and allows custom workflows and parallel vector/matrix computations.

In this paper, we investigate the data analysis requirements of Molecular Dynamics (MD) applications. MD simulations are significant consumers of supercomputing cycles, producing immense amounts of data. A typical MD simulation of physical system of $O(100k)$ atoms over μsecs can produce from $O(10)$ to $O(1000)$ GBs of data [8]. In addition to being the prototypical HPC application, there is increasingly a need for the analysis to be integrated with simulations and drive the next stages of execution (analysis-driven-simulation) [9]. The attempt does not focus on which approach is better (Big Data vs HPC), but how to provide the "best of both" to a diverse set of applications [4, 10].

We investigate three task-parallel frameworks and their suitability for implementing MD trajectory analysis algorithms. In addition to Spark and Dask, two Big Data frameworks, we investigate RADICAL-Pilot [11], a Pilot-Job [12] framework designed for implementing task-parallel applications on HPC. We utilize MPI4py [13] to provide MPI equivalent implementations of the algorithms. The task-parallel implementations performance and scalability compared to MPI consists the basis of our analysis. MD trajectories are time series of atoms/particles positions and velocities, which are analyzed using different statistical methods to infer certain properties, e.g. the relationship between distinct trajectories, snapshots of a trajectory etc. As a result, they can be considered as a representative set of scientific datasets that are organized as time series and their analysis algorithms. Many of these algorithms can be expressed using simple task abstractions or MapReduce [14, 15]. Thus, the selected frameworks are promising candidates for MD analysis applications.

The paper makes the following contributions: i) it characterizes and explains the behaviour of different MDAnalysis algorithms on these frameworks, and ii) provides a conceptual basis for comparing the abstraction, capabilities and performance of these frameworks.

This paper is organized as follows: Section 2 discusses the Molecular Dynamics analysis algorithms under investigation, and provides a brief characterization based on the Big Data Ogres classification ontology [16]. Section 3, provides a description of the different frameworks that were used for evaluation. Section 4 provides a description of the implementation of the MD algorithms on top of RADICAL-Pilot, Spark and Dask, as well as a performance evaluation and a discussion of findings. Section 5 reviews different MD analysis frameworks in particular with respect to their ability to support scalable analytics of large volume MD trajectories. The paper concludes with a summary and discussion of future work in section 6.

2 MOLECULAR DYNAMICS ANALYSIS APPLICATIONS

Some of the commonly used algorithms in the analysis of MD trajectories are Root Mean Square Deviation (RMSD), Pairwise Distances (PD), and Sub-setting [17]. Two more advanced algorithms are Path Similarity Analysis (PSA) [18] and "Leaflet Identification" [19]. RMSD is used to identify the deviation of atom positions between frames produced by a MD simulation. PD and PSA methods calculate distances based on different metrics either between atoms or

trajectories. Sub-setting methods are used to isolate parts of interest of MD simulation. Leaflet identification provides information about groups of lipids in space by identifying the lipid leaflets in lipid bilayer. All these methods, in some way, read and process the whole physical system generated via simulations. The analysis done reduces the data to either a number or matrix.

We discuss in more detail two of these methods and their implementations in MDAnalysis [19, 20]. Specifically, we discuss a Path Similarity Analysis algorithm using the Hausdorff distance and a Leaflet Identification method. In addition, we explore the applications' Ogres Facets and Views [16], which will provide a more systematic characterization.

Big Data Ogres [16] are organized into four classification dimensions, called *views*. The possible features of a view are called *facets*. A combination of facets from all views defines an Ogre. The four views are: (1) execution, (2) data source & style, (3) processing and (4) problem architecture. The execution view describes aspects, such as I/O, memory, compute ratios, whether computations are iterative, and the 5 V's of Big Data (Volume, Velocity, Value, Variety and Veracity). The data source & style view discusses the way input data are collected, stored and accessed. The processing view describes algorithms and kernels used for computation. The problem architecture view, describes the application architecture needed to support the application.

2.1 MDAnalysis

MDAnalysis is a Python library [19, 20] that provides a comprehensive environment for filtering, transforming and analyzing MD trajectories in all commonly used file formats. It provides a common object-oriented API to trajectory data and leverages existing libraries in the scientific Python software stack, such as NumPy [21] and SciPy [22].

2.1.1 Path Similarity Analysis (PSA): Hausdorff Distance. Path Similarity Analysis (PSA) [18] is used to quantify the similarity between trajectories taking into account their full atomic detail. The underlying idea is to compute pair-wise distances (for instance, using the Hausdorff metric [23]) between members of an ensemble of trajectories (Algorithm 1) and cluster the trajectories based on their distance matrix. The PSA approach was used to compare different path sampling methods [18] and is a general approach to cluster paths that exist in high-dimensional spaces.

Each trajectory is represented as a two dimensional array. The first dimension correspond to the time frames of the trajectory, the second to the N atom positions, in 3-dimensional space, represented as a $3N$ configuration space vector.

Algorithm 1 describes the PSA algorithm with the Hausdorff metric over multiple trajectories. We apply a 2-dimensional data partitioning over the output matrix to parallelize, which is shown in algorithm 2. Our Hausdorff metric calculation is based on a naive algorithm. Recently, an algorithm was introduced that uses early break to speedup execution [24] although we are not aware of a parallel implementation of this algorithm.

The algorithm is embarrassingly parallel and uses linear algebra kernels for calculations. It has complexity $O(n^2)$ (problem architecture & processing views). Input data volume is medium to large while the output is small. Specific execution environments, such as HPC nodes, and Python arithmetic libraries, e.g., NumPy, are

Algorithm 1 Path Similarity Algorithm: Hausdorff Distance

```

1: procedure HAUSDORFFDISTANCE( $T_1, T_2$ ) ▷  $T_1$  and  $T_2$  are a set of 3D points
2:   List  $D_1, D_2$ 
3:   for  $\forall frame_1$  in  $T_1$  do
4:     for  $\forall frame_2$  in  $T_2$  do
5:       Append in  $D_1$   $d_{RMS}(frame_1, frame_2)$ 
6:     end for
7:    $D_{t1}$  append  $\min(D_1)$ 
8:   end for
9:   for  $\forall frame_2$  in  $T_2$  do
10:    for  $\forall frame_1$  in  $T_1$  do
11:      Append in  $D_2$   $d_{RMS}(frame_2, frame_1)$ 
12:    end for
13:     $D_{t2}$  append  $\min(D_2)$ 
14:  end for
15:  return  $\max(\max(D_{t1}), \max(D_{t2}))$ 
16: end procedure
17:
18: procedure PSA( $Traj$ ) ▷  $Traj$  is a set of trajectories
19:   for  $\forall T_1$  in  $Traj$  do
20:     for  $\forall T_2$  in  $Traj$  do
21:        $D(T_1, T_2) = \text{HausdorffDistance}(T_1, T_2)$ 
22:     end for
23:   end for
24:   return  $D$ 
25: end procedure

```

Algorithm 2 Two Dimensional Partitioning

- 1: Initially, there are N^2 distances, where N is the number of trajectories. Each distance defines a computation task.
- 2: Map the initial set to a smaller set with $k = N/n_1$ elements, where n_1 is a divisor of N , by grouping n_1 by n_1 elements together.
- 3: Execute over the new set with k^2 tasks. Each task is the comparisons between n_1 and n_1 elements of the initial set. They are executed serially.

