

UniFaaS: Programming across Distributed Cyberinfrastructure with Federated Function Serving

Yifei Li¹, Ryan Chard², Yadu Babuji³, Kyle Chard^{3,2}, Ian Foster^{2,3}, Zhuozhao Li^{1,*}

¹Dept. of Computer Science and Engineering, Southern University of Science and Technology, Shenzhen, China

²Data Science and Learning Division, Argonne National Laboratory, Lemont, IL, USA

³Dept. of Computer Science, University of Chicago, Chicago, IL, USA

Abstract—Modern scientific applications are increasingly decomposable into individual functions that may be deployed across distributed and diverse cyberinfrastructure such as supercomputers, clouds, and accelerators. Such applications call for new approaches to programming, distributed execution, and function-level management. We present UniFaaS, a parallel programming framework that relies on a federated function-as-a-service (FaaS) model to enable composition of distributed, scalable, and high-performance scientific workflows, and to support fine-grained function-level management. UniFaaS provides a unified programming interface to compose dynamic task graphs with transparent wide-area data management. UniFaaS exploits an observe-predict-decide approach to efficiently map workflow tasks to target heterogeneous and dynamic resources. We propose a dynamic heterogeneity-aware scheduling algorithm that employs a delay mechanism and a re-scheduling mechanism to accommodate dynamic resource capacity.

Our experiments show that UniFaaS can efficiently execute workflows across computing resources with minimal scheduling overhead. We show that UniFaaS can improve the performance of a real-world drug screening workflow by as much as 22.99% when employing an additional 19.48% of resources and a montage workflow by 54.41% when employing an additional 47.83% of resources across multiple distributed clusters, in contrast to using a single cluster.

Index Terms—federated cyberinfrastructure, federated function serving

I. INTRODUCTION

The rapid adoption of hardware accelerators and exponential increases in data volumes have spurred a major transformation in the nature of programming. Developers increasingly decompose previously monolithic workflows into many individual, often lightweight tasks [1]. These tasks may be individual functions or executables developed in different programming languages, with varied resource requirements—from single-core functions and multi-core simulation codes, to accelerator-based computations. When workflows are decomposed in this way, it is then feasible, and indeed often desirable, to execute different tasks in different locations: for example, where accelerators or resources are available, where software is installed, near data, or with the best cost-performance efficiency.

Unfortunately, such fine-grain workflows are not well supported by current research cyberinfrastructure (CI) that is managed by batch schedulers designed for scheduling large

batch jobs and typically have long (and unpredictable) queue times. As researchers often have access to a *resource pool* containing several computing resources (e.g., institutional clusters, supercomputers, and accelerators), one may want to exploit these resources concurrently to amortize queue times (though perhaps at the cost of additional data transfers) for faster scientific analyses. For example, it may take only minutes to transfer gigabytes and hours to transfer terabytes [2], but days to obtain hundreds of nodes on an oversubscribed supercomputer.

As a result, researchers increasingly use multiple resources *together* across federated CI (e.g., supercomputers, clouds, accelerators, and local compute) to run a single workflow—a landscape referred to as *cross-facility computing* [3]. Such a trend towards fine-grained and distributed scientific workflows drives new system requirements. To name just a few (more are described in §II): 1) *programmability* to allow users to construct programs able to execute on diverse and distributed CI with minimal barriers; 2) automatic *fine-grained* task-level management in terms of resource provisioning and heterogeneous environment management; and 3) enabling *flexible* task execution on dynamic resources as the resource availability across federated CI may change at runtime.

Myriad systems have been developed to address the needs of users writing and deploying workflows [4]. While these systems address some of the requirements above, none address the entire set of requirements. For example, many workflow systems [5], [6] rely on domain-specific languages to support workflow composition. Python-based computing frameworks such as Ray [7], Parsl [8], and Dask [9] support development of parallel programs on a single computer but cannot easily adapt to federated CI. Pegasus [10] leverages a static configuration-based model to compose scientific workflows as directed-acyclic graphs (DAGs) but requires deployment of HTCondor [11] which is rarely supported on HPC systems and cannot support fine-grained task-level management.

While in principle one can develop modern workflows by writing batch scripts or combining various workflow systems for federated CI, this approach in practice not only imposes a significant development and management burden (e.g., fault tolerance, data movement, and resource management) on researchers, but is also not flexible, scalable, or portable across computers, reducing efficiency and reproducibility that are

*Zhuozhao Li is the corresponding author. Email: lizz@sustech.edu.cn.

crucial for science [12].

We present UniFaaS, a general-purpose, parallel programming framework that adapts a federated function-as-a-service (FaaS) model to enable developers to compose distributed, scalable, and high-performance scientific workflows that span federated CI. The FaaS model, first implemented by commercial cloud providers [13], can benefit modern scientific applications from several perspectives. FaaS reduces the burden of managing computing resources and instead exposes a function-based API that allows users to manage and schedule scientific workflows at a fine-grained function level. Further, it enables the use of containers to both simplify the management of complicated environments and reduce portability challenges. Finally, FaaS, by definition, abstracts the resource pool and thus enables resources to be dynamically added (or removed) during execution.

We implement UniFaaS upon *funcX* [14], [15], a *federated* FaaS platform that extends the FaaS model and enables users to invoke functions on arbitrary resources. UniFaaS provides a *unified* programming interface to express task parallelism and compose dynamic dependency graphs, which can be further deployed across distributed resources seamlessly. UniFaaS implements a data manager to *transparently* manage data transfers across computers on behalf of users, using widely-used transfer mechanisms such as Globus [16] and rsync.

It is non-trivial to achieve high performance in cross-facility computing, since the resources across federated CI are heterogeneous and each resource may have dynamic capacity during execution due to scheduled downtimes and use by others.

In this paper, we explore an observe-predict-decide approach to improve the performance: UniFaaS monitors the key characteristics (e.g., input data sizes and environments) of tasks on different computers and predicts the task performance (e.g., transfer and execution time) using common performance models. We propose a dynamic heterogeneity-aware scheduling algorithm that employs a delay mechanism and a re-scheduling mechanism to accommodate the workflow and resource dynamics. UniFaaS supports elasticity—it can automatically scale various resources based on workflow characteristics. Furthermore, the modular design of UniFaaS allows users to easily plug in any appropriate schedulers or data transfer mechanisms for their workflows.

To the best of our knowledge, UniFaaS is the first framework that adapts the convenient federated FaaS model to enable monolithic workflows to be broken into small schedulable function units that can be flexibly executed across a resilient, federated resource pool. In this paper, we explore the feasibility of adapting the FaaS paradigm and aim to design an extensible system to address the unique challenges of programmability, scheduling, and data management (more in §II) in cross-facility computing. Together, these features enable productive parallel programming and simple, yet performant, execution management on federated CI.

The contributions of our work are as follows:

- We design UniFaaS, a general-purpose parallel programming framework that simplifies workflow structuring and enables fine-grained function-level management of workflows in federated environments.
- We propose a dynamic heterogeneity-aware scheduler and a re-scheduling mechanism to handle heterogeneous and dynamic computing environments.
- Our evaluation shows that UniFaaS can deploy workflows across up to 16 computing resources with high performance.
- We discuss lessons learned when using UniFaaS in real use cases and identify areas for further improvement.

The rest of this paper is as follows. §II presents general requirements for modern scientific workflows. §III presents the programming model of UniFaaS. §IV describes the UniFaaS system architecture, optimizations, and implementations. §V evaluates UniFaaS’s performance. §VII discusses the lessons learned with UniFaaS. §VIII discusses related work. §IX concludes the paper with future remarks.

II. MOTIVATION AND BACKGROUND

In this section, we highlight the key requirements for modern scientific workflows deployed across federated CI by describing a real-world use case. We note that these characteristics are shared by many use cases, for example in modern AI-driven simulations and data-driven workflows [17]–[26]. Further, recent work has shown that such distribution can also improve energy consumption of workflows [27].

The COVID-19 pandemic highlighted the need for discovery of effective therapeutics among a near-infinite search space of small molecules. One common way to accelerate the design and development of antiviral treatments is to use machine learning (ML) models to computationally screen small molecules rapidly, much faster than is possible with wetlab studies. A typical drug screening pipeline [28] involves multiple computational stages, such as molecular simulations, feature computations, fingerprinting, etc., each with distinct computational needs. Each stage relies on different toolkits and methods (e.g., simulation, machine learning), requires different types of resources (e.g., CPUs, accelerators) and different amounts of those resources, and has diverse task features (e.g., parallelism, dependency, duration, input size). Computers also vary significantly in characteristics, for example, some may be powerful but have long queue times while others may have fewer resources but are immediately available. An application that needs, for example, to screen millions of molecules quickly must be able to adapt to these different tasks and resource characteristics.

These modern use cases motivate the need for deployment in federated environments and exhibit the following key requirements.

- 1) **Fine-grained management:** workflows usually consist of tasks with diverse requirements on resources, environments, software dependencies, data locations, etc., and thus require fine-grained task-level management automatically to reduce the burden on users.

- 2) **Portability:** scientific applications require portability in different dimensions: first, a workflow may run on various computing resources with different environments for reproducing and sharing [28]; second, a stage may be portable to different computations, e.g., different ML models or different methods.
- 3) **Elasticity:** workflows may have different resource demands at different stages, and thus require elasticity, i.e., automatic resource provisioning.
- 4) **Programmability:** one should be able to use simple and straightforward languages to express parallelism and compose workflows with dynamic dependency graphs.
- 5) **Data dependency:** handling complex data dependencies across CI is complicated. Abstracting data movements between resources is crucial to simplify development.
- 6) **Dynamic execution:** a workflow may be dynamically changing based on the prior computation results and its functions may execute dynamically on various resources based on resource availability and status.
- 7) **Performance:** given the heterogeneous and dynamic (resource availability may vary) nature of modern CI, tasks require to be scheduled to appropriate resources with high performance.

Existing systems only partially satisfy these requirements and in general are not designed to manage computation across federated CI. We are thus motivated to develop UniFaaS. UniFaaS leverages *funcX* as the main execution backend. *funcX* is a federated function serving fabric and supports executing function tasks on arbitrary computing resources in a FaaS manner. UniFaaS leverages two main components in *funcX*: **endpoint** and **client**.

An **endpoint** represents a computing resource that can execute tasks using the FaaS model, i.e., the endpoint can elastically launch multiple *worker* processes (or containers) and assign a task to a worker to be performed. One can deploy the endpoint software on any computer and integrate it into the execution fabric as an available resource.

The *funcX* **client** provides a secure means to interact with the cloud-hosted *funcX* web service to submit function tasks to specific endpoints, track task states, and retrieve task results. When a function task is submitted to UniFaaS, UniFaaS employs internally the *funcX* client to dispatch the task to the remote endpoint and retrieve the results.

We leverage *funcX* as the execution backend for two important reasons. First, *funcX* allows one to easily add/remove a resource via its flexible endpoint software. Second, the FaaS model can inherently satisfy requirements 1–3: fine-grain function management, portability, and elasticity. However, *funcX* offers only an independent task execution interface. *There are still research challenges (e.g., programmability, data, dynamic execution, and performance) toward a simple programming model across a resilient, federated resource pool.* In this paper, we discuss how we design UniFaaS to satisfy the requirements and solve the challenges of adapting the FaaS model for programming federated scientific workflows.

III. PROGRAMMING WITH UNIFAAS

To satisfy the programmability requirement, UniFaaS adopts a Python-based programming model and allows users to compose workflows with dynamic dependency graphs.

A. Programming Model

Functions and tasks: A *function* can be a pure Python function or a function call to other software. Each function to be executed on remote computing resources must be decorated with `@function`. A *task* represents an invocation request of a function. The invocation to a decorated function does not return results immediately, but instead returns a `Future` object, indicating that the function instance is being executed asynchronously. A function may accept Python objects or `Future` objects of other tasks as input arguments. This mechanism is commonly used in existing Python-based systems such as Dask [9], Parsl [8], and Ray [7].

Data: A function may be invoked with Python objects and files as input arguments. Python objects are directly serialized when passed as input arguments. There is a hard limit on the size (10 MB) of a Python object that can be passed among resources in *funcX*. Any object with a size larger than the limit must be serialized and invoked as a `RemoteFile` object.

UniFaaS enables users to read and write files via the `RemoteFile` object in a function. When a task needs files located on a remote resource, UniFaaS supports transferring the files via two mechanisms: Globus [16] and rsync. The `RemoteFile` object includes two subclasses: `GlobusFile` and `RsyncFile` (more details in §IV-E). We will use `GlobusFile` below to illustrate the programming model.

```

1 @function
2 def compute_fingerprint(GlobusFile: mol_file):
3     from rdkit import *
4     import GlobusFile
5
6     mol_path = mol_file.get_remote_file_path()
7     molecule = open(mol_path).readline()
8     fp = AllChem.GetMorganFingerprint(
9         Chem.MolFromSmiles(molecule), 2)
10
11     out_file = GlobusFile.create("fp.txt")
12     out_path = out_file.get_remote_file_path()
13     open(out_path, 'w').write(fp)
14     return out_file

```

Listing 1. An example function with `GlobusFile`. The function is to compute the fingerprint of a molecule in SMILES format.

Listing 1 shows a `compute_fingerprint` function that accepts a `mol_file` parameter of type `GlobusFile`. The path of the remote file can be retrieved via the `get_remote_file_path()` method and further read/write operations on the file can be done via Python’s built-in I/O libraries. One needs to call the `GlobusFile.create` method to create a new file on the remote compute resource, which returns a `GlobusFile` object that can be recognized and managed by UniFaaS.

B. Dynamic Task Graph

UniFaaS represents a workflow as a directed acyclic graph (DAG), where each node indicates a task and each edge

indicates a dependency between two tasks. UniFaaS allows passing `future` objects as function arguments to construct task graphs. Specifically, when a `future` F is passed to a task T , UniFaaS adds an edge from the task that is associated with F to T . A task can execute only when all its dependencies are ready. Through `future` passing, UniFaaS supports constructing a task graph *dynamically*, which means that the graph can change during execution time.

C. Configuration

Each `funcX` endpoint is assigned a universally unique identifier (UUID) when it is deployed. UniFaaS provides a `Config` interface that allows one to specify endpoints as computing resources by their UUIDs, as shown in Listing 2. One can also configure other parameters such as scheduling strategy, the maximum number of retries, and file transfer type. The `Config` interface is separated from the programming interface. In other words, one can write a workflow once and deploy it on different sets of endpoints by simply updating the UUIDs specified in the workflow configuration, enabling *write once, run anywhere*.

```

1 config = Config(
2     executors=[
3         Executor(label="Cluster1",
4                 endpoint="6156af-...-54e93"),
5         Executor(label="Cluster2",
6                 endpoint="9c2344-...-7ff98"),
7     ],
8     scheduling_strategy="LOCALITY",
9     max_transfer_retries=3,
10    file_transfer_type="Globus")

```

Listing 2. An example of the `Config` interface.

IV. ARCHITECTURE AND IMPLEMENTATION

UniFaaS comprises five system components: monitors, profilers, scheduler, data manager, and task executor, as shown in Figure 1. All components are extensible to any appropriate alternatives. In this section, we describe the design of each component and discuss how UniFaaS can satisfy the requirements presented in §II.

A. Overview

The performance of tasks at different stages of a federated workflow may vary significantly due to the heterogeneity of both tasks and hardware. Fortunately, decomposing a workflow into functions enables fine-grained task performance prediction. Scientific workflows are often run repeatedly on similar sets of computers and have predictable performance [29]. These characteristics together allow us to exploit an observe-predict-decide approach to improve the performance of federated workflows. We briefly describe the execution flow of a workflow with UniFaaS as follows:

- 1) Deploy `funcX` endpoint software on the accessible resources and supply the endpoint UUIDs in the `Config`.
- 2) Create the functions, decorate them with `@function`, and express task graphs. The DAG generator analyzes the task dependencies and constructs a DAG.

- 3) The profilers load task characteristics from the local database (if any), build performance models, and predict execution time and data transfer time for tasks when needed by the scheduler.
- 4) The scheduler creates a schedule for tasks based on the information provided by the profilers and monitors.
- 5) Once the scheduling decisions are made, the data manager performs data transfers in advance when possible, and the task executor submits tasks to endpoints.
- 6) When a task runs, the task monitor tracks the characteristics and logs them into the local database after the task completes. The profilers are updated accordingly based on the latest runs.

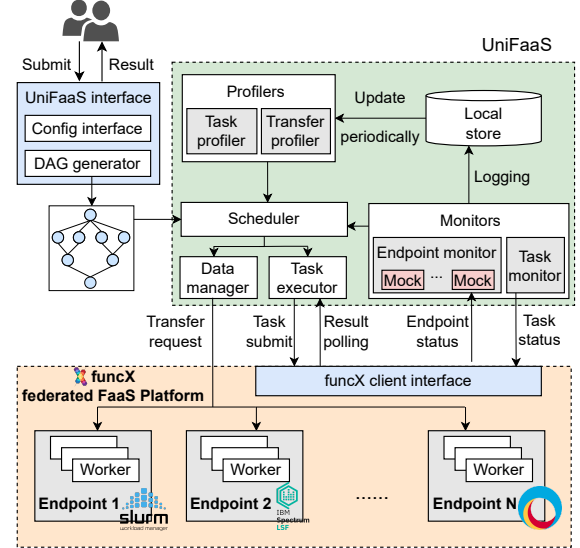


Fig. 1. UniFaaS architecture.

B. Monitors

Task monitor: This component is responsible for monitoring task execution information, such as task states (e.g., running, completed, failed, etc.), task characteristics (e.g., CPU utilization) and task completion times on various computers. The monitored information is streamed into a local database and the profilers. The local database can be treated as historical knowledge—UniFaaS allows a user to start a workflow by loading an existing database so that the profilers can pre-build performance models for the scheduler. However, the task execution status may vary over time, resulting in inaccurate scheduling decisions. Therefore, real-time task information is also fed into the profilers to update the performance models.

Endpoint monitor: UniFaaS's scheduler requires endpoint information, including hardware configuration and real-time endpoint status (e.g., available and pending tasks), to make proper scheduling decisions. While `funcX` provides RESTful interfaces for polling endpoint status, it is only updated periodically (e.g., every minute) but not in real-time, which may lead to inaccurate scheduling decisions. Moreover, frequently polling for endpoint status updates may increase the load on the `funcX` services, and hence is not practical.