Algorithm 3 Leaflet Finder Algorithm

```

1: procedure LEAFLETFINDER( $Atoms, Cutoff$ ) ▷  $Atoms$  is a set of 3D points that represent the position of atoms in space.  $Cutoff$  is an Integer Number
2:   Graph  $G = (V = Atoms, E = \emptyset)$ 
3:   for  $\forall atom$  in  $Atoms$  do
4:      $N = \{a \in V : d(a, atom) \leq Cutoff\}$ 
5:     Add edges  $[(atom, a) : a \in N]$  in  $G$ 
6:   end for
7:    $C = \text{ConnectedComponents}(G)$ 
8:   return  $C$ 
9: end procedure

```

used (execution view). Input data are produced by HPC simulations, and are typically stored on HPC storage systems, such as parallel filesystem like Lustre (data source & style view).

2.1.2 Leaflet Finder. Algorithm 3 describes the Leaflet Finder algorithm as presented in Ref. [19]. It solves the problem to assign particles to one of two curved but locally approximately parallel sheets, provided that the nearest neighbor inter-particle distance between particles in a sheet is smaller than the distance between sheets. In biomolecular simulations of lipid membranes, which consist of a double layer of lipid molecules, it is used to identify the lipids in the outer and inner leaflets (sheets) from trajectory information. The algorithm consists of two stages: a) a graph connecting particles based on threshold distance (cutoff) is constructed, and b) the connected components of the graph are computed determining the lipids located on the outer leaflet and those on the inner leaflet.

The application comprises of multiple stages with different complexities: The first stage identifies neighboring atoms. There are different implementation alternatives: (1) computing the distance between all atoms ($O(n^2)$); (2) utilizing a tree-based nearest neighbor (Construction: $O(n \log n)$, Query: $O(\log n)$). In both alternatives

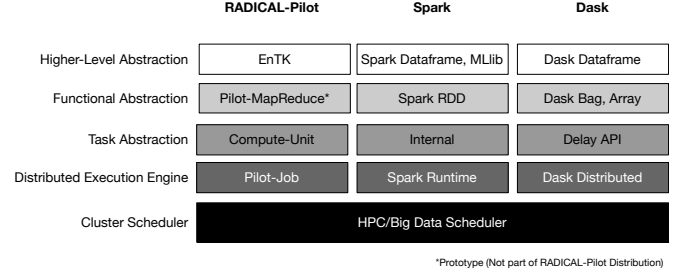


Figure 1: Architecture of RADICAL-Pilot, Spark and Dask: The frameworks share common architectural components for managing cluster resource, managing task. Spark and Dask offer several high-level abstractions inspired by MapReduce.

the input data volume is medium size and the output of this stage is smaller than the input. The complexity of connected components is: $O(|V| + |E|)$ (V : Vertices, E : Edges), i.e. it greatly depends on the characteristics of the graph (in particular its sparsity).

The application typically uses HPC nodes as the execution environment, NumPy arrays (execution view). It uses matrices to represent the physical system and the distance matrix. The output data representation is a graph. The Leaflet Finder can be efficiently implemented using the MapReduce abstraction (problem architecture view). Furthermore, it uses graph algorithms and linear algebra kernels (processing view facets). The data source & style view facets are the same as the PSA algorithm.

3 BACKGROUND OF EVALUATED FRAMEWORKS

The landscape of frameworks for data-intensive applications is manifold [1, 10] and has been extensively studied in the context of scientific [25] applications. In this section, we investigate the suitability of frameworks such as Spark [6], Dask [7] and RADICAL-Pilot [11], for molecular dynamics data analytics.

MapReduce [2] and its open source Hadoop implementation combined a simple functional API with a powerful distributed computing engine that exploits data locality to allow optimized I/O performance. Although, MapReduce is inefficient for interactive workloads and iterative machine learning algorithms [6, 26]. Spark and Dask provide richer APIs, caching and other capabilities critical for analytics applications. RADICAL-Pilot is a Pilot-Job framework [12] that supports task-level parallelism on HPC resources. It successfully adds a parallelization level on top of HPC MPI-based applications.

As described in [10], these frameworks typically comprise of distinct layers, e.g., to access the cluster scheduler, framework-level scheduling, and higher-level abstractions. On top of these low-level resource management capabilities various higher-level abstractions can be provided, e.g., MapReduce-inspired functional APIs. These then provide the foundation for analytics abstractions, such as Dataframes, Datasets and Arrays. Figure 1 visualizes the different components of RADICAL-Pilot, Spark and Dask. In the following, we describe each framework in detail.

3.1 Spark

Spark [6] extends MapReduce [2] providing a rich set of transformations on top of the Resilient Distributed Dataset (RDD) abstraction [27]. RDDs are cached in memory making Spark well suitable for iterative applications that need to cache a set of data across multiple stages. PySpark provides a Python API to Spark.

A Spark job is compiled of multiple stages; a stage is a set of parallel tasks executed without any communication (e.g., map) and an action (e.g., reduce). Each action defines new stage. The DAGScheduler is responsible for translating the workflow specified via RDD transformations and actions to an execution plan. Spark's distributed execution engine handles the low-level details of task execution based on this plan. The execution of a Spark job is triggered by actions.

Spark can read data from different sources, such as HDFS, blob storage, parallel and local filesystems. While Spark caches loaded data in memory, it offloads to disk when an executor does not have enough free memory to hold all the data of its partition. Persisted RDDs remain in memory, unless specified to use the disk either complementary or as a single target. In addition, Spark writes to disk data that are used in a shuffle. As a result, it allows quick access to those data when transmitted to another executor. Finally, Spark provides a set of actions that allow to write text files, Hadoop sequence files or object files to the local filesystem, HDFS or any filesystem that supports Hadoop. Apart from RDDs, Spark supports higher-level data abstractions for processing structured data, such as dataframes, Spark-SQL, datasets, and data streams.

3.2 Dask

Dask [7] provides a Python-based parallel computing library, which is designed to inter-operate and parallelize native Python code written for NumPy and Pandas. In contrast to Spark, Dask also provides a lower-level task API (the delayed API) that allows the user to construct arbitrary graphs. Being written in Python does not require to translate data types from one language to another in contrast to PySpark for example, which needs to move data from Python's interpreter to Java/Scala and vice versa.

In addition to the low-level task API, Dask offers three higher-level abstractions: Bags, Arrays and Dataframes. Dask Arrays are collection of independent NumPy arrays organized as a grid. Dask Bags are similar to Spark RDDs and are used to analyze semi-structured data, like JSON files. Dask Dataframe is a distributed collection of Pandas dataframes that can be analyzed in parallel.

Furthermore, Dask offers three schedulers: multithreading, multiprocessing and distributed. The multithreaded and multiprocessing schedulers can be used only on a single node and the parallel execution is done through threads or processes respectively. The distributed scheduler creates a cluster with a scheduling process and multiple worker processes. A client process creates and communicates a DAG to the scheduler. The scheduler assigns tasks to workers.