To address this problem, we propose a *local mocking mechanism*. Specifically, the endpoint monitor creates a *mock*

endpoint object for each endpoint listed in the `Config` interface. A mock endpoint serves as a proxy to the corresponding genuine endpoint and has the same attributes as the genuine endpoint, including hardware information, task queues, the number of busy and idle workers, etc. When initializing a mock endpoint, the endpoint monitor communicates with the *funcX* service to retrieve initial information. When submitting a task to a genuine endpoint, a mock task is pushed into the task queue of the mock endpoint and the number of idle workers is decreased. Similarly, the mock task is popped out of the queue after the task is completed. Meanwhile, to ensure the accuracy of the mock information, the endpoint monitor synchronizes the mock objects with the *funcX* service periodically. While polling-based approaches may have varying latency depending on the network latency, the overhead of this local mocking mechanism is negligible, which enables the scheduler to obtain real-time endpoint status for accurate scheduling decisions.

C. Profilers

UniFaaS uses an execution profiler and a transfer profiler to predict the execution and transfer time. These profilers are deployed on separate threads in the UniFaaS client. They periodically learn or update models based on real-time monitoring information and historical data from prior runs and do not interfere with the main scheduling loop.

Execution profiler: When UniFaaS is initialized, the execution profiler trains an initial performance model for each function based on historical data. When a workflow runs, the profiler updates the models periodically if there is new task information collected from the monitor. We apply the random forest regression algorithm [30], [31] as the default execution performance model. The model takes the input size, number of cores, CPU frequency, and RAM size of the endpoint to run on as inputs, and estimates the execution time and output data size. We note that UniFaaS’s modular design means that users can easily extend it to other appropriate performance models such as XGBoost [32] and Bayesian linear regression [33].

Transfer profiler: Data transfer time is primarily determined by the data size and the network conditions between endpoints. Prior work [2] shows that the data transfer time is relatively predictable across federated CI. In the absence of recent data, the transfer profiler can send probing file transfers to measure the network bandwidth between endpoints when UniFaaS is initialized. The current implementation uses a polynomial regression model that uses bandwidth, data size, and the maximum number of concurrent transfers (set by the data manager in §IV-E) to predict the data transfer time.

D. Scheduler

The scheduler maps workflow tasks to heterogeneous endpoints, with the objective of minimizing the workflow completion time (i.e., makespan), which has been demonstrated to be NP-hard [34]. As aforementioned, the scheduler needs to consider that both the workflow DAGs and resource capacity may vary during execution in cross-facility computing,

which brings a unique challenge on how to make scheduling decisions at the most appropriate time. Resource capacity in this paper represents the number of workers in an endpoint and each function is mapped to a worker. The resources allocated (e.g., CPUs, GPUs, memory, and disks) per worker is configurable on the *funcX* endpoint. The scheduler can consider various resource dimensions when making decisions.

In cross-facility computing, data staging may incur significant delay, thus it is often desirable to make task scheduling decisions as early as possible (e.g., immediately after tasks are submitted), allowing overlapping data staging and computation to reduce workflow makespan. However, early scheduling decisions may result in poor performance in the case of dynamic workflows or resource capacity. To address this challenge, we propose three scheduling algorithms, capacity-aware, locality-aware, and heterogeneity-aware, to support various scenarios, from static to dynamic workflow DAGs and resource capacity.

Capacity-aware scheduling (*Capacity* in short): *Capacity* assigns workflow tasks to endpoints based on the capacity of each endpoint, i.e., the number of tasks scheduled to an endpoint is proportional to its total computing capacity. Assume we have a set of N endpoints $\mathbb{EP} = \{e_1, e_2, \dots, e_N\}$ and the capacity of the endpoints are $\mathbb{C} = \{c_1, c_2, \dots, c_N\}$ (measured by the number of workers on an endpoint). Suppose we have a set of M tasks. The number of tasks scheduled to the endpoint e_i can be computed by:

$$M_i = M * \frac{c_i}{\sum_{i=1}^N c_i}. \quad (1)$$

After determining $M_i, i \in [1, N]$, the scheduler searches the candidate tasks for each endpoint in a depth-first search (DFS) order, with the intention to consider data locality and allocate tasks on the same path to the same endpoint, reducing data transferred across endpoints.

After partitioning the DAG based on endpoint capacity, tasks are placed into the corresponding data staging queues between endpoint pairs. Upon the completion of data staging, they are immediately dispatched to remote endpoints, as shown in Figure 2.

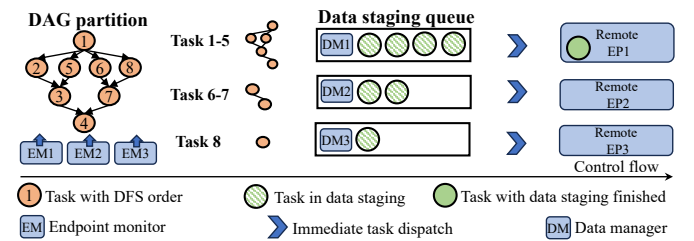


Fig. 2. An illustration of *Capacity*. EPs 1-3 have 5, 2, and 1 workers, respectively. According to the capacity and DFS order, tasks 1–5, tasks 6–7, and task 8 are assigned to endpoint 1, 2, and 3, respectively.

Note that the scheduling decisions of *Capacity* are generated *offline*, i.e., immediately after a workflow DAG is submitted and formed. Therefore, *Capacity* is most suitable for workflows that have static DAGs and run on endpoints with similar hardware performance and static resource capacity.

Locality-aware scheduling (*Locality* in short): Unlike *Capacity*, *Locality* only assigns tasks to an endpoint when there are available resources on the endpoint. Specifically, When assigning a task, *Locality* examines the data distribution for all dependencies of the task on all the endpoints. Based on the data distribution, it computes the amount of data transferred if placed on a specific endpoint and selects the endpoint that leads to the least amount of transfer, denoted as *locality selection* in Figure 3. Upon completion of the data staging, tasks are immediately dispatched to remote endpoints.

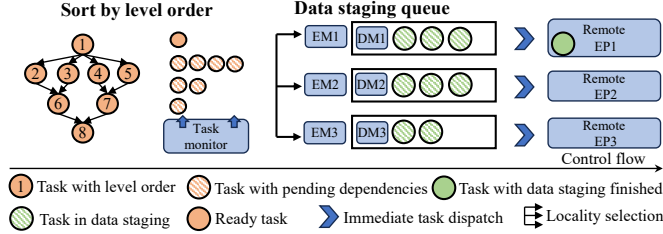


Fig. 3. An illustration of *Locality*. When an endpoint monitor detects idle resources, it performs locality selection for the next ready task and immediately dispatches the task to the target endpoints.

While *Locality* is similar to *Capacity* in that they both intend to reduce data transfers across federated CI, *Locality* produces real-time scheduling decisions whenever there are idles sources and only considers current task states (e.g., data distribution and priority). Therefore, *Locality* is applicable to workflows with dynamic DAGs, as well as workflows that run on endpoints with dynamic resource capacity.

Dynamic heterogeneity-aware scheduling (DHA in short):

DHA assumes that all task information and data transfer information is known a priori, via user input or from the profilers. DHA involves two stages when scheduling, task prioritization and endpoint selection. DHA first calculates the priority of all tasks and then selects an appropriate endpoint for every task in order based on the priority, as shown in Figure 4.

In detail, the priority of a task t_i is recursively computed based on the following equation:

$$\text{priority}(t_i) = \overline{d_i} + \overline{w_i} + \max_{t_j \in \text{succ}(t_i)} \text{priority}(t_j), \quad (2)$$

where $\overline{d_i}$ denotes the average data staging time of task t_i over all the endpoints, $\overline{w_i}$ is the average execution time of task t_i over all the endpoints, and $\text{succ}(t_i)$ is the set of immediate successors (if there is any) of task t_i . The priority calculation is inspired by the heterogeneous earliest finish time (HEFT) [35] algorithm. The recursive calculation ensures that predecessor nodes are assigned to endpoints ahead of successor nodes.

Submitted tasks are scheduled based on the priority. Once all the dependencies of a task are complete (i.e., ready task), DHA selects the target endpoint in a heterogeneity-aware manner that minimizes the completion time. After endpoint selection, DHA implements a *delay scheduling* mechanism to delay the task dispatch to the target endpoint. Specifically, the endpoint selection allows the task to instantiate the data staging immediately whenever a dependency of the task is complete. However,

the task dispatch is delayed until the target endpoint has idle resources. In other words, this mechanism allows tasks with data staging completed to wait in the UniFaaS client queue.

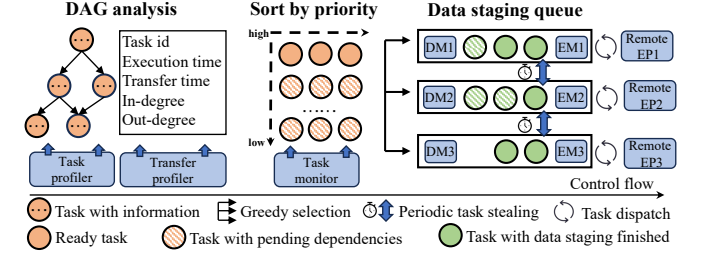


Fig. 4. An illustration of DHA. DHA involves prioritization and endpoint selection. Tasks are not dispatched until the target endpoint has idle resources.

To adapt to dynamic resource capacity during execution, we propose a *re-scheduling* mechanism: whenever the resource capacity changes, DHA periodically recomputes the scheduling decisions of pending tasks (including those with data staging completed) considering several factors (e.g., data movement cost and execution time on different clusters, and available resources) and performs task stealing if necessary. With the *delay* mechanism, the pool of tasks that can be re-scheduled is expanded. The optimization goal of the re-scheduling is to maximize the utilization of dynamic resources while reducing the completion time of the corresponding tasks.

Unlike *Capacity* and *Locality* that either makes offline or real-time scheduling decisions, DHA is a hybrid algorithm between them, i.e., it generates scheduling decisions offline but only dispatches tasks until there are idle resources. The delay and re-scheduling mechanism are key enablers for DHA to support dynamic workflow DAGs and dynamic resource capacity. DHA is applicable to any scenario but requires workflows to have full knowledge (e.g., graph structure, network bandwidth, and task characteristics) to have optimal performance.

TABLE I
SUMMARY OF THE SCHEDULING ALGORITHMS.

	Capacity	Locality	DHA
Scheduling type	Offline	Real-time	Hybrid
Dynamic DAG supported	✗	✓	✓
Dynamic resource supported	✗	✓	✓
Knowledge required	✗	✗	✓

Summary: We summarize the features of the three scheduling algorithms in Table I. The UniFaaS scheduler is designed to be extensible and configurable. Users may extend the scheduler simply to adapt to their use cases.

E. Data Manager

Workflows often involve complex inter-task data dependencies. While UniFaaS can automatically pass small objects among tasks (via *futures*), there may exist large files to stage across federated CI. Some tools support programmatic data transfers but require real-time authentication of computers and pre-knowledge of specific locations, which misaligns with the federated workflows—a task may run on various resources and the location is unknown when programming.

UniFaaS resolves this via a data manager that i) provides a shim layer with `RemoteFile` and `RemoteDirectory` objects, which allows users to wrap data and supports interfaces to perform read/write operations similar to regular files and directories; ii) stages data *transparently* for users using the built-in transfer mechanisms, when a task with data dependencies is scheduled; and iii) monitors the progress of data transfers and retries failed transfers.

Currently, the data manager supports Globus [16] and rsync, two widely used mechanisms for data transfers across different resources. The data manager also enables concurrent data transfers, and the number of threads can be configured based on the connection limit of resources or transfer mechanisms.

F. Task Executor

UniFaaS leverages *funcX* as the backend to run tasks in a FaaS manner. Users must therefore deploy *funcX* endpoints on desired resources to run their tasks. The task executor asynchronously submits tasks and polls results via the *funcX* service. The task executor wraps the task with data staging finished as a *funcX* task, submits the task to the corresponding endpoint via the *funcX* client, and records the task’s returned future. The task executor has a separate thread to poll the results of all pending tasks via the *funcX* client. Upon completion of a task (i.e., when its result is retrieved), the task executor updates the corresponding future and also streams the task execution status such as the execution time and location into the task monitor as described in §IV-B.

G. Fault Tolerance

UniFaaS implements several fault tolerance mechanisms.

Data transfer retry: Data transfers may fail due to network conditions, especially for large data volumes. UniFaaS will retry failed transfers several times (configurable). If all retries fail, the corresponding tasks will be marked as failed.

Task reassignment: A task may fail for various systematic reasons, e.g., data transmission failure, endpoint disconnection, or incorrect runtime environment. For a failed task, UniFaaS attempts to execute the task again on an endpoint according to the scheduler. If it fails again, UniFaaS reassigns it to the endpoint with the highest success rate based on prior runs. If it fails on all endpoints, UniFaaS returns an error message.

H. Optimizations

UniFaaS applies several performance optimizations to ensure the efficient and robust execution of federated workflows.

Multi-endpoint elasticity: The resource requirements of a workflow often vary at different stages. Each *funcX* endpoint itself can dynamically scale: spawning more workers when there are more tasks than workers and killing idle workers when there is no incoming task for a certain period of time. However, each endpoint does not have a global view of workflows and may make suboptimal scaling decisions. With a full view of workflows, UniFaaS can perform multi-endpoint scaling in advance based on the characteristics of workflows.

UniFaaS implements a `Scaling` interface which allows users to implement their own multi-endpoint scaling logic

(e.g., preferences for certain endpoints). The default multi-endpoint scaling strategy is straightforward: if the number of pending tasks in a workflow is more than the number of workers, UniFaaS scales out the workers on all the endpoints; the scale-in logic is left to the endpoints to decide, as each endpoint can scale in if there are idle resources (which may indicate that the endpoint is less preferred by the UniFaaS scheduler). Such an approach that *scales out aggressively but scales in conservatively* works well in most use cases, since killing idle resources is generally easier than requesting resources on federated CI.

Batching: Workflows may be composed of thousands of tasks. To amortize the communication and computation costs for managing tasks and endpoints across various components, UniFaaS implements batching mechanisms at several levels, including task submission, result retrieval, endpoint status polling, and performance prediction, whenever possible.

I. Implementation

We implement a prototype of UniFaaS in Python, which is based on *Parsl* [8], a parallel programming library widely used in science. UniFaaS is open source on GitHub.¹

V. EVALUATION

We evaluate UniFaaS’s performance in terms of several system metrics including latency, scalability, elasticity, and scheduler overhead.

A. Testbeds

To evaluate UniFaaS’s ability to manage scientific workflows across federated CI, we deploy *funcX* endpoints on the following heterogeneous clusters, and submit workflows via UniFaaS on a local workstation. The hardware of these clusters and computers is listed in Table II.

- **Taiyi** is a 2.5-petaflops supercomputer that was in the TOP500 list until recently.
- **Qiming** is a 0.3-petaflops academic supercomputer.
- **Dept. cluster** is a department cluster used primarily for teaching and research.
- **Lab cluster** is a local compute cluster.

TABLE II
HARDWARE OF THE HETEROGENEOUS TESTBED.

Name	CPU	RAM (GB)	# nodes
Taiyi	2*Xeon Gold 6148@2.4GHz	192	815
Qiming	2*Xeon E5-2690@2.6GHz	64	230
Dept. cluster	2*Xeon Platinum 8260@2.4GHz	770	26
Lab cluster	2*Xeon Gold 5320@2.2GHz	128	2
Workstation	Core i5-9400@2.9Ghz	16	1

B. Latency

We measure the latency incurred by each component in UniFaaS, by running a “hello world” task with a 1 MB input file to transfer, and timing the latency across each component. The endpoint is deployed on Qiming and the average of 20

¹<https://github.com/SUSTech-HPCLab/UniFaaS>

runs is reported in Figure 5. The task takes around 1,087 ms to execute. Each UniFaaS component results in only minimal latency. For example, the profilers predict job characteristics and transfer time within 2 ms (included in the scheduling). The local mocking overhead, included in the submission stage, is 0.08 ms. The majority of the latency is from unavoidable data transfer, as well as task dispatching to the remote endpoint and result polling that are highly relevant to the network delay. Nonetheless, the execution times of data analysis tasks in scientific workflows often vary from minutes to hours, in which case this level of overhead is acceptable.

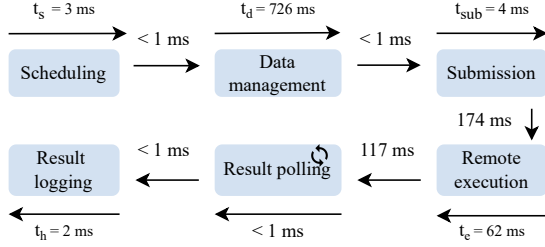


Fig. 5. UniFaaS latency breakdown.

C. Scalability

We evaluate the strong and weak scaling of UniFaaS. Strong scaling measures performance for a fixed number of tasks, as the number of workers increases; weak scaling measures performance for increasing numbers of workers when the average number of tasks *per worker* is fixed. It has been previously demonstrated that a single *funcX* endpoint can scale up to 130,000 workers [14], [15]; hence, we focus here on evaluating UniFaaS scalability when deploying tasks across multiple endpoints—more endpoints lead to higher task scheduling and submission overheads. In this experiment, each endpoint has 24 workers and all endpoints are deployed on Qiming. We create two types of tasks: 1 s and 5 s compute-intensive CPU stress (i.e., while loop) tasks.

Figure 6 shows the strong and weak scaling performance from 1 to 16 endpoints. In the strong scaling case, we measure the performance for a) 100,000 $\times$ 1 s tasks and b) 20,000 $\times$ 5 s tasks. In the weak scaling case, each worker runs, on average, either a) 260 $\times$ 1 s tasks or b) 52 $\times$ 5 s tasks, yielding the same workloads in total for strong and weak scaling on 16 endpoints. We note that 16 endpoints is a sufficient number for many of our scientific use cases, but not a limit of UniFaaS. The results show that the scalability for 5 s tasks is close to the ideal for up to 12 endpoints; with yet longer-duration tasks, we would expect to see good scaling for yet larger numbers of endpoints. The completion time keeps decreasing until six endpoints for 1 s tasks and 12 endpoints for 5 s tasks. The performance of 100,000 $\times$ 1 s tasks is worse than that of 20,000 $\times$ 5 s tasks. This is primarily because a larger number of 1 s tasks suffer from higher network latency and scheduling overheads, causing worse nonlinear scaling.