3.3 RADICAL-Pilot

RADICAL-Pilot [11] is a Pilot system that implements the pilot paradigm as outlined in Ref. [28]. RADICAL-Pilot is implemented in Python and provides a well defined API and usage modes. Although RP is vehicle for research in scalable computing, it also supports

	RADICAL-Pilot	Spark	Dask
Languages	Python	Java, Scala, Python, R	Python
Task Abstraction	Task	Map-Task	Delayed
Functional Abstraction	-	RDD API	Bag
Higher-Level Abstractions	EnTK [30]	Dataframe, ML Pipeline, MLlib [32]	Dataframe, Arrays for block computations
Resource Management	Pilot-Job	Spark Execution Engines	Dask Distributed Scheduler
Scheduler	Individual Tasks	Stage-oriented DAG	DAG
Shuffle	-	hash/sort-based shuffle	hash/sort-based shuffle
Limitations	no shuffle, filesystem-based communication	high overheads for Python tasks (serialization)	Dask Array can not deal with dynamic output shapes

Table 1: Frameworks Comparison: Dask and Spark are designed for data-related task, while RADICAL-Pilot focuses on compute-intensive tasks.

production grade science. Currently, it is being used by applications drawn from diverse domains, ranging from earth sciences and biomolecular sciences to high-energy physics. RP can be used as a runtime system by third party workflow or workload management systems [29–31]. In 2017, RADICAL-Pilot was used to support more than 100M core-hours on US DOE, NSF resources (Blue Waters and XSEDE), and European supercomputers (Archer and SuperMUC).

RADICAL-Pilot allows concurrent task execution on HPC resources. The user can define a set of Compute-Units (CU) - the abstraction used to define a task along with its dependencies - which are submitted to RADICAL-Pilot. RADICAL-Pilot schedules these CUs to be executed under the acquired resources. RADICAL-Pilot uses the existing environment of the resource to execute tasks. Any data communication between tasks is done via the use of an underlying shared filesystem, e.g., Lustre. Task execution coordination and communication is done through a database (MongoDB).

3.4 Comparison of Frameworks

Table 1 summarizes the properties of these frameworks with respect to abstractions and runtime properties provided to create and execute parallel data applications.

API and Abstractions. RADICAL-Pilot provides a low-level API for executing tasks onto resources. While this API can be used to implement high-level capabilities, e.g. MapReduce [33], they are not provided out-of-the box. Both Spark and Dask provide such capabilities. Dask's API is generally lower level than Spark's, e.g., it allows specifying arbitrary task graphs. Although, data partition size is automatically decided, in many cases it is necessary to fine-tune parallelism by specifying the number of partitions.

Another important aspect is the availability of high-level abstractions. High-level tools for RADICAL-Pilot, such as Ensemble Toolkit (EnTK) [30], are designed for workflows involving compute-intensive tasks. Spark and Dask already offer a set of high-level data-oriented abstractions, such as Dataframes and ML APIs.

Scheduling. Both Spark and Dask create a Direct Acyclic Graph (DAG) based on operations over data, which is then executed using their execution engine. Spark jobs are separated into stages. Once all tasks in a stage are completed, the scheduler executes the next stage.

Dask's DAGs are represented by a tree where each node is a task. Leaf tasks do not depend on other task for the execution. Dask tasks are executed when their dependencies are satisfied, starting from

leaf tasks. When a task is reached with unsatisfied dependencies, the scheduler executes the dependent task first. Dask’s scheduler does not rely on synchronization points that Spark’s stage-oriented scheduler introduces. RADICAL-Pilot does provide a DAG and requires the execution order and synchronization to be explicitly specified by the user.

Suitability for MDAnalysis Algorithms. Trajectory analysis methods are often embarrassingly. So, they are ideally suited for task management and functional MapReduce APIs. PSA-like methods typically require a single pass over the data and return a set of values that correspond to a relationship between frames or trajectories. They can be expressed as a bag of tasks using a task management API or a map-only application in a MapReduce-style API.

Leaflet Finder is more complex and requires two stages: a) the edge discovery stage, and b) the connected components stage. It is possible to implement Leaflet Finder with a simple task-management API, although the MapReduce programming model allows more efficient implementation with a map for computing and filtering distances and a reduce phase for finding the components. The shuffling required between the map and reduce phase is medium as the number of edges is a fraction of the input data.

4 EXPERIMENTS AND DISCUSSION

In this section, we characterize the performance of RADICAL-Pilot, Spark and Dask compared to MPI4py. In section 4.1 we evaluate the task throughput using a synthetic workload. In sections 4.2 and 4.3 we evaluate the performance of two algorithms from the MDAnalysis library: PSA and Leaflet Finder using different real-world datasets. We investigate: (1) what capabilities and abstractions of the frameworks are needed to efficiently express these algorithms, (2) what architectural approaches can be used to implement these algorithms with these frameworks, and (3) the performance trade-offs of these frameworks.

The experiments were executed on the XSEDE Supercomputers: Comet and Wrangler. SDSC Comet is a 2.7 PFlop/s cluster with 24 Haswell cores/node and 128 GB memory/node (6,400 nodes). TACC Wrangler has 24 Haswell hyperthreading enabled cores/node and 128 GB memory/node (120 nodes); it is optimized for data-intensive computing. Experiments were carried using RADICAL-Pilot and Pilot-Spark [34] extension, which allows to efficiently manage Spark on HPC resources through a common resource management API. We utilize a set of custom scripts to start the Dask cluster. We used RADICAL-Pilot 0.46.3, Spark 2.2.0, Dask 0.14.1 and Distributed 1.16.3. The data presented are means over multiple runs; error bars represent the standard deviation of the sample. We employed up to 10 nodes in Comet and Wrangler.

4.1 Frameworks Evaluation

As data-parallelism often involves a large number of short-running tasks, task throughput is a critical metric to assess the different frameworks. To evaluate the throughput we use zero workload tasks (/bin/hostname). We submit an increasing number of such tasks to RADICAL-Pilot, Spark and Dask and measure the execution time.

For RADICAL-Pilot, all tasks were submitted as bulk CUs. The RADICAL-Pilot backend database was running on the same node to

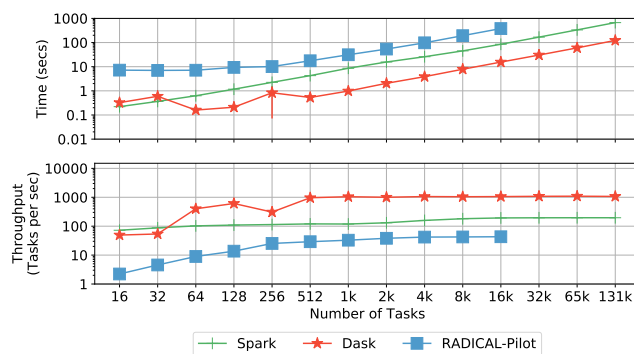


Figure 2: Task Throughput by Framework (Single Node): Time/Throughput executing a given number of zero-workload tasks on Wrangler. Dask performs best; Dask and Spark have very small delays for few tasks. RADICAL-Pilot offers the smallest throughput.

avoid large communication latencies. For Spark, we created an RDD with as many partitions as the number of tasks – each partition is mapped to a task by Spark. For Dask, we created tasks using delayed functions that were executed by the Distributed scheduler. We used Wrangler and Comet for this experiment.

The results are shown in Figure 2. Dask needed the least time to schedule and execute the assigned tasks, followed by Spark and RADICAL-Pilot. Dask and Spark quickly reach their maximum throughput, which is sustained while the number of tasks is increased. RADICAL-Pilot showed the worst throughput and scalability mainly due to some architectural limitations, e.g., the reliance on MongoDB to communicate between Client and Agent. Thus, we were not able to scale RADICAL-Pilot to 32k or more tasks.