D. Multi-endpoint Elasticity

To demonstrate the multi-endpoint elasticity of UniFaaS, we deployed three endpoints on Qiming (EP1), Dept. cluster

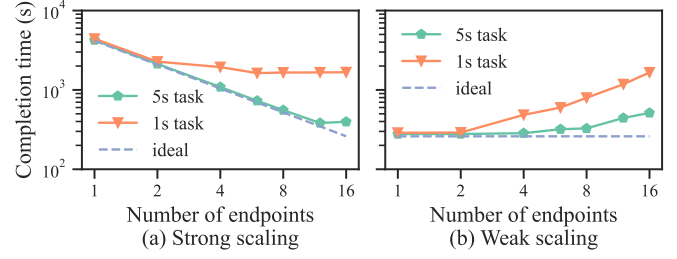


Fig. 6. Strong and weak scaling of UniFaaS.

(EP2), and Lab cluster (EP3), respectively. Each endpoint can scale up and down in terms of nodes and each node has 20 workers, whereas the maximum number of workers of EP1, EP2, and EP3 is set to 100, 40, and 20, respectively. We create three types of compute-intensive stress tasks: 30 s tasks (task1 on EP1), 15 s tasks (task2 on EP2), and 10 s tasks (task3 on EP3). Each endpoint runs a distinct task duration since we want to show that each endpoint can scale independently.

Figure 7 shows the number of pending tasks and active workers versus time. At $t = 10$, we run 50 $\times$ task1, 20 $\times$ task2, and 10 $\times$ task3. Consequently, EP1 scales up to 60 workers, while EP2 and EP3 scale up to 20 workers. At around $t = 50$, since EP3 has been idle for more than 30 seconds (configured maximum idle interval), EP3 returns all the workers. At $t = 70$, we run 200 $\times$ task1, 80 $\times$ task2, and 40 $\times$ task3 on the corresponding endpoints. This time all the endpoints scale up to the maximum number of workers (i.e., 100, 40, and 20). After all the tasks are complete, each endpoint scales down to zero workers. We repeat the above process twice and observe that each endpoint can scale up and down promptly and independently, as expected. It is worth mentioning that the performance of elasticity in practice is subject to the queuing delays of batch schedulers.

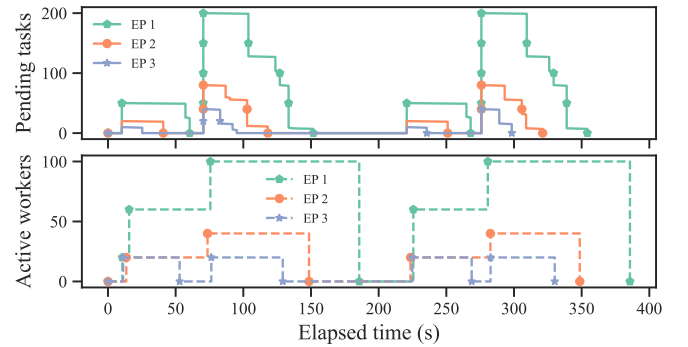


Fig. 7. Number of tasks and workers over time. Top: number of pending tasks. Bottom: number of active workers.

E. Scheduler Overhead

We measure the scheduler overhead (including the time to predict task characteristics if needed) when scheduling a drug screening workflow in Figure 8. Note that this overhead experiment is conducted on the Workstation and better performance may be achievable with a more powerful server. Table III shows the overhead per task of different algorithms. We see that all the algorithms have only a modest overhead.

DHA involves predicting task characteristics and prioritizing the tasks in the DAG, resulting in a slight increase in the overhead.

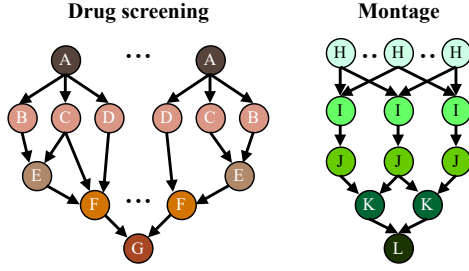


Fig. 8. DAG structures of the drug screening and the montage workflows. Each character represents a distinct task type. The drug screening workflow consists of 24,001 functions. The total computation time is 1,447 hours with an average of 220 seconds per task. The total size of the input, intermediate, and output data is 480.64 GB. The montage workflow consists of 11,340 functions. The total computation time is 108 hours with an average of 6.4 seconds per task. The total size of the input, intermediate, and output data is 673.49 GB.

TABLE III
OVERHEAD OF DIFFERENT ALGORITHMS.

Scheduling algorithm	Overhead (s)
Capacity	1.72×10^{-4}
Locality	3.00×10^{-3}
DHA	3.46×10^{-3}

VI. CASE STUDIES

We use two open-source workflows, a drug screening workflow [36] and a montage workflow [37], to study the characteristics of different scheduling strategies and evaluate UniFaaS's ability to manage workflows across multiple resources.

A. Static resource capacity

In this experiment, we use two workflows shown in Figure 8. When executing the drug screening workflow, we deploy 2000 workers (50 nodes) on Taiyi (EP1), 384 workers (16 nodes) on Qiming (EP2), 48 workers (2 nodes) on Dept. cluster (EP3), and 52 workers (2 nodes) on Lab cluster (EP4). When executing the montage workflow, we deploy 120 workers (4 nodes) on EP1, 240 workers (10 nodes) on EP2, 48 workers (2 nodes) on EP3, and 52 workers (2 nodes) on EP4. All workers are launched before the experiment, and the number of workers is *static* during this experiment. For DHA, we assume full knowledge can be retrieved from the profilers.

TABLE IV
RESULTS FOR STATIC RESOURCE CAPACITY.

Workflow	Experiment	Makespan (s)	Transfer size (GB)
Drug	Capacity	3,240	4.86
	Locality	3,882	53.46
	DHA	2,898	44.94
	Baseline: Only Taiyi	3,763	0
	Baseline: Only Qiming	1,994	0
Montage	Capacity	1,027	2.57
	Locality	1,055	13.35
	DHA	909	18.27
	Baseline: Only Taiyi	1,027	0
	Baseline: Only Qiming	1,994	0

Analysis: Table IV shows the makespan of the workflows under different scheduling algorithms. The makespan

is defined as the completion time of the workflow, including scheduling overhead, polling latency, etc. For both workflows, DHA outperforms *Capacity* and *Locality* in terms of makespan, since DHA can leverage knowledge such as DAG structure and task characteristics. *Capacity* results in the smallest data movement because it is designed for static DAGs and can use certain knowledge like DAG structures to schedule. *Locality* operates without any prior knowledge, thereby minimizing data transfer size to the best extent possible.

To further analyze the reason for the performance variance, we plot in Figure 9 the worker utilization of different scheduling algorithms. For both workflows, DHA achieves consistent high worker utilization. The worker utilization of both *Locality* and *Capacity* gradually decreases, exhibiting a long-tail pattern. However, the root causes of these declines in worker utilization are different. The makespan of *Locality* is severely impacted by the data staging, since *Locality* makes scheduling decisions in real-time and cannot hide the data staging delays, resulting in a longer makespan. This is proven by Figure 10, which shows the number of tasks in data staging over time. *Capacity* makes scheduling decisions offline and hence the data staging can be done immediately after each task's dependencies are completed, overlapping the data staging with the computation. Figure 11 demonstrates that *Capacity* can evenly distribute tasks to endpoints based on the number of workers. DHA is heterogeneity-aware and tends to allocate more tasks to Taiyi with more advanced hardware.

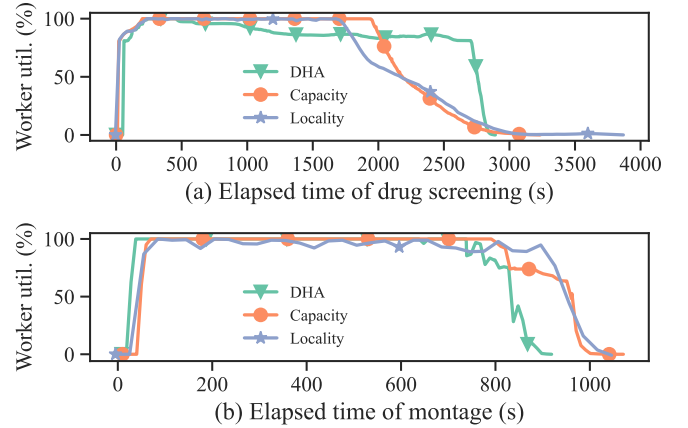


Fig. 9. Worker utilization over time under static resource capacity.

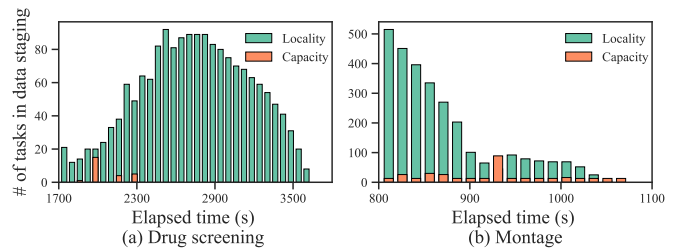


Fig. 10. Locality results in more tasks in data staging state compared with Capacity.

To demonstrate the potential of deploying workflows across federated CI with UniFaaS, we compare the three scheduling

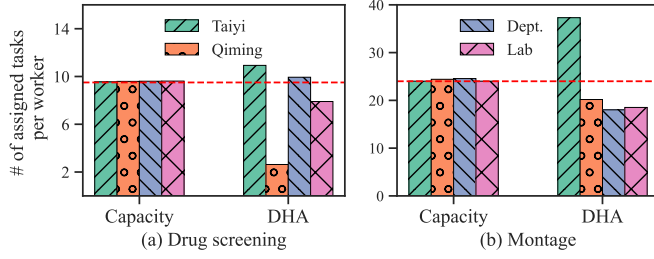


Fig. 11. Workload distribution of `Capacity` and `DHA`. `DHA` prefer `Taiyi`, a higher performance cluster.

algorithms with a baseline, which runs on *only* `Taiyi` (2000 workers) for the drug screening workflow and *only* `Qiming` (240 workers) for the montage workflow. As shown in Table IV, the makespan is improved by up to 22.99% (54.41%) with an additional 19.48% (47.83%) of the workers while executing the drug screening (montage) workflow.

The results demonstrate that *UniFaaS can effectively deploy workflows across federated CI*, even though cross-facility computing may incur a significant amount of data transfers.

B. Dynamic resource capacity

We study the effectiveness of `DHA` under dynamic resource capacity. We execute the drug screening workflow with 12,001 functions and the montage workflow with 11,340 functions. The captions of Figure 12 and Figure 13 show how the amount of resources varies over time (in terms of the number of workers) on the endpoints for the two workflows.

TABLE V
RESULTS FOR DYNAMIC RESOURCE CAPACITY.

Workflow	Experiment	Makespan (s)	Transfer size (GB)
Drug screening	Capacity	3,610	3.26
	Locality	2,130	43.61
	DHA	1,666	33.01
	DHA without re-sched.	2,183	39.47
Montage	Capacity	2,671	2.48
	Locality	1,360	14.18
	DHA	1,257	31.05
	DHA without re-sched.	1,868	29.62

Analysis: Table V shows the makespan of the workflows with dynamic resource capacity for different scheduling algorithms. In both workflows, `DHA` attains the lowest makespan because the re-scheduling mechanism can promptly respond to resource variations and effectively balance the workload across endpoints. As a result, `DHA` improves the makespan by up to 32% compared to `DHA` without re-scheduling. `Locality` reduces the makespan by more than 41% when compared with `Capacity` because the real-time nature of `Locality` enables it to dynamically adapt based on the current state of resources, resulting in more efficient utilization of available capacity, while operates as an offline scheduler. Figure 12(a) and Figure 13(a) imply that `Capacity` fails to balance the tasks to more capable endpoints, resulting in long-tail latency due to the bottleneck endpoints. Figure 12(b) and Figure 13(b) show that `DHA` can quickly re-schedule tasks when there is resource variation.

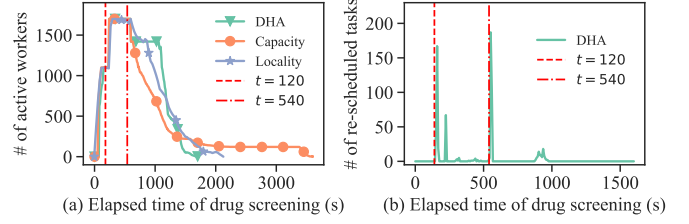


Fig. 12. Running the drug screening workflow under dynamic capacity. Initially, 400, 600, 48, and 52 workers are deployed on EP1, EP2, EP3 and EP4 respectively. At $t = 120$, EP2's worker count increases by 600. At $t = 540$, EP1's worker count decreases by 280.

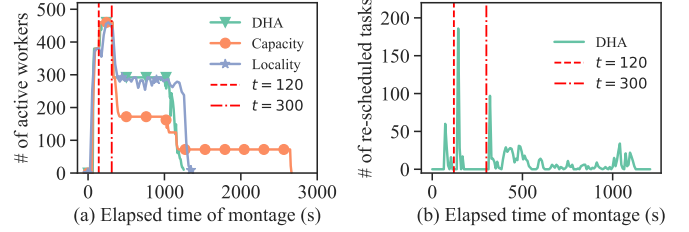


Fig. 13. Running the montage workflow under dynamic capacity. Initially, 40, 240, 48, and 52 workers are deployed on EP1, EP2, EP3 and EP4 respectively. At $t = 120$, EP1's worker count increases by 80. At $t = 300$, EP2's worker count decreases by 168.

VII. LESSONS LEARNED

We discuss our experiences with the case studies using `UniFaaS` and traditional approaches.

Traditional approaches versus UniFaaS: We attempted to compare the performance of `UniFaaS` with traditional systems such as `Pegasus` [10] in the evaluation. However, it was challenging since traditional systems are not designed for those requirements or do not run workflows in a manner like `UniFaaS` does. For example, `Pegasus` is coupled with `HTCondor`, which needs to open a port on each computing resource and establish a direct connection to the submit host to listen for incoming jobs. However, this is rarely allowed by HPC clusters and requires privileged access.

Further, traditional systems launch tasks of workflows at the job level and are not suitable for workflows with dynamic DAGs and dynamic resource capacity. `UniFaaS` differs in that it allows one to programmatically compose a workflow and map the workflow to a resilient resource pool at the function level. `UniFaaS` enables construction of workflows to run across federated cyberinfrastructure using Python, while traditional systems often require domain-specific languages or specifications to describe workflows.

In summary, `UniFaaS` differs from traditional approaches in terms of both *programmability* and *task management*. Such a FaaS-based approach creates new opportunities for considering workflows in which small schedulable units can be flexibly placed across *dynamic* federated computers.

Applicability to federated workflows: We utilized two supercomputers, `Taiyi` and `Qiming`, in our evaluation. `Taiyi`, boasting new generation hardware, usually has longer queue times than `Qiming`. In general, users run workflows on either `Taiyi` or `Qiming`, since it is a burden to manage

workflows across two heterogeneous clusters using traditional approaches. UniFaaS provides the ability to easily explore tradeoffs between hardware performance and queue time, using Taiyi and Qiming, as well as the Dept. cluster and Lab cluster. Based on our observations, the tasks of the workflow use Qiming preferably when Taiyi is busy, and prefer Taiyi when its resources are available.

Programmability and Portability: We developed and debugged the drug screening and montage workflows in \$V locally, and finally ran them across several computing resources. These workflows were originally designed to run a single cluster. We made them executable across distributed CI with minimal effort using UniFaaS: i) decorate each stage as functions; ii) replace the original file I/O operations with `RemoteFile` object; and iii) deploy the resource pool and specify the endpoints' UUID. During the development, the main computation code remained unchanged. Additionally, Qiming was updated from the PBS to the LSF scheduler during our experiments. With UniFaaS, we transitioned the workflow to the updated Qiming by merely replacing the *funcX* endpoint with one configured for LSF. These experiences demonstrate that UniFaaS simplifies the utilization of diverse resources.

VIII. RELATED WORK

FaaS. FaaS platforms are widely offered by most cloud providers [38]–[40]. There are also many open-source FaaS frameworks (e.g., OpenWhisk [41], KNIX [42], and DFaaS [43]) that allow users to deploy on-premise and for different scenarios (e.g., IoT). The success of the FaaS model in clouds motivates us to adopt and extend the FaaS model for modern science workflows. We thus build UniFaaS upon *funcX* [14], [15], a specialized FaaS platform for federated CI.

Several papers [44]–[48] focus on migrating DAG-like applications to use the FaaS model. For instance, PyWren [45] and NumPyWren [46] leverage FaaS for specific types of applications such as MapReduce and linear algebra; Wukong [47] is a serverless parallel programming framework that relies on AWS Lambda. These works are tied to centralized (cloud-hosted) FaaS platforms (e.g., Lambda) or are limited to just one endpoint. Our novelty lies in the use of the FaaS model as a way to distribute computation across federated CI. We focus on the unique challenges of such federated environments (e.g., scheduling and data transfer).

Workflow management systems. While there are many workflow management systems developed [5], [6], [10], [49]–[51], none aim to address the needs of fine-grain task execution across federated CI specifically, to the best of our knowledge. For instance, Pegasus [10] similarly aims to bridge distributed and diverse CI. However, Pegasus relies on a static DAG model and requires HTCondor [11] as the broker to interact with different cluster schedulers. UniFaaS instead leverages a Python-based dynamic DAG model and the flexibility of *funcX* endpoints allows one to simply run on arbitrary resources. Python parallel frameworks (e.g., Dask [9], Parsl [8], Ray [7]) support construction of parallel programs with Python functions and simple deployment on various types of clusters

(e.g., supercomputers and clouds). While UniFaaS uses a similar programming interface, these frameworks are primarily designed for deployment on a single cluster.

Workflow performance modeling and scheduling. Estimating runtimes and other task characteristics of workflows is a well-studied area [30], [33], [52]–[56]. UniFaaS relies on several performance models to predict task characteristics and transfer performance. The profilers of UniFaaS are designed to be extensible to any of these models.

Workflow management systems [5], [10], [57] rely on efficient scheduling algorithms to map workflows to target resources. Many papers propose workflow scheduling algorithms for different scenarios. For example, prior papers [35], [58]–[60] focus on scheduling tasks efficiently to heterogeneous processors. BaRRS [61] leverages task graph partitioning and data replication to reduce the data transfers among different resources. UniFaaS addresses a unique scheduling problem in the federated FaaS scenario, where workflow DAGs could be dynamic and resources may be added (or removed) during execution. However, there are common experiences and techniques in these papers that we can draw lessons from or further integrate into UniFaaS. For example, the priority calculation in the DHA algorithm is adopted from HEFT [35].