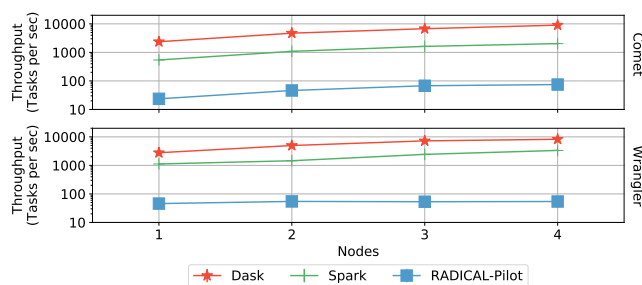


Figure 3: Task Throughput by Framework (Multiple Nodes): Task throughput for 100k zero-workload tasks on different numbers of nodes for each framework. Dask has the largest throughput, followed by Spark and RADICAL-Pilot. Wrangler and Comet show a comparable performance with Comet slightly outperforming Wrangler.

Figure 3 illustrate the throughput when scaling to multiple nodes measured by submitting 100k tasks. Dask’s throughput on all three resources increases almost linearly to the number of nodes. Spark’s throughput is an order of magnitude lower than Dask’s. RADICAL-Pilot’s throughput plateaus at below 100 task/sec.

4.2 Path Similarity Analysis: Hausdorff Distance

The PSA algorithm is embarrassingly parallel and can be implemented using simple task-level parallelism or a map only MapReduce application. The input data, i. e. a set of trajectory files, is equally distributed over the cores, generating one task per core. Each task reads its respective input files in parallel, executes and writes the result to a file.

For RADICAL-Pilot we define a Compute-Unit for each task and execute them using a Pilot-Job. For Spark, we create an RDD with one partition per task. The tasks are executed in a map function. In Dask, the tasks are defined as delayed functions. In MPI, each task is executed by a process.

The experiments were executed on Comet and Wrangler. The dataset used consists of three different atom count trajectories: small (3341 atoms/frame), medium (6682 atoms/frame) and large (13364 atoms/frame), and 102 frames. We used 128 and 256 trajectories of each size.

Figure 4 shows the runtime for 128 and 256 trajectories on Wrangler. Figure 5 compares the execution times on Comet and Wrangler for 128 large trajectories. We see that the frameworks have similar performance on both systems. Furthermore, we see that Wrangler gives smaller speedup than Comet. Although, we used the same number of logical cores, we see that utilizing half the nodes due to hyperthreading results to smaller speedup.

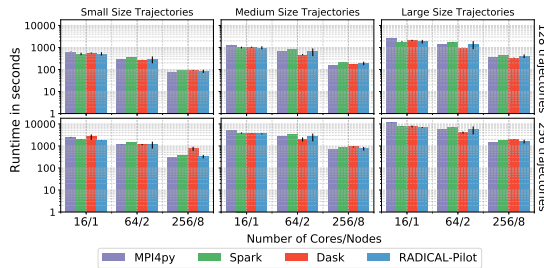


Figure 4: Hausdorff Distance on Wrangler using RADICAL-Pilot, Spark and Dask: Runtimes over different number of cores, trajectory sizes and number of trajectories. All frameworks scaled by a factor of 6 from 16 to 256 cores.

MPI4py, RADICAL-Pilot, Spark and Dask have similar performance when used to execute algorithms that are embarrassingly parallel. All frameworks achieved similar speedups as the number of cores increased, which are lower than MPI4py. Although, the frameworks' overheads are comparably low in relation to the overall runtime, they were significant to impact their speedup. In summary, all three frameworks provide appropriate abstractions and runtime performance, compared to MPI, for embarrassingly parallel algorithms. In this case aspects such as programmability and integrate-ability are the most important considerations, e. g., both RADICAL-Pilot and Dask are native Python frameworks making the integration with MDAnalysis easier and more efficient than with other frameworks, which are based on other languages.

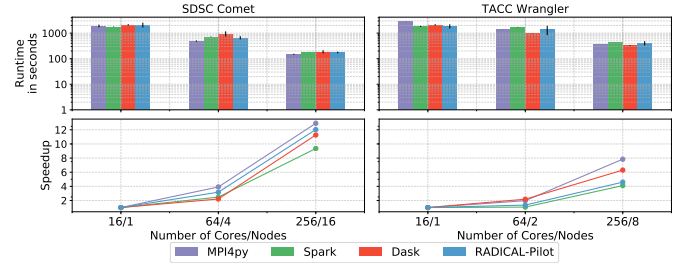


Figure 5: Hausdorff Distance on Comet and Wrangler: Runtime and Speedup for 128 large trajectories.

4.3 Leaflet Finder

In this section, we investigate four different approaches for implementing the Leaflet Finder algorithm using RADICAL-Pilot, Spark, Dask, and MPI4py (see Table 2):

- 1) **Broadcast and 1-D Partitioning:** The physical system is broadcast and partitioned through a data abstraction. Use of RDD API (broadcast), Dask Bag API (scatter), and MPI Bcast to distribute data to all nodes. A map function calculates the edge list using `cdist` from SciPy [22] – realized as a loop iteration for MPI. The list is collected to the master process (gathered to rank 0) and the connected components are calculated.
- 2) **Task API and 2-D Partitioning:** Data management is done without using the data-parallel API. The framework is only used for task scheduling. The data is pre-partitioned in 2-D partitions and passed to a map function that calculates the edge list using `cdist`. The list is collected and the connected components are calculated. In MPI, data are partitioned in the same manner and a loop calculates the edge list, which is then gathered to rank 0.
- 3) **Parallel Connected Components:** Data are managed as in approach 2. Each map task performs edge list and connected components computations. The reduce phase joins the calculated components into one, when there is at least one common node.
- 4) **Tree-based Nearest Neighbor and Parallel-Connected Components (Tree-Search):** This approach is different to approach 3 only on the way edge discovery in the map phase is implemented. A tree-structure containing all atoms is created which is then used to query for adjacent atoms.

We use four physical systems with 131k, 262k, 524k, and 4M atoms with 896k, 1.75M, 3.52M, and 44.6M edges in their graphs. Experimentation was conducted on Wrangler where we utilized up to 256 cores. Data partitioning results into 1024 partitions for each approach, thus 1024 map tasks. Due to memory limitations from using `cdist` – uses double precision floating point – Approach 3 data partitioning of the 4M atom dataset resulted to 42k tasks for both Spark and MPI4py.

Figure 6 shows the runtimes for all datasets for Spark, Dask and MPI4py. RADICAL-Pilot's performance is illustrated in Figure 8. We continue by analyzing the performance of each architectural approach and used framework in detail.

4.3.1 Broadcast and 1-D Partitioning. Approach 1 utilizes a broadcast to distribute the data to all nodes, which is supported by Spark, Dask and MPI. All nodes maintain a complete copy of the dataset. Each map task computes the pairwise distance on its

	Broadcast and 1-D (Approach 1)	Task API and 2-D (Approach 2)	Parallel Connected Components (Approach 3)	Tree-Search (Approach 4)
Data Partitioning	1D	2D	2D	2D
Map	Edge Discovery via Pairwise Distance	Edge Discovery via Pairwise Distance	Edge Discovery via Pairwise Distance and Partial Connected Components	Edge Discovery via Tree-based Algorithm and Partial Connected Components
Shuffle	Edge List ($O(E)$)	Edge List ($O(E)$)	Partial Connected components ($O(n)$)	Partial Connected components ($O(n)$)
Reduce	Connect Components	Connected Components	Joined Connected Components	Joined Connected Components

Table 2: MapReduce Operations used by Leaflet Finder

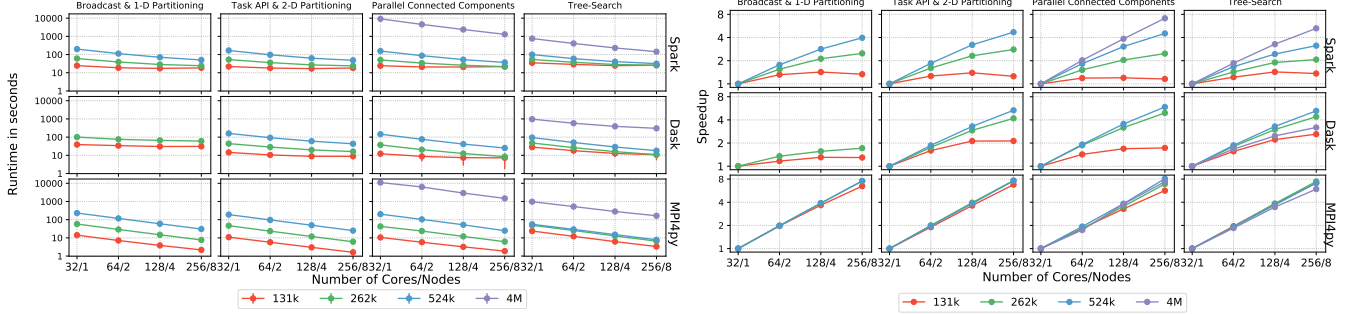


Figure 6: Leaflet Finder: Performance of Different Architectural Approaches for Spark & Dask: Runtimes and Speedups for different system sizes over different number of cores for all approaches and frameworks.

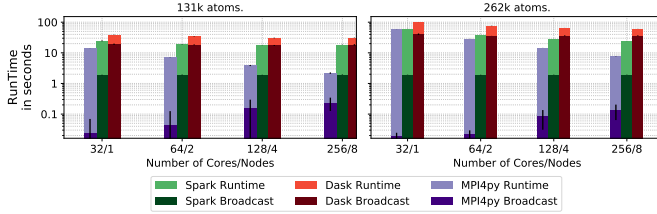


Figure 7: Broadcast and 1-D Partitioned Leaflet Finder (Approach 1): Runtime for multiple system sizes on different number of cores for Spark, Dask and MPI4py.

partition. We use 1-D partitioning. Figure 7 shows the detailed results: as expected the usage of a broadcast has severe limitations for Spark and Dask. MPI's broadcast is a fraction of the overall execution time and significantly smaller than Spark's and Dask's. MPI's broadcast times increase linearly as the number of processes increases, while Spark's and Dask's remain relatively constant for each dataset, due to more elaborate broadcast algorithms compared to. Broadcast times are about 3% – 15% of the edge discovery time for Spark, 40% – 65% for Dask, and < 1% – 10% for MPI4py. Spark offer a more efficient communication subsystem compared to Dask. In addition, Dask broadcast partitions the dataset to a list where each element represents a value from the initial dataset. This lead to not being able to broadcast the 524k atom dataset. Nevertheless, the limited scalability of this approach due to the need to transmit the entire dataset renders it only usable for small datasets. It shows the worst performance and scaling of all approaches for Spark, Dask and MPI4py.

Furthermore, this approach only scales up to 262k atoms for Dask, and 524k atoms for Spark and MPI4py on Wrangler. Spark's performance is comparable to MPI4py for the 262k, and 524k datasets. It also shows better performance for the smallest core count in the 524k case. Dask is at least two times slower than our MPI implementation.

4.3.2 Task-API and 2-D Partitioning. Approach 2 attempts to overcome the limitations of approach 1, especially broadcasting and

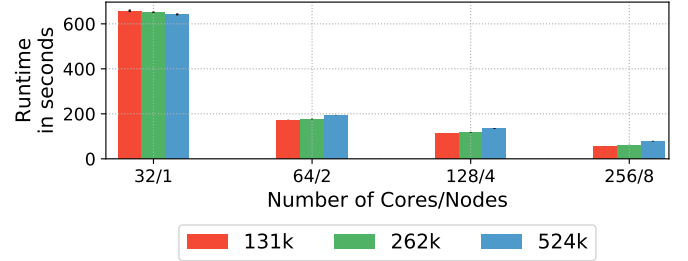


Figure 8: RADICAL-Pilot Task API and 2-D Partitioned Leaflet Finder (Approach 2): Runtime for multiple system sizes over different number of cores using RADICAL-Pilot. RADICAL-Pilot is in the overheads since execution times for the pairwise distance are similar despite the system size.

1-D partitioning. A 2-D block partitioning is essential, as it evenly distributes the compute and more efficiently utilizes the available memory. 2-D partitioning is not well supported by Spark and Dask. Spark's RDDs are optimized for data-parallel applications with 1-D partitioning. While Dask's array supports 2-D block partitioning, it was not used for this implementation. We return the adjacency list of the graph instead of an array to fully use the capabilities of the abstraction. Thus, each task works on a 2-D pre-partitioned part of the input data.

The runtimes of approach 2 are shown in Figure 6 for Spark, Dask, MPI4py and Figure 8 for RADICAL-Pilot. As expected this approach overcomes the limitations of approach 1 and can easily scale to larger datasets (e.g., 524k atoms) while improving the overall runtime. Dask's execution time was smaller by at least a factor of two. However, we were not able to scale this implementation to the 4M dataset, due to memory requirements of `cdist`. For RADICAL-Pilot we observed significant task management overheads (see also section 4.1). This is a limitation of RADICAL-Pilot with respect to managing large numbers of tasks. This is particularly visible when

the scenario was run on a single node with 32 cores. As more resources for the RADICAL-Pilot agent become available, i.e. more than 64 cores, the performance improves dramatically.

Furthermore, Spark and Dask did not scale as well as MPI, which achieved linear speedups of ~ 8 when using 256 cores. Spark and Dask achieved maximum speedups of 4.5 and ~ 5 respectively. Despite this fact, both frameworks had similar performance on 32 cores for the 262k and 524k datasets.

4.3.3 Parallel Connected Components. Another important aspect is the communication between the edge discovery and the connected components stage. For the 524k atoms dataset the output of the edge discovery phase is ≈ 100 MB. To reduce the amount of data that needs to be shuffled, we refined the algorithm to compute the graph components on the partial dataset in the map phase. The partial components are then merged in a reduce phase. This reduces the amount of shuffle data by more than 50% (e.g., to 12MB for Spark and 48MB for Dask). Figure 6 shows the improvements in the runtimes, by $\sim 20\%$ for Spark and Dask, but not MPI4py. Further, we were able to run very large datasets, such as the 4M dataset, using this architectural approach using Spark and MPI4py. Dask was restarting its worker processes because their memory utilization was reaching 95%.

Spark, and Dask have comparable performance with MPI on 32 cores, which utilizes a single node on Wrangler. However, the MPI4py implementation scales almost linearly for all datasets, Spark and Dask cannot, reaching a maximum of ~ 5 for the three smaller datasets. In addition, Spark is able to scale almost linearly for the 4M atoms dataset providing comparable performance to MPI4py.

4.3.4 Tree-Search. A bottleneck of approaches 1, 2 and 3 is the edge discovery via the naive calculation of the distances between all pairs of atoms. In approach 4, we replace the pairwise distance function with a tree-based, nearest neighbor search algorithm, in particular BallTree [36]. The algorithm: (1) constructs a tree, and (2) queries for neighboring atoms. Using the tree-search the computational complexity can be reduced from n^2 to $\log$. For our implementation, we use the BallTree as offered by Scikit-Learn [37].