IX. CONCLUSION

UniFaaS adopts a federated FaaS model that enables users to focus on *what* to do in the functions and *when* to invoke the functions tasks of workflows, without considering the management of execution. UniFaaS provides a unified, function-based programming interface for expressing task parallelism and composing task graphs, as well as providing transparent data management. Internally, UniFaaS monitors task characteristics, creates performance models for tasks on different endpoints, and decides where to dispatch tasks to achieve high performance in large-scale, heterogeneous, and dynamic environments. We demonstrated that UniFaaS introduces only minimal latency overhead and that UniFaaS can support a workflow to deploy efficiently across up to 16 different endpoints. In the future, we intend to incorporate additional data transfer profilers into UniFaaS, which will consider various factors such as communication patterns and network bandwidth. We will investigate more comprehensive scheduling algorithms and explore the coordination of these algorithms with multi-endpoint elasticity to enhance resource utilization and performance.

ACKNOWLEDGMENT

This work was supported in part by the National Natural Science Foundation of China Grant No. 62202216, the Guangdong Basic and Applied Basic Research Foundation Grant No. 2023A1515010244, the Shenzhen Science and Technology Program Grant 20231121101752002, the DOE contract DE-AC02-06CH11357, and the U.S. National Science Foundation Grant 2004894. This work was also supported by Center for Computational Science and Engineering at Southern University of Science and Technology.

REFERENCES

- [1] T. Shaffer, Z. Li, B. Tovar, Y. Babuji, T. Dasso, Z. Surma, K. Chard, I. Foster, and D. Thain, "Lightweight function monitors for fine-grained management in large scale Python applications," in *IEEE International Parallel and Distributed Processing Symposium*. IEEE, 2021, pp. 786–796.
- [2] Z. Liu, P. Balaprakash, R. Kettimuthu, and I. Foster, "Explaining wide area data transfer performance," in *26th International Symposium on High-Performance Parallel and Distributed Computing (HPDC)*. ACM, 2017, p. 167–178.
- [3] R. F. da Silva, R. M. Badia, V. Bala, D. Bard, P.-T. Bremer, I. Buckley, S. Caino-Lores, K. Chard, C. Goble, S. Jha *et al.*, "Workflows community summit 2022: A roadmap revolution," *arXiv preprint arXiv:2304.00019*, 2023.
- [4] A. Alsaadi, L. Ward, A. Merzky, K. Chard, I. Foster, S. Jha, and M. Turilli, "Radical-Pilot and Parsl: Executing heterogeneous workflows on HPC platforms," *arXiv preprint arXiv:2105.13185*, 2021.
- [5] P. Di Tommaso, M. Chatzou, E. W. Floden, P. P. Barja, E. Palumbo, and C. Notredame, "Nextflow enables reproducible computational workflows," *Nature Biotechnology*, vol. 35, no. 4, pp. 316–319, 2017.
- [6] M. Wilde, M. Hategan, J. M. Wozniak, B. Clifford, D. S. Katz, and I. Foster, "Swift: A language for distributed parallel scripting," *Parallel Computing*, vol. 37, no. 9, pp. 633–652, 2011.
- [7] P. Moritz, R. Nishihara, S. Wang, A. Tumanov, R. Liaw, E. Liang, M. Elibol, Z. Yang, W. Paul, M. I. Jordan *et al.*, "Ray: A distributed framework for emerging AI applications," in *13th USENIX Symposium on Operating Systems Design and Implementation (OSDI)*, 2018, pp. 561–577.
- [8] Y. Babuji, A. Woodard, Z. Li, D. S. Katz, B. Clifford, R. Kumar, L. Lacinski, R. Chard, J. M. Wozniak, I. Foster *et al.*, "ParSl: Pervasive parallel programming in Python," in *28th International Symposium on High-Performance Parallel and Distributed Computing*, 2019, pp. 25–36.
- [9] M. Rocklin, "Dask: Parallel computation with blocked algorithms and task scheduling," in *14th Python in Science Conference*, vol. 130. SciPy Austin, TX, 2015, p. 136.
- [10] E. Deelman, K. Vahi, G. Juve, M. Rynge, S. Callaghan, P. J. Maechling, R. Mayani, W. Chen, R. F. Da Silva, M. Livny *et al.*, "Pegasus, a workflow management system for science automation," *Future Generation Computer Systems*, vol. 46, pp. 17–35, 2015.
- [11] D. Thain, T. Tannenbaum, and M. Livny, "Distributed computing in practice: The Condor experience," *Concurrency and Computation: Practice and Experience*, vol. 17, no. 2–4, pp. 323–356, 2005.
- [12] R. Madduri, K. Chard, M. D'Arcy, S. C. Jung, A. Rodriguez, D. Sulakhe, E. Deutsch, C. Funk, B. Heavner, M. Richards *et al.*, "Reproducible big data science: A case study in continuous fairness," *PLoS one*, vol. 14, no. 4, p. e0213013, 2019.
- [13] E. Jonas, J. Schleier-Smith, V. Sreekanti, C.-C. Tsai, A. Khandelwal, Q. Pu, V. Shankar, K. Carreira, K. Krauth, N. Yadwadkar, J. Gonzalez, R. A. Popa, I. Stoica, and D. A. Patterson, "Cloud programming simplified: A Berkeley view on serverless computing," *arXiv preprint arXiv:1902.03383*, 2019.
- [14] Z. Li, R. Chard, Y. Babuji, B. Galewsky, T. J. Skluzacek, K. Nagaitsev, A. Woodard, B. Blaiszik, J. Bryan, D. S. Katz *et al.*, "funcX: Federated function as a service for science," *IEEE Transactions on Parallel and Distributed Systems*, vol. 33, no. 12, pp. 4948–4963, 2022.
- [15] R. Chard, Y. Babuji, Z. Li, T. Skluzacek, A. Woodard, B. Blaiszik, I. Foster, and K. Chard, "Funcx: A federated function serving fabric for science," in *29th International Symposium on High-performance Parallel and Distributed Computing*, 2020, pp. 65–76.
- [16] K. Chard, S. Tuecke, and I. Foster, "Efficient and secure transfer, synchronization, and sharing of big data," *IEEE Cloud Computing*, vol. 1, no. 3, pp. 46–55, 2014.
- [17] D. Balouek-Thomert, E. G. Renart, A. R. Zamani, A. Simonet, and M. Parashar, "Towards a computing continuum: Enabling edge-to-cloud integration for data-driven workflows," *The International Journal of High Performance Computing Applications*, vol. 33, no. 6, pp. 1159–1174, 2019.
- [18] M. AbdelBaky, M. Zou, A. R. Zamani, E. Renart, J. Diaz-Montes, and M. Parashar, "Computing in the continuum: Combining pervasive devices and services to support data-driven applications," in *2017 IEEE 37th International Conference on Distributed Computing Systems (ICDCS)*, 2017, pp. 1815–1824.
- [19] D. Rosendo, A. Costan, P. Valduriez, and G. Antoniu, "Distributed intelligence on the edge-to-cloud continuum: A systematic literature review," *Journal of Parallel and Distributed Computing*, vol. 166, pp. 71–94, 2022.
- [20] A. Clyde, S. Galanie, D. W. Kneller, H. Ma, Y. Babuji, B. Blaiszik, A. Brace, T. Bretin, K. Chard, R. Chard, L. Coates, I. Foster, D. Hauner, V. Kertesz, N. Kumar, H. Lee, Z. Li, A. Merzky, J. G. Schmidt, L. Tan, M. Titov, A. Trifan, M. Turilli, H. Van Dam, S. C. Chennubhotla, S. Jha, A. Kovalevsky, A. Ramanathan, M. S. Head, and R. Stevens, "High-throughput virtual screening and validation of a SARS-CoV-2 main protease noncovalent inhibitor," *Journal of Chemical Information and Modeling*, vol. 62, no. 1, pp. 116–128, 2022.
- [21] J. G. Pauloski, V. Hayot-Sasson, L. Ward, N. Hudson, C. Sabino, M. Baughman, K. Chard, and I. Foster, "Accelerating communications in federated applications with transparent object proxies," in *Proceedings of the International Conference for High Performance Computing, Networking, Storage and Analysis*, 2023.
- [22] L. Ward, G. Sivaraman, J. G. Pauloski, Y. Babuji, R. Chard, N. Dandu, P. C. Redfern, R. S. Assary, K. Chard, L. A. Curtiss *et al.*, "Colmena: Scalable machine-learning-based steering of ensemble simulations for high performance computing," in *2021 IEEE/ACM Workshop on Machine Learning in High Performance Computing Environments (MLHPC)*. IEEE, 2021, pp. 9–20.
- [23] R. Chard, Z. Li, K. Chard, L. Ward, Y. Babuji, A. Woodard, S. Tuecke, B. Blaiszik, M. J. Franklin, and I. Foster, "Dlhub: Model and data serving for science," in *2019 IEEE International Parallel and Distributed Processing Symposium (IPDPS)*. IEEE, 2019, pp. 283–292.
- [24] Z. Li, R. Chard, L. Ward, K. Chard, T. J. Skluzacek, Y. Babuji, A. Woodard, S. Tuecke, B. Blaiszik, M. J. Franklin *et al.*, "Dlhub: Simplifying publication, discovery, and use of machine learning models in science," *Journal of Parallel and Distributed Computing*, vol. 147, pp. 64–76, 2021.
- [25] T. J. Skluzacek, R. Wong, Z. Li, R. Chard, K. Chard, and I. Foster, "A serverless framework for distributed bulk metadata extraction," in *30th International Symposium on High-Performance Parallel and Distributed Computing*, 2021, pp. 7–18.
- [26] A. E. Woodard, A. Trisovic, Z. Li, Y. Babuji, R. Chard, T. Skluzacek, B. Blaiszik, D. S. Katz, I. Foster, and K. Chard, "Real-time hep analysis with funcx, a high-performance platform for function as a service," in *EPJ Web of Conferences*, vol. 245. EDP Sciences, 2020, p. 07046.
- [27] Carbon Aware Workflow Scheduling. <https://github.com/AK2000/caws>. Accessed February 2024.
- [28] A. A. Saadi, D. Alfe, Y. Babuji, A. Bhati, B. Blaiszik, A. Brace, T. Bretin, K. Chard, R. Chard, A. Clyde, P. Coveney, I. Foster, T. Gibbs, S. Jha, K. Keipert, D. Kranzlmüller, T. Kurth, H. Lee, Z. Li, H. Ma, G. Mathias, A. Merzky, A. Partin, A. Ramanathan, A. Shah, A. Stern, R. Stevens, L. Tan, M. Titov, A. Trifan, A. Tsaris, M. Turilli, H. Van Dam, S. Wan, D. Wifling, and J. Yin, "IMPECCABLE: Integrated Modeling Pipeline for COVID Cure by Assessing Better LEads," in *50th International Conference on Parallel Processing*, ser. ICPP 2021, 2021.
- [29] T. Shu, Y. Guo, J. Wozniak, X. Ding, I. Foster, and T. Kurc, "Bootstrapping in-situ workflow auto-tuning via combining performance models of component applications," in *International Conference for High Performance Computing, Networking, Storage and Analysis*, 2021, pp. 1–15.
- [30] T.-P. Pham, J. J. Durillo, and T. Fahringer, "Predicting workflow task execution time in the cloud using a two-stage machine learning approach," *IEEE Transactions on Cloud Computing*, vol. 8, no. 1, pp. 256–268, 2017.
- [31] A. Singh, A. Rao, S. Purawat, and I. Altintas, "A machine learning approach for modular workflow performance prediction," in *12th Workshop on Workflows in Support of Large-scale Science*, 2017, pp. 1–11.
- [32] T. Chen and C. Guestrin, "XGBoost: A scalable tree boosting system," in *22nd ACM International Conference on Knowledge Discovery and Data Mining*, 2016, pp. 785–794.
- [33] J. Bader, F. Lehmann, L. Thamsen, J. Will, U. Leser, and O. Kao, "Lotaru: Locally estimating runtimes of scientific workflow tasks in heterogeneous clusters," in *34th International Conference on Scientific and Statistical Database Management*, 2022.
- [34] J. D. Ullman, "NP-complete scheduling problems," *Journal of Computer and System sciences*, vol. 10, no. 3, pp. 384–393, 1975.
- [35] H. Topcuoglu, S. Hariri, and M.-Y. Wu, "Performance-effective and low-complexity task scheduling for heterogeneous computing," *IEEE Transactions on Parallel and Distributed Systems*, vol. 13, no. 3, pp. 260–274, 2002.