Figure 6 illustrates the performance of the implementation. For small datasets, i.e., 131k and 262k atoms, approach 3 is faster than the tree-based approach, since the number of points is too small. For the large datasets, the tree approach is faster. In addition, the tree algorithm has a smaller memory footprint than `cdist`. This allowed to scale to larger problems, e.g., a 4M atoms and 44.6M edges dataset without changing the total number of tasks.

Dask shows better scaling than Spark for 131k, 262k, and 524k atoms. This is not the case for the 4M atoms, indicating that Dask’s communication layer is not able to scale as well as Spark’s. Spark shows similar performance with MPI4py for the largest dataset due to minimal shuffle traffic. Thus, MPI’s efficient communication does not become relevant.

4.4 Conceptual Framework and Discussion

In this section we provide a conceptual framework that allows application developers to carefully select a framework according to their requirements (e.g., computational and I/O characteristics).

	RADICAL-Pilot	Spark	Dask
Task Management			
Low Latency	-	o	+
Throughput	-	+	++
MPI/HPC Tasks	+	o	o
Task API	+	o	++
Large Number of Tasks	-	++	++
Application Characteristics			
Python/native Code	++	o	+
Java	o	++	o
Higher-Level Abstraction	-	++	+
Shuffle	-	++	+
Broadcast	-	++	+
Caching	-	++	o

Table 3: Decision Framework: Criteria and Ranking for Framework Selection. - : Unsupported or low performance
+: Supported, ++: Major Support, and o: Minor support.

Thus, it is important to understand both the properties of the application and Big Data frameworks: Table 3 illustrates the criteria of the conceptual framework and ranks the three frameworks.

4.4.1 Application Perspective. We have shown that it is possible to implement MD trajectory data analysis applications using all three frameworks, as well as using MPI4py. The performance critically depends on implementation aspects, such as the computational complexity, and the amount of data that needs to be shuffled across the network. For embarrassingly parallel applications, such as the path similarity analysis, with coarse grained tasks, the choice of the framework does not have a large influence on the performance (Figures 4 and 5). In addition, the performance difference compared to MPI4py was not significant (Figures 4 and 5). Thus, other aspects, such as programmability and integrate-ability become more important.

For fine-grained data parallelism, a Big Data framework, such as Spark and Dask, clearly outperforms RADICAL-Pilot (Figures 6, 8). If some coupling is introduced, i.e. communication between the tasks is required, e.g., a reduce, the usage of Spark becomes advantageous (Approaches 3 & 4). MPI4py outperformed Dask, and Spark, despite both frameworks scaling for the larger datasets. Especially Spark was able to provide linear speedup for approach 3 of the Leaflet Finder algorithm (Figure 6). Integrating with frameworks that provide higher level abstractions provides scalable solutions for more complex algorithms. However, integrating Spark with other tools needs to be carefully considered: the integration of Python tools, e.g. MDAnalysis, often causes overheads due to the frequent need for serialization and copying data between the Python and Java space.

4.4.2 Framework Perspective. RADICAL-Pilot is particularly suited for HPC applications, e.g., ensembles (up to 50k tasks) of parallel MPI applications, as shown in Ref. [11, 38]. It shows scalability limitations with supporting large numbers of short-running tasks as often found in data-intensive workloads. The file staging implementation of RADICAL-Pilot is not suitable for supporting the data exchange patterns, i.e. shuffling, required for these applications. However, executing MPI and Spark applications alongside on the same resource makes RADICAL-Pilot particularly suitable when different programming paradigms need to be combined.

Dask provides a highly flexible, low-latency task management and excellent support for parallelization of Python libraries. Based on Figures 2 and 3, we established that Dask has higher throughput.

However, we see that Spark provides better speedups for the largest datasets compared to Dask (Figure 6). Thus, Dask’s communication subsystem has some weaknesses that are particularly visible in the broadcast and shuffle performance, since the broadcast (Leaflet Finder approach 1) and shuffle (Leaflet Finder approaches 2- 4) performance is significantly worse for larger problems compared to Spark. Spark needs to be particularly considered for shuffle-intensive applications. Its in-memory caching mechanism is particularly suited for iterative algorithms that maintain a static set of points in-memory and conduct multiple passes on that set.

5 RELATED WORK

MD analysis algorithms were until recently executed serially and parallelization was not straightforward. During the last years several frameworks emerged providing parallel algorithms for analyzing MD trajectories. Some of those frameworks are CPPTraj [35], HiMach [15], Pteros 2.0 [39], MDTraj [40], and nMoldyn-3 [41]. We compare these frameworks with our approach over the parallelization techniques used. Any

CPPTRAJ [35] provides several analysis algorithms parallelized through MPI and OpenMP. MPI is being used to parallelize the execution over the frame of a single trajectory or each trajectory in an ensemble of trajectories. OpenMP is used to parallelize the execution of compute intensive algorithms.

HiMach [15] was developed by D. E. Shaw Research group to provide a parallel analysis framework for MD simulations, extends Google’s MapReduce. HiMach API defines trajectories, does per frame data acquisition (Map) and cross-frame analysis (Reduce). HiMach’s runtime is responsible to parallelize and distribute Map and Reduce phases to resources. Data transfers are done through a communication protocol created specifically for HiMach.

Pteros-2.0 [39] is a open-source library that is used for modeling and analyzing MD trajectories, providing a plugin for each supported algorithm. The execution is done by a user defined driver application, which setups trajectory I/O and frame dispatch for analysis. It offers a C++ and Python API. Pteros 2.0 parallelizes computational intensive algorithms by using OpenMP and Multithreading. As a result, it is bounded to execute on a single node, making any analysis execution highly dependent on memory size. Through RADICAL-Pilot, Spark and Dask, we avoided the need to recompile every time there is a change to the underlying resource, ensuring the application’s execution.

MDTraj [40] is a Python package for analyzing MD trajectories. It links MD data and Python statistical and visualization software. MDTraj proposes parallelizing the execution by using the parallel package of IPython as a wrapper along with an out-of-core trajectory reading method. Our approach support of data analysis frameworks allows data parallelization on any level of the execution, not only in data read.

nMoldyn-3 [41] parallelizes the execution through a Master-Worker architecture. The master or client defines analysis tasks, submits them to a task manager, which then are executed by the worker process. In addition, it provides adaptability allowing on-the-fly addition of resources, and execution fault tolerance when worker processes disconnect.

In contrast, our approach utilizes more general purpose frameworks for parallelization. Because the used frameworks provide higher level abstractions, e.g machine learning, any integration with other data analysis methods can be fast and easier. In addition, resource acquisition and management is done transparently.

6 CONCLUSION AND FUTURE WORK

In this paper, we investigated the use of different programming abstractions and frameworks for the implementation of a range of algorithms for molecular dynamics trajectory analysis. We conducted an in-depth analysis of the application characteristics and assessed the architectures of RADICAL-Pilot, Spark and Dask. We provide a conceptual framework that enables application developers to qualitatively evaluate Big Data frameworks with respect to their application requirements. Our benchmarks enable them to quantitatively assess framework performance as well as the expected performance of different implementation alternatives. In addition, we provided a comparison with the respective MPI implementations.

While the task abstractions provided by all frameworks are well-suited for implementing all use cases, the high-level MapReduce programming model provided by Spark and Dask provides several advantages: it is easier to use and efficiently support common data exchange patterns, such as the shuffling of data between the map and reduce stage. In our benchmarks Spark outperforms Dask in communication-intensive tasks, such as broadcasts and shuffles, as shown in Fig 6. Further, the in-memory RDD abstraction performs well for iterative algorithms (such as many machine learning algorithms [32]). Dask provides more versatile low-level and high-level APIs and integrates better with pythonic frameworks. RADICAL-Pilot does not provide a MapReduce API, but is well suited for coarse-grained task-level parallelism [11, 38] and for cases where HPC and analytics framework need to be integrated. We also identified severe limitation in Dask and Spark: while both frameworks provide some support for linear algebra – both provide abstractions for distributed array – these proved not flexible enough for implementing the all-pairs patterns efficiently and required significant workarounds in the implementation and utilization of out-of-framework functions to read and partition the input data (Table 2). Although, none of these frameworks outperformed MPI, their scaling capabilities along with their high-level APIs create a strong case on utilizing them for data analytics of HPC applications.

In the future, we will further improve the performance of the presented algorithms, e. g., by reducing the memory and computation footprint, data transfer sizes between stages, by optimizing filesystem usage. To better support PyData tools in RADICAL-Pilot, we plan to extend the Pilot-Abstraction to support Dask and other Big Data frameworks. Further, we will refine the RADICAL-Pilot task execution engine to meet the requirement of data analytics application and devise task execution strategies that can mitigate issues occurring at large scale, such as stragglers. Another area of research, is dynamic resource management and the ability to dynamically scale the resource pool (e. g., by adding or removing nodes) to meet the requirements of a specific application stage.

Acknowledgements We thank Andre Merzky for useful discussions. This work is funded by NSF 1443054 and 1440677. Computational resources were provided by NSF XRC awards TG-MCB090174 and TG-MCB130177.