- [36] Y. Babuji, B. Blaiszik, T. Brettin, K. Chard, R. Chard, A. Clyde, I. Foster, Z. Hong, S. Jha, Z. Li *et al.*, “Targeting SARS-CoV-2 with AI and HPC-enabled lead generation: A first data release,” *arXiv preprint arXiv:2006.02431*, 2020.
- [37] G. B. Berriman, E. Deelman, J. C. Good, J. C. Jacob, D. S. Katz, C. Kesselman, A. C. Laity, T. A. Prince, G. Singh, and M.-H. Su, “Montage: a grid-enabled engine for delivering custom science-grade mosaics on demand,” in *Optimizing scientific return for astronomy through information technologies*, vol. 5493. SPIE, 2004, pp. 221–232.
- [38] Amazon Lambda. <https://aws.amazon.com/lambda>. Accessed February 2024.
- [39] Azure Functions. <https://azure.microsoft.com/en-us/services/functions/>. Accessed February 2024.
- [40] Google Cloud Functions. <https://cloud.google.com/functions/>. Accessed February 2024.
- [41] Apache OpenWhisk. <http://openwhisk.apache.org/>. Accessed February 2024.
- [42] KNIX MicroFunctions. <https://github.com/knix-microfunctions/knix>. Accessed February 2024.
- [43] M. Ciavotta, D. Motterlini, M. Savi, and A. Tundo, “DFaaS: Decentralized function-as-a-service for federated edge computing,” in *IEEE 10th International Conference on Cloud Networking (CloudNet)*. IEEE, 2021, pp. 1–4.
- [44] M. Malawski, A. Gajek, A. Zima, B. Balis, and K. Figiela, “Serverless execution of scientific workflows: Experiments with HyperFlow, AWS Lambda and Google Cloud Functions,” *Future Generation Computer Systems*, vol. 110, pp. 502–514, 2020.
- [45] E. Jonas, Q. Pu, S. Venkataraman, I. Stoica, and B. Recht, “Occupy the cloud: Distributed computing for the 99%,” in *Symposium on Cloud Computing*. ACM, 2017, p. 445–451.
- [46] V. Shankar, K. Krauth, K. Vodrahalli, Q. Pu, B. Recht, I. Stoica, J. Ragan-Kelley, E. Jonas, and S. Venkataraman, “Serverless linear algebra,” in *11th ACM Symposium on Cloud Computing*. ACM, 2020, p. 281–295.
- [47] B. Carver, J. Zhang, A. Wang, A. Anwar, P. Wu, and Y. Cheng, “Wukong: A scalable and locality-enhanced framework for serverless parallel computing,” in *11th ACM Symposium on Cloud Computing*. ACM, 2020, p. 1–15.
- [48] A. Mahgoub, K. Shankar, S. Mitra, A. Klimovic, S. Chaterji, and S. Bagchi, “SONIC: Application-aware data passing for chained serverless applications,” in *USENIX Annual Technical Conference (ATC)*, 2021, pp. 285–301.
- [49] A. Jain, S. P. Ong, W. Chen, B. Medasani, X. Qu, M. Kocher, M. Brafman, G. Petretto, G.-M. Rignanes, G. Hautier, D. Gunter, and K. A. Persson, “FireWorks: A dynamic workflow system designed for high-throughput applications,” *Concurrency and Computation: Practice and Experience*, vol. 27, no. 17, pp. 5037–5059, 2015.
- [50] Luigi. <https://github.com/spotify/luigi>. Accessed February 2024.
- [51] Apache Airflow. <https://airflow.apache.org/>. Accessed February 2024.
- [52] M. H. Hilman, M. A. Rodriguez, and R. Buyya, “Task runtime prediction in scientific workflows using an online incremental learning approach,” in *IEEE/ACM 11th International Conference on Utility and Cloud Computing (UCC)*, 2018, pp. 93–102.
- [53] M. R. Wyatt, S. Herbein, T. Gamblin, A. Moody, D. H. Ahn, and M. Tauber, “PRIONN: Predicting runtime and io using neural networks,” in *47th International Conference on Parallel Processing*, 2018.
- [54] R. F. da Silva, G. Juve, E. Deelman, T. Glatard, F. Desprez, D. Thain, B. Tovar, and M. Livny, “Toward fine-grained online task characteristics estimation in scientific workflows,” in *8th Workshop on Workflows in Support of Large-Scale Science*, 2013.
- [55] F. Nadeem, D. Alghazzawi, A. Mashat, K. Fakeeh, A. Almalaise, and H. Hagra, “Modeling and predicting execution time of scientific workflows in the grid using radial basis function neural network,” *Cluster Computing*, vol. 20, pp. 2805–2819, 2017.
- [56] J. Yu, M. Gao, Y. Li, Z. Zhang, W. H. Ip, and K. L. Yung, “Workflow performance prediction based on graph structure aware deep attention neural network,” *Journal of Industrial Information Integration*, vol. 27, p. 100337, 2022.
- [57] M. Bux, J. Brandt, C. Lipka, K. Hakimzadeh, J. Dowling, and U. Leser, “SAASFEE: Scalable scientific workflow execution engine,” *Proceedings of the VLDB Endowment*, vol. 8, no. 12, pp. 1892–1895, 2015.
- [58] Y. Dai and X. Zhang, “A synthesized heuristic task scheduling algorithm,” *The Scientific World Journal*, vol. 2014, 2014.
- [59] J. G. Barbosa and B. Moreira, “Dynamic scheduling of a batch of parallel task jobs on heterogeneous clusters,” *Parallel Computing*, vol. 37, no. 8, pp. 428–438, 2011.
- [60] R. Kumar, M. Baughman, R. Chard, Z. Li, Y. Babuji, I. Foster, and K. Chard, “Coding the computing continuum: Fluid function execution in heterogeneous computing environments,” in *Heterogeneity in Computing Workshop*, 2021.
- [61] I. Casas, J. Taheri, R. Ranjan, L. Wang, and A. Y. Zomaya, “A balanced scheduler with data reuse and replication for scientific workflows in cloud computing systems,” *Future Generation Computer Systems*, vol. 74, pp. 168–178, 2017.