REFERENCES

- [1] S. Kamburugamuve, P. Wickramasinghe, S. Ekanayake†, and G. C. Fox, “Anatomy of machine learning algorithm implementations in mpi, spark, and flink,” in *Technical Report*, Indiana University, Bloomington, 2017.
- [2] J. Dean and S. Ghemawat, “MapReduce: Simplified Data Processing on Large Clusters,” in *OSDI’04: Proceedings of the 6th conference on Symposium on Operating Systems Design & Implementation*. Berkeley, CA, USA: USENIX Association, 2004, pp. 137–150.
- [3] X. Lu, D. Shankar, S. Guhani, and D. K. Panda, “High-performance design of apache spark with rdma and its benefits on various workloads,” December 2016.
- [4] G. Fox, J. Qiu, S. Jha, S. Kamburugamuve, and A. Luckow, “Hpc-abds high performance computing enhanced apache big data stack,” in *Proceedings of Workshop on Scalable Computing For Real-Time Big Data Applications (SCRAMBL’15)*. Shenzhen, China: 15th IEEE/ACM International Symposium on Cluster, Cloud and Grid Computing, 2015.
- [5] S. Kamburugamuve, G. Fox, P. Wickramasinghe, G. Kannan, and V. Abeykoon, “Twister:net - communication library for big data processing in hpc and cloud environments,” 03 2018.
- [6] M. Zaharia, M. Chowdhury, M. J. Franklin, S. Shenker, and I. Stoica, “Spark: Cluster computing with working sets,” in *Proceedings of the 2Nd USENIX Conference on Hot Topics in Cloud Computing*, ser. HotCloud’10. Berkeley, CA, USA: USENIX Association, 2010, pp. 10–10.
- [7] M. Rocklin, “Dask: Parallel computation with blocked algorithms and task scheduling,” in *Proceedings of the 14th Python in Science Conference*, K. Huff and J. Bergstra, Eds., 2015, pp. 130 – 136.
- [8] T. Cheatham and D. Roe, “The impact of heterogeneous computing on workflows for biomolecular simulation and analysis,” *Computing in Science Engineering*, vol. 17, no. 2, pp. 30–39, 2015.
- [9] V. Balasubramanian, I. Bethune, A. Shkurti, E. Breitmoser, E. Hruska, C. Clementi, C. Laughton, and S. Jha, “Extasy: Scalable and flexible coupling of md simulations and advanced sampling techniques,” in *2016 IEEE 12th International Conference on e-Science (e-Science)*, Oct 2016, pp. 361–370.
- [10] S. Jha, J. Qiu, A. Luckow, P. K. Mantha, and G. C. Fox, “A tale of two data-intensive paradigms: Applications, abstractions, and architectures,” *Proceedings of 3rd IEEE International Congress of Big Data*, vol. abs/1403.1528, 2014.
- [11] A. Merzky, M. Turilli, M. Maldonado, M. Santcroos, and S. Jha, “Using Pilot Systems to Execute Many Task Workloads on Supercomputers,” 2018, <http://arxiv.org/abs/1512.08194>.
- [12] A. Luckow, M. Santcroos, A. Merzky, O. Weidner, P. Mantha, and S. Jha, “P*: A model of pilot-abstractions,” *IEEE 8th International Conference on e-Science*, pp. 1–10, 2012, <http://dx.doi.org/10.1109/eScience.2012.6404423>.
- [13] L. Dalcín, R. Paz, and M. Storti, “Mpi for python,” *Journal of Parallel and Distributed Computing*, vol. 65, no. 9, pp. 1108 – 1115, 2005. [Online]. Available: <http://www.sciencedirect.com/science/article/pii/S0743731505000560>
- [14] M. Khoshlessan, I. Paraskevakos, S. Jha, and O. Beckstein, “Parallel Analysis in MDAnalysis using the Dask Parallel Computing Library,” in *Proceedings of the 16th Python in Science Conference*, K. Huff, D. Lippa, D. Niederhut, and M. Pacer, Eds., Austin, TX, 2017, pp. 64–72. [Online]. Available: http://conference.scipy.org/proceedings/scipy2017/mahzad_khoshlessan.html
- [15] T. Tu, C. A. Rendleman, D. W. Borhani, R. O. Dror, J. Gullingsrud, M. O. Jensen, J. L. Klepeis, P. Maragakis, P. Miller, K. A. Stafford, and D. E. Shaw, “A scalable parallel framework for analyzing terascale molecular dynamics simulation trajectories,” in *2008 SC - International Conference for High Performance Computing, Networking, Storage and Analysis*, Nov 2008, pp. 1–12.
- [16] G. C. Fox, S. Jha, J. Qiu, and A. Luckow, “Towards an understanding of facets and exemplars of big data applications,” in *Proceedings of Beowulf’14*. Annapolis, MD, USA: ACM, 2014.
- [17] C. Mura and C. E. McAnany, “An introduction to biomolecular simulations and docking,” *Molecular Simulation*, vol. 40, no. 10–11, pp. 732–764, 2014. [Online]. Available: <http://dx.doi.org/10.1080/08927022.2014.935372>
- [18] S. L. Seyler, A. Kumar, M. F. Thorpe, and O. Beckstein, “Path similarity analysis: A method for quantifying macromolecular pathways,” *PLoS Comput Biol*, vol. 11, no. 10, pp. 1–37, 10 2015. [Online]. Available: <http://dx.doi.org/10.1371/journal.pcbi.1004568>
- [19] N. Michaud-Agrawal, E. J. Denning, T. B. Woolf, and O. Beckstein, “Mdanalysis: A toolkit for the analysis of molecular dynamics simulations,” *Journal of Computational Chemistry*, vol. 32, no. 10, pp. 2319–2327, 2011. [Online]. Available: <http://dx.doi.org/10.1002/jcc.21787>
- [20] Richard J. Gowers, Max Linke, Jonathan Barnoud, Tyler J. E. Reddy, Manuel N. Melo, Sean L. Seyler, Jan Domański, David L. Dotson, Sébastien Buchoux, Ian M. Kenney, and Oliver Beckstein, “MDAnalysis: A Python Package for the Rapid Analysis of Molecular Dynamics Simulations,” in *Proceedings of the 15th Python in Science Conference*, Sebastian Benthall and Scott Rostrup, Eds., 2016, pp. 98 – 105.
- [21] S. Van Der Walt, S. C. Colbert, and G. Varoquaux, “The numpy array: a structure for efficient numerical computation,” *Computing in Science & Engineering*, vol. 13, no. 2, pp. 22–30, 2011.
- [22] E. Jones, T. Oliphant, P. Peterson *et al.*, “SciPy: Open source scientific tools for Python,” 2001–. [Online]. Available: <http://www.scipy.org/>
- [23] D. P. Huttenlocher, G. A. Klanderman, and W. J. Rucklidge, “Comparing images using the hausdorff distance,” *IEEE Transactions on Pattern Analysis and Machine Intelligence*, vol. 15, no. 9, pp. 850–863, 1993.
- [24] A. A. Taha and A. Hanbury, “An efficient algorithm for calculating the exact hausdorff distance,” *IEEE Transactions on Pattern Analysis and Machine Intelligence*, vol. 37, no. 11, pp. 2153–2163, Nov 2015.
- [25] S. Jha, D. S. Katz, A. Luckow, N. Chue Hong, O. Rana, and Y. Simmhan, “Introducing distributed dynamic data-intensive (d3) science: Understanding applications and infrastructure,” *Concurrency and Computation: Practice and Experience*, pp. e4032–n/a, 2017, e4032 cpe.4032. [Online]. Available: <http://dx.doi.org/10.1002/cpe.4032>
- [26] J. Ekanayake, H. Li, B. Zhang, T. Gunaratne, S.-H. Bae, J. Qiu, and G. Fox, “Twister: A runtime for iterative mapreduce,” in *Proceedings of the 19th ACM International Symposium on High Performance Distributed Computing*, ser. HPDC ’10. New York, NY, USA: ACM, 2010, pp. 810–818. [Online]. Available: <http://doi.acm.org/10.1145/1851476.1851593>
- [27] M. Zaharia, M. Chowdhury, T. Das, A. Dave, J. Ma, M. McCauley, M. J. Franklin, S. Shenker, and I. Stoica, “Resilient distributed datasets: A fault-tolerant abstraction for in-memory cluster computing,” in *Proceedings of the 9th USENIX Conference on Networked Systems Design and Implementation*, ser. NSDI’12. Berkeley, CA, USA: USENIX Association, 2012, pp. 2–2. [Online]. Available: <http://dl.acm.org/citation.cfm?id=2228298.2228301>
- [28] M. Turilli, M. Santcroos, and S. Jha, “A comprehensive perspective on pilot-jobs,” *ACM Computing Surveys (accepted, in press)*, *arXiv preprint arXiv:1508.04180v3*, 2017, <https://arxiv.org/abs/1508.04180>.
- [29] M. Turilli, A. Merzky, V. Balasubramanian, and S. Jha, “A building blocks approach towards domain specific workflow systems,” *Short Paper (IEEE/ACM CCGrid 2018)*, 2018, <http://arxiv.org/abs/1609.03484>.
- [30] V. Balasubramanian, M. Turilli, W. Hu, M. Lefebvre, W. Lei, G. Cervone, J. Tromp, and S. Jha, “Harnessing the Power of Many: Extensible Toolkit for Scalable Ensemble Applications,” *IPDPS 2018 (accepted)*, 2018, <https://arxiv.org/abs/1710.08491>.
- [31] M. Turilli, Y. N. Babuji, A. Merzky, M. T. Ha, M. Wilde, D. S. Katz, and S. Jha, “Evaluating distributed execution of workloads,” in *2017 IEEE 13th International Conference on e-Science (e-Science)*, Oct 2017, pp. 276–285.
- [32] X. Meng, J. Bradley, B. Yavuz, E. Sparks, S. Venkataraman, D. Liu, J. Freeman, D. Tsai, M. Amde, S. Owen, D. Xin, R. Xin, M. J. Franklin, R. Zadeh, M. Zaharia, and A. Talwalkar, “Mllib: Machine learning in apache spark,” *Journal of Machine Learning Research*, vol. 17, no. 34, pp. 1–7, 2016. [Online]. Available: <http://jmlr.org/papers/v17/15-237.html>
- [33] P. K. Mantha, A. Luckow, and S. Jha, “Pilot-MapReduce: an extensible and flexible MapReduce implementation for distributed data,” in *Proceedings of third international workshop on MapReduce and its Applications*, ser. MapReduce ’12. New York, NY, USA: ACM, 2012, pp. 17–24. [Online]. Available: <https://raw.githubusercontent.com/saga-project/radical.wp/master/publications/pdf/pilot-mapreduce2012.pdf>
- [34] A. Luckow, I. Paraskevakos, G. Chantzalexioiu, and S. Jha, “Hadoop on hpc: Integrating hadoop and pilot-based dynamic resource management,” *2016 IEEE International Parallel and Distributed Processing Symposium Workshops (IPDPSW)*, pp. 1607–1616, 2016.
- [35] D. R. Roe and I. Thomas E. Cheatham, “Ptraj and cpptraj: Software for processing and analysis of molecular dynamics trajectory data,” *Journal of Chemical Theory and Computation*, vol. 9, no. 7, pp. 3084–3095, 2013, pMID: 26583988. [Online]. Available: <http://dx.doi.org/10.1021/ct400341p>
- [36] S. M. Omohundro, “Five balltree construction algorithms,” Tech. Rep., 1989.
- [37] “Scikit-Learn: Nearest Neighbors,” <http://scikit-learn.org/stable/modules/neighbors.html>, 2016.
- [38] A. Merzky, M. Turilli, M. Maldonado, and S. Jha, “Design and performance characterization of radical-pilot on titan,” *in preparation*, 2018, <https://arxiv.org/abs/1801.01843>.
- [39] S. O. Yesylevskyy, “Pteros 2.0: Evolution of the fast parallel molecular analysis library for c++ and python,” *Journal of Computational Chemistry*, vol. 36, no. 19, pp. 1480–1488, 2015. [Online]. Available: <http://dx.doi.org/10.1002/jcc.23943>
- [40] R. McGibbon, K. Beauchamp, M. Harrigan, C. Klein, J. Swails, C. Hernández, C. Schwantes, L.-P. Wang, T. Lane, and V. Pande, “Mdtraj: A modern open library for the analysis of molecular dynamics trajectories,” *Biophysical Journal*, vol. 109, no. 8, pp. 1528 – 1532, 2015. [Online]. Available: <http://www.sciencedirect.com/science/article/pii/S0006349515008267>
- [41] K. Hinsen, E. Pellegrini, S. Stachura, and G. R. Kneller, “nmoldyn 3: Using task farming for a parallel spectroscopy-oriented analysis of molecular dynamics simulations,” *Journal of Computational Chemistry*, vol. 33, no. 25, pp. 2043–2048, 2012. [Online]. Available: <http://dx.doi.org/10.1002/jcc.23035